

PENERAPAN ALGORITMA NAIVE BAYES UNTUK KLASIFIKASI SENTIMEN ULASAN PENGGUNA APLIKASI PIDI UMKM DENGAN PEMBOBOTAN FITUR TF-IDF

Abdulah Haris, Tukino, Agustia Hananto, Fitria nurapriani

Sistem Informasi, Universitas Buana Perjuangan Karawang

Jl. Ronggo Waluyo Sinarbaya, Puseurjaya, Telukjambe Timur, Karawang, Jawa Barat, Indonesia

Si22.abdulahharis@mhs.ubpkarawang.ac.id

ABSTRAK

Seiring pesatnya pertumbuhan aplikasi digital, ulasan pengguna menjadi sumber informasi berharga untuk memahami sentimen dan meningkatkan kualitas layanan. Aplikasi Pidi UMKM, sebagai platform yang mendukung usaha mikro, kecil, dan menengah, memerlukan analisis mendalam terhadap ulasan penggunanya untuk mengidentifikasi perbaikan area dan fitur yang diminati. Permasalahan yang dihadapi adalah volume ulasan yang beragam sehingga sulit untuk dijelaskan secara manual, menghambat menghilangkan sentimen secara efisien. Tujuan penelitian ini adalah menerapkan metode klasifikasi Naive Bayes untuk mengotomatisasi proses analisis sentimen ulasan pengguna aplikasi Pidi UMKM. Metode yang digunakan meliputi tahap pengumpulan data ulasan sebanyak 103 data, pra-pemrosesan teks (seperti tokenizing, stop word, dan stemming), pembentukan fitur vektor menggunakan TF-IDF, dan selanjutnya klasifikasi menggunakan algoritma Naive Bayes. Hasil penelitian ini menunjukkan bahwa metode Naive Bayes mampu mengklasifikasikan ulasan dengan akurasi sebesar 65% dalam membedakan sentimen positif, negatif, dan netral, sehingga memberikan wawasan yang signifikan bagi pengembang aplikasi Pidi UMKM untuk membuat keputusan berbasis data.

Kata kunci: Pidi UMKM, Klasifikasi, Naive Bayes, Analisis Sentimen, Ulasan, Pembelajaran Mesin

1. PENDAHULUAN

Di era digital yang kompetitif ini, aplikasi Pidi UMKM hadir sebagai platform strategis dalam mendukung pertumbuhan dan pengembangan usaha mikro, kecil, dan menengah. Dengan peningkatan jumlah pengguna secara signifikan, aplikasi ini menerima volume ulasan yang semakin beragam. Ulasan tersebut tidak hanya mencerminkan opini, tetapi juga persepsi dan pengalaman nyata dari pengguna. Informasi ini sangat berharga untuk memahami sentimen pengguna, mengidentifikasi kekuatan maupun kelemahan aplikasi, serta menyusun strategi pengembangan berdasarkan kebutuhan riil. Dalam konteks persaingan yang semakin ketat di ranah aplikasi digital, kemampuan untuk memahami dan merespons sentimen pengguna secara cepat dan tepat menjadi kunci untuk mempertahankan eksistensi dan mendorong pertumbuhan aplikasi Pidi UMKM.

Namun demikian, mengelola dan menganalisis ribuan hingga puluhan ribu ulasan secara manual merupakan pekerjaan yang tidak hanya memakan waktu dan tenaga, tetapi juga sangat rentan terhadap bias interpretasi manusia[1]. Hambatan ini berdampak langsung pada keterlambatan menghilangkan masalah, ketidaktepatan dalam memahami ekspektasi pengguna, serta potensi kesalahan dalam pengambilan keputusan. Oleh karena itu, diperlukan sebuah metode otomatis yang mampu mengekstrak informasi penting dari ulasan pengguna dan secara akurat mengklasifikasikan sentimennya[2].

