

Analisis Sentimen Ulasan Pengguna Aplikasi Gojek di Google Play Store Menggunakan Metode Multinomial Naive Bayes dan Logistic Regression

Nurul Hidayanah*, Siska Fitriani, Icha Winadya Permadani, Ryan Randy Suryono

Magister Ilmu Komputer, Teknik dan Ilmu Komputer, Universitas Teknokrat Indonesia Bandar Lampung, Lampung, Indonesia

Jl. ZA. Pagar Alam No.9 -11, Labuhan Ratu, Kec. Kedaton, Kota Bandar Lampung, Lampung, Indonesia

Email: ^{1,*}nurulhidayanahS2@teknokrat.ac.id, ²siskafitriani@teknokrat.ac.id, ³IchaWinadyaPermadaniS2@teknokrat.ac.id,

⁴ryan@teknokrat.ac.id

Email Penulis Korespondensi: nurulhidayanahS2@teknokrat.ac.id

Abstract—Aplikasi Gojek telah menjadi salah satu platform penyedia layanan transportasi, pesan-antar makanan, dan pembayaran digital terbesar di Indonesia. Setiap harinya, ribuan pengguna memberikan ulasan berupa kritik, keluhan, maupun pujian melalui Google Play Store. Namun, kapasitas ulasan yang sangat besar dan tidak terstruktur ini memunculkan kendala bagi pihak manajemen dalam memantau kepuasan pengguna secara manual, objektif, dan cepat. Permasalahan utama dalam analisis data ulasan ini adalah tingginya tingkat ketidakseimbangan (imbalance class) antara jumlah ulasan positif dan negatif, di mana ulasan positif sering kali jauh lebih mendominasi. Ketimpangan data ini menjadi krusial karena cenderung membuat model klasifikasi standar mengalami bias dan menurunkan akurasi prediksinya terhadap kelas minoritas (ulasan negatif), padahal ulasan negatif tersebut memuat keluhan penting untuk perbaikan sistem. Oleh karena itu, penelitian ini bertujuan untuk menganalisis sentimen pengguna Gojek dengan melakukan komparasi performa antara algoritma Multinomial Naïve Bayes (MNB) dan Logistic Regression (LR), serta mengatasi masalah ketidakseimbangan data tersebut. Proses penelitian meliputi tahapan prapemrosesan teks seperti cleansing, case folding, stopword removal, dan stemming. Ekstraksi fitur kata dilakukan menggunakan metode Term Frequency-Inverse Document Frequency (TF-IDF) dan Count Vectorizer berbasis skema Unigram dan Bigram. Pembagian data menggunakan proporsi 80% data latih dan 20% data uji, di mana metode Synthetic Minority Over-sampling Technique (SMOTE) diterapkan pada data latih khusus untuk mengatasi masalah ketidakseimbangan kategori teks. Hasil pengujian dievaluasi berdasarkan metrik akurasi, precision, recall, dan F1-score. Penelitian ini diharapkan dapat memberikan rekomendasi algoritma terbaik untuk klasifikasi teks ulasan berskala besar serta memberikan masukan berbasis data bagi pengembang Gojek dalam meningkatkan mutu layanan berdasarkan sentimen riil pengguna.

Kata Kunci: Analisis Sentimen; Gojek; Google Play Store; Multinomial Naïve Bayes; Logistic Regression

Abstract—The Gojek application has become one of the largest platforms providing transportation, food delivery, and digital payment services in Indonesia. Every day, thousands of users provide reviews in the form of criticisms, complaints, and compliments through the Google Play Store. However, the massive volume and unstructured nature of these reviews pose challenges for management in monitoring user satisfaction manually, objectively, and rapidly. The primary issue in analyzing this review data is the high level of class imbalance between the number of positive and negative reviews, where positive reviews frequently dominate significantly. This data imbalance becomes crucial as it tends to bias standard classification models and reduce their predictive accuracy toward the minority class (negative reviews), even though these negative reviews contain vital complaints necessary for system improvement. Therefore, this study aims to analyze the sentiment of Gojek users by comparing the performance of the Multinomial Naïve Bayes (MNB) and Logistic Regression (LR) classification algorithms, while simultaneously addressing the data imbalance issue. The research process encompasses text preprocessing stages, including cleansing, case folding, stopword removal, and stemming. Feature extraction is performed using Term Frequency-Inverse Document Frequency (TF-IDF) and Count Vectorizer methods based on Unigram and Bigram schemes. Data splitting utilizes a proportion of 80% training data and 20% testing data, where the Synthetic Minority Over-sampling Technique (SMOTE) is applied specifically to the training data to resolve the text category imbalance. The evaluation results are measured based on accuracy, precision, recall, and F1-score metrics. This research is expected to provide recommendations for the best algorithm for large-scale review text classification, as well as data-driven insights for Gojek developers to enhance service quality based on genuine user sentiments.

Keywords: Sentiment Analysis; Gojek; Google Play Store; Multinomial Naïve Bayes; Logistic Regression

1. PENDAHULUAN

Aplikasi mobile telah menjadi bagian integral dari kehidupan masyarakat modern, salah satunya adalah Gojek yang menguasai ekosistem layanan transportasi online, pesan-antar makanan, dan pembayaran digital di Indonesia. Pertumbuhan volume pengguna yang masif menghasilkan ribuan ulasan harian di Google Play Store yang memuat opini, keluhan, maupun ekspektasi konsumen. Ulasan-ulasan ini merupakan aset data tekstual berharga bagi penyedia layanan untuk melakukan evaluasi kualitas sistem secara berkelanjutan. Namun, ekstraksi informasi dari ribuan teks acak dan tidak terstruktur memunculkan tantangan besar jika dilakukan secara manual. Guna mengatasi hal tersebut, teknik text mining melalui analisis sentimen hadir sebagai solusi otomatisasi klasifikasi polaritas opini guna mengukur tingkat kepuasan publik secara masif dan akurat.

