

## ALGORITMA NAIVE BAYES UNTUK MENINGKATKAN AKURASI MODEL KLASIFIKASI ANALISIS SENTIMEN PADA APLIKASI QUORA

Emha Caraka<sup>1\*</sup>, Nana Suarna<sup>2</sup>, Irfan Ali<sup>3</sup>, Dodi Solihudin<sup>4</sup>.

Program Studi Teknik Informatika<sup>124</sup>  
Program Studi Rekayasa Perangkat Lunak<sup>3</sup>  
STMIK IKMI Cirebon

Website : <https://ikmi.ac.id> Email : [info@ikmi.ac.id](mailto:info@ikmi.ac.id)  
[caraka822@gmail.com](mailto:caraka822@gmail.com)<sup>1\*</sup>, [st\\_nana@yahoo.com](mailto:st_nana@yahoo.com)<sup>2</sup>, [irfanaali0.0@gmail.com](mailto:irfanaali0.0@gmail.com)<sup>3</sup>,  
[dodisolihudin@gmail.com](mailto:dodisolihudin@gmail.com)<sup>4</sup>

(\*) Corresponding Author : [caraka822@gmail.com](mailto:caraka822@gmail.com)  
Published : 30 Desember 2025

**Abstract**— This research aims to improve the accuracy of sentiment analysis classification model on user reviews of Quora application in Google Play Store using Naive Bayes and K-Means Clustering algorithms. The dataset consists of 5000 Indonesian reviews collected through scraping, processed by preprocessing, TF-IDF weighting, clustering, and classification using Python on Google Colab. Model evaluation was conducted using accuracy, precision, recall, and F1-score metrics to ensure optimal performance. The evaluation results showed a model accuracy of 97%, with precision and recall reaching 95% for negative sentiment and 99% for positive sentiment, respectively. The Naive Bayes algorithm proved to be effective in handling large and diverse amounts of data, although challenges such as data imbalance and stylistic variations remain. The analysis process also involved a clustering stage using the K-Means algorithm to help group the review data into two sentiment categories, namely positive and negative, thus minimising the need for manual labelling. The combination of these algorithms is able to provide consistent and representative results in processing complex text data. This research makes an important contribution to the development of machine learning-based sentiment analysis methods and helps application developers better understand user needs and satisfaction. In addition, the results of this study can be a reference for other researchers to optimise the Naive Bayes algorithm in the context of sentiment analysis on community-based platforms such as Quora. The developed model is expected to be able to provide more accurate information to support data-based decision making.

**Keywords:** Naive Bayes, K-Means Clustering, Analisis Sentimen, Quora, Google Play Store.

**Intisari**— Penelitian ini bertujuan meningkatkan akurasi model klasifikasi analisis sentimen pada ulasan pengguna aplikasi Quora di *Google Play Store* menggunakan algoritma *Naive Bayes* dan *K-Means Clustering*. Dataset terdiri dari 5000 ulasan berbahasa Indonesia yang dikumpulkan melalui *scraping*, diolah dengan *preprocessing*, pembobotan *TF-IDF*, *clustering*, dan klasifikasi menggunakan *Python* di *Google Colab*. Evaluasi model dilakukan menggunakan metrik akurasi, presisi, *recall*, dan *F1-score* untuk memastikan kinerja yang optimal. Hasil evaluasi menunjukkan akurasi model sebesar 97%, dengan presisi dan *recall* masing-masing mencapai 95% untuk sentimen negatif dan 99% untuk sentimen positif. Algoritma Naive Bayes terbukti efektif dalam menangani data berjumlah besar dan beragam, meskipun tantangan seperti ketidakseimbangan data dan variasi gaya bahasa tetap ada. Proses analisis ini juga melibatkan tahapan *clustering* menggunakan algoritma *K-Means* untuk membantu mengelompokkan data ulasan ke dalam dua kategori sentimen, yaitu positif dan negatif, sehingga meminimalkan kebutuhan pelabelan manual. Kombinasi algoritma ini mampu memberikan hasil yang konsisten dan representatif dalam pengolahan data teks yang kompleks. Penelitian ini memberikan kontribusi penting dalam pengembangan metode analisis sentimen berbasis pembelajaran mesin dan membantu pengembang aplikasi memahami kebutuhan serta kepuasan pengguna secara lebih baik. Selain itu, hasil penelitian ini dapat menjadi referensi bagi peneliti lain untuk mengoptimalkan algoritma *Naive Bayes* dalam konteks analisis sentimen pada platform berbasis komunitas seperti Quora. Model yang dikembangkan diharapkan mampu memberikan informasi yang lebih akurat untuk mendukung pengambilan keputusan berbasis data.

