

SENTIMENT ANALISIS REVIEW APLIKASI RUANG GURU PADA GOOGLE PLAYSTORE DENGAN METODE NAÏVE BAYES

Christoper Andrian, Yo Ceng Giap

Teknik Informatika, Universitas Buddhi Dharma
Jl. Imam Bonjol, Tangerang, Banten, Indonesia
christoper11a@gmail.com

ABSTRAK

Platform pendidikan online, seperti Ruang Guru, telah mendorong pengguna untuk berbagi pengalaman mereka melalui ulasan di platform seperti Google Play Store. Namun, ketidaksesuaian antara peringkat bintang dan ulasan tertulis sering kali menyamarkan perasaan sebenarnya dari pengguna. Masalah ini menjadi tantangan dalam memahami umpan balik pengguna secara menyeluruh, yang dapat memberikan wawasan lebih dalam tentang kepuasan dan kebutuhan mereka terkait aplikasi tersebut. Studi ini menggunakan analisis sentimen untuk mengevaluasi umpan balik pengguna pada aplikasi Ruang Guru, dengan memanfaatkan metode klasifikasi Naïve Bayes untuk mengkategorikan ulasan menjadi sentimen positif, netral, dan negatif. Penelitian ini melibatkan pengumpulan 10.000 ulasan dari Google Play Store, diikuti dengan penerapan teknik pra-pemrosesan data seperti tokenisasi dan penghapusan kata stop sebelum klasifikasi. Tiga pembagian data pelatihan-pengujian yang berbeda (70:30, 80:20, dan 90:10) dievaluasi untuk menentukan konfigurasi yang paling optimal. Pembagian 90:10 menghasilkan kinerja terbaik, dengan akurasi 83,52%, presisi 71,44%, recall 65,84%, dan skor F1 67,83%. Analisis menunjukkan dominasi sentimen positif dalam ulasan. Metode ini memberikan wawasan berharga bagi pengembang, memungkinkan mereka untuk menangani kekhawatiran pengguna dan meningkatkan fungsionalitas aplikasi. Selain itu, studi ini menyoroti potensi signifikan analisis sentimen dalam meningkatkan pengalaman pengguna aplikasi teknologi pendidikan dan memberikan rekomendasi untuk menangani data yang tidak seimbang guna memperbaiki efektivitas model.

Kata kunci : Analisis Sentimen , Aplikasi Ruang Guru, Google Playstore, Text Mining

1. PENDAHULUAN

Perkembangan teknologi yang pesat di bidang pendidikan telah mendorong dukungan yang luas terhadap strategi pembelajaran secara daring, karena media daring menyediakan kemudahan, akses ke informasi terbaru, dan kemudahan akses. Saat ini, berbagai situs web pembelajaran dan aplikasi daring terus berkembang pesat di Indonesia, mulai dari aplikasi gratis hingga berbayar, salah satunya adalah aplikasi Ruang Guru. [1]. Dan juga pembelajaran online ini menjadi sebuah alternatif dimanah pada masa era pandemi covid-19, yang dimanah aktivitas kegiatan pembelajaran tatap muka yang secara langsung dihentikan sementara waktu dan dilakukan kegiatan pembelajaran online secara daring yang dilakukan di rumah, sebagai bentuk untuk pencegahan penyebaran covid-19[2].

Ruang Guru adalah perusahaan teknologi pendidikan terbesar di Indonesia yang menyediakan layanan pembelajaran online melalui aplikasi Ruang Guru. Ruang Guru memiliki 22 juta pengguna aplikasi dan 300.000 guru mentor yang menyediakan layanan di lebih dari Pendidikan dasar hingga pendidikan tingkat lanjut. Selain itu, Ruang Guru juga menawarkan materi pembelajaran tambahan seperti platform media kelas virtual, sistem simulasi ujian online, diskusi konten video berbasis pelanggan, dan pasar layanan yang disesuaikan. [3]. Aplikasi ruang guru dapat di akses oleh pengguna melalui *google playstore* maupun di *App IOS*. Dengan segala fitur baru yang diperkenalkan untuk aplikasi pendidikan ini.

