

MENGOPTIMALKAN KEPUASAN PENGGUNA: ANALISIS SENTIMEN REVIEW APLIKASI GRAB DI INDONESIA

Aldi Suryana¹, Ade Irma Purnamasari², Irfan Ali³

^{1,2} Teknik Informatika, STMIK IKMI Cirebon

³ Rekayasa Perangkat Lunak, STMIK IKMI Cirebon

Jl. Perjuangan No.10 B Majasem, Kec. Kesambi, Kota Cirebon Jawa Barat 45135, Indonesia

aldisur123@gmail.com

ABSTRAK

Sebagai perusahaan terkemuka di sektor transportasi dan teknologi di Asia Tenggara, Grab, didirikan pada 2012 di Malaysia sebagai platform pemesanan taksi melalui aplikasi. Pertumbuhan pesat perusahaan ini menciptakan kebutuhan untuk strategi analisis sentimen guna meningkatkan layanan dan kepuasan pengguna. Penelitian ini menggunakan pendekatan eksperimental dan kuantitatif dengan fokus pada ulasan aplikasi Grab di Google Play Store. Penggunaan teknik web scraping melalui Google Play Scraper mengumpulkan 1000 data ulasan. Hasil analisis sentimen menggunakan metode klasifikasi Naïve Bayes dimana metode ini merupakan metode pengklasifikasian probabilistik sederhana yang menghitung sekumpulan probabilistik dengan menjumlahkan frekuensi dan kombinasi nilai dari dataset yang diberikan menunjukkan tingkat akurasi sebesar 87%, dengan presisi 86%, dan recall 97%. Tahapan preprocessing, seperti case folding, tokenizing, stopwords, dan stemming, memastikan kualitas data sebelum klasifikasi. Evaluasi menggunakan Confusion Matrix di Google Colaboratory menunjukkan kinerja memuaskan, mengidentifikasi sentiment masyarakat pengguna aplikasi Grab. Tingkat akurasi 87% memberikan gambaran rinci tentang kemampuan model dalam menangani variasi sentimen, termasuk presisi, recall, dan F1-score. Penerapan metode klasifikasi Naïve Bayes dalam analisis sentimen aplikasi Grab terbukti efektif, memberikan wawasan berharga untuk pengembangan layanan yang lebih baik.

Kata kunci : *Grab, Google Play Store, Sentiment analysis, App reviews*

1. PENDAHULUAN

Sebagai perusahaan terkemuka di sektor transportasi dan teknologi di Asia Tenggara, Grab telah mengalami pertumbuhan yang signifikan sejak pendiriannya. Didirikan pada tahun 2012 di Malaysia, Grab awalnya dikenal sebagai platform pemesanan taksi melalui aplikasi. Namun, seiring berjalannya waktu, perusahaan ini terus mengembangkan model bisnisnya dengan menawarkan layanan yang beragam. Moda transportasi adalah jenis kendaraan yang dipilih individu untuk mencapai destinasi dari satu tempat ke tempat lain [1]. Salah satu moda transportasi daring yang menjadi perbincangan luas di kalangan masyarakat adalah ojek online [2]. Di Indonesia sendiri, Grab sebagai perusahaan ojek daring terkemuka menarik sejumlah besar pelanggan di negara ini [3].

Pemahaman mendalam terhadap permasalahan-permasalahan ini menjadi kunci dalam merancang strategi analisis sentimen yang efektif untuk meningkatkan layanan dan kepuasan pengguna aplikasi Grab. Berkembangnya jumlah pengguna smartphone dan kendaraan telah mengakibatkan peningkatan signifikan dalam jumlah pengguna aplikasi Grab. Kondisi ini menyebabkan adanya tuntutan yang lebih tinggi terhadap kualitas layanan aplikasi. Di Google Playstore, setiap pengguna memiliki kemampuan untuk memberikan penilaian dan ulasan terhadap aplikasi tertentu. Ulasan dan informasi mengenai suatu produk disimpan dalam bentuk teks, sehingga text mining menjadi solusi yang relevan dalam ekstraksi informasi berbasis teks. Salah

satu aspek dari proses ekstraksi informasi teks adalah Analisis Sentimen, yang menitikberatkan pada evaluasi sentimen atau pendapat dalam teks tersebut [4]. Menghadapi permasalahan yang muncul, diperlukan tindakan solutif dalam bentuk analisis terhadap saran dan keluhan yang dikirimkan kepada perusahaan ojek daring Grab [2]. Semakin banyak orang yang mengungkapkan pandangan mereka terhadap layanan Grab di ruang publik, maka semakin meningkat pula tingkat sentimen yang dapat mencakup rasa puas dan ketidakpuasan terhadap pelayanan yang disediakan. Analisis sentimen merupakan proses yang menguraikan pandangan, penilaian, evaluasi, sikap, dan emosi seseorang terhadap suatu topik, layanan, produk, individu, organisasi, atau kegiatan tertentu. [1] Analisis sentimen adalah proses klasifikasi teks yang bertujuan untuk mengkategorikan teks atau dokumen yang mengandung pendapat menjadi pendapat yang bersifat positif, negatif, atau netral. [3] Karena itu, diperlukan analisis terhadap pendapat-pendapat tersebut dalam penelitian ini agar dapat digunakan sebagai indikator penilaian dari perspektif pelanggan terkait kualitas layanan jasa transportasi online. Analisis sentimen dikenal sebagai aspek yang menjelaskan pandangan, penilaian, evaluasi, sikap, dan emosi individu terhadap suatu topik, layanan, produk, individu, organisasi, atau kegiatan tertentu [1].