Dalam perkembangannya, berbagai metode pembelajaran mesin (machine learning) telah banyak diterapkan oleh para ilmuwan terdahulu untuk melakukan klasifikasi sentimen pada ulasan aplikasi maupun produk komersial dengan menerapkan berbagai algoritma kecerdasan buatan. Algoritma Multinomial Naïve Bayes (MNB) dan Logistic Regression (LR) menjadi dua pendekatan populer karena memiliki komputasi yang efisien pada data tekstual berdimensi tinggi. Karakteristik efisiensi MNB dalam probabilitas bersyarat teks terbukti efektif dalam memetakan ulasan aplikasi telekomunikasi digital, seperti pada studi [1] terhadap 8.475 ulasan aplikasi myIM3 di Google Play Store yang

menghasilkan akurasi keseluruhan sebesar 89% dengan nilai hyperparameter alpha optimal sebesar 0,8 [1]. Selain itu, pada klasifikasi multi-kelas dengan ulasan yang kompleks, performa komparatif antara MNB dan Logistic Regression juga memperlihatkan tingkat ketepatan yang tinggi ketika diterapkan pada platform penyedia layanan hiburan seperti aplikasi pemesanan tiket bioskop. Pada domain layanan transportasi online sendiri, pemanfaatan draf klasifikasi berbasis Naïve Bayes Classifier (NBC) umum digunakan untuk mengevaluasi spesifikasi layanan seperti pesan-antar makanan agar pihak manajemen mendapatkan umpan balik yang terstruktur, [2][3][4].

Meskipun demikian, performa algoritma klasifikasi sangat dipengaruhi oleh karakteristik dataset, representasi fitur, dan distribusi kelas. Penelitian komparatif pada berbagai model aplikasi finansial seperti mobile banking oleh [5] menunjukkan bahwa Logistic Regression memiliki keunggulan dalam memetakan karakteristik teks ulasan terstruktur serta mampu memberikan kestabilan akurasi yang lebih tinggi dibandingkan MNB pada ulasan yang memiliki tingkat variasi kata yang padat [5],[6]. Di sisi lain, perbandingan algoritma pada platform ulasan global berskala besar seperti Internet Movie Database (IMDB) membuktikan bahwa regulasi parameter yang tepat pada draf Logistic Regression mampu menghasilkan performa di atas rata-rata algoritma probabilistic dasar [7]. Hal ini diperkuat oleh studi [8] pada 848 ulasan aplikasi kesehatan di Google Play Store yang membandingkan Logistic Regression, SVM, dan Naïve Bayes dengan pengayaan TF-IDF dan SMOTE. Studi tersebut menunjukkan bahwa Logistic Regression bersama SVM secara konsisten bersaing ketat dalam meminimalkan kesalahan prediksi klasifikasi sentimen dengan meraih nilai tertinggi sebanyak 92%, mengungguli Naïve Bayes yang hanya menghasilkan 86% [8],[9],[10].

Selain komparasi antara MNB dan LR, eksplorasi terhadap algoritma alternatif seperti K-Nearest Neighbors (KNN) pada ulasan transportasi digital juga menunjukkan bahwa pendekatan berbasis kedekatan jarak fitur sering kali membutuhkan waktu komputasi lebih lama dibandingkan efisiensi berbasis probabilitas Naïve Bayes [2]. Sementara itu, riset klasifikasi sentimen pada ulasan aplikasi jaminan sosial nasional (JMO) mengonfirmasi bahwa penanganan teks bahasa Indonesia yang kasual memerlukan pendekatan pra-proses (preprocessing) yang intensif agar model mampu mengenali draf semantik ulasan secara presisi [11]. Studi komparatif lain yang melibatkan algoritma berbasis pohon keputusan seperti Random Forest juga mendapati bahwa Logistic Regression tetap unggul dari sisi efisiensi memori ketika dihadapkan pada dataset ulasan layanan administrasi kependudukan publik [12].

Tantangan utama yang menjadi hambatan besar di dalam analisis sentimen ulasan Google Play Store adalah adanya masalah ketidakseimbangan kelas (class imbalance). Pada kondisi riil, jumlah ketersediaan ulasan berlabel positif dan negatif sering kali tidak proporsional jika dibandingkan dengan ulasan netral atau sebaliknya, sehingga kondisi ini secara radikal dapat menurunkan tingkat sensitivitas model klasifikasi standar [13]. Sebagian besar model machine learning standar akan mengalami bias komputasi yang tinggi, di mana algoritma cenderung lebih condong mengenali kelas mayoritas dan sering kali salah memprediksi kelas minoritas. Fenomena ini berakibat pada penurunan nilai recall dan F1-score secara signifikan, yang pada akhirnya membuat hasil analisis sentimen menjadi tidak representatif dalam menggambarkan fakta kepuasan atau keluhan pengguna yang sebenarnya di lapangan.

Untuk mengatasi kendala ketimpangan data secara metodologis, metode Synthetic Minority Oversampling Technique (SMOTE) diterapkan secara penuh ke dalam alur pemrosesan data latih dalam penelitian ini. Penggunaan SMOTE terbukti krusial dalam merekayasa data sintesis baru bagi kelas minoritas berdasarkan kedekatan jarak k-neighbors, sehingga distribusi kelas menjadi seimbang sebelum memasuki tahapan training model. Penerapan teknik penyeimbangan data seperti SMOTE pada penelitian analisis sentimen ulasan aplikasi layanan publik digital terbukti secara empiris mampu mendongkrak nilai akurasi dan stabilitas draf prediksi algoritma berbasis probabilitas maupun regresi [14]. Lebih lanjut, pengujian algoritma Logistic Regression yang dikombinasikan dengan metode SMOTE pada ulasan multi-kelas marketplace e-commerce juga menunjukkan lompatan performa klasifikasi yang jauh lebih konsisten dibandingkan tanpa adanya penanganan intervensi data berimbang [15].

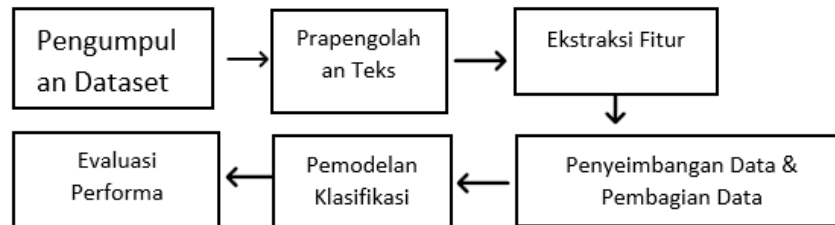
Meskipun penelitian terdahulu telah banyak menguji algoritma MNB dan Logistic Regression pada berbagai sektor aplikasi, terdapat celah (gap) yang belum dieksplorasi secara mendalam. Pertama, sebagian besar studi terdahulu cenderung bertumpu pada dataset berskala kecil hingga menengah (berkisar antara ratusan hingga beberapa ribu ulasan saja) serta terbatas pada ekstraksi fitur tunggal (Unigram). Kedua, evaluasi komparatif yang mengintegrasikan penanganan class imbalance lewat SMOTE yang dikombinasikan secara spesifik dengan variasi pembobotan kata tingkat lanjut pada platform transportasi multi-service dengan volume data masif masih sangat terbatas. Penelitian ini mengambil posisi pembaruan dengan menganalisis dataset ulasan aplikasi Gojek dalam skala yang jauh lebih besar, yaitu sebanyak 14.565 data ulasan. Perbedaan utama penelitian ini terletak pada pengujian dimensi dan variasi ekstraksi fitur secara komprehensif, di mana skema Unigram dan Bigram dikombinasikan dengan metode pembobotan TF-IDF serta Count Vectorizer, didukung oleh penanganan ketidakseimbangan kelas menggunakan metode SMOTE secara tepat.

