

PENGUNAAN ALGORITMA NAIVE BAYES UNTUK ANALISIS SENTIMEN PADA ULASAN APLIKASI *THREADS* DI *GOOGLE PLAY STORE*

USE OF NAIVE BAYES ALGORITHM FOR SENTIMENT ANALYSIS ON THREADS APP REVIEWS ON GOOGLE PLAY STORE

R. Bagus Bambang Sumantri¹, Dede Yusuf², Asharryadi Noegroho³

¹Program Studi Sistem Informasi Universitas Harapan Bangsa, ²Program Studi Informatika, ³Program Studi Bisnis Digital Universitas Al Irsyad Cilacap
e-mail :bagus100486@gmail.com

Abstrak

Kemajuan teknologi informasi mendorong pengguna untuk memberikan ulasan terhadap aplikasi digital, termasuk *Threads*, platform berbasis teks yang dirilis oleh Meta Platforms Dengan menggunakan Algoritma *Naive Bayes*, penelitian ini bertujuan menganalisa ulasan pengguna aplikasi *Threads* di *Google Play Store*. Tahapan penelitian meliputi pengumpulan data melalui teknik *scraping*, pemrosesan awal data atau pra-proses, pembobotan menggunakan Teknik TF-IDF, serta klasifikasi sentimen dengan membagi data ke dalam set pelatihan (*training*) dan pengujian (*testing*). Hasil analisis memperlihatkan bahwa banyak pembahasan yang menggunakan memiliki sentimen positif, dengan tingkat akurasi klasifikasi sebesar 84%. Untuk sentimen negatif, *precision*, *recall*, dan *f1-score* masing-masing mencapai 74%, 65%, dan 69%. Sementara itu, sentimen positif memiliki *precision* sebesar 84%, *recall* 89%, dan *f1-score* 86%. Meskipun model berhasil mengklasifikasikan sentimen dengan baik, masih terdapat kesalahan identifikasi pada ulasan negatif, yang disebabkan oleh faktor seperti bahasa ambigu dan sentimen campuran. Penelitian ini bertujuan untuk memberikan informasi yang baik untuk pengembang agar bisa meningkatkan kualitas aplikasi *Threads* pada pengalaman user. Untuk penelitian selanjutnya, disarankan membandingkan algoritma lain seperti k-Nearest Neighbour atau kernel RBF serta menerapkan metode *k-fold cross validation* untuk meningkatkan keakuratan model.

Kata kunci: Analisis sentimen, *Naive Bayes*, *Threads*, ulasan pengguna, *Google Play Store*.

Abstract

Advances in information technology encourage users to review digital applications, including Threads, a text-based platform released by Meta Platforms Using the Naive Bayes Algorithm, this study aims to analyze user reviews of Threads applications on the Google Play Store. The research stages include data collection through scraping techniques, data initial processing or pre-processing, weighting using TF-IDF techniques, and sentiment classification by dividing data into training and testing sets. The results of the analysis show that many of the discussions that use have positive sentiments, with a classification accuracy rate of 84%. For negative sentiment, precision, recall, and f1-score reached 74%, 65%, and 69%, respectively. Meanwhile, positive sentiment has a precision of 84%, a recall of 89%, and an f1-score of 86%. Although the model managed to classify sentiment well, there were still misidentification in negative reviews, caused by factors such as ambiguous language and mixed sentiment. This research aims to provide good information for developers to improve the quality of Threads applications in the user experience. For further research, it is recommended to compare other algorithms such as k-Nearest Neighbour or RBF kernel and apply the k-fold cross validation method to improve the accuracy of the model.

Keywords: Sentiment analysis, *Naive Bayes*, *Threads*, user reviews, *Google Play Store*.

1. PENDAHULUAN

Perkembangan teknologi informasi telah menghasilkan perubahan signifikan dalam pola komunikasi manusia, khususnya dengan hadirnya aplikasi yang berfokus pada media sosial. Salah satu aplikasi yang menonjol adalah *Threads*, sebuah platform komunikasi yang dikembangkan oleh *Meta Platforms, Inc.* Aplikasi ini memungkinkan pengguna untuk berbagi konten dengan lingkaran sosial yang lebih kecil, seperti teman dekat atau kelompok diskusi tertentu. Namun, ulasan di *Google Play Store* menunjukkan bahwa aplikasi ini tidak luput dari kritik pengguna, seperti keluhan terkait performa, fitur, dan pengalaman antarmuka (2).