Perubahan sentimen dari waktu ke waktu juga menjadi perhatian, karena tren yang berubah-ubah dapat menciptakan kesulitan dalam memahami perkembangan persepsi pengguna terhadap aplikasi. Isu teknis dan keluhan pengguna, bersama dengan

kemungkinan adanya sentimen palsu atau tidak relevan, juga dapat mempengaruhi analisis sentimen keseluruhan. Penelitian ini mengajukan penggunaan metode machine learning sebagai titik pembandingan terhadap metode yang diusulkan. Metode machine learning yang digunakan melibatkan Multinomial Naïve Bayes [3]. penelitian pertama dilakukan Algoritma Klasifikasi Naïve Bayes dianggap sebagai metode yang memiliki potensi yang baik untuk melakukan klasifikasi data terstruktur dibandingkan dengan metode pembagian terstruktur lainnya, terutama dalam hal akurasi dan efisiensi komputasi. [1] Meskipun Naive Bayes memiliki kemudahan penggunaan, tingkat akurasi yang tinggi dapat dicapai asalkan dataset telah dibersihkan secara optimal dengan memastikan keberadaan kelas negatif dan positif yang baik [5]. NBC (Naive Bayes Classifier) sendiri merupakan metode pengklasifikasian probabilistik sederhana yang menghitung sekumpulan probabilistik dengan menjumlahkan frekuensi dan kombinasi nilai dari dataset yang diberikan, digunakan sebagai pendekatan untuk menentukan sentimen aplikasi di Google Play melalui evaluasi komentar yang disampaikan oleh para pengguna [6].

Dalam rangka menghadapi pertumbuhan yang signifikan sejak pendiriannya, Grab, sebagai perusahaan terkemuka di sektor transportasi dan teknologi di Asia Tenggara, mengarahkan fokus penelitian ini pada penerapan metode Naive Bayes pada ulasan pelanggan. Tujuan utama dari pendekatan ini adalah untuk melakukan analisis sentimen pelanggan terhadap layanan Grab yang tercermin dalam ulasan di Google Play Store. Dengan mengidentifikasi aspek-aspek kunci yang memengaruhi sentimen keseluruhan, perusahaan berharap dapat memahami lebih baik area yang perlu diperbaiki atau dikembangkan. Selain itu, penelitian ini bertujuan untuk memberikan pemahaman yang lebih mendalam tentang tingkat kepuasan pengguna terhadap berbagai aspek aplikasi. Hasil analisis sentimen diharapkan dapat memberikan dasar yang kuat untuk perbaikan layanan, membantu perusahaan mengambil keputusan strategis yang lebih tepat, dan meningkatkan akurasi analisis sentimen melalui optimalisasi metode Naive Bayes. Dengan demikian, pendekatan ini diharapkan dapat memberikan kontribusi positif terhadap pengembangan Grab dalam menjaga dan meningkatkan kepercayaan pengguna di pasar transportasi daring di Asia Tenggara.

2. TINJAUAN PUSTAKA

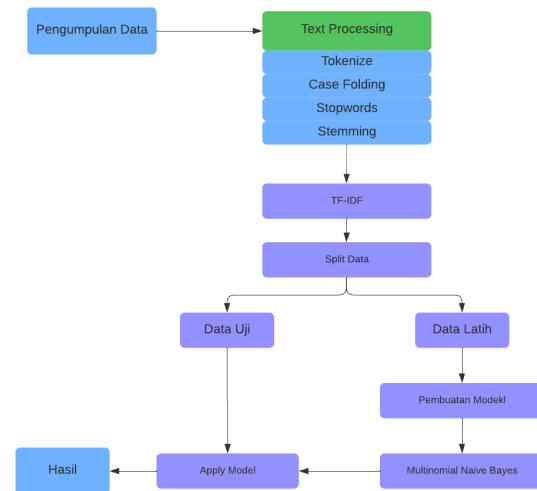
Dapat di simpulkan dari tabel ringkasan literatur review bahwa algoritma Naïve Bayes dan support vektor machine berguna untuk klasifikasi sentiment. Selain itu beberapa jurnal menggunakan tambahan algoritma seperti naïve bayes classifier dan support vector machine pun berperan sebagai pendukung klasifikasi yang dilakukan.

2.1. Penelitian Terdahulu

Penelitian ini menyoroti tantangan dalam sistem pelayanan kesehatan Indonesia, khususnya terkait jadwal pemeriksaan dan antrian pasien. Aplikasi konsultasi dokter, seperti Klik Dokter, Halodoc, dan Alodokter, menjadi solusi alternatif, dan persepsi masyarakat diungkap melalui analisis sentimen Twitter. Hasilnya menunjukkan Klik Dokter mendapat komentar negatif terendah (98,57%), diikuti oleh Halodoc (82,86%) dan Alodokter (62,86%). Studi ini memberikan wawasan kritis terhadap kekurangan sistem dan kualitas layanan. aplikasi konsultasi dokter, memberikan dasar untuk perbaikan lebih lanjut di sektor kesehatan [6]. Penelitian membandingkan kinerja metode klasifikasi Support Vector Machine (SVM) dan *Naïve Bayes* untuk menganalisis ulasan aplikasi berbahasa Indonesia di Google Play Store. SVM menunjukkan akurasi lebih tinggi (81,46%) dibandingkan *Naïve Bayes* (75,41%) melalui pengujian 3-folds cross-validation, menunjukkan keefektifan SVM dalam mengklasifikasikan sentimen ulasan aplikasi. Ini berpotensi meningkatkan metode analisis sentimen di masa depan [7]. Penelitian ini menggunakan metode *Naïve Bayes Classifier* (NBC) untuk mengklasifikasikan sentimen pengguna Twitter terhadap layanan BMKG Nasional. Dengan pertumbuhan pengguna Twitter yang terus meningkat, platform ini menjadi saluran utama bagi masyarakat untuk menyampaikan kritik dan saran. Penggunaan algoritma Naïve Bayes dalam klasifikasi, dilakukan dengan Python 3.74, menunjukkan hasil uji akurasi sebesar 69.97%. Meskipun memberikan kemudahan dalam melihat opini pengguna, penelitian ini mungkin menghadapi beberapa permasalahan terkait kompleksitas analisis sentimen pada kalimat komentar yang memerlukan pendekatan lebih lanjut.[8] Penelitian ini menganalisis sentimen ulasan Halodoc di Google Play Store, sebuah layanan kesehatan digital populer. Proses analisis melibatkan preprocessing data, pembagian kelas positif dan negatif dengan validasi silang 10-fold, dan menggunakan algoritma Naïve Bayes Classifier. Hasilnya menunjukkan akurasi 81.68% dan AUC 0.756, tergolong baik. Namun, permasalahan mungkin muncul pada analisis sentimen bahasa kompleks atau review dengan konten ambigu [9]. Penelitian analisis sentimen ulasan Halodoc di Google Play Store menggunakan Naïve Bayes Classifier. Proses melibatkan tahap preprocessing data, pembagian kelas positif dan negatif dengan validasi silang 10-fold. Hasilnya menunjukkan akurasi 81.68% dan AUC 0.756, tergolong baik. Namun, mungkin muncul permasalahan pada analisis sentimen bahasa kompleks atau review ambigu [10]. Penelitian analisis sentimen opini masyarakat terhadap aplikasi Gojek menggunakan metode text mining. Kepuasan pengguna diukur melalui analisis sentimen terhadap ulasan pengguna. Metode melibatkan pembobotan kata untuk menentukan sentimen positif atau negatif dengan algoritma Naive Bayes. Hasilnya, akurasi