Analisis sentimen atau klasifikasi opini merupakan salah satu pendekatan dalam mesin pembelajaran yang dirancang untuk mengenali polaritas emosi baik positif, negatif, maupun netral

dari suatu teks. Di antara berbagai algoritma yang tersedia, Naive Bayes merupakan metode yang populer dan terbukti efektif untuk klasifikasi teks. Algoritma ini dikenal karena kemudahannya, efisiensinya dalam proses komputasi, serta kemampuannya dalam menghasilkan kinerja yang baik pada data teks berdimensi tinggi. Penerapannya dalam klasifikasi ulasan aplikasi telah memberikan kontribusi yang signifikan dalam mengungkap pola sentimen dan kecenderungan opini pengguna, sehingga memberikan dasar yang kuat bagi pengambilan keputusan pengembangan produk[3].

Beberapa penelitian terdahulu juga telah membuktikan keefektifan metode klasifikasi dalam konteks analisis sentimen ulasan aplikasi. Misalnya, penelitian yang dilakukan oleh Aini dkk. Menunjukkan bahwa sebagian besar pengguna aplikasi Dana memberikan ulasan dengan sentimen positif[4]. Selain itu, berbagai penelitian lainnya juga menunjukkan keberhasilan penerapan algoritma pembelajaran mesin seperti Naive Bayes dalam mengekstraksi dan mengidentifikasi sentimen pada ulasan produk digital. Temuan-temuan ini memperkuat potensi besar dari pendekatan data mining dalam menggali wawasan strategi dari data ulasan pengguna yang dapat dimanfaatkan untuk evaluasi kinerja aplikasi serta merancang perbaikan secara tepat sasaran.

Mengingat kompleksitas serta volume besar ulasan pengguna di aplikasi Pidi UMKM, penelitian ini bertujuan untuk menerapkan metode Naive Bayes dalam proses klasifikasi sentimen ulasan. Hasil dari klasifikasi ini diharapkan dapat memberikan solusi otomatis dan efisien dalam mengekstraksi wawasan

dari opini pengguna, sehingga mendukung proses pengambilan keputusan berdasarkan data dalam pengembangan aplikasi Pidi UMKM ke depannya..

2. TINJAUAN PUSTAKA

2.1. Data Mining dan Analisis Sentimen

Data mining adalah proses penemuan pola atau pengetahuan yang tersembunyi dari data dalam jumlah besar yang berguna untuk mendukung pengambilan keputusan[5]. Dalam dunia bisnis aplikasi digital, data mining sering digunakan untuk menganalisis tren penggunaan, perilaku konsumen, dan evaluasi produk atau layanan. Salah satu implementasi penting dari data mining adalah analisis sentimen (juga dikenal sebagai *opinion mining*), yaitu proses pengelompokan data teks ke dalam beberapa kategori berdasarkan polaritas emosi (positif, negatif, atau netral). Dalam konteks ulasan aplikasi, sentimen digunakan untuk mengelompokkan ulasan berdasarkan sentimen yang dijelaskan pengguna terhadap fitur atau kinerja aplikasi[6]. Dengan mengetahui sentimen pengguna, pengembang aplikasi dapat menentukan strategi peningkatan kualitas dan manajemen fitur secara lebih tepat.

2.2. Algoritma Naïve Bayes

Naïve Bayes adalah algoritma klasifikasi probabilistik yang bekerja berdasarkan Teorema Bayes dengan asumsi independensi antar fitur. Algoritma ini termasuk dalam kelompok *supervised learning* dan telah banyak digunakan untuk berbagai keperluan klasifikasi, seperti klasifikasi teks, deteksi spam, hingga klasifikasi citra. Keunggulan *Naïve Bayes* terletak pada efisiensi implementasi, efisiensi komputasi, dan kinerja yang baik pada dataset berukuran kecil maupun besar. Penelitian oleh Afriansyah dkk. menerapkan *Naïve Bayes* untuk mengklasifikasikan rasa buah Pisang Raja berdasarkan fitur warna RGB dan berhasil mencapai akurasi rata-rata sebesar 86,66%[7]. Hal ini menunjukkan bahwa *Naïve Bayes* memiliki potensi besar untuk diterapkan dalam konteks klasifikasi sentimen ulasan aplikasi.

