

Comparison of the Performance of K-Nearest Neighbor and Naive Bayes Algorithms for Sentiment Analysis of PinjamYuk Application User Reviews Using SMOTE and TF-IDF

Lindya Rossita Handoko ¹, Amelia Faza ², Mula Agung Barata ³, Afril Efan Pajri ⁴

Teknik Informatika, Universitas Nahdlatul Ulama Sunan Giri Bojonegoro

lindiarosita343@gmail.com ¹, fazaa4928@gmail.com ², mula.ab26@gmail.com ³, afril@unugiri.ac.id ⁴

Article Info

Article history:

Received 2026-06-19

Revised 2026-07-24

Accepted 2026-08-12

Keyword:

*K-Nearest Neighbor,
Naive Bayes,
Sentiment Analysis,
SMOTE.*

ABSTRACT

The growth of online lending services has driven the increasing adoption of digital financial applications, providing users with convenient access to financial services. This growing number of users has generated a large volume of reviews on the Google Play Store, which can serve as a valuable source for understanding users' perspectives on the benefits and performance of these applications. This study aims to compare the performance of the K-Nearest Neighbor (KNN) and Naive Bayes algorithms in classifying the sentiment of user reviews of the PinjamYuk application. The study uses a secondary dataset obtained from Kaggle, consisting of 500 user reviews of the PinjamYuk application on the Google Play Store during the 2023–2024 period. The reviews were categorized into three sentiment classes: positive, neutral, and negative, based on their rating scores. Because the class distribution in the dataset was imbalanced, the Synthetic Minority Oversampling Technique (SMOTE) was applied to balance the classes before the classification process. The research procedure included data preprocessing, feature extraction using Term Frequency–Inverse Document Frequency (TF-IDF), data balancing using SMOTE, and KNN parameter optimization using GridSearchCV. The models were evaluated using accuracy, precision, recall, F1-score, confusion matrix, stratified K-fold cross-validation, and the McNemar test. The results show that the Naive Bayes algorithm outperformed KNN. The stratified K-fold cross-validation results yielded an average accuracy of 81.0% for Naive Bayes and 80.8% for KNN. Furthermore, the McNemar test produced a p-value of 0.014 ($p < 0.05$), indicating that the performance difference between the two algorithms was statistically significant. These findings demonstrate that the Naive Bayes algorithm is more effective for analyzing user sentiment in reviews of the PinjamYuk application.



This is an open access article under the [CC-BY-SA](https://creativecommons.org/licenses/by-sa/4.0/) license.

I. PENDAHULUAN

Perkembangan teknologi digital telah mengubah berbagai aspek kehidupan masyarakat, termasuk dalam sektor keuangan. Perkembangan teknologi informasi telah menciptakan beragam inovasi dalam layanan keuangan digital yang lebih mudah dijangkau oleh masyarakat luas [1]. Layanan ini menyediakan proses pencairan dana yang cepat dengan syarat yang cukup sederhana sehingga menjadi

pilihan bagi masyarakat yang memerlukan dana secara cepat. Fenomena peminjaman online semakin berkembang di masyarakat, tidak hanya sebagai solusi cepat untuk keuangan tetapi juga menggambarkan krisis dalam literasi finansial dan pengendalian diri [2]. Hal ini mendorong perkembangan pesat platform pinjaman daring di Indonesia sehingga semakin banyak orang menggunakan layanan ini untuk memenuhi kebutuhan keuangan.

Layanan pinjaman online yang semakin meningkat mendorong pengguna untuk memberikan berbagai ulasan terhadap aplikasi, baik berupa komentar maupun penilaian yang mencerminkan pengalaman mereka selama menggunakan layanan. Ulasan pengguna dapat dimanfaatkan sebagai sumber informasi bagi pengembang untuk melakukan evaluasi terhadap kualitas layanan maupun aplikasi [3]. Namun, jumlah ulasan yang terus bertambah menyebabkan proses analisis manual menjadi kurang efektif, sehingga diperlukan pendekatan analisis sentimen untuk mengolah data ulasan secara otomatis dan memperoleh informasi yang lebih komprehensif [4]. Oleh karena itu, analisis sentimen merupakan salah satu metode yang dapat digunakan untuk memahami opini publik mengenai kualitas aplikasi pinjaman online.

Oleh karena itu, salah satu cara untuk mengetahui opini publik tentang kualitas aplikasi pinjaman online adalah dengan menggunakan analisis sentimen. Karena jumlah orang yang menggunakan aplikasi pinjaman online meningkat, semakin banyak data yang dikumpulkan, sehingga analisis manual menjadi kurang efektif. Salah satu komponen penambangan teks, atau penambangan teks, adalah analisis sentimen, yang dapat digunakan untuk menganalisis, mengumpulkan, dan membagi pendapat pengguna ke dalam kategori tertentu, seperti positif, negatif, dan netral [5]. Teknik ini memungkinkan analisis opini dilakukan secara otomatis pada kumpulan data teknis yang besar. Analisis sistematis mengenai kekhawatiran masyarakat terkait layanan pinjaman daring sangatlah penting, karena hal ini berperan krusial bagi perlindungan konsumen dan pengembangan kebijakan keuangan yang lebih responsive [6]. Selain itu, karakteristik ulasan aplikasi pinjaman online yang umumnya menggunakan bahasa informal, singkatan, kata tidak baku, bahkan kalimat bernada sarkastik menjadikan proses klasifikasi sentimen lebih menantang sehingga diperlukan metode klasifikasi yang sesuai.

Berdasarkan penelitian sebelumnya, algoritma K-Nearest Neighbor (KNN) telah banyak dimanfaatkan untuk menangani permasalahan klasifikasi dan prediksi di sektor keuangan [7]. Penelitian tentang prediksi kredit bermasalah menunjukkan bahwa algoritma KNN berhasil mencapai tingkat akurasi 93,14% dan menunjukkan kinerja yang superior dibandingkan Naive Bayes dalam mengenali nasabah yang berpotensi mengalami kredit macet [8]. Di samping itu, penelitian mengenai data status kredit nasabah Bank Perkreditan Rakyat (BPR) menunjukkan bahwa penggunaan *Synthetic Minority Oversampling Technique* (SMOTE) pada algoritma KNN dapat meningkatkan kemampuan klasifikasi data tidak seimbang dengan cara meningkatkan nilai sensitivitas dan G-mean [9]. Berdasarkan hasil tersebut, algoritma KNN dan Naive Bayes dipilih sebagai model perbandingan karena keduanya adalah algoritma klasifikasi yang mudah, memiliki kompleksitas komputasi yang relatif rendah, mudah diimplementasikan, serta banyak digunakan sebagai *baseline* pada penelitian

analisis sentimen. Sementara itu, algoritma seperti *Support Vector Machine* (SVM), Random Forest, maupun IndoBERT tidak digunakan karena penelitian ini berfokus pada evaluasi performa model klasifikasi konvensional yang lebih ringan sehingga hasilnya dapat dibandingkan secara langsung dengan berbagai penelitian terdahulu.