Seiring dengan perkembangan aplikasi Ruang guru orang-orang sedang membahas kelebihan dan kekurangan aplikasi Ruang Guru dalam ulasan di Google Play Store. Setiap ulasan di Google Play Store memiliki rating dari 1 hingga 5, namun rating yang diberikan oleh pengguna sering kali tidak sesuai dengan ulasan tentang aplikasi Ruang Guru. Oleh karena itu, hal ini tidak cukup untuk menggambarkan kelebihan dan kekurangan aplikasi Ruang Guru. Setiap perangkat lunak atau aplikasi *mobile* memiliki kelebihan dan keterbatasan, yang dapat menyebabkan ulasan yang bervariasi dari pengguna, mulai dari kepuasan hingga ketidakpuasan terhadap aplikasi.[4]

Tujuan penelitian ini meinformasi yang diperoleh dari analisis sentimen dengan teknik *Teks Mining* dapat memberikan umpan balik berharga bagi pengembang untuk memprioritaskan peningkatan dan menangani keluhan pengguna, sehingga meningkatkan fungsionalitas aplikasi dan pengalaman pengguna. Analisis semacam ini sangat penting karena ulasan pengguna memberikan wawasan detail tentang kegunaan aplikasi dan berfungsi sebagai acuan bagi pengembang untuk mengevaluasi produk mereka [5].

Pada penelitian ini, untuk memperoleh mengenai opini pengguna terhadap aplikasi ruang guru telah diterapkan metode analisa yang mengekstrak informasi penting dari Data tidak terstruktur atau ulasan pengguna aplikasi ruang guru, sehingga penelitian ini dapat mengidentifikasi sentimen pengguna Google Play Store terhadap aplikasi Ruang

Guru, serta metode pemrosesan data yang digunakan dalam analisis sentimen pada studi ini dengan teknik *Teks Mining* [6]. Analisis ini proses yang akan menentukan apakah opini yang disampaikan pada ulasan komentar *review* atau sebuah pada teks kalimat bersifat positif, negatif atau netral, yang membantu dalam mengambil keputusan, oleh karena itu hal ini mencerminkan tren yang berguna mengenai sikap dan emosi[7]. Pada analisis sentimen pada teknik *Teks Mining* dengan klasifikasinya dibantu dengan metode algoritma *Naïve Bayes* pada proses mengklasifikasi komentar *review*. Algoritma *Naïve Bayes* merupakan algoritma klasifikasi berdasarkan dari probabilitas dimanah algoritma ini mempelajari dari data lampau untuk memprediksikan data baru. Dan klasifikasi *Naïve Bayes* merupakan teknik *machine learning* untuk mengklasifikasi pada teks, serta model yang sederhana dan efisien pada model klasifikasinya [8].

2. TINJAUAN PUSTAKA

2.1. Penelitian Terdahulu

Pada penelitian terdahulu yang sebagai acuan di penelitian ini, untuk melakukan proses analisis klasifikasi yang dilakukan oleh Giovani, dkk, pada penelitian ini menggunakan tiga algoritma klasifikasi teks, yaitu *Naïve Bayes*, *Support Vector Machine* (SVM), dan *K-Nearest Neighbor* (KNN) sebagai analisa model klasifikasi sentimen dan *dataset* pada label data untuk analisa model tidak diketahui dalam pembahasan penelitian ini, dan hasil dari penelitian pengujian 3 model klasifikasi tersebut model SVM dengan nilai akurasi 78,55%, KNN nilai akurasinya 77,21% dan *Naïve Bayes* nilai akurasi 67,32% [3].

Pada penelitian selanjutnya yang dilakukan oleh penelitian Syafrizal, A, dkk berdasarkan penelitiannya menggunakan dua algoritma klasifikasi teks, yaitu *Naïve Bayes Classifier* (NBC) dan *K-Nearest Neighbor* (KNN) sebagai perbandingannya analisa model klasifikasi sentimen dan *dataset* pada label data menggunakan pelabelan secara manual yang bantuan ahli pakar bahasa dalam pelabelan datanya untuk analisa model. Hasil dari penelitian pengujian 2 model klasifikasi tersebut model *Naïve Bayes Classifier* (NBC) memiliki performa lebih baik dibandingkan KNN dengan akurasi 77,69% untuk model NBC dan KNN nilai akurasi 76,4% [9].