Pada platform distribusi seperti *Google Play Store*, meningkatnya jumlah pengguna dan ulasan menimbulkan masalah yang signifikan untuk menganalisis sentimen ulasan pengguna. Masalah utama terletak pada bagaimana melakukan analisis sentimen secara efektif dan akurat, terutama untuk aplikasi seperti *Threads*. Kesulitan ini termasuk kompleksitas bahasa alami, variasi dalam gaya penulisan, dan konteks khusus ulasan. Mengingat semakin banyak pengguna yang bergantung pada ulasan sebagai referensi utama saat memilih aplikasi, sangat penting untuk mengetahui pendapat pengguna tentang aplikasi tersebut. Akibatnya, Penelitian ini mengusulkan penggunaan Algoritma *Naive Bayes* sebagai solusi yang efektif. Selain itu, tujuan dari penelitian ini adalah untuk mendukung pengembangan Teknik analisis sentimen yang lebih canggih dalam bidang informatika, penelitian ini juga mengusulkan penggunaan Algoritma *Naive Bayes*.

Metode penting untuk memahami pendapat publik tentang aplikasi ini adalah melihat ulasan pengguna. Dalam memproses data ulasan yang sangat besar, *Algoritma Naive Bayes*, merupakan salah satu metode dalam klasifikasi yang paling umum digunakan dalam proses pemrosesan bahasa alami (NLP), telah terbukti mampu menghasilkan hasil yang efektif dan akurat. Penelitian sebelumnya menunjukkan bahwa *algoritma* ini dapat mencapai akurasi tinggi, misalnya pada analisis ulasan aplikasi lain seperti *KAI Access* dan *by.U*, dengan hasil mencapai tingkat akurasi hingga 88% (2)

Penelitian ini bertujuan mengeksplorasi lebih lanjut potensi *Algoritma Naive Bayes* dalam menganalisa pembahasan aplikasi *Threads*. Fokus utama adalah untuk mengidentifikasi akar permasalahan dari ulasan negatif dan memberikan rekomendasi perbaikan yang relevan. Untuk bisa mengembangkan aplikasi dan memberikan kontribusi sangat berharap pada penelitian kali ini, serta penelitian ini juga akan meningkatkan pengalaman pengguna, serta memperluas pemahaman tentang penerimaan dan penyempurnaan aplikasi yang berfokus pada pengalaman pengguna.

Penelitian ini meliputi serangkaian langkah yang dimulai dengan pengumpulan dataset ulasan, diikuti oleh proses *preprocessing*, dan diakhiri dengan penerapan *Algoritma Naive Bayes*

untuk mengklasifikasikan sentimen dari ulasan yang ada. Hasil dari penelitian ini diharapkan dapat memberikan kontribusi signifikan dalam memahami analisis sentimen ulasan pengguna, serta memberikan wawasan mendalam mengenai tingkat kepuasan pengguna *Threads*. Selain itu, temuan penelitian ini dapat berfungsi sebagai pedoman praktis bagi para pengembang aplikasi. Selanjutnya, hasil penelitian ini juga dapat dijadikan acuan penting agar terus berkembangnya metode analisis sentimen yang lebih baik dari masa sebelumnya.(12)

2. TINJAUAN PUSTAKA

2.1. Penelitian terdahulu

Fokus penelitian sebelumnya adalah pelaksanaan pemilihan umum (PEMILU), terutama menjelang pilkada, yang menyebabkan banyak perdebatan di masyarakat. Akibatnya, penelitian ini memeriksa tweet yang membahas masalah publik terkait pemilihan pemimipaan daerah 2020 saat pandemi COVID-19. (8. Selanjutnya dalam penelitian lain yang menjadi rujukan membahas layanan transportasi online. Untuk memahami bagaimana masyarakat memandang layanan transportasi online, penelitian ini menggunakan analisis sentimen. Data yang digunakan diperoleh dari media sosial Twitter yang berkaitan dengan layanan *GrabId* dan *GojekIndonesia*. Dalam penelitian ini, metode analisis sentimen yang diterapkan meliputi *Naive Bayes Classifier* serta *Support Vector Machine*.(1).