Algoritma K-Nearest Neighbour (KNN) dapat memberikan hasil yang baik dalam analisis sentimen terhadap ulasan aplikasi Bank Aladin [10]. Selain itu, penelitian yang dilakukan oleh Mustaqim dkk. menggunakan algoritma KNN untuk analisis sentimen terhadap ulasan aplikasi PosPay, yang menghasilkan hasil klasifikasi yang memuaskan [11]. Namun, penelitian tersebut masih memiliki beberapa batasan. Sebagian besar penelitian pertama hanya menerapkan satu algoritma klasifikasi, sehingga belum memberikan gambaran tentang performa relatif antaralgoritma dalam domain layanan keuangan. Kedua, penelitian sebelumnya tidak mengevaluasi pengaruh penerapan SMOTE pada berbagai algoritma klasifikasi dalam dataset ulasan yang memiliki distribusi kelas yang tidak seimbang. Selain itu, penelitian sebelumnya lebih banyak fokus pada aplikasi perbankan digital dan layanan pembayaran, sementara analisis sentimen pada aplikasi pinjaman online, khususnya PinjamYuk, masih cukup jarang dilakukan. Karakteristik ulasan aplikasi pinjaman online sebenarnya lebih bervariasi karena mencakup keluhan tentang pengajuan pinjaman, suku bunga, penagihan, dan kualitas layanan yang memerlukan analisis tersendiri. Sebagai akibatnya, belum ada bukti empiris yang menunjukkan algoritma mana yang memberikan kinerja terbaik dalam klasifikasi sentimen ulasan aplikasi pinjaman online setelah penerapan penyeimbangan data dengan menggunakan SMOTE.

Mengacu pada berbagai kekurangan yang masih ada dalam studi sebelumnya, penelitian ini menganalisis efektivitas algoritma *K-Nearest Neighbor* (KNN) dan Naive Bayes dalam mengkategorikan sentimen dari ulasan pengguna aplikasi PinjamYuk. Proses klasifikasi menggunakan pembobotan *Term Frequency-Inverse Document Frequency* (TF-IDF) dan menerapkan *Synthetic Minority Oversampling Technique* (SMOTE) untuk menangani ketidakseimbangan distribusi kelas. Tidak seperti penelitian terdahulu yang biasanya menekankan pada satu algoritma klasifikasi atau menggunakan objek penelitian aplikasi perbankan digital atau dompet digital, penelitian ini berfokus pada ulasan aplikasi pinjaman online yang memiliki karakteristik bahasa serta proporsi kelas sentimen yang bervariasi. Dengan pendekatan ini, penelitian ini diharapkan bisa menambah bukti empiris tentang efektivitas penerapan SMOTE dalam proses klasifikasi sentimen dan sekaligus memberikan referensi untuk pengembangan metode analisis sentimen di bidang teknologi finansial. Evaluasi kinerja tiap model dilakukan dengan memakai metrik akurasi, presisi, recall, dan F1-score untuk menentukan algoritma yang memiliki kinerja terbaik.

Fokus utama penelitian ini adalah analisis menyeluruh terhadap dua algoritma klasifikasi yang digunakan dalam aplikasi pinjaman online yang menggabungkan metode TF-IDF dan SMOTE. Penelitian ini membandingkan hasil klasifikasi berdasarkan akurasi, presisi, recall, dan F1-score, serta melakukan uji McNemar untuk mengukur konsistensi model. Proses evaluasi ini memberikan landasan empiris yang lebih kuat untuk memilih algoritma analisis sentimen terbaik untuk hasil aplikasi pinjaman online.

II. METODE

A. Dataset

Data penelitian memanfaatkan dataset sekunder yang diambil dari platform Kaggle, yang merupakan hasil dari pengambilan data ulasan pengguna aplikasi PinjamYuk di Google Play Store untuk periode 2023–2024. Dataset ini mencakup atribut-atribut berikut: nama pengguna, penilaian, tanggal review, dan ulasan pengguna. Selanjutnya, proses seleksi data dilakukan dengan menghapus ulasan yang duplikat, kosong, atau terlalu pendek, sehingga menghasilkan sekitar 500 ulasan yang memenuhi syarat untuk dijadikan dataset penelitian. Dataset yang terpilih selanjutnya diproses melalui analisis berbasis rating, preprocessing, ekstraksi fitur dengan metode *Term Frequency-Inverse Document Frequency* (TF-IDF), analisis kelas menggunakan *Synthetic Minority Oversampling Technique* (SMOTE), dan klasifikasi dengan algoritma *K-Nearest Neighbor* (KNN) serta Naive Bayes.

B. Pelabelan Data

Dalam penelitian ini, analisis sentimen dilakukan secara otomatis berdasarkan penilaian pengguna di Google Play Store. Peringkat 1–2 menunjukkan sentimen negatif, peringkat 3 menunjukkan sentimen netral, dan peringkat 4–5 menunjukkan sentimen positif. Peringkat ini dipilih karena mencerminkan tingkat kepuasan pengguna terhadap aplikasi, sehingga prosesnya konsisten dan mudah dilakukan. Berikut adalah kriteria yang digunakan dalam penelitian ini:

TABEL I
DATASET SENTIMEN

Rating	Label Sentimen
1-2	Negatif
3	Netral
4-5	Positif

Tabel I menunjukkan bahwa ulasan dengan peringkat 1 dan 2 termasuk dalam kategori sentimen negatif karena ulasan tersebut menurunkan kepuasan pengguna terhadap aplikasi. Rating 3 tergolong dalam kategori sentimen netral karena mencerminkan evaluasi yang cenderung biasa atau tidak menunjukkan kecenderungan emosi yang kuat. Sementara itu, peringkat 4 dan 5 dianggap positif karena membuktikan tingkat kepuasan pengguna yang tinggi terhadap aplikasi.

C. Preprocessing

Sebelum memulai proses klasifikasi, langkah-langkah prapemrosesan dilakukan untuk meningkatkan mutu data teks. Ini mencakup penyederhanaan kata dengan case folding, penghapusan karakter tidak penting, pengubahan kata menjadi urutan kata, penghapusan stopword untuk menghilangkan kata-kata umum yang tidak relevan, serta mengubah kata-kata yang tidak relevan menjadi kata kunci. Setelah melewati tahapan di atas, diperoleh data yang lebih teratur, dan sebagian besar dari data tersebut dapat digunakan dalam proses pengambilan fitur dengan metode *Term Frequency-Inverse Document Frequency* (TF-IDF), serta dalam proses klasifikasi menggunakan algoritma *K-Nearest Neighbor* (KNN) dan Naive Bayes.

D. TF-IDF

Selanjutnya, metode *Term Frequency-Inverse Document Frequency* (TF-IDF) digunakan untuk mengubah teks menjadi format numerik. Metode ini menilai setiap kata berdasarkan frekuensi kemunculannya dalam dokumen dan relevansinya terhadap keseluruhan dokumen. Teks diubah menjadi vektor numerik selama proses prapemrosesan ini [12]. Tujuan TF-IDF adalah untuk membandingkan Frekuensi Istilah (TF) dengan Frekuensi Dokumen Terbalik (IDF), sehingga kata-kata yang sering muncul di satu dokumen tetapi tidak di dokumen lain diberi nilai yang lebih tinggi [13].