**Kata Kunci:** Naive Bayes, K-Means Clustering, Analisis Sentimen, Quora, Google Play Store.

## PENDAHULUAN

Dalam era teknologi informasi yang terus berkembang, analisis sentimen menjadi salah satu alat penting untuk memahami persepsi pengguna terhadap aplikasi atau layanan tertentu. Sebagai salah satu platform berbasis komunitas, aplikasi Quora menyediakan berbagai ulasan pengguna pada *google play store* yang dapat memberikan wawasan berharga bagi pengembang aplikasi dalam meningkatkan kualitas layanan. Namun, analisis ulasan ini menghadapi berbagai tantangan, seperti volume data yang besar, variasi bahasa, dan bias data, sehingga diperlukan model klasifikasi yang andal dan akurat. Penelitian ini menerapkan algoritma *Naive bayes* sebagai alat utama analisis. Algoritma *Naive Bayes* sendiri [1] adalah algoritma klasifikasi yang digunakan untuk meramalkan probabilitas dan statistik.

Algoritma *Naive Bayes* sendiri telah terbukti menjadi salah satu metode yang efektif untuk analisis sentimen. Sebagai contoh, penelitian [2] menunjukkan bahwa *Naive Bayes* mampu mencapai akurasi hingga 83% pada analisis ulasan Shopee di *Google Play Store* dengan pembagian data tertentu. Penelitian lain [3] mengungkapkan bahwa algoritma ini mencapai akurasi 70,33% dalam analisis sentimen ulasan aplikasi Digital Korlantas POLRI. Meski demikian, berbagai penelitian mengindikasikan adanya ruang untuk peningkatan akurasi, khususnya dalam menghadapi data dengan distribusi tidak seimbang dan gaya bahasa yang beragam.

Penelitian ini bertujuan untuk mengetahui seberapa efektif Algoritma *Naive Bayes* dapat menganalisis volume data dari ulasan pengguna untuk mengidentifikasi sentimen positif atau negatif. Dengan memanfaatkan [4] *k-means clustering* yang biasa digunakan pada implementasi data *mining* untuk mengkluster atau mengelompokkan data menjadi beberapa kelompok data. Dan dalam penelitian ini, digunakan untuk mengelompokkan ulasan menjadi positif dan negatif. Selain itu, penelitian ini juga berupaya untuk mengatasi ketidakakuratan bahasa yang dipakai oleh pengguna aplikasi Quora di *google play store* dalam memberi ulasan menggunakan algoritma *Naive Bayes*. Selanjutnya, penelitian ini bertujuan untuk meningkatkan akurasi algoritma *Naive Bayes* dalam mengklasifikasikan sentimen ulasan pengguna aplikasi Quora di *Google Play Store*.

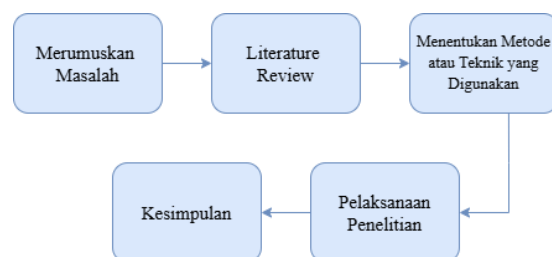
Melalui pendekatan kuantitatif, penelitian ini melibatkan berbagai tahap, termasuk pengumpulan data melalui *scraping*, *preprocessing* menggunakan teknik *Natural Language Processing (NLP)*, pembobotan *TF-IDF*, *clustering*, *splitting* data dan

klasifikasi. [5] Hasil evaluasi model akan diukur menggunakan metrik akurasi, presisi, *recall*, dan *F1-score* untuk memastikan keandalan analisis.

Penelitian ini diharapkan dapat memberikan kontribusi signifikan dalam pengembangan metode analisis sentimen. Yang dimana analisis sentimen digunakan untuk menemukan informasi berharga yang dibutuhkan dari data yang tidak terstruktur [6], sehingga diharapkan pada penelitian ini dapat diketahui sentimen pengguna. Tidak hanya pada aplikasi Quora tetapi juga pada platform lain yang membutuhkan pendekatan serupa. Hasil penelitian ini juga diharapkan mampu memberikan wawasan yang lebih mendalam bagi pengembang aplikasi dalam memahami kebutuhan dan kepuasan pengguna.

