

ANALISIS SENTIMEN TINGKAT KEPUASAN APLIKASI WORDPRESS MENGUNAKAN METODE K-NEAREST NEIGHBOOR DAN NAIVE BAYES

SENTIMENT ANALYSIS OF WORDPRESS APPLICATION SATISFACTION LEVEL USING K-NEAREST NEIGHBOR AND NAIVE BAYES METHODS

Moch. Siddiq Hamid ^{1*}

Ade Eviyanti ²

Hindarto Hindarto ³

^{1,2,3}Teknik Informatika, Fakultas Sains dan Teknologi, Universitas Muhammadiyah Sidoarjo
¹211080200085@umsida.ac.id, ²eviyantiade@umsida.ac.id, ³hindarto@umsida.ac.id

Abstrak

Kepuasan pengguna mencerminkan emosi saat membandingkan layanan yang diterima dengan harapan, sehingga memahami kepuasan pengguna menjadi penting bagi pengembangan aplikasi. Penelitian ini bertujuan untuk menilai tingkat kepuasan individu terhadap penggunaan aplikasi WordPress di platform Google Play mengidentifikasi area yang perlu ditingkatkan. Analisis sentimen dengan algoritma KNN dan Naïve bayes sebagai metode yang digunakan untuk mengekstrak informasi dari 5.000 ulasan pengguna yang diunduh dari Google Play Store. Hasil penelitian menunjukkan mayoritas ulasan memiliki sentimen positif, dengan Naïve Bayes memberikan hasil yang lebih baik dibandingkan KNN, mencapai akurasi 88%, presisi 89,45%, recall 88%, dan F1-Score 83% pada pembagian data 90:10. Awan kata dari ulasan positif menampilkan kata-kata seperti "mantap", "bagus", "membantu", "aplikasi", dan "baik", yang mencerminkan kepuasan pengguna terhadap kemudahan dan manfaat aplikasi, sedangkan ulasan negatif menampilkan kata-kata seperti "sulit", "coba", dan "gagal" yang mengindikasikan kendala teknis dan ketidakpuasan pengguna. Penelitian ini menyimpulkan bahwa aplikasi WordPress telah memberikan pengalaman yang memuaskan bagi sebagian besar pengguna, namun beberapa area teknis perlu diperbaiki. Hasil penelitian ini akan memberikan informasi penting bagi pengembang aplikasi dalam upaya meningkatkan kualitas layanan dan citra aplikasi.

Kata Kunci: Analisis Sentimen, KNN, Naive Bayes, Tingkat kepuasan, WordPress

Abstract

User satisfaction reflects emotions when comparing services received with expectations, so understanding user satisfaction is important for app development. This research aims to evaluate user satisfaction with WordPress apps on the Google Play Store and identify areas for improvement. Sentiment analysis with KNN and Naïve bayes algorithms as the method used to extract information from 5,000 user reviews downloaded from Google Play Store. The results showed the majority of reviews had positive sentiments, with Naïve Bayes providing better results than KNN, achieving 88% accuracy, 89.45% precision, 88% recall, and 83% F1-Score on a 90:10 data split. The word cloud of positive reviews featured words such as "great", "good", "helpful", "app", and "good", reflecting user satisfaction with the ease and benefits of the app, while negative reviews featured words such as "difficult", "try", and "fail" indicating technical difficulties and user dissatisfaction. This study concludes that WordPress apps have provided a satisfactory experience for most users, but some technical areas need improvement. The results of this study will provide valuable information for app developers in efforts to improve service quality and the app's reputation.

keywords: Sentiment analysis, KNN Naïve bayes, Satisfaction level, Wordpress

1. Pendahuluan

Di masa perkembangan teknologi digital ini, aplikasi sekarang digunakan sebagai media utama untuk memfasilitasi berbagai aktivitas, baik personal maupun profesional. Salah satu aspek penting yang memengaruhi keberhasilan suatu aplikasi adalah tingkat kepuasan pengguna. Tingkat kepuasan pengguna merupakan tingkat emosional yang dirasakan pengguna ketika membandingkan pelayanan atau hasil yang dicapai dengan harapannya [1]. Ketika pengguna merasa bahwa aplikasi berfungsi dengan baik, mudah digunakan, dan memenuhi atau bahkan melampaui ekspektasi mereka, maka tingkat kepuasan mereka akan tinggi. Hal ini sering kali

dituangkan dalam bentuk ulasan atau komentar pada aplikasi. Ekspresi emosi manusia, baik berupa kepuasan maupun kekesalan, kerap diungkapkan melalui tulisan, termasuk komentar yang diposting pengguna di platform distribusi aplikasi, seperti Google Play Store[2].

Namun, ulasan pengguna yang berisi opini dan persepsi terhadap aplikasi sering kali tidak terstruktur dan sulit untuk dianalisis secara manual. Dalam hal ini, analisis sentimen menjadi solusi yang efektif. Analisis sentimen adalah metode untuk mengumpulkan dan menganalisis ulasan pengguna untuk menilai perasaan dan pendapat mereka tentang aplikasi yang akan disimpan dalam teks [3].

