

EFISIENSI ANGGARAN DALAM WACANA PUBLIK: ANALISIS SENTIMEN PLATFORM X DENGAN NAÏVE BAYES

Rizky Ananda Hafika, Alvin Hafiz, Stefen Agus Waruwu,
Muhammad Yazid Noor, Yuda Advis Ambrosius Sitohang, Kana Saputra

Ilmu Komputer, Universitas Negeri Medan
Jl. William Iskandar Ps. V, Kenangan Baru, Kec. Percut Sei Tuan
Kabupaten Deli Serdang, Sumatera Utara, Indonesia
rizkyanandahafika@gmail.com

ABSTRAK

Efisiensi anggaran merupakan isu krusial dalam pengelolaan keuangan negara yang sering menjadi sorotan publik, khususnya di media sosial seperti Twitter. Opini masyarakat yang tersebar di platform tersebut mencerminkan tingkat kepercayaan terhadap keterbukaan dan akuntabilitas pemerintah. Namun, data opini yang bersifat tidak terstruktur menimbulkan tantangan dalam analisis sentimen publik secara menyeluruh. Penelitian ini bertujuan untuk mengevaluasi opini masyarakat mengenai efisiensi anggaran menggunakan algoritma Naïve Bayes. Sebanyak 1.610 tweet dikumpulkan dan diproses melalui tahap preprocessing yang meliputi pembersihan data, case folding, tokenisasi, normalisasi, penghapusan stopword, dan stemming. Setelah preprocessing, data diberi label secara manual untuk klasifikasi sentimen. Selanjutnya, dilakukan ekstraksi fitur menggunakan metode TF-IDF dan pembagian data menjadi 80% untuk pelatihan dan 20% untuk pengujian model. Hasil analisis menunjukkan bahwa opini publik cenderung negatif dengan distribusi 58,4% negatif, 23,5% netral, dan 18,1% positif. Model klasifikasi menghasilkan akurasi sebesar 73,29%, dengan nilai F1-score tertinggi pada sentimen negatif (0,82) dan terendah pada sentimen netral (0,06), yang mengindikasikan adanya ketidakseimbangan data. Penelitian ini menyimpulkan bahwa algoritma Naïve Bayes cukup efektif dalam mengidentifikasi sentimen negatif, namun memerlukan perbaikan dalam mendeteksi sentimen netral, antara lain melalui teknik penyeimbangan data dan eksplorasi algoritma lain di masa mendatang.

Kata kunci : Analisis Sentimen, Efisiensi Anggaran, Naïve Bayes, TF-IDF, Twitter

1. PENDAHULUAN

Efisiensi adalah hubungan antara sumber daya yang digunakan dalam proses produksi dan hasil yang dihasilkan, baik berupa barang maupun jasa. Secara sistematis, efisiensi dapat dihitung sebagai kuantitas output per unit input, atau sebagai perbandingan input dan output. Mirip dengan produktivitas, efisiensi adalah kemampuan suatu program, organisasi, atau kegiatan untuk menghasilkan output maksimum dengan sumber daya yang paling sedikit [1].

Tingkat kepercayaan publik terhadap akuntabilitas dan keterbukaan pemerintah dalam mengelola keuangan negara sering tercermin dalam diskusi tentang efisiensi anggaran. Meningkatnya jumlah data opini di situs media sosial seperti Platform X menjadikan analisis sentimen sebagai alat yang berguna untuk mengetahui bagaimana perasaan masyarakat umum terhadap kebijakan anggaran.

Sebelumnya dikenal sebagai Twitter, Platform X adalah platform media sosial global yang memungkinkan pengguna untuk bebas berbagi pemikiran mereka. Wawasan mendalam dapat diperoleh dari analisis data yang dihasilkan di X dengan menggunakan teknik penambangan data. Di sini, mudah untuk melacak opini publik mengenai peristiwa atau tren terkini [2].

Analisis sentimen, sebuah teknik untuk menemukan dan menginterpretasikan sentimen atau opini dalam teks, digunakan untuk memahami opini pengguna terhadap sebuah aplikasi. Karena analisis

sentimen dapat mengungkapkan seberapa puas atau tidak puasnya pengguna terhadap sebuah program, evaluasi pengguna merupakan sumber informasi yang sangat baik [3].

Tujuan dari penelitian ini adalah untuk menggunakan algoritma Naïve Bayes, sebuah teknik berbasis probabilitas yang efektif dalam kategorisasi teks, untuk mengkaji opini publik terkait efisiensi anggaran. Dengan mengidentifikasi opini publik, apakah itu netral, negatif, atau positif, pendekatan ini akan mampu menjelaskan lebih lanjut tentang bagaimana perasaan masyarakat umum tentang kebijakan ekonomi pemerintah. Diharapkan bahwa temuan penelitian ini akan membantu para pembuat kebijakan lebih memahami dinamika opini publik dan meningkatkan taktik komunikasi dan transparansi anggaran di masa mendatang.