Penelitian selanjutnya yang dilakukan oleh Taufiqurrahman, T., dkk berdasarkan penelitiannya bertujuan menganalisis dan membandingkan performa dua algoritma klasifikasi teks, yakni *Naïve Bayes* (NB) dan *K-Nearest Neighbor* (KNN) sebagai klasifikasi sentimen dan *dataset* pada label datanya menggunakan pelabelan *vader* untuk pelabelan data untuk digunakan analisa model. Hasil dari penelitian perbandingan 2 model klasifikasi tersebut model *Naïve Bayes* (NB) memiliki performa lebih baik dibandingkan KNN dengan nilai akurasi 75% dan KNN nilai akurasi 73% [10].

Berdasarkan penelitian sebelumnya model klasifikasi *Naïve Bayes* (NB) menunjukkan memiliki model yang cukup baik dalam klasifikasi sentimen dan berdasarkan pembahasan yang telah didapatkan pada penelitian terdahulu maka, penulis menerapkan algoritma klasifikasi *Naïve Bayes* dalam melakukan analisis sentimen.

2.2. Text Mining

Text Mining adalah penambangan dan pengolahan data yang bertujuan untuk mencari nilai informasi yang tersembunyi baik dari sumber data informasi semi-struktur ataupun informasi data yang tidak struktur [11].

Text mining juga dikenal sebagai teknik yang bertujuan untuk mengidentifikasi pola, tren, dan hubungan yang ada dalam pada data teks, pada tekniknya menggunakan eksplorasi dan analisis yang pada hasilnya yang dapat digunakan untuk pengambilan keputusan yang lebih baik[12].

2.3. Klasifikasi

Klasifikasi merupakan teknik yang digunakan untuk proses mengelompokkan data dengan mengacu pada karakteristik tertentu, pada proses pengelompokan di dalam *data mining* ataupun *text mining* [13].

2.4. Naïve Bayes

Naïve Bayes adalah algoritma yang sering dipakai dengan teknik data *mining* dan juga merupakan teori *Bayes*, algoritma *Naïve Bayes* ini juga sebagai strategi *classifier* probabilitas yang dapat menentukan dengan pertambahan suatu nilai iterasi dari sebuah indeks informasi yang sudah ada dan diberikan, dengan mengasumsikan seluruh atributnya sehingga mereka tidak saling ketergantungan dari nilai yang memiliki variabel kelas [14].

Rumus umum pada *naïve bayes* dirumuskan sebagai berikut :

$$P(H|X) = \frac{P(X|H)P(H)}{P(X)} \quad (1)$$

Keterangan :

X : Data dengan kelas yang belum diketahui.

H : Hipotesis data pada suatu kelas spesifik.

$P(H|X)$: Probabilitas hipotesis dari kondisi (*posteriori probability*).

$P(H)$: Probabilitas hipotesis awalan kelas (*prior probability*).

$P(X|H)$: Probabilitas hipotesis kondisi berdasarkan hipotesis kelasnya(probabilitas *likelihood*).

$P(X)$: Probabilitas keseluruhan data X (probabilitas *evidence*).

Pada kaitan antara *Naïve Bayes* dengan klasifikasi terdapat dua tahapan, di mana untuk

tahapan pertama dilakukan yaitu, proses tahapan pembelajaran pada sebuah teks dokumen yang kategorinya sudah diketahui. Pada tahapan kedua yaitu proses pengujian pada klasifikasi teks dokumen yang kategorinya tidak diketahui. Pada formulasinya *Naïve Bayes* untuk klasifikasi algoritma akan mencari probabilitas tinggi data dokumen dari semua kategori yang diuji (*Vmap*). Dengan diterapkan *teorema Bayes* pada persamaan berikut :

$$V_{MAP} = \underset{V_j \in V}{\operatorname{argmax}} \prod_{i=1}^n P(X_i|V_j)P(V_j) \quad (2)$$

Keterangan :

V_j = Kategori ulasan atau dokumen, $j = 1, 2, 3, \dots, n$.
diartikan sebagai

j_1 = Kategori sentimen negatif,

j_2 = Kategori sentimen positif, dan

j_3 = Kategori sentimen netral.

$P(x_i|V_j)$ = Probabilitas x_i pada kategori.

