

ANALISIS SENTIMEN OPINI PUBLIK TERHADAP KEBIJAKAN BARU SKRIPSI PADA MEDIA SOSIAL TWITTER MENGGUNAKAN METODE NAÏVE BAYES

Refandah Puspitasari¹, Aries Dwi Indriyanti²

^{1,2} Jurusan Teknik Informatika/Program Studi Sistem Informasi, Universitas Negeri Surabaya

¹refandah.20004@mhs.unesa.ac.id

²ariesdwi@unesa.ac.id

Abstrak— Transformasi kebijakan pada Pendidikan Tinggi Indonesia, fokusnya pada perubahan persyaratan penyelesaian tugas akhir untuk mahasiswa program sarjana. Sebelumnya, skripsi merupakan kewajiban untuk gelar sarjana, namun Kebijakan Menteri Pendidikan, Kebudayaan, Riset, dan Teknologi Republik Indonesia Nomor 53 Tahun 2023 telah mengubahnya. Kebijakan tersebut memicu berbagai pendapat yang pro dan kontra di Masyarakat, karena tidak lagi memerintahkan skripsi sebagai tugas akhir, melainkan memperbolehkan alternatif seperti prototipe atau proyek. Respons masyarakat terhadap perubahan ini tercermin dalam unggahan di media sosial, khususnya platform Twitter (yang saat ini disebut X), yang merupakan wadah populer untuk berbagi pendapat. Tujuan dari penelitian ini yaitu melakukan analisis sentimen pada opini publik terhadap kebijakan baru skripsi pada media sosial twitter menggunakan metode *Naïve Bayes*. Dataset yang digunakan adalah hasil crawling pada media sosial twitter dengan memasukan keyword yang berhubungan dengan kebijakan baru skripsi. Dataset berjumlah 470 tweets dengan label positif berjumlah 135 data, label negatif berjumlah 191 data, dan label netral berjumlah 144 data. Dapat disimpulkan bahwa Masyarakat cenderung beropini negatif yang artinya kurang setuju dengan adanya kebijakan baru skripsi. Berdasarkan pengujian yang dilakukan dengan metode *naïve bayes* dengan porsi perbandingan data latih dan data uji yang berbeda yaitu dengan rasio 90%:10% dan 80%:20%, dari pembagian rasio tersebut sama sama memiliki akurasi terbaik sebesar 76%. Perbedaan dari pembagian rasio tersebut hanya pada Precision yang memiliki nilai lebih tinggi pada rasio 90:10 yaitu 77%.

Kata Kunci— Skripsi, Twitter, Analisis Sentimen, Naïve Bayes, Crawling

I. PENDAHULUAN

Pendidikan tinggi adalah tahap krusial yang berperan dalam memajukan kemampuan akademik dan profesional seseorang. Tahap ini bertujuan untuk peserta didik mempersiapkan diri agar mereka dapat menjadi individu yang mampu mengaplikasikan, memajukan dan mewujudkan ilmu pengetahuan, teknologi, dan seni dengan baik. Menurut Peraturan Menteri Pendidikan dan Kebudayaan Republik Indonesia Nomor 3 Tahun 2020 tentang Standar Nasional Pendidikan Tinggi, mahasiswa menyelesaikan program sarjana diwajibkan untuk menyusun skripsi atau laporan tugas akhir serta mengunggahnya ke situs web perguruan tinggi.

Saat ini Pendidikan tinggi mengalami transformasi penting melalui implementasi kebijakan – kebijakan baru yang bertujuan untuk mengatasi berbagai tantangan serta memajukan kualitas dan relevansi Pendidikan tinggi. Salah satu perubahan signifikan dalam kebijakan baru dalam Kemendikbudristek Republik Indonesia Nomor 53 Tahun 2023 tentang penjaminan mutu Pendidikan tinggi adalah bahwa mahasiswa program sarjana tidak lagi diwajibkan untuk menyelesaikan skripsi. Berdasarkan kebijakan tersebut, tugas akhir bagi mahasiswa program sarjana dapat berupa prototipe, proyek, atau bentuk lainnya yang bisa dikerjakan secara individu atau dalam kelompok.

