

ANALISIS SENTIMEN DI MEDIA SOSIAL TERHADAP KEBIJAKAN EFISIENSI ANGGARAN MENGGUNAKAN METODE RANDOM FOREST

Bryan Ifan Etikamena, Karina Auliasari, Mira Orisa
Teknik Informatika, Institut Teknologi Nasional Malang
Jalan Raya Karanglo km 2 Malang, Indonesia
ifanetikamena09@gmail.com

ABSTRAK

Kebijakan efisiensi anggaran merupakan langkah yang diambil oleh pemerintah untuk mengatur dan mengelola keuangan negara secara lebih efektif dan hemat, dengan memangkas atau mengalihkan anggaran pada bidang tertentu agar penggunaan dana negara menjadi lebih optimal. Kebijakan ini menuai beragam opini dari masyarakat yang diungkapkan melalui media sosial. Untuk memahami opini masyarakat publik secara lebih sistematis dan menyeluruh, penelitian ini mengimplementasikan metode *Random Forest* dalam mengklasifikasikan analisis sentimen terhadap opini publik di media sosial X (Twitter) mengenai kebijakan efisiensi anggaran. Data sentimen dikumpulkan dari *tweet* selama Januari–Februari 2025 sebanyak 1.589 data. Kemudian dilakukan *pre-processing text* dan pembobotan TF-IDF pada data tersebut sebelum dilakukan klasifikasi menggunakan *Random Forest*. Berdasarkan hasil evaluasi, model terbaik diperoleh pada rasio pelatihan 85:15 dengan akurasi sebesar 76,57%, *precision* rata-rata 79,01%, *recall* rata-rata 73,35%, dan *F1-score* rata-rata 75,11%. Hasil dari penelitian ini diharapkan dapat membantu pemerintah dalam memahami pandangan masyarakat dan menjadi dasar dalam merumuskan kebijakan yang lebih tepat sasaran.

Kata kunci : Analisis Sentimen, Efisiensi Anggaran, Media Sosial, Machine Learning, Random Forest, TF-IDF

1. PENDAHULUAN

Kebijakan efisiensi anggaran merupakan salah satu strategi pemerintah dalam mengelola keuangan negara agar lebih efektif dan tepat guna. Kebijakan ini bertujuan untuk memaksimalkan manfaat dari setiap pengeluaran publik, namun tidak jarang menimbulkan pro dan kontra di kalangan masyarakat. Opini publik terhadap kebijakan tersebut sering kali disampaikan melalui media sosial, yang kini menjadi ruang terbuka bagi masyarakat untuk mengekspresikan pandangan mereka.

Seiring meningkatnya penggunaan media sosial seperti X (Twitter), media ini menjadi sumber data yang sangat kaya untuk menggambarkan persepsi publik secara *real-time*. Dalam konteks ini, analisis sentimen menjadi metode yang relevan untuk mengolah dan memahami opini masyarakat. Melalui klasifikasi teks menjadi sentimen positif, negatif, maupun netral, analisis sentimen memungkinkan pengambilan kesimpulan yang lebih akurat terhadap respons masyarakat.

Penelitian ini menggunakan algoritma *Random Forest* untuk melakukan klasifikasi sentimen. *Random Forest* dipilih karena kemampuannya dalam menangani data teks yang kompleks serta menghasilkan performa klasifikasi yang tinggi, seperti yang ditunjukkan dalam berbagai studi sebelumnya. Dengan pendekatan ini, diharapkan dapat diperoleh gambaran umum mengenai tanggapan masyarakat terhadap kebijakan efisiensi anggaran [1].

Hasil dari penelitian ini diharapkan dapat memberikan masukan bagi pemerintah dalam merumuskan kebijakan yang lebih responsif terhadap opini publik, serta meningkatkan kesadaran

masyarakat akan pentingnya partisipasi dalam menyampaikan pendapat melalui media sosial.

2. TINJAUAN PUSTAKA

2.1. Penelitian Terdahulu

Menurut Muhammad Yusril Aldean dalam penelitiannya yang berjudul "Analisis Sentimen Masyarakat Terhadap Vaksinasi Covid-19 di Twitter Menggunakan Metode *Random Forest Classifier* (Studi Kasus Vaksin Sinovac)", penelitian ini bertujuan untuk mengetahui sentimen masyarakat terhadap vaksinasi COVID-19 khususnya vaksin Sinovac di media sosial Twitter, menggunakan pendekatan algoritma *Random Forest*. Dari hasil pengujian, model menghasilkan akurasi sebesar 79% dalam mengklasifikasikan sentimen positif dan negatif. Penelitian ini diharapkan dapat memberikan informasi kepada pemerintah dan masyarakat mengenai persepsi publik terhadap program vaksinasi, serta menjadi dasar guna mendukung pengambilan keputusan yang lebih optimal di tengah situasi pandemi [2].