2.2. Aplikasi *Threads*

Aplikasi *Threads* merupakan aplikasi yang serumpun dengan *WhatsApp*, *Instagram*, dan *Facebook*, yang merupakan bagian dari *Meta*, perusahaan induk, meluncurkan sebuah platform media sosial yang berbasis teks dan disebut dengan Aplikasi *Threads*. Aplikasi ini didesain sebagai pelengkap *Instagram* untuk memfasilitasi percakapan berbasis teks antara pengguna. Diluncurkan pada Juli 2023, *Threads* mendapatkan perhatian luas sebagai pesaing Twitter, menawarkan fitur yang memungkinkan pengguna untuk berbagi ide, berita, dan opini secara singkat melalui teks, dengan dukungan gambar, video, serta interaksi seperti *like*, *reply*, dan *repost*.

2.3. Algoritma *Naive Bayes*

Metode *Naive Bayes* merupakan teknik yang digunakan untuk menghitung kemungkinan bahwa sebuah dokumen dengan kata-kata tertentu mencerminkan jenis sentimen tertentu. Metode ini sangat efektif dalam klasifikasi teks, terutama dalam analisis sentimen, di mana tujuannya adalah untuk menentukan apakah sebuah dokumen atau teks memiliki sentimen netral, positif atau negatif, berdasarkan kata-kata yang terdapat di dalamnya.(12). Formula teoritis dari metode *Naive Bayes* disajikan di bawah ini:

$$P(X|H) = \frac{P(X|H) \times P(H)}{P(X)} \dots\dots\dots (1)$$

Keterangan:

X merupakan Sampel Data dengan kategori (label) yang tidak diketahui.

H adalah Hipotesis yang menyatakan bahwa X adalah data dengan kategori (label) C.

$P(H|X)$ adalah Probabilitas bahwa hipotesis tersebut benar (valid) untuk sampel data X yang telah teliti.

$P(X|H)$ adalah Probabilitas sampel data X, jika diasumsikan bahwa hipotesis tersebut benar (valid).

$P(H)$ adalah Probabilitas hipotesis H.

$P(X)$ adalah Probabilitas dari sampel data yang telah diamati

2.4. Analisis Sentimen

Analisis sentimen bertujuan untuk mengolah perbedaan pendapat mengenai produk, layanan, atau lembaga tertentu, dan merupakan gabungan antara *text mining* dan *data mining*. Produk, layanan, atau lembaga tertentu menjadi fokus analisis perbedaan pendapat dari konsumen atau ahli, yang bertujuan untuk mengolah informasi tersebut (10). Proses ini melibatkan serangkaian langkah, mulai dari pengolahan data hingga penarikan kesimpulan, serta kadang kala juga mencakup penilaian terhadap tingkat emosi yang tercermin dalam teks, seperti sentimen positif, negatif, atau netral. Keunggulan utama dari analisis sentimen terletak pada kemampuannya untuk memahami respons atau persepsi audiens serta pelanggan terhadap layanan, peristiwa tertentu, produk, atau merk. Hal ini sangat mendukung pengambilan keputusan yang lebih berbasis informasi bagi organisasi. Analisis sentimen memberikan pemahaman yang lebih mendalam tentang bagaimana individu bereaksi terhadap suatu topik atau situasi. Dengan kemampuan untuk mengenali ekspresi emosional seperti kemarahan, kesedihan, atau kegembiraan, analisis ini tidak hanya dapat mengidentifikasi sentimen sebagai positif, negatif, atau netral (6).