Bentuk tanggapan maupun opini pro dan kontra disampaikan Masyarakat melalui unggahan pada media sosial twitter yang saat ini berubah nama menjadi X. Oleh karena itu, pengumpulan data memanfaatkan media sosial twitter atau X pada opini Masyarakat terhadap ditetapkannya kebijakan baru yang tertuang dalam Peraturan Kemendikbudristek Nomor 53 Tahun 2023 tentang penjaminan mutu Pendidikan tinggi. Dengan adanya tanggapan masyarakat yang bersifat positif, negatif, dan netral, dilakukan penelitian dalam menganalisis sentimen Masyarakat untuk memahami tanggapan dan opini masyarakat terhadap kebijakan baru yang telah ditetapkan. Analisis sentimen merupakan identifikasi pendapat atau kecenderungan opini seseorang pada suatu masalah atau objek, dalam mengetahui pandangan atau opini tersebut cenderung negatif atau positif [1].

Penelitian ini akan menggunakan *metode naïve bayes* dalam melakukan analisis sentimen. Algoritma ini umum digunakan dalam analisis sentimen karena dikenal memiliki tingkat akurasi yang tinggi, terutama dalam mengklasifikasikan teks. Metode *naïve bayes* dikenal dengan algoritma dengan kesederhanaanya dalam melakukan pengklasifikasian, meskipun algoritma *naïve bayes* memiliki kesederhanaan dalam pengklasifikasian, namun tingkat akurasi yang dihasilkannya cukup tinggi [2].

Berdasarkan penelitian terdahulu yang pernah melakukan analisis sentimen, diantaranya yaitu Rahmad harun, R. Ishak dan S. S. menganalisis sentimen opini publik pengguna Twitter terkait kenaikan harga BBM dengan menggunakan algoritma *Naïve Bayes*, hasil dari penelitian tersebut mendapatkan akurasi sebesar 85% [3]. Penelitian lain yaitu Penerapan pembobotan TF-IDF dengan metode *Naïve Bayes* untuk menganalisis sentimen masyarakat terkait isu kenaikan BIPIH, hasil dari penelitian tersebut mendapatkan akurasi

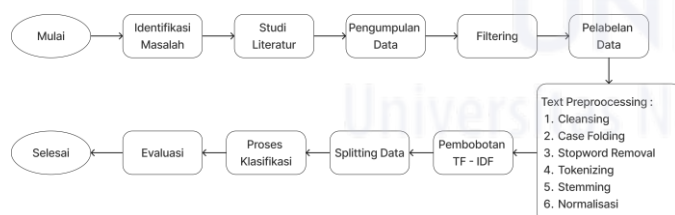
sebesar 89% [4]. Penelitian terdahulu berikutnya menggunakan metode *Naïve Bayes*, *Decision Tree*, dan *Random Forest* dalam menganalisis studi kasus kampanye anti LGBT, dari penelitian tersebut menghasilkan bahwa *naïve bayes* memperoleh akurasi tertinggi yaitu 83,43%, dari metode *decision tree* dan *random forest* yang memiliki tingkat akurasi yang cukup rendah [5].

Berdasarkan penelitian terdahulu yang telah dijabarkan, dapat dibuktikan bahwa metode *naïve bayes* memberikan tingkat akurasi terbaik. Oleh karena itu untuk melakukan analisis sentimen opini Masyarakat terhadap kebijakan baru pada tugas akhir untuk mahasiswa pada media sosial twitter menggunakan metode *naïve bayes*. Dari analisis ini menghasilkan output yang diharapkan dapat memberikan informasi kepada Masyarakat, program studi maupun perguruan tinggi mengenai tanggapan Masyarakat yang diunggah pada media sosial twitter tentang kebijakan baru yang telah ditetapkan, apakah cenderung pada sentimen positif, negatif ataupun netral, serta dapat menjadi bahan pertimbangan oleh perguruan tinggi maupun prodi untuk merancang standar kelulusan mahasiswa pada kebijakan baru ini.

II. METODOLOGI PENELITIAN

Metode yang digunakan dalam penelitian dengan metode penelitian kuantitatif, yang merupakan pendekatan penelitian yang objektif yang melibatkan pengumpulan dan analisis data menggunakan teknik statistik (Hermawan, 2005). Penelitian akan berfokus pada opini Masyarakat terhadap ditetapkan kebijakan baru yang tertuang dalam Peraturan Menteri Pendidikan, Kebudayaan, Riset dan Teknologi Republik Indonesia Nomor 53 Tahun 2023 tentang penjaminan mutu Pendidikan tinggi.