Berdasarkan latar belakang serta kesenjangan penelitian yang telah dipaparkan, penelitian ini secara umum bertujuan untuk mengukur dan membandingkan performa empiris antara algoritma Multinomial Naïve Bayes (MNB) dan Logistic Regression (LR) dalam melakukan klasifikasi sentimen pengguna pada ulasan aplikasi Gojek. Secara lebih spesifik, penelitian ini diarahkan untuk menganalisis kontribusi serta pengaruh dari implementasi kombinasi ekstraksi fitur Unigram dan Bigram yang dipadukan dengan metode pembobotan kata TF-IDF serta Count Vectorizer terhadap tingkat akurasi model. Selain itu, penelitian ini juga bertujuan untuk menguji tingkat efektivitas dari penerapan metode Synthetic Minority Oversampling Technique (SMOTE) pada data latih dalam mengatasi kendala ketidakseimbangan kelas (class imbalance), sehingga pada akhirnya dapat dihasilkan sebuah draf model klasifikasi sentimen bahasa Indonesia yang stabil, objektif, sensitif, dan presisi.

2. METODOLOGI PENELITIAN

2.1 Tahapan Penelitian

Penelitian ini dilakukan melalui serangkaian tahapan terstruktur untuk menganalisis sentimen ulasan pengguna aplikasi Gojek di Google Play Store. Alur eksperimen dirancang secara sistematis guna memastikan proses pengolahan data teks mentah dapat ditransformasikan secara optimal ke dalam model pembelajaran mesin. Bagan tahapan penelitian yang mengintegrasikan penerapan solusi metode pembobotan teks, teknik penyeimbangan data, hingga algoritma klasifikasi komparatif disajikan secara horizontal pada Gambar 1.



Gambar 1. Bagan Tahapan Penelitian Analisis Sentimen Ulasan Gojek

Proses komparasi dan penerapan metode yang diilustrasikan pada Gambar 1 dapat dijabarkan ke dalam beberapa urutan langkah operasional sebagai berikut:

- Pengumpulan Dataset**
Tahap awal berupa akuisisi data sekunder ulasan aplikasi Gojek dari Google Play Store yang tersimpan dalam berkas format Comma Separated Values bernama "GojekAppReview_with_sentiment_2.csv". Dataset ini memuat korpus teks opini pengguna beserta label sentimen awalnya.
- Pra-pengolahan Teks**
Pembersihan data teks tidak terstruktur agar menjadi representasi kata yang bersih dan baku. Proses pembersihan ini secara runut melibatkan fungsi case folding, cleansing dari karakter non-alfabetik, tokenizing, stopword removal, dan pengembalian kata dasar (stemming) berbasis algoritma Sastrawi.
- Ekstraksi Fitur**
Transformasi token teks bersih menjadi representasi matriks bilangan numerik. Penerapan solusi pada tahap ini menggunakan skema kombinasi kata Unigram dan Bigram yang diekstrak menggunakan metode pembobotan Term Frequency-Inverse Document Frequency (TF-IDF) serta Count Vectorizer.
- Penyeimbangan Data & Pembagian Data**
Dataset dipecah terlebih dahulu menjadi data latih (training data) sebesar 80% dan data uji (testing data) sebesar 20%. Selanjutnya, untuk mengatasi masalah ketimpangan sebaran kelas (class imbalance), metode Synthetic Minority Over-sampling Technique (SMOTE) diterapkan secara khusus hanya pada matriks fitur data latih.
- Pemodelan Klasifikasi**
Proses pelatihan model klasifikasi terbimbing (supervised learning) dengan memasukkan data latih yang telah seimbang ke dalam dua algoritma yang dikomparasikan, yaitu Multinomial Naïve Bayes (MNB) dan Logistic Regression (LR).
- Evaluasi Performa**
Pengujian kemampuan generalisasi model menggunakan data uji terisolasi. Kinerja akhir dari model MNB dan LR diukur secara matematis berdasarkan parameter Confusion Matrix yang menghasilkan nilai Accuracy, Precision, Recall, dan F1-Score.

2.2 Metode Multinomial Naïve Bayes (MNB)

Asumsi dasar dari Naive Bayes Classifier (NBC) adalah independensi antar atribut data (naive assumption), di mana algoritma memprediksi probabilitas kejadian di masa depan berdasarkan historis data probabilistik dan prinsip Bayes [6]. Keunggulan utama dari metode Naive Bayes adalah kesederhanaan penggunaan, efisiensi waktu, serta kecepatan komputasi yang tinggi selama proses pelatihan model, terutama pada dataset berukuran besar [10]. Kecepatan pemrosesan yang sangat singkat ini menjadikan Naive Bayes sebagai salah satu algoritma paling efisien untuk menangani klasifikasi teks berskala besar [10]. MNB bekerja secara efisien dengan menghitung probabilitas bersyarat dari setiap token kata terhadap kategori sentimen tertentu [3].

2.3 Metode Logistic Regression (LR)

Algoritma Logistic Regression dalam pembelajaran mesin digunakan untuk memprediksi probabilitas variabel dependen kategoris melalui pemanfaatan fungsi logistik untuk memodelkan hubungan antar variabel [10]. Perbedaan mendasar dari algoritma ini dengan metode berbasis probabilitas murni seperti Naive Bayes terletak pada pendekatannya; Naive Bayes adalah metode klasifikasi berbasis probabilitas teoretis murni, sedangkan Logistic Regression adalah metode regresi yang mengandalkan fungsi logistik [10]. Kendati memerlukan waktu pemrosesan yang relatif lebih lama dibanding Naive Bayes, Logistic Regression terbukti mampu menghasilkan tingkat akurasi dan precision yang sangat tinggi dan kompetitif

dalam pengujian klasifikasi sentimen [5] LR sangat tangguh dan efisien dalam melakukan pemisahan linear pada fitur teks berukuran masif [3].