2.3. Fitur Representasi dan Ekstraksi Data Teks Ulasan

Dalam proses klasifikasi, fitur representasi sangat menentukan akurasi model. Pada data berbasis teks atau naratif seperti ulasan pengguna, teknik representasi fitur seperti TF-IDF (Term Frekuensi – Inverse Document Frekuensi) digunakan untuk mengekstraksi bobot dari setiap kata berdasarkan frekuensi kemunculan dalam dokumen dan keseluruhan korpus[8]. Penggunaan TF-IDF banyak diaplikasikan pada klasifikasi dokumen dan telah terbukti efektif dalam mengidentifikasi kata-kata penting untuk proses klasifikasi[9]. Data ulasan yang terdiri dari teks opini pengguna dapat diubah menjadi representasi numerik menggunakan TF-IDF, sehingga

dapat dimasukkan ke dalam algoritma klasifikasi seperti *Naïve Bayes* untuk dilakukan pelatihan model.

2.4. Penelitian Terkait

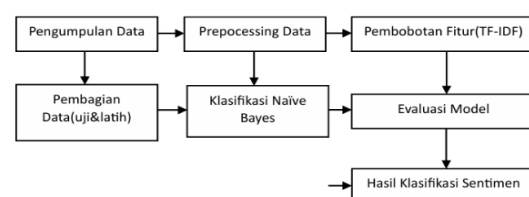
Berbagai penelitian terdahulu telah menerapkan kombinasi teknik data mining dan algoritma *Naïve Bayes* dalam klasifikasi teks. Widodo dkk. melakukan klasifikasi rasa buah jeruk peras menggunakan sensor tegangan dan diameter buah, menghasilkan akurasi hingga 80%[10]. Subagja dkk. Algoritma klasifikasi yang digunakan pada penelitian ini adalah *Naïve Bayes*, Penelitian ini menghasilkan nilai akurasi sebesar 57%[11]. Pada bidang kesehatan, *Naïve Bayes* juga berhasil digunakan dalam klasifikasi penyakit berdasarkan gejala pasien. Meskipun telah banyak digunakan dalam berbagai domain, penerapan algoritma *Naïve Bayes* untuk klasifikasi sentimen ulasan aplikasi Pidi UMKM secara spesifik masih sangat terbatas. Oleh karena itu, penelitian ini hadir untuk mengisi kekosongan tersebut dengan pendekatan klasifikasi sentimen ulasan berbasis data mining, guna mendukung efisiensi analisis dan evaluasi strategi aplikasi.

3. METODE PENELITIAN

Bagian ini menjelaskan metodologi yang digunakan dalam penelitian, mulai dari tahapan pengumpulan data, prapemrosesan data, hingga implementasi algoritma klasifikasi *Naïve Bayes*. Pendekatan penelitian ini bertujuan untuk mengklasifikasikan sentimen ulasan pengguna aplikasi Pidi UMKM.

3.1. Tahapan Penelitian

Tahapan penelitian ini mengadopsi kerangka umum dalam proses data mining dan machine learning untuk klasifikasi teks, yang dijelaskan pada Gambar 1.



Gambar 1. Tahapan Penelitian

3.2. Pengumpulan Data

Data yang digunakan dalam penelitian ini adalah ulasan pengguna aplikasi Pidi UMKM yang diperoleh dari platform Google Play. Setiap ulasan data berupa teks yang berisikan komentar, rating, dan pengalaman pengguna terkait aplikasi Pidi UMKM. Seluruh data ulasan ini kemudian dilabeli secara manual ke dalam tiga kategori sentimen: positif, negatif, dan netral, yang menjadi dasar untuk pelatihan dan pengujian model klasifikasi.