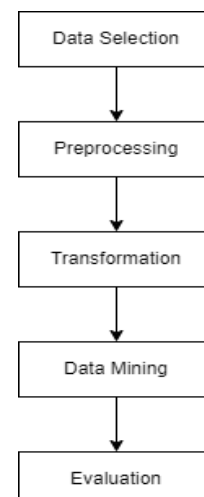
mencapai 68% dari 800 data latih dan uji. Meskipun tinggi, penelitian mungkin dihadapi kompleksitas analisis sentimen pada bahasa yang kompleks atau review ambigu [11]. Penelitian analisis sentimen ulasan pengguna aplikasi Sambara, inovasi pembayaran pajak kendaraan di Jawa Barat, menggunakan algoritma Naïve Bayes dan seleksi fitur Chi Square. Data dari Play Store bahasa Indonesia periode April-Desember 2021 dibagi menjadi data latih dan uji (80:20). Hasilnya menunjukkan akurasi tertinggi 94.4%, specificity 93.5%, dan recall 95.3%. Potensial masalah melibatkan kompleksitas analisis sentimen dalam bahasa Indonesia dan variasi interpretasi subjektif sentimen dalam ulasan pengguna [12]. Penelitian analisis sentimen masyarakat terhadap kinerja DPR melalui Twitter menggunakan Naive Bayes Classifier pada 1546 tweet. Proses mencakup pengumpulan data, preprocessing, dan klasifikasi. Hasilnya menunjukkan 75% sentimen positif, 79% sentimen netral, dan 82% sentimen negatif, dengan akurasi keseluruhan 80%. Potensial masalah melibatkan interpretasi sentimen subjektif, variasi ekspresi bahasa di media sosial, atau bias dalam pemilihan data [13]. Penelitian ini membandingkan analisis sentimen ulasan aplikasi Zoom Cloud Meetings di Google Play Store menggunakan metode Naïve Bayes dan Support Vector Machine (SVM). Data diambil melalui web scraping, diproses dengan RapidMiner, dan dimodelkan dengan Naïve Bayes dan SVM menggunakan 10 fold cross-validation. Evaluasi menunjukkan bahwa model SVM memiliki akurasi lebih baik (81.22%) dan AUC (0.886) daripada Naïve Bayes (74.37% dan 0.659). Studi sebelumnya juga menunjukkan keunggulan SVM dalam analisis sentimen. Kesimpulannya, SVM lebih unggul dalam menganalisis sentimen ulasan pengguna Zoom Cloud Meetings di Google Play Store [14]. Penelitian ini menerapkan algoritma Naïve Bayes untuk menganalisis sentimen ulasan aplikasi MyPertamina di Google Play Store. Tujuannya adalah mengklasifikasikan sentimen ulasan sebagai positif atau negatif dan mengevaluasi akurasi, presisi, dan recall. Data diambil melalui web scraping, dilabeli, dibersihkan, dan diklasifikasikan dengan Naïve Bayes. Hasil penelitian menunjukkan akurasi 77.42%, presisi 49.98%, dan recall 76.87%. Metodologi penelitian mencakup web scraping, pelabelan dataset, preprocessing, implementasi Naïve Bayes, dan evaluasi [15]. Artikel ini membandingkan ulasan pengguna Google Meet dan Zoom Cloud Meeting dengan algoritma Naïve Bayes. Penelitian melibatkan analisis sentimen, data mining, dan penggunaan teknik seleksi fitur serta optimasi untuk meningkatkan akurasi algoritma. Metodologi penelitian mencakup selection, pre-processing, transformation, data mining, dan interpretation/evaluation. Hasil menunjukkan algoritma K-Nearest Neighbor dengan feature optimasi (SMOTE dan PSO) menghasilkan akurasi dan AUC lebih tinggi daripada Naïve Bayes, khususnya setelah penambahan SMOTE dan PSO pada

2.1. Gambar Metode penelitian Teknik analisis data dan Text processing



Gambar 1. Tahapan metode penelitian.

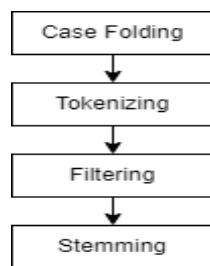
Dapat dilihat Penelitian ini memanfaatkan pendekatan eksperimental untuk mengevaluasi dampak variabel independen terhadap variabel dependen. Pendekatan kuantitatif diadopsi untuk menganalisis fenomena yang sedang diselidiki. Sumber data kuantitatif berasal dari ulasan aplikasi Grab di Google Play Store. Proses pengumpulan data ulasan dilakukan menggunakan aplikasi Google Play Scraper melalui teknik web scraping. Pada tahap pembobotan kata, diterapkan metode Text Mining untuk mendapatkan nilai bobot kata-kata dalam data yang digunakan. Metode Term Frequency - Inverse Document Frequency (TF-IDF) digunakan untuk menghitung nilai bobot tersebut. Setelah proses pelabelan, sentimen diklasifikasikan menggunakan metode klasifikasi teks. Algoritma Multinomial Naïve Bayes digunakan untuk mengklasifikasikan sentimen menjadi positif, negatif, atau netral. Penggunaan label otomatis dengan menyesuaikan skor pada ulasan Google Play Store dapat mempercepat proses pelabelan.



Gambar 2. teknik analisis data menggunakan

2.2. Knowledge Discovery in Databases (KDD).

Dalam penelitian ini, metode analisis data yang digunakan adalah *Knowledge Discovery in Database* (KDD). KDD adalah proses analisis data terstruktur untuk memperoleh informasi baru dan menemukan pola dari data besar dan kompleks. Data mining merupakan inti dari proses KDD dengan menggunakan algoritma tertentu untuk mengeksplorasi data, membangun model, dan menemukan pola yang belum diketahui. Proses text mining dengan KDD memiliki tahapan yang sama dengan tahapan data mining. Seperti 1. *Data Selection*, 2. *Preprocessing*, 3. *Transformation*, 4. *Data Mining*, 5. *Evaluation*.