2. TINJAUAN PUSTAKA

2.1. Penelitian Terdahulu

Sebuah penelitian sebelumnya oleh Muhammad Irawan Adi Hartono, Feni Rosalia, dan Tabah Maryanah menggunakan metodologi deskriptif kualitatif dan kuantitatif untuk menyelidiki dampak teknologi digital terhadap efisiensi anggaran. Metode ini juga dianggap dapat meningkatkan akuntabilitas dan transparansi serta mempercepat proses pengadaan barang dan jasa. Studi ini juga mengidentifikasi sejumlah kesulitan, seperti periode implementasi yang singkat, ketersediaan produk yang terbatas dalam e-

katalog, dan kurangnya pemasok di luar Pulau Jawa [4].

Selain itu, sebuah penelitian oleh Adi Kusuman dan Agung Nugroho menunjukkan kemampuan analisis sentimen berbasis teks dari platform media sosial seperti Twitter dalam mengukur opini populer. Metode Naive Bayes sangat populer karena kemudahan penggunaan dan tingkat akurasi klasifikasi sentimen yang tinggi. Selain itu, yang juga penting adalah prosedur prapemrosesan data termasuk pelipatan huruf, tokenizing, penyaringan, dan stemming [5].

2.2. Efisiensi Anggaran dalam Wacana Publik

Efisiensi adalah perbandingan antara jumlah sumber daya atau masukan yang digunakan dalam suatu kegiatan dan jumlah keluaran, seperti barang dan jasa, yang dihasilkan. Organisasi sektor publik biasanya dievaluasi berdasarkan sejumlah kriteria terkait efisiensi. Tingkat efisiensi yang dicapai meningkat seiring dengan rasio masukan terhadap keluaran [6]. Gagasan ini dapat dibandingkan dengan produktivitas, yang menyatakan bahwa suatu organisasi, program, atau kegiatan dikatakan efisien jika dapat memaksimalkan produksi sekaligus memanfaatkan sumber daya yang tersedia dengan sebaik-baiknya. Efisiensi dalam pengelolaan anggaran mengacu pada penggunaan dana yang bijaksana dan tepat sasaran [7].

Sementara itu, baik di tingkat daerah maupun di dalam instansi pemerintah, anggaran memainkan peran penting dalam perencanaan dan pengendalian [8]. Manajemen dapat memperkirakan sasaran pengeluaran agar sesuai dengan rencana yang ditetapkan dengan memeriksa anggaran. Lebih jauh lagi, anggaran dapat digunakan sebagai alat peramalan untuk memproyeksikan keadaan keuangan di masa mendatang [9].

Sebagai alat untuk merencanakan dan mengelola keuangan pemerintah, anggaran publik sangatlah penting. Anggaran membantu dalam menetapkan tujuan yang harus dicapai dalam perencanaan, dan menjamin bahwa uang publik yang disahkan oleh legislatif didistribusikan dan digunakan sesuai dengan tujuan yang telah ditetapkan dalam pengelolaan [10].

2.3. Analisis Sentimen

Salah satu metode pemrosesan bahasa alami untuk menentukan pikiran dan perasaan seseorang tentang subjek atau objek tertentu adalah analisis sentimen. Data dianalisis dalam prosedur ini untuk menentukan kecenderungan opini terhadap suatu objek, yang kemudian dikategorikan sebagai positif atau negatif [11]. Ketika kumpulan dokumen teks dengan opini tentang suatu objek disediakan, penambangan opini dapat digunakan untuk mengidentifikasi apakah opini tersebut netral, negatif, atau positif dengan mengekstraksi karakteristik dan elemen dari objek yang dibahas dalam setiap dokumen. Lebih jauh, studi ulasan dan pengelolaan

subjektivitas, sentimen, dan opini dalam suatu teks juga dikaitkan dengan penambangan opini [12].

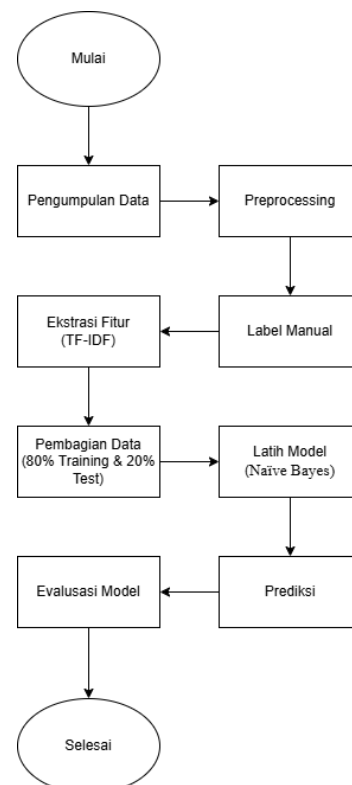
2.4. Platform X atau Twitter

Twitter adalah situs jejaring sosial terkenal yang memungkinkan pengguna berbagi dan me-retweet apa pun hingga sepanjang 140 karakter. Sekitar 300 juta orang telah mendaftar sebagai pengguna Twitter sejauh ini, dan lebih dari 500 juta akun baru diikuti setiap hari [13]. Twitter memudahkan komunikasi dan interaksi antar penggunanya. Selain itu, pengguna dapat mengakses tweet populer dengan lebih cepat dan mudah melalui fungsi Top Trending pada platform ini [14].

