

ANALISIS SENTIMEN PUBLIK TERHADAP PENJUALAN IPHONE 16 DAN KEBIJAKAN TKDN DI INDONESIA

ABSTRACT

The Domestic Component Level (TKDN) policy on Apple products in Indonesia has sparked various opinions on social media, particularly on platform X. This study analyzes public sentiment toward the policy using machine learning methods. Data was collected through crawling techniques, processed through preprocessing, and extracted using TF-IDF. The Random Forest algorithm was applied to classify opinions into negative, neutral, and positive categories. Experimental results show that the model achieved 91% accuracy, with 91% precision, 91% recall, and an F1-score of 89%. Confusion matrix analysis revealed that the model tends to classify sentiment as neutral more frequently, which may be due to class imbalance in the dataset. To improve performance, future research can explore class balancing techniques or more complex models