

Analisis Sentimen Terhadap Kinerja Kepolisian Indonesia Menggunakan Metode Multinomial Naive Bayes, Long Short-Term Memory, dan Lexicon-Based

Finsa Nurpandi¹, Fietri Setiawati Sulaeman², Aditya Hermawan³

Program Studi Teknik Informatika

Fakultas Teknik

Universitas Suryakencana

finsa@unsur.ac.id¹, fietrisetiawati@gmail.com², adityahh.rm@gmail.com³

Abstract

According to the Indonesian Survey Institute, the public trust index in the Indonesian Police decreased to 53% as of October 2022. The level of public satisfaction and trust in police can be measured by examining how the police handle each case that occurs. This research focuses on analyzing and measuring the public's sentiment toward the police's performance among Twitter users in Indonesia. After scraping (collecting) data from Twitter, out of the 14,000 collected tweets, only 12,222 tweets were for analysis after the preprocessing stage. Then, the data was labeled into two categories, negative sentiment and positive sentiment using InSet Lexicon. The dataset is divided into three, namely 80% for data training, 20% for data testing, and 15% for data validation. The results showed that 53% (6,478) tweets have a negative sentiment, while 47% (5,744) tweets have a positive sentiment. The topics frequently appear and discussed related to such as the shooting case of Brigadier J, the Kanjuruhan tragedy, the security of the G20 Bali Summit, securing Christmas 2022 and securing New Year 2023. Subsequently, classification was performed on the test data using the Multinomial Naïve Bayes and Long Short-Term Memory (LSTM) methods. In the testing phase, a confusion matrix used to calculate accuracy, recall, precision, and f1-score. The results showed that the Multinomial Naïve Bayes method achieved an accuracy score of 78%, precision of 78.5%, recall of 77%, and an f1-score of 77%, with an average score of 77.6%. Similarly, the LSTM method resulted in an equal accuracy, precision, recall, and f1-score of 87%.

Keywords: Sentiment Analysis, Inset Lexicon, LSTM, Multinomial Naive Bayes, Twitter

Abstrak

Menurut Lembaga Survei Indonesia indeks kepercayaan masyarakat terhadap Kepolisian Indonesia menurun ke angka 53% per Oktober 2022. Tingkat kepuasan dan kepercayaan masyarakat terhadap kinerja kepolisian tersebut dapat diukur dengan melihat bagaimana kepolisian menyelesaikan setiap kasus yang terjadi. Penelitian ini bertujuan untuk menganalisis dan mengukur sentimen masyarakat terhadap kinerja kepolisian di kalangan pengguna Twitter di Indonesia. Setelah melakukan pengumpulan (scraping) data Twitter, dari 14.000 data cuitan yang terkumpul, hanya 12.222 data yang dapat dianalisis setelah melalui tahap preprocessing. Kemudian, data tersebut dilabeli ke dalam dua kategori yaitu sentimen negatif dan sentimen positif dengan menggunakan InSet Lexicon. Dataset dibagi menjadi tiga yaitu data training sebanyak 80%, data testing sebanyak 20%, dan data validation sebanyak 15%. Hasil penelitian menunjukkan bahwa 53% (6.478) data memiliki sentimen negatif, sementara 47% (5.744) data memiliki sentimen positif. Dari klasifikasi tersebut didapat bahwa topik atau kasus yang sering muncul dan dibahas diantaranya adalah kasus penembakan terhadap Brigadir J, tragedi kanjuruhan, pengamanan KTT G20 Bali, pengamanan Natal 2022 serta Tahun Baru 2023. Selanjutnya, dilakukan pengklasifikasian pada data uji dan pengujian menggunakan metode Multinomial Naïve Bayes dan Long Short-Term Memory (LSTM). Pada tahap pengujian tersebut menggunakan confusion matrix untuk menghitung nilai accuracy, recall, precision, dan f1-score. Dari hasil pengujian menunjukkan bahwa metode Multinomial Naïve Bayes menghasilkan skor accuracy sebesar 78%, precision 78.5%, recall 77%, dan f1-score sebesar 77% dengan rata-rata skor sebesar 77.6%. Kemudian untuk metode LSTM menghasilkan skor accuracy, precision, recall, dan f1-score sama besar yaitu 87%.

Kata kunci: Analisis Sentimen, Inset Lexicon, LSTM, Multinomial Naive Bayes, Twitter

I. PENDAHULUAN

Menurut datareportal tahun 2022, dikatakan bahwa terdapat sekitar 5.34 miliar pengguna internet di dunia, dan 204.7 juta diantaranya adalah pengguna internet di Indonesia. Dari jumlah tersebut, sebanyak 191.4 juta

adalah user yang menggunakan internet untuk mengakses jejaring sosial. Beberapa jejaring sosial yang banyak digunakan diantaranya adalah Twitter atau sekarang disebut X, Whatsapp, Facebook, Instagram, Youtube, Telegram, dan lain-lain [1]. Berdasarkan laporan We Are Social menunjukkan, Indonesia menjadi negara dengan pengguna Twitter terbesar kelima di dunia dengan 18.45

juta pengguna. Jumlah tersebut setara dengan 4.23% dari total pengguna Twitter di dunia, dan bahkan jumlah tersebut naik 31.3% dibandingkan tahun sebelumnya [2].

Dengan adanya jejaring sosial Twitter, masyarakat dapat dengan bebas mengutarakan pendapatnya mengenai suatu hal, salah satunya tentang kinerja kepolisian di Indonesia. Kepolisian Negara Republik Indonesia merupakan instrumen negara yang berperan dalam memelihara keamanan dan ketertiban masyarakat, menegakkan hukum serta memberikan perlindungan, pengayoman dan pelayanan kepada masyarakat dalam rangka terpeliharanya keamanan dalam negeri sesuai dengan Pasal 5 UU No.2 Tahun 2002 [3]. Berdasarkan hasil survei Lembaga Survei Indonesia (LSI), survei indeks kepercayaan masyarakat terhadap Polri per Oktober 2022 merosot di angka 53 persen, dimana Polri pernah mendapat penilaian tertinggi kepercayaan publik pada tahun sebelumnya yaitu sebesar 80.2 persen. Kepolisian Republik Indonesia (Polri) didera rentetan permasalahan yang susul menyusul [4]. Beberapa rapor merah yang paling mendapatkan sorotan publik serta menurunkan kepercayaan masyarakat terhadap Polri, adalah kasus penembakan terhadap Brigadir J oleh mantan Kepala Divisi Profesi dan Pengamanan (Divpropam) Polri, tragedi Kanjuruhan di Kabupaten Malang, dan kasus peredaran gelap narkoba yang dilakukan mantan Kapolda Sumatera Barat. Meski banyak mendapat rapor merah, beberapa catatan positif masyarakat terhadap kinerja Polri ikut meningkatkan kepercayaan publik, seperti pengamanan mudik Idul Fitri yang berjalan lancar dan kondusif, pengamanan KTT G20 di Bali, hingga pengamanan Natal 2022 dan Tahun Baru 2023.