2.5. Naïve Bayes

Teknik klasifikasi yang dikenal sebagai Naïve Bayes didasarkan pada Teorema Bayes. Metode ini mengklasifikasikan data menggunakan konsep statistik dan probabilitas. Dengan menggunakan pola atau pengalaman dari data masa lalu, Naïve Bayes memperkirakan probabilitas kejadian di masa mendatang [15]. Kategori pembelajaran terawasi mencakup algoritma Naïve Bayes, yang menggunakan data pelatihan awal sebagai panduan saat membuat keputusan atau mengklasifikasikan data [16]. Data skala ordinal cocok untuk algoritma Naïve Bayes. Meskipun variabel dalam data ordinal ini dapat diurutkan dan memiliki nilai yang direpresentasikan oleh simbol, jaraknya tidak dapat diukur atau dijumlahkan [17].

3. METODE PENELITIAN



Gambar 1. Alur penelitian

Untuk mengklasifikasikan sentimen dalam wacana publik mengenai efisiensi anggaran, model yang digunakan dalam penelitian ini menggunakan pendekatan kuantitatif dan teknik analisis sentimen berbasis Naïve Bayes. Data dikumpulkan dari platform X (sebelumnya Twitter) dengan menggunakan teknik crawling.

3.1. Pengumpulan Data

Prosedur web scraping yang menggunakan strategi perayapan (crawling) terhadap postingan di platform X dengan kata kunci “efisiensi anggaran” digunakan untuk mendapatkan data. Sebanyak 1.610 titik data dengan 15 properti berhasil dikumpulkan. Namun, data utama untuk analisis sentimen dalam penelitian ini terbatas pada kolom “full_text”. Teks lengkap dari setiap tweet terdapat pada kolom ini dan akan melalui proses preprocessing tambahan. Untuk merampingkan data dan meningkatkan kinerja pemrosesan, atribut yang tidak berhubungan dengan analisis sentimen akan dihapus sebelum tahap persiapan dimulai.

3.2. Tahap Preprocessing

Agar data siap untuk diproses oleh model Naïve Bayes, dilakukan tahap preprocessing yang meliputi:

- Cleansing: Membersihkan teks dengan menghapus karakter khusus, angka, tanda baca, serta URL yang tidak diperlukan.
- Case Folding: Untuk mempertahankan keseragaman format, semua teks diubah menjadi huruf kecil.
- Tokenisasi: Memisahkan teks menjadi unit kata yang lebih kecil untuk mempermudah analisis.
- Normalisasi: Mengonversi kata tidak baku atau singkatan ke dalam bentuk yang lebih standar agar lebih mudah dipahami.
- Stopword Removal: Menyingkirkan kata-kata umum yang tidak memiliki nilai penting dalam analisis sentimen.
- Stemming: Mengembalikan kata ke bentuk dasarnya guna menyederhanakan struktur teks dan mengurangi variasi kata.

3.3. Label Manual

Data yang telah terkumpul selanjutnya diberi label secara manual. Setiap cuitan dikategorikan ke dalam sentimen positif, negatif, atau netral berdasarkan makna yang terkandung dalam teks. Labeling manual dilakukan oleh peneliti untuk memastikan akurasi dalam proses pelabelan.

3.4. Ekstraksi Fitur (TF-IDF)

Setelah proses preprocessing selesai, fitur teks diekstraksi menggunakan metode TF-IDF. Metode ini mengonversi teks ke dalam bentuk numerik sehingga dapat digunakan dalam model pembelajaran mesin.

3.5. Pembagian data

Data yang telah melalui proses pengolahan kemudian dibagi menjadi dua bagian utama:

- 80% Sebagai data training
- 20% Sebagai data testing

Pembagian ini bertujuan untuk melatih model secara optimal sebelum dilakukan evaluasi.

3.6. Latih Model

Model yang digunakan dalam penelitian ini adalah Naïve Bayes, yang merupakan salah satu metode klasifikasi berbasis probabilitas. Model ini dilatih menggunakan data training untuk memahami pola sentimen dalam teks.

3.7. Prediksi

Setelah model Naïve Bayes selesai dilatih dengan data training, tahap berikutnya adalah melakukan prediksi terhadap data uji. Pada proses ini, model yang telah dipelajari diterapkan untuk mengklasifikasikan sentimen dari teks yang belum pernah dianalisis sebelumnya.

3.8. Evaluasi Model

Setelah pelatihan, tahap evaluasi model menilai seberapa baik kinerja model dalam mengidentifikasi data secara akurat. Dengan membandingkan hasil proyeksi model dengan label nyata dari data uji, evaluasi dilakukan dengan menggunakan ukuran-ukuran termasuk akurasi, presisi, recall, dan skor F1. Recall menunjukkan kapasitas model untuk mengidentifikasi semua sampel positif, presisi menunjukkan proporsi prediksi yang akurat, dan akurasi menunjukkan persentase prediksi yang benar. Tujuan evaluasi ini adalah untuk memastikan bahwa model beroperasi sebagaimana mestinya. Tindakan korektif, termasuk prapemrosesan data yang lebih baik atau penyesuaian parameter, dapat dilakukan untuk meningkatkan akurasi model jika hasil evaluasi tidak memadai.