$V_j P(V_j)$ = Probabilitas dari V_j

Untuk $P(v_j)$ dan $P(x_i|V_j)$ pada perhitungan di saat pelatihan dirumuskan pada persamaan sebagai berikut:

$$P(V_j) = \frac{|\text{docs } j|}{|\text{docs } n|} \quad (3)$$

$$P(X_i|V_j) = \frac{nk+1}{n+|\text{kosakata}|} \quad (4)$$

Keterangan :

$|\text{docs } j|$ = jumlah ulasan atau dokumen setiap kategori j .

$|\text{docs } n|$ = jumlah ulasan atau dokumen dari semua kategori.

nk = jumlah frekuensi dari setiap kemunculan kata.

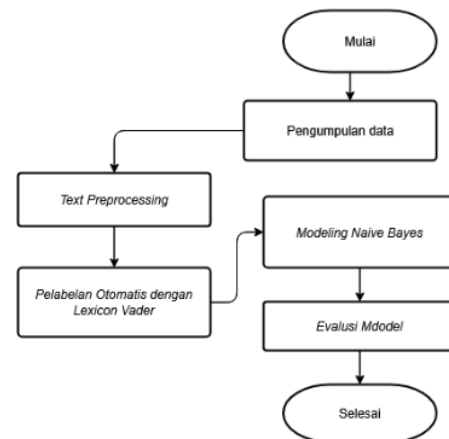
n = jumlah frekuensi setiap kategori dari kemunculan kata.

2.5. Confusion Matrix

Confusion Matrix adalah suatu metode yang dipakai untuk mengetahui ukuran atau melakukan perhitungan akurasi pada konsep data mining, *confusion matrix* berfungsi untuk memperoleh nilai *precision*, *recall*, dan *accuracy*. Hasil nilai *confusion matrix* biasanya dilambangkan dalam satuan persen (%) [5].

3. METODE PENELITIAN

Penelitian yang dilakukan untuk mengklasifikasi ini menerapkan pada proses tahapan penelitian dan akan melalui beberapa tahapan yang di tunjukan sebagai berikut :



Gambar 1. Diagram Tahapan Penelitian

3.1. Pengumpulan data

Tahapan ini dalam pengumpulan data menggunakan *library python* yaitu web scraper dengan menggunakan *url*, dimanah memudahkan untuk pengambilan data teks komentar dari *google Playstore* pada aplikasi ruang guru. Data yang diambil sebanyak 10.000 *dataset*.

3.2. Textpreprocessing

Text preprocessing merupakan suatu tahapan sebelum *dataset* di masukan sebelum ke dalam model, yang dimanah melakukan proses *text preprocessing* dalam data teks untuk menjadi data bersih sebelum dilakukan proses lebih lanjut [12]. Proses pembersihan data meliputi, yaitu : (1) *transform case* untuk sebuah mengubah semua data teks menjadi huruf kecil semua untuk guna menyesuaikan semua data seragam, (2) *remove emoji & karakter simbol* dimanah menghapus semua emoji dan karakter simbol yang terdapat pada komentar *dataset*, (3) *Tokenizing* merupakan sebuah proses yang membagi sebuah teks menjadi perkata – kata. (4) *remove stopword* untuk menghilangkan kata – kata yang tidak memiliki makna , (5) *Stemming* bertujuan untuk mengubah kata – kata dalam teks menjadi bentuk kata dasar [15]. Hasil dari *text preprocessing* dari 10.000 *dataset* menjadi data bersih sebanyak 9.700.

Tabel 1. Proses *Text Processing*

Data sebelum	Data sesudah
Aplikasinya bagus 🍌	['aplikasi', 'bagus']
Bisa gampang buat memahami materi yang disekolah	['gampang', 'memahami', 'materi', 'sekolah']

3.3. Pelabelan otomatis *lexicon vader*

Lexicon Vader suatu model yang diciptakan oleh C.J.Hutto dan Eric Gibert pada tahun 2014 yang dimanah untuk memudahkan untuk melabeli data secara otomatis, *Vader* merupakan pendekatan leksikal yang dipakai untuk sebagai model analisis intensitas