Twitter juga seringkali digunakan untuk mengungkapkan emosi mengenai suatu hal, baik memuji ataupun memaki dalam bentuk emosi [5]. Emosi ini dapat dikenali dengan analisa opini atau sentimen (opinion analysis atau sentiment analysis) [6]. Analisis sentimen (sentiment analysis) atau biasa disebut opinion mining ini merupakan bidang studi untuk menganalisis pendapat, evaluasi penilaian, sikap, dan emosi terhadap suatu produk, layanan, organisasi, individu, isu sosial, dan hal lainnya. Analisis sentimen ini dilakukan dengan proses text mining untuk memperoleh hasil sentimennya. Text Mining memiliki tujuan untuk menemukan pola pada data agar dapat dimanfaatkan untuk memprediksi dan mengambil keputusan di masa depan. Teks yang akan dianalisa merupakan teks dengan data yang tidak terstruktur, sehingga data inilah yang harus diperbaiki agar teks tersebut menjadi data yang lebih terstruktur berdasarkan tahapan-tahapan yang ada [7]. Dalam menemukan pola tersebut, penggunaan metode klasifikasi machine learning seperti Multinomial Naïve Bayes, Long Short-Term Memory (LSTM), dan metode lainnya dapat dijadikan acuan untuk penentuan pola serta pengujian hasil akhir dengan bantuan dari Lexicon-Based sebagai salah satu metode untuk pengklasifikasian.

Metode Multinomial Naive Bayes (MNB) merupakan salah satu metode probabilistic reasoning yang memanfaatkan nilai probabilitas dari setiap data yang digunakan. MNB merupakan metode supervised learning, sehingga setiap data perlu diberikan label sebelum dilakukan training [8]. Algoritma MNB bekerja pada

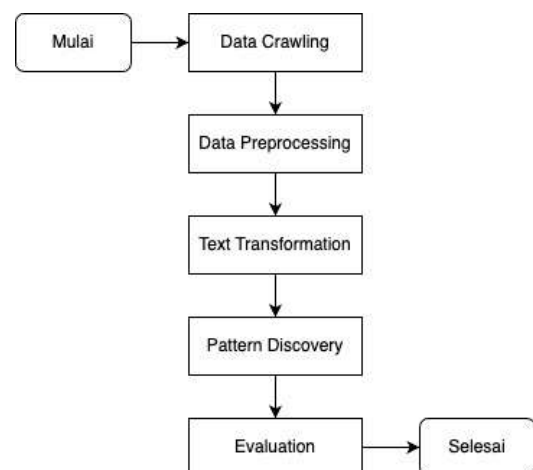
konsep *term frequency* yang berarti berapa kali kata yang ditentukan muncul dalam sebuah dokumen sehingga meningkatkan frekuensi istilah dalam model ini tetapi pada saat yang bersamaan, kata ini juga dapat berupa kata penghenti yang berpotensi tidak menambah makna pada dokumen, namun memiliki frekuensi istilah yang tinggi [9]. Model MNB menjelaskan tentang dua fakta yaitu apakah kata tersebut muncul pada sebuah dokumen atau tidak serta frekuensinya kemunculan dalam dokumen [10].

Metode LSTM merupakan metode yang mengimplementasikan konsep neural network atau saraf jaringan yang mampu mengatasi ketergantungan jangka panjang (long term dependencies) pada inputannya. Sebuah cell dalam LSTM menyimpan sebuah nilai atau keadaan (cell state), baik untuk periode waktu yang panjang atau singkat [11]. LSTM terdiri dari tiga gates, yaitu: Input Gate, Forget Gate, dan Output Gate. Gate yang memiliki fungsi untuk menentukan apakah informasi layak disimpan atau dilupakan [12]. LSTM mempunyai memory block yang akan menentukan nilai mana yang akan dipilih sebagai keluaran yang relevan terhadap masukan yang diberikan [13].

Metode Lexicon-Based berbasis *Dictionary Based Approach* [14] dapat digunakan untuk pengklasifikasian sentimen berdasarkan kamus data yang digunakan dan dapat menghasilkan sebuah pengklasifikasian kedalam klasifikasi data bersentimen negatif ataupun positif [15]. Hasil klasifikasi dari metode ini dapat digunakan untuk pengujian dengan menggunakan metode Multinomial Naive Bayes dan LSTM.

II. METODE PENELITIAN

Penelitian ini menggunakan data yang bersumber dari cuitan pengguna sosial Twitter yang mengacu pada tagar #Polri yang diambil dari 1 Juni Hingga 30 Desember Tahun 2022, sehingga didapat sekitar 14.000 data. Metode penelitian ini menggunakan model tahapan Text Mining. Pada prosesnya, data akan dibagi menjadi dua yaitu data training sebanyak 80%, data testing sebanyak 20%. Sedangkan untuk proses validasi, akan menggunakan 15% dari data yang tersedia. Tahapan pada *Text Mining* untuk melakukan analisis sentimen dapat dilihat pada gambar di bawah ini [16].



Gambar 1: Tahapan Alur Text Mining

Proses data cleaning akan melakukan penghapusan karakter atau kata-kata yang tidak akan digunakan pada tahap preprocessing. Tahap ini akan melakukan penghapusan *mentions*, *hashtag*, *hyperlink*, *number*, dan *whitespace*. Selanjutnya adalah melakukan pengecekan terhadap duplikasi data, jika terdapat duplikasi maka akan dilakukan penghapusan data atau baris data tersebut.

Terdapat beberapa tahapan pada proses *preprocessing*, dimulai dari *Casefolding* yaitu merubah bentuk kalimat yang sudah dibersihkan *lowercase*. *Tokenizing* merupakan proses merubah kalimat yang telah di *casefolding* ke dalam token-token atau menjadi perkata yang dipisahkan oleh delimiter spasi. *Stopwords* melakukan proses penghapusan kata yang dinilai tidak cukup penting berdasarkan *list stopwords* bahasa Indonesia yang digunakan. *Stemming* merupakan proses akhir yang akan mengubah kata hasil *stopwords* ke dalam bentuk *root* atau kata dasar.