Menurut Tanti Cahya Herdiyani dalam penelitiannya yang berjudul "Sentiment Analysis Terkait Pemindahan Ibu Kota Indonesia Menggunakan Metode *Random Forest* Berdasarkan *Tweet* Warga Negara Indonesia", penelitian dilakukan melalui tahapan pengumpulan data menggunakan Twitter API dan *scraping* manual, dilanjutkan dengan *preprocessing* teks, pembobotan menggunakan TF-IDF, dan klasifikasi dengan algoritma *Random Forest*. Berdasarkan pengujian terhadap 1.639 *tweet*, model berhasil mencapai akurasi sebesar 76%, dengan *reacall* sebesar 70%, negatif 35%, *precision* sebesar 69%, dan *f1-score* sebesar 69 %. Hasil penelitian ini diharapkan

dapat memberikan gambaran mengenai persepsi publik terhadap isu pemindahan ibu kota dan menjadi bahan pertimbangan dalam perumusan kebijakan pemerintah [3].

2.2. Analisis sentimen

Analisis sentimen merupakan metode dalam *text mining* yang bertujuan untuk menggali dan mengidentifikasi emosi maupun pendapat subjektif dari suatu teks ke dalam klasifikasi tertentu seperti positif, netral, atau negatif. Dalam implementasinya, analisis sentimen sering digunakan pada data teks yang tidak terstruktur seperti yang di temukan pada komentar atau unggahan di media sosial, ulasan produk, berita daring, dan dokumen digital lainnya [4].

2.3. Scraping

Web scraping merupakan proses otomatisasi untuk mengekstrak data dari halaman *web*, seperti artikel berita, ulasan produk, komentar media sosial, atau forum diskusi. Metode ini memungkinkan pengambilan data teks berskala besar secara efisien dan sistematis oleh peneliti, yang kemudian dapat dijadikan sebagai bahan utama untuk analisis lebih lanjut, termasuk dalam analisis sentimen [5].

2.4. Pre-Processing

Pre-processing merupakan tahapan untuk membersihkan dan menyiapkan data teks sebelum dianalisis agar informasi penting dapat diekstraksi secara efisien. Beberapa tahapan yang dilakukan dalam proses ini antara lain: *case folding*, pembersihan data (*cleaning text*), tokenisasi, normalisasi, penghapusan kata-kata umum (*stopword removal*) serta *stemming* [6].

2.5. TF-IDF

TF-IDF (*Term Frequency-Inverse Document Frequency*) merupakan teknik pembobotan kata dalam *machine mining* yang digunakan untuk menilai tingkat kepentingan suatu kata dalam sebuah dokumen relatif terhadap seluruh kumpulan dokumen. Teknik ini membantu untuk menurunkan bobot kata-kata yang umum dan menaikkan bobot kata-kata yang unik pada suatu dokumen [7].

Tahapan TF-IDF meliputi:

a. Term Frequency (TF)

Menunjukkan frekuensi kemunculan suatu kata dalam sebuah dokumen. Rumus:

$$TF = \frac{(\text{Jumlah kemunculan kata dalam dokumen})}{(\text{Total kata dalam dokumen})} \quad (1)$$

b. Document Frequency (DF)

Banyaknya dokumen yang memuat kata tersebut.

c. Inverse Document Frequency (IDF)

Digunakan untuk menilai seberapa unik sebuah kata dengan menggunakan rumus:

$$IDF = \log \left(\frac{N}{DF} \right) \quad (2)$$

d. TF-IDF

Perhitungan bobot dilakukan dengan mengalikan nilai TF dan IDF untuk mendapatkan bobot akhir.

$$TF - IDF = TF \times IDF \quad (3)$$

TF-IDF digunakan secara luas dalam analisis sentimen karena membantu mengidentifikasi kata-kata yang memiliki makna penting untuk proses klasifikasi.

2.6. Random Forest

Random Forest merupakan metode klasifikasi berbasis *ensemble learning* yang membangun banyak pohon keputusan (*decision tree*) dan menggabungkan hasil prediksi dari masing-masing pohon. Metode ini merupakan pengembangan dari teknik *bagging* yang dirancang untuk meningkatkan akurasi dan mengurangi *overfitting* [8].

Setiap pohon dibangun menggunakan *subset* data hasil teknik *bootstrap*, dan pada setiap *node* hanya *subset* acak dari fitur yang dipertimbangkan untuk pemisahan. Hasil akhir dari model *Random Forest* adalah prediksi berdasarkan suara terbanyak dari pohon-pohon tersebut.

2.7. Confusion Matrix

Confusion matrix adalah representasi berbentuk tabel yang menyajikan gambaran evaluasi kinerja model klasifikasi dalam *machine learning*. Struktur *confusion matrix* dibagi menjadi empat bagian penting yaitu *True Positive* (TP), *False Positive* (FP), *True Negative* (TN), dan *False Negative* (FN) [9].

Matriks ini menjadi dasar perhitungan metrik evaluasi lain seperti:

a. Akurasi

Akurasi merupakan total prediksi benar dibagi jumlah total prediksi.

b. Presisi

Rasio antara TP dan (TP + FP), mengukur ketepatan prediksi positif.

c. Recall

Rasio antara TP dan (TP + FN), menunjukkan seberapa efektif model menangkap kelas positif.

d. F1-score

Rata-rata harmonik dari presisi dan *recall*, cocok untuk data yang tidak merata.

Confusion matrix penting dalam konteks analisis sentimen karena mampu menunjukkan ketepatan model dalam mengenali sentimen positif, negatif, maupun netral secara menyeluruh dan adil.