emosi dan analisa suasana hati dan juga dari keunikannya dimanah tidak mempersiapkan data latih untuk menggunakannya [16]. Fitur *Vader* ini memiliki skor positif, negatif atau netral dalam bentuk *compound*. *Compound* merupakan matriks yang untuk menghitung semua skor yang dinormalisasi dari -1 hingga +1. Pada penelitian ini menggunakan tiga buah label yaitu positif yang ditentukan dari nilai *compound* yang lebih dari 0, label negatif pada hasil nilai kurang dari 0 dan netral didapatkan dari nilai 0 yang dihasilkan *compound* [17]. Pada tahapan labeling otomatis ini data teks yang ingin kita label dengan *lexicon vader* ini harus melakukan transformasi dengan mengubah data teks bahas menjadi bahasa inggris terlebih dahulu, karena *lexicon vader* mendukung kumpulan data kamus yang hanya mendukung penggunaan bahasa inggris. Hasil dari pelabelan otomatis *Lexicon Vader* pada *dataset* 9.700 menghasilkan data pada 3 label sebanyak yaitu positif 6.726 data, 2.214 data berlabel netral dan sebanyak 760 data berlabel negatif.

3.4. Modeling naïve bayes

Pada Pengujian model pada *naive bayes* menggunakan pengujian pembagian *dataset* dengan dibaginya menjadi data *train* dan data uji, pada pembagian data *train* dan data uji ini membagi data banding dengan tiga rasio pembagian data yaitu 70:30, 80:20 dan 90:10. Rasio pembagian data ini untuk mengetahui pelatihan model *naive bayes*, pada mengukur evaluasi terbaik model. Pembagian dan hasil dari pembagian data tersebut akan ditampilkan pada tabel 2.

Tabel 2. Pembagian data *train* dan data uji

Data rasio	Data <i>train</i>	Data uji
70:30	6789	2911
80:20	7759	1941
90:10	8729	971

Pengujian model menggunakan library Python text blob untuk mengimplementasikan model Naïve Bayes ini, dan implementasinya ditunjukkan pada contoh Gambar 2.



Gambar 2. Penerapan Naïve bayes

3.5. Evaluasi model

Dan setelah melakukan pelatihan model maka model klasifikasi di ukur dengan melakukan perhitungan akurasi dengan menggunakan *Confusion Matrix* untuk mengetahui seberapa baik pengklasifikasi kinerja model untuk mengenali fitur

dari kategori kelas yang berbeda [18]. Pada hasilnya *confusion matrix* menghasilkan hasil seperti akurasi, presisi, *recall* dan F1-score, *matrix* memiliki terdiri dari empat hasil dari prediksinya yaitu :

- True negative* (TN) : Total dari data yang benar negatif dan hasil prediksi sebagai data di kategori negatif.
- True Positive* (TP) : Total dari data yang benar positif dan hasil prediksi sebagai data di kategori positif.
- False Positive* (FP) : Total dari data yang benar negatif dan hasil prediksi data yang di kategori positif.
- False Negative* (FN) : Total dari data yang benar positif dan hasil prediksi data yang di kategori negatif.

Berikut merupakan gambaran analisa pada *confusion matrix* :

- Akurasi

Merupakan hasil analisa untuk pengukuran model pada klasifikasi sebagai persentase jumlah data yang diprediksikan secara benar oleh algoritma.

Rumus :

$$\text{akurasi} = \frac{(TP+TN)}{(TP+TN+FP+FN)} * 100\% \quad (5)$$

- Presisi

Merupakan sebuah hasil dari ketepatan data yang diprediksi untuk kategori yang benar dari model klasifikasi dalam pengujian.

Rumus :

$$\text{Pr e s i s i} = \frac{TP}{(TP+FP)} * 100\% \quad (6)$$

- Recall

Merupakan nilai yang mengukur hasil data yang terklasifikasi dengan benar pada kategorinya.

Rumus :

$$\text{R e c a l l} = \frac{TP}{(TP+FN)} * 100\% \quad (7)$$

4. HASIL DAN PEMBAHASAN

4.1. Hasil modeling

Pada hasil proses model klasifikasi menggunakan *naive bayes* dihasilkan prediksi dari sentimen analisis yang di tampilkan pada tabel berikut :

Tabel 3. Hasil Klasifikasi Data uji Pada Model

Data rasio	Data uji	Hasil Prediksi		
		Positif	Netral	Negatif
70:30	2911	2117	650	144
80:20	1941	1421	436	84
90:10	971	711	214	46

Hasil klasifikasi sentimen pada komentar *review* aplikasi ruang guru menunjukkan bahwa memiliki, hasil sentimen positif yang cukup tinggi oleh pengguna aplikasi yang hasilnya terdapat pada di tabel 4.