Adapun tahapan pada penelitian ini terdapat beberapa yang harus dilalui untuk mencapai tujuan penelitian. Alur penelitian berupa alur pengerjaan yang membahas mengenai gambaran umum yang dilakukan peneliti dalam tahap pengerjaan dari awal hingga akhir.



Gbr. 1 Alur Penelitian

A. Identifikasi Masalah

Tahapan awal akan mengidentifikasi masalah pada sebuah permasalahan mengenai topik yang dipilih oleh peneliti yaitu terkait opini publik terhadap kebijakan baru skripsi yang telah ditetapkan dengan menggunakan metode *naïve bayes*. Dengan menganalisis sentimen akan memanfaatkan metode *naïve bayes* dalam mengetahui opini maupun pandangan Masyarakat terhadap kebijakan baru ini dan juga dapat

menjadi bahan pertimbangan universitas maupun kaprodi dalam merancang standar kelulusan mahasiswa.

B. Studi Literatur

Studi literatur melibatkan pembelajaran dari literatur yang relevan dengan konsep dan metode guna untuk memecahkan masalah yang ada dalam penelitian ini. Sumber literatur yang digunakan mencakup jurnal, tesis, dan referensi terkait yang relevan dengan masalah penelitian.

C. Pengumpulan Data

Penelitian ini mengumpulkan data dari opini Masyarakat yang diunggah pada media sosial twitter, dengan memanfaatkan Teknik *Crawling*. *Crawling* dilakukan menggunakan python dengan tools Google Colab. Data yang digunakan diambil dari memasukan *keyword* pencarian yang berhubungan dengan kebijakan baru skripsi. *Keyword* yang digunakan adalah “hapus skripsi”, “kebijakan tugas akhir”, “tidak wajib skripsi” dan “skripsi dihapus”. *Crawling* dilakukan dari bulan Agustus 2023 sampai dengan Oktober 2023

D. Filtering

Filtering merupakan proses untuk menghapus tweet-tweet yang duplikat menjadi satu. Hal ini diperlukan karena dalam data yang diperoleh dari proses *crawling* sering kali terdapat *duplicate* akibat dari retweet atau repost tweet. Dengan menerapkan *filtering*, data tweet dapat dibersihkan dari duplikasi sehingga tidak ada tweet yang berulang. Pada tahap *filtering* juga menghapus data yang *unrelated*, terdapat data tweet yang tidak sesuai dengan topik yang dicari untuk penelitian ini. *Filtering* akan dilakukan secara manual.

E. Labeling

Setelah didapatkan data dari *crawling* twitter dan dilakukan *filtering*, tahap selanjutnya adalah pelabelan data dengan label komentar untuk tweet opini Masyarakat terhadap kebijakan baru skripsi menggunakan label positif, negatif dan netral. Melabelkan data dilakukan dengan cara manual untuk membuat data latih dan melakukan penandaan pada data uji untuk mengukur keakuratan dari hasil sentimen yang dihasilkan secara otomatis.

F. Text Preprocessing

Tahap *text preprocessing* bertujuan untuk memproses data mentah agar siap digunakan. Tahap *text preprocessing* meliputi langkah-langkah berikut:

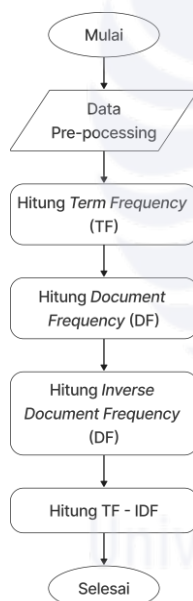
1. Tahap *cleansing* adalah tahap membersihkan data pada dokumen dari kata-kata yang tidak dibutuhkan untuk meminimalisir gangguan atau noise, seperti hashtag, *mention*, angka serta tanda baca agar text menjadi lebih bersih
2. Tahap *case folding* membantu menyederhanakan teks dan menjadi seragam serta konsisten dalam penggunaan huruf besar dan kecil, sehingga

mempermudah proses pencocokan dan perbandingan teks.