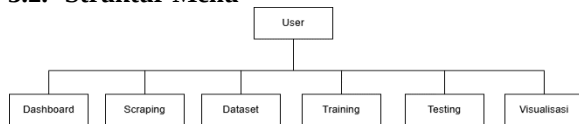
3. METODE PENELITIAN

3.1. Analisis Masalah

Kebijakan efisiensi anggaran dari pemerintah sering menimbulkan beragam tanggapan di masyarakat, baik yang mendukung maupun yang merasa dirugikan. Media sosial menjadi wadah utama masyarakat untuk menyuarakan pendapat tersebut. Sayangnya, opini yang tersebar belum dimanfaatkan secara maksimal karena volumenya yang besar dan sulit dianalisis secara manual. Padahal, informasi ini

bisa menjadi masukan penting dalam pengambilan kebijakan. Oleh karena itu, dibutuhkan sistem yang mampu menganalisis sentimen masyarakat secara otomatis dan efisien. Penelitian ini bertujuan membangun sistem analisis sentimen menggunakan metode *Random Forest* agar pemerintah dapat memahami opini publik secara lebih objektif dan berbasis data.

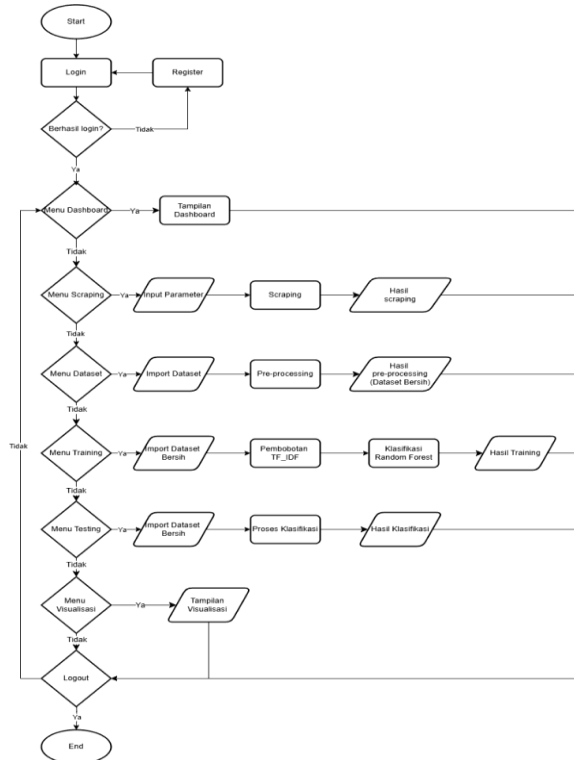
3.2. Struktur Menu



Gambar 1. Struktur Menu

Struktur menu sistem dirancang untuk memudahkan pengguna dalam mengakses fitur-fitur utama. Setelah berhasil *login*, pengguna secara otomatis dibawa ke halaman *dashboard*. Pengguna dapat melakukan pengumpulan data sentimen melalui menu *scraping*. Pada menu *Dataset*, pengguna dapat mengunggah *file dataset* berformat *.csv* atau *.xlsx*. Data yang diunggah kemudian melalui proses *pre-processing* untuk menghasilkan data yang bersih dan siap digunakan. *Dataset* yang sudah bersih ini dapat digunakan untuk melatih model menggunakan metode *Random Forest* melalui menu *Training*, dan diuji pada menu *Testing* untuk mengklasifikasikan sentimen. Hasil klasifikasi kemudian divisualisasikan dalam bentuk grafik melalui menu *visualisasi* agar lebih mudah dipahami oleh pengguna.

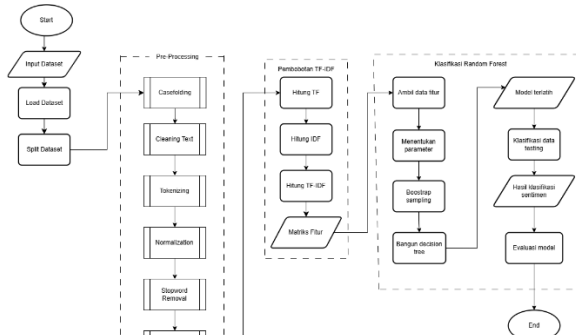
3.3. Flowchart Sistem



Gambar 2. Flowchart Sistem

Pada gambar 2 ditampilkan *flowchart* sistem. Pengguna yang belum memiliki akun harus registrasi terlebih dahulu, kemudian *login* untuk masuk ke *dashboard*. Dari *dashboard*, pengguna dapat mengakses menu *dataset* untuk mengunggah dan melakukan *pre-processing* data. *Dataset* yang telah dibersihkan digunakan untuk melatih model *Random Forest* dan juga untuk analisis pada menu *testing*. Hasil klasifikasi ditampilkan di halaman hasil, lalu divisualisasikan dalam bentuk diagram pada menu *visualisasi*. Setelah semua proses selesai, pengguna dapat keluar dari sistem.