4. HASIL DAN PEMBAHASAN

4.1. Penjelasan Singkat Mengenai Data

Kata kunci “efisiensi anggaran” dalam bahasa Indonesia digunakan dalam proses crawling dari platform Twitter (Platform X) untuk mengumpulkan data untuk penelitian ini. Crawling dilakukan dengan menggunakan alat tweet-harvest yang memungkinkan pengumpulan data tweet berdasarkan kata kunci tertentu. Data yang berhasil dikumpulkan berjumlah 1610 tweet yang diambil dalam periode waktu terkini dengan pengaturan pencarian di tab “LATEST.”

Dataset yang diperoleh ini terdiri dari berbagai atribut, termasuk informasi terkait waktu pengiriman tweet, ID pengguna, lokasi, dan sebagainya. Namun, hanya satu kolom yang digunakan untuk menganalisis dalam penelitian ini: kolom full_text, yang berisi teks penuh dari tweet-tweet yang diposting oleh pengguna Twitter. Kolom ini menjadi fokus utama analisis

sentimen karena berisi tanggapan publik dan pendapat tentang "efisiensi anggaran".

Setelah pengumpulan data, data mentah yang tersedia kemudian diproses lebih lanjut melalui beberapa langkah preprocessing sebelum digunakan untuk analisis sentimen.

4.2. Pre-processing

Sebelum melakukan analisis sentimen, data yang dikumpulkan dari Twitter memerlukan beberapa langkah preprocessing untuk memastikan bahwa data yang digunakan dalam model klasifikasi memiliki kualitas yang baik. Langkah-langkah ini termasuk membersihkan, memotong casefold, tokenisasi, normalisasi, menghapus stopword, dan stemming.

```

                                full_text \
0 makan gratis asbunnya si gendut yg akhirnya ba...
1 @sweetchese30 @starfess Maap ya emang YG lagi ...
2 @agusqwerty234 @halimah_ir16969 @Dandhy_Lakson...

                                clean text \
0 makan gratis asbunnya si gendut yg akhirnya ba...
1 Maap ya emang YG lagi efisiensi Itu Doyoung Ju...
2 Menurut saya yg parah sih Danantara amp efisie...

                                case_folding \
0 makan gratis asbunnya si gendut yg akhirnya ba...
1 maap ya emang yg lagi efisiensi itu doyoung ju...
2 menurut saya yg parah sih danantara amp efisie...

                                tokenized \
0 [makan, gratis, asbunnya, si, gendut, yg, akhi...
1 [maap, ya, emang, yg, lagi, efisiensi, itu, do...
2 [menurut, saya, yg, parah, sih, danantara, amp...

                                normalized \
0 [makan, gratis, asbunnya, si, gendut, yang, ak...
1 [maap, ya, emang, yang, lagi, efisiensi, itu, ...
2 [menurut, saya, yang, parah, sih, danantara, a...

                                stopword_removed \
0 [makan, gratis, asbunnya, si, gendut, efisiensi...
1 [maap, ya, emang, efisiensi, doyoung, junghwan]
2 [parah, sih, danantara, amp, efisiensi, danant...

                                stemmed
0 [makan, gratis, asbunnya, si, gendut, efisiensi...
1 [maap, ya, emang, efisiensi, doyoung, junghwan]
2 [parah, sih, danantara, amp, efisiensi, danant...

```

Gambar 2. Proses pre-processing teks dalam analisis sentiment

Setelah langkah-langkah pre-processing ini, teks yang telah dibersihkan dan disiapkan kemudian digunakan dalam tahap selanjutnya, yaitu proses labeling dan ekstraksi fitur.

4.3. Labeling

Setelah proses preprocessing selesai, langkah selanjutnya adalah labeling, yaitu memberi label pada setiap tweet berdasarkan sentimen yang terkandung dalam teks. Labeling ini dilakukan secara manual, di mana setiap tweet dianalisis untuk menentukan apakah sentimennya bersifat positif, netral, atau negatif.

Pemberian label dilakukan dengan membaca setiap tweet secara cermat dan mengidentifikasi apakah isi tweet tersebut menunjukkan pandangan

yang positif, negatif, atau netral terhadap topik yang dibahas, yaitu efisiensi anggaran.

- Label Positif (1) diberikan untuk tweet yang mengungkapkan pendapat atau opini yang mendukung atau memberikan pandangan baik terhadap efisiensi anggaran.
- Label Negatif (-1) diberikan untuk tweet yang menunjukkan ketidaksetujuan atau kritik terhadap efisiensi anggaran, mengarah pada pandangan buruk atau negatif.
- Label Netral (0) diberikan untuk tweet yang tidak mengungkapkan pendapat yang jelas atau yang cenderung bersifat informatif tanpa ekspresi emosional yang jelas terkait efisiensi anggaran.

Pemberian label ini dilakukan dengan mempertimbangkan konteks kalimat dan kata-kata yang digunakan dalam tweet. Beberapa tweet mungkin memiliki nuansa ambigu atau campuran antara positif dan negatif, sehingga penilaiannya dilakukan dengan hati-hati dan berdasarkan interpretasi makna yang jelas dari keseluruhan isi tweet.