Pada tahap Text Transformation merupakan transformasi data serta pembentukan atribut yang mengacu pada proses untuk mendapatkan representasi dokumen yang diharapkan. Perubahan data menjadi numerikal atau angka untuk membantu mesin melakukan perhitungan pada model yang digunakan. Pada penelitian ini akan menggunakan aturan Tokenizer dan pembobotan TF-IDF untuk model yang dibangun. Persamaan yang digunakan oleh metode ini adalah seperti di bawah ini.

$$w_{i,j} = tf_{i,j} \times \log \left(\frac{N + 1}{df_i + 1} \right) + 1 \quad (1)$$

Pattern discovery akan menjelaskan metode serta model yang digunakan untuk pengklasifikasian dan menemukan pola pada keseluruhan dataset yang telah diolah. Pada proses ini menggunakan metode Inset Lexicon untuk pengklasifikasian, Multinomial Naive Bayes dan LSTM untuk menemukan pola serta mengetahui hasil akurasi model. Multinomial Naive Bayes dapat diformulasikan menggunakan persamaan (2) [17]:

$$P(p|n) \propto P(p) \prod_{1 \leq k \leq n} P(t_k|p) \quad (2)$$

Dimana $P(t_k|p)$ merupakan probabilitas munculnya dokumen teks (t_k). N adalah jumlah dokumen dan p adalah polaritas. Untuk menghitung nilai polaritas dapat menggunakan persamaan berikut:

$$P(t_k|p) = \frac{\text{count}(t_k|p) + 1}{\text{count}(t_p) + |V|} \quad (3)$$

Berikut di bawah ini adalah persamaan-persamaan yang digunakan pada LSTM:

$$f_t = \sigma(W_f[h_{t-1}, x_t] + b_f) \quad (4)$$

$$i_t = \sigma(W_i[h_{t-1}, x_t] + b_i) \quad (5)$$

$$o_t = \sigma(W_o[h_{t-1}, x_t] + b_o) \quad (6)$$

$$\tilde{c}_t = \tanh(W_c[h_{t-1}, x_t] + b_c) \quad (7)$$

$$c_t = f_t * c_{t-1} + i_t * \tilde{c}_t \quad (8)$$

$$h_t = o_t * \tanh(c^t) \quad (9)$$

Tahap terakhir adalah evaluasi yang akan ditampilkan dalam bentuk visualisasi.

III. HASIL DAN PEMBAHASAN

Hasil penelitian analisis sentimen ini dilakukan data collection yang menghasilkan data sebanyak 14.000 tweet untuk menganalisis sentimen pengguna twitter dengan variabel kinerja kepolisian. Kata pencarian yang digunakan adalah dengan menggunakan hashtag #polisi yang dikumpulkan dari 1 Juni hingga 30 Desember 2022.

3.1 Data Preprocessing

Tahapan data preprocessing dilakukan dengan menggunakan lima tahapan yang telah dijelaskan pada tahapan metode penelitian. Berikut hasil proses *data preprocessing* pada suatu tweet hasil proses *data crawling*.

Tabel 1: Hasil Text Cleaning

Text	Text Cleaning
Saya mengucapkan terima kasih kepada seluruh personel Polri dimanapun bertugas, tetaplah semangat dalam memberikan pelayanan terbaik kepada masyarakat, sehingga kita bisa menjadi Polisi yang tegas, humanis, dekat dan dicintai masyarakat. https://t.co/Xqb23gckuA #Polri	Saya mengucapkan terima kasih kepada seluruh personel Polri dimanapun bertugas tetaplah semangat dalam memberikan pelayanan terbaik kepada masyarakat sehingga kita bisa menjadi Polisi yang tegas humanis dekat dan dicintai masyarakat

Tabel 2: Hasil Casefolding

Text	Casefolding
Saya mengucapkan terima kasih kepada seluruh personel Polri dimanapun bertugas tetaplah semangat dalam memberikan pelayanan terbaik kepada masyarakat sehingga kita bisa menjadi Polisi yang tegas humanis	saya mengucapkan terima kasih kepada seluruh personel polri dimanapun bertugas tetaplah semangat dalam memberikan pelayanan terbaik kepada masyarakat sehingga kita bisa menjadi polisi yang tegas humanis

dekat dan dicintai masyarakat	dekat dan dicintai masyarakat
-------------------------------	-------------------------------

Tabel 3: Hasil Tokenizing

Text	Tokenizing
saya mengucapkan terima kasih kepada seluruh personel polri dimanapun bertugas tetaplah semangat dalam memberikan pelayanan terbaik kepada masyarakat sehingga kita bisa menjadi polisi yang tegas humanis dekat dan dicintai masyarakat	['saya', 'mengucapkan', 'terima', 'kasih', 'kepada', 'seluruh', 'personel', 'polri', 'dimanapun', 'bertugas', 'tetaplah', 'semangat', 'dalam', 'memberikan', 'pelayanan', 'terbaik', 'kepada', 'masyarakat', 'sehingga', 'kita', 'bisa', 'menjadi', 'polisi', 'yang', 'tegas', 'humanis', 'dekat', 'dan', 'dicintai', 'masyarakat']

Tabel 4: Stopwords

Tokenizing	Stopwords
['saya', 'mengucapkan', 'terima', 'kasih', 'kepada', 'seluruh', 'personel', 'polri', 'dimanapun', 'bertugas', 'tetaplah', 'semangat', 'dalam', 'memberikan', 'pelayanan', 'terbaik', 'kepada', 'masyarakat', 'sehingga', 'kita', 'bisa', 'menjadi', 'polisi', 'yang', 'tegas', 'humanis', 'dekat', 'dan', 'dicintai', 'masyarakat']	['terima', 'kasih', 'personel', 'polri', 'dimanapun', 'bertugas', 'tetaplah', 'semangat', 'pelayanan', 'terbaik', 'masyarakat', 'polisi', 'humanis', 'dicintai', 'masyarakat']

Tabel 5: Stemming

Tokenizing	Stopwords
['saya', 'mengucapkan', 'terima', 'kasih', 'kepada', 'seluruh', 'personel', 'polri', 'dimanapun', 'bertugas', 'tetaplah', 'semangat', 'dalam', 'memberikan', 'pelayanan', 'terbaik', 'kepada', 'masyarakat', 'sehingga', 'kita', 'bisa', 'menjadi', 'polisi', 'yang', 'tegas', 'humanis', 'dekat', 'dan', 'dicintai', 'masyarakat']	['terima', 'kasih', 'personel', 'polri', 'mana', 'tugas', 'tetap', 'semangat', 'layan', 'baik', 'masyarakat', 'polisi', 'humanis', 'cinta', 'masyarakat']

3.2 Text Transformation

Pengklasifikasian (kelas) sentimen didasarkan atas kamus InSet Lexicon dimana didalam kamus tersebut terdapat kata serta bobot dari setiap kata tersebut. Kemudian perhitungannya dengan menjumlahkan bobot atau score dari setiap kata yang telah melewati tahap preprocessing. Sehingga besar atau kecilnya score tersebut ditentukan oleh setiap kata yang ada dalam dokumen.