Gambar 3. Text Processing

Preprocessing text merupakan tahapan selanjutnya setelah tahap scraping, pada tahap ini data akan di samakan seluruh bentuk dan format nya agar pada tahap selanjutnya data dapat diproses dan diolah. Tahap preprocessing text terdiri dari 4 tahapan yaitu, *case folding*, *tokenizing*, *filtering*, dan *stemming*.

3. METODE PENELITIAN

Dapat dilihat Penelitian ini memanfaatkan pendekatan eksperimental untuk mengevaluasi dampak variabel independen terhadap variabel dependen. Pendekatan kuantitatif diadopsi untuk menganalisis fenomena yang sedang diselidiki. Sumber data kuantitatif berasal dari ulasan aplikasi Grab di Google Play Store. Proses pengumpulan data ulasan dilakukan menggunakan aplikasi Google Play Scraper melalui teknik web scraping. Pada tahap pembobotan kata, diterapkan metode Text Mining untuk mendapatkan nilai bobot kata-kata dalam data yang digunakan. Metode Term Frequency - Inverse Document Frequency (TF-IDF) digunakan untuk menghitung nilai bobot tersebut. Setelah proses pelabelan, sentimen diklasifikasikan menggunakan metode klasifikasi teks. Algoritma Multinomial Naïve Bayes digunakan untuk mengklasifikasikan sentimen menjadi positif, negatif, atau netral. Penggunaan label otomatis dengan menyesuaikan skor pada ulasan Google Play Store dapat mempercepat proses pelabelan. Gambar 1 menunjukkan urutan metodologi penelitian yang telah diterapkan oleh penelitian.

3.1. Data Selection

Data yang digunakan dalam penelitian ini diperoleh dari ulasan pengguna aplikasi Grab di Google Play Store. Pengumpulan data dilakukan dengan menggunakan alat Google Scraper melalui

platform Google Colab. Sebanyak 1000 data diambil untuk analisis.

Populasi dan Sample data sentimen pengguna Grab yang diambil dari ulasan Play Store mencakup semua ulasan yang diberikan oleh pengguna terhadap aplikasi tersebut. Sebagai ilustrasi, sebagian ulasan dalam kumpulan ini dapat berisi pengalaman positif seperti "Aplikasi sangat membantu, pengemudi ramah," atau pengalaman negatif seperti "Pelayanan pelanggan kurang responsif." Dengan pertumbuhan pengguna yang terus berlanjut, populasi ulasan ini terus berkembang seiring berjalannya waktu. Untuk melakukan analisis sentimen dengan efisien, kita sering menggunakan sampel yang merupakan representasi kecil dari populasi keseluruhan. Sebagai contoh, sampel dapat terdiri dari beberapa ulasan yang mencakup variasi sentimen, termasuk positif, negatif, dan netral. Melalui analisis sampel ini, kita dapat memahami pandangan umum pengguna terhadap aplikasi Grab. Sampel data set yang di scrapper memiliki atribut awal seperti. *revieweId*, *userName*, *userImagecontent*, *score*, *thumbsUpCount*, *reviewCreatedVersion*, *at*, *replyContent*, *appVersion*.

3.2. Text Processing

Text Preprocessing merupakan tahapan selanjutnya setelah tahap scraping, pada tahap ini data akan di samakan seluruh bentuk dan format nya agar pada tahap selanjutnya data dapat diproses dan diolah. Tahap preprocessing text terdiri dari 4 tahapan yaitu, *case folding*, *tokenizing*, *filtering*, dan *stemming*[16].

3.3. Tokenizing

Tokenizing adalah proses pemotongan teks menjadi bagian-bagian yang lebih kecil disebut token, proses ini dilakukan setelah melakukan tahapan *case folding*.

3.4. Case Folding

Case folding merupakan proses perubahan pada seluruh huruf besar menjadi huruf kecil. Pada proses ini contohnya karakter 'A'-'Z' yang terdapat pada data diubah kedalam karakter 'a'-'z'. tabel 3 menunjukan bahwa proses *case folding* dimana pada data sebelumnya terdapat huruf kapital setelah proses *case folding* berubah menjadi huruf kecil.

3.5. Stopwords

Stopwords adalah kata-kata umum yang secara sering muncul dalam suatu bahasa, namun umumnya memberikan sedikit nilai informasi dalam analisis teks karena keberadaannya yang luas dan merata di banyak dokumen. Dalam pemrosesan teks atau analisis *Natural Language Processing* (NLP), kata-kata ini sering kali diabaikan atau dihapus karena kontribusi mereka yang jarang memberikan makna yang signifikan atau berarti terhadap pemahaman isi teks. Proses ini membantu menyederhanakan teks dan memusatkan perhatian pada kata-kata yang lebih bermakna dan informatif.

3.6. Stemming

Stemming merupakan langkah dalam pemrosesan teks yang melibatkan penghilangan afiks (baik akhiran maupun awalan) dari sebuah kata dengan maksud menghasilkan bentuk dasarnya atau kata dasar. Fungsinya adalah untuk menyederhanakan kata-kata ke bentuk dasarnya sehingga kata-kata yang memiliki akar kata yang sama dapat lebih mudah diidentifikasi dalam proses analisis teks atau pemrosesan bahasa alami. *Stemming* sendiri merupakan proses pemetaan dan penguraian bentuk dari suatu kata untuk melakukan suatu *stemming* proses ini menggunakan *library Sastrawi* untuk melakukan proses yang dilakukan.

3.7. Pembobotan TF-IDF

Proses ini melibatkan manipulasi data agar dapat dipersiapkan sesuai dengan kebutuhan analisis tertentu yang akan dilakukan. TF digunakan untuk mengukur sejauh mana sebuah kata muncul dalam suatu dokumen dengan menghitung rasio kemunculannya dibagi dengan jumlah total kata dalam dokumen tersebut. Sementara itu, IDF menilai tingkat keunikan atau informativitas sebuah kata dengan membandingkan frekuensinya di seluruh koleksi dokumen. Kata-kata yang muncul lebih jarang di seluruh koleksi dapat memiliki nilai IDF yang lebih tinggi. TF-IDF, hasil perkalian dari TF dan IDF, memberikan bobot pada setiap kata dalam dokumen. Dimana melakukan TF-IDF menggunakan *tfidfTransform* dan *TF IDF Vectorizer* dimana Dengan *Tfidftransformer* penulis akan menghitung jumlah kata secara sistematis menggunakan *CountVectorizer* dan kemudian menghitung nilai *Inverse Document Frekuensi (IDF)* dan baru kemudian menghitung skor *Tf-idf*. disini penulis menggunakan *TF IDF Vectorizer* dimana *tf idf* ini melakukan ketiga langkah sekaligus. Di balik terpal, ia menghitung jumlah kata, nilai IDF, dan skor *Tf-idf*, semuanya menggunakan kumpulan data yang sama.

