



# Pemodelan Analisis Sentimen Ulasan Pengguna Aplikasi Info Bmkg Menggunakan Pendekatan Multinomial Naïve Bayes

**Moh. Syaogi<sup>1\*</sup>, Nur Ariesanto Ramdhan<sup>2</sup>, Otong Saeful Bachri<sup>3</sup>, Bambang Irawan<sup>4</sup>**

<sup>1\*,2,3,4</sup>Program Studi Teknik Informatika, Universitas Muhadi Setiabudi, Indonesia

\*Email : [mohammadsyaogi301204@gmail.com](mailto:mohammadsyaogi301204@gmail.com)

## **Abstract:**

*Info BMKG is one of several digital platforms that have been pushed by the fast evolution of IT to replace traditional methods of providing public services. Reviews on the Play Store can be used to determine user perceptions and levels of satisfaction with the application. Manual analysis is laborious and inefficient due to the high number of evaluations. Consequently, the purpose of this research is to use the Naive Bayes algorithm to categorize evaluations of the Info BMKG app as either positive or negative in order to do sentiment analysis. Using a web scraping approach, a total of 5,000 user evaluations were obtained for the study data. Next, the data underwent text preprocessing, word weighting using the TF-IDF technique, and sentiment classification with the Multinomial Naive Bayes algorithm. There was an 80:20 split between the dataset's training and testing sets. The experimental findings show that the Naive Bayes algorithm achieves an accuracy of 87.83% on the testing data when it comes to classifying user review emotions.*

**Keywords:** *BMKG info; Sentiment Analysis; Naive Bayes; User Reviews*

## **1. PENDAHULUAN**

Perkembangan teknologi digital telah mendorong pemanfaatan aplikasi berbasis mobile sebagai media penyampaian informasi kepada masyarakat secara lebih cepat dan efisien (Winoto dkk., 2024). Badan Meteorologi, Klimatologi, dan Geofisika (BMKG) sebagai lembaga pemerintah menyediakan aplikasi Info BMKG untuk memudahkan masyarakat dalam mengakses informasi cuaca, iklim, gempa bumi, dan peringatan dini bencana secara cepat dan real-time melalui perangkat mobile (Firmansyah & Puspitasari, 2021).

Google Play Store tidak hanya berfungsi sebagai platform distribusi aplikasi berbasis Android, tetapi juga sebagai media bagi pengguna untuk memberikan penilaian dan ulasan berdasarkan pengalaman mereka dalam menggunakan suatu aplikasi (Lubis dkk., 2023). Ulasan pengguna tersebut mencerminkan tingkat kepuasan, keluhan, serta saran yang sangat berguna bagi pengembang dalam meningkatkan kualitas aplikasi (Mauliddiyah dkk., 2024). Aplikasi Info BMKG memiliki jumlah pengguna yang besar dan menerima ribuan ulasan dengan beragam opini, baik positif maupun negatif (Khotimah dkk., 2025).

Namun, menganalisis semua ulasan secara manual sangat merepotkan dan memakan banyak waktu. Oleh karena itu, diperlukan sistem otomatis yang dapat mengklasifikasikan pandangan pengguna dengan benar dan cepat (Pramuji dkk., 2024). Analisis sentimen, sebuah metode untuk mengkategorikan pandangan atau emosi dalam teks ke dalam kategori sentimen tertentu seperti sentimen positif dan negatif, adalah salah satu metodologi yang sering digunakan untuk mengevaluasi opini pengguna (Nurrochmah dkk., 2025).

Algoritma Naïve Bayes sering digunakan dalam analisis sentimen berbasis teks karena memiliki proses yang sederhana, efisien, dan mampu memberikan performa yang baik meskipun menggunakan data latih dalam jumlah terbatas. Dalam analisis sentimen berbasis teks, banyak penelitian telah menunjukkan bahwa algoritma Naïve Bayes memiliki kinerja yang kompetitif dibandingkan dengan teknik klasifikasi lainnya seperti Support Vector Machine (Supian dkk., 2024).

Penelitian ini bertujuan untuk memanfaatkan konteks yang telah disebutkan di atas untuk melakukan analisis sentimen pada ulasan aplikasi Info BMKG yang ditemukan di Google Play Store menggunakan metode Naive Bayes. Wawasan tentang bagaimana pengguna mempersepsikan aplikasi tersebut diharapkan dapat diperoleh dari temuan penelitian ini (Ramadhan & Juardi, 2025).