Tabel 1 . Data Ulasan Asli

Username	Rating	Review Text	Date
Ahmad Nursyamsi	5	Susah belanja nya	2025-05-21T03:50:05
Digital Trader	5	sistem akulaku patut di tiru. Memberikan bunga yg rendah kepada usernya & merchant di berbagai wilayah seperti jaringan sumber daya yg tak terbatas.	2025-05-20T04:59:34
Moch Jamhari	5	ayo di tingkatkan lagi dengan UI system yang simple tidak berat	2025-05-02T10:53:38
Ronni Abdurrahman	1	tidak bisa chat penjual, ktika di chat langsung keluar aplikasi	2025-04-24T03:31:16
Dafa Hakim	5	Oke	2025-03-08T10:11:33

3.3. Data Prapemrosesan (Pemrosesan Awal)

Data ulasan yang masih mentah seringkali mengandung noise atau informasi yang tidak relevan

untuk proses klasifikasi. Oleh karena itu, diperlukan tahapan prapemrosesan untuk membersihkan dan menyiapkan data.

Tabel 2. Proses Preprocessing Data

Cleansing	Case_folding	Normalisasi	Tokenize	Stopword	Stemming
Susah belanja nya	Susah belanja nya	susah belanja ya	['susah', 'belanja', 'ya']	['susah', 'belanja', 'ya']	susah belanja ya
sistem akulaku patut di tiru Memberikan bunga yg rendah kepada usernya merchant di berbagai wilayah seperti jaringan sumber daya yg tak terbatas	sistem akulaku patut di tiru memberikan bunga yg rendah kepada usernya merchant di berbagai wilayah seperti jaringan sumber daya yg tak terbatas	sistem akulaku patut di tiru memberikan bunga yang rendah kepada usernya merchant di berbagai wilayah seperti jaringan sumber daya yang tak terbatas	['sistem', 'akulaku', 'patut', 'di', 'tiru', 'memberikan', 'bunga', 'yang', 'rendah', 'kepada', 'usernya', 'merchant', 'di', 'berbagai', 'wilayah', 'seperti', 'jaringan', 'sumber', 'daya', 'yang', 'tak', 'terbatas']	['sistem', 'akulaku', 'patut', 'tiru', 'bunga', 'rendah', 'usernya', 'merchant', 'wilayah', 'jaringan', 'sumber', 'daya', 'terbatas']	sistem akulaku patut tiru bunga rendah usernya merchant wilayah jaring sumber daya batas
ayo di tingkatkan lagi dengan UI system yang simple tidak berat	ayo di tingkatkan lagi dengan ui system yang simple tidak berat	ayo di tingkatkan lagi dengan ui system yang simple tidak berat	['ayo', 'di', 'tingkatkan', 'lagi', 'dengan', 'ui', 'system', 'yang', 'simple', 'tidak', 'berat']	['ayo', 'tingkatkan', 'ui', 'system', 'simple', 'berat']	ayo tin gkat ui system simple berat
tidak bisa chat penjual ktika di chat langsung keluar aplikasi	tidak bisa chat penjual ktika di chat langsung keluar aplikasi	tidak bisa chat penjual ketika di chat langsung keluar aplikasi	['tidak', 'bisa', 'chat', 'penjual', 'ketika', 'di', 'chat', 'langsung', 'keluar', 'aplikasi']	['chat', 'penjual', 'chat', 'langsung', 'aplikasi']	chat jual chat langsung aplikasi
Oke	Oke	oke	['oke']	['oke']	oke

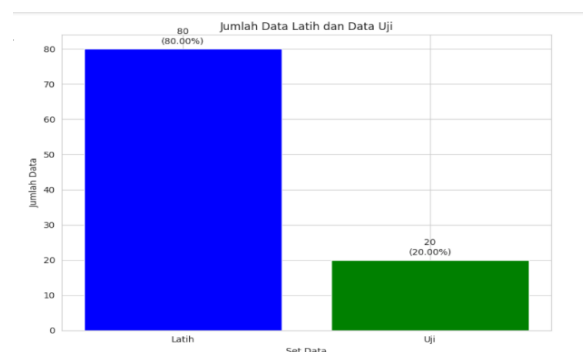
3.4. Fitur Pembobotan (TF-IDF)

Setelah peninjauan data melalui proses prapemrosesan, tahapan selanjutnya adalah pembobotan fitur menggunakan metode TF-IDF (Term Frekuensi – Inverse Document Frekuensi). TF-IDF digunakan untuk mengubah data teks menjadi representasi numerik yang dapat dipahami oleh pembelajaran mesin algoritma. Metode ini memberikan pada setiap kata (term) dalam setiap dokumen (ulasan) berdasarkan seberapa sering kata tersebut muncul dalam dokumen (Term Frekuensi) bobot dan seberapa jarang kata tersebut muncul dalam seluruh kumpulan dokumen (Inverse Document Frekuensi).