2.5. Google Collaboratory

Proyek dengan kebutuhan komputasi tinggi, seperti *deep learning* dan *machine learning*, sangat diuntungkan oleh akses ke sumber daya *google cloud*, termasuk GPU (*Graphic Procces Unit*) dan TPU (*Tensor Procces Unit*), yang merupakan salah satu fitur unggulan *Colab*. Tanpa perlu konfigurasi atau instalasi tambahan, pengguna dapat menulis dan menjalankan kode *Python* menggunakan *Jupyter Notebook* dalam lingkungan berbasis *cloud*. *Google Colaboratory*, yang lebih dikenal sebagai *Google Colab*, merupakan platform pengembangan gratis yang disediakan oleh *Google*. *Google Collabs* juga Memberikan akses langsung kepada kolaborator, *Colab* memungkinkan pengguna untuk berbagi *notebook* secara online, sehingga memfasilitasi kolaborasi tim. Pengguna juga dapat dengan mudah menyimpan, mengelola, dan mengakses *notebook* mereka berkat integrasi yang mulus dengan *Google Drive*. Ini menjadi fitur yang sangat

membantu *Proyek Python* lainnya dan pengembangan model *machine learning* menjadikan *Google Colaboratory* sebagai favorit di kalangan praktisi *data science*, pelajar, dan peneliti. (5).

2.6. Confusion matrix

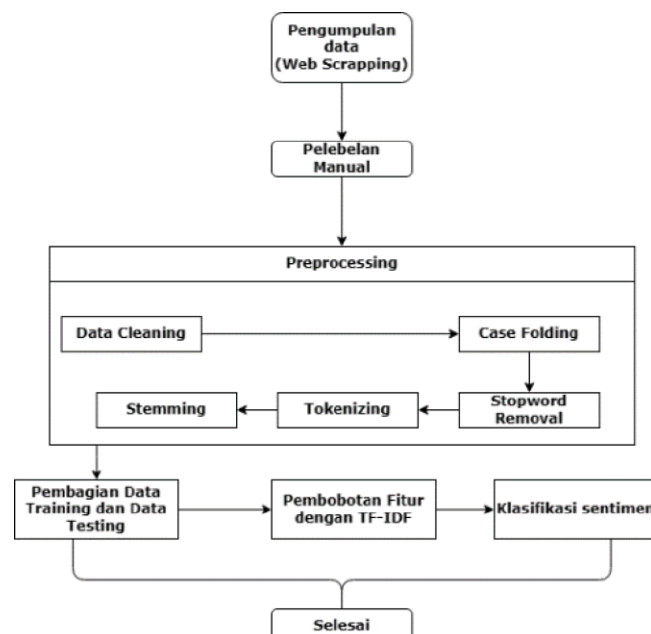
Confusion matrix adalah tools dipakai untuk mengevaluasi kinerja prototype dengan mencatat jumlah prediksi yang benar serta yang salah. Matriks ini dapat diterapkan pada klasifikasi dengan dua kelas atau lebih (7). Dalam penelitian ini, indikator seperti *recall*, *precision*, *accuracy*, dan *F1-score* digunakan agar bisa melihat tingkat akurasi dari hasil prediksi. Data yang dihasilkan oleh struktur *confusion matrix* ditampilkan pada Tabel di bawah.

Tabel 1. *Confusion Matrix*

	<i>Actual Positif (1)</i>	<i>Actual Negatif (0)</i>
<i>Prediksi Positif (1)</i>	TP	FP
<i>Prediksi Negatif (0)</i>	FN	TN

3. METODE PENELITIAN

Penelitian ini melibatkan proses klasifikasi sentimen dari pembahasan yang sudah dijadikan satu, yang kemudian dijalankan menggunakan *Algoritma Naive Bayes* untuk menghasilkan tingkat akurasi sentimen. Pengolahan data dilakukan melalui *platform Google Colab*, yang digunakan untuk mengakses data ulasan.



Gambar 1. Tahapan Penelitian

Berikut ini merupakan tahapan penelitiannya:

1) Mengumpulkan Data

Tahapan pertama dimulai dengan pengumpulan data; pada data ini dikumpulkan dari *server Google*. Proses pengambilan data ini dibantu oleh *Google Colab*, alat yang memungkinkan pengumpulan data dalam jumlah besar. Studi ini menggunakan dataset ulasan aplikasi *Threads*, yang diperoleh dari *Play Store*.