3.4. Flowchart Metode



Gambar 3. Flowchart Metode

Flowchart metode pada gambar 3 menggambarkan alur analisis sentimen menggunakan metode *Random Forest* dan *TF-IDF*. Proses diawali dengan unggah *dataset*, dilanjutkan *pre-processing* seperti *casefolding*, *tokenizing*, *normalisasi*, hingga *stemming*. Setelah itu dilakukan pembobotan *TF-IDF* dan data diformat sesuai kebutuhan model. Model *Random Forest* kemudian dilatih melalui pembentukan pohon keputusan. Hasil dari model digunakan untuk mengklasifikasikan data *testing* menjadi sentimen positif, negatif, atau netral, lalu dievaluasi untuk mengukur performa

4. HASIL DAN PEMBAHASAN

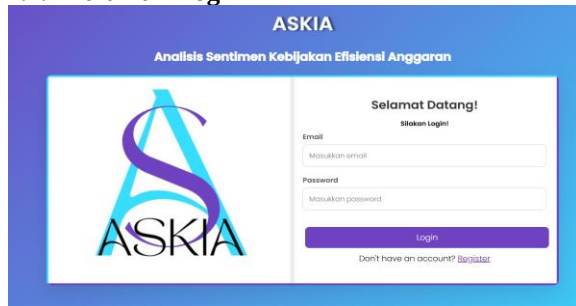
4.1. Halaman Register

Gambar 4. Halaman Register

Pada gambar 4 ditampilkan halaman *register* yang berfungsi untuk melakukan pendaftaran akun. Pada halaman ini *user* akan diminta untuk mengisi *form email*, *username* dan *password*. Setelah

pendaftaran berhasil user dapat melakukan proses login.

4.2. Halaman Login



Gambar 5. Halaman Login

Pada gambar 5 ditampilkan halaman login yang terdiri dari email dan password. Pada bagian login ini, user akan diminta untuk mengisi email dan password yang telah terdaftar dengan benar. Apabila belum terdaftar maka dapat mendaftar dulu pada halaman register.

4.3. Halaman Dashboard



Gambar 6. Halaman Dashboard

Pada gambar 6 ditampilkan halaman dashboard dari web yang sudah dibuat. Dashboard tersebut berisikan tampilan jumlah dataset yang digunakan beserta dengan pembagiannya. Selain itu terdapat juga penjelasan singkat mengenai analisis sentimen, topik yang dibahas serta penjelasan fungsi-fungsi dari setiap menu yang ada.

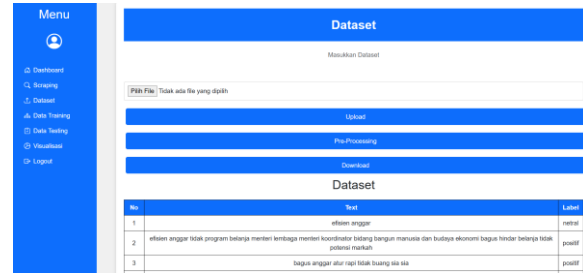
4.4. Scraping



Gambar 7. Halaman Scraping

Pada gambar 7 ditampilkan menu scraping yang digunakan untuk mencari dan mengambil data sentimen dari sosial media x berdasarkan parameter yang dimasukkan. Kemudian file yang berisikan dataset sentimen tersebut dapat diunduh.

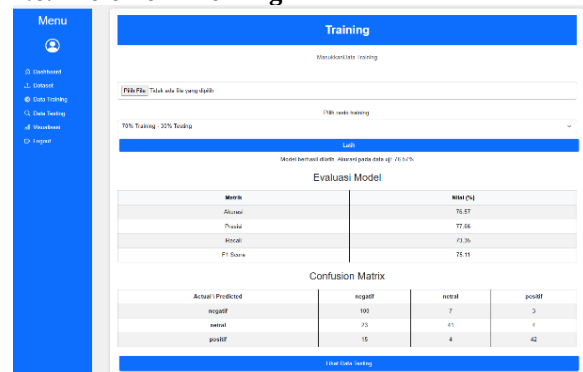
4.5. Halaman Dataset



Gambar 8. Halaman Dataset

Pada gambar 8 ditampilkan menu dataset dari web yang sudah dibuat. Menu dataset tersebut bisa digunakan untuk mengunggah file dataset, kemudian file yang diunggah dapat dilakukan preprocessing serta hasil dari preprocessing tersebut dapat diunduh. Kemudian hasil dari file dataset yang sudah dilakukan preprocessing juga akan ditampilkan dalam bentuk tabel.

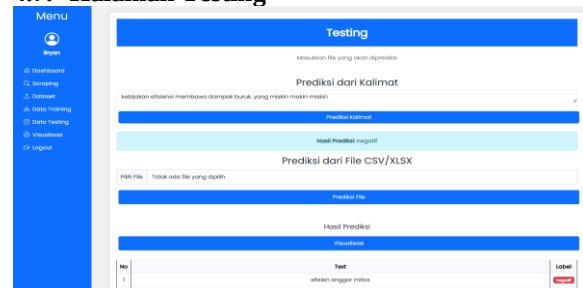
4.6. Halaman Training



Gambar 9. Halaman Training

Pada gambar 9 ditampilkan menu training dari web yang sudah dibuat. Menu training tersebut bisa digunakan untuk mengunggah file dataset yang sudah dibersihkan (preprocessing), agar digunakan untuk melatih model. Terdapat juga dropdown pembagian dataset ke dalam subset training dan testing. Kemudian setelah dilakukan pelatihan, maka akan ditampilkan hasil evaluasi model dan juga confusion matrix. Terdapat juga button lihat data testing yang digunakan untuk menampilkan hasil testing dari dataset yang sudah dilatih sebelumnya.