Dari total 1610 tweet yang dikumpulkan, label diberikan pada semua tweet tersebut. Proses labeling dilakukan secara manual, yang memerlukan ketelitian untuk memastikan bahwa setiap tweet diberi label dengan benar. Setelah proses labeling selesai, jumlah tweet yang telah dilabeli berdasarkan sentimen adalah sebagai berikut:

- Positif (1): 292 tweet
- Netral (0): 378 tweet
- Negatif (-1): 940 tweet

Jumlah tweet yang dilabeli ini akan digunakan untuk melatih model klasifikasi dan menjadi dasar untuk analisis sentimen lebih lanjut.

Proses labeling ini tidak tanpa tantangan. Beberapa tweet mungkin bersifat ambigu, seperti tweet yang menyampaikan informasi yang tidak jelas apakah itu mendukung atau menentang topik efisiensi anggaran. Dalam hal ini, keputusan untuk memberi label dilakukan berdasarkan interpretasi yang paling mendekati makna yang dimaksudkan oleh penulis tweet.

4.4. Ekstraksi Fitur (TF-IDF)

Pendekatan TF-IDF digunakan untuk ekstraksi fitur mengikuti prosedur pelabelan. Membuat representasi numerik dari teks untuk digunakan dalam model klasifikasi adalah tujuan dari teknik ini. Pendekatan ini menentukan tingkat kepentingan sebuah kata dalam sebuah dokumen dengan membandingkan frekuensi kemunculannya di dalam dokumen dengan tingkat kelangkaannya di seluruh korpus.

Hasil ekstraksi fitur TF-IDF menghasilkan dimensi matriks yang menunjukkan hubungan antara tweet dan kata-kata di dalamnya. Matriks yang dihasilkan memiliki ukuran 1610 x 4655, di mana: 1610 adalah jumlah tweet yang ada dalam dataset, dan

4655 adalah jumlah fitur (kata-kata unik) yang dihasilkan setelah proses ekstraksi fitur menggunakan TF-IDF.

```
Kata-kata dengan TF-IDF tertinggi:
asbunnya    0.519897
gendut      0.418231
gratis      0.364256
si          0.349173
program     0.326429
makan       0.309442
jalan       0.305080
efisiensi   0.069294
pager       0.000000
padu        0.000000
Name: 0, dtype: float64
```

Gambar 3. Daftar kata-kata dengan nilai TF-IDF tertinggi

Kata-kata dengan nilai TF-IDF tertinggi, seperti "asbunnya", "gendut", "gratis", dan "si", memiliki pengaruh lebih besar dalam representasi dokumen pertama dibandingkan dengan kata-kata lain. Kata-kata tersebut dianggap lebih signifikan dalam mendefinisikan makna dari tweet tertentu, karena memiliki frekuensi yang lebih tinggi dalam tweet tersebut atau lebih jarang muncul di seluruh dataset.

4.5. Pembagian Data

Setelah ekstraksi fitur TF-IDF selesai, dataset dipisahkan menjadi dua bagian: data pelatihan dan data uji. Tujuan dari bagian ini adalah untuk melatih model klasifikasi menggunakan sejumlah data tertentu dan mengevaluasi kinerja model menggunakan data yang tidak digunakan untuk pelatihan. Melalui bagian ini, kemampuan model untuk menggeneralisasi data baru juga dapat dinilai.

80% dari dataset digunakan untuk pelatihan, dan 20% sisanya digunakan untuk pengujian. Tujuan dari pembagian ini adalah untuk menyediakan data yang cukup untuk pelatihan model dan menjamin bahwa data pengujian cukup representatif untuk menilai kinerja model dengan cara yang tidak bias.

Jumlah data yang digunakan untuk pelatihan dan pengujian setelah pembagian adalah sebagai berikut:

- Jumlah Data Latih: 1288 tweet
- Jumlah Data Uji: 322 tweet

Model klasifikasi diajarkan menggunakan data pelatihan, dan akurasi serta kinerjanya dinilai menggunakan data uji setelah pelatihan.

4.6. Penerapan Naïve Bayes

Model klasifikasi Naïve Bayes kemudian digunakan untuk memeriksa sentimen tweet setelah data dipisahkan menjadi data latih dan data uji. Salah satu pendekatan yang paling populer untuk klasifikasi teks, khususnya untuk analisis sentimen, adalah model ini. Asumsi independensi fakta dalam model ini dimungkinkan oleh prinsip probabilitas Bayes.

```
from sklearn.naive_bayes import MultinomialNB

nb_model = MultinomialNB()

nb_model.fit(X_train, y_train)

print("Model Naïve Bayes berhasil dilatih!")
```

Gambar 4. Implementasi model Naïve Bayes menggunakan *scikit-learn*

Dalam aplikasi ini, model Multinomial Naïve Bayes dari SciKit-Learn digunakan. Kode ini melatih model dengan menggunakan data yang telah dikumpulkan dan diberi label sebelumnya. Metode `fit()` digunakan untuk melakukan penelitian ini. Metode ini mengajarkan model untuk mengajarkan hubungan antara fitur (TF-IDF) dan sentimen label (positif, negatif, atau netral).