## BAHAN DAN METODE

Penelitian ini menggunakan pendekatan kuantitatif dengan desain eksploratif untuk meningkatkan akurasi model klasifikasi sentimen pada ulasan pengguna aplikasi Quora di *Google Play Store*. Pendekatan kuantitatif ini diterapkan melalui proses *text mining*, yang meliputi tahapan *information extraction (IE)*, *information retrieval (IR)*, *natural language processing (NLP)*, *clustering*, *categorization*, *visualization* dan *evaluating classification*. Teknik ini telah terbukti efektif dalam memenuhi kebutuhan proses klasifikasi dokumen dengan berbagai jenis konten [7]. Algoritma *Naive Bayes* digunakan untuk klasifikasi sentimen, yang digabungkan dengan *clustering K-Means* untuk mengatasi tantangan dalam analisis sentimen berbasis komunitas. Tahapan tersebut dilakukan, untuk mendapatkan pola-pola penting yang dapat meningkatkan akurasi model klasifikasi. Tahapan metode penelitian dapat dilihat pada gambar di bawah ini.



Sumber data penelitian ini tergolong ke sumber data sekunder. Karena, penelitian ini berasal dari ulasan pengguna aplikasi yang ada pada *google play store*. Ketersediaan data untuk penelitian ini mencakup dataset dengan jumlah sebanyak 5000 *record*. Dataset ini mencakup informasi-informasi terkait *username*, *score*, *at*, dan

*content* yang diperlukan untuk melakukan analisis sentimen klasifikasi menggunakan algoritma *Naive Bayes*. Dengan jumlah data sebanyak 5000 *record* tersebut, diharapkan penelitian dapat memberikan peluang untuk meningkatkan akurasi klasifikasi.

Teknik pengumpulan data melibatkan tahap *NLP*, termasuk pembersihan teks, normalisasi, *stopword removal*, tokenisasi, dan stemming menggunakan pustaka *Sastrawi*. Data kemudian direpresentasikan secara numerik menggunakan metode *Term Frequency-Inverse Document Frequency (TF-IDF)*. Analisis data dilakukan dalam dua tahap: pertama, *clustering* menggunakan algoritma *K-Means* untuk mengelompokkan ulasan berdasarkan pola sentimen menjadi kategori positif dan negatif. Kedua, klasifikasi sentimen diterapkan menggunakan *Naive Bayes*. Evaluasi model dilakukan dengan metrik akurasi, *presisi*, *recall*, dan *F1-score*. Semua analisis didukung oleh *platform Google Colab* untuk memanfaatkan sumber daya komputasi yang efisien, dengan tujuan menghasilkan model klasifikasi yang akurat dan dapat diandalkan. Visualisasi pengambilan data, bisa dilihat pada Gambar 2 visualiasi pengambilan data di bawah ini.



Gambar 2 visualiasi pengambilan data

## HASIL DAN PEMBAHASAN

Penelitian ini bertujuan untuk meningkatkan akurasi model klasifikasi sentimen pada ulasan pengguna aplikasi Quora di *Google Play Store* menggunakan algoritma *Naive Bayes* dan *K-Means*. Berikut adalah hasil dan pembahasan menggunakan *Text Mining* yang diperoleh berdasarkan tahapan penelitian:

### 1. Information Extraction (IE)

Proses *scraping* menggunakan pustaka *google-play-scraper* menghasilkan dataset berisi 5000 ulasan pengguna. Proses *scraping* data ini juga memiliki kemampuan untuk mengubah data yang tidak terstruktur menjadi terstruktur, sehingga dapat disimpan dalam basis data [8]. Data ini mencakup elemen penting seperti teks ulasan (*content*) berbahasa Indonesia, skor, dan tanggal publikasi ulasan. Hasil *scraping* ini adalah data mentah atau data yang masih belum diolah. Proses serta hasil

dari *scraper* bisa dilihat pada gambar dan tabel di bawah ini.

```

from google_play_scraper import Sort, reviews

result, continuation_token = reviews(
    'com.quora.android',
    lang='id',
    country='id',
    sort=Sort.MOST_RELEVANT,
    count=5000,
    filter_score_with=None
)
  
```

Gambar 1 code scraping

Tabel 1. hasil scraping

| No  | Content                                           | Skor | Tanggal    |
|-----|---------------------------------------------------|------|------------|
| 1   | Algoritmanya jelek banget! Pasti yang muncul...   | 2    | 2024-12-16 |
| 2   | Aplikasi berkualitas, cocok untuk saya yang me... | 5    | 2024-07-14 |
| ... | ...                                               | ...  | ...        |
| 49  | bagus banget, banyak pengetahuan baru...          | 5    | 2019-02-20 |
| 50  | Banyak sekali pengetahuan baru yg saya...         | 5    | 2019-01-30 |