## **2. TINJAUAN PUSTAKA**

### **2.1 Analisis Sentimen**

Analisis sentimen merupakan metode pengolahan data teks yang bertujuan untuk mengklasifikasikan opini atau emosi pengguna ke dalam kategori sentimen tertentu, seperti positif dan negatif, guna mengetahui kecenderungan sikap atau penilaian pengguna kepada suatu produk atau layanan (Mauliddiyah dkk., 2024).

### **2.2 Ulasan Pengguna**

Ulasan pengguna merupakan bentuk evaluasi yang diberikan setelah menggunakan suatu aplikasi dan mencerminkan pengalaman serta tingkat kepuasan pengguna. Pada platform Google Play Store, ulasan pengguna menjadi salah satu sumber informasi yang bernilai bagi pengembang untuk mengidentifikasi kebutuhan dan permasalahan pengguna (Lubis dkk., 2023).

### **2.3 Analisis Sentimen Pada Ulasan Aplikasi**

Analisis sentimen ulasan aplikasi digunakan untuk mengelompokkan opini pengguna secara otomatis sehingga memudahkan pengembang dalam memahami kecenderungan sentimen yang dominan. Pendekatan ini relevan diterapkan pada aplikasi layanan publik yang memiliki jumlah pengguna dan ulasan yang besar (Khotimah dkk., 2025).

## 2.4 Text Mining

*Text mining* merupakan proses untuk mengekstraksi dan memperoleh informasi penting dari sekumpulan dokumen teks dengan memanfaatkan teknik analisis tertentu, yang termasuk ke dalam bagian dari data mining. Tujuan dari proses ini untuk mengolah teks yang tidak terstruktur atau semi-terstruktur agar dapat menghasilkan informasi yang bernilai dan bermanfaat (Effendi & Ernawati, 2025).

## 2.5 Algoritma Naive Bayes

Algoritma Naive Bayes adalah teknik klasifikasi probabilistik yang menggunakan Teorema Bayes dengan premis bahwa fitur-fitur bersifat independen. Dengan menggunakan kemungkinan kemunculan kata, metode ini mengklasifikasikan teks ke dalam kategori sentimen untuk analisis sentimen (Jiwangga dkk., 2024).

## 3. METODOLOGI PENELITIAN

Gambar 1 menunjukkan langkah-langkah pendekatan penelitian yang digunakan untuk studi ini. Dimulai dengan pengumpulan data ulasan pengguna aplikasi, diagram alir secara metodis menguraikan tahapan penelitian, termasuk pelabelan data, pra-pemrosesan teks, pembobotan kata, pembagian data, klasifikasi sentimen menggunakan algoritma Naive Bayes, evaluasi kinerja model, dan penarikan kesimpulan.



Gambar 1. Diagram Alur Penelitian

### 3.1 Pengumpulan Data

Pengumpulan data dilakukan melalui teknik *web scraping* untuk memperoleh ulasan pengguna aplikasi Info BMKG yang tersedia pada platform Google Play Store. Pengumpulan data dilakukan dengan menggunakan bahasa pemrograman Python melalui *library google-play-scraper* pada Google Colab, yang memungkinkan ekstraksi informasi ulasan secara otomatis. Dalam penelitian analisis sentimen terhadap ulasan aplikasi, teknik *web scraping* sering digunakan karena kemampuannya dalam mengakuisisi data teks ulasan, rating, dan metadata lainnya secara masif dari platform tersebut (Khotimah dkk., 2025). Dataset yang digunakan dalam penelitian ini berjumlah 5.000 ulasan terbaru yang diambil pada tanggal 4 Januari 2026, yang meliputi konten

teks ulasan (*review text*), tanggal ulasan, serta nilai rating yang diberikan oleh pengguna, sehingga dapat merepresentasikan opini pengguna terkini terhadap aplikasi Info BMKG.