3.5. Pembagian Data

Data ulasan yang telah diolah dan dibobotkan kemudian dibagi menjadi dua subset, yaitu data pelatihan dan data pengujian, sebagaimana ditunjukkan pada Gambar 2. Pengujian data ini tidak pernah digunakan selama proses pelatihan model, sehingga dapat memberikan gambaran tujuan

mengenai kemampuan generalisasi model terhadap data baru. Pembagian data dengan rasio 80% untuk pelatihan dan 20% untuk pengujian merupakan praktik umum dalam mesin pembelajaran untuk memastikan bahwa model tidak hanya mampu mengenali pola pada data pelatihan, tetapi juga memiliki kinerja yang baik saat memprediksi data yang belum pernah dilihat sebelumnya.

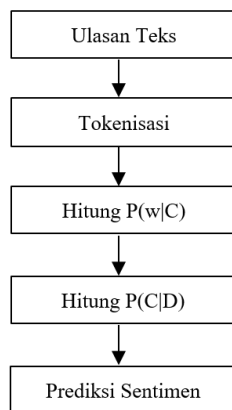


Gambar 2. Data Latih dan Data Uji

3.6. Klasifikasi Naive Bayes

Algoritma Naive Bayes akan diterapkan pada data yang telah siap. Naive Bayes bekerja berdasarkan prinsip probabilitas Bayes, dengan asumsi independensi antar fitur (kata). Model akan menghitung probabilitas kemunculan setiap kata untuk setiap kelas sentimen (positif, negatif, netral) dari data pelatihan. Selanjutnya, pada saat pengujian, model akan memprediksi sentimen kelas untuk setiap ulasan baru berdasarkan probabilitas gabungan dari kata-kata yang terkandung di dalamnya.

Untuk menggambarkan alur pemrosesan data dengan algoritma Naive Bayes, berikut disajikan flowchart proses klasifikasi:



Gambar 3. Alur Klasifikasi Naive Bayes

3.7. Model Evaluasi

Untuk mengukur kinerja model klasifikasi Naive Bayes, dilakukan evaluasi menggunakan metrik yang umum dalam klasifikasi, seperti:

- a. Akurasi (Akurasi): Rasio jumlah prediksi benar akan diukur untuk menunjukkan seberapa tepat model mengklasifikasikan sentimen ulasan

$$\frac{TP + TN}{TP + TN + FP + FN}$$

Dimana:

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

- b. Presisi: Presisi menunjukkan proporsi prediksi positif yang benar-benar relevan.

Rumus:

$$\text{Presisi} = \frac{TP}{TP + FP}$$

- c. Recall: ukuran seberapa baik model dapat menemukan semua data positif.

Rumus:

$$\text{Recall} = \frac{TP}{TP + FN}$$

- d. F1-Score rata-rata harmonik dari presisi dan recall, digunakan untuk menyeimbangkan keduanya.

Rumusnya:

$$F1 - \text{Score} = 2X \frac{\text{Presisi} \times \text{Recall}}{\text{Presisi} + \text{Recall}}$$

4. HASIL DAN PEMBAHASAN

4.1. Hasil Analisis Ulasan Pengguna

Data ulasan pengguna yang diperoleh dari Google Play Store selama periode 2025 menunjukkan pola menarik antara rating yang diberikan dan isi ulasan (Tabel 3). Sebagian besar pengguna memberikan skor tertinggi (rating 5), tetapi isi komentarnya tidak selalu mencerminkan sentimen positif. Contohnya, komentar “susah belanja ya” yang ditulis oleh pengguna dengan rating 5, mengindikasikan adanya pengalaman negatif, meskipun rating yang diberikan menunjukkan sebaliknya.