### 2. Information Retrieval (IR)

Seperti yang terlihat pada tabel 1, hasil data mentah dari tahapan *scraping* menyertakan data-data yang tidak perlukan. Data-data yang telah didapatkan akan disaring(ekstraksi) kembali untuk diambil data-data pentingnya saja. Data penting yang dimaksud yaitu adalah data *content*(ulasan). seperti yang dapat dilihat pada gambar dan tabel di bawah ini.

```
my_df = sorted_df[['content']]
```

Gambar 2 code ekstraksi

Tabel 2. hasil ekstraksi

| No | Content                                                                                                                                                                                                                                                                                                                                                                                                                                          |
|----|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| 1  | Bisa bertanya apapun<br>Algoritmanya jelek banget! Pasti yang muncul di feed itu konten personal, generik, basi (mirip konten facebook anying), gak fresh, gak menarik, dan gak spesifik. Malah banyak meme juga padahal meme mah cukup di Facebook aja. Padahal gw udah mencet x (saya tidak ingin melihat ini) berkali-kali, tapi tetep aja. Saya ke Quora untuk konten informatif, fresh, edukatif, ilmiah, menghibur, berbeda, dan spesifik! |
| 2  | saya baru daftar, kok pas masuk lagi ke apknya ada tulisan "akun anda telah dilarang" padahal baru daftar dan tidak melakukan apapun                                                                                                                                                                                                                                                                                                             |
| 3  | Terimakasih Quaraa 🙏🙏                                                                                                                                                                                                                                                                                                                                                                                                                            |
| 4  | Sangat bermanfaat                                                                                                                                                                                                                                                                                                                                                                                                                                |

### 3. Natural Language Processing (NLP)

Setelah data berhasil diekstraksi, pada langkah ini akan dilakukan tahap *text preprocessing* yang

meliputi: *Cleaning Text*, *Case Folding*, *Stopword Removal*, Tokenisasi dan *Stemming* [9].

#### A. Cleaning Text dan Case Folding

Langkah pertama *text preprocessing* akan dilakukan *cleaning text* dan juga *case folding*. *cleaning text* berfungsi untuk menghapus simbol, angka, atau karakter yang tidak relevan. Serta *case folding* berfungsi untuk mengubah semua teks menjadi huruf kecil (*lowercase*) [10]. Hasil proses ini, bisa dilihat pada tabel di bawah ini.

Tabel 3. hasil cleaning text dan case folding

| No | Content                                                                                                                                                                                                                                                                                                                                                                                                                  |
|----|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| 1  | bisa bertanya apapun algoritmanya jelek banget pasti yang muncul di feed itu konten personal generik basi mirip konten facebook anying gak fresh gak menarik dan gak spesifik malah banyak meme juga padahal meme mah cukup di facebook aja padahal gw udah mencet x saya tidak ingin melihat ini berkali-kali tapi tetep aja saya ke quora untuk konten informatif fresh edukatif ilmiah menghibur berbeda dan spesifik |
| 2  | saya baru daftar kok pas masuk lagi ke apknya ada tulisan akun anda telah dilarang padahal baru daftar dan tidak melakukan apapun                                                                                                                                                                                                                                                                                        |
| 3  | terimakasih quaraa                                                                                                                                                                                                                                                                                                                                                                                                       |
| 4  | sangat bermanfaat                                                                                                                                                                                                                                                                                                                                                                                                        |

#### B. Stopword Removal

Menghapus kata-kata tidak berarti atau kata umum yang tidak memberikan kontribusi signifikan, seperti "dan", "atau", "di" dan lain-lain [11]. Hasil proses ini, bisa dilihat pada tabel di bawah ini.

Tabel 4. hasil stopwords removal

| No | Content                                                                                                                                                                                                                                                                             |
|----|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| 1  | bisa apapun algoritmanya jelek banget muncul feed konten personal generik basi konten facebook anying fresh menarik spesifik banyak meme meme mah cukup facebook gw udah mencet x tidak berkali-kali tetep quora konten informatif fresh edukatif ilmiah menghibur berbeda spesifik |
| 2  | baru daftar masuk apknya tulisan akun dilarang baru daftar tidak apapun                                                                                                                                                                                                             |
| 3  | terimakasih quaraa                                                                                                                                                                                                                                                                  |
| 4  | bermanfaat                                                                                                                                                                                                                                                                          |

#### C. Tokenizing

Memecah teks menjadi unit kata per-kata atau menjadi unit-unit yang lebih kecil yang disebut dengan "token" [12]. Seperti contoh: "aplikasi bagus bermanfaat" menjadi "aplikasi", "bagus", "bermanfaat". Hasil proses ini, bisa dilihat pada tabel di bawah ini.