3.8. Split Data

Teknik split data memainkan peran krusial dalam manajemen dataset untuk keperluan pelatihan dan pengujian model. Dataset umumnya dipecah menjadi dua bagian utama: dataset pelatihan dan dataset pengujian. Dataset pelatihan berfungsi sebagai materi pembelajaran bagi model, memungkinkan algoritma untuk memahami pola dan keterhubungan dalam data. Sebaliknya, dataset pengujian digunakan untuk menguji sejauh mana model mampu melakukan prediksi pada data yang belum pernah dilihat

3.9. Apply model

Proses data mining, merujuk pada proses ekstraksi pola atau pengetahuan yang berharga dari sejumlah besar data dengan memanfaatkan berbagai metode, algoritma, dan teknik analisis. Fokus utama dari kegiatan data mining adalah untuk mengungkapkan hubungan atau pola yang masih

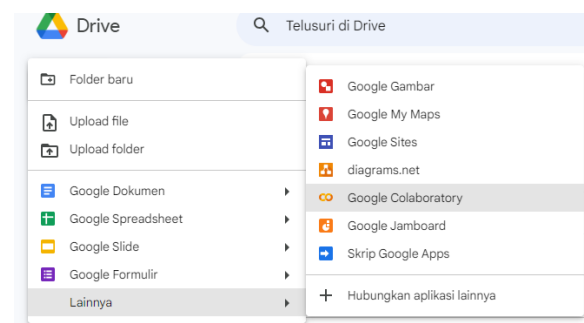
tersembunyi dalam data, yang dapat memberikan wawasan dan informasi berharga. Tujuan akhirnya adalah untuk mendapatkan pemahaman yang lebih dalam dan memanfaatkan potensi nilai dari data yang ada. Dalam penelitian penulis menggunakan metode *naïve bayes* untuk menentukan *accuracy* sentimen ulasan positif dan negatif pada pengguna aplikasi grab.

Setelah penyelesaian proses klasifikasi, langkah berikutnya adalah tahap evaluasi. Tahap ini merupakan fase akhir dalam penelitian ini, dilakukan untuk menilai sejauh mana keefektifan klasifikasi menggunakan algoritma *naïve bayes*. Evaluasi dilakukan menghitung hasil akhir nilai akurasi *multinomialNB*. Pada hasil analisis *multinomialNB* didapatkan nilai akurasi sebesar 87%.

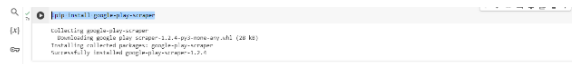
Dan evaluasi *Confusion Matrix* di mana data dibagi menjadi dua bagian, yakni data pengujian (*testing*) dan data pelatihan (*training*). Jumlah data *testing* sebanyak 183 data, data tersebut kemudian di proses menjadi *Confusion matrix* dapat dilihat melalui tabel di bawah ini. setelah didapatkan nilai *True negative (TN)* sebanyak 177 data klasifikasi yang dinyatakan benar kelas positif. Lalu terdapat 6 data klasifikasi dimana kelas positif yang benar tetapi diklasifikasikan sebagai *false negative (FP)*. untuk kelas negative, terdapat 24 data klasifikasi yang seharusnya masuk ke dalam kelas negative tetapi termasuk kedalam kelas *false negative (FN)* dan yang terakhir terdapat 67 klasifikasi data yang sudah benar nilai kelasnya sebagai *true negative (TN)*.

4. HASIL DAN PEMBAHASAN

Hasil penelitian merujuk pada kesimpulan atau temuan yang diperoleh dari suatu studi atau penyelidikan ilmiah yang dilakukan oleh peneliti. Hal ini mencakup data, informasi, dan analisis yang dihasilkan dari proses penelitian yang dilakukan secara sistematis dan metodologis berdasarkan fakta melalui usaha pikiran dalam mengelola dan menganalisis objek penelitian untuk memecahkan suatu permasalahan atau menguji suatu hipotesis sehingga terbentuk prinsip-prinsip umum. Pada Langkah ini penulis melakukan pengumpulan data dengan menggunakan bahasa pemrograman Python melalui alat Google Colaboratory dan paket *google-play-scraper*. metode analisis data yang digunakan adalah *Knowledge Discovery in Database (KDD)*.

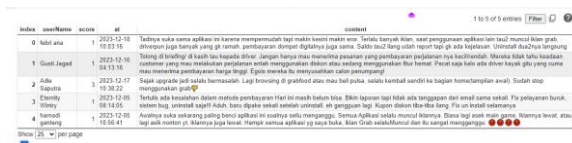


Gambar 4. Google colaboratory



Gambar 5. google-play-scraper

Penulis melakukan scraping data melalui website google play store, dengan cara masuk ke menu search ketikan aplikasi yang ingin di cari data ulasan komentarnya yang ingin dijadikan sampel penelitian. Disini penulis menggunakan data ulasan aplikasi grab.di dapat 1000 data ulasan komentar, Selanjutnya karena data yang didapatkan banyak sekali memiliki kolom atribut, kita filter kembali sehingga di dapatkan kolom username, at dan content.