Berdasarkan penelitian sebelumnya [2],[3],[4],[5], menunjukkan bahwa analisis sentimen dapat membantu mengatasi beberapa permasalahan seperti mengidentifikasi persepsi publik, meningkatkan efisiensi analisis data, dan mendukung perbaikan produk atau layanan berdasarkan feedback pengguna. Dalam mengevaluasi kepuasan terhadap suatu aplikasi oleh pengguna, dapat digunakan metode analisis sentimen, misalnya dalam survei untuk menyelidiki sentimen aplikasi Wordpress yang diisi oleh pengguna melalui kolom komentar.

Namun, meskipun sudah banyak penelitian tentang analisis sentimen, masih terdapat gap dalam memilih metode yang paling akurat dan sesuai untuk mengevaluasi ulasan pengguna terhadap aplikasi tertentu, seperti Wordpress. Wordpress adalah aplikasi resmi yang dirancang untuk membantu pengguna dalam mengelola situs web atau blog berbasis Wordpress langsung dari perangkat smartphone.

Penelitian ini dilakukan untuk mengisi gap tersebut dengan menganalisis umpan balik pelanggan pada aplikasi Wordpress yang tersedia di platform Google Play Store analisis sentiment serta menggunakan 2 metode yaitu KNN dan Naive Bayes dalam mengklasifikasikan komentar menjadi positif dan negatif. Pengklasifikasi Naive Bayes (NB) adalah Merupakan algoritma klasifikasi yang memanfaatkan teorema Bayes untuk menghitung probabilitas suatu data termasuk ke dalam kelas tertentu, dengan asumsi bahwa setiap fitur atau prediktor bersifat independen satu sama lain (asumsi "naive") [6] . Pada algoritma k-nearest neighbor (KNN), suatu titik data akan diklasifikasikan berdasarkan kedekatannya dengan titik-titik data lainnya yang ada dalam dataset [7]

Pada temuan sebelumnya yang dilakukan oleh Aluisius Dwiki Adhi Putra, Safitri Juanita dengan judul ' Analisis Sentimen Pada Ulasan Pengguna Aplikasi Bibit Dan Bareksa Dengan Algoritma KNN' menunjukkan bahwa nilai akurasi, presisi, dan recall model klasifikasi untuk aplikasi Bibit dan Bareksa masing-masing sebesar 85,14%, 91,91%, dan 76,44%, menggunakan algoritma K-Nearest Neighbors dengan pembagian data 60:40 pada dataset ulasan pengguna. Sebaliknya, angka untuk Bareksa adalah 81,70%, 87,15%, dan 75,73%[8].

Penelitian yang dilakukan oleh Moh Khoirul Insan, Umi Hayati, dan Odi Nurdiawan dengan judul 'Analisis Sentimen Aplikasi BRImo pada Ulasan Pengguna di Google Play Menggunakan Algoritma Naive Bayes' menunjukkan bahwa pendekatan Naive Bayes untuk klasifikasi menghasilkan tingkat akurasi 84,52%, presisi 82,51%, dan recall 87,62%[9].

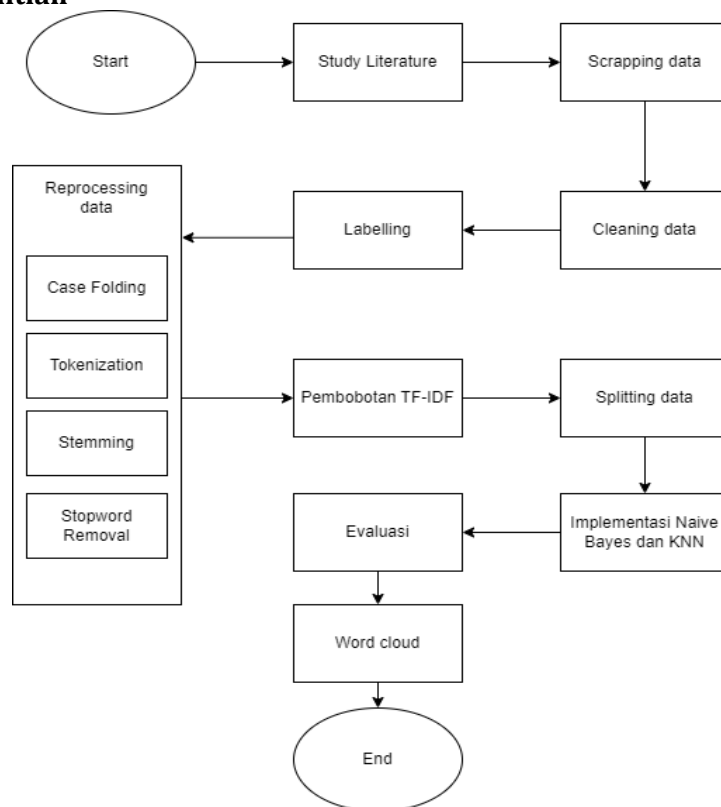
Berdasarkan tinjauan terhadap penelitian sebelumnya, Algoritme KNN dan Naive Bayes telah menunjukkan keampuannya dalam menganalisis sentimen ulasan pengguna di berbagai aplikasi, seperti *Bibit*, *Bareksa*, dan *BRImo*. Algoritma KNN menunjukkan keunggulan dalam menangani data dengan pola non-linear, sementara Naive Bayes unggul dalam memanfaatkan asumsi independensi fitur untuk menghasilkan hasil klasifikasi yang stabil.