4.2. Hasil modeling

Hasil dari prediksi diimplementasikan ke tabel *confusion matrix* untuk mempermudah mengukur

kinerja dan akurasi model klasifikasi, berikut tabel *confusion matrix* setiap skenario :

Tabel 4. *Confusion matrix* 70:30

Label		Prediksi			Total
		Negatif	Netral	Positif	
Aktual	Negatif	65	41	122	228
	Netral	7	492	166	665
	Positif	72	117	1829	2018

Tabel 5. *Confusion Matrix* 80:20

Label		Prediksi			Total
		Negatif	Netral	Positif	
Aktual	Negatif	43	28	81	152
	Netral	3	331	109	443
	Positif	38	77	1231	1346

Tabel 6. *Confusion Matrix* 90:10

Label		Prediksi			Total
		Negatif	Netral	Positif	
Aktual	Negatif	22	11	43	76
	Netral	3	170	49	222
	Positif	21	33	619	673

Berdasarkan pada tabel *confusion matrix* tersebut, didapatkan hasil akurasi, presisi, *recall*, dan *F1-score*, dalam perhitungan evaluasi dari setiap skenario, dapat dilihat pada tabel berikut :

Tabel 7. Hasil *Confusion Matrix* 70:30

Label	Presisi	Recall	F1-score
Negatif	45,13%	28,51%	34,94%
Netral	75,69%	73,91%	74,82%
Positif	86,39%	90,63%	88,45%
Akurasi			81,96%
Macro avg	69,07%	64,35%	66,07%
Weighted avg	80,71%	81,73%	81,14%

Tabel 8. Hasil *Confusion Matrix* 80:20

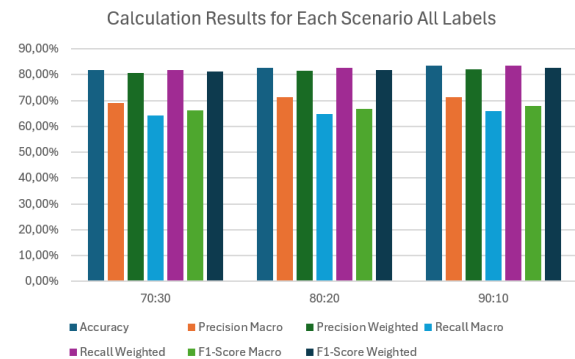
Label	Presisi	Recall	F1-score
Negatif	51,19%	28,29%	36,44%
Netral	75,92%	74,72%	75,31%
Positif	86,63%	91,46%	88,98%
Akurasi			82,68%
Macro avg	71,25%	64,82%	66,91%
Weighted avg	81,41%	82,69%	81,75%

Tabel 9. Hasil *Confusion Matrix* 90:10

Label	Presisi	Recall	F1-score
Negatif	47,83%	28,95%	36,06%
Netral	79,44%	76,58%	75,31%
Positif	87,06%	91,98%	89,45%
Akurasi			83,52%
Macro avg	71,44%	65,84%	67,83%
Weighted avg	82,24%	83,52%	82,64%

Pada gambar 3 menampilkan hasil perhitungan untuk berbagai skenario evaluasi model, menggunakan beberapa metrik kinerja, yang dianalisis pada berbagai pembagian data pelatihan-pengujian (70:30, 80:20, 90:10). Metrik yang digunakan meliputi

Akurasi, Presisi Makro, Recall Makro, dan F1-score Makro.



Gambar 3. Penerapan Naïve bayes

Dari gambar 3, terlihat bahwa Akurasi tetap relatif stabil di semua pembagian data, dengan peningkatan kecil saat berpindah dari pembagian 70:30 ke 90:10. Sementara itu, Presisi Makro dan F1-score Makro menunjukkan hasil serupa di berbagai skenario, dengan Presisi Makro sedikit lebih unggul daripada metrik lain di pembagian 70:30 dan 80:20. Recall Weighted dan F1-Score Weighted menunjukkan variasi yang lebih besar di seluruh pembagian data, dengan F1-Score Weighted umumnya berkinerja sedikit lebih baik daripada Recall Weighted di sebagian besar skenario. Secara keseluruhan, grafik ini menggambarkan bagaimana metrik kinerja model berubah dengan rasio data pelatihan-pengujian yang berbeda, menyoroti bahwa meskipun Akurasi tetap relatif stabil, metrik lain seperti Presisi dan Recall dapat bervariasi tergantung pada keseimbangan data, dan memberikan wawasan berharga untuk penyempurnaan dan optimasi model.