Setelah proses pelatihan selesai, model Naïve Bayes digunakan untuk memprediksi sentimen pada data. Hasil dari proses pelatihan menunjukkan bahwa model telah berhasil dikembangkan dan digunakan untuk menganalisis data yang belum dianalisis.

4.7. Evaluasi Model

Setelah model Naïve Bayes dilatih, langkah selanjutnya adalah mengevaluasi kinerja model dengan menggunakan data uji yang sebelumnya dipisahkan. Seberapa baik model dapat memprediksi perubahan pada data yang belum pernah dilihat sebelumnya adalah tujuan dari evaluasi ini.

```
Akurasi Model: 0.7329
Laporan Klasifikasi:
```

	precision	recall	f1-score	support
-1	0.71	0.97	0.82	185
0	1.00	0.03	0.06	67
1	0.82	0.79	0.80	70
accuracy			0.73	322
macro avg	0.84	0.59	0.56	322
weighted avg	0.79	0.73	0.66	322

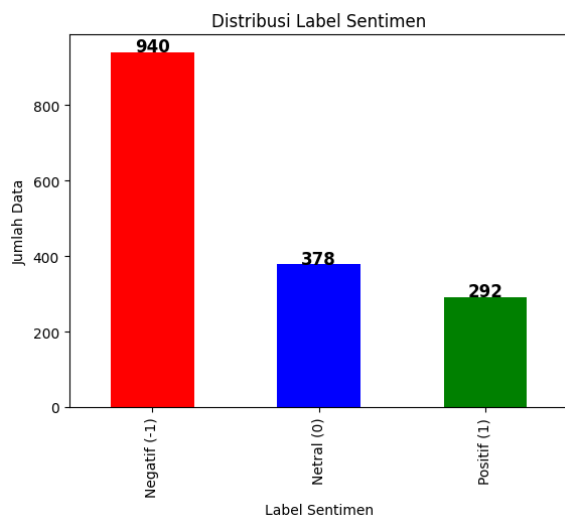
Gambar 5. Hasil evaluasi model Naïve Bayes

Dengan akurasi 73%, presisi 79%, recall 73%, dan skor F1 66%, model Naïve Bayes menunjukkan kemampuan yang cukup untuk mengklasifikasikan data secara keseluruhan, tetapi masih memiliki kelemahan dalam mengenali kelas tertentu, terutama kelas 0, yang memiliki recall sangat rendah (3%). Hal ini mengindikasikan bahwa model kesulitan mendeteksi kelas 0 dengan baik, meskipun mampu mengklasifikasikan kelas -1 dan 1 dengan cukup akurat.

4.8. Visualisasi

Untuk lebih memahami kinerja model Naïve Bayes dalam analisis sentimen terhadap tweet yang dikumpulkan, beberapa visualisasi telah dibuat. Visualisasi ini membantu dalam mengevaluasi performa model dan distribusi kelas sentimen dalam data yang digunakan. Berikut adalah penjelasan mengenai visualisasi yang digunakan:

4.9. Grafik Batang Labeling Manual



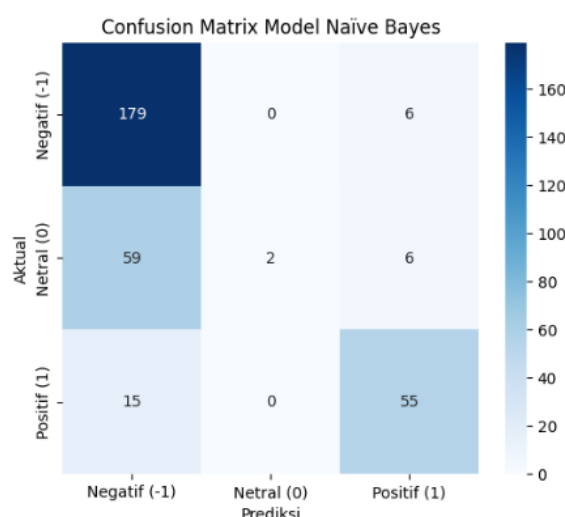
Gambar 6. Grafik distribusi label sentimen

Grafik distribusi kelas sentimen menunjukkan jumlah tweet berdasarkan sentimen yang telah dilabeli:

- Positif (1): 292 tweet
- Netral (0): 378 tweet
- Negatif (-1): 940 tweet

Dari grafik ini, terlihat bahwa kelas negatif mendominasi dataset, yang menciptakan ketidakseimbangan kelas. Ketidakseimbangan ini dapat memengaruhi kinerja model, karena model cenderung lebih sering memprediksi kelas yang lebih dominan (negatif), sementara kelas positif dan netral yang lebih sedikit mungkin kurang terprediksi dengan akurat.