1. Sentimen Positif

Pengklasifikasian kedalam sentimen positif ditandai dengan total score yang lebih dari sama dengan ($\geq$) 0. Contoh perhitungan score untuk sentimen positif dapat dilihat pada Tabel berikut:

Tabel 6: Perhitungan Score untuk Sentimen Positif

Tweet (raw)	Saya mengucapkan terima kasih kepada seluruh personel Polri dimanapun bertugas, tetaplah semangat dalam memberikan pelayanan terbaik kepada masyarakat, sehingga kita bisa menjadi Polisi yang tegas, humanis, dekat dan dicintai masyarakat. #Polri
preprocessing	terima kasih personel polri mana tugas tetap semangat layan baik masyarakat polisi humanis cinta masyarakat
Pembobotan	Terima = 2; kasih = 2; semangat = 4, -1; layan = 2, -2; baik = 3, -1; cinta = 3, -1; tugas = -1;
Total Score	$score = 2 + 2 + 3 + 0 + 2 + 2 + (-1) = 12$

2. Sentimen Negatif

Pengklasifikasian ke dalam sentimen negatif ditandai dengan total score yang kurang dari 0. Contoh perhitungan untuk sentimen negatif dapat dilihat pada tabel di bawah ini.

Tabel 7: Contoh Perhitungan Score untuk Sentimen Negatif

Tweet (raw)	Polri boleh buat iklan tapi masyarakat butuhkan bukan iklan saja tapi kenyataan sekedar saran jangan dikit dikit di bilang oknum aku saja klo ada kesalahan dan perbaiki dengan maksimal bukan sekedar iklan babat habis namanya oknum termasuk atasannya
preprocessing	polri iklan masyarakat butuh iklan nyata saran dikit dikit bilang oknum aku klo salah baik maksimal dar iklan babat habis nama oknum atas

Pembobotan	Nyata = 4, -2; saran = 2, -3; baik = 3; maksimal = 3; habis = 3, -3; butuh = -5; dikit (2) = -4; salah = -4; baik = -1; atas = -4
Total Score	$score = (4) + (-2) + (2) + (-3) + (3) + (3) + (3) + (-3) + (-5) + (-8) + (-4) + (-1) + (-4) = -15$

Metode pembobotan yang digunakan yaitu TF-IDF dengan bantuan library sklearn yaitu Tfidfvectorizer. Metode ini berfungsi untuk mencari representasi nilai dari kumpulan data. Berikut contoh data yang dihitung dengan TF-IDF sebagai berikut:

Tabel 8: Dokumen/Tweet

Tweet	Text
A	terima kasih personel polri mana tugas tetap semangat layan baik masyarakat polisi humanis cinta masyarakat
B	polri iklan masyarakat butuh iklan nyata saran dikit dikit bilang oknum aku klo salah baik maksimal dar iklan babat habis nama oknum atas
C	biar polri laku masyarakat peran awal

Selanjutnya adalah tahapan pembobotan pada setiap kata dimana menghitung jumlah kemunculan atau disebut dengan term frequency pada setiap data yang dapat dilihat pada gambar di bawah ini.

Term	TF			DF	IDF = $\log((N+1)/(df+1))+1$			Bobot (w) = $tf * idf$		
	A	B	C		A	B	C	A	B	C
terima	0,07	0,00	0,00	1	1,699	1,699	1,699	0,113	0,000	0,000
kasih	0,07	0,00	0,00	1	1,699	1,699	1,699	0,113	0,000	0,000
personel	0,07	0,00	0,00	1	1,699	1,699	1,699	0,113	0,000	0,000
polri	0,07	0,04	0,17	3	1,637	1,637	1,637	0,109	0,071	0,273
tugas	0,07	0,00	0,00	1	1,699	1,699	1,699	0,113	0,000	0,000
semangat	0,07	0,00	0,00	1	1,699	1,699	1,699	0,113	0,000	0,000
layan	0,07	0,00	0,00	1	1,699	1,699	1,699	0,113	0,000	0,000
masyarakat	0,07	0,04	0,17	3	1,637	1,637	1,637	0,109	0,071	0,273
polisi	0,07	0,00	0,00	1	1,699	1,699	1,699	0,113	0,000	0,000
humanis	0,07	0,00	0,00	1	1,699	1,699	1,699	0,113	0,000	0,000
cinta	0,07	0,00	0,00	1	1,699	1,699	1,699	0,113	0,000	0,000
iklan	0,00	0,04	0,00	1	1,699	1,699	1,699	0,000	0,074	0,000
butuh	0,00	0,04	0,00	1	1,699	1,699	1,699	0,000	0,074	0,000
nyata	0,00	0,04	0,00	1	1,699	1,699	1,699	0,000	0,074	0,000
saran	0,00	0,04	0,00	1	1,699	1,699	1,699	0,000	0,074	0,000
dikit	0,00	0,04	0,00	1	1,699	1,699	1,699	0,000	0,074	0,000
oknum	0,00	0,04	0,00	1	1,699	1,699	1,699	0,000	0,074	0,000
salah	0,00	0,04	0,00	1	1,699	1,699	1,699	0,000	0,074	0,000
baik	0,07	0,04	0,00	2	1,653	1,653	1,653	0,110	0,072	0,000
maksimal	0,00	0,04	0,00	1	1,699	1,699	1,699	0,000	0,074	0,000
babat	0,00	0,04	0,00	1	1,699	1,699	1,699	0,000	0,074	0,000
habis	0,00	0,04	0,00	1	1,699	1,699	1,699	0,000	0,074	0,000
biar	0,00	0,00	0,17	1	1,699	1,699	1,699	0,000	0,000	0,283
laku	0,00	0,00	0,17	1	1,699	1,699	1,699	0,000	0,000	0,283
peran	0,00	0,00	0,17	1	1,699	1,699	1,699	0,000	0,000	0,283
awal	0,00	0,00	0,17	1	1,699	1,699	1,699	0,000	0,000	0,283

Gambar 2: Sample Perhitungan Pembobotan TF-IDF

3.3 Pattern Discovery

Model yang digunakan pada tahapan ini adalah Multinomial Naive Bayes (MNB) yang digunakan untuk pengklasifikasian dan pengujian data yang digunakan dalam hal ini adalah *data training* dan *data testing*. Pengujian ini dimaksudkan untuk menilai performa dari model yang telah dibuat.

3.3.1 Model Multinomial Naive Bayes

1. Data Training MNB

Sebuah dokumen training sudah diklasifikasikan secara manual dan sudah dilakukan proses preprocessing seperti pada tabel di bawah ini.