Begitu pula komentar seperti “sistem akulaku patut tiru bunga rendah...” atau “ayo tingkat ui system” bersifat ambigu dan fragmentaris, sehingga menyulitkan algoritma untuk mengklasifikasikan sentimen secara tepat. Kondisi ini menampilkan adanya pemetaan semantik antara angka rating dan opini verbal yang diberikan pengguna.

Tabel 3. Hasil Preprocessing

Date	Username	Rating	Review Text
2025-05-21T03:50:05	Ahmad Nursyamsi	5	susah belanja ya
2025-05-20T04:59:34	Digital Trader	5	sistem akulaku patut tiru bunga rendah usernya merchant wilayah jaring sumber daya batas
2025-05-02T10:53:38	Moch Jamhari	5	ayo tingkat ui system simple berat
2025-04-24T03:31:16	Ronni Abdurrahman	1	chat jual chat langsung aplikasi
2025-03-08T10:11:33	Dafa Hakim	5	oke

4.2. Implementasi Algoritma Naive Bayes

Model Naive Bayes diimplementasikan untuk melakukan klasifikasi sentimen terhadap ulasan pengguna aplikasi Pidi UMKM. Implementasi proses dimulai dengan tahapan preprocessing data, seperti pembersihan teks (cleansing), tokenisasi, stopword removal, dan normalisasi kata. Setelah data dibersihkan, dilakukan proses pembobotan fitur menggunakan metode TF-IDF (Term Frekuensi-

Inverse Document Frekuensi) untuk mengubah teks ulasan menjadi representasi numerik.

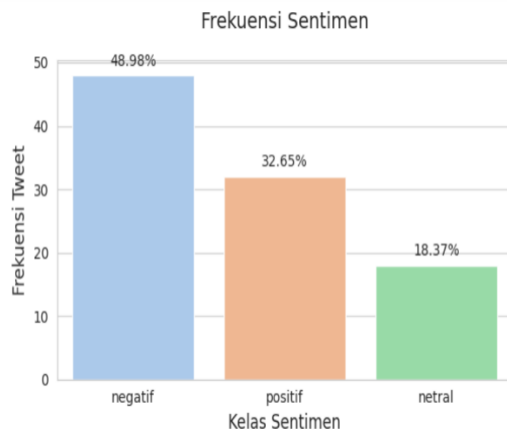
Setelah data siap, dilakukan pembagian data menjadi data latih dan data uji, biasanya dalam proporsi 80:20. Selanjutnya, data latih digunakan oleh model Naive Bayes untuk menghitung probabilitas kemunculan kata tertentu pada masing-masing kelas sentimen (positif, negatif, dan netral). Perhitungan ini dilakukan berdasarkan prinsip teorema Bayes, dengan

asumsi bahwa setiap kata bersifat independen satu sama lain.

Setelah proses pelatihan selesai, model diuji dengan data uji yang belum pernah dilihat sebelumnya. Untuk setiap ulasan pada uji data, model perhitungan:

- Probabilitas posterior setiap kelas sentimen berdasarkan kata-kata yang muncul, lalu
- Memilih kelas dengan probabilitas tertinggi sebagai hasil klasifikasi.

Arsitektur penerjemahan ini divisualisasikan pada Gambar 4.



Gambar 4. Frekuensi Hasil Klasifikasi Sentimen

yaitu diagram batang frekuensi hasil klasifikasi sentimen. Gambar tersebut menunjukkan bahwa:

- Sentimen negatif mendominasi ulasan dengan proporsi sebesar 48.98%,
- Diikuti oleh sentimen positif sebesar 32.65%, dan
- Sentimen netral sebesar 18,37%.

Distribusi ini memberikan indikasi bahwa sebagian besar pengguna menyampaikan keluhan atau pengalaman negatif terhadap aplikasi Pidi UMKM, sehingga temuan ini dapat menjadi dasar dalam perumusan strategi perbaikan aplikasi ke depan..