2) Pelabelan dilakukan secara Manual

Proses pelabelan data dilakukan secara manual karena data yang diperoleh belum dilengkapi dengan nilai sentimen negatif atau positif, proses pelabelan data bisa dilakukan secara manual untuk menentukan apakah data tersebut termasuk dalam kategori sentimen negatif atau positif.

3) Praproses

Praproses adalah proses selanjutnya dalam analisis sentiment yang dilakukan secara manual, yang artinya, menghapus kesalahan dan ketidaksesuaian dari data belum jadi agar data menjadi lebih benar dan siap untuk identifikasi. Ini adalah proses tindakan yang dilakukan dalam tahap ini:

a. Case Folding

Proses ini mengubah semua karakter dalam teks menjadi huruf kecil. Ini berfungsi untuk mempermudah proses pengolahan dan pencarian teks, mengingat ada beberapa teks yang tidak konsisten saat menggunakan huruf kapital (3).

b. Pembersihan

Pada tahap ini merupakan proses menghilangkan informasi yang tidak penting atau tidak berarti, serta untuk meringkas data teks. Tahap ini bertujuan untuk memastikan data yang akan diolah memiliki kualitas yang lebih baik dengan format yang sesuai (11).

c. Tokenization

Tokenization adalah proses memecah string teks menjadi kata-kata yang terpisah. Secara umum, proses ini memisahkan beberapa karakter pada teks menjadi unit kata dengan memanfaatkan pemisah tertentu(9).

d. Normalisasi

Langkah ini bertujuan menyesuaikan kata-kata dalam teks agar sesuai dengan kaidah bahasa Indonesia yang baku. Normalisasi menjamin kata-kata yang digunakan memiliki makna yang sesuai dengan KBBI(9).

e. Stopword Removal

Tahap ini membersihkan teks dari kata-kata yang tidak mempunyai makna khusus, seperti kata penghubung atau kata pengganti. Meski sering digunakan, kata-kata ini tidak memberikan informasi penting dalam analisis (4).

f. Stemming

Stemming adalah langkah untuk menghapus afiks dari kata sehingga diperoleh bentuk dasar. Hasil stemming tetap mematuhi aturan yang berlaku, karena proses ini memanfaatkan pustaka Sastrawi yang telah di sesuaikan dengan kaidah Bahasa Indonesia yang baik dan benar. (7).

4) Pemecahan Data Latih dan Data Uji.

Data set dipisahkan menjadi dua segmen, yaitu Data Latih dan Data Uji, yang merupakan proses memisahkan dataset menjadi dua segmen terpisah. Satu segmen digunakan untuk melatih model pembelajaran mesin (data latih), sementara segmen lainnya dipakai untuk menguji kinerja *prototype* yang sudah diuji (data uji).

5) Pembobotan TF-IDF

Algoritma TF-IDF (*Term Frequency-Inverse Document Frequency*) merupakan metode statistik yang dipakai untuk menilai sejauh mana suatu istilah berkontribusi dalam sebuah dokumen dalam tumpukan dokumen. Sasaran utama dari algoritma ini yaitu untuk memberikan nilai pada keyword berdasarkan frekuensi kemunculannya dalam suatu dokumen (*Term Frequency*) dan sejauh mana kata tersebut bersifat unik dalam keseluruhan kumpulan dokumen (*Inverse Document Frequency*).

6) Pengelompokan

Pada tahap pengelompokan, proses dilakukan dengan membangun model pembelajaran mesin memakai 80% data latih dan 20% data sample yang diambil dilakukan secara random dari seluruh dataset untuk melakukan validasi silang dan menciptakan nilai perkiraan guna mengukur tingkat akurasi.