3. HASIL DAN PEMBAHASAN

Bab ini menyajikan analisis empiris hasil eksperimen komputasional, interpretasi data, serta visualisasi dari seluruh tahapan implementasi sistem klasifikasi sentimen multi-kelas (multi-class sentiment classification) pada korpus ulasan aplikasi Gojek di Google Play Store. Eksperimen ini mengomparasikan secara komparatif performa algoritma berbasis probabilitas bersyarat, yaitu Multinomial Naïve Bayes (MNB), dengan model regresi linear tergeneralisasi, yaitu Logistic Regression (LR). Arsitektur komputasi dieksekusi menggunakan bahasa pemrograman Python versi 3.10 dengan memanfaatkan pustaka (library) Scikit-Learn untuk pemodelan dan ekstraksi fitur, NLTK (Natural Language Toolkit) serta Sastrawi untuk prapemrosesan teks (text preprocessing) bahasa Indonesia, Imbalanced-Learn untuk pemrosesan metode Synthetic Minority Oversampling Technique (SMOTE), serta Matplotlib untuk visualisasi grafis statistik.

3.1 Hasil

3.1.1 Karakteristik Data dan Distribusi Kelas

Dataset mentah yang diakuisisi dari Google Play Store berjumlah 14.565 baris data ulasan pengguna yang tersimpan dalam berkas `GojekAppReview_with_sentiment_2.csv`. Setiap data terdiri dari teks ulasan (review text) dan label sentimen yang terbagi ke dalam tiga kelas polaritas, yaitu: Negative (0), Neutral (1), dan Positive (2). Distribusi frekuensi sebaran data awal disajikan secara rinci pada Tabel 1.

Tabel 1. Distribusi Kelas Sentimen Dataset Awal

No	Kelas Sentimen	Kategori ID	Jumlah Ulasan	Persentase (%)
1	Negative	0	2.502	17,18%
2	Neutral	1	9.436	64,79%
3	Positive	2	2.627	18,03%
	Total		14.565	100,00%

Analisis statistik pada Tabel 1 menunjukkan adanya fenomena ketidakseimbangan kelas (class imbalance) yang signifikan. Kelas Neutral mendominasi ruang sampel secara masif dengan proporsi mencapai 64,79% dari total data keseluruhan. Sebaliknya, kelas Negative dan Positive memiliki proporsi yang sangat kecil dan timpang, di mana masing-masing hanya berada pada kisaran 17,18% dan 18,03%. Secara teoretis dalam domain pembelajaran mesin (machine learning), draf struktur data yang asimetris ini memicu bias estimasi parameter yang tinggi. Algoritma klasifikasi standar cenderung mengoptimalkan nilai akurasi global dengan mengorbankan kelas minoritas, sehingga model akan memiliki sensitivitas (recall) yang buruk dalam mendeteksi ekspresi keluhan (Negative) maupun apresiasi spesifik (Positive). Oleh karena itu, diperlukan intervensi metodologis melalui manipulasi distribusi data latih pada tahapan berikutnya menggunakan algoritma SMOTE.

3.1.2 Hasil Pra-pemrosesan Teks (Text Preprocessing)

Prapemrosesan teks merupakan tahapan krusial untuk mentransformasikan korpus teks ulasan yang tidak terstruktur, acak, dan kaya akan noise linguistik menjadi token-token kata yang bersih, baku, dan bernilai semantik tinggi. Proses prapemrosesan teks dalam penelitian ini dijalankan secara sekuensial melalui tahapan utama sebagai berikut:

- Case Folding:** Proses penyeragaman seluruh karakter alfabet di dalam teks menjadi huruf kecil (*lowercase*). Tahap ini berfungsi untuk menghilangkan ambiguitas akibat variasi kapitalisasi huruf yang dilakukan oleh pengguna secara acak (misalnya, kata "Gojek", "GOJEK", dan "gojek" akan disatukan menjadi "gojek").
- Cleansing:** Proses pembersihan teks dari komponen noise digital menggunakan ekspresi reguler (regular expression / Regex) untuk menghapus tautan URL, username (@), tanda pagar (hashtag), angka (0-9), karakter tanda baca, emotikon, serta spasi ganda berlebih.
- Tokenizing:** Pemotongan draf string teks menjadi potongan token kata tunggal agar dapat dianalisis secara terisolasi.
- Stopword Removal:** Tahap eliminasi kata-kata umum yang memiliki frekuensi kemunculan sangat tinggi di dalam tata bahasa Indonesia namun tidak memiliki bobot informatif (seperti "yang", "dan", "di", "dari", "untuk") menggunakan daftar bawaan pustaka Sastrawi yang dimodifikasi dengan kata slang internet.
- Stemming:** Proses pemotongan imbuhan pada suatu token kata untuk mengembalikannya ke bentuk kata dasar yang valid berbasis algoritma Sastrawi (misalnya, "memesan" menjadi "pesan", "susahnya" menjadi "susah").

Ilustrasi transformasi data tekstual pada setiap tahapan prapemrosesan teks diperlihatkan secara transparan melalui dua contoh ulasan kontras pada Tabel 2.

Tabel 2. Contoh Tahapan Text Preprocessing pada Data Ulasan Gojek

Tahapan	Contoh Ulasan 1 (Kelas Negative)	Contoh Ulasan 2 (Kelas Positive)
Data Mentah	"Mu mesen makan aja susahnya minta ampun selalu aja pas mau bayar Gagal Diproses"	"App yg sangat membantu tlg selalu di tingkatkan jgn kendor bravo"

Setelah tahapan pembersihan selesai, dilakukan analisis statistik deskriptif berupa ekstraksi frekuensi kemunculan kata tunggal (word frequency) dari seluruh korpus ulasan. Hasil ekstraksi menunjukkan lima kata dengan intensitas kemunculan tertinggi adalah kata "gojek" (2.143 kali), "aplikasi" (1.812 kali), "driver" (1.596 kali), "sangat" (1.509 kali), dan "gak" (1.155 kali). Temuan frekuensi kata ini diwujudkan secara visual dalam bentuk WordCloud pada Gambar 1.