Penelitian ini bertujuan untuk mengimplementasikan kedua algoritma tersebut pada ulasan pengguna aplikasi Wordpress dan membandingkannya, sehingga dapat memberikan perbandingan performa yang lebih mendalam. Dengan mengevaluasi akurasi, precision, recall, dan F1 score dari

Diharapkan bahwa kedua algoritme dan penelitian ini akan menawarkan perspektif baru tentang seberapa baik algoritma bekerja dalam analisis sentimen. Analisis dilakukan menggunakan platform Google Colab dengan metode Naïve Bayes dan KNN untuk membandingkan akurasi dengan tools tersebut [10].

Temuan analisis ini diharapkan dapat memberikan wawasan yang mendalam tentang persepsi pengguna, membantu para pengembang dalam meningkatkan kualitas dan kegunaan aplikasi WordPress.

2. Metode Penelitian



Gambar 1. Alur Penelitian

Sepuluh langkah diambil Dalam penelitian ini untuk mencapai temuan yang dibutuhkan, seperti yang diilustrasikan pada Gambar 1. Penjelasan lebih lanjut mengenai proses penelitian akan diuraikan berikut ini:

Studi literature

Pada tahapan pertama, studi literatur pengumpulan refrensi yang meliputi buku, jurnal ilmiah dari dalam dan luar negeri, serta pencarian informasi melalui internet [11].

Scrapping data

Pada tahapan kedua, Scrapping data merupakan proses otomatisasi mengumpulkan data dari berbagai sumber, seperti web, basis data, aplikasi bisnis, atau sistem lama [12]. Web scraping memiliki berbagai manfaat praktis, seperti memungkinkan pengumpulan data harga dari berbagai situs untuk perbandingan, pengumpulan alamat email untuk pemasaran (dengan memperhatikan hukum privasi), serta memantau topik trending di media sosial untuk strategi pemasaran [13].

Pada penelitian ini data dikumpulkan menggunakan tools google colab dan Teknik scrapping data, didapat data dengan total 5000 data berhasil diperoleh dari Google Playstore. Atribut dan pada tabel 2 merupakan tampilan data.

Tabel 2. Data awal			
reviewId	content	rating	at
20e50699- cef2-4ef4- b71b- 4ff3dc25805e	Sangat enak menggunak an WordPr	5	2024-10-16
956259d0-a992- 4864-9e1b- 1a05e70d0fa0	Mantap ..bu at perkenlkn usaha	5	2024-10-16

Setelah data sudah tampil seperti pada gambar 1 , selanjutnya akan akan dilakukan penghapusan fitur yang tidak d yang tidak dibutuhkan sehingga menyisakan content dan score untuk penelitian lebih lanjut yang ditujunkkan pada tabel 3:

Tabel 3. Data setelah penghapusan fitur	
content	rating
Sangat enak menggunakan WordPress	5
Mantap ..buat perkenlkn usaha	3

Cleaning data

Pada tahapan ketiga, proses cleaning data merupakan proses untuk memfilter tribut-atribut tidak memberikan dampak signifikan terhadap penelitian, serta menyaring komentar-komentar tidak relevan pada penelitian [14]. Setelah fitur dihapus, tahap berikutnya adalah pembersihan data, untuk hasilnya terdapat pada tabel 4.

Tabel 4. Cleaning data	
Sebelum cleaning data	Setelah cleaning data
Aneh sekali, harus dipaksa install aplikasi lain agar bisa membuat web lebih dari satu.	Aneh sekali harus dipaksa install aplikasi lain agar bisa membuat web lebih dari satu

Labelling data

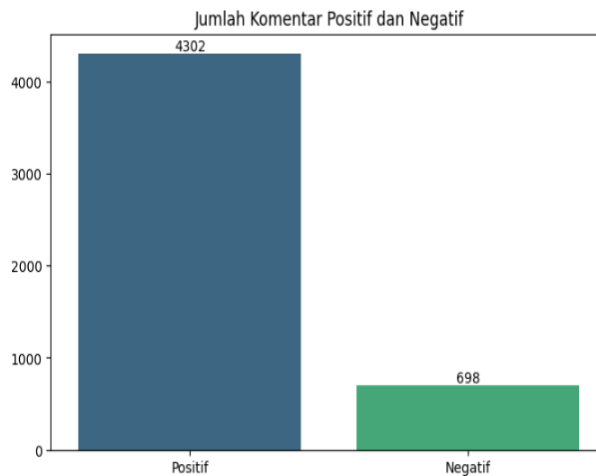
Pada tahap keempat, ulasan diklasifikasikan berdasarkan rating, di mana rating 1, 2, dan 3 digolongkan sebagai sentimen negatif, sementara rating 4 dan 5 digolongkan sebagai sentimen positif [15].