### 3.2 Pelabelan Data

Tahap pelabelan data dilakukan setelah proses pengumpulan ulasan selesai untuk membentuk dataset yang siap digunakan dalam pelatihan (*training*) dan pengujian (*testing*) model. Pelabelan dilakukan dengan mengaitkan nilai rating yang diberikan pengguna dengan kategori sentimen. Sentimen positif didefinisikan sebagai ulasan dengan peringkat 4 atau 5, sedangkan sentimen negatif didefinisikan sebagai ulasan dengan nilai 1 atau 2. Karena bersifat netral dan dapat memengaruhi hasil, ulasan dengan peringkat 3 tidak digunakan (Mola dkk., 2024). Pendekatan pelabelan otomatis berdasarkan nilai rating ini merupakan metode yang umum digunakan dalam penelitian analisis sentimen karena memberikan label yang konsisten dan mudah direplikasi untuk skenario klasifikasi biner (A'la, 2022).

### 3.3 Preprocessing Data

Tahapan *preprocessing* dilakukan untuk mengonversi data teks mentah ke dalam bentuk yang lebih terstruktur sehingga dapat diolah oleh sistem dengan baik (Saputro & Hermawan, 2021). Proses *preprocessing* atau prapemrosesan merupakan tahapan penting dalam analisis sentimen karena kualitas hasil klasifikasi sangat dipengaruhi oleh kualitas data yang digunakan. Data yang telah diproses dengan baik akan menghasilkan representasi teks yang lebih optimal untuk proses klasifikasi. Tahapan *preprocessing* yang diterapkan meliputi:

#### 3.3.1 Cleaning

Tahapan *cleaning* dilakukan untuk membersihkan data dari duplikasi, nilai kosong (*missing value*), serta ulasan yang bersifat tidak relevan atau mengandung *noise*. Selain itu, karakter atau atribut yang tidak berkontribusi terhadap proses klasifikasi, seperti simbol, angka acak, maupun tautan yang sering muncul dalam teks ulasan, dihapus agar data menjadi lebih bersih dan relevan untuk dianalisis (Firmansyah & Puspitasari, 2021).

#### 3.3.2 Case Folding

Konten ulasan diubah menjadi huruf kecil. Kata-kata yang terdengar mirip tetapi memiliki ejaan yang berbeda akan distandarisasi melalui pendekatan ini.

#### 3.3.3 Slangword Normalization

*Slangword Normalization* dilakukan dengan mengonversi kata tidak baku atau singkatan (*slang*) yang terdapat pada ulasan pengguna ke dalam bentuk kata baku atau kata dengan makna sebenarnya. Proses ini bertujuan untuk menyeragamkan penulisan kata sehingga variasi kata dengan makna yang sama tidak dianggap sebagai kata yang berbeda oleh sistem.

#### 3.3.4 Tokenizing

Tahap *tokenizing* dilakukan dengan memecah teks ulasan menjadi unit-unit yang lebih kecil berupa kata atau token. Tahapan ini memungkinkan setiap kata pada ulasan diidentifikasi dan dianalisis secara terpisah pada tahap pemrosesan selanjutnya.

### 3.3.5 Stopword Removal

*Stopword removal*, untuk melakukan ini, kami mengecualikan kata penghubung, kata ganti, dan istilah umum lainnya yang tidak memberikan kontribusi signifikan terhadap penentuan emosi. Tahapan ini bertujuan guna mengurangi kata-kata yang tidak informatif sehingga hanya kata-kata yang relevan terhadap konteks ulasan yang diproses oleh model.

### 3.3.6 Stemming

*Stemming* merupakan tahap akhir preprocessing yang dilakukan dengan menghilangkan imbuhan seperti *prefiks* dan *sufiks*, serta mengubah kata ke dalam bentuk dasarnya. Tahapan ini dilakukan untuk menyatukan berbagai variasi kata yang memiliki kesamaan makna sehingga dapat meningkatkan kinerja proses klasifikasi sentimen.

## 3.4 Pembobotan Kata

Pembobotan kata dilakukan menggunakan pendekatan TF-IDF, yang merupakan singkatan dari *Term Frequency-Inverse Document Frequency*. Untuk menentukan bobot suatu kata, teknik ini mempertimbangkan baik frekuensi kemunculannya—seberapa sering kata tersebut muncul dalam sebuah dokumen—dan frekuensi kemunculannya dalam dokumen—seberapa jarang kata tersebut muncul dalam keseluruhan kumpulan dokumen (Andriani & Wibowo, 2021).