3. Tahap *Stopword removal* adalah tahap yang akan menghapus kata-kata yang tidak memberikan pengaruh yang signifikan, seperti kata penghubung, kata depan, dan kata sambung.
4. Tahap *tokenizing* adalah tahap mengubah kalimat menjadi kata-kata. Tahap *tokenizing* yaitu dilakukan pemotongan terhadap kalimat yang akan diubah menjadi kata terpisah pada dataset.
5. Tahap *stemming* adalah proses mengidentifikasi dan mengubah kata yang memiliki imbuhan menjadi bentuk dasar.
6. Tahap normalisasi adalah proses memperbaiki ejaan yang salah atau terdapat kata yang tidak baku agar menjadi bentuk baku dengan kamus normalisasi

G. Pembobotan TF-IDF

Pada TF-IDF digunakan menilai pentingnya kata-kata dalam sebuah dokumen. *Term frequency* menghitung sering munculnya kata tersebut dalam dokumen tertentu dan *inverse document frequency* mengukur sebaran kata tersebut di seluruh koleksi dokumen untuk menyesuaikan bobotnya. Gbr. 2 adalah alur dari TF-IDF.



Gbr. 2 Alur Pembobotan TF-IDF

Berikut adalah rumus dari TF-IDF :

$$TF - IDF = TF \times IDF$$

Keterangan :

TF : Frekuensi dari *term* dan dokumen

IDF : Inversi Frekuensi Dokumen dari *Term*

Untuk menghitung nilai *Inverse Document Frequency* (IDF) dalam proses pembobotan dengan rumus berikut

$$IDF = \log\left(\frac{D}{DF}\right)$$

Keterangan :

D = Jumlah dokumen

DF = Dokumen frekuensi

H. Splitting Data

Pada pengujian metode *naïve bayes* akan dilakukan splitting data. Data akan dipisahkan menjadi data latih untuk melatih model dan data uji untuk mengukur kinerja model.

I. Klasifikasi

Proses klasifikasi dilakukan dengan metode *naïve bayes*, di mana penentuan sentimen dilakukan dengan menghitung probabilitas dokumen data uji berdasarkan data latih

J. Evaluasi

Evaluasi pada penelitian ini menggunakan *confusion matrix* untuk mengevaluasi hasil uji dan mengukur tingkat akurasi dari proses klasifikasi yang dilakukan oleh sistem. Selain itu, dalam implementasinya *confusion matrix* digunakan dalam perhitungan *accuracy*, *precision*, *recall*, dan *F1-score*.

III. HASIL DAN PEMBAHASAN

Pada penelitian analisis sentimen terhadap kebijakan baru skripsi dilakukan menggunakan metode *naïve bayes*. Penelitian ini bertujuan untuk mengidentifikasi sentimen masyarakat yang diunggah di media sosial twitter terkait kebijakan baru yang telah ditetapkan, apakah cenderung pada sentimen positif, negatif ataupun netral, serta dapat menjadi bahan pertimbangan oleh perguruan tinggi maupun prodi untuk merancang standar kelulusan mahasiswa pada kebijakan baru ini. Dataset yang akan dianalisis merupakan opini masyarakat yang diunggah di media sosial twitter, dikumpulkan menggunakan teknik *crawling*. Proses *crawling* dilakukan dengan python menggunakan tools Google Colab.

A. Hasil Pengumpulan Data

Dataset penelitian menggunakan hasil dari *crawling* dengan memasukkan kata *keyword* terkait dengan kebijakan baru skripsi. *Keyword* yang digunakan adalah “hapus skripsi”, “kebijakan tugas akhir”, “tidak wajib skripsi” dan “skripsi dihapus”. *Crawling* dilakukan dari bulan Agustus 2023 sampai dengan Oktober 2023 dengan hasil yang telah dikumpulkan sejumlah 1883.