4. HASIL DAN PEMBAHASAN

4.1. Collecting Data

Dalam penelitian ini, proses pengumpulan data dilakukan dengan mengambil ulasan dari pengguna *Threads* yang terdapat di *Google Play Store*, menggunakan *Google Colab*. Proses pengambilan data, yang juga dikenal sebagai "*scraping*", menghasilkan 1000 ulasan yang diperoleh dari *Google Play Store* dan diambil dari 7 Desember sampai 10 Desember 2024 untuk *web scraping* penulis menggunakan *Google Colaboratory* dengan Bahasa pemrograman *Python*.

```

from google_play_scraper import Sort, reviews

result, continuation_token = reviews(
    'com.instagram.barcelona',
    lang='id', # defaults to 'en'
    country='id', # defaults to 'us'
    sort=Sort.MOST_RELEVANT, # defaults to Sort.MOST_RELEVANT
    count=1000, # defaults to 100
    filter_score_with=None # defaults to None(means all score)
)

```

Gambar 2. Proses Scraping

Gambar di atas terdapat script sebagai berikut :

from google_play_scraper import sort, reviews: Mengimpor modul *sort* dan *reviews* dari library *google_play_scraper*, fungsi *reviews* dipanggil untuk mengambil komentar. Hasil pengambilan data disimpan dalam variabel *hasil*, dan jika terdapat komentar tambahan yang dapat diambil, informasi token lanjutan disimpan dalam variabel *token_lanjutan*. Dataset yang telah dikumpulkan kemudian disimpan dan diubah formatnya menjadi file CSV dengan jumlah komentar yang sesuai dengan yang diharapkan.

Tabel 2. Content dan Sentimen

Content	Sentimen
Jelek	Negatif
Keren	Positif

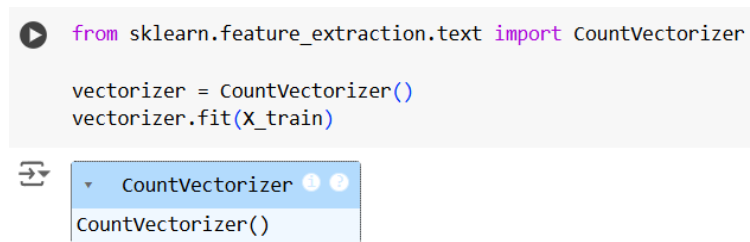
4.2. Praproses

Tabel 3. Preprocessing dan Content

Praproses	Konten
Data Mentah	sangat bermanfaat sekali
Case Folding	sangat bermanfaat sekali
Stopword	sangat bermanfaat sekali
Tokenisasi	sangat bermanfaat sekali
Pengembalian Kata Dasar	sangat bermanfaat sekali

4.3. Pembobotan

Algoritma *TF-IDF* (*Term Frequency-Inverse Document Frequency*) merupakan metode statistik yang berfungsi untuk mengukur tingkat kepentingan sebuah kata dalam sebuah dokumen di dalam suatu kumpulan dokumen atau korpus. Algoritma ini bertujuan memberikan nilai kepada setiap kata berdasarkan frekuensi kemunculannya dalam dokumen tertentu (*Term Frequency*) dan keunikannya dibandingkan dengan keseluruhan kumpulan dokumen (*Inverse Document Frequency*).



```
from sklearn.feature_extraction.text import CountVectorizer

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

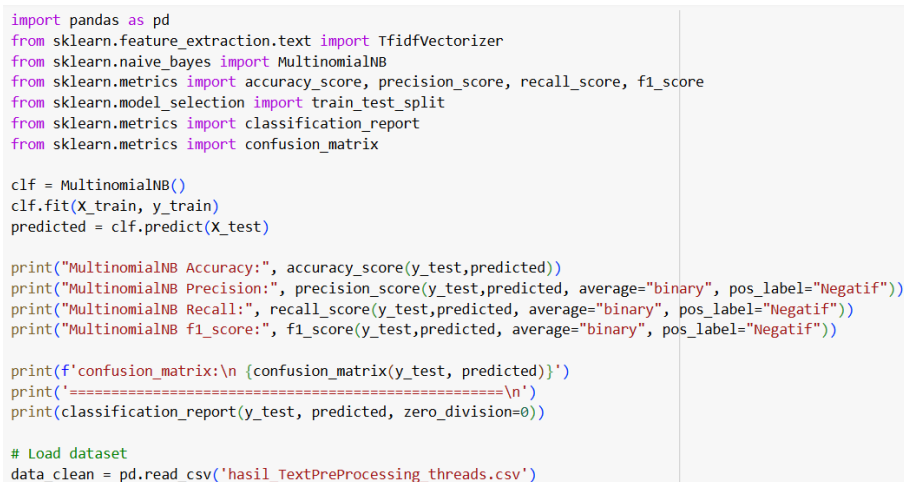
CountVectorizer

CountVectorizer()

Gambar 3. Konversi Teks Menjadi Vector

4.4. Naive Bayes Classifier

Sebagai data pengujian, 20% yang digunakan, sementara 80% data digunakan sebagai data pelatihan dalam proses pembuatan model pembelajaran mesin yang dilakukan pada tahap ini. Untuk *cross-validation* dan memperoleh nilai prediksi yang tepat, pemilihan data acak dilakukan pada data set. Tahap klasifikasi digambarkan di sini menggunakan *code algoritma Naive Bayes di Google Collabs*.