Dalam penelitian ini, proses pemodelan fitur dilakukan menggunakan paket *TfidfVectorizer* dari *Scikit-Learn* (Python) dengan parameter `max_features = 5000` dan `ngram_range = (1,2)`. Meskipun `ngram_range = (1,2)` memungkinkan penggunaan unigram dan bigram dalam proses ekstraksi fitur, `max_features` berfungsi untuk menentukan jumlah fitur yang digunakan. Transformasi TF-IDF menghasilkan representasi numerik untuk setiap kelas pengguna, yang kemudian digunakan dalam tahap penyeimbangan kelas menggunakan metode SMOTE serta proses klasifikasi menggunakan algoritma *K-Nearest Neighbor* (KNN) dan Naive Bayes.

E. Synthetic Minority Oversampling Technique (SMOTE)

Setelah mentransformasi data menggunakan pendekatan *Term Frequency-Inverse Document Frequency* (TF-IDF), kumpulan data tersebut dibagi menjadi data pelatihan dan data uji, dengan 80% dialokasikan ke kumpulan data pelatihan dan 20% ke kumpulan data uji. Data tersebut dibagi menggunakan parameter `random_state = 42` untuk memastikan hasil penelitian dapat direproduksi, sementara parameter `stratify` digunakan untuk menetapkan proporsi setiap kelas sentiment positif, negatif, dan netral baik pada data pelatihan maupun data pengujian. Data pengujian digunakan untuk mengevaluasi model, sedangkan data pelatihan digunakan dalam proses pembangunan model klasifikasi.

Selain itu, penelitian ini menggunakan *Synthetic Minority Oversampling Technique* (SMOTE) untuk mengatasi

masalah ketidakseimbangan distribusi kelas. Metode SMOTE bekerja dengan menghasilkan sampel sintetis untuk kelas minoritas berdasarkan fitur-fitur data, sehingga distribusi setiap kelas menjadi lebih seimbang [14]. Dalam penelitian ini, SMOTE diterapkan hanya pada data pelatihan setelah proses pembagian data selesai dilakukan. Langkah ini bertujuan untuk mencegah kebocoran data, guna memastikan bahwa evaluasi model tetap objektif dan tidak dipengaruhi oleh informasi dari data pengujian. Data yang diperoleh dari proses penyeimbangan ini kemudian digunakan untuk membangun model klasifikasi menggunakan algoritma *K-Nearest Neighbour* (KNN) dan Naive Bayes [15].

F. Klasifikasi Menggunakan K-Nearest Neighbor (KNN)

Metode klasifikasi *K-Nearest Neighbor* (KNN) adalah algoritma yang mendasarkan klasifikasinya pada prinsip kesamaan antar data. Metode ini mengelompokkan data ke dalam kategori yang relevan berdasarkan mayoritas K terdekat yang ada dalam analisis data [16]. Kedekatan antara data dihitung dengan metode Jarak Euclidean seperti yang ditunjukkan pada persamaan [17]. Metode Euclidean Distance digunakan untuk menghitung jarak antar data dalam algoritma *K-Nearest Neighbor* (KNN). Euclidean Distance merupakan jarak lurus antara dua titik data dalam ruang multidimensi dan dinyatakan dalam Persamaan (1)

$$d(x,y) = \sqrt{\sum_{i=1}^n (x_i - y_i)^2}$$

Keterangan:

$d(x, y)$ = jarak antara data x dan data y

x_i = nilai atribut ke- i pada data x

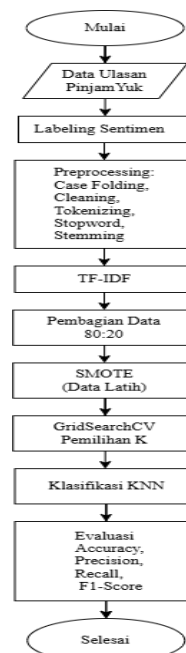
y_i = nilai atribut ke- i pada data y

n = jumlah atribut atau fitur yang digunakan dalam perhitungan jarak.

Dalam penelitian ini, data hasil pembobotan TF-IDF yang telah melalui proses penyeimbangan kelas menggunakan metode SMOTE digunakan sebagai masukan dalam proses klasifikasi. Untuk memperoleh parameter terbaik, dilakukan optimasi menggunakan GridSearchCV dengan teknik 3-fold cross validation dan metrik evaluasi accuracy. Parameter yang diuji adalah jumlah tetangga $n_neighbors$ sebesar 3, 5, 7, dan 9. Algoritma *K-Nearest Neighbor* (KNN) menggunakan pengaturan bawaan Scikit-Learn, yaitu metode pembobotan tetangga *uniform* serta metrik jarak *Minkowski* ($p = 2$) yang ekuivalen dengan Euclidean Distance.

Berdasarkan hasil optimasi GridSearchCV, parameter optimalnya adalah $n_neighbors = 3$, dengan akurasi validasi silang rata-rata sekitar 90,74 persen. Parameter-parameter ini kemudian digunakan untuk membangun model *K-Nearest Neighbour* (KNN), yang diterapkan pada data uji untuk mengklasifikasikan sentimen ke dalam

kategori positif, negatif, dan netral. Kinerja model tersebut kemudian dievaluasi menggunakan metrik akurasi, presisi, recall, dan F1-score.



Gambar 1. Alur Proses Klasifikasi Menggunakan Algoritma K-Nearest Neighbor (KNN)

Gambar 1 mengilustrasikan proses klasifikasi kalimat menggunakan algoritma KNN, yang dimulai dengan tahap persiapan dan transformasi data menggunakan metode TF-IDF. Dataset dibagi menjadi data latih dan data uji dengan rasio 80:20. Selanjutnya, teknik SMOTE diterapkan pada data latih untuk mengatasi ketidakseimbangan distribusi kelas. Hasil analisis data digunakan dalam proses optimasi parameter menggunakan GridSearchCV untuk menentukan nilai K yang paling optimal. Model KNN yang telah dibuat digunakan untuk memprediksi sentimen pada data uji dan dievaluasi menggunakan metrik akurasi, presisi, recall, dan F1-score.

G. Klasifikasi Menggunakan Naive Bayes (NB)

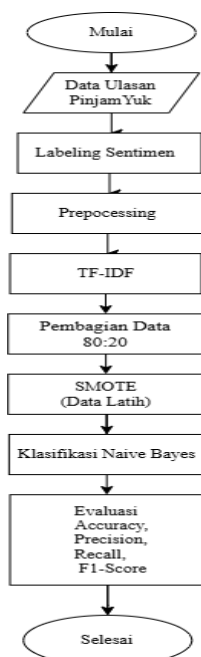
Metode Naive Bayes adalah metode klasifikasi statistik yang digunakan untuk meningkatkan keterlibatan siswa di dalam kelas. Metode ini didasarkan pada teorema [18]. Metode ini mengidentifikasi kelas suatu data dengan merujuk pada nilai probabilitas tertinggi yang dihasilkan dari perhitungan untuk setiap kelas. Probabilitas suatu data tergolong dalam kelas tertentu dihitung dengan memanfaatkan Teorema Bayes yang dinyatakan dalam Persamaan (2).

$$P(C|X) = \frac{P(X|C) \times P(C)}{P(X)}$$

Keterangan:

$P(C | X)$ = probabilitas data X termasuk ke dalam kelas C
 $P(X | C)$ = probabilitas kemunculan data X pada kelas C
 $P(C)$ = probabilitas awal dari kelas C
 $P(X)$ = probabilitas kemunculan data X

Dalam penelitian ini, data yang diperoleh dari pembobotan TF-IDF yang telah menjalani tahap penyeimbangan kelas dengan metode SMOTE digunakan sebagai masukan dalam proses klasifikasi. Model Naive Bayes dibangun dengan memanfaatkan data pelatihan, lalu digunakan untuk meramalkan sentimen pada data uji untuk mengkategorikan sebagai sentimen positif, netral, atau negatif. Model ini digunakan untuk meningkatkan performa klasifikasi terhadap umpan balik pengguna aplikasi PinjamYuk dengan memanfaatkan metrik berikut: akurasi, presisi, F1-score, dan recall.