Berdasarkan hasil confusion matrix untuk setiap skenario, model mencapai akurasi stabil sebesar 82% di semua skenario. Namun, model pada skenario 90:10 mencapai F1-score rata-rata makro dan presisi yang relatif tinggi dibandingkan skenario lain, menjadikannya yang paling seimbang dalam mengenali semua label. Skenario terbaik untuk analisis sentimen dalam model klasifikasi Naïve Bayes adalah skenario 90:10 (90% data pelatihan dan 10% data uji), dengan akurasi 83,52%, rata-rata presisi makro 71,44%, recall 65,84%, dan skor F1 67,84%.

5. KESIMPULAN DAN SARAN

Berdasarkan penelitian yang telah dilakukan didapatkan kesimpulan yaitu dari hasil pengambilan data, dan *preprocessing* mendapatkan total data 9700 data. Dilakukan pelabelan data dengan bantuan *lexicon vader* mendapatkan hasil pelabelan positif 6.726 , netral sebanyak 2.232 dan negatif sebanyak 760. Bahwa pelabelan menunjukkan sentimen pada komentar *review* lebih banyak positif, dan dilakukan pengujian klasifikasi sentimen dengan *naive bayes* menghasilkan model baik untuk klasifikasi naive bayes pada skenario 90:10 dimana mendapatkan nilai

akurasi sebesar 83,52% , dapat disimpulkan bahwa *naive bayes* model yang baik untuk mengklasifikasi sentimen komentar. Namun ada saran dalam penanganan data tidak seimbang antar label untuk meningkatkan model disarankan dengan SMOTE ataupun teknik *oversampling* yang membuat meningkatkan klasifikasi antar label pada model.