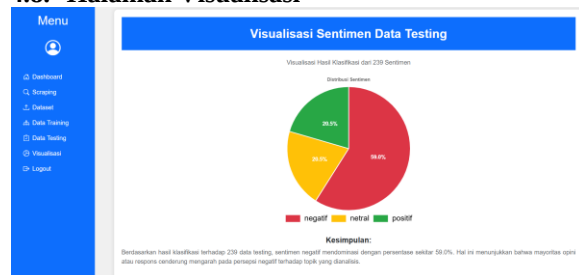
4.7. Halaman Testing



Gambar 10. Halaman Testing

Pada gambar 10 ditampilkan menu *testing*. Menu tersebut bisa digunakan untuk mengunggah *file* maupun menuliskan kalimat yang akan diprediksi. Hasil prediksi tersebut akan ditampilkan dalam bentuk teks dan tabel yang berisikan *text* sentimen dan label hasil prediksinya. Selain itu terdapat juga *button* visualisasi yang digunakan untuk menampilkan hasil *testing* dalam bentuk grafik

4.8. Halaman Visualisasi



Gambar 11. Halaman Visualisasi

Pada gambar 11 ditampilkan menu visualisasi dari *web* yang sudah dibuat. Menu visualisasi tersebut bisa digunakan untuk menampilkan grafik dari data *testing* pada menu *testing*. Terdapat jumlah sentimen, persebaran sentimen positif, negatif dan netral yang dibedakan dengan warna. Sentimen negatif ditandai dengan warna merah, sentimen positif dengan warna hijau, dan sentimen netral dengan warna kuning.

4.9. Pengumpulan Data

Data dalam penelitian ini bersumber dari media sosial X menggunakan *library Tweet-Harvests*. Proses pengambilan data dilakukan dengan memasukkan kata kunci kebijakan efisiensi anggaran dalam rentang waktu bulan Januari hingga Februari 2025. Setelah kata kunci dimasukkan, sistem akan menjalankan proses *scraping* untuk mengumpulkan *tweet* yang sesuai. Hasil dari proses *crawling* tersebut kemudian disimpan dalam format *file .csv*. *Dataset* yang berhasil didapatkan dan dipakai dalam penelitian ini sebanyak 1.589 baris data, yang berisi sentimen terkait kebijakan efisiensi anggaran.

4.10. Pelabelan Data

Proses pelabelan data dilakukan secara manual oleh *annotator* untuk menentukan sentimen dari setiap data yang telah dikumpulkan. *Dataset* yang digunakan pada tahap pelatihan model terdiri dari 1.589 baris data, yang berasal dari *tweet* di media sosial X. Setiap baris data dilengkapi dengan anotasi atau label sentimen, berupa positif, negatif, dan netral, sesuai dengan isi dan konteks cuitan. Hasil pelabelan *dataset* terdiri dari 405 data berlabel positif, 732 data berlabel negatif, dan 452 data berlabel netral.

4.11. Preprocessing Data

Langkah awal yang harus dilakukan sebelum proses klasifikasi adalah melakukan proses *pre-*

processing teks. Tahapan *pre-processing* data yang dilakukan adalah seperti berikut ini:

a. Casefolding

Mengonversi seluruh karakter dalam teks menjadi huruf kecil untuk menghindari duplikasi kata akibat perbedaan kapitalisasi [10].

Tabel 1. Casefolding

Input	Output
@BetaEpsilonPhi Negara ini punya duid ga sih??.. Tekor diakhir tahun anggaran dan Efisiensi di awal anggaran nunjukan spt lagi sakit.. Liat berita: Pemberitaan ttg penggunaan dana pribadi tuk kegiatan negara.. Gaji tunjangan dan honor di daerah yg tertunda atau tdk dibayarkan	@betaepsilonphi negara ini punya duid ga sih??.. tekor diakhir tahun anggaran dan efisiensi di awal anggaran nunjukan spt lagi sakit.. liat berita: pemberitaan ttg penggunaan dana pribadi tuk kegiatan negara.. gaji tunjangan dan honor di daerah yg tertunda atau tdk dibayarkan

b. Cleaning Text

Menghapus atau menghilangkan karakter, simbol, angka, tanda baca, URL, serta karakter yang tidak relevan [11].

Tabel 2. Cleaning Text

Input	Output
@betaepsilonphi negara ini punya duid ga sih??.. tekor diakhir tahun anggaran dan efisiensi di awal anggaran nunjukan spt lagi sakit.. liat berita: pemberitaan ttg penggunaan dana pribadi tuk kegiatan negara.. gaji tunjangan dan honor di daerah yg tertunda atau tdk dibayarkan	negara ini punya duid ga sih tekor diakhir tahun anggaran dan efisiensi di awal anggaran nunjukan spt lagi sakit liat berita pemberitaan ttg penggunaan dana pribadi tuk kegiatan negara gaji tunjangan dan honor di daerah yg tertunda atau tdk dibayarkan

c. Tokenizing.