```
import pandas as pd
from sklearn.feature_extraction.text import TfidfVectorizer
from sklearn.naive_bayes import MultinomialNB
from sklearn.metrics import accuracy_score, precision_score, recall_score, f1_score
from sklearn.model_selection import train_test_split
from sklearn.metrics import classification_report
from sklearn.metrics import confusion_matrix

clf = MultinomialNB()
clf.fit(X_train, y_train)
predicted = clf.predict(X_test)

print("MultinomialNB Accuracy:", accuracy_score(y_test, predicted))
print("MultinomialNB Precision:", precision_score(y_test, predicted, average="binary", pos_label="Negatif"))
print("MultinomialNB Recall:", recall_score(y_test, predicted, average="binary", pos_label="Negatif"))
print("MultinomialNB f1_score:", f1_score(y_test, predicted, average="binary", pos_label="Negatif"))

print(f'confusion_matrix:\n {confusion_matrix(y_test, predicted)}')
print('=====')
print(classification_report(y_test, predicted, zero_division=0))

# Load dataset
data_clean = pd.read_csv('hasil_TextPreProcessing_threads.csv')
```

Gambar 4. Tahapan Klasifikasi Naive Bayes

Script Python pada Gambar diatas. digunakan untuk melakukan training model klasifikasi *Naive Bayes* dengan memanfaatkan representasi fitur teks menggunakan metode TF-IDF.

4.5. Visualisasi kata

Gambar 5 menunjukkan *keyword* yang paling sering muncul pada komentar *thread* dalam penelitian ini: "produk", "dan", "sangat", "aplikasi", "buruk". Jika kata yang paling besar ukuran kata dalam kotak tersebut, semakin banyak kata yang ditemukan.



Gambar 5. Visualisasi Kata Positif dan Negatif

5. KESIMPULAN

Hasil yang di dapat dalam penelitian ini menunjukkan bahwa pengguna aplikasi *Threads* cenderung memberikan ulasan positif. Pengolahan ulasan dengan *algoritma Naive Bayes* menghasilkan tingkat akurasi sebesar 84%. Pada analisis sentimen negatif, diperoleh precision sebesar 74%, *recall* 65%, dan *f1-score* 69%. Sementara itu, untuk sentimen positif, nilai precision tercatat sebesar 84%, *recall* 89%, dan *f1-score* 86%. Meskipun demikian, terdapat kemungkinan yang cukup besar untuk kesalahan dalam mengidentifikasi ulasan negatif, yang mungkin disebabkan oleh penggunaan bahasa yang tidak jelas (ambigu) atau adanya berbagai sentiment dalam satu ulasan. Hasil analisis sentimen menunjukkan bahwa mayoritas ulasan bersentimen positif, sehingga menunjukkan kepuasan pengguna terhadap fitur aplikasi *Threads*. Untuk pengembangan penelitian selanjutnya, disarankan untuk membandingkan klasifikasi menggunakan kernel lain seperti RBF, sigmoid, atau polinomial, serta mencoba algoritma lain seperti *k-Nearest Neighbour*. Selain itu, metode *k-fold cross validation* dapat diterapkan agar bisa membandingkan tingkat akurasi percobaan yang berbeda.

DAFTAR PUSTAKA