Gambar 2. Alur Proses Klasifikasi Menggunakan Algoritma Naive Bayes

Gambar 2 menggambarkan langkah-langkah klasifikasi dengan algoritma Naive Bayes, dimulai dari analisis dan pengolahan data dengan metode TF-IDF. Data dibagi menjadi data latih dan data uji dengan proporsi 80:20. Selanjutnya, metode SMOTE digunakan pada data untuk mengatasi ketidakseimbangan dalam distribusi kelas. Model Naive Bayes dikembangkan dengan memanfaatkan data yang diperoleh dari proses penelitian. Model yang dirancang kemudian digunakan untuk meramalkan sentimen pada data uji dan menilai F1-score, akurasi, presisi, serta recall.

H. Evaluasi Model

Penilaian terhadap algoritma *K-Nearest Neighbor* (KNN) dan Naive Bayes dalam mengkategorikan masukan pengguna ke dalam tiga kelompok: positif, negatif, dan

netral. Model dievaluasi dengan menggunakan matriks kebingungan serta metrik akurasi, presisi, recall, dan F1-score. Skor akurasi digunakan untuk menilai kemampuan model dalam mengklasifikasikan seluruh data uji dengan tepat. Sementara recall menunjukkan kemampuan model dalam mengidentifikasi semua data yang termasuk dalam tiap kategori, presisi digunakan untuk menilai ketepatan prediksi di dalam masing-masing kategori. F1-score berperan sebagai metrik yang menggabungkan presisi dan recall untuk menyeimbangkan output model.

Selain uji analisis data, penelitian ini juga menggunakan 5 kali lipat Stratified K-Fold Cross Validation untuk menilai kestabilan model kerja. Metode itu menjaga proporsi setiap kategori sentimen di setiap lipatan, sehingga hasil evaluasi lebih mencerminkan kondisi sebenarnya dan mengurangi dampak variasi dalam Pembagian data latih dan data uji.

Selain itu, penelitian ini menerapkan uji McNemar dengan tingkat signifikansi $\alpha = 0,05$ untuk menentukan apakah perbedaan kinerja antara algoritma *K-Nearest Neighbour* (KNN) dan Naive Bayes bersifat signifikan secara statistik. Uji tersebut dilakukan berdasarkan hasil prediksi kedua model pada data uji, guna memastikan apakah perbedaan kinerja yang diamati disebabkan oleh kemampuan masing-masing algoritma atau hanya dipengaruhi oleh variasi pada data uji.

III. HASIL DAN PEMBAHASAN

A. Distribusi Data Sentimen

Kumpulan data yang digunakan dalam penelitian ini mencakup 500 pengguna aplikasi PinjamYuk yang telah melalui proses klasifikasi sentimen ke dalam tiga kategori: positif, negatif, dan netral. Distribusi data sebelum penerapan *Synthetic Oversampling Minority* (SMOTE).

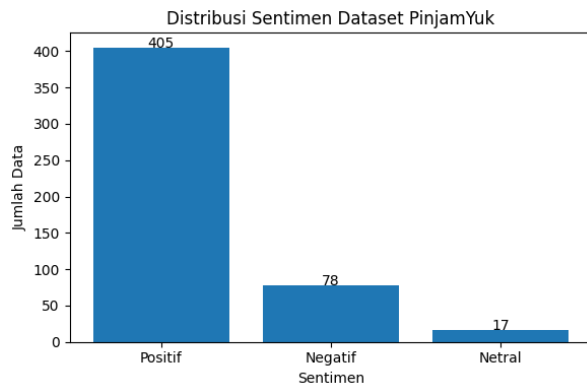
TABEL II
DISTRIBUSI DATA SENTIMEN SEBELUM SMOTE

Sentimen	Jumlah Data
Positif	405
Negatif	78
Netral	17
Total	500

Berdasarkan Tabel II, distribusi data sentimen menunjukkan ketidakseimbangan kelas yang ada. Sentimen positif mendominasi dataset dengan total 405 ulasan (81%), sementara sentimen negatif mencapai 78 ulasan (15,6%) dan sentimen netral hanya 17 ulasan (3,4%). Selisih jumlah data yang signifikan antara kelas dapat membuat model klasifikasi cenderung memprediksi kelas yang lebih banyak dan mengurangi kemampuan model untuk mengidentifikasi kelas yang lebih sedikit.

Dalam penelitian ini, analisis sentimen didasarkan pada peringkat pengguna dari Google Play Store. Peringkat ini dipilih karena dapat memberikan ilustrasi awal tentang sentimen pengguna. Namun, penggunaan peringkat sebagai dasar evaluasi memiliki kelemahan karena peringkat tidak

selalu mendukung teknik ulasan. Dalam beberapa kasus, pengguna mungkin memberikan peringkat tinggi berdasarkan komentar negatif atau positif, yang dapat menyebabkan perbedaan antara label sentimen dan ulasan. Kondisi ini dapat memengaruhi proses pembelajaran model dan hasil klasifikasi yang diperoleh.



Gambar 3. Distribusi Data Sentimen Sebelum SMOTE

Gambar 3 menampilkan distribusi jumlah data pada setiap kategori sentimen. Dibandingkan dengan kelas netral dan negatif, jumlah titik data pada kelas sentimen positif jauh lebih tinggi. Hal ini menunjukkan bahwa dataset tersebut memiliki distribusi kelas yang tidak seimbang. Penelitian ini menggunakan metode SMOTE dalam analisis data untuk menambah jumlah titik data pada kelas minoritas, sehingga menghasilkan distribusi data yang lebih seimbang sebelum dilakukan klasifikasi menggunakan algoritma *K-Nearest Neighbour* (KNN) dan Naive Bayes.

B. Hasil Penyeimbangan Data Menggunakan SMOTE

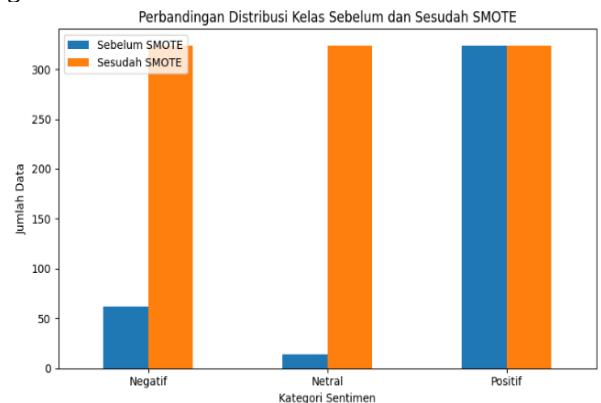
Ketidakseimbangan distribusi kelas dalam suatu dataset dapat memengaruhi kinerja klasifikasi model, karena model cenderung lebih memperhatikan kelas mayoritas dibandingkan kelas minoritas. Oleh karena itu, penelitian ini menerapkan metode *Synthetic Minority Oversampling* (SMOTE) yang memanfaatkan data pelatihan untuk menghasilkan data sintesis dengan distribusi kelas yang seimbang, sehingga distribusi data menjadi lebih merata.