4.3. Model Hasil Akurasi

Klasifikasi model kinerja evaluasi dilakukan untuk mengetahui seberapa baik algoritma Naive Bayes dalam memprediksi sentimen ulasan. Metrik evaluasi utama yang digunakan adalah akurasi, presisi, recall, dan f1-score. Hasil evaluasi kinerja model disajikan dalam Classification Report pada gambar 5.

```

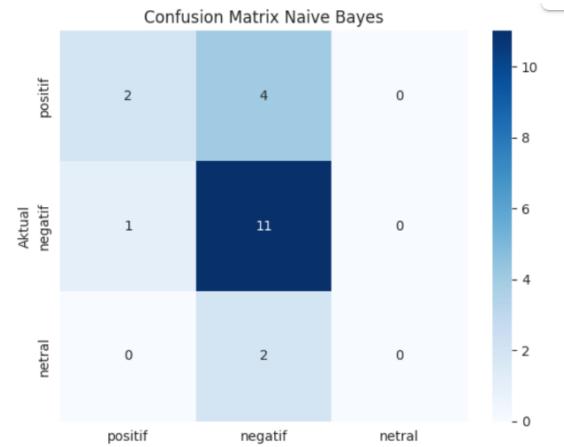
=== Classification Report Naive Bayes ===

```

	precision	recall	f1-score	support
negatif	0.67	1.00	0.80	12
netral	0.00	0.00	0.00	2
positif	1.00	0.33	0.50	6
accuracy			0.70	20
macro avg	0.56	0.44	0.43	20
weighted avg	0.70	0.70	0.63	20

Gambar 5. Clasification Refort

Untuk memberikan pemahaman yang lebih mendalam mengenai kinerja klasifikasi pada setiap kelas sentimen, Confusion Matrix disajikan pada Gambar 6.



Gambar 6. Clasification Matric

Berdasarkan Gambar 5 dan Gambar 6, model klasifikasi Naive Bayes yang dibangun berhasil mencapai akurasi sebesar 65% pada pengujian data (20 data). Nilai akurasi ini menunjukkan bahwa 65% dari total ulasan pada data pengujian berhasil diklasifikasikan dengan benar oleh model.

Meskipun akurasi yang dicapai cukup baik, terdapat beberapa catatan penting dari hasil evaluasi model. Pada kelas "netral", model tidak berhasil melakukan klasifikasi dengan baik, yang ditunjukkan oleh nilai precision, recall, dan f1-score sebesar 0.00. Hal ini mengindikasikan bahwa model kesulitan dalam mengenali pola pada kelas tersebut. Di sisi lain, model menunjukkan kinerja yang lebih baik pada kelas "negatif" dengan nilai f1-score sebesar 0.76 dan recall yang tinggi (0.92), yang berarti model mampu mengidentifikasi sebagian besar ulasan negatif dengan benar.

Kelas "positif" memiliki nilai precision yang cukup baik (0.67), tetapi recall yang rendah (0.33), menunjukkan bahwa model sering kali melewatkan ulasan positif yang seharusnya terdeteksi. Dengan demikian, meskipun akurasi secara keseluruhan mencapai 65%, masih terdapat ruang untuk perbaikan, terutama dalam meningkatkan kinerja model pada kelas "netral" dan "positif".

5. KESIMPULAN DAN SARAN

Penelitian ini berhasil mengimplementasikan algoritma Naive Bayes dengan pembobotan TF-IDF untuk klasifikasi sentimen ulasan aplikasi Pidi UMKM, mencapai akurasi 60% pada 20 data pengujian, meskipun variasi ditemukan kinerja antar kelas sentimen (terutama rendah pada kelas netral karena keterbatasan data) yang dipengaruhi oleh jumlah data, ke akurasi kelas, karakteristik bahasa, dan jumlah fitur unik.