Tabel 5. hasil tokenizing

| No | Content            |
|----|--------------------|
| 1  | ['bisa', 'apapun'] |

|   |                                                                                                                                                                                                                                                                                                                                                                                         |
|---|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| 1 | ['algoritmanya', 'jelek', 'banget', 'muncul', 'feed', 'konten', 'personal', 'generik', 'basi', 'konten', 'facebook', 'anying', 'fresh', 'menarik', 'spesifik', 'banyak', 'meme', 'meme', 'mah', 'cukup', 'facebook', 'gw', 'udah', 'mencet', 'x', 'tidak', 'berkali-kali', 'tetep', 'quora', 'konten', 'informatif', 'fresh', 'edukatif', 'ilmiah', 'menghibur', 'berbeda', 'spesifik'] |
| 2 | ['baru', 'daftar', 'masuk', 'apknya', 'tulisan', 'akun', 'dilarang', 'baru', 'daftar', 'tidak', 'apapun']                                                                                                                                                                                                                                                                               |
| 3 | ['terimakasih', 'quaraa']                                                                                                                                                                                                                                                                                                                                                               |
| 4 | ['bermanfaat']                                                                                                                                                                                                                                                                                                                                                                          |

#### D. Stemming

Proses *stemming* yaitu untuk mengubah kata yang terdapat imbuhan menjadi kata dasar [13]. Tahap *stemming* ini kita dapat menggunakan *library Python* Sastrawi untuk mengubah ke kata dasar bahasa Indonesia. Hasil dari proses ini, bisa dilihat pada tabel di bawah ini.

Tabel 5. hasil stemming

| No | Content                                                                                                                                                                                                                                                                 |
|----|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| 1  | bisa apa algoritmanya jelek banget muncul feed konten personal generik basi konten facebook anying fresh tarik spesifik banyak meme meme mah cukup facebook gw udah mencet x tidak berkali-kali tetep quora konten informatif fresh edukatif ilmiah hibur beda spesifik |
| 2  | baru daftar masuk apknya tulis akun larang baru daftar tidak apa                                                                                                                                                                                                        |
| 3  | terimakasih quaraa                                                                                                                                                                                                                                                      |
| 4  | manfaat                                                                                                                                                                                                                                                                 |

#### 4. Clustering

Pada tahap ini, data hasil dari *text preprocessing* akan diberikan bobot menggunakan *tf-idf* terlebih dahulu. Menurut buku Pengantar Metode Analisis Sentimen, menerangkan bahwa *Naive Bayes* mampu menentukan apakah ulasan positif atau negatif dalam bentuk angka melalui penerapan metode *term frequency-inverse document frequency (TF-IDF)*. Setiap kata dalam ulasan diberi bobot yang dihitung menggunakan frekuensi kemunculan kata dalam dokumen [14]. Setelah data teks diubah menjadi numerik, algoritma *naive bayes* kemudian dapat menghitung probabilitas setiap kelas (positif atau negatif) untuk setiap ulasan. Pada tahap pembobotan ini, menggunakan *tf-idf* vectorizer dari *library sklearn* untuk mengubah data teks menjadi representasi numerik dengan membatasi jumlah fitur (kata atau istilah unik) maksimum yang digunakan sebanyak 4700. *Code* proses *tf-idf* bisa dilihat pada gambar di bawah ini.

```
from sklearn.feature_extraction.text import TfidfVectorizer
tfidf = TfidfVectorizer(max_features=4700)
X = tfidf.fit_transform(df['text_steaming']).toarray()
```

Gambar 3 code tf-idf

Selanjutnya, algoritma *k-Means clustering* digunakan untuk mengelompokkan data berdasarkan kesamaan fitur numerik ke dalam *cluster* dan sentimen [15], yaitu 0 untuk negatif dan 1 untuk positif. Pengkodean dan hasil *clustering* bisa dilihat pada gambar dan tabel berikut ini.

```
from sklearn.cluster import KMeans
# Clustering dengan K-Means,
# kita set jumlah cluster menjadi 2 (positif dan negatif)
kmeans = KMeans(n_clusters=2, random_state=0, n_init=10)

df['Cluster'] = kmeans.fit_predict(X)

cluster_to_sentiment = {0: "Negatif", 1: "Positif"}
df['Sentiment'] = df['Cluster'].map(cluster_to_sentiment)
df[['text_steamingdo', 'Cluster', 'Sentiment']].head(5)
```