Gambar 6. data mentah yang di ambil

Pada Langkah selanjutnya penulis melakukan pelabelan terlebih dahulu dengan memanfaatkan score sebagai acuan label. Dimana score kurang dari 3 akan diberi ulasan negative sedangkan score 4 dan score 5 diberi ulasan positive Setelah itu, dilakukan pemberian nilai terhadap label positif sebesar 1 dan label negatif sebesar 0. Selanjutnya kita cari jumlah baris data yang bersifat null dengan menggunakan fitur `my_df.isnull().sum()`. terdapat kolom label sebanyak 87 yang memiliki nilai kosong di karenakan score 3 adalah ambigu penulis tidak memberi label positif atau negative. Lalu di lakukan penghapusan kolom yang kosong dan save ke dalam bentuk csv.

```
[ ] #mencari jumlah baris data yang bernilai null
#terdapat kolom label memiliki nilai kosong
my_df.isnull().sum()

content      0
score        0
Label        87
dtype: int64
```

Gambar 7. penghapusan nilai yang kosong

Setelah tahap scraping, selanjutnya text processing, pada tahap ini data akan di samakan seluruh bentuk dan format nya agar pada tahap selanjutnya data dapat diproses dan diolah. Tahap preprocessing text terdiri dari 4 tahapan yaitu, *case folding*, *tokenizing*, *filtering*, dan *stemming*. Lalu masuk tahap case folding dimana data memasuki tahap perubahan pada seluruh huruf kapital menjadi huruf kecil. Selanjutnya pada taha *tokenizing* yaitu tahap pemotongan teks menjadi bagian-bagian kecil yang di sebut token. Tahap Filtering biasa di sebut dengan *stopwords removal* Stopwords adalah kata-kata umum yang secara sering muncul dalam suatu bahasa, namun umumnya memberikan sedikit nilai informasi dalam analisis teks karena keberadaannya yang luas dan merata di banyak dokumen. Dalam pemrosesan teks atau analisis Natural Language Processing (NLP),

kata-kata ini sering kali diabaikan atau dihapus karena kontribusi mereka yang jarang memberikan makna yang signifikan atau berarti terhadap pemahaman isi teks. Proses ini membantu menyederhanakan teks dan memusatkan perhatian pada kata-kata yang lebih bermakna dan informatif. Proses akhir dari *text processing* adalah *stemming*, dimana proses ini merupakan langkah dalam pemrosesan teks yang melibatkan penghilangan afiks (baik akhiran maupun awalan) dari sebuah kata dengan maksud menghasilkan bentuk dasarnya atau kata dasar. Fungsinya adalah untuk menyederhanakan kata-kata ke bentuk dasarnya sehingga kata-kata yang memiliki akar kata yang sama dapat lebih mudah diidentifikasi dalam proses analisis teks atau pemrosesan bahasa alami. *Stemming* sendiri merupakan proses pemeataan dan penguraian dalam bentuk dari suatu kata untuk melakukan suatu *stemming* proses ini menggunakan *library Sastrawi*.

Split Data, teknik split data memainkan peran krusial dalam manajemen dataset untuk keperluan pelatihan dan pengujian model. Dataset umumnya dipecah menjadi dua bagian utama: dataset pelatihan dan dataset pengujian. Dataset pelatihan berfungsi sebagai materi pembelajaran bagi model, memungkinkan algoritma untuk memahami pola dan keterhubungan dalam data. Sebaliknya, dataset pengujian digunakan untuk menguji sejauh mana model mampu melakukan prediksi pada data yang belum pernah dilihat sebelumnya. Ini membantu dalam mengevaluasi kinerja model serta mengukur kemampuannya untuk menggeneralisasi. Ada berbagai teknik pembagian data, termasuk pembagian acak, pembagian berdasarkan proporsi, dan validasi silang, yang dapat diadopsi sesuai dengan karakteristik data dan tujuan analisis. Pemilihan metode pembagian data yang tepat menjadi faktor kunci dalam memastikan hasil evaluasi yang objektif dan membangun model pembelajaran mesin yang dapat diandalkan. Penulis sendiri membagi data menjadi data training dan testing dengan $\text{test_size} = 0.2$ dan $\text{random_state} = 0$.

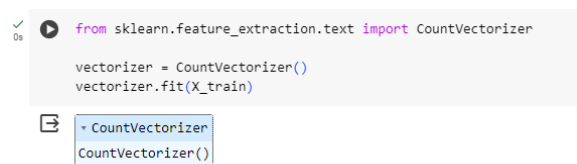
```
print(X_train.shape)
print(y_train.shape)
print(X_test.shape)
print(y_test.shape)
```

```
(637,)
(637,)
(274,)
(274,)
```

Gambar 8. Spilt Data

Pembobotan TF-IDF, Proses ini melibatkan manipulasi data agar dapat dipersiapkan sesuai dengan kebutuhan analisis tertentu yang akan dilakukan. TF digunakan untuk mengukur sejauh mana sebuah kata muncul dalam suatu dokumen dengan menghitung rasio kemunculannya dibagi dengan jumlah total kata dalam dokumen tersebut. Sementara itu, IDF menilai tingkat keunikan atau informativitas sebuah kata dengan membandingkan frekuensinya di seluruh

koleksi dokumen. Kata-kata yang muncul lebih jarang di seluruh koleksi dapat memiliki nilai IDF yang lebih tinggi. TF-IDF, hasil perkalian dari TF dan IDF, memberikan bobot pada setiap kata dalam dokumen. Dimana melakukan TF-IDF menggunakan `TfidfTransformer` dan `TFIDF Vectorizer`. Dimana Dengan `TfidfTransformer` penulis akan menghitung jumlah kata secara sistematis menggunakan `CountVectorizer` dan kemudian menghitung nilai Inverse Document Frekuensi (IDF) dan baru kemudian menghitung skor `Tf-idf`. disini penulis menggunakan `TFIDF Vectorizer`, dimana `tfidf` ini melakukan ketiga langkah sekaligus. Di balik terpal, ia menghitung jumlah kata, nilai IDF, dan skor `Tf-idf`, semuanya menggunakan kumpulan data yang sama.