Persamaan TF-IDF dapat dituliskan sebagai berikut:

$$W_{dt} = TF_{dt} \times IDF_{ft} \quad (1)$$

Keterangan:

- $W_{dt}$ : bobot suatu kata ke-t pada dokumen ke-d
- $TF_{dt}$ : menyatakan frekuensi kemunculan kata tertentu dalam satu dokumen
- $IDF_{ft}$ : *Inverse Document Frequency*

$$\log \left( \frac{N}{df} \right)$$

## 3.5 Pemodelan Algoritma

Algoritma Naive Bayes, yang merupakan pendekatan klasifikasi berbasis probabilitas yang didasarkan pada Teorema Bayes, digunakan untuk melakukan prosedur pemodelan. Algoritma ini beroperasi dengan premis bahwa semua fitur data sepenuhnya terpisah satu sama lain. Kemudian, data dibagi 80:20, dengan 20% digunakan untuk pengujian dan 80% untuk pelatihan. Model probabilistik dilatih menggunakan data pelatihan untuk menentukan kemungkinan terjadinya setiap kelas. Kemudian menggunakan data uji, model tersebut digunakan untuk memprediksi kelas tersebut (Putri dkk., 2023).

Persamaan dasar Teorema Bayes yang digunakan untuk menghitung probabilitas bersyarat adalah sebagai berikut:

$$P(H | X) = \frac{P(X|H) \cdot P(H)}{P(X)} \quad (2)$$

Keterangan:

- X : Bukti atau fitur yang diamati
- H : Hipotesis atau kelas yang diprediksi

- $P(H | X)$ : Probabilitas posterior, yaitu probabilitas hipotesis setelah mempertimbangkan bukti yang tersedia
- $P(H)$ : Probabilitas awal (*prior*) dari hipotesis
- $P(X | H)$ : Peluang kemunculan bukti X dengan asumsi bahwa hipotesis H benar
- $P(X)$ : Probabilitas total dari bukti X

### 3.6 Evaluasi Model

Ukuran umum yang digunakan untuk mengevaluasi kinerja model klasifikasi meliputi skor F1, recall, akurasi, dan presisi. Untuk lebih memeriksa distribusi prediksi akurat dan tidak akurat untuk setiap kelas, digunakan *confusion matrix*. Informasi lebih lanjut mengenai kinerja model dalam klasifikasi sentimen dapat ditemukan dalam *confusion matrix*, yang menunjukkan hubungan antara prediksi model dan klasifikasi aktual, termasuk True Positive (TP), False Positive (FP), False Negative (FN), dan True Negative (TN).

Tabel 1. Confusion Matrix

| Kelas<br>Prediksi | Positif<br>(Aktual)    | Negatif<br>(Aktual)    |
|-------------------|------------------------|------------------------|
| Positif           | True Positive<br>(TP)  | False Positive<br>(FP) |
| Negatif           | False Negative<br>(FN) | True Negative<br>(TN)  |

Dengan keterangan sebagai berikut:

1. TP : Data yang diprediksi kelas positif dan sesuai dengan kelas aktual.
2. TN : Data yang diprediksi sebagai kelas negatif dan sesuai dengan kelas sebenarnya.
3. FP : Data yang diprediksi sebagai kelas positif, namun pada kondisi aktual termasuk ke dalam kelas negatif.
4. FN : Data yang diprediksi sebagai kelas negatif, tetapi pada kenyataannya berada pada kelas positif.

## 4. HASIL DAN PEMBAHASAN

### 4.1 Hasil Preprocessing Data

Tahap ini bertujuan untuk menyiapkan teks ulasan agar dapat diproses lebih baik oleh algoritma. Proses ini meliputi pembersihan teks (*cleaning*), *case folding*, normalisasi, *tokenizing*, *stopword removal*, hingga *stemming*. Setelah seluruh tahapan diterapkan, data menjadi lebih terorganisir serta siap masuk ke tahap analisis.