Tabel 9: Contoh kasus data training

Dokumen	Text	Polarity
Tweet A	terima kasih personel polri mana tugas tetap semangat layan baik masyarakat polisi humanis cinta masyarakat	P
Tweet B	polri iklan masyarakat butuh iklan nyata saran dikit dikit bilang oknum aku klo salah baik maksimal dari iklan babat habis nama oknum atas	N
Tweet C	biar polri laku masyarakat peran awal	P
...	...	...

Dari tabel di atas, dibuat sebuah model prior probabilities untuk menemukan probabilitas dari setiap polarity atau kelas yang ada, dengan menghitung total kemunculan kelas yang dibagi dengan total keseluruhan data.

Tabel 10: Hasil perhitungan Prior Probabilities

Polarity/Kelas	P(class)
Positif (P)	$4/10 = 0,4$
Negatif (N)	$6/10 = 0,6$

Setiap term yang ditemukan akan dihitung probabilitas kemunculannya untuk setiap kelas, kemudian dibagi dengan total term yang ditemukan. Tabel 11 merupakan hasil model perhitungan klasifikasinya. Nilai probabilitas yang dicari adalah nilai yang terbesar dan dijadikan sebagai kelas term tersebut. Sebagai contoh untuk term "masyarakat", dimana kelas negatif menunjukkan 0,0298 sedangkan kelas positif 0,0615. Maka didapatkan term "masyarakat" akan diklasifikasikan pada kelas positif karena lebih besar dibandingkan nilai kelas negatif.

Tabel 11: Model Perhitungan Klasifikasi

Term	Negatif	Positif
terima	$1/23 = 0,0298$	$1/21 = 0,0307$
polri	$2/23 = 0,0447$	$2/21 = 0,0461$
semangat	$0/23 = 0,0149$	$1/21 = 0,0307$
masyarakat	$1/23 = 0,0298$	$3/21 = 0,0615$
baik	$1/23 = 0,0298$	$1/21 = 0,0307$
...	...	...

2. Data Testing MNB

Perbedaan dari proses training dan testing tidak jauh berbeda, hanya ketika proses telah selesai, data testing

akan dihitung dengan nilai probabilitas akhir. Sebagai contoh, data testing seperti terdapat pada tabel di bawah ini.

Tabel 12: Contoh kasus data testing

Dokumen	Text	Polarity
Tweet A	Polri berikan rasa aman cinta warga binaan sambangi masyarakat	?
Tweet B	selagi fadlizon tidak cocok berarti polri udah di rel yang benar	?

Pada tabel 12, terdapat dua dokumen yang perlu ditentukan kelasnya. Nilai setiap term yang ada pada tiap dokumen akan memiliki nilai probabilitasnya masing-masing, apabila term tersebut sudah melalui proses training sebelumnya. Sehingga pada data uji term yang dihitung hanya term yang telah melalui tahapan training. Berikut adalah contoh perhitungan probabilitas penentuan kelas pada setiap dokumen.

Tabel 13: Perhitungan probabilitas kelas

Tweet A	$P(\text{tweet A} \text{positif})$ $= P(\text{positif}) * P(\text{polri}) * P(\text{cinta}) * P(\text{masyarakat})$ $= \frac{4}{10} * 0.0461 * 0.0307 * 0.0615$ $= \mathbf{0.00003481}$
	$P(\text{tweet A} \text{negatif})$ $= P(\text{negatif}) * P(\text{polisi}) * P(\text{cinta}) * P(\text{masyarakat})$ $= \frac{6}{10} * 0.0461 * 0.0307 * 0.0615$ $= \mathbf{0.00005223}$
Tweet B	$P(\text{tweet B} \text{positif})$ $= P(\text{positif}) * P(\text{polri})$ $= \frac{4}{10} * 0.0461$ $= \mathbf{0.01844}$
	$P(\text{tweet B} \text{negatif})$ $= P(\text{negatif}) * P(\text{polri})$ $= \frac{6}{10} * 0.0461$ $= \mathbf{0.02766}$

Pada hasil perhitungan, dapat ditentukan bahwa nilai probabilitas terbesar pada Tweet A adalah sebesar 0,00005223, sehingga nilai tersebut akan membuat Tweet A diklasifikasikan ke dalam kelas Negatif. Sedangkan pada Tweet B, nilai probabilitas terbesar adalah 0,02766,

sehingga Tweet B akan diklasifikasikan ke dalam kelas Negatif.

3.3.2 Model Long Short-Term Memory

1. Preprocessing LSTM

Tahapan ini digunakan untuk mempersiapkan dan mengubah data inputan sehingga sesuai dengan kebutuhan model dan memaksimalkan performa model.

1.1 Tokenizer

Tahapan ini akan mengubah format teks menjadi token-token yang dipisahkan oleh spasi. Berikut di bawah ini adalah hasil tokenizer:

Tabel 14: Hasil Tokenizer LSTM

Dokumen	Text
Tweet A	['terima'], ['kasih'], ['personel'], ['polri'], ['tugas'], ['semangat'], ['layan'], ['baik'], ['masyarakat'], ['humanis'], ['cinta'], ['masyarakat']
Tweet B	['polri'], ['iklan'], ['masyarakat'], ['butuh'], ['iklan'], ['nyata'], ['saran'], ['dikit'], ['dikit'], ['oknum'], ['salah'], ['baik'], ['maksimal'], ['iklan'], ['habis'], ['oknum']
Tweet C	['biar'], ['polri'], ['laku'], ['masyarakat'], ['peran'], ['awal']
...	...

1.2 Dictionary

Setelah tahap tokenizing selesai, selanjutnya adalah merubah hasil tersebut ke dalam format khusus kamus data atau *dictionary* (key & value). Setiap kata diberi nomor indeks yang sesuai dengan kamus data. Berikut di bawah ini adalah tabel hasil dictionary.

Tabel 15: Hasil Dictionary

{`terima`:1, `kasih`:2, `personel`:3, `polri`:4, `tugas`:5, `semangat`:6, `layan`:7, `baik`:8, `masyarakat`:9, `humanis`:10, `cinta`:11, `iklan`:12, `butuh`:13, `nyata`:14, `saran`:15, `dikit`:16, `oknum`:17, `salah`:18, `baik`:19, `maksimal`:20, `habis`:21, `biar`:22, `laku`:23, `peran`:24, `awal`:25}

1.3 Sequence

Pada tahap ini objek hasil *Tokenizer* digunakan untuk mengubah teks menjadi urutan angka berdasarkan kamus kata yang telah dibangun. Metode ini mengambil teks sebagai input dan mengonversinya menjadi urutan angka yang mewakili kata-kata dalam teks tersebut sesuai dengan kamus yang telah dibuat sebelumnya. Urutan angka ini kemudian dapat diberikan sebagai input kepada jaringan saraf. Berikut di bawah ini adalah tabel hasil sequence.