B. Filtering

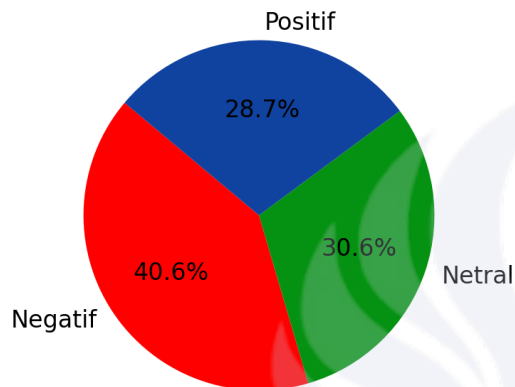
Setelah dataset terkumpul tahap selanjutnya adalah tahap filtering. Sebelum data di berikan label, data akan melalui tahap filtering yaitu menghapus data yang *unrelated* dan menghapus tweet yang duplikat atau ganda sehingga hanya tersisa satu. Data *unrelated* yang dimaksud adalah data yang tidak sesuai dengan topik penelitian ini. Pada awal *crawling* diperoleh jumlah data sebanyak 1883 tweets. Setelah

dilakukan filtering menjadi 470 tweets yang siap dilakukan labelling.

C. Labelling

Dataset yang terkumpul akan dilakukan proses pelabelan, di mana setiap tweet akan diberi label untuk mengklasifikasikannya sebagai sentimen positif, negatif, atau netral.

Proses *labelling* dilakukan secara manual. Hasil dataset yang sudah diberikan label, diperoleh label positif berjumlah 135 data, label negatif berjumlah 191 data, dan label netral berjumlah 144 data. Pada Gbr.3 adalah visualisasi hasil dari labelling, diperoleh 40,6% sentimen negatif, 28,7% sentimen positif dan 30,6% sentimen netral.



Gbr. 3 Persentase Hasil Labelling

D. Hasil Text Preprocessing

Setelah data sudah diberi label akan melalui tahap *text preprocessing*, yang merupakan proses data mentah akan diubah menjadi data yang siap digunakan.

TABEL I
HASIL TEXT PREPROCESSING

Tweet	
Skripsi dihapus kurang setuju, karena semua pelajaran saat kuliah cuman materi skripsi saja yg paling mengerti wkwk	
Text Preprocessing	Hasil
Cleansing	Skripsi dihapus kurang setuju karena semua pelajaran saat kuliah cuman materi skripsi saja yg paling mengerti wkwk
Case Folding	skripsi dihapus kurang setuju karena semua pelajaran saat kuliah hanya materi skripsi saja yang paling mengerti wkwk
Stopword Removal	skripsi dihapus setuju pelajaran kuliah materi skripsi mengerti
Tokenizing	['skripsi', 'dihapus', 'setuju', 'pelajaran', 'kuliah', 'materi', 'skripsi', 'mengerti']
Stemming	['skripsi', 'hapus', 'tuju', 'ajar', 'kuliah', 'materi', 'skripsi', 'erti']

Normalisasi	skripsi hapus setuju ajar kuliah materi skripsi arti
-------------	--

E. Hasil Pengujian

Proses klasifikasi dengan metode *naïve bayes*, dilakukan pembagian rasio 2 kali yaitu rasio 90:10 dan 80:10. pembagian rasio untuk membandingkan kinerja model pada 2 skenario pembagian rasio yang berbeda. Tabel II adalah pembagian rasio yang telah dilakukan.

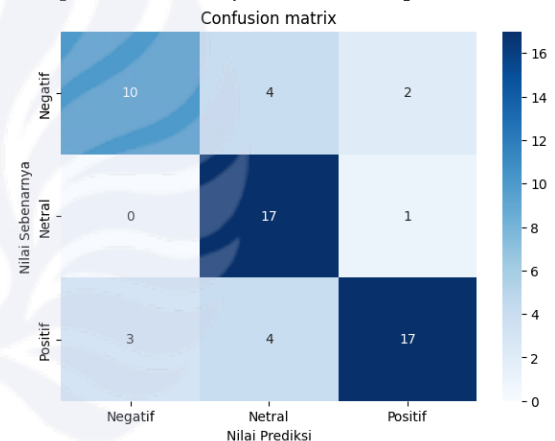
TABEL II
PEMBAGIAN RASIO

Rasio	Parameter
90 : 10	test_size = 0,1
80 : 20	test_size = 0,2

Hasil percobaan dilakukan dengan membagi data uji dan data latih dalam rasio sebagai berikut :

1. Rasio 90 : 10

Setelah dilakukan percobaan dengan pembagian rasio 90 : 10 didapatkan hasil *confusion matrix* seperti Gbr. 4.