Setelah melalui tahap reprocessing akan dibaeri label. Label klasifikasi dibagi menjadi dua kategori: positif dan negatif, yang ditentukan berdasarkan rating atau jumlah bintang yang diberikan pada setiap komentar. Komentar yang mendapat nilai 1 hingga 3 akan diklasifikasikan sebagai negatif, sedangkan komentar yang mendapat nilai 4 hingga 5 akan diklasifikasikan sebagai positif. Pada tabel 5 menampilkan informasi mengenai labelling

Tabel 5. Labelling

Content	rating	Sentimen
Hati hati kalo tidak mau di hack login akun nya	1	Negatif
Mantap buat perkenlkn usaha	5	Positif
Sangat mendukung	4	Positif
Sangat enak menggunakan WordPress	5	Positif

Setelah proses pelabelan, distribusi data berdasarkan label positif dan negatif, pada gambar 2 menggambarkan sebaran data dalam kategori sentimen tersebut.



Gambar 2. Distribusi data positif dan negatif

Processing data

Pada tahapan kelima, yaitu pemrosesan data, terdapat empat langkah pemrosesan data yang akan dilanjutkan ke pembobotan TF-IDF yaitu :

pertama, case folding berfungsi untuk mengubah komentar yang ada di Google Playstore dan Appstore menjadi huruf kecil untuk memudahkan analisis teks secara konsisten [16].

Teknik tokenisasi membagi teks menjadi kata-kata atau token sesuai dengan karakter baris baru, tab, dan spasi [17].

Ketiga, stemming memiliki tujuan yaitu mengubah bentuk kata yang terinfleksi atau terderivasi (misalnya, variasi bentuk kata) menjadi bentuk dasar yang umum [18].

Keempat, pada tahap ini terdapat proses pencarian kata yang frekuensi kemunculannya lebih sering dibanding kata lain namun tidak ada arti atau makna terhadap suatu kalimat [19].

Setelah tahap pelabelan data selesai, dilakukan proses reprocessing atau pemrosesan ulang data melalui empat tahapan utama yaitu Langkah pertama adalah case folding, tokenization, stemming, stopwords removal. Tahap reprocessing ada di tabel 6.

Tabel 6. Reprocessing data

Tahap	Komen
Data Awal	Aneh sekali, harus dipaksa install aplikasi lain agar bisa membuat web lebih dari satu.
Case Folding	aneh sekali harus dipaksa install aplikasi lain agar bisa membuat web lebih dari satu
Tokenization	["aneh", "sekali", "harus", "dipaksa", "install", "aplikasi", "lain", "agar", "bisa", "membuat", "web", "lebih", "dari", "satu"]
Stemming	["aneh", "sekali", "harus", "paksa", "install", "aplikasi", "lain", "agar", "bisa", "buat", "web", "lebih", "dari", "satu"]
Stopword removal	["aneh", "paksa", "install", "aplikasi", "buat", "web"]

Pembobotan TF-IDF

Pada tahapan keenam, Untuk menentukan signifikansi sebuah kata di dalam koleksi dokumen, pembobotan TF-IDF digunakan. Metode untuk menentukan tingkat relevansi kata dalam dokumen disebut TF-IDF (Term Frequency-Inverse Document Frequency) [20]. Pembobotan TF-IDF terbagi menjadi dua tahap, yaitu:

Kata yang sering muncul dalam suatu dokumen akan mempengaruhi perhitungan maka dengan itu digunakanlah TF, yang dapat menunjukkan apakah kata tersebut dianggap penting dalam konteks dokumen itu. Rumus perhitungannya terdapat pada persamaan 1.

$$TF(t, d) = \frac{f_{t,d}}{\sum_{t' \in d} f_{t',d}} \quad (1)$$

IDF digunakan untuk menilai signifikansi sebuah kata di seluruh koleksi teks, dengan rumus perhitungannya terdapat pada persamaan 2.

$$IDF(t) = \log \left(\frac{N}{df(t)} \right) \quad (2)$$

TF-IDF gabungan dari 2 rumus

$$TF-IDF(t,d) = TF(t,d) \times IDF(t) \quad (3)$$

Splitting data

Pada tahapan ketujuh, pembagian data merupakan kegiatan memisahkan suatu data menjadi 2 bagian yaitu: kolom X sebagai fitur data pelatihan dan data uji, serta kolom Y untuk label pelatihan dan uji. Besaran data uji dapat diatur menggunakan parameter `test_size`, dan untuk menguji keseluruhan dataset, dapat digunakan sampel data [21].