Tabel 16: Hasil sequence tokenizer LSTM

Text	Sequences
['terima'], ['kasih'], ['personel'], ['polri'], ['tugas'], ['semangat'], ['layan'], ['baik'], ['masyarakat'], ['humanis'], ['cinta'], ['masyarakat']	[1, 2, 3, 4, 5, 6, 7, 8, 9, 10, 11, 9]
['polri'], ['iklan'], ['masyarakat'], ['butuh'], ['iklan'], ['nyata'], ['saran'], ['dikit'], ['dikit'], ['oknum'], ['salah'], ['baik'], ['maksimal'], ['iklan'], ['habis'], ['oknum']	[4, 12, 9, 13, 12, 14, 15, 16, 16, 17, 18, 19, 20, 12, 21, 17]
['biar'], ['polri'], ['laku'], ['masyarakat'], ['peran'], ['awal']	[22, 4, 23, 9, 24, 25]

1.4 Pad Sequence

Pad Sequences digunakan untuk mengisi urutan angka (didepan atau dibelakang) dengan padding atau mengatur panjang agar seragam dengan maksimum panjang yang telah ditentukan sebelumnya.

Tabel 17: Hasil Pad Sequence

Text	Pad Sequence (maxlen=8)
terima kasih personel polri mana tugas tetap semangat layan baik masyarakat polisi humanis cinta masyarakat	[1 2 3 4 5 6 7 8]
polri iklan masyarakat butuh iklan nyata saran dikit dikit bilang oknum aku klo salah baik maksimal dar iklan babat habis nama oknum atas	[4 12 9 13 12 14 15 16]
biar polri laku masyarakat peran awal	[0 0 0 22 4 23 9 24]

Pada hasil diatas terdapat **maxlen=8**, yang berarti maksimum panjang dari setiap teks yang ada adalah sebanyak 8, sehingga teks yang kurang dari 8 akan ditambahkan *padding* (angka 0) agar teks tersebut memiliki ukuran panjang yang sama.

2. Pengujian Model

Dalam pengujian atau pengukuran performa model klasifikasi terdapat beberapa cara, tetapi yang sering atau umum digunakan adalah dengan menghitung nilai parameter pada metode *confusion matrix*. Metode *confusion matrix* akan memberi tahu seberapa baik model yang dibuat dengan memperhitungkan keempat parameter serta *building blocks* yang digunakan. Contoh perhitungan pengujian model dengan asumsi telah melewati proses *data training* dan *testing*:

Tabel 18: Pengujian Model

Predicted Values	Actual Values	
	Positive	Negative
Positive	True Positive (TP) TP = 65	False Positive (FP) FP = 7
Negative	False Negative (FN) FN = 4	True Negative (TN) TN = 38

Setelah mendapatkan nilai masing-masing dari *True Positive*, *False Positive*, *False Negative*, dan *True Negative*, maka *performance matrices* dapat langsung dihitung.

2.1 accuracy

$$accuracy = \frac{65 + 38}{65 + 38 + 7 + 4} = 0.90$$

Dari perhitungan accuracy, maka teks yang diprediksi positif ataupun negatif dari total keseluruhan teks yang ada adalah sebesar 90%.

2.2 precision

$$recision = \frac{65}{65 + 7} = 0.90$$

Dari perhitungan precision, maka teks yang diklasifikasikan positif dari keseluruhan teks yang diprediksi adalah sebesar 90%.

2.3 recall

$$recall = \frac{65}{65 + 4} = 0.94$$

Dari perhitungan accuracy, maka teks yang diprediksi positif dibandingkan dengan teks yang diklasifikasikan positif adalah sebesar 94%.

2.4 fl-score

$$accuracy = 2 * \frac{0.90 * 0.94}{0.90 + 0.94} = 0.92$$

Dari perhitungan accuracy, maka teks yang benar diprediksi positif maupun negatif dari total keseluruhan teks yang ada adalah sebesar 92%.

3.4 Evaluation

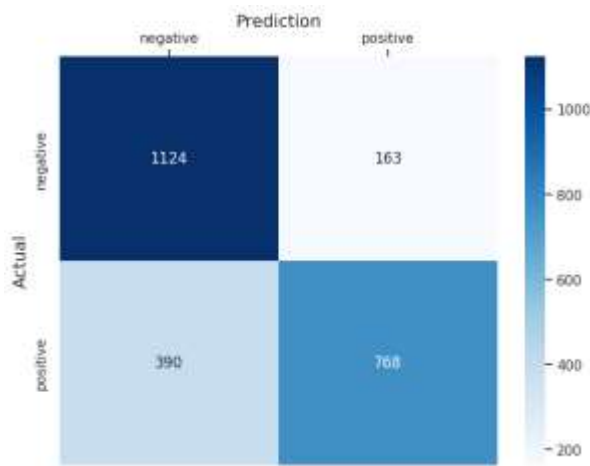
Hasil akurasi merupakan hasil yang diperoleh saat menggunakan model MNB dan LSTM, didalamnya terdapat nilai akurasi yang didapat dari masing-masing model, serta *confusion matrix* yang terbentuk.

3.4.1 Akurasi Multinomial Naive Bayes

Accuracy: 0.7738241308793457				
Train Accuracy: 0.8132351437046129				
	precision	recall	f1-score	support
0	0.74	0.87	0.80	1287
1	0.82	0.66	0.74	1158
accuracy			0.77	2445
macro avg	0.78	0.77	0.77	2445
weighted avg	0.78	0.77	0.77	2445

Gambar 3: Akurasi model Multinomial Naive Bayes

Berdasarkan pada Gambar 3, hasil akurasi model dengan modum MNB didapat sebesar 78% untuk akurasi, 78,5% untuk presisi, 77% untuk recall, dan 77% untuk f1-score. Pada confusion matrix model MNB, dapat dilihat pada Gambar 4. Hasil dari confusion matrix terdapat actual negative dengan predictive negatif sebesar 1.124 yang berarti jumlah data yang bersentimen negative sebesar 1.124, diprediksi benar negative (true negative).