Memecah teks menjadi unit terkecil, biasanya berupa kata [12].

Tabel 3. Tokenizing

Input	Output
negara ini punya duid ga sih tekor diakhir tahun anggaran dan efisiensi di awal anggaran nunjukan spt lagi sakit liat berita pemberitaan ttg penggunaan dana pribadi tuk kegiatan negara gaji tunjangan dan honor di daerah yg tertunda atau tdk dibayarkan	'negara', 'ini', 'punya', 'duid', 'ga', 'sih', 'tekor', 'diakhir', 'tahun', 'anggaran', 'dan', 'efisiensi', 'di', 'awal', 'anggaran', 'nunjukan', 'spt', 'lagi', 'sakit', 'liat', 'berita', 'pemberitaan', 'ttg', 'penggunaan', 'dana', 'pribadi', 'tuk', 'kegiatan', 'negara', 'gaji', 'tunjangan', 'dan', 'honor', 'di', 'daerah', 'yg', 'tertunda', 'atau', 'tdk', 'dibayarkan'

d. Normalisasi.

Mengubah kata tidak baku menjadi bentuk baku, seperti “gw” menjadi “saya” [13].

Tabel 4. Normalisasi

Input	Output
negara ini punya duit ga sih tekor diakhir tahun anggaran dan efisiensi di awal anggaran nunjukan spt lagi sakit liat berita pemberitaan ttg penggunaan dana pribadi tuk kegiatan negara gaji tunjangan dan honor di daerah yg tertunda atau tdk dibayarkan	negara ini punya uang tidak sih rugi di akhir tahun anggaran dan efisien di awal anggaran nunjukkan seperti lagi sakit lihat berita pemberitaan tentang penggunaan dana pribadi untuk giat negara gaji tunjangan dan honor di daerah yang tertunda atau tidak dibayarkan

e. Stopword Removal

Menghilangkan kata-kata umum yang tidak memberikan informasi yang signifikan [14].

Tabel 5. Stopword Removal

Input	Output
negara ini punya uang tidak sih rugi di akhir tahun anggaran dan efisien di awal anggaran nunjukkan seperti	negara uang tidak rugi akhir tahun anggaran efisien awal anggaran nunjukkan sakit

Input	Output
lagi sakit lihat berita pemberitaan tentang penggunaan dana pribadi untuk giat negara gaji tunjangan dan honor di daerah yang tertunda atau tidak dibayarkan	pemberitaan penggunaan dana pribadi giat negara gaji tunjangan honor daerah tertunda tidak dibayarkan

f. Stemming

Mengubah kata berimbuhan ke bentuk dasar, seperti kata “memakan” diubah menjadi “makan” [15].

Tabel 6. Stemming

Input	Output
negara uang tidak rugi akhir tahun anggaran efisien awal anggaran nunjukkan sakit pemberitaan penggunaan dana pribadi giat negara gaji tunjangan honor daerah tertunda tidak dibayarkan	negara uang tidak rugi akhir tahun anggar efisien awal anggar nunjuk sakit berita guna dana pribadi giat negara gaji tunjang honor daerah tunda tidak bayar

4.12. Pengujian Blackbox

Pengujian *blackbox* merupakan teknik untuk menguji perangkat lunak untuk menilai fungsi sistem tanpa perlu tahu bagaimana sistem bekerja di dalamnya (tidak melihat kode program), tapi hanya fokus pada *input* dan *output*.

Tabel 7. Pengujian Blackbox

No	Fitur yang Diuji	Deskripsi Pengujian	Input	Ekspektasi	Status
1	Login Pengguna	Menguji validasi login dengan data yang benar dan salah	Email & Password	- Login berhasil jika data valid - Tampilkan pesan <i>error</i> jika tidak valid	Lulus
2	Register Pengguna	Menguji validasi register dengan data yang benar dan salah	Email, Username & Password	- Register berhasil jika data valid - Tampilkan pesan <i>error</i> jika tidak valid	Lulus
3	Form input Kata Kunci Pencarian	Menguji apakah form dapat menerima input text lebih dari satu kata	Text	Text yang dimasukkan berhasil digunakan sebagai kata kunci dalam melakukan scraping	Lulus
4	Form input Sejak Tanggal	Menguji apakah form dapat menerima input berupa date dengan tanggal yang tidak melebihi tanggal pada form “Hingga Tanggal”	Tanggal	- Form berhasil menerima input dengan tipe data date - Tanggal yang dipilih berhasil tidak melebihi tanggal pada form “Hingga Tanggal”	Lulus
5	Form input Hingga Tanggal	Menguji apakah form dapat menerima input berupa date dengan tanggal yang tidak lebih awal dari tanggal pada form “Sejak Tanggal”	Tanggal	- Form berhasil menerima input dengan tipe data date - Tanggal yang dipilih berhasil tidak lebih awal dari tanggal pada form “Hingga Tanggal”	Lulus
6	Form input Jumlah Tweet	Menguji apakah form dapat menerima input number diantara angka 10 sampai 500	Number	Form berhasil menerima masukan number	Lulus
7		Menguji apakah form dapat menerima input number selain diantara angka 10 sampai 500	Number	Muncul peringatan bahwa number yang dimasukkan tidak berada diantara angka 10 sampai 500	Lulus
8	Form Format File	Menguji apakah sistem dapat mengunduh file berdasarkan format file yang dipilih (csv & excel)	Dataset hasil scraping	- Sistem berhasil mengunduh file dengan format csv - Sistem berhasil mengunduh file dengan format excel	Lulus