Gbr. 4 Hasil Confusion Matrix Rasio 90:10

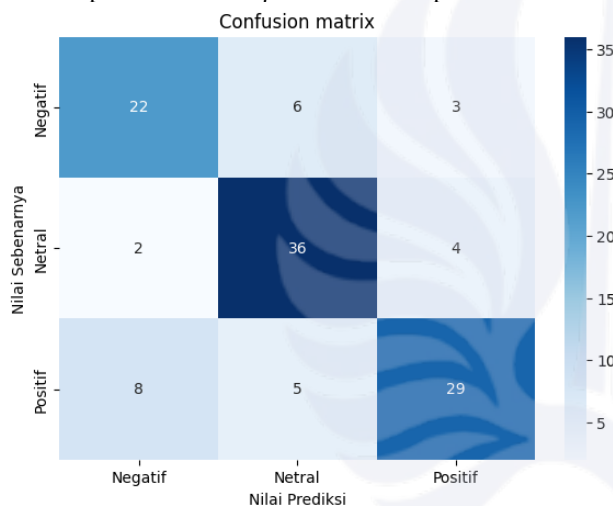
Hasil implementasi perhitungan akurasi, presisi, recall, dan nilai F1-score.

Accuracy:	0.7586206896551724			
Precision:	0.7749602122015915			
Recall:	0.7586206896551724			
F1 Score:	0.7553878524760365			
	<u>Precision</u>	<u>recall</u>	<u>f1-score</u>	<u>support</u>
Negatif	0.77	0.71	0.69	16
Netral	0.68	0.86	0.79	18
Positif	0.85	0.69	0.77	24
accuracy			0.76	58
macro avg	0.77	0.76	0.75	58
weighted avg	0.77	0.76	0.76	58

Gbr. 5 Hasil Perhitungan Confusion Matrix Rasio 90:10

2. Rasio 80 :20

Setelah dilakukan percobaan dengan pembagian rasio 90 : 10 didapatkan hasil *confusion matrix* seperti Gbr. 6.



Gbr. 6 Hasil Confusion Matrix Rasio 80:20

Hasil implementasi perhitungan akurasi, presisi, recall, dan nilai F1-score.

Accuracy:	0.7565217391304347			
Precision:	0.7592699660807894			
Recall:	0.7565217391304347			
F1 Score:	0.7552964432300044			
	<u>precision</u>	<u>recall</u>	<u>f1-score</u>	<u>support</u>
Negatif	0.69	0.71	0.70	31
Netral	0.77	0.86	0.81	42
Positif	0.81	0.69	0.74	42
accuracy			0.76	115
macro avg	0.75	0.75	0.75	115
weighted avg	0.76	0.76	0.76	115

Gbr. 7 Hasil Perhitungan Confusion Matrix Rasio 90:10

Hasil perhitungan akurasi, presisi, recall, dan nilai F1-score dari metode *naïve bayes* dari 2 kali pengujian perbandingan rasio.

TABEL III

HASIL PERHITUNGAN ACCURACY, PRECISSION, RECALL DAN F1-SCORE

Rasio	Hasil Evaluasi			
	<i>Accuracy</i>	<i>Precision</i>	<i>Recall</i>	<i>F1-Score</i>
90 : 10	76%	77%	76%	75%
80 : 20	76%	75%	76%	75%

Dari Tabel III menunjukan bahwa pembagian rasio 90 : 10 dengan rasio 80 : 20 sama sama memiliki akurasi terbaik sebesar 76%. Perbedaan dari pembagian rasio tersebut hanya pada *Precision* yang memiliki nilai lebih tinggi pada rasio 90:10 yaitu 77%. Hasil dari percobaan 2 pembagian rasio pada metode *naïve bayes* telah diambil rata-ratanya. Dalam percobaan ini menghasilkan akurasi mencapai 76%.