Untuk meningkatkan kinerja model klasifikasi sentimen di masa mendatang, disarankan untuk menambah volume dan keragaman dataset (terutama untuk kelas minoritas), menerapkan teknik penanganan ke ukuran data, mengeksplorasi metode rekayasa fitur yang lebih canggih seperti N-gram atau word embedding, serta membandingkan kinerja dengan algoritma klasifikasi lain seperti SVM atau deep learning..

DAFTAR PUSTAKA

- [1] Z. Kastrati, A. S. Imran, And A. Kurti, “Weakly Supervised Framework For Aspect-Based Sentiment Analysis On Students’ Reviews Of Moocs,” *Ieee Access*, Vol. 8, Pp. 106799–106810, 2020, Doi: 10.1109/Access.2020.3000739.
- [2] R. Obiedat *Et Al.*, “Sentiment Analysis Of Customers’ Reviews Using A Hybrid Evolutionary Svm-Based Approach In An Imbalanced Data Distribution,” *Ieee Access*, Vol. 10, Pp. 22260–22273, 2022, Doi: 10.1109/Access.2022.3149482.
- [3] R. Raman, V. Kumar, B. G. Pillai, D. Rabadiya, P. K. Bhoyar, And N. Kumbhojkar, “Analyzing Sentiments In Motorku X App Reviews Using The Naive Bayes Classifier Approach,” In *2024 Second International Conference On Data Science And Information System (Icdsis)*, Ieee, May 2024, Pp. 1–5. Doi: 10.1109/Icdsis61070.2024.10594505.
- [4] A. Nurian And B. Nurina Sari, “Analisis Sentimen Ulasan Pengguna Aplikasi Google Play Menggunakan Naïve Bayes,” *Jurnal Informatika Dan Teknik Elektro Terapan*, Vol. 11, No. 3, Pp. 2830–7062, Doi: 10.23960/Jitet.V11i3%20s1.3348.
- [5] J. Olufemi Ogunleye, “The Concept Of Data Mining,” 2022. Doi: 10.5772/Intechopen.99417.
- [6] A. Verma And R. Vashisth, “Sentiment Analysis Of Google Play And App Store Reviews Of Threads,” In *2024 11th International Conference On Reliability, Infocom Technologies And Optimization (Trends And Future Directions) (Icrito)*, Ieee, Mar. 2024, Pp. 1–5. Doi: 10.1109/Icrito61523.2024.10522237.
- [7] M. Afriansyah, J. Saputra, V. Yoga Pudya Ardhana, Y. Sa, And U. Qamarul Huda Badaruddin, “Algoritma Naive Bayes Yang Efisien Untuk Klasifikasi Buah Pisang Raja Berdasarkan Fitur Warna,” *Hal. 236 Journal Of Information Systems Management And Digital Business (Jismdb)*, Vol. 1, No. 2, 2024.
- [8] J. E. Br Sinulingga And H. C. K. Sitorus, “Analisis Sentimen Opini Masyarakat Terhadap Film Horor Indonesia Menggunakan Metode Svm Dan Tf-Idf,” *Jurnal Manajemen Informatika (Jamika)*, Vol. 14, No. 1, Pp. 42–53, Feb. 2024, Doi: 10.34010/Jamika.V14i1.11946.
- [9] Y. Sergio, V. Putranta, B. Rahayudi, And W. Purnomo, “Analisis Sentimen Masyarakat Terhadap Kebijakan Penghapusan Subsidi Bbm Pada Media Sosial Twitter Menggunakan Algoritma Naive Bayes Classifier Dengan Ekstraksi Fitur N-Gram Tf-Idf,” 2023. [Online]. Available: [Http://J-Ptiik.Ub.Ac.Id](http://J-Ptiik.Ub.Ac.Id)
- [10] P. Sistem *Et Al.*, “Menggunakan Diameter Buah Dan Nilai Sensor Tegangan Berbasis Arduino Uno Dengan Metode Naïve Bayes,” 2023. [Online]. Available
- [11] R. A. Subagja *Et Al.*, “Klasifikasi Ulasan Aplikasi Jenius Pada Google Play Store Menggunakan Algoritma Naive Bayes,” Vol. 3, P. 2021