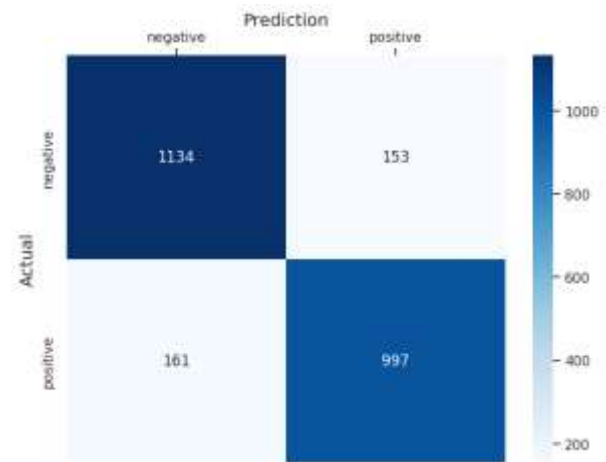
Gambar 4: Confusion Matrix model MNB

3.4.2 Akurasi LSTM

Akurasi model LSTM didapat sebesar 87% untuk akurasi, 87% untuk presisi, 87% untuk recall, dan 87 untuk f1-score seperti tercantum pada gambar di bawah 5. Untuk confusion matrix dapat dilihat pada Gambar 6, dengan actual positive dengan predictive positive sebesar 997.

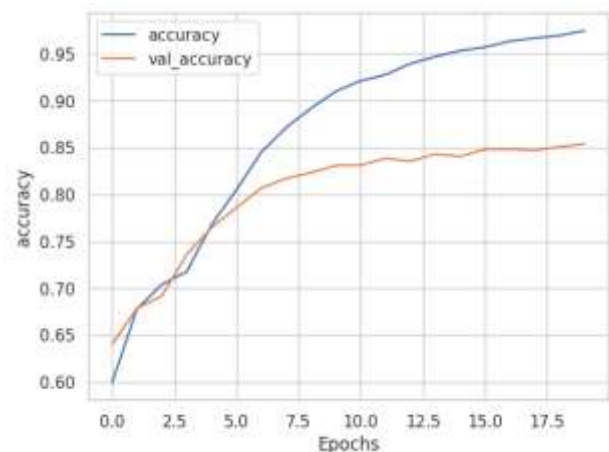
Test loss: 0.3548421561717987				
Test accuracy: 0.8715746402740479				
	precision	recall	f1-score	support
0	0.88	0.88	0.88	1287
1	0.87	0.86	0.86	1158
accuracy			0.87	2445
macro avg	0.87	0.87	0.87	2445
weighted avg	0.87	0.87	0.87	2445

Gambar 5: Akurasi Model LSTM



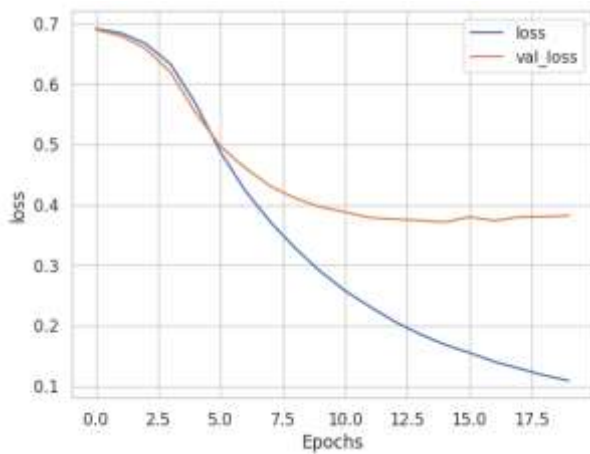
Gambar 6: Confusion Matrix model LSTM

Grafik pada gambar di bawah ini menggambarkan proses train dan test terhadap nilai akurasi. Grafik tersebut memiliki model yang cukup bagus ditandai dengan ketika epoch bertambah, nilai akurasi dengan val_accuracy berjalan tidak terlalu jauh.



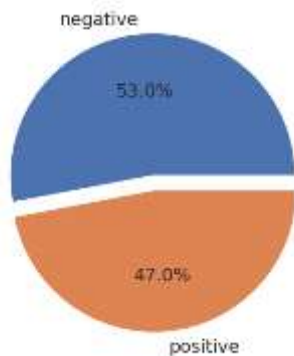
Gambar 7: Grafik nilai akurasi model LSTM

Pada grafik selanjutnya yang ditampilkan pada Gambar 9, menggambarkan proses train dan test terhadap nilai loss. Grafik tersebut memiliki model yang cukup bagus ditandai dengan ketika epoch bertambah, nilai loss dengan val_loss berjalan tidak terlalu jauh. Tetapi masih dapat diperbaiki agar menghasilkan nilai val_loss yang lebih kecil.



Gambar 8: Grafik nilai loss model LSTM

Berdasarkan perbandingan sentimen dari total 12.222 tweet, setelah dilakukan preprocessing, terdapat 53% bersentimen negatif dan 47% bersentimen positif yang berarti terdapat 6.478 data bersentimen negatif dan 5.744 data bersentimen positif. Perbandingan sentimen dapat dilihat pada gambar di bawah ini beserta dengan wordcloud tentang banyaknya kata yang muncul.



Gambar 9: Perbandingan Sentimen



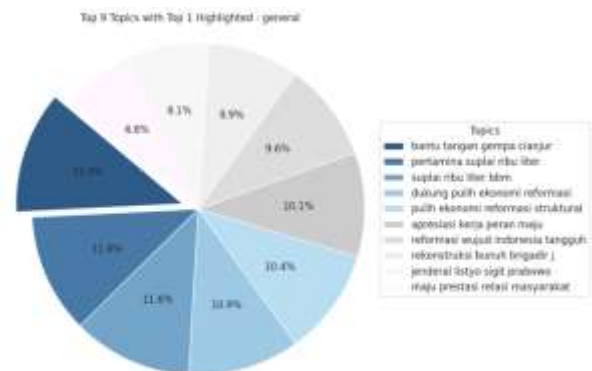
Gambar 10: Wordcloud untuk keseluruhan data



Gambar 11: Wordcloud untuk data bersentimen positif



Gambar 12: Wordcloud data bersentimen negatif



Gambar 13: Persentase kemunculan topik keseluruhan

IV. KESIMPULAN

Pada penelitian ini, berdasarkan pada hasil proses *crawling*, didapatkan data dengan jumlah 14.000 data. Data didapat dari aplikasi *microblogging* Twitter dengan menggunakan hashtag #polri. Data tersebut dibagi menjadi tiga jenis data, yaitu 80% untuk *data training*, 20% untuk *data testing*, dan 15 untuk *data validation*. Penggunaan Metode Multinomial Naïve Bayes dan Long Short-Term Memory digunakan untuk melakukan perbandingan tingkat akurasi dari kedua metode tersebut. Sedangkan Lexicon-Based digunakan untuk melakukan optimalisasi pada setiap proses tahapan text mining. Metode pengujian menggunakan model Multinomial Naive Bayes yang mampu menghasilkan nilai akurasi sebesar 78%, precision sebesar 78.5%, recall sebesar 77%, dan f1-score sebesar 77%. Dengan model Long Short-Term Memory mampu menghasilkan accuracy, precision, recall, dan f1-score sebesar 87%. Hasil analisa sentimen masyarakat terhadap kinerja Kepolisian diperoleh berdasarkan hasil interpretasi dan evaluasi akhir dengan persentase sebesar 53% atau 6.478 data memiliki sentimen negatif dan 47% atau 5.744 data bersentimen positif.