DAFTAR PUSTAKA

- [1] E. F. Baharsyah, A. Armanto, T. H. B. Aviani, dan C. Wulandari, "Analisis Sentimen Pengguna Terhadap Aplikasi Belajar Online Ruang Guru Pada Ulasan Google Play Store Menggunakan Algoritma Naïve Bayes dan Support Vector Machine," *Innov. J. Soc. Sci. Res.*, vol. 4, no. 3, hal. 2965–2979, 2024.
- [2] F. Firman dan S. Rahayu, "Pembelajaran Online di Tengah Pandemi Covid-19," *Indones. J. Educ. Sci.*, vol. 2, no. 2, hal. 81–89, 2020, doi: 10.31605/ijes.v2i2.659.
- [3] A. P. Giovani, A. Ardiansyah, T. Haryanti, L. Kurniawati, dan W. Gata, "Analisis Sentimen Aplikasi Ruang Guru Di Twitter Menggunakan Algoritma Klasifikasi," *J. Teknoinfo*, vol. 14, no. 2, hal. 115, 2020, doi: 10.33365/jti.v14i2.679.
- [4] K. Lukoff, U. Lyngs, S. Gueorguieva, E. S. Dillman, A. Hiniker, dan S. A. Munson, "From Ancient Contemplative Practice to the App Store," in *Proceedings of the 2020 ACM Designing Interactive Systems Conference*, New York, NY, USA: ACM, Jul 2020, hal. 1551–1564. doi: 10.1145/3357236.3395444.
- [5] D. S. Putri, N. Sulistiyowati, dan A. Voutama, "Analisis Sentimen dan Pemodelan Ulasan Aplikasi AdaKami Menggunakan Algoritma SVM dan KNN," *J. Sensi*, vol. 9, no. 2, hal. 209–225, Agu 2023, doi: 10.33050/sensi.v9i2.2914.
- [6] Raksaka Indra Alhaqq, I Made Kurniawan Putra, dan Yova Ruldeviyani, "Analisis Sentimen terhadap Penggunaan Aplikasi MySAPK BKN di Google Play Store," *J. Nas. Tek. Elektro dan Teknol. Inf.*, vol. 11, no. 2, hal. 105–113, Mei 2022, doi: 10.22146/jnteti.v11i2.3528.
- [7] N. Kewsuwun dan S. Kajornkasirat, "A sentiment analysis model of agritech startup on Facebook comments using naive Bayes classifier," *Int. J. Electr. Comput. Eng.*, vol. 12, no. 3, hal. 2829–2838, 2022, doi: 10.11591/ijece.v12i3.pp2829-2838.
- [8] D. Pratmanto, R. Rousyati, F. F. Wati, A. E. Widodo, S. Suleman, dan R. Wijianto, "App Review Sentiment Analysis Shopee Application in Google Play Store Using Naive Bayes Algorithm," *J. Phys. Conf. Ser.*, vol. 1641, no. 1, 2020, doi: 10.1088/1742-6596/1641/1/012043.
- [9] S. Syafrizal, M. Afdal, dan R. Novita, "Analisis Sentimen Ulasan Aplikasi PLN Mobile Menggunakan Algoritma Naïve Bayes Classifier dan K-Nearest Neighbor," *MALCOM Indones. J. Mach. Learn. Comput. Sci.*, vol. 4, no. 1, hal. 10–19, 2023, doi: 10.57152/malcom.v4i1.983.
- [10] H. Taufiqurrahman, F. Tri Anggraeny, dan M. Muharrom Al Haromainy, "Perbandingan Algoritma Naïve Bayes Dan K-Nearest Neighbor Pada Analisis Sentimen Ulasan Aplikasi Mypertamina," *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 7, no. 6, hal. 3934–3939, 2024, doi: 10.36040/jati.v7i6.7801.
- [11] A. Firdaus, W. I. Firdaus, P. Studi, T. Informatika, M. Digital, dan P. N. Sriwijaya, "Text Mining Dan Pola Algoritma Dalam Penyelesaian Masalah Informasi: (Sebuah Ulasan)," vol. 13, no. 1, hal. 66–78, 2021.
- [12] R. Christian dan I. Fenriana, "PERANCANGAN APLIKASI ANALISIS SENTIMEN PADA ULASAN MASKAPAI PENERBANGAN DI TRIPADVISOR MENGGUNAKAN METODE NAÏVE BAYES," *akselerator J. Sains Terap. Dan Teknol.*, hal. 25–39, 2023.
- [13] F. V. Sunata, S. Hariyanto, D. S. Dwi Putra, dan H. Wijaya, "Penerapan Data Mining Untuk Merekomendasikan Seri Produk NAS Kepada Calon Konsumen Toko Storage Menggunakan Algoritma Multinomial Naïve Bayes," *Algor*, vol. 4, no. 1, hal. 44–55, 2022, doi: 10.31253/algor.v4i1.1498.
- [14] S. H. Ramadhani dan M. I. Wahyudin, "Analisis Sentimen Terhadap Vaksinasi Astra Zeneca pada Twitter Menggunakan Metode Naïve Bayes dan K-NN," *J. JTIK (Jurnal Teknol. Inf. dan Komunikasi)*, vol. 6, no. 4, hal. 526–534, 2022, doi: 10.35870/jtik.v6i4.530.
- [15] N. A. Pasaribu dan Sriani, "The Shopee Application User Reviews Sentiment Analysis Employing Naïve Bayes Algorithm," *Int. J. Softw. Eng. Comput. Sci.*, vol. 3, no. 3, hal. 194–204, 2023, doi: 10.35870/ijsecs.v3i3.1699.
- [16] Y. Asri, W. N. Suliyanti, D. Kuswardani, dan M. Fajri, "Pelabelan Otomatis Lexicon Vader dan Klasifikasi Naive Bayes dalam menganalisis sentimen data ulasan PLN Mobile," *PETIR J. Pengkaj. dan Penerapan Tek. Inform.*, vol. 15, no. 2, hal. 264–275, 2022, doi: 10.33322/petir.v15i2.1733.
- [17] N. Chintalapudi, G. Battineni, M. Di Canio, G. G. Sagaro, dan F. Amenta, "Text mining with sentiment analysis on seafarers' medical documents," *Int. J. Inf. Manag. Data Insights*, vol. 1, no. 1, hal. 100005, Apr 2021, doi: 10.1016/j.jjime.2020.100005.
- [18] D. Normawati dan S. A. Prayogi, "Implementasi Naïve Bayes Classifier Dan Confusion Matrix Pada Analisis Sentimen Berbasis Teks Pada Twitter," *J. Sains Komput. Inform. (J-SAKTI)*, vol. 5, no. 2, hal. 697–711, 2021.