IV. KESIMPULAN

Hasil dari penelitian, dapat ditarik kesimpulan yang menunjukan hasil sebagai berikut :

1. Proses klasifikasi dimulai dengan memberikan label secara manual, kemudian dilanjutkan dengan tahap preprocessing untuk memproses teks tweet. Setelah itu, dilakukan ekstraksi fitur untuk menghitung frekuensi kemunculan kata menggunakan metode TF-IDF, yang membantu menentukan tingkat pentingnya kata dalam dokumen. Setelah itu, Model menggunakan klasifikasi *naïve bayes* untuk mengkategorikan dataset dengan label positif, negatif, dan netral. Selanjutnya, kinerja model yang telah dibuat dievaluasi dengan menggunakan *confusion matrix*. Setelah dilakukan dua kali percobaan pengujian dengan metode *naïve bayes* dengan perbedaan porsi perbandingan data *training* dan data *testing* yaitu sebesar 90%:10% dan 80%:20%, didapatkan hasil bahwa kedua rasio tersebut memiliki tingkat akurasi yang sama yaitu 76%. Perbedaan dari kedua pembagian rasio hanya pada *precision* yang memiliki nilai lebih tinggi pada rasio 90:10 yaitu 77%. Dengan tingkat akurasi tersebut, metode *naïve Bayes* menunjukkan kinerja yang baik dalam menganalisis sentimen opini publik terhadap kebijakan baru skripsi di media sosial Twitter.
2. Berdasarkan penelitian yang telah dilakukan pada pengujian Analisis sentimen opini publik terhadap kebijakan baru skripsi pada media sosial twitter menggunakan metode *naïve bayes* menggunakan dataset sejumlah 470 data diperoleh 40,6% sentimen negatif, 28,7% sentimen positif dan 30,6% sentimen netral. Dalam penelitian ini didapatkan bahwa Masyarakat

cenderung beropini negatif yang artinya kurang setuju dengan adanya kebijakan baru skripsi.

V. SARAN

Hasil dari penelitian, penulis memberikan saran untuk mengembangkan penelitian dalam masa yang akan datang sebagai berikut:

1. Dalam penelitian selanjutnya, diharapkan para peneliti menggunakan berbagai sumber data dari platform media sosial lainnya, seperti Facebook, Instagram, YouTube, TikTok, dan sebagainya. Tujuannya adalah untuk memperoleh variasi data yang berbeda sehingga hasilnya bisa dibandingkan dengan penelitian sebelumnya. Selain itu, para peneliti juga diharapkan dapat mengembangkan dataset dengan topik yang beragam.
2. Pada penelitian berikutnya, disarankan untuk mempertimbangkan penggunaan dua atau lebih metode atau algoritma yang berbeda. Hal ini bertujuan untuk membandingkan dan menentukan metode atau algoritma

yang paling efektif atau optimal dalam menyelesaikan masalah yang serupa.

REFERENSI

- [1] Nurzahputra, A., & Muslim, A. Analisis Sentimen pada Opini Mahasiswa Menggunakan Natural Language Processing. In Seminar Nasional Ilmu Komputer, 2016
- [2] Handayani, F., Feddy, D., & Pribadi, S. (n.d.). Implementasi Algoritma Naive Bayes Classifier dalam Pengklasifikasian Teks Otomatis Pengaduan dan Pelaporan Masyarakat melalui Layanan Call Center 110.
- [3] Harun, R., Ishak, R., & Panna, S. S. Analisis Sentimen Opini Publik Pengguna Twitter Terhadap Kenaikan Harga BBM Menggunakan Algoritma Naive Bayes. Copyright @BALOK, 2(1), 2023
- [4] Wati, R., Ernawati, S., & Rachmi, H. Pembobotan TF-IDF Menggunakan Naive Bayes pada Sentimen Masyarakat Mengenai Isu Kenaikan BIPIH. Jurnal Manajemen Informatika (JAMIKA), 13(1), 84–93. <https://doi.org/10.34010/jamika.v13i1.9424>, 2023
- [5] Fitri, V. A., Andreswari, R., & Hasibuan, M. A. Sentiment analysis of social media Twitter with case of Anti-LGBT campaign in Indonesia using Naive Bayes, decision tree, and random forest algorithm. Procedia Computer Science, 161, 765–772. <https://doi.org/10.1016/j.procs.2019.11.181>, 2019.