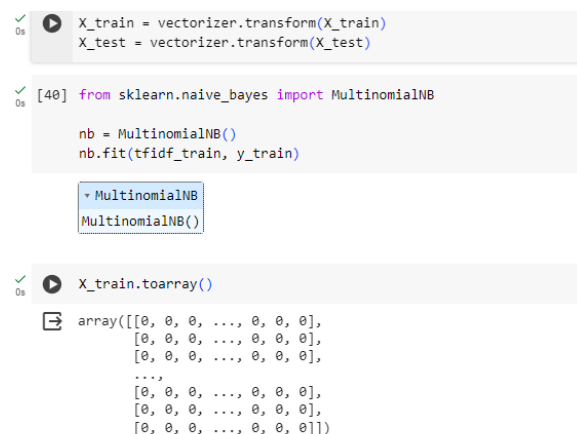


```
from sklearn.feature_extraction.text import CountVectorizer

vectorizer = CountVectorizer()
vectorizer.fit(X_train)
```

Gambar 9. TF-IDF

Klasifikasi *Naïve Bayes*, merupakan Proses data mining, merujuk pada proses ekstraksi pola atau pengetahuan yang berharga dari sejumlah besar data dengan memanfaatkan berbagai metode, algoritma, dan teknik analisis. Fokus utama dari kegiatan data mining adalah untuk mengungkapkan hubungan atau pola yang masih tersembunyi dalam data, yang dapat memberikan wawasan dan informasi berharga. Tujuan akhirnya adalah untuk mendapatkan pemahaman yang lebih dalam dan memanfaatkan potensi nilai dari data yang ada. Dalam penelitian penulis menggunakan metode *naïve bayes* untuk menentukan accuracy sentimen ulasan positif dan negatif pada pengguna aplikasi grab.



```
X_train = vectorizer.transform(X_train)
X_test = vectorizer.transform(X_test)

[40] from sklearn.naive_bayes import MultinomialNB

nb = MultinomialNB()
nb.fit(tfidf_train, y_train)

X_train.toarray()
```

Gambar 10. Klasifikasi Naïve Bayes

Setelah penyelesaian proses klasifikasi, langkah berikutnya adalah tahap Apply model. Tahap ini merupakan fase akhir dalam penelitian ini, dilakukan untuk menilai sejauh mana keefektifan klasifikasi

menggunakan algoritma *naïve bayes*. Apply model dilakukan menghitung hasil akhir nilai akurasi `multinomialNB`. Pada hasil analisis `multinomialNB` didapatkan nilai akurasi sebesar 89%.

	precision	recall	f1-score	support
Negatif	0.88	0.97	0.92	183
Positif	0.92	0.74	0.82	91
accuracy			0.89	274
macro avg	0.90	0.85	0.87	274
weighted avg	0.89	0.89	0.89	274

Gambar 11. Hasil Nilai Akurasi

Dan evaluasi *Confusion Matrix* di mana data dibagi menjadi dua bagian, yakni data pengujian (testing) dan data pelatihan (training). Jumlah data testing sebanyak 183 data, data tersebut kemudian di proses menjadi *Confusion matrix* dapat dilihat melalui tabel di bawah ini. Setelah didapatkan nilai *True negative (TN)* sebanyak 177 data klasifikasi yang dinyatakan benar kelas positif. Lalu terdapat 6 data klasifikasi dimana kelas positif yang benar tetapi diklasifikasikan sebagai *false negative (FP)*. Untuk kelas negative, terdapat 24 data klasifikasi yang seharusnya masuk ke dalam kelas negative tetapi termasuk kedalam kelas *false negative (FN)* dan yang terakhir terdapat 67 klasifikasi data yang sudah benar nilai kelasnya sebagai *true negative (TN)*.

Tabel 1. Confusion Matriks

		Kelas sebenarnya	
		positif	Negatif
Kelas prediksi	positif	177	6
	negatif	24	67

Pembahasan dari hasil penelitian yang sudah dilakukan yakni mengenai analisis sentimen pengguna aplikasi grab pada ulasan Google play store yang diambil melalui tahap scraping sebanyak 1000 data pada ulasan google play store. Dari banyaknya data yang di ambil dan telah melalui tahapan-tahapan preprocessing seperti *case folding*, *tokenizing*, *stopwords* dan *stemming* melalui tahap klasifikasi menggunakan algoritma *Naïve Bayes* serta evaluasi *Confusion Matrix* menggunakan *Google Colaboratory*. Dapat disimpulkan dari hasil yang di dapat bahwa, penerapan metode klasifikasi *naïve bayes* dalam klasifikasi ulasan pengguna aplikasi grab tergolong cukup baik dengan hasil 89,05% dengan tingkat akurasi (*Accuracy*) sebesar 89% presisi (*Precision*) sebesar 88%, dan sebesar 97% untuk tingkat keberhasilan *recall* nya. Hasil ini dapat dinilai bahwa algoritma klasifikasi *naïve bayes* cukup baik dalam pemrosesan data ulasan pada google playstore, dikarenakan hasil persentasenya mencapai 89% berdasarkan nilai tersebut ulasan masyarakat mengenai aplikasi grab tergolong cukup baik.

Hasil evaluasi dari pemrosesan menggunakan metode *naïve bayes* memiliki kinerja yang cukup memuaskan dalam mengidentifikasi sentimen

masyarakat pengguna aplikasi grab. Dengan tingkat akurasi yang di dapat sebesar 87% serta memberikan gambaran lebih terperinci mengenai kemampuan model dalam menangani variasi sentimen melibatkan presisi, recall, dan F1-score. Kinerja klasifikasi yang berhasil dari model ini menunjukkan kemampuannya yang efektif dalam mengidentifikasi serta memahami sentimen positif pengguna. Sementara itu, keakuratan dalam mengklasifikasikan sentimen negatif tercermin dari keseimbangan antara recall dan presisi. Evaluasi ini memberikan landasan untuk penelitian berikutnya, seperti meningkatkan akurasi model dan menyesuaikan model dengan perubahan tren sentimen.

5. KESIMPULAN DAN SARAN

Berdasarkan hasil penelitian analisis sentimen aplikasi Grab pada ulasan Google Play Store dengan metode Naïve Bayes, dapat disimpulkan bahwa penerapan metode klasifikasi ini memberikan hasil yang memuaskan. Pengambilan 1000 data ulasan melalui proses scraping dan tahap preprocessing seperti case folding, tokenizing, stopwords, dan stemming telah berhasil menghasilkan model yang efektif dalam mengidentifikasi sentimen positif dan negatif dari ulasan pengguna. Dari evaluasi yang dilakukan, dapat dilihat bahwa algoritma Naïve Bayes mampu mengklasifikasikan ulasan dengan tingkat akurasi yang tinggi, mencapai 89%. Hal ini menunjukkan bahwa model yang dibangun cukup baik dalam memahami dan mengklasifikasikan sentimen ulasan pengguna aplikasi Grab.