Implementasi KNN dan Naive Bayes

Tahap kedelapan adalah proses penerapan algoritma yang digunakan yaitu KNN dan Naive Bayes:

Menemukan kumpulan atau kelompok k objek dalam data pelatihan yang paling mirip dengan objek dalam data pengujian adalah cara kerja algoritma KNN. Sistem klasifikasi ini diperlukan untuk mencari informasi yang relevan [22]. Dalam data mining, algoritma KNN dikenal melakukan klasifikasi dan regresi [23]. Rumus KNN terdapat pada persamaan 4.

$$d(p, q) = \sqrt{\sum_{i=1}^n (p_i - q_i)^2} \quad (4)$$

Klasifikasi Naive Bayes adalah algoritma dengan cara pengaplikasiannya menggunakan teori Bayes untuk mengklasifikasikan data. [24]. Naive Bayes merupakan teknik klasifikasi yang Persamaan 5 menunjukkan rumus untuk Naive Bayes. sederhana, tetapi tetap menunjukkan kinerja dan akurasi yang sangat baik. [25]. Untuk rumus Naive bayes terdapat pada persamaan 5.

$$P(C_i|X) = \frac{P(C_i|X) \cdot P(C_i)}{P(X)} \quad (5)$$

Evaluasi

Pada tahapan kesembilan, dilakukan penilaian terhadap algoritma yang diterapkan yaitu naive bayes dan KNN. Matriks confusion, sebuah teknik evaluasi yang dibuat khusus untuk model klasifikasi, digunakan untuk membandingkan kedua algoritme tersebut. Proses ini berfungsi membandingkan hasil prediksi yang dihasilkan oleh model dengan data atau nilai yang sebenarnya, sehingga memungkinkan untuk mengevaluasi seberapa akurat prediksi yang dibuat oleh model [26].

Melalui penggunaan confusion matrix, peneliti dapat memperoleh wawasan lebih mendalam tentang berbagai aspek kinerja model, termasuk jumlah prediksi yang benar dan salah. Selain itu, evaluasi metrik lain seperti akurasi, presisi, recall, dan F1-score dapat dihitung, proses ini sangat penting untuk mengetahui efektivitas masing-masing algoritma yang diterapkan. Pada tabel 1 digambarkan confusion matriks dalam melakukan prediksi.

Tabel 1. Confusion Matrix			
		Prediksi	
		Positive	Negative
Aktual	Positive	True Positive (TP)	False Negative(FN)
	Negative	False Negative (FP)	True Negative(TN)

Setelah menemukan nilai dari keempat istilah yang ada setelah itu dapat dilakukan untuk menghitung akurasi, presisi recall dan f1 score yang terdapat pada rumus berikut:

Akurasi menilai terkait prediksi suatu data dengan membandingkan jumlah yang diprediksi apakah sudah sesuai dengan kelas aslinya terhadap total keseluruhan prediksi yang dihasilkan oleh model klasifikasi [27]. Pada persamaan 6 merupakan rumus akurasi.

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (6)$$

Presisi adalah ukuran yang menunjukkan seberapa banyak data yang diklasifikasikan sebagai positif benar-benar positif [28]. Untuk rumus recall terdapat di persamaan 7.

$$Precision = \frac{TP+TN}{TP+FP} \quad (7)$$

Persentase dokumen yang berhasil ditemukan selama proses pencarian informasi dikenal dengan istilah recall [29].

$$Recall = \frac{TP}{TP+FN} \quad (8)$$

Harmonic mean dari presisi dan recall, yang dikenal sebagai F1-Score, digunakan untuk memberikan gambaran seimbang tentang performa model, terutama ketika ada ketidakseimbangan kelas dalam data

$$F1\ Score = 2 \times \frac{Recall \times Precision}{recall+Precision} \quad (9)$$

Wordcloud

Selain itu, ada wordcloud di tahap terakhir ini. wordcloud untuk menggambarkan secara grafis istilah-istilah yang paling sering muncul di artikel yang dianalisis. [30].

3. Hasil

Setelah data melalui tahap pemrosesan pada tahap selanjutnya akan dilakukan pembobotan TF-IDF yang hasilnya akan ditampilkan pada gambar 3 yang dilakukan menggunakan ekstraksi fitur menggunakan modul TfidfVectorizer dari library Scikit-Learn.

bagus	0.072376
good	0.061777
mantap	0.055382
sangat	0.053901
bantu	0.039805
aplikasi	0.021600
mudah	0.021020
keren	0.018576
coba	0.017248
baik	0.016423
wordpress	0.015980
dulu	0.015015
suka	0.014366
buat	0.013921

Gambar 3. Hasil pembobotan TF-IDF

Sebelum mengklasifikasikan data menggunakan teknik machine learning, data terlebih dahulu dipisah menjadi data uji dan data latih. Evaluasi dilaksanakan dengan memanfaatkan confusion matrix yang mengukur aspek-aspek seperti akurasi, presisi, recall, dan skor F1 untuk KNN dan Naïve Bayes.