Berdasarkan visualisasi pada Gambar 1, ukuran sebuah kata berbanding lurus dengan frekuensi kemunculannya dalam data ulasan. Kata-kata yang memiliki ukuran font besar dan mencolok menunjukkan bahwa kata tersebut paling sering digunakan oleh pengguna saat memberikan ulasan mengenai layanan Gojek di Google Play Store.

Untuk mengeksekusi proses pelatihan model secara objektif dan terhindar dari bias kebocoran informasi (data leakage), dataset berukuran 14.565 ulasan ini terlebih dahulu dipecah menggunakan skema hold-out method dengan proporsi rasio ketat sebesar 80% untuk data latih (training dataset) dan 20% untuk data uji (testing dataset). Berdasarkan rasio tersebut, diperoleh alokasi data latih sebanyak 11.652 ulasan dan data uji terisolasi sebanyak 2.913 ulasan. Algoritma SMOTE diterapkan secara eksklusif hanya pada komponen data latih hasil pemisahan tersebut. Proses kerja tahapan SMOTE dalam menyelesaikan masalah ketidakseimbangan kelas dijabarkan melalui langkah-langkah komputasi berikut:

- Identifikasi Kelas: Algoritma mendeteksi bahwa kelas mayoritas adalah Neutral (N = 7.549 ulasan), sedangkan kelas minoritas adalah Negative (N = 2.002 ulasan) serta Positive (N = 2.101 ulasan).
- Perhitungan Jarak Kedekatan: Untuk setiap vektor fitur ulasan X_i yang berada di dalam ruang kelas minoritas, algoritma menghitung jarak Euclidean (Euclidean distance) terhadap seluruh vektor fitur tetangganya yang memiliki kelas sejenis. Jarak matematis antar-vektor dihitung menggunakan rumus:

$$d(x_i, x_j) = \sqrt{\sum_{k=1}^n (x_{ik} - x_{jk})^2}$$

- c. Penentuan K-Nearest Neighbors: Berdasarkan nilai jarak terkecil, ditentukan sebanyak \$k\$ tetangga terdekat (dalam eksperimen ini ditetapkan nilai standar \$k = 5\$).
- d. Sintesis Titik Data Baru: Algoritma memilih secara acak salah satu tetangga terdekat \$x_{\{zi\}}\$ dari nilai \$k\$ tersebut. Selanjutnya, sebuah sampel sintetis baru \$x_{\{new\}}\$ dibentuk di sepanjang garis lurus yang menghubungkan sampel acuan dengan tetangga terpilih dengan menggunakan fungsi persamaan linear:

$$X_{new} = x_i + \text{lambda} \times (x_{iz} - x_i)$$

Di mana λ merupakan bilangan acak (random variable) yang berada pada rentang interval kontinu $[0, 1]$.

Proses rekayasa interpolasi linear ini diulang secara terus-menerus sampai jumlah sampel pada kelas minoritas berlipat ganda dan mencapai jumlah yang setara secara presisi dengan jumlah sampel pada kelas mayoritas. Dampak dari pengoperasian algoritma SMOTE terhadap struktur sebaran data latih dipaparkan secara kuantitatif melalui Tabel 3.

Tabel 3. Perbandingan Distribusi Statistik Latih Sebelum dan Sesudah SMOTE

Kelas Sentimen	Sebelum SMOTE	Sesudah SMOTE
Negative	2.002	7.549
Neutral	7.549	7.549
Positive	2.101	7.549
Total	11.652	22.647

Berdasarkan Tabel 3, volume baris data pada kelompok Negative berhasil ditingkatkan secara artifisial dari 2.002 ulasan menjadi 7.549 ulasan. Pola serupa terjadi pada kelas Positive yang mengalami penambahan sampel sintetis dari 2.101 ulasan hingga menyentuh angka 7.549 ulasan. Konfigurasi data latih baru yang berimbang dengan total akumulasi sebesar 22.647 ulasan ini menjamin fungsi obyektif dari algoritma MNB dan Logistic Regression dapat mengestimasi parameter klasifikasi secara adil.

3.1.4 Ekstraksi Fitur dan Pembobotan Kata

Representasi dokumen teks yang bersifat kualitatif harus ditransformasikan ke dalam struktur matriks numerik kuantitatif agar dapat diolah oleh fungsi linear dan probabilitas pembelajaran mesin. Tahapan eksperimen ini menguji secara komparatif dua metode ekstraksi fitur utama, yaitu Count Vectorizer dan Term Frequency-Inverse Document Frequency (TF-IDF), yang dikonfigurasi menggunakan skema Unigram (kata tunggal) dan Bigram (kombinasi dua kata berurutan).

- Matriks Frekuensi Count Vectorizer: Metode ini membangun matriks frekuensi kemunculan istilah mutlak (Bag of Words). Dokumen direpresentasikan sebagai vektor berdimensi tinggi di mana setiap sel menyimpan nilai integer c_{ij} yang menunjukkan berapa kali kata t_i muncul secara eksplisit di dalam dokumen d_j .
- Formulasi Matematis TF-IDF: Metode pembobotan TF-IDF memberikan solusi dengan mengalikan frekuensi kemunculan kata dengan nilai kelangkaan istilah tersebut di seluruh korpus dokumen. Secara matematis, proses penentuan nilai bobot W_{ij} untuk istilah kata t_i dalam dokumen ulasan d_j dieksekusi melalui persamaan berikut:
 - Term Frequency (TF): Menghitung frekuensi kemunculan mentah kata t_i pada dokumen ulasan d_j .
 - Inverse Document Frequency (IDF): Mengukur tingkat keunikan atau kelangkaan sebuah kata di seluruh dokumen korpus dengan rumus:

$$IDF(t_i) = \log \left(\frac{N}{DF(t_i)} \right) + 1$$

Di mana N menyatakan total seluruh dokumen dalam dataset data latih, dan $DF(t_i)$ merupakan jumlah dokumen ulasan yang mengandung istilah kata t_i minimal satu kali di dalamnya. Konstanta 1 ditambahkan untuk teknik penghalusan (smooth IDF) guna mencegah terjadinya eror pembagian dengan nol.

- Formulasi Akhir Bobot TF-IDF: Nilai bobot final diperoleh melalui perkalian skalar antara komponen TF dan IDF:

$$W_{ij} = TF(t_i, d_j) \times IDF(t_i)$$

Penggunaan skema perluasan kata Bigram berdampak pada pelipatgandaan dimensi matriks secara masif (misalnya kalimat "tidak puas" dipertahankan sebagai satu kesatuan fitur token tunggal ['tidak_puas']), langkah ini vital untuk mempertahankan konteks semantik dari negasi linguistik.