V. REFERENSI

- [1] Simon Kemp, "DIGITAL 2022: INDONESIA." Accessed: May 11, 2023. [Online]. Available: <https://datareportal.com/reports/digital-2022-indonesia>
- [2] Monavia Ayu Rizaty, "Pengguna Twitter di Indonesia Capai 18,45 Juta Pada Tahun 2022." Accessed: May 11, 2023. [Online]. Available: <https://dataindonesia.id/internet/detail/pengguna-twitter-di-indonesia-capai-1845-juta-pada-2022>
- [3] M. Hidayatullah, S. Alam, and I. Jaelani, "Sentiment Analysis of Police Performance On Twitter Users Using Naïve Bayes Method," vol. 2, no. 2, p. 2021, Accessed: May 11, 2024. [Online]. Available: <https://journal.institutpendidikan.ac.id/index.php/ristec/article/view/96>
- [4] Lembaga Survey Indonesia, "RILIS LSI 20 OKTOBER 2022." Accessed: May 11, 2023. [Online]. Available: <https://www.lsi.or.id/post/rilis-survei-lsi-20-oktober-2022>
- [5] F. Agung, J. Ayomi, and K. E. Dewi, "ANALISIS EMOSI PADA MEDIA SOSIAL TWITTER MENGGUNAKAN METODE MULTINOMIAL NAÏVE BAYES DAN SYNTHETIC MINORITY OVERSAMPLING TECHNIQUE," *KOMPUTA: Jurnal Ilmiah Komputer dan Informatika*, vol. 12, no. 2, 2023, [Online]. Available: <https://ojs.unikom.ac.id/index.php/komputa/article/view/9454>
- [6] A. Deviyanto, M. R. Didik Wahyudi, and T. Informatika UIN Sunan Kalijaga Yogyakarta Jl Marsda Adi Sucipto No, "PENERAPAN ANALISIS SENTIMEN PADA PENGGUNA TWITTER MENGGUNAKAN METODE K-NEAREST NEIGHBOR," *Jurnal Informatika Sunan Kalijaga*, vol. 3, no. 1, pp. 1–13, 2018, [Online]. Available: <https://ejournal.uin-suka.ac.id/saintek/JISKA/article/view/31-01>
- [7] M. R. Faisal, D. Kartini, and T. H. Saragih, *Belajar Data Science: Text Mining Untuk Pemula I*. 2022. [Online]. Available: <https://www.researchgate.net/publication/359619425>
- [8] A. Sabrani, I. W. Gede Putu Wirarama Wedashwara, and F. Bimantoro, "METODE MULTINOMIAL NAÏVE BAYES UNTUK KLASIFIKASI ARTIKEL ONLINE TENTANG GEMPA DI INDONESIA (Multinomial Naïve Bayes Method for Classification of Online Article About Earthquake in Indonesia)." [Online]. Available: <http://jtika.if.unram.ac.id/index.php/JTIKA/article/view/87>
- [9] A. M. Kibriya, E. Frank, B. Pfahringer, and G. Holmes, "Multinomial Naive Bayes for Text Categorization Revisited," 2004. [Online]. Available: <http://kdd.ics.uci.edu/>
- [10] Yuyun, Nurul Hidayah, and Supriadi Sahibu, "Algoritma Multinomial Naïve Bayes Untuk Klasifikasi Sentimen Pemerintah Terhadap Penanganan Covid-19 Menggunakan Data Twitter," *Jurnal RESTI (Rekayasa Sistem dan Teknologi Informasi)*, vol. 5, no. 4, pp. 820–826, Aug. 2021, doi: 10.29207/resti.v5i4.3146.
- [11] D. Sintia Amelia, N. Cahyana Aminuallah, and S. Informasi, "TEKS DAN ANALISIS SENTIMEN PADA CHAT GRUP WHATSAPP MENGGUNAKAN LONG SHORT TERM MEMORY (LSTM)," *Jurnal Teknologi Terkini*, vol. 3, no. 2, p. 1, 2023, [Online]. Available: <http://teknologiterkini.org/index.php/terkini/article/view/354>
- [12] M. Z. Rahman, Y. A. Sari, and N. Yudistira, "Analisis Sentimen Tweet COVID-19 menggunakan Word Embedding dan Metode Long Short-Term Memory (LSTM)," vol. 5, no. 11, pp. 5120–5127, 2021, [Online]. Available: <https://j-ptiik.ub.ac.id/index.php/j-ptiik/article/view/10188>
- [13] L. Wiranda and M. Sadikin, "PENERAPAN LONG SHORT TERM MEMORY PADA DATA TIME SERIES UNTUK MEMPREDIKSI PENJUALAN PRODUK PT. METISKA FARMA," *Janapati*, vol. 8, no. 3, 2019, [Online]. Available: <https://ejournal.undiksha.ac.id/index.php/janapati/article/view/19139>
- [14] R. Arief and K. Imanuel, "ANALISIS SENTIMEN TOPIK VIRAL DESA PENARI PADA MEDIA SOSIAL TWITTER DENGAN METODE LEXICON BASED," *Jurnal Ilmiah MATRIK*, vol. 21, no. 3, 2019.
- [15] R. Mahendrajaya, G. A. Buntoro, and M. B. Setyawan, "ANALISIS SENTIMEN PENGGUNA GOPAY MENGGUNAKAN METODE LEXICON BASED DAN SUPPORT VECTOR MACHINE," *Komputek*, vol. 3, no. 2, 2019, [Online]. Available: <https://studentjournal.umpo.ac.id/index.php/komputek/article/view/270>
- [16] A. K. Fauziyyah and D. H. Gautama, "ANALISIS SENTIMEN PANDEMI COVID19 PADA STREAMING TWITTER DENGAN TEXT MINING PYTHON," *Jurnal Ilmiah SINUS*, vol. 18, no. 2, p. 31, Jul. 2020, doi: 10.30646/sinus.v18i2.491.
- [17] G. Singh, B. Kumar, L. Gaur, and A. Tyagi, "Comparison between Multinomial and Bernoulli Naïve Bayes for Text Classification," in *International Conference on Automation, Computational and Technology Management (ICACTM)*, 2019. [Online]. Available: <https://ieeexplore.ieee.org/abstract/document/8776800>