TABEL III
DISTRIBUSI DATA SENTIMEN SETELAH SMOTE

Sentimen	Jumlah Data
Positif	324
Negatif	324
Netral	324
Total	972

Tabel III menunjukkan bahwa jumlah data pada kelas-kelas minoritas telah meningkat, sehingga distribusi data di antara kelas-kelas tersebut menjadi lebih merata dibandingkan sebelum penerapan SMOTE. Diharapkan hal ini dapat membantu model memahami karakteristik masing-masing kategori sentimen secara lebih efektif dan

mengurangi kecenderungan model untuk memprediksi kelas yang dominan.



Gambar 4. Distribusi Data Sentimen Setelah Penerapan SMOTE

Gambar 4 menunjukkan perbandingan distribusi data setelah metode SMOTE diterapkan. Terlihat jelas bahwa jumlah data pada setiap kategori sentimen menjadi lebih proporsional dibandingkan sebelumnya. Analisis data yang diperoleh digunakan sebagai dasar untuk proses klasifikasi dengan algoritma *K-Nearest Neighbor* (KNN) dan Naive Bayes.

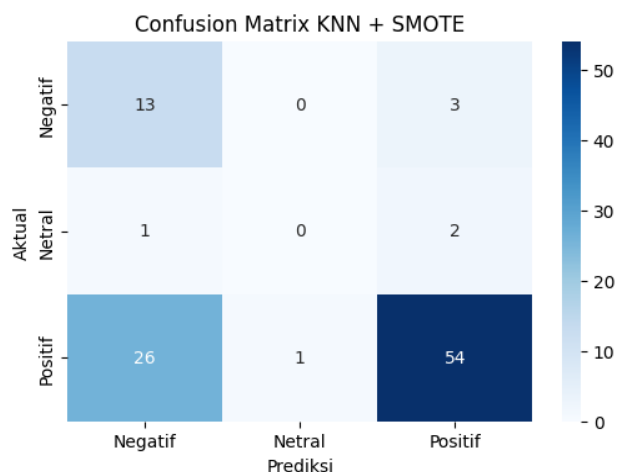
C. Hasil Klasifikasi Menggunakan K-Nearest Neighbor (KNN)

Klasifikasi sentimen dilakukan dengan menggunakan algoritma *K-Nearest Neighbour* (KNN) setelah data dianalisis menggunakan metode *Synthetic Minority Oversampling* (SMOTE). Parameter-parameter tersebut kemudian dioptimalkan menggunakan *GridSearchCV* untuk mendapatkan konfigurasi model terbaik.

TABEL IV
HASIL EVALUASI ALGORITMA KNN MENGGUNAKAN SMOTE

Metrik	Nilai
Akurasi	0,67
Presisi	0,79
Recall	0,67
F1-Score	0,70

Tabel IV menunjukkan bahwa algoritma KNN yang dipadukan dengan metode SMOTE mencapai nilai akurasi sebesar 67%. Tingkat presisi 79% menunjukkan bahwa sejumlah besar data yang diklasifikasikan ke dalam kategori sentimen benar-benar termasuk dalam kategori ini. Angka recall sekitar 67% menunjukkan kemampuan model untuk mengambil data yang sebelumnya hilang di sebagian besar kelas sentimen. Demikian pula, skor F1 minimal 70% menunjukkan adanya keseimbangan yang cukup baik antara presisi dan recall.



Gambar 5. Confusion Matrix Algoritma KNN Menggunakan SMOTE

Gambar 5 menunjukkan bahwa model KNN mampu mengklasifikasikan 13 contoh sentimen negatif secara akurat dari total 16 contoh. Pada kelas sentimen netral, model tersebut tidak menunjukkan kinerja yang baik, karena semua contoh netral masih diklasifikasikan ke dalam kelas yang berbeda. Sementara itu, pada kelas sentimen positif, 54 dari 81 contoh berhasil diprediksi secara akurat. Hasil penelitian menunjukkan bahwa penerapan SMOTE mampu meningkatkan efektivitas model dalam mengenali kelompok minoritas, khususnya dalam konteks sentimen negatif. Namun, kinerja model pada kelas netral tetap buruk akibat jumlah data awal yang sangat sedikit, yang berarti algoritma KNN belum mampu mempelajari pola sentimen netral secara optimal.

TABEL V
CLASSIFICATION REPORT ALGORITMA KNN MENGGUNAKAN SMOTE

Kelas	Presisi	Recall	F1-Score	Support
Negatif	0,33	0,81	0,46	16
Netral	0,00	0,00	0,00	3
Positif	0,92	0,67	0,77	81

Tabel V menunjukkan bahwa kinerja algoritma K-Nearest Neighbour (KNN) bervariasi tergantung pada kategori sentimen. Presisi sebesar 0,92, recall sebesar 0,67, dan F1-score sebesar 0,77 menunjukkan bahwa sejumlah besar prediksi positif dilakukan secara akurat. Untuk sentimen negatif, presisi sekitar 0,33, recall sekitar 0,81, dan skor F1 sekitar 0,46. Fakta bahwa recall lebih tinggi daripada presisi menunjukkan bahwa model berhasil mengidentifikasi sebagian besar data negatif, meskipun masih ada data dari kategori lain yang diidentifikasi sebagai negatif. Sebaliknya, kategori sentimen netral memiliki nilai presisi, recall, dan F1-score yang sama, yaitu 0,00. Hasil ini menunjukkan bahwa model KNN masih menghadapi kesulitan. Situasi ini

diduga disebabkan oleh kelangkaan data sentimen netral dalam dataset awal; akibatnya, meskipun SMOTE telah diterapkan, algoritma KNN tidak mampu mempelajari karakteristik kelas tersebut secara optimal.

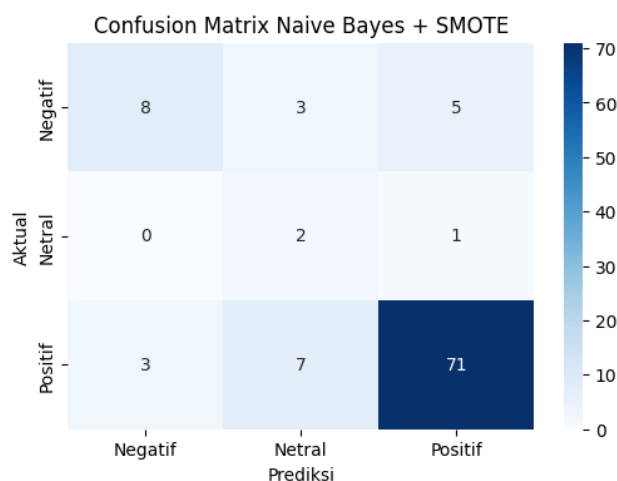
D. Hasil Klasifikasi Menggunakan Naive Bayes

Selain menerapkan algoritma *K-Nearest Neighbor* (KNN), penelitian ini juga memanfaatkan algoritma Naive Bayes untuk mengklasifikasikan data dari aplikasi PinjamYuk. Setelah proses klasifikasi, data dianalisis dengan menggunakan metode *Synthetic Minority Oversampling Technique* (SMOTE) untuk menangani ketidakseimbangan kelas. Selanjutnya, model Naive Bayes dibuat dengan memanfaatkan data latih dan uji untuk menilai kinerja klasifikasi yang dihasilkan.