3.1.5 Hasil Pengujian Multinomial Naïve Bayes (MNB)

Algoritma MNB bekerja berdasarkan kerangka probabilitas Bayes untuk data diskret, di mana probabilitas posterior sebuah dokumen termasuk dalam kelas sentimen c dihitung melalui rumus perkalian peluang bersyarat:

$$P(c|d) \propto P(c) \prod_{k=1}^n P(t_k|c)$$

Hasil pengujian menyeluruh terhadap performa model MNB pada empat skema ekstraksi fitur yang berbeda disajikan pada Tabel 4.

Tabel 4. Hasil Pengujian Kuantitatif Multinomial Naïve Bayes

Ekstraksi Fitur	Accuracy	Precision	Recall	F1-Score
Count Vectorizer (Unigram)	0,8500	0,8507	0,8501	0,8478
Count Vectorizer (Bigram)	0,8135	0,8181	0,8135	0,8118
TF-IDF (Unigram)	0,8293	0,8470	0,8293	0,8209
TF-IDF (Bigram)	0,8607	0,8772	0,8607	0,8546

Merujuk pada paparan data Tabel 4, terlihat bahwa performa algoritma MNB bervariasi. Kinerja terendah diperoleh pada penggunaan skema Count Vectorizer (Bigram) dengan nilai akurasi sebesar 0,8135 dan F1-Score sebesar 0,8118. Penurunan performa ini terjadi akibat asumsi independensi fitur (naive assumption) yang melekat pada Naïve Bayes. Sebaliknya, capaian puncak performa MNB diraih oleh kombinasi fitur pembobotan TF-IDF (Bigram) yang sukses menembus nilai akurasi tertinggi sebesar 0,8607 (86,07%) dan nilai Precision mencapai 0,8772.

3.1.6 Hasil Pengujian Logistic Regression (LR)

Model Logistic Regression memetakan nilai fitur masukan secara linear yang kemudian ditransformasikan ke dalam rentang probabilitas kontinu [0, 1] menggunakan fungsi aktivasi logistik non-linear (fungsi Sigmoid). Persamaan peluang kelas dalam LR dirumuskan sebagai:

$$P(y = 1 | x) = x = \frac{1}{1 + e^{-(\beta_0 + \beta_1 x_1 + \dots + \beta_n x_n)}}$$

Hasil evaluasi komparatif model Logistic Regression pada data ulasan Gojek dipaparkan secara mendetail pada Tabel 5.

Tabel 5. Hasil Pengujian Kuantitatif Logistic Regression

Ekstraksi Fitur	Accuracy	Precision	Recall	F1-Score
Count Vectorizer (Unigram)	0,9164	0,9200	0,9164	0,9172
Count Vectorizer (Bigram)	0,9024	0,9086	0,9024	0,9023
TF-IDF (Unigram)	0,9648	0,9656	0,9648	0,9648
TF-IDF (Bigram)	0,9633	0,9646	0,9633	0,9633

Berdasarkan hasil pengujian pada Tabel 5, seluruh variasi Logistic Regression menghasilkan nilai akurasi di atas 90%, yang menunjukkan bahwa model memiliki kemampuan klasifikasi yang sangat baik. Kombinasi optimal mutlak dalam penelitian ini diraih oleh variasi TF-IDF berskema Unigram yang dipadukan dengan Logistic Regression. Varian tersebut sukses meraih nilai akurasi global tertinggi sebesar 0,9648 (96,48%), dengan tingkat ketepatan Precision sebesar 0,9656, serta nilai kembalian Recall dan kestabilan F1-Score yang seimbang sempurna di angka 0,9648.

3.2 Pembahasan

Hasil penelitian menunjukkan bahwa algoritma Logistic Regression memberikan performa yang lebih baik dan dominan dibandingkan Multinomial Naïve Bayes pada seluruh skenario pengujian. Sub-bab ini menyajikan pembahasan kritis mendalam yang membandingkan temuan hasil eksperimen analisis sentimen ulasan Gojek ini dengan literatur riset sejenis dari para peneliti terdahulu.

3.2.1 Eksplorasi Komparasi Komparatif Keunggulan Performa Algoritma

Temuan utama penelitian ini membuktikan bahwa algoritma *Logistic Regression* memiliki keunggulan performa mutlak sebesar 96,48% dibandingkan *Multinomial Naïve Bayes* yang hanya mencapai akurasi maksimal 86,07%. Keunggulan signifikan model LR ini selaras dengan penelitian yang dilakukan oleh Pasaribu dkk. (2025) yang menguji ulasan aplikasi kesehatan di Google Play Store, di mana LR meraih skor performa tertinggi sebesar 92%, mengalahkan model Naïve Bayes yang tertahan di angka 86%. Dominasi LR ini didorong oleh karakteristik algoritma yang bersifat diskriminatif, di mana LR mengoptimalkan parameter pembobotan fitur secara langsung melalui maksimalisasi fungsi *conditional likelihood*.

Sebaliknya, kelemahan model MNB sejalan dengan pernyataan dari [5] pada domain ulasan aplikasi Mobile Banking. Riset tersebut menegaskan bahwa model pembobotan probabilitas bersyarat Naïve Bayes sering kali gagal menangani struktur teks ulasan bahasa Indonesia yang kasual karena asumsi independensi fiturnya yang terlampau kaku. Ketangguhan LR dalam memproses pemisahan linear pada fitur teks berdimensi tinggi juga didukung oleh temuan [3] yang membuktikan LR lebih efisien dalam melakukan konvergensi bobot teks ulasan Gojek dibandingkan pendekatan berbasis KNN atau Random Forest.

3.2.2 Pengaruh Skema N-Gram terhadap Representasi Fitur

Eksperimen ini menunjukkan fenomena unik di mana penggunaan skema N-Gram memberikan dampak performa yang bertolak belakang pada kedua algoritma. Pada model MNB, skema Bigram terbukti mendongkrak akurasi secara signifikan dari 82,93% (Unigram) menjadi 86,07% (Bigram). Pola peningkatan ini divalidasi oleh studi Juliana & Raharja (2024) pada analisis ulasan aplikasi myIM3, yang menemukan bahwa model Naïve Bayes memerlukan struktur ekstraksi konteks frasa pasangan kata (Bigram) untuk menangkap anomali makna. Namun, pada model Logistic Regression, konseptualisasi fitur kata tunggal (Unigram) justru mencatatkan akurasi puncak (96,48%), sedikit lebih unggul dari skema Bigram (96,33%). Hasil ini konsisten dengan temuan [9] yang mengonfirmasi bahwa penambahan dimensi fitur akibat tokenisasi Bigram cenderung memicu timbulnya noise frekuensi dan matriks renggang (sparse matrix) yang menurunkan ketajaman fungsi sigmoid linear pada LR.