Pembagian data dilakukan dengan sejumlah rasio yang berbeda, seperti 90:10, 80:20, 70:30, 60:40, dan 50:50, menyebabkan variasi dalam hasil confusion matrix. Anda dapat melihat hasil-hasil tersebut dalam Tabel 7 untuk akurasi, Tabel 8 untuk presisi, Tabel 9 untuk recall, dan Tabel 10 untuk skor F1.

Tabel 7. Tabel accuracy

Splitting Data (%)	Akurasi	
	Naive bayes	KNN
90 : 10	88%	87.2%
80 : 20	87%	86.3%
70 : 30	86.46%	85.8%
60 : 40	86.3%	85.7%
50 : 50	85.56%	84.64%

Tabel 8. Tabel precision

Splitting Data (%)	Precision	
	Naive bayes	KNN
90 : 10	89.45%	76.03%
80 : 20	87.3 %	88.1%
70 : 30	88.3%	77.3%
60 : 40	88.1%	79%
50 : 50	88.3%	77.1%

Tabel 9. Tabel recall

Splitting Data (%)	Recall	
	Naive bayes	KNN
90 : 10	88%	87.2%
80 : 20	87%	86.3%
70 : 30	86.46%	85.8%
60 : 40	86.3%	85.7%
50 : 50	85.56%	84.64%

Tabel 10. Tabel F1-Score

Splitting Data (%)	F1-Score	
	Naive bayes	KNN
90 : 10	83%	81.2%
80 : 20	81.8%	80%
70 : 30	80.6%	79.4%
60 : 40	80.4%	79.3%
50 : 50	79.28%	78.77%

Data diklasifikasikan menggunakan KNN dan naïve bayes, diikuti dengan visualisasi data menggunakan wordcloud untuk mengidentifikasi kata-kata yang sering digunakan dalam ulasan (Gambar 4 dan 5).



Gambar 4. Word cloud positif

Digunakan algoritma KNN dan Naïve Bayes pada penelitian ini yang nantinya akan dibandingkan performa kedua algoritma dalam mengklasifikasikan sentimen positif dan negatif.

Hasil analisis menunjukkan bahwa algoritma Naïve Bayes lebih efektif dibandingkan KNN dalam mengklasifikasikan sentimen, terutama ketika proporsi data pelatihan lebih besar dibandingkan data pengujian. Sebagian besar ulasan memiliki sentimen positif, mencerminkan aspek-aspek seperti kemudahan penggunaan dan manfaat aplikasi, sementara ulasan negatif menyoroti masalah teknis yang perlu diperbaiki. Temuan memberikan masukan terkait aplikasi oleh pengguna aplikasi untuk terus mengembangkan aplikasinya.

Hasil dari word cloud memperlihatkan kata-kata positif seperti "mantap," "good," "membantu," "aplikasi," dan "baik" (gambar 4), yang mencerminkan aspek-aspek dari aplikasi WordPress yang dihargai oleh pengguna, seperti kemudahan penggunaan dan manfaatnya. Di sisi lain, kata-kata negatif seperti "susah," "coba," dan "gagal" (gambar 5) menggambarkan masalah yang dihadapi pengguna, dengan "susah" dan "gagal" merujuk pada kendala teknis, sedangkan "coba" menggambarkan kekecewaan setelah mencoba fitur tertentu. Meskipun sebagian besar ulasan bersifat positif, temuan ini memberikan wawasan berharga bagi pengelola aplikasi WordPress untuk meningkatkan pengalaman pengguna dan memperbaiki citra aplikasi.

Penelitian ini memiliki beberapa keterbatasan, termasuk terbatasnya eksplorasi model klasifikasi yang digunakan serta fokus pada ulasan dalam satu bahasa. Untuk pengembangan penelitian, disarankan untuk menggunakan model yang lebih kompleks seperti Support Vector Machine (SVM), Random Forest, atau model deep learning seperti Long Short-Term Memory (LSTM) guna membandingkan kinerja dengan algoritma yang telah digunakan. Selain itu, penelitian lanjutan dapat mengkaji analisis sentimen multibahasa untuk mengidentifikasi perbedaan persepsi berdasarkan negara atau bahasa, serta memantau tren perubahan sentimen pengguna dari waktu ke waktu untuk memahami bagaimana persepsi pengguna terhadap aplikasi berkembang.