TABEL VI
HASIL EVALUASI NAIVE BAYES MENGGUNAKAN SMOTE

Metrik	Nilai
Akurasi	0,81
Presisi	0,87
Recall	0,81
F1-Score	0,83

Berdasarkan Tabel VI, algoritma Naive Bayes mencapai akurasi sekitar 81%. Angka akurasi ini menunjukkan bahwa sebagian besar data yang diprediksi termasuk dalam kategori tertentu memang benar-benar termasuk dalam kategori tersebut. Angka recall sekitar 81% menunjukkan kemampuan model untuk mengidentifikasi data yang sebelumnya telah dikumpulkan dan termasuk dalam setiap kategori. Nilai F1-score sekitar 83 persen menunjukkan keseimbangan yang baik antara recall dan precision, sehingga memungkinkan model untuk mencapai kinerja klasifikasi yang optimal.



Gambar 6. Confusion Matrix Algoritma Naive Bayes Menggunakan SMOTE

Gambar 6 menunjukkan bahwa algoritma Naive Bayes berhasil mengklasifikasikan sebagian besar data dalam setiap kategori sentimen secara akurat. Kesalahan klasifikasi

masih terdeteksi pada beberapa data sentimen negatif dan netral, namun jumlahnya lebih rendah dibandingkan dengan algoritma *K-Nearest Neighbour* (KNN). Hal ini menunjukkan bahwa Naive Bayes lebih efektif dalam menganalisis pola kata dalam representasi TF-IDF, sehingga menghasilkan prediksi yang lebih stabil untuk semua kategori sentimen. Penggunaan SMOTE juga berkontribusi pada peningkatan kemampuan model dalam mengidentifikasi kelas minoritas dengan menciptakan distribusi data yang lebih seimbang sebelum tahap klasifikasi.

TABEL VII
CLASSIFICATION REPORT ALGORITMA NAIVE BAYES MENGGUNAKAN SMOTE

Kelas	Presisi	Recall	F1-Score	Support
Negatif	0,73	0,50	0,59	16
Netral	0,17	0,67	0,27	3
Positif	0,92	0,88	0,90	81

Berdasarkan Tabel VII, algoritma Naive Bayes menunjukkan kinerja yang lebih baik daripada algoritma *K-Nearest Neighbour* (KNN) di hampir semua kategori. Pada kelas positif, model ini mencapai presisi sebesar 0,92, recall sebesar 0,88, dan F1-score sebesar 0,90. Hal ini menunjukkan bahwa Naive Bayes dapat secara konsisten mendeteksi sentimen positif dengan tingkat akurasi yang tinggi. Untuk kelas sentimen negatif, hasil yang diperoleh adalah: presisi sebesar 0,73, recall sebesar 0,50, dan skor F1 sebesar 0,59. Hal ini menunjukkan bahwa model tersebut sangat efisien dalam mendeteksi sentimen negatif, meskipun masih ada beberapa data yang tidak dapat dikategorikan secara akurat. Presisi, recall, dan skor F1 untuk kelas netral masing-masing adalah 0,17, 0,67, dan 0,27. Nilai recall yang tinggi menunjukkan bahwa sebagian besar data netral dapat dikenali dengan baik, sedangkan nilai presisi yang rendah menunjukkan bahwa data dari kelas lain telah diprediksi sebagai netral.

Secara keseluruhan, algoritma Naive Bayes menghasilkan hasil klasifikasi yang lebih baik daripada KNN. Hal ini disebabkan karena algoritma Naive Bayes menggunakan pendekatan probabilistik untuk meminimalkan kemungkinan munculnya setiap kata dalam sebuah kalimat. Penggunaan TF-IDF sebagai representasi fitur menghasilkan vektor berdimensi tinggi dan jarang, sehingga memungkinkan Naive Bayes memanfaatkan distribusi probabilitas setiap fitur secara lebih efisien.

E. Perbandingan Kinerja Algoritma KNN dan Naive Bayes Menggunakan SMOTE

Setelah menerapkan metode *Synthetic Minority Oversampling Technique* (SMOTE) untuk menentukan algoritma dengan kinerja terbaik, hasil klasifikasi dari algoritma *K-Nearest Neighbor* (KNN) dan Naive Bayes dipertimbangkan.

TABEL VIII
PERBANDINGAN HASIL KLASIFIKASI KNN DAN NAIVE BAYES MENGGUNAKAN SMOTE

Algoritma	Akurasi	Presisi	Recall	F1-Score
KNN + SMOTE	0,67	0,79	0,67	0,70
Naive Bayes + SMOTE	0,81	0,87	0,81	0,83

Tabel VIII mengungkapkan bahwa algoritma Naive Bayes yang diterapkan dengan metode SMOTE menunjukkan performa lebih unggul daripada algoritma *K-Nearest Neighbor* (KNN) di semua metrik evaluasi. Algoritma Naive Bayes memperoleh akurasi 81%, presisi 87%, recall 81%, serta skor F1 83%. Sebaliknya, algoritma KNN memperoleh akurasi 67%, presisi 79%, recall 67%, dan skor F1 sebesar 70%. Hasil penelitian ini menunjukkan bahwa algoritma Naive Bayes lebih unggul daripada algoritma *K-Nearest Neighbour* (KNN) dalam mengklasifikasikan sentimen pengguna pada aplikasi PinjamYuk.

Hal ini terjadi karena Naive Bayes menerapkan pendekatan probabilistik yang memperhitungkan frekuensi setiap kata pada masing-masing kelas sentimen, sehingga lebih efektif dalam memanfaatkan representasi TF-IDF yang memiliki dimensi tinggi dan bersifat jarang. Sebaliknya, algoritma *K-Nearest Neighbor* (KNN) mengidentifikasi kelas berdasarkan seberapa dekat jarak antara dokumen. Pada data teks yang memiliki banyak fitur, perhitungan jarak jadi kurang mewakili sehingga kemampuan KNN dalam membedakan kelas sentimen berkurang. Keadaan itu tampak pada confusion matrix KNN yang masih menunjukkan kesalahan dalam klasifikasi, terutama pada kategori sentimen netral, sementara Naive Bayes dapat memberikan prediksi yang lebih stabil di semua kategori sentimen.

Selain itu, bobot TF-IDF menghasilkan representasi data sebagai vektor dengan banyak fitur dan kebanyakan nilainya adalah nol. Dalam situasi tersebut, algoritma Naive Bayes masih dapat menunjukkan kinerja yang baik karena proses klasifikasinya didasarkan pada perhitungan kemungkinan munculnya setiap fitur pada setiap kelas sentimen. Sebaliknya, algoritma *K-Nearest Neighbor* bergantung pada pengukuran jarak antara vektor untuk mengidentifikasi kelas dari suatu data. Di ruang fitur berdimensi tinggi, selisih jarak antara data cenderung semakin mendekat sehingga kapabilitas algoritma untuk membedakan tetangga terdekat menurun. Keadaan tersebut mengakibatkan kinerja Naive Bayes lebih unggul dibandingkan KNN dalam mengklasifikasikan sentimen ulasan dalam penelitian ini.