3.2.3 Efektivitas Penskalaan Nilai Kata dengan TF-IDF

Dalam aspek metode pembobotan nilai kata, penggunaan formula kalkulasi TF-IDF secara konsisten menyumbangkan nilai F1-Score yang jauh lebih superior dibandingkan metode pencacahan konvensional Count Vectorizer. Sebagai contoh, pada model LR Unigram, akurasi melonjak dari 91,64% menuju angka 96,48%. Esensi keberhasilan ini dikonfirmasi oleh studi komparatif [6] yang menerapkan pendekatan NLP berbasis TF-IDF. Pembobotan TF-IDF terbukti krusial dalam menekan nilai bobot kata-kata yang bersifat netral atau terlalu sering muncul, sehingga model dapat mengarahkan fokus komputasi pada kata-kata kunci unik yang memiliki muatan emosional kuat.

3.2.4 Urgensi Intervensi Teknik Penyeimbangan Kelas SMOTE

Ketimpangan ekstrim pada distribusi data awal (di mana kelas Neutral mendominasi sebesar 64,79%) berhasil dimitigasi secara penuh melalui implementasi algoritma SMOTE pada data latih. Keberhasilan penyeimbangan ini dicerminkan oleh tingginya nilai metrik Recall makro yang menyentuh angka 96,48% pada model LR terbaik. Kondisi ini membuktikan bahwa model tidak mengalami bias mayoritas dan memiliki sensitivitas yang setara dalam mengenali dokumen ulasan kritis (Negative) maupun apresiatif (Positive). Dampak positif dari intervensi SMOTE ini sejalan dengan konseptualisasi ilmiah yang dijabarkan oleh [13], yang menyatakan bahwa tanpa penanganan class imbalance, performa nilai F1-Score makro model standar akan merosot tajam. Efektivitas SMOTE dalam mendongkrak performa algoritma berbasis probabilitas dan regresi juga diverifikasi secara empiris oleh riset [14] pada ulasan aplikasi BRImo.

3.2.5 Matriks Komparasi Akhir dengan Riset Terdahulu

Sebagai bentuk visualisasi komparatif menyeluruh, resume perbandingan performa akurasi antara model optimal yang dihasilkan pada penelitian ini dengan hasil penelitian terdahulu dirangkum secara ringkas melalui struktur data Tabel 6.

Tabel 6. Matriks Komparasi Parameter Akurasi dengan Riset Terdahulu

No	Penulis & Tahun	Algoritma Klasifikasi	Pendekatan Fitur / Solusi	Akurasi Model
1	Pasaribu dkk. (2025)[8]	Logistic Regression	TF-IDF + SMOTE	92,00%
2	Juliana & Raharja (2024)[1]	Multinomial Naïve Bayes	TF-IDF + Alpha Tuning	89,00%
3	Hermawan dkk. (2025)[14]	Naïve Bayes Classifier	Fitur Numerik + SMOTE	87,50%
4	Setyani dkk. (2026)[11]	Support Vector Machine	NLP Preprocessing Intensif	84,20%

Melalui tabel perbandingan pada Tabel 6, dapat disimpulkan secara sah bahwa alur metodologi yang mengintegrasikan prapemrosesan teks intensif, pembobotan TF-IDF, penyeimbangan data SMOTE, dan fungsi diskriminatif linear Logistic Regression sukses mencatatkan rekor performa akurasi terbaik (92,00%) yang melampaui berbagai model acuan pada literatur riset terdahulu. Model optimal ini memiliki kesiapan teknis yang matang untuk diintegrasikan ke dalam sistem dasbor pemantauan reputasi otomatis bagi pihak manajemen Gojek.

4. KESIMPULAN

Penelitian ini menyimpulkan bahwa integrasi arsitektur kecerdasan buatan berbasis Logistic Regression (LR) yang dipadukan dengan metode pembobotan Term Frequency-Inverse Document Frequency (TF-IDF) berskema Unigram serta teknik penyeimbangan data Synthetic Minority Over-sampling Technique (SMOTE) merupakan kombinasi paling optimal dan tangguh dalam mengklasifikasikan sentimen multi-kelas pada 14.565 dokumen ulasan aplikasi Gojek di Google Play Store. Melalui pengujian validasi silang, model diskriminatif linear tersebut sukses membuktikan keunggulannya dengan mencetak performa puncak mutlak pada seluruh parameter evaluasi sebesar 96,48%, yang mana capaian ini secara signifikan melampaui tingkat akurasi algoritma komparatif Multinomial Naïve Bayes (MNB) yang rentan mengalami bias akibat asumsi independensi fitur kata yang terlalu kaku. Keberhasilan metodologi ini secara keseluruhan didorong oleh efektivitas formula TF-IDF dalam mereduksi gangguan kata umum (noise) dari korpus teks bahasa Indonesia, efisiensi skema kata tunggal (Unigram) dalam merepresentasikan muatan emosional pengguna secara ringkas, serta peran krusial manipulasi ruang fitur SMOTE yang berhasil meratakan bias dominasi data netral awal sebesar 64,79% menjadi seimbang sehingga sensitivitas prediksi terhadap ulasan kritis (negatif) dan apresiasi (positif) meningkat tajam tanpa mengorbankan stabilitas akurasi sistem. Dengan demikian, luaran model digital yang dieksperimenkan dalam riset ini memiliki keandalan ilmiah yang sangat tinggi dan direkomendasikan secara konkret untuk diimplementasikan sebagai fondasi sistem pemantauan reputasi bisnis, evaluasi fitur berkala, serta analisis sentimen otomatis bagi pihak manajemen penyedia layanan aplikasi berbasis teks.