4.10. Confusion Matrix

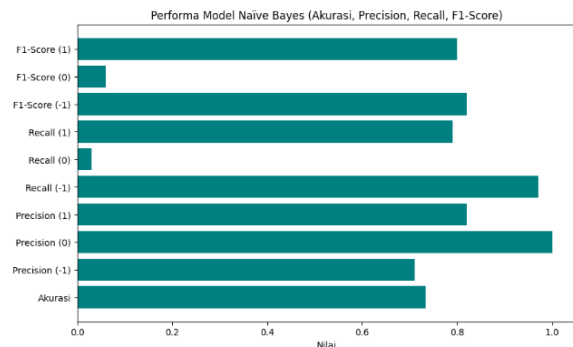


Gambar 7. Confusion matrix model Naïve Bayes

Menurut matrix kekacauan, model Naive Bayes menemukan kelas Negatif (-1) dengan 179 prediksi yang benar dari 185 sampel. Namun, model kesulitan

mengenali kelas Netral (0), karena hanya 2 dari 67 sampel yang diklasifikasikan dengan benar, sementara sebagian besar salah diklasifikasikan sebagai Negatif. Untuk kelas Positif (1), model cukup baik dengan 55 dari 70 sampel diprediksi benar, meskipun beberapa salah diklasifikasikan sebagai Negatif. Secara keseluruhan, model cenderung mengklasifikasikan data sebagai Negatif, sehingga perlu perbaikan terutama dalam mengenali kelas Netral.

4.11. Grafik Batang Performa Model



Gambar 8. Grafik performa model Naïve Bayes

Nilai untuk setiap kelas sentimen negatif (-1), netral (0), dan positif (1) ditunjukkan pada grafik batang ini untuk masing-masing metrik model Naive Bayes (Akurasi, Ketepatan, Recall, dan Skor F1).

Dari grafik, terlihat bahwa Precision untuk kelas Netral (0) sangat tinggi. Namun, nilai Recall dan F1 kelas ini sangat rendah, yang menunjukkan bahwa model hampir tidak dapat mengenali data Netral dengan baik. Sebaliknya, nilai Negatif (-1) memiliki nilai Recall yang sangat tinggi, yang menunjukkan bahwa model mampu menangkap hampir semua sampel Negatif, tetapi nilai Precision-nya sedikit lebih rendah. Kelas Positif (1) memiliki keseimbangan yang cukup baik antara Precision, Recall, dan F1-Score, tetapi tidak setinggi kelas Negatif.

Akurasi model cukup tinggi, tetapi ketidakseimbangan dalam performa antar kelas menunjukkan bahwa model lebih cenderung mengklasifikasikan data sebagai Negatif, sehingga perlu perbaikan terutama dalam mengenali kelas Netral.

5. KESIMPULAN DAN SARAN

Studi ini telah berhasil menganalisis sentimen publik terhadap efisiensi anggaran di platform Twitter, dan menemukan bahwa sentimen publik cenderung negatif, dengan distribusi 58,4% negatif, 23,5% netral, dan 18,1% positif. Ini menunjukkan banyak perspektif kritis tentang efisiensi anggaran. Dalam penerapannya, algoritma Naive Bayes digunakan untuk membuat model klasifikasi sentimen melalui proses analisis dan pelabelan manual. Selain itu, algoritma ini juga digunakan untuk mengekstraksi fitur menggunakan TF-IDF dan membangun model klasifikasi sentimen melalui proses preprocessing teks yang mencakup

perbaikan, case folding, tokenisasi, normalisasi, penghapusan kata tunggal, dan stemming. Data dibagi menjadi 20% untuk pelatihan dan 80% untuk pengujian sebelum model dilatih dan diuji. Evaluasi model menunjukkan akurasi sebesar 73,29%, di mana F1-score tertinggi terdapat pada kelas negatif (0,82), sementara kelas netral memiliki F1-score terendah (0,06) Karena ketidakseimbangan data, kategori negatif memiliki tingkat recall dan keakuratan yang lebih tinggi dibandingkan kategori lainnya, yang menunjukkan bahwa model lebih mampu mengidentifikasi sentimen negative. Berdasarkan hasil ini, penelitian merekomendasikan penanganan ketidakseimbangan data melalui teknik seperti oversampling atau undersampling untuk meningkatkan akurasi pada kategori positif dan netral, serta eksplorasi model lain seperti Support Vector Machine (SVM) atau penggunaan word embedding guna meningkatkan performa klasifikasi di masa depan.