Selain itu, penerapan *Synthetic Minority Oversampling Technique* (SMOTE) terbukti efektif dalam menangani ketidakseimbangan kelas dengan meningkatkan jumlah data pada kelas yang minoritas. Penyebaran data yang lebih seimbang membantu kedua algoritma dalam menemukan pola sentimen negatif dan netral dengan lebih efisien. Namun, berdasarkan semua metrik evaluasi seperti akurasi,

presisi, recall, dan skor F1, algoritma Naive Bayes terus menunjukkan kinerja yang lebih baik daripada KNN. Oleh karena itu, algoritma ini lebih cocok untuk klasifikasi sentimen ulasan pengguna aplikasi PinjamYuk dalam penelitian ini.

F. Hasil Stratified K-Fold Cross Validation

Dalam menilai stabilitas kinerja model, penelitian ini juga menggunakan Validasi Silang K-Fold Terstratifikasi dengan k-folds pada kumpulan data yang telah melalui proses pembobotan TF-IDF dan penyeimbangan kelas menggunakan metode SMOTE. Metode ini dipilih karena dapat menjaga proporsi masing-masing kelas sentimen pada setiap fold, sehingga evaluasi menjadi lebih konsisten dibandingkan dengan hanya melakukan satu pembagian antara data latih dan data uji.

TABEL IX
HASIL EVALUASI STRATIFIED K-FOLD CROSS VALIDATION

Algoritma	Rata-rata Accuracy	Standar Deviasi
KNN + SMOTE	0,808	0,004
Naive Bayes + SMOTE	0,810	0,000

Tabel IX menunjukkan bahwa algoritma *K-Nearest Neighbor* (KNN) memperoleh akurasi rata-rata 0,808 dengan deviasi standar 0,004, sedangkan algoritma Naive Bayes mencapai akurasi rata-rata 0,810 dengan deviasi standar yang hampir nol. Standar deviasi yang rendah menunjukkan bahwa kedua algoritma tersebut menghasilkan kinerja yang stabil pada setiap fold. Walaupun tidak ada perbedaan signifikan dalam akurasi rata-rata antara kedua algoritma tersebut, Naive Bayes menunjukkan performa yang lebih baik dan lebih stabil dibandingkan KNN. Penelitian ini mendukung studi yang mengindikasikan bahwa algoritma Naive Bayes lebih tepat untuk mengklasifikasikan sentimen dari masukan pengguna aplikasi PinjamYuk.

G. Uji Signifikansi Menggunakan McNemar

Untuk memastikan bahwa perbedaan performa antara algoritma *K-Nearest Neighbor* (KNN) dan Naive Bayes tidak dipengaruhi oleh variasi dalam data pengujian, penelitian ini menerapkan uji McNemar pada hasil prediksi kedua model menggunakan data uji. Uji ini dilaksanakan pada tingkat signifikansi 5% ($\alpha = 0,05$) untuk mengetahui apakah terdapat perbedaan signifikan dalam kinerja klasifikasi antara kedua algoritma.

TABEL X
HASIL UJI SIGNIFIKANSI PERFORMA MODEL MENGGUNAKAN MCNEMAR

Parameter	Nilai
Statistik McNemar	6,04
p-value	0,014
Keputusan	Berbeda signifikan ($p < 0,05$)

Tabel X menunjukkan hasil statistik McNemar sebesar 6,04 dengan nilai p sebesar 0,014. Nilai p yang lebih kecil

dari tingkat signifikansi 0,05 mengindikasikan adanya perbedaan kinerja yang signifikan secara statistik antara algoritma *K-Nearest Neighbor* (KNN) dan Naive Bayes dalam klasifikasi sentimen dari diskusi pengguna aplikasi PinjamYuk. Temuan tersebut menunjukkan bahwa algoritma Naive Bayes memiliki keunggulan dalam hal akurasi, presisi, recall, serta F1-score yang lebih tinggi dibandingkan KNN, yang didukung oleh hasil analisis statistik. Dengan pendekatan ini, algoritma Naive Bayes menunjukkan hasil yang jauh lebih baik dibandingkan algoritma *K-Nearest Neighbor* (KNN) pada dataset yang digunakan dalam penelitian ini.

H. Keterbatasan Penelitian

Penelitian ini memiliki beberapa keterbatasan yang perlu diperhatikan saat menafsirkan hasilnya. Pertama, kumpulan data yang digunakan hanya terdiri dari 500 ulasan pengguna terhadap aplikasi PinjamYuk; oleh karena itu, data tersebut tidak memberikan gambaran menyeluruh mengenai beragamnya pendapat pengguna. Kedua, penelitian ini dilakukan pada satu subjek saja yaitu aplikasi PinjamYuk yang berarti hasilnya tidak dapat digeneralisasikan ke aplikasi pinjaman daring lainnya dengan karakteristik ulasan atau perilaku pengguna yang berbeda. Ketiga, proses pelabelan sentimen didasarkan pada peringkat yang diberikan pengguna di Google Play Store. Meskipun pendekatan ini memudahkan proses pelabelan yang konsisten dan mudah direplikasi, pelabelan berbasis peringkat tidak selalu sepenuhnya mencerminkan sentimen yang terkandung dalam isi ulasan. Pengguna mungkin memberikan peringkat yang tidak konsisten dengan komentar yang mereka tulis, dan situasi ini berpotensi memengaruhi proses pembelajaran model serta akurasi hasil klasifikasi.

Selain itu, studi ini hanya membandingkan dua metode klasifikasi *K-Nearest Neighbour* (KNN) dan Naive Bayes dengan menggunakan representasi fitur TF-IDF dan teknik penyeimbangan SMOTE. Oleh karena itu, penelitian selanjutnya harus membandingkan hasil studi ini dengan algoritma lain seperti *Support Vector Machines* (SVM), Random Forest, atau model berbasis transformer. Untuk meningkatkan generalisasi hasil penelitian ini, juga diperlukan penggunaan dataset yang lebih luas dan beragam yang dikumpulkan dari berbagai aplikasi pinjaman online.

IV. KESIMPULAN

Tujuan penelitian ini adalah untuk membandingkan kemampuan algoritma *K-Nearest Neighbor* (KNN) dan Naive Bayes dalam mengkategorikan ulasan pengguna aplikasi PinjamYuk. Untuk mencapai tujuan tersebut, penelitian ini akan menggunakan metode penyeimbangan data *Synthetic Minority Oversampling* (SMOTE) dan pembobotan *Term Frequency-Inverse Document Frequency* (TF-IDF). Hasil pengujian menunjukkan bahwa algoritma Naive Bayes mengalahkan KNN secara kinerja. Keunggulan ini didasarkan pada hasil validasi silang Stratified K-Fold,

yang menunjukkan bahwa model beroperasi dengan stabil, serta hasil uji McNemar, yang menunjukkan perbedaan kinerja statistik yang signifikan antara kedua algoritma. Akibatnya, algoritma Naive Bayes dianggap lebih efisien untuk menganalisis sentimen dalam ulasan aplikasi pinjaman online yang memanfaatkan representasi fitur TF-IDF pada dataset dengan distribusi kelas yang tidak seimbang. Dengan menerapkan berbagai algoritma klasifikasi dan menggunakan dataset yang lebih beragam, penelitian ini diharapkan dapat menjadi landasan bagi penelitian lebih lanjut mengenai pengembangan analisis sentimen di industri fintech.