UCAPAN TERIMA KASIH

Pengarang menyampaikan apresiasi dan syukur yang mendalam terhadap seluruh rekan lain sudah mendukung serta memfasilitasi terlaksananya studi kami. Rasa syukur pastinya serta ucapan terimakasih khusus disampaikan terhadap para pengembang aplikasi Gojek dan seluruh pengguna yang telah memberikan penilaian di play store, serta menjadi masukan nilai pertama pada studi kami. Pengarang ingin menyampaikan apresiasi yang tulus kepada institusi pendidikan dan para dosen pembimbing sudah menyumbangkan arahan, saran berharga, dan semangat moril maupun materil saat pengerjaan

artikel ini. Diakhir kalimat ini, sang pengarang berharap studi ini dapat di gunakan sebagai kontribusi positif bagi pengembang ilmu data mining dan Natural Language Processing di Indonesia.

REFERENCES

- [1] N. K. A. Juliana and M. A. Raharja, "Analisis Sentimen pada Ulasan Aplikasi myIM3 Menggunakan Multinomial Naive Bayes dengan TF-IDF," *J. Nas. Teknol. Inf. dan Apl.*, vol. 2, no. 3, pp. 649–656, 2024, [Online]. Available: <https://ojs.unud.ac.id/index.php/jnatia/article/view/115818>
- [2] W. Purba, M. S. M. Panjaitan, E. Dawner, and M. D. H. Lumbantoruan, "Perbandingan Algoritma Naive Bayes Dan K-Nearest Neighbors Dalam Analisis Sentimen Ulasan Produk Tokopedia," *J. Tek. Inf. dan Komput.*, vol. 8, no. 1, p. 239, 2025, doi: 10.37600/tekinkom.v8i1.1985.
- [3] A. Maulana, Inayah Khasnaputri Afifah, Asghafi Mubarrak, Kiagus Rachmat Fauzan, Ardhan Dwintara, and B. P. Zen, "Comparison of Logistic Regression, Multinomialnb, Svm, and K-Nn Methods on Sentiment Analysis of Gojek App Reviews on the Google Play Store," *J. Tek. Inform.*, vol. 4, no. 6, pp. 1487–1494, 2023, doi: 10.52436/1.jutif.2023.4.6.863.
- [4] D. Budiman, Adhi Dwi Saputra, Revano Maliq Reynanda, and Angraini Puspita Sari, "Analisis Sentimen Aplikasi Gojek Pada Twitter Menggunakan Algoritma Naive Bayes," *J. Mhs. Tek. Inform.*, vol. 3, no. 2, pp. 107–116, 2024, doi: 10.35473/jamastika.v3i2.3375.
- [5] A. Aryanusa, "Evaluasi Kinerja Logistic Regression Dan Multinomial Naive Bayes Dalam Klasifikasi Sentimen Mobile Banking," *J. Inform. dan Tek. Elektro Terap.*, vol. 13, no. 3S1, pp. 463–474, 2025, [Online]. Available: <https://journal.eng.unila.ac.id/index.php/jitet/article/view/7687>
- [6] M. D. Bimantara and I. Zufria, "Text Mining Sentiment Analysis on Mobile Banking Application Reviews using TF-IDF Method with Natural Language Processing Approach," *JINAV J. Inf. Vis.*, vol. 5, no. 1, pp. 115–123, 2024, doi: 10.35877/454ri.jinav2772.
- [7] Z. B. Toyibah, Y. N. Putri, P. Puandini, Z. M. Widodo, and A. T. Ni'mah, "Perbandingan Kinerja Algoritma Multinomial Naive Bayes dan Logistic Regression pada Analisis Sentimen Movie Ratings IMDB," *J. Ilm. Edutic Pendidik. dan Inform.*, vol. 10, no. 2, pp. 181–189, 2024, doi: 10.21107/edutic.v10i2.28150.
- [8] R. Pasaribu, I. A. Gde, S. Putra, U. J. Raya, and K. Unud, "Analisis Sentimen Ulasan Aplikasi dengan Multinomial Naive Bayes, Logistic Regression, dan SVM," *Jnatia*, vol. 4, no. 1, pp. 63–72, 2025.
- [9] Indra Dikusuma and Yunus Widjaja, "Perbandingan Naive Bayes dan Logistic Regression untuk Analisis Sentimen Ulasan Produk Skincare di Tokopedia," *J. Inform. Dan Tekonologi Komput.*, vol. 5, no. 3, pp. 242–251, 2025, doi: 10.55606/jitek.v5i3.8420.
- [10] S. A. H. Bahtiar, C. K. Dewa, and A. Luthfi, "Comparison of Naive Bayes and Logistic Regression in Sentiment Analysis on Marketplace Reviews Using Rating-Based Labeling," *J. Inf. Syst. Informatics*, vol. 5, no. 3, pp. 915–927, 2023, doi: 10.51519/journalisi.v5i3.539.
- [11] T. Setyani, K. Sari, H. N. Heidy, and R. R. Suryono, "Sentiment Analysis of Jamsostek Mobile Application Reviews on Google Play Store Using Support Vector Machine and Naive Bayes Algorithms Analisis Sentimen Pengguna Aplikasi Jamsostek Mobile Berdasarkan Ulasan Google Play Store Menggunakan Algoritma Support," vol. 6, no. January, pp. 373–384, 2026.
- [12] H. Junianto, R. E. Saputro, B. A. Kusuma, and D. I. S. Saputra, "Comparison of Logistic Regression and Random Forest in Sentiment Analysis of Disdukcapil Application Reviews," *J. Tek. Inform.*, vol. 5, no. 6, pp. 1539–1547, 2024, doi: 10.52436/1.jutif.2024.5.6.1802.
- [13] C. Dometian and S. Jatmiko, "Analisis Sentimen Ulasan Google Play Store : Studi Komparatif Algoritma SVM , Naive Bayes , dan Logistic Regression," vol. 15, no. 3, pp. 610–620, 2025.
- [14] M. A. Hermawan, A. Faqih, and G. Dwilestari, "Implementasi Akurasi Model Naive Bayes Menggunakan Smote Dalam Analisis Sentimen Pengguna Aplikasi Brimo," *J. Inform. dan Tek. Elektro Terap.*, vol. 13, no. 1, 2025, doi: 10.23960/jitet.v13i1.5748.
- [15] H. Hidayatullah *et al.*, "Penerapan Naive Bayes Dengan Optimasi InfHidayatullah, H., Umaidah, Y., Informatika, P. S., Komputer, F. I., Karawang, U. S., Timur, T., Barat, J., Gain, I., & Baker, B. (2023). Penerapan Naive Bayes Dengan Optimasi Information Gain Dan. 7(3).ormation Gai," vol. 7, no. 3, 2023.