Dataset awal terdiri dari 5.000 ulasan, namun setelah proses preprocessing beberapa ulasan memiliki nilai kosong (*missing value*). Setelah menghapus *missing value*, tersisa 4.187 ulasan yang siap digunakan. Untuk memastikan distribusi kelas yang adil, data kemudian dibagi secara bertingkat 80:20 menjadi set pelatihan dan set pengujian. Hasil pembagian data adalah sebagai berikut:

- Data Training: 3349 ulasan
- Data Testing: 838 ulasan

Pembagian secara stratified memastikan bahwa proporsi sentimen positif dan negatif pada data training dan testing tetap seimbang, sehingga model dapat dilatih dan diuji secara adil.

```
data_filtered = data.dropna(subset=['Category']).reset_index(drop=True)

X = data_filtered['filtered_text']
y = data_filtered['Category']

X_train, X_test, y_train, y_test = train_test_split(
    X, y, test_size=0.2, random_state=42, stratify=y
)

print("Train:", X_train.shape[0], "Test:", X_test.shape[0])

Train: 3349 Test: 838
```

Gambar 2. Proses Split Data

#### 4.2 Hasil Pembobotan Kata TF-IDF

Tahap berikutnya adalah melakukan pembobotan kata menggunakan metode TF-IDF. Parameter yang digunakan pada proses ini meliputi:

- max\_features = 4000
- ngram\_range = (1, 2)
- min\_df = 2

Setelah dilakukan proses *fitting*, sistem menghasilkan total fitur yang digunakan sebanyak:

```
# 1 Import library
from sklearn.feature_extraction.text import TfidfVectorizer
from sklearn.naive_bayes import MultinomialNB
from sklearn.metrics import accuracy_score, classification_report, confusion_matrix
import seaborn as sns
import matplotlib.pyplot as plt

# 2 TF-IDF vectorization (ngram 1-2, max_features tuning)
vectorizer = TfidfVectorizer(max_features=4000, ngram_range=(1,2), min_df=2)
X_train_tfidf = vectorizer.fit_transform(X_train)
X_test_tfidf = vectorizer.transform(X_test)

print("Dimensi fitur (TF-IDF):", X_train_tfidf.shape)

Dimensi fitur (TF-IDF): (3349, 2420)
```

Gambar 3. Dimensi Fitur TF-IDF

#### 4.3 Hasil Pelatihan Model Naïve Bayes

Model dilatih menggunakan algoritma Multinomial Naïve Bayes dengan parameter alpha = 0.3, yang berfungsi sebagai *smoothing* untuk mengatasi masalah frekuensi kata rendah. Hasil evaluasi pada data training adalah sebagai berikut:

```
# 3 Inisialisasi dan latih model Naïve Bayes
mnb = MultinomialNB(alpha=0.3)
mnb.fit(X_train_tfidf, y_train)

# 4 Evaluasi pada data training
train_pred = mnb.predict(X_train_tfidf)
train_acc = accuracy_score(y_train, train_pred)
print("Akurasi Data Training: {:.2f}%".format(train_acc * 100))

Akurasi Data Training: 92.06%
```

Gambar 4. Akurasi Data Training

Tingginya nilai akurasi pada data training menunjukkan bahwa model dapat mempelajari pola kata dan distribusi sentimen dengan baik.

#### 4.4 Hasil Pengujian Model

Setelah proses pelatihan, model diuji dengan menggunakan 838 ulasan pada data testing. Hasil evaluasi yang diperoleh adalah:

```
# 5 Evaluasi pada data testing
test_pred = mnbc.predict(X_test_tfidf)
test_acc = accuracy_score(y_test, test_pred)
print("Akurasi Data Testing: {:.2f}%".format(test_acc * 100))
Akurasi Data Testing: 87.83%
```

Gambar 5. Akurasi Data Testing

Akurasi testing yang tinggi dan tidak berbeda jauh dari data training menandakan bahwa model tidak mengalami *overfitting* serta mampu melakukan generalisasi dengan baik pada data baru.