6. Referensi

- [1] E. Hasibuan and E. A. Heriyanto, "ANALISIS SENTIMEN PADA ULASAN APLIKASI AMAZON SHOPPING DI GOOGLE PLAY STORE MENGGUNAKAN NAIVE BAYES CLASSIFIER," *JTS*, vol. 1, no. 3, 2022.
- [2] N. C. Agustina, D. Herlina Citra, W. Purnama, C. Nisa, and A. Rozi Kurnia, "MALCOM: Indonesian Journal of Machine Learning and Computer Science The Implementation of Naïve Bayes Algorithm for Sentiment Analysis of Shopee Reviews on Google Play Store Implementasi Algoritma Naive Bayes untuk Analisis Sentimen Ulasan Shopee pada Goo," vol. 2, pp. 47–54, 2022.
- [3] S. Rahayu, Y. MZ, J. E. Bororing, and R. Hadiyat, "Implementasi Metode K-Nearest Neighbor (K-NN) untuk Analisis Sentimen Kepuasan Pengguna Aplikasi Teknologi Finansial FLIP," *Edumatic J. Pendidik. Inform.*, vol. 6, no. 1, pp. 98–106, Jun. 2022, doi: 10.29408/edumatic.v6i1.5433.
- [4] T. A. Sari, E. Sinduningrum, and F. Noor Hasan, "KLIK: Kajian Ilmiah Informatika dan Komputer Analisis Sentimen Ulasan Pelanggan Pada Aplikasi Fore Coffee Menggunakan Metode Naïve Bayes," *Media Online*, vol. 3, no. 6, pp. 773–779, 2023, doi: 10.30865/klik.v3i6.884.
- [5] M. N. Muttaqin and I. Kharisudin, "Analisis Sentimen Pada Ulasan Aplikasi Gojek Menggunakan Metode Support Vector Machine dan K Nearest Neighbor," *UNNES J. Math.*, vol. 10, no. 2, pp. 22–27, 2021, [Online]. Available: <http://journal.unnes.ac.id/sju/index.php/ujm>

- [6] O. Peretz, M. Koren, and O. Koren, "Naive Bayes classifier – An ensemble procedure for recall and precision enrichment," *Eng. Appl. Artif. Intell.*, vol. 136, no. PB, p. 108972, 2024, doi: 10.1016/j.engappai.2024.108972.
- [7] P. Hou, L. Zhou, and Y. Yang, "Density clustering method based on k-nearest neighbor propagation," *J. Phys. Conf. Ser.*, vol. 2858, no. 1, 2024, doi: 10.1088/1742-6596/2858/1/012041.
- [8] A. Asro'i and H. Februariyanti, "Analisis Sentimen Pengguna Twitter Terhadap Perpanjangan PpkM Menggunakan Metode K-Nearest Neighbor," *J. Khatulistiwa Inform.*, vol. 10, no. 1, pp. 17–24, 2022, doi: 10.31294/jki.v10i1.12624.
- [9] M. K. Insan, U. Hayati, and O. Nurdiawan, "Analisis Sentimen Aplikasi Brimo Pada Ulasan Pengguna Di," *J. Mhs. Tek. Inform.*, vol. 7, no. 1, pp. 478–483, 2023.
- [10] Y. Yuliska and K. U. Syaliman, "Peningkatan Akurasi K-Nearest Neighbor Pada Data Index Standar Pencemaran Udara Kota Pekanbaru," *IT J. Res. Dev.*, vol. 5, no. 1, pp. 11–18, Jul. 2020, doi: 10.25299/itjrd.2020.vol5(1).4680.
- [11] N. Nurfaizah and S. R. Hidayat, "Sentimen Analisis Pengguna Produk Ponsel Menggunakan Algoritma Naïve Bayes," *J. Inf. Syst. Manag.*, vol. 6, no. 1, pp. 10–14, 2024, doi: 10.24076/joism.2024v6i1.1625.
- [12] R. Merdiansah and A. Ali Ridha, "Sentiment Analysis of Indonesian X Users Regarding Electric Vehicles Using IndoBERT," *J. Ilmu Komput. dan Sist. Inf. (JIKOMSI)*, vol. 7, no. 1, pp. 221–228, 2024.
- [13] I. S. H. Almaqbali, F. M. A. Al Khufairi, M. S. Khan, A. Z. Bhat, and I. Ahmed, "Web Scrapping: Data Extraction from Websites," *J. Student Res.*, pp. 1–4, 2020, doi: 10.47611/jsr.vi.942.
- [14] N. Y. Pradipta and H. Soetanto, "Sentiment Classification of General Election 2024 News Titles on Detik. com Online Media Website Using Multinomial Naive Bayes Method," *J. Appl. Sci. Eng. ...*, vol. 6, no. 1, 2024, [Online]. Available: <https://ascijournal.eu/index.php/asci/article/view/2754%0Ahttps://ascijournal.eu/index.php/asci/article/download/2754/1833>
- [15] N. Cahyono and Anggista Oktavia Praneswara, "Analisis Sentimen Ulasan Aplikasi TikTok Shop Seller Center di Google Playstore Menggunakan Algoritma Naive Bayes," *Indones. J. Comput. Sci.*, vol. 12, no. 6, pp. 3925–3940, 2023, doi: 10.33022/ijcs.v12i6.3473.
- [16] R. Kosasih and A. Alberto, "Analisis Sentimen Produk Permainan Menggunakan Metode TF-IDF Dan Algoritma K-Nearest Neighbor," *InfoTekJar J. Nas. Inform. dan Teknol. Jar.*, vol. 6, no. 1, pp. 134–139, 2021, [Online]. Available: <https://doi.org/10.30743/infotekjar.v6i1.3893>
- [17] W. G. S. Parwita, "A document recommendation system of stemming and stopword removal impact: A web-based application," *J. Phys. Conf. Ser.*, vol. 1469, no. 1, 2020, doi: 10.1088/1742-6596/1469/1/012050.
- [18] A. Özçift, K. Akarsu, F. Yumuk, and C. Söylemez, "Advancing natural language processing (NLP) applications of morphologically rich languages with bidirectional encoder representations from transformers (BERT): an empirical case study for Turkish," *Automatika*, vol. 62, no. 2, pp. 226–238, 2021, doi: 10.1080/00051144.2021.1922150.
- [19] O. Manullang, C. Prianto, and N. H. Harani, "Analisis Sentimen Untuk Memprediksi Hasil Calon Pemilu Presiden Menggunakan Lexicon Based Dan Random Forest," *J. Ilm. Inform.*,