DAFTAR PUSTAKA

- [1] L. S. Gunawan and C. S. T. Kansil, "Dampak Risiko Bagi Konsumen dalam Praktik Merugikan Pinjaman Online Ilegal," *J. Manaj. Pendidik. Dan Ilmu Sos.*, vol. 6, no. 1, pp. 461–469, 2024.
- [2] A. P. Setiawan, D. Christensen, N. Ramadhani, S. M. Ridwan, and S. N. Azizah, "Fenomena Pinjaman Online sebagai Bentuk Kenakalan Terselubung di Kalangan Generasi Z: Studi Perspektif Sosial dan Kaitannya dengan Tujuan Pembangunan Berkelanjutan (SDGs)," *J. Humaniora, Ekon. Syariah dan Muamalah*, vol. 3, no. 2, pp. 23–29, 2025, doi: 10.38035/jhesm.v3i2.323.
- [3] A. R. Putra and D. E. Ratnawati, "Analisis Sentimen Berbasis Aspek pada Aplikasi Mobile Menggunakan Naive Bayes berdasarkan Ulasan Pengguna Playstore (Studi Kasus : Jconnect Mobile)," *J. Teknol. Inf. dan Ilmu Komput.*, vol. 12, no. 2, pp. 293–300, 2025, doi: 10.25126/jtiik.2025127556.
- [4] F. A. Irawan, A. R. Atmadja, and A. Wahana, "Analisis Sentimen Ulasan Aplikasi Bank Digital Menggunakan Algoritma Naive Bayes," *J. Comput. Sci. Inf. Technol.*, vol. 4, no. 2, pp. 60–68, 2024, doi: 10.47065/explorer.v4i2.1181.
- [5] D. Munandar, M. Afdal, Z. Zarnelly, and R. Novita, "Analisis Sentimen Ulasan Pengguna Aplikasi Mobile Banking Menggunakan Algoritma K-Nearest Neighbor," *J. Teknol. Sist. Inf. dan Apl.*, vol. 7, no. 3, pp. 1309–1318, 2024, doi: 10.32493/jtsi.v7i3.41409.
- [6] K. Miranda and R. R. Suryono, "Analisis Sentimen Pinjaman Online: Studi Komparatif Algoritma Naive Bayes, Decision Tree, dan KNN," *J. Pendidik. Inform.*, vol. 9, no. 2, pp. 372–381, 2025, doi: 10.29408/edumatic.v9i2.30142.
- [7] S. Ainun Mardiah Hasibuan, "Klasifikasi Kelayakan Peminjaman Nasabah Menggunakan Algoritma K-Nearest Neighbor (K-NN)," *J. Mach. Learn. Comput. Sci. J.*, vol. 4, no. 10, pp. 1525–1532, 2024.
- [8] A. W. Septiano Anggun Pratama, "Perbandingan Klasifikasi Data Mining Untuk Prediksi Kredit Macet Dengan Menggunakan Metode Naive Bayes Dan K – Nearest Neighbors (Studi Kasus: Bank BUMN Di Daerah Jawa Barat)," *Simetris J. Tek. Mesin, Elektro dan Ilmu Komput.*, vol. 15, no. 2, pp. 1–13, 2024, doi: 10.24176/simet.v15i2.12321.
- [9] S. R. Ardhana, T. Widiarini, and B. A. Saputra, "Klasifikasi Menggunakan Algoritma K-Nearest Neighbor pada Imbalance Class Data dengan SMOTE. (Studi Kasus: Nasabah Bank Perkreditan Rakyat 'X')," *Indones. J. Appl. Stat.*, vol. 6, no. 2, p. 152, 2024, doi: 10.13057/ijas.v6i2.79389.
- [10] T. R. D. Putri1, M. R. Sanjaya, M. A. Firdaus, and D. R. Indah, "Implementasi Algoritma K-Nearest Neighbors dalam Analisis Sentimen Ulasan Aplikasi Bank Aladin," *J. Mach. Learn. Comput. Sci.*, vol. 6, no. April, pp. 608–618, 2026.
- [11] K. Mustaqim and I. N. D. Fatia Amalia Maresti, "Analisis Sentimen Ulasan Aplikasi PosPay untuk Meningkatkan Kepuasan Pengguna dengan Metode K-Nearest Neighbor (KNN)," *Edumatic J. Pendidik. Inform.*, vol. 8, no. 1, pp. 11–20, 2024, doi: 10.29408/edumatic.v8i1.24779.
- [12] R. N. Muhammad Iqbal, Afdal, "Implementasi Algoritma Support Vector Machine Untuk Analisa Sentimen Data Ulasan Aplikasi Pinjaman Online di Google Play Store," *Indones. J. Mach. Learn. Comput. Sci.*, vol. 4, no. 4, pp. 1244–1252, 2024, doi: 10.57152/malcom.v4i4.1435.
- [13] M. Rival Afandi, M. Afdal, Rice Novita and Fakultas, "Analisis Sentimen Masyarakat Terhadap Pinjaman Online di Twitter Menggunakan Algoritma Naive Bayes Classifier dan K-Nearest Neighbor," *Build. Informatics, Technol. Sci.*, vol. 6, no. 2, pp. 596–605, 2024, doi: 10.47065/bits.v6i2.5300.
- [14] S. Rifki Fahrezi Putra Ichsanasyah, "Perbandingan Algoritma Naive Bayes Dan Random Forest Dalam Klasifikasi Sentimen Ulasan Pengguna Kredivo Di Play Store," *J. Teknol. Inf. dan Ilmu Komput.*, vol. 13, no. 2, pp. 297–308, 2026.
- [15] M. D. Firmansyah, R. R. Setiawan, and Y. Irawan, "Analisis Perbandingan Kinerja Algoritma SVM, Naive Bayes Dan KNN Dalam Klasifikasi Sentimen Ulasan Aplikasi Pinterest Dengan SMOTE Dan PSO," *J. Teknol. dan Sist. Inf. Univrab*, vol. 11, no. 1, pp. 639–652, 2026.
- [16] I. M. T. Dian Tri Wilujeng, Mohamat Fatekurohman, "Analisis Risiko Kredit Perbankan Menggunakan Algoritma K-Nearest Neighbor dan Nearest Weighted K-Nearest Neighbor," *Indones. J. Appl. Stat.*, vol. 5, no. 2, pp. 142–148, 2023.
- [17] K. Nika Nur Cahayana, Ulfi Saidata Aesyi, "Model Analisis Sentimen pada Ulasan Pengguna Mobile Banking Menggunakan Kombinasi K-Means dan Naive Bayes," *Angkasa J. Ilm. Bid. Teknol.*, vol. 17, no. 1, p. 44, 2025, doi: 10.28989/angkasa.v17i1.2500.
- [18] N. Hafid and L. Khikmah, "Perbandingan Metode Algoritma C4 . 5 , Naive Bayes , dan Logistic Regression untuk Penentuan Kelayakan Penerima Kredit Comparison of C4 . 5 , Naive Bayes , and Logistic Regression Algorithm Methods for Determining Credit Recipient Eligibility," *Teknol. J. Ilm. Teknol. Sist. Inf.*, vol. 14, no. 2, pp. 85–93, 2024.