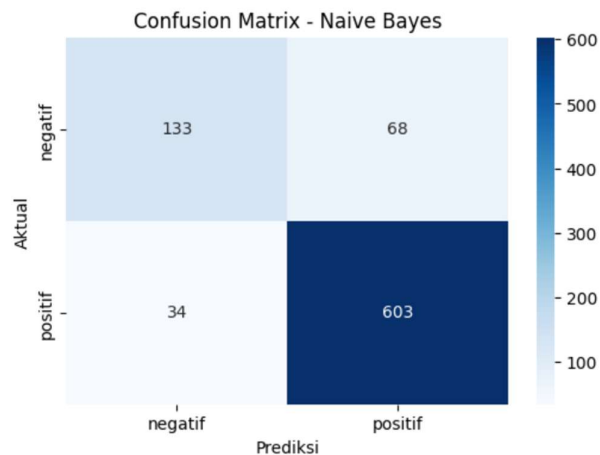
Selain itu, *classification report* berikut juga menunjukkan performa model pada masing-masing kelas:

```
=== Classification Report (Data Testing) ===
              precision    recall  f1-score   support

negatif      0.80      0.66      0.72      201
positif      0.90      0.95      0.92      637

accuracy          0.88      838
macro avg      0.85      0.80      0.82      838
weighted avg   0.87      0.88      0.87      838
```

Gambar 6. Classification Report (Data Testing)



Gambar 7. Confusion Matrix

## 5. KESIMPULAN

Untuk menganalisis nada ulasan Info BMKG, penelitian ini menggunakan metode Naive Bayes. Kesimpulan berikut dapat ditarik dari data tersebut:

1. Tahap preprocessing yang mencakup *cleaning*, *case folding*, *slangword normalization*, *tokenizing*, *stopword removal*, dan *stemming* berhasil mempersiapkan data sehingga layak digunakan pada tahap pemodelan. Dari 5.000 ulasan awal, sebanyak 4.187 ulasan dapat digunakan setelah menghapus data yang memiliki nilai kosong (*missing value*).
2. Model Naïve Bayes yang dilatih menggunakan pembagian data sebesar 80% data pelatihan dan 20% data pengujian melalui *stratified split* menunjukkan performa yang sangat baik. Model memperoleh akurasi sebesar 92,06% pada data pelatihan dan 87,83% pada data pengujian, yang menunjukkan bahwa model memiliki kemampuan generalisasi yang kuat.
3. Hasil evaluasi memperlihatkan bahwa model memiliki performa yang tinggi pada sentimen positif, dengan nilai *precision* 0,90, *recall* 0,95, dan *F1-score* 0,92. Performa pada sentimen negatif lebih rendah dengan nilai *F1-score* 0,72 karena jumlah datanya jauh lebih sedikit dibandingkan ulasan positif.

Secara keseluruhan, algoritma Naïve Bayes terbukti memiliki efektivitas yang baik dalam mengklasifikasikan sentimen ulasan pengguna Info BMKG dengan tingkat akurasi yang tinggi. Hasil analisis sentimen menunjukkan bahwa sebagian besar ulasan pengguna tergolong dalam kategori sentimen positif.

### UCAPAN TERIMA KASIH (OPSIONAL)

Penulis mengucapkan terima kasih kepada pihak-pihak yang telah memberikan dukungan dalam pelaksanaan penelitian ini, sehingga penelitian ini dapat diselesaikan dengan baik.

### REFERENSI

- A'la, F. Y. (2022). Indonesian Sentiment Analysis towards MyPertamina Application Reviews by Utilizing Machine Learning Algorithms. *Journal of Informatics, Information System, Software Engineering and Applications*, 8106, 80–91. <https://doi.org/10.20895/INISTA.V5I1.838>
- Andriani, N., & Wibowo, A. (2021). Implementasi Text Mining Klasifikasi Topik Tugas Akhir Mahasiswa Teknik Informatika Menggunakan Pembobotan TF-IDF dan Metode Cosine Similarity Berbasis Web. *Seminar Nasional Mahasiswa Ilmu Komputer Dan Aplikasinya (SENAMIKA)*, September, 130–137.
- Effendi, P. A., & Ernawati, T. (2025). ANALISIS SENTIMEN ULASAN APLIKASI GAME HAY DAY MENGGUNAKAN ALGORITMA RANDOM FOREST. *JITET (Jurnal Informatika Dan Teknik Elektro Terapan)*, 13(3), 1–8.
- Firmansyah, Z., & Puspitasari, N. F. (2021). ANALISIS SENTIMEN MASYARAKAT TERHADAP VAKSINASI COVID-19 BERDASARKAN OPINI PADA TWITTER MENGGUNAKAN ALGORITMA NAIVE BAYES. *JURNAL TEKNIK INFORMATIKA*, 14(2), 171–178.
- Jiwangga, A. T., Ramdhan, N. A., Hamid, A., Premana, A., Bhakti, R. M. H., Studi, P., Informatika, T., Teknik, F., Setiabudi, U. M., & Tengah, P. J. (2024). Analisis Sentimen pada Ulasan Tempat Wisata Taman Pancasila Kota Tegal Menggunakan Metode Naive Bayes. *QISTINA: Jurnal Multidisiplin Indonesia*, 3(2), 1220–1227.