No	Fitur yang Diuji	Deskripsi Pengujian	Input	Ekspektasi	Status
9	Scraping	Menguji apakah sistem bisa melakukan scraping berdasarkan parameter yang dimasukkan	Parameter scraping di menu scraping	Sistem berhasil mengumpulkan dataset sentimen dari media sosial sesuai dengan parameter yang dimasukkan	Lulus
10	Upload File di Halaman Dataset	Menguji apakah sistem bisa menerima file excel/CSV	File excel berisikan data sentimen	Data berhasil diunggah dan bisa ditampilkan	Lulus
11	Pre-processing Dataset	Menguji apakah sistem bisa melakukan pre-processing pada dataset	File excel yang berisikan data sentimen	File berhasil di lakukan pre-processing dan ditampilkan	Lulus
12	Upload File di Halaman Training	Menguji apakah sistem bisa menerima file excel/CSV	File hasil pre-processing	Data berhasil diunggah	Lulus
13	Proses Training	Menguji apakah sistem bisa melakukan proses training	File hasil pre-processing	Proses Training berhasil dilakukan dan ditampilkan hasil evaluasi mode serta confusion matrix	Lulus
14	Fitur prediksi kalimat	Menguji apakah sistem bisa melakukan prediksi kalimat	Text	Text berhasil diinput dan diprediksi	Lulus
15	Upload File di Halaman Testing	Menguji apakah sistem bisa menerima file excel/CSV	File hasil pre-processing	Data berhasil diunggah	Lulus
16	Proses Testing	Menguji apakah sistem bisa melakukan proses testing	File hasil pre-processing	Proses Testing berhasil dilakukan dan ditampilkan hasil prediksi dan visualisasi	Lulus
17	Tampilan Visualisasi	Menguji apakah sistem bisa menampilkan visualisasi dalam bentuk pie chart	Data hasil traning dan testing	Berhasil menampilkan hasil dari training dan testing dalam bentuk pie chart	Lulus

4.13. Evaluasi Klasifikasi Model

Pembagian data *split* terbaik adalah 85% untuk *training* dan 15% untuk *testing*, karena menghasilkan performa klasifikasi tertinggi dibandingkan dengan rasio lainnya. Berdasarkan *confusion matrix*, model berhasil mengklasifikasikan 183 dari 239 data uji dengan benar.

Tabel 8. Confusion Matrix

Actual\Predicted	Negatif	Netral	Positif
Negatif	100	7	3
Netral	23	41	4
Positif	15	4	42

Hasil perhitungan akurasi serta presisi, *recall* dan *f1-score* tiap kelas.

1. Akurasi

$$Akurasi = \frac{100 + 41 + 42}{239} = \frac{183}{239} = 76,57\%$$

2. Kelas Negatif

Tabel 9. Evaluasi Model Kelas Negatif

Matrix	Nilai
Precision	72,46%
Recall	90,91%
F1 Score	80,65%

3. Kelas Netral

Tabel 10. Evaluasi Model Kelas Netral

Matrix	Nilai
Precision	78,85%
Recall	60,29%
F1 Score	68,33%

4. Kelas Positif

Tabel 11. Evaluasi Model Kelas Positif

Matrix	Nilai
Precision	85,71%
Recall	68,85%
F1 Score	76,36%

Berdasarkan hasil evaluasi model dengan data *split* 85% untuk *training* dan 15% untuk *testing*, diperoleh tingkat akurasi sebesar 76,57% dengan nilai *F1-score* pada kelas positif sebesar 76,36%, kelas negatif sebesar 80,65%, dan kelas netral sebesar 68,33%. Maka dapat disimpulkan bahwa hasil penelitian yang mengimplementasikan metode *Random Forest* ini menghasilkan model dengan tingkat akurasi dan keselarasan yang cukup tinggi dalam mengklasifikasikan sentimen, terutama untuk kategori negatif dan positif, meskipun masih terdapat sedikit kesulitan dalam mengenali sentimen netral secara konsisten.

5. KESIMPULAN DAN SARAN

Penerapan metode *Random Forest* berhasil digunakan untuk mengklasifikasikan sentimen masyarakat terhadap kebijakan efisiensi anggaran yang diambil dari media sosial X (Twitter), melalui tahapan *preprocessing* dan pembobotan kata menggunakan TF-IDF. Performa terbaik diperoleh ketika model dilatih dan diuji dengan rasio pembagian *dataset* 85:15, menghasilkan akurasi sebesar 76,57%, presisi rata-rata 79,01%, *recall* rata-rata 73,35%, dan *F1-score* rata-rata 75,11%. Sistem berbasis *web* yang

dibangun dengan Flask memudahkan pengguna dalam mengunggah *dataset*, melatih model, menguji data baru, serta menampilkan hasil analisis dalam bentuk visualisasi yang informatif. Dari hasil analisis, ditemukan bahwa sentimen negatif mendominasi tanggapan masyarakat terhadap kebijakan tersebut, menunjukkan adanya ketidakpuasan yang bisa menjadi bahan pertimbangan bagi pemerintah.