Gambar 4 code clustering

| No | text_steamindo                                                                                                                                                                        | cluster | sentiment          |
|----|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------|--------------------|
| 1  | bisa apa<br>algoritmanya jelek banget<br>muncul feed konten<br>personal generik basi<br>konten facebook anying<br>fresh tarik spesifik banyak                                         | 0       | Negatif<br>Negatif |
| 2  | meme meme mah cukup<br>facebook gw udah mencet x<br>tidak berkali kali tetep<br>quora konten informatif<br>fresh edukatif ilmiah hiburan<br>beda spesifik<br>baru daftar masuk apknya | 0       | Negatif            |
| 3  | tulis akun larang baru<br>daftar tidak apa                                                                                                                                            | 0       |                    |
| 4  | terimakasih quaraa                                                                                                                                                                    | 0       | Negatif            |
| 5  | manfaat                                                                                                                                                                               | 1       | Positif            |

## 5. Categorization

Setelah hasil *cluster* didapatkan, selanjutnya kita masuk ke tahap *data split*. Yang dimana, data dibagi menjadi data latih (*training*) dan data uji (*testing*). Data *training* digunakan untuk membentuk model klasifikasi yang akan digunakan untuk prediksi kelas, sedangkan data *testing* digunakan untuk mengukur performa dari model yang telah dibentuk data *training* [16]. Data akan dibagi dengan *test\_size* = 0.3 (70:30) dan *random state*-nya 42. Pengkodean ini, bisa dilihat pada gambar di bawah ini.

[illegible]

Gambar 5 code split data

Hasil *spliting* data dari pembagian *test\_size* = 0.3 (70:30) menghasilkan data *train* sebanyak 3500 dan data *test* sebanyak 1500. Hasil *spliting* bisa dilihat pada Gambar 7 hasil *splitting* data di bawah ini.

```
Jumlah Data Train: 3500
Jumlah Data Test: 1500
```

Data Latih (Train):

[illegible]

Data Uji (Test):

[illegible]

Gambar 6 hasil splitting data

Setelah ini, kita akan melakukan pemodelan algoritma *Naive Bayes* dengan parameter ( $\alpha=0.1$ ) serta *data set* yang digunakan hasil dari *resampling* menggunakan *smote (oversampling)*. Seperti yang dapat dilihat pada gambar di bawah ini.

```
nb_model = MultinomialNB(alpha=0.1)
nb_model.fit(X_resampled, y_resampled)
```

Gambar 7 pemodelan *Naive Bayes*

## 6. Visualization

Hasil klasifikasi dapat divisualisasikan menggunakan *wordcloud*. Visualisasi ini bertujuan untuk memberikan gambaran kata-kata yang paling sering muncul dalam ulasan pengguna berdasarkan kategori sentimen, baik positif maupun negatif [17]. Hasil dari visualisasi ini, dibagi menjadi 3, yaitu *worldcloud* gabungan (positif dan negatif) *woldcloud* negatif dan terakhir *worldcloud* positif. Hasil visualisasi tersebut bisa dilihat pada 3 gambar di bawah ini



Gambar 8 worldcloud gabungan



Gambar 9 woldcloud negatif


$$\begin{bmatrix} 910 & 46 \\ 6 & 885 \end{bmatrix}$$

| Classification Report: |      | precision | recall | f1-score | support |
|------------------------|------|-----------|--------|----------|---------|
| Negatif                | 0.99 | 0.95      | 0.97   | 956      |         |
| Positif                | 0.95 | 0.99      | 0.97   | 891      |         |
| accuracy               |      |           | 0.97   | 1847     |         |
| macro avg              | 0.97 | 0.97      | 0.97   | 1847     |         |
| weighted avg           | 0.97 | 0.97      | 0.97   | 1847     |         |

```
cm = confusion_matrix(y_test, y_pred, labels=["Negatif", "Positif"])

# Menghitung metrik evaluasi manual
accuracy = accuracy_score(y_test, y_pred)
precision = precision_score(y_test, y_pred, pos_label="Positif")
recall = recall_score(y_test, y_pred, pos_label="Positif")
f1 = f1_score(y_test, y_pred, pos_label="Positif")

# Menampilkan metrik evaluasi utama
print(f"\nMultinomialNB Accuracy: {accuracy}")
print(f"\nMultinomialNB Precision: {precision}")
print(f"\nMultinomialNB Recall: {recall}")
print(f"\nMultinomialNB f1_score: {f1}")

# Menampilkan Confusion Matrix
print("Confusion Matrix:")
print(cm)
print("\n")

# Menampilkan Classification Report
print("\nClassification Report:")
print(classification_report(y_test, y_pred))
```

Hasil penelitian ini mendukung pengembangan teknologi pemrosesan bahasa alami, khususnya dalam analisis sentimen dengan dataset besar. Pendekatan ini tidak hanya meningkatkan akurasi model, tetapi juga memberikan wawasan penting bagi pengembang aplikasi dalam memahami kebutuhan pengguna. Dengan demikian, hasil ini dapat berkontribusi pada peningkatan kualitas layanan dan pengalaman pengguna secara keseluruhan.