DAFTAR PUSTAKA

- [1] R. A. Hapsari and Z. W. Asis, "Analisis Efektivitas dan Efisiensi Anggaran Corporate Social Responsibility," *J. Ekon. dan Bisnis*, vol. 2, no. 1, pp. 1–6, 2023, doi: 10.57151/jeko.v2i1.67.
- [2] C. Tertentu and D. I. Pilpres, "ANALISIS SENTIMEN MASYARAKAT DI PLATFORM X TERHADAP PENGGUNAAN BANSOS UNTUK MEMENANGKAN SALAH SATU," vol. 8, no. 5, pp. 9941–9947, 2024.
- [3] "Eskiyaturrofikoh" and R. R. 'Suryono, "Analisis Sentimen Aplikasi X Pada Google Play Store Menggunakan Algoritma Naïve Bayes Dan Support Vector Machine (Svm)," *JUPI(Jurnal Ilm. Penelit. dan Pembelajaran Inform.*, vol. 9, no. 3, pp. 1408–1419, 2024, [Online]. Available: <https://www.jurnal.stkippgritulungagung.ac.id/index.php/jupi/article/view/5392>
- [4] M. Irawan, A. Hartono, F. Rosalia, and T. Maryanah, "Implementasi E-Procurement Logistik Sebagai Efisiensi Anggaran Pada Pemilihan Kepala Daerah Kabupaten Pesawaran Tahun 2020," *J. MODERAT*, vol. 8, no. 3, pp. 500–516, 2022.
- [5] A. Kusuma and A. Nugroho, "Analisa Sentimen Pada Twitter Terhadap Kenaikan Tarif Dasar Listrik Dengan Metode Naïve Bayes," *J. Ilm. Teknol. Inf. Asia*, vol. 15, no. 2, p. 137, 2021, doi: 10.32815/jitika.v15i2.557.
- [6] Wakhid Yuliyanto, S. Wahyuningsih, R. Kurniasih, and A. Waluyo, "Pengukuran Kinerja Melalui Pendekatan 'Value For Money' Pada Pelaksanaan Anggaran Dinas 'X' di Sektor Publik," *J. E-Bis*, vol. 7, no. 1, pp. 233–245, 2023, doi: 10.37339/e-bis.v7i1.1183.
- [7] Y. Pratama and F. Pikri, "Efisiensi Dan Efektivitas Anggaran Belanja Pada Rumah Sakit Umum Daerah Cicalengka Kabupaten Bandung," *Minist. J. Birokrasi dan Pemerintah. Drh.*, vol. 2, no. 2, pp. 75–86, 2020, doi: 10.15575/jbpd.v2i2.9385.
- [8] S. Mohammad, D. P. E. Saerang, and H. Gamaliel, "Analisis Efektivitas Pendapatan dan Efisiensi Belanja Pada LRA BPKAD Kota Tidore Kepulauan," *J. LPPM Bid. EkoSosBudKum (Ekonomi, Sos. Budaya, dan Hukum)*, vol. 7, no. 3, pp. 241–250, 2023.
- [9] R. Anwar, Y. Yuniarsih, A. P. Depeda, E. C. Tambunan, and R. Tina, "Penggunaan Analisis Anggaran Sebagai Alat Perencanaan Dan Pengendalian Keuangan Dalam Perusahaan," *J. Educ. Lang. Res.*, vol. 1, no. 8, pp. 1083–1096, 2022.
- [10] Gt. Indriani Puspitasari, "Efisiensi Dan Efektivitas Realisasi Anggaran, Optimalisasi Dan Kinerja Keuangan," *Kindai*, vol. 18, no. 3, pp. 444–455, 2023, doi: 10.35972/kindai.v18i3.913.
- [11] O. Manullang and C. Prianto, "Analisis Sentimen dalam Memprediksi Hasil Pemilu Presiden dan Wakil Presiden : Systematic Literature Review," *J. Inform. Dan Teknol. Komput.*, vol. 4, no. 2, pp. 104–113, 2023, [Online]. Available: <https://ejurnalunsam.id/index.php/jicom/>
- [12] D. Alita and A. R. Isnain, "Pendeteksian Sarkasme pada Proses Analisis Sentimen Menggunakan Random Forest Classifier," *J. Komputasi*, vol. 8, no. 2, pp. 50–58, 2020, doi: 10.23960/komputasi.v8i2.2615.
- [13] M. Kholilullah, M. Martanto, and U. Hayati, "Analisis Sentimen Pengguna Twitter(X) Tentang Piala Dunia Usia 17 Menggunakan Metode Naive Bayes," *JATI (Jurnal Mhs. Tek. Inform.*, vol. 8, no. 1, pp. 392–398, 2024, doi: 10.36040/jati.v8i1.8378.
- [14] Rizqi Nandadita Pamungkas, D. Permadi, and I. D. Florina, "Strategi Humor Gibran Rakabuming dalam Komunikasi Politik di Media Sosial X (Twitter)," *J. Pemerintah. dan Polit.*, vol. 9, no. 3, pp. 175–182, 2024, doi: 10.36982/jpg.v9i3.4057.
- [15] B. Purba and R. Syahputra, "Implementasi metode Naive Bayes Classifier pada Evaluasi Kepuasan Mahasiswa terhadap Pembelajaran Daring," *InfoTekjar J. Nas. Inform. dan Teknol. Jar.*, vol. 6, no. 1, pp. 85–91, 2021, [Online]. Available: <https://doi.org/10.30743/infotekjar.v6i1.4352>
- [16] Meriyam Yunita & Tri Widodo, "Sistem Pakar Diagnosa Gangguan jiwa Menggunakan Metode Naïve Bayes Berbasis Web," *Informatika*, vol. 4, no. 1, pp. 166–174, 2021.
- [17] H. D. Wijaya and S. Dwiasnati, "Implementasi Data Mining dengan Algoritma Naïve Bayes pada Penjualan Obat," *J. Inform.*, vol. 7, no. 1, pp. 1–7, 2020, doi: 10.31311/ji.v7i1.6203.