DAFTAR PUSTAKA

- [1] D. A. Fitri and Damayanti, "Komparasi Algoritma Random Forest Classifier dan Support Vector Machine Untuk Sentimen Masyarakat Terhadap Pinjaman Online di Media Sosial," vol. 9, no. 4, pp. 2018–2029, 2024.
- [2] M. Y. Aldean, P. Paradise, and N. A. Setya Nugraha, "Analisis Sentimen Masyarakat Terhadap Vaksinasi Covid-19 di Twitter Menggunakan Metode Random Forest Classifier (Studi Kasus: Vaksin Sinovac)," *J. Informatics, Inf. Syst. Softw. Eng. Appl.*, vol. 4, no. 2, pp. 64–72, 2022, doi: 10.20895/inista.v4i2.575.
- [3] T. C. Herdiyani and A. U. Zailani, "Sentiment Analysis Terkait Pemindahan Ibu Kota Indonesia Menggunakan Metode Random Forest Berdasarkan Tweet Warga Negara Indonesia," *J. Teknol. Sist. Inf.*, vol. 3, no. 2, pp. 154–165, 2022, doi: 10.35957/jtsi.v3i2.2920.
- [4] D. S. Prasetyo, K. Auliasari, and Y. A. Pranoto, "Analisis Sentimen Pada Komentar Media Sosial Terkait Isu Joki Dengan Menggunakan Metode LSTM," *J. Comput. Eng. Syst. Sci.*, vol. 10, no. 1, pp. 39–52, 2025, [Online]. Available: <https://jurnal.unimed.ac.id/2012/index.php/cess>
- [5] D. F. Setiawan, T. Tristiyanto, and A. Hijriani, "Aplikasi Web Scraping Deskripsi Produk," *J. Teknoinfo*, vol. 14, no. 1, p. 41, 2020, doi: 10.33365/jti.v14i1.498.
- [6] D. Rifaldi, Abdul Fadlil, and Herman, "Teknik Preprocessing Pada Text Mining Menggunakan Data Tweet 'Mental Health,'" *Decod. J. Pendidik. Teknol. Inf.*, vol. 3, no. 2, pp. 161–171, 2023, doi: 10.51454/decode.v3i2.131.
- [7] R. Deswandi Yahya, S. Adi Wibowo, and N. Vendyansyah, "Analisis Sentimen Untuk Deteksi Ujaran Kebencian Pada Media Sosial Terkait Pemilu 2024 Menggunakan Metode Support Vector Machine," *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 8, no. 2, pp. 1182–1189, 2024, doi: 10.36040/jati.v8i2.9076.
- [8] M. L. Suliztia, "Penerapan Analisis Random Forest Pada Prototype Sistem Prediksi Harga Kamera Bekas Menggunakan Flask," *Fak. Mat. Dan Ilmu Pengetah. Alam*, pp. 1–107, 2020.
- [9] T. Firmansyah, R. Kurniawan, and A. T. Hidayat, "Klasifikasi Tingkat Kematangan Roasting Biji Kopi Berbasis Deep Learning dengan Arsitektur MobileNet," no. February, 2025, doi: 10.47065/josh.v6i2.6811.
- [10] N. Widhiyanta, I. Muhandhis, R. S. Jannah, and L. A. Wulansari, "Analisis Sentimen Ulasan Produk Moisturizer Skintific di Toko Pedia Menggunakan Support Vector Machine," vol. 18, no. 1, pp. 129–142, 2025.
- [11] N. Putu, K. Pradnyawati, D. Pramana, I. G. Ngurah, and A. Kusuma, "Analisis Sentimen Ulasan Pengguna GlobalXtreme di Google Review Menggunakan Metode SVM (Support Vector Machine)," vol. 2, no. 1, pp. 829–834, 2025.
- [12] G. V. Simbolon and L. Albinur, "Analisis Serntimen Pengguna Aplikasi Loket X Menggunakan Metode KNN, SVM, DAN Naive Bayes," *J. Inov. Glob.*, vol. 2, no. 3, pp. 543–551, 2024.
- [13] S. A. Putri, N. W. Kencana, and A. Khoirudin, "Analisis Sentimen Media Sosial X terhadap Gerakan Muhammadiyah Menggunakan Algoritma Naïve Bayes," vol. 6, no. 1, pp. 9–17, 2025.
- [14] R. T. Adek, Z. Fitri, and S. C. Siegar, "Analisis Sentimen Komentar Pada Saluran Youtube Beauty Vlogger Berbahasa Indonesia Menggunakan Metode Support Vector Machine," vol. 5, no. 2, pp. 164–175, 2025, doi: 10.35957/algoritme.v5i2.9692.
- [15] R. J. Kasim and E. Utami, "Penerapan Algoritma Boyer Moore yang di Modifikasi untuk stemmer Bahasa Indonesia," *JUPI (Jurnal Ilm. Penelit. dan Pembelajaran Inform.)*, vol. 9, no. 3, pp. 1657–1667, 2024, doi: 10.29100/jipi.v9i3.5449.