vol. 11, no. 02, pp. 159–169, 2023, doi: 10.33884/jif.v11i02.7987.

- [20] A. H. Dani, E. Y. Puspaningrum, and R. Mumpuni, “Studi Performa TF-IDF dan Word2Vec Pada Analisis Sentimen Cyberbullying,” *Router J. Tek. Inform. dan Terap.*, vol. 2, no. 2, pp. 94–106, 2024, [Online]. Available: <https://doi.org/10.62951/router.v2i2.76>
- [21] J. Ipmawati, S. Saifulloh, and K. Kusnawi, “Analisis Sentimen Tempat Wisata Berdasarkan Ulasan pada Google Maps Menggunakan Algoritma Support Vector Machine,” *MALCOM Indones. J. Mach. Learn. Comput. Sci.*, vol. 4, no. 1, pp. 247–256, 2024, doi: 10.57152/malcom.v4i1.1066.
- [22] F. T. Admojo and Ahsanawati, “Klasifikasi Aroma Alkohol Menggunakan Metode KNN,” *Indones. J. Data Sci.*, vol. 1, no. 2, pp. 34–38, 2020, doi: 10.33096/ijodas.v1i2.12.
- [23] Q. A. A’yuniyah and M. Reza, “Penerapan Algoritma K-Nearest Neighbor Untuk Klasifikasi Jurusan Siswa Di Sma Negeri 15 Pekanbaru,” *Indones. J. Inform. Res. Softw. Eng.*, vol. 3, no. 1, pp. 39–45, 2023, doi: 10.57152/ijirse.v3i1.484.
- [24] S. A. Utiahman and A. M. M. Pratama, “Analisis Perbandingan KNN, SVM, Decision Tree dan Regresi Logistik Untuk Klasifikasi Obesitas Multi Kelas,” *KLIK Kaji. Ilm. Inform. dan Komput.*, vol. 4, no. 6, pp. 3137–3146, 2024, doi: 10.30865/klik.v4i6.1871.
- [25] H. F. Putro, R. T. Vulandari, and W. L. Y. Saptomo, “Penerapan Metode Naive Bayes Untuk Klasifikasi Pelanggan,” *J. Teknol. Inf. dan Komun.*, vol. 8, no. 2, 2020, doi: 10.30646/tikomsin.v8i2.500.
- [26] M. Yusuf, “Analisis Sentimen Data Twitter Terhadap Bakal Calon Presiden Republik Indonesia 2024 Dengan Metode Backpropagation,” 2022, [Online]. Available: <http://repo.palcomtech.ac.id/id/eprint/1662/>
- [27] F. Aziz, P. Ishak, and S. Abasa, “Klasifikasi Depresi Menggunakan Support Vector Machine: Pendekatan Berbasis Data Text Mining,” *J. Pharm. Appl. Comput. Sci.*, vol. 2, no. 2, pp. 33–38, 2024, doi: 10.59823/jopacs.v2i2.53.
- [28] S. A. Pratiwi, A. Fauzi, S. Arum, P. Lestari, and Y. Cahyana, “KLIK: Kajian Ilmiah Informatika dan Komputer Prediksi Persediaan Obat Pada Apotek Menggunakan Algoritma Decision Tree,” *Media Online*, vol. 4, no. 4, pp. 2381–2388, 2024, doi: 10.30865/klik.v4i4.1681.
- [29] D. P. Wijaya, L. D. Murti, M. R. Rachman, D. Arsip, and K. Bandung, “Recall dan Precision pada Online Public Access Catalog (OPAC) Dinas Arsip dan Perpustakaan Kota Bandung Didik,” vol. 24, no. 1, 2022.
- [30] M. Pirnau *et al.*, “Content Analysis Using Specific Natural Language Processing Methods for Big Data,” *Electron.*, vol. 13, no. 3, pp. 1–22, 2024, doi: 10.3390/electronics13030584.
- [31] M. Iqbal, A. Davy Wiranata, R. Suwito, and R. Faiz Ananda, “Perbandingan Algoritma Naïve Bayes, KNN, dan Decision Tree terhadap Ulasan Aplikasi Threads dan Twitter,” *Media Online*, vol. 4, no. 3, pp. 1799–1807, 2023, doi: 10.30865/klik.v4i3.1402.