<https://ruangjurnal.or.id>

Berdasarkan penelitian yang difokuskan pada analisis sentimen ulasan pengguna aplikasi Quora menggunakan algoritma *Naive Bayes* dan *k-means clustering*, diperoleh beberapa kesimpulan penting. Pertama, algoritma *Naive Bayes* terbukti efektif dalam mengklasifikasikan sentimen positif dan negatif dari ulasan pengguna. Dengan dataset sebanyak 5000 ulasan, model mencapai akurasi 97%, dengan *precision* yang sangat baik pada sentimen negatif (99%) dan positif (95%), serta *F1-score* konsisten sebesar 97% untuk kedua kategori. Evaluasi menggunakan *confusion matrix* menunjukkan kemampuan algoritma dalam menangani data ulasan yang beragam. Kedua, langkah *preprocessing* data seperti *Cleaning Text*, *case folding*, *stopword removal*, *tokenizing*, dan *stemming* berhasil meningkatkan kualitas data analisis dengan mengurangi dampak ketidakakuratan bahasa, termasuk penggunaan bahasa slang, kesalahan ejaan, dan variasi bahasa informal. Ketiga, penggunaan teknik pembobotan *TF-IDF* dan metode *oversampling* meningkatkan akurasi model, terutama dalam menangani ketidakseimbangan kelas. Hasil menunjukkan algoritma bekerja konsisten pada kedua sentimen, mencerminkan keberhasilan pendekatan yang diterapkan. Kesimpulan ini mendukung rumusan masalah dan tujuan penelitian, sekaligus menegaskan efektivitas algoritma *Naive Bayes* dalam meningkatkan kualitas analisis sentimen pada aplikasi berbasis komunitas seperti Quora. Untuk penelitian lanjutan, disarankan memperluas dataset agar mencakup lebih banyak variasi bahasa serta mengeksplorasi algoritma lain seperti *SVM*, *RNN*, atau *Random Forest* guna meningkatkan akurasi. Selain itu, mempertimbangkan analisis sentimen netral dapat memberikan wawasan lebih mendalam. Pengembang aplikasi diharapkan memanfaatkan hasil analisis ini untuk menyusun strategi peningkatan kualitas layanan secara tepat sasaran.

#### REFERENCE

- [1] M. Y. Siregar, A. Davy Wiranata, and R. A. Saputra, "KLIK: Kajian Ilmiah Informatika dan Komputer Analisis Sentimen Pada Ulasan Pengguna Aplikasi Streaming Vidio Menggunakan Metode Naïve Bayes," *Media Online*, vol. 4, no. 5, pp. 2419–2429, 2024, doi: 10.30865/klik.v4i5.1787.
- [2] N. Agustina, D. H. Citra, W. Purnama, C. Nisa, and A. R. Kurnia, "Implementasi Algoritma Naive Bayes untuk Analisis Sentimen Ulasan Shopee pada Google Play Store," *MALCOM Indones. J. Mach. Learn. Comput. Sci.*, vol. 2, no. 1, pp. 47–54, 2022, doi: 10.57152/malcom.v2i1.195.
- [3] S. H. W. Putra and D. Febriawan, "Analisis Sentimen Ulasan Aplikasi Digital Korlantas POLRI Menggunakan Naïve Bayes pada Google Play Store," *KLIK Kaji. Ilm. Inform. dan Komput.*, vol. 4, no. 4, pp. 1962–1971, 2024, doi: 10.30865/klik.v4i4.1600.
- [4] N. N. Hasanah and A. S. Purnomo, "Implementasi Data Mining Untuk Pengelompokan Buku Menggunakan Algoritma K-Means Clustering (Studi Kasus : Perpustakaan Politeknik LPP Yogyakarta)," *J. Teknol. Dan Sist. Inf. Bisnis*, vol. 4, no. 2, pp. 300–311, 2022, doi: 10.47233/jteksis.v4i2.499.
- [5] A. Christian *et al.*, "ANALISIS MACHINE LEARNING UNTUK PREDIKSI PENYAKIT PARU-PARU MENGGUNAKAN Citra paru," vol. 7, pp. 122–131, 2025.
- [6] A. P. Giovani, A. Ardiansyah, T. Haryanti, L. Kurniawati, and W. Gata, "Analisis Sentimen Aplikasi Ruang Guru Di Twitter Menggunakan Algoritma Klasifikasi," *J. Teknoinfo*, vol. 14, no. 2, p. 115, 2020, doi: 10.33365/jti.v14i2.679.
- [7] F. Fathonah and A. Herliana, "Penerapan Text Mining Analisis Sentimen Mengenai Vaksin Covid - 19 Menggunakan Metode Naïve Bayes," *J. Sains dan Inform.*, vol. 7, no. 2, pp. 155–164, 2021, doi: 10.34128/jsi.v7i2.331.
- [8] H. Z. Muflih, A. R. Abdillah, and F. N. Hasan, "Analisis Sentimen Ulasan Pengguna Aplikasi Ajaib Menggunakan Metode Naïve Bayes," *KLIK Kaji. Ilm. Inform. dan Komput.*, vol. 4, no. 3, pp. 1613–1621, 2023, doi: 10.30865/klik.v4i3.1303.
- [9] S. Juanita, "Analisis Sentimen Persepsi Masyarakat Terhadap Pemilu 2019 Pada Media Sosial Twitter Menggunakan Naive Bayes," *J. Media Inform. Budidarma*, vol. 4, no. 3, p. 552, 2020, doi: 10.30865/mib.v4i3.2140.
- [10] Y. Firmansyah, R. Kurniawan, and Y. A. Wijaya, "Analisis Data Sentimen Pemain Game Role-Playing Game (RPG) Honkai Star Rail dengan Algoritma Naive Bayes," *J. Inform. dan Rekayasa Perangkat Lunak*, vol. 6, no. 1, pp. 127–135, 2024.
- [11] S. I. Nurhafida and F. Sembiring, "Analisis Sentimen Aplikasi Novel Online Di Google Play Store Menggunakan Algoritma Support Vector Machine (SVM)," *J. Sains Komput. Inform. (J-SAKTI)*, vol. 6, no. 1, pp. 317–327, 2022.
- [12] D. Nurwahidah, G. Dwilestari, N. Dienwati