1. Dwianto, E., & Sadikin, M. (2021). Analisis Sentimen Transportasi Online pada Twitter Menggunakan Metode Klasifikasi Naive Bayes dan Support Vector Machine. *Format: Jurnal Ilmiah Teknik Informatika*, 10(1), 94. <https://doi.org/10.22441/format.2021.v10.i1.009>
2. Fransiska, S., Rianto, R., & Gufroni, A. I. (2020). Sentiment Analysis Provider By.U on Google Play Store Reviews with TF-IDF and Support Vector Machine (SVM) Method. *Scientific Journal of Informatics*, 7(2), 203–212. <https://journal.unnes.ac.id/nju/sji/article/view/25596>
3. Friska Aditia Indriyani, Ahmad Fauzi, & Sutan Faisal. (2023). Analisis sentimen aplikasi tiktok menggunakan algoritma Naive bayes dan support vector machine. *TEKNOSAINS: Jurnal Sains, Teknologi Dan Informatika*, 10(2), 176–184. <https://doi.org/10.37373/tekno.v10i2.419>
4. Hendriyanto, M. D., Ridha, A. A., & Enri, U. (2022). Analisis Sentimen Ulasan Aplikasi Mola Pada Google Play Store Menggunakan Algoritma Support Vector Machine. *INTECOMS: Journal of Information Technology and Computer Science*, 5(1), 1–7. <https://doi.org/10.31539/intecom.v5i1.3708>
5. Insan, M. K., Hayati, U., & Nurdiawan, O. (2023). Analisis Sentimen Aplikasi Brimo Pada Ulasan Pengguna Di. *Jurnal Mahasiswa Teknik Informatika*, 7(1), 478–483.
6. M Walid, F. H. (2022). Klasifikasi Kemandirian Siswa SMA/MA Double Track Menggunakan Metode Naive Bayes. *Jurnal ICT: Information Communication & Technology*, 2022.
7. Maulana, R., Voutama, A., & Ridwan, T. (2023). Analisis Sentimen Ulasan Aplikasi MyPertamina pada Google Play Store menggunakan Algoritma NBC. *Jurnal Teknologi Terpadu*, 9(1), 42–48. <https://doi.org/10.54914/jtt.v9i1.609>
8. Muzaki, A., & Witanti, A. (2021). Sentiment Analysis of the Community in the Twitter To the 2020 Election in Pandemic Covid-19 By Method Naive Bayes Classifier. *Jurnal Teknik Informatika (Jutif)*, 2(2), 101–107. <https://doi.org/10.20884/1.jutif.2021.2.2.51>
9. Ramadhan, R. (2020). Analisis Sentimen Pada Ulasan Aplikasi Di Google Play Store Dengan K-Nearest Neighbor. *Jurnal Ekonomi Volume 18, Nomor 1 Maret 201*, 2(1), 41–49.
10. Romadloni, N. T., Santoso, I., & Budilaksono, S. (2019). Perbandingan Metode Naive Bayes, Knn Dan Decision Tree Terhadap Analisis Sentimen Transportasi Krl Commuter Line. *Jurnal IKRA-ITH Informatika: Jurnal Komputer Dan Informatika*, 3(2), 1–9.
11. Safryda Putri, D., & Ridwan, T. (2023). Analisis Sentimen Ulasan Aplikasi Pospay Dengan Algoritma Support Vector Machine. *Jurnal Ilmiah Informatika*, 11(01), 32–40. <https://doi.org/10.33884/jif.v11i01.6611>
12. Sanubari, F. D., Enri, U., & Susilawati. (2023). Analisis Sentimen Terhadap Perubahan Rute Krl Commuter Jabodetabek Menggunakan Algoritme Support Vector Machine (Svm). *Jurnal Ilmiah Wahana Pendidikan*, 2023(15), 155–163. <https://doi.org/10.5281/zenodo.8206986>