Hasil analisis juga menunjukkan bahwa terdapat keseimbangan antara presisi dan recall untuk kedua kelas sentimen, positif dan negatif, yang menunjukkan kemampuan model dalam mengklasifikasikan dengan baik. Dengan demikian, penelitian ini memberikan kontribusi yang berarti dalam pemahaman terhadap sentimen pengguna terhadap aplikasi Grab serta memberikan landasan untuk pengembangan layanan yang lebih baik dan peningkatan kualitas platform bisnis.

Namun, untuk meningkatkan keakuratan model dan mengantisipasi perubahan tren sentimen di masa depan, disarankan untuk terus melakukan pemantauan dan penyesuaian model, serta mengintegrasikan faktor-faktor baru yang mungkin mempengaruhi sentimen pengguna. Selain itu, penelitian lanjutan dapat memperluas cakupan data dan menguji berbagai metode klasifikasi lainnya untuk mendapatkan pemahaman yang lebih komprehensif tentang sentimen pengguna.

DAFTAR PUSTAKA

- [1] F. N. Hasan and M. Dwijayanti, "Analisis Sentimen Ulasan Pelanggan Terhadap Layanan Grab Indonesia Menggunakan Multinomial Naïve Bayes Classifier," *J. Linguist. Komputasional*, vol. 4, no. 2, pp. 52–58, 2021, doi: <https://doi.org/10.26418/jlk.v4i2.61>.
- [2] S. Mandasari, B. H. Hayadi, and R. Gunawan, "Analisis Sentimen Pengguna Transportasi Online Terhadap Layanan Grab Indonesia Menggunakan Multinomial Naive Bayes Classifier," *J-SISKO TECH (Jurnal Teknol. Sist. Inf. dan Sist. Komput. TGD)*, vol. 5, no. 2, p. 118, 2022, doi: 10.53513/jsk.v5i2.5635.
- [3] D. R. Alghifari, M. Edi, and L. Firmansyah, "Implementasi Bidirectional LSTM untuk Analisis Sentimen Terhadap Layanan Grab Indonesia," *J. Manaj. Inform.*, vol. 12, no. 2, pp. 89–99, 2022, doi: 10.34010/jamika.v12i2.7764.
- [4] rahayu deny danar dan alvi furwanti Alwie, A. B. Prasetyo, R. Andespa, P. N. Lhokseumawe, and K. Pengantar, "Tugas Akhir Tugas Akhir," *J. Ekon. Vol. 18, Nomor 1 Maret 201*, vol. 2, no. 1, pp. 41–49, 2020.
- [5] A. Rahman, E. Utami, and S. Sudarmawan, "Sentimen Analisis Terhadap Aplikasi pada Google Playstore Menggunakan Algoritma Naïve Bayes dan Algoritma Genetika," *J. Komtika (Komputasi dan Inform.)*, vol. 5, no. 1, pp. 60–71, 2021, doi: 10.31603/komtika.v5i1.5188.
- [6] M. T. Anjasmoros, Istiadi, and F. Marisa, "Seminar Nasional Hasil Riset Prefix-RTR ANALISIS SENTIMEN APLIKASI GO-JEK MENGGUNAKAN METODE SVM DAN NBC (STUDI KASUS: KOMENTAR PADA PLAY STORE)," *Conf. Innov. Appl. Sci. Technol. (CIASTECH 2020)*, no. Ciastech, pp. 489–498, 2020.
- [7] L. Budi and A. Mude, "Perbandingan Metode Klasifikasi Support Vector Machine dan Naïve Bayes untuk Analisis Sentimen pada Ulasan Tekstual di Google Play Store," vol. 12, no. 2, pp. 154–161, 2020.
- [8] D. Darwis, N. Siskawati, and Z. Abidin, "Penerapan Algoritma Naive Bayes untuk Analisis Sentimen Review Data Twitter BMKG Nasional," vol. 15, no. 1, pp. 131–145.
- [9] A. Hendra, "Analisis Sentimen Review Halodoc Menggunakan Naïve Bayes Classifier," vol. 6, no. 2, pp. 78–89, 2021.
- [10] K. D. Indarwati, H. Februariyanti, U. Stikubank, P. Studi, and S. Informasi, "Analisis sentimen terhadap kualitas pelayanan aplikasi go-jek menggunakan metode naive bayes classifier," no. 1.
- [11] M. Irfan and D. Widiyanto, "Analisis Sentimen Pada Ulasan Aplikasi Samsat Mobile Jawa Barat (Sambara) Menggunakan algoritma Naive bayes Classifier dengan seleksi fitur Chi Square," pp. 337–346, 2022.
- [12] D. D. Putri, G. F. Nama, and W. E. Sulistiono, "ANALISIS SENTIMEN KINERJA DEWAN PERWAKILAN RAKYAT (DPR) PADA TWITTER MENGGUNAKAN METODE NAIVE BAYES CLASSIFIER," vol. 10, no. 1, pp. 34–40, 2022.
- [13] N. Herlinawati, Y. Yuliani, S. Faizah, W. Gata,

- and S. Samudi, “Analisis Sentimen Zoom Cloud Meetings di Play Store Menggunakan Naïve Bayes dan Support Vector Machine,” *CESS (Journal Comput. Eng. Syst. Sci.)*, vol. 5, no. 2, p. 293, 2020, doi: 10.24114/cess.v5i2.18186.
- [14] L. Rahmawati and D. B. Santoso, “Implementasi Metode Naïve Bayes Untuk Klasifikasi Ulasan Aplikasi E-Commerce Tokopedia,” *INTECOMS J. Inf. Technol. Comput. Sci.*, vol. 6, no. 1, pp. 116–124, 2023, doi: 10.31539/intecom.v6i1.5515.
- [15] F. Setya Ananto and F. N. Hasan, “Implementasi Algoritma Naïve Bayes Terhadap Analisis Sentimen Ulasan Aplikasi MyPertamina pada Google Play Store,” *J. ICT Inf. Commun. Technol.*, vol. 23, no. 1, pp. 75–80, 2023, [Online]. Available: <https://ejournal.ikmi.ac.id/index.php/jict-ikmi>
- [16] A. I. Tanggraeni and M. N. N. Sitokdana, “Analisis Sentimen Aplikasi E-Government pada Google Play Menggunakan Algoritma Naïve Bayes,” *JATISI (Jurnal Tek. Inform. dan Sist. Informasi)*, vol. 9, no. 2, pp. 785–795, 2022, doi: 10.35957/jatisi.v9i2.1835.