- Nuris, and R. Narasati, "Analisis Sentimen Data Ulasan Pengguna Aplikasi Google Kelas Pada Google Play Store Menggunakan Algoritma Naïve Bayes," *JATI (Jurnal Mhs. Tek. Inform.*, vol. 7, no. 6, pp. 3673–3678, 2024, doi: 10.36040/jati.v7i6.8245.
- [13] V. A. Permadi, "Analisis Sentimen Menggunakan Algoritma Naive Bayes Terhadap Review Restoran di Singapura," *J. Buana Inform.*, vol. 11, no. 2, pp. 141–151, 2020, doi: 10.24002/jbi.v11i2.3769.
- [14] D. Purnamasari *et al.*, *Pengantar Metode Analisis Sentimen*. 2023.
- [15] Anita Rahayu, "Perbedaan Klasifikasi dan Clustering," Binus University. Accessed: Dec. 26, 2024. [Online]. Available: [https://binus.ac.id/malang/2022/09/perbedaan-klasifikasi-dan-clustering/#:~:text=Klasifikasi adalah memprediksi klasifikasi sebuah,\(kebutuhan dan perilaku konsumsi\)](https://binus.ac.id/malang/2022/09/perbedaan-klasifikasi-dan-clustering/#:~:text=Klasifikasi adalah memprediksi klasifikasi sebuah,(kebutuhan dan perilaku konsumsi).).
- [16] Y. R. R. Iqbar Sajidan Widiyanto and M. A. K. Komara, "Analisis sentimen e-wallet gopay, shopeepay, dan ovo menggunakan algoritma naive bayes," vol. 12, no. 3, 2024.
- [17] A. Y. R. Shinta Nilam Sari Muslim, Firman Nurdiansyah, "PERBANDINGAN ALGORITMA NAIVE BAYES DAN KNN DALAM ANALISIS SENTIMEN ULASAN PENGGUNA APLIKASI CAPCUT," *J. Inform. dan Teknol.*, vol. 2, no. 3, pp. 61–69, 2024.
- [18] D. Surya Sayogo, B. Irawan, and A. Bahtiar, "Analisis Sentimen Ulasan Instagram Di Google Play Store Menggunakan Algoritma Naïve Bayes," *JATI (Jurnal Mhs. Tek. Inform.*, vol. 7, no. 6, pp. 3314–3319, 2024, doi: 10.36040/jati.v7i6.8178.
- [19] D. F. Zhafira, B. Rahayudi, and I. Indriati, "Analisis Sentimen Kebijakan Kampus Merdeka Menggunakan Naive Bayes dan Pembobotan TF-IDF Berdasarkan Komentar pada Youtube," *J. Sist. Informasi, Teknol. Informasi, dan Edukasi Sist. Inf.*, vol. 2, no. 1, pp. 55–63, 2021, doi: 10.25126/justsi.v2i1.24.