- 
- Khotimah, K., Dikananda, A. R., & Rifa'i, A. (2025). ANALISIS SENTIMEN ULASAN APLIKASI PINTU DI GOOGLE PLAY STORE MENGGUNAKAN ALGORITMA NAÏVE BAYES. *JITET (Jurnal Informatika Dan Teknik Elektro Terapan)*, 13(1).
- Lubis, S. K., Dar, M. H., & Nasution, F. A. (2023). Analisis Sentimen Ulasan Pengguna Aplikasi pada Google Play Store Menggunakan Algoritma Support Vector Machine. *INFORMATIKA*, 11(2).
- Mauliddiyah, S., Hidayat, M. N. F., & Rizal, F. (2024). ANALISIS SENTIMENT ULASAN APLIKASI PEMBELAJARAN DUOLINGGO DI PLAY STORE MENGGUNAKAN DISTILBERT. *Jurnal TEKINKOM*, 7, 502–511. <https://doi.org/10.37600/tekinkom.v7i1.1395>
- Mola, S. A. S., Baun, D. L. B., Nunes, I. O., & Sani, M. M. A. R. (2024). ANALISIS SENTIMEN APLIKASI HALO BCA DI GOOGLE PLAY STORE MENGGUNAKAN METODE NAIVE BAYES , SUPPORT VECTOR MACHINE DAN RANDOM FOREST. *HOAQ: JURNAL TEKNOLOGI INFORMASI*, 15(c), 69–79.
- Nurrochmah, D. S., Rahaningsih, N., Dana, R. D., & Rohmat, C. L. (2025). PENERAPAN ALGORITMA NAIVE BAYES DALAM ANALISIS SENTIMEN ULASAN APLIKASI KITALULUS DI GOOGLE PLAY STORE. *Jurnal Informatika Terpadu*, 11(1), 1–11.
- Pramuji, M. C., Purnamasari, R., & Eliskar, Y. (2024). Analisis Sentimen Menggunakan Algoritma Naive Bayes Classifier Pada Ulasan Aplikasi PLN Mobile di Google Play Store. *E-Proceeding of Engineering*, 11(6), 5700–5706.
- Putri, K. S., Setiawan, I. R., & Pambudi, A. (2023). ANALISIS SENTIMEN TERHADAP BRAND SKINCARE LOKAL MENGGUNAKAN NAÏVE BAYES CLASSIFIER. *Technologia*, 14(3), 227–232.
- Ramadhan, W. P., & Juardi, D. (2025). ANALISIS SENTIMEN ULASAN APLIKASI BTN MOBILE MENGGUNAKAN ALGORITMA NAIVE BAYES. *JITET (Jurnal Informatika Dan Teknik Elektro Terapan)*, 13(1).
- Saputro, T. H., & Hermawan, A. (2021). The Accuracy Improvement of Text Mining Classification on Hospital Review through The Alteration in The Preprocessing Stage. *International Journal of Computer and Information Technology*, 10(4), 140–146.
- Supian, A., Revaldo, B. T., Marhadi, N., & Efrizoni, L. (2024). Perbandingan Kinerja Naïve Bayes dan SVM pada Analisis Sentimen Twitter Ibukota Nusantara. *Jurnal Ilmiah Informatika (JIF)*.
- Winoto, D., Aditia, V. D., Sorisa, C., Priskila, R., & Pranatawijaya, V. H. (2024). ANALISIS SENTIMEN PADA ULASAN PENGGUNA TERHADAP APLIKASI PEMBELAJARAN BAHASA DUOLINGO : MENGGUNAKAN ALGORITMA NAÏVE BAYES DAN K-NEAREST NEIGHBOR. *JATI (Jurnal Mahasiswa Teknik Informatika)*, 8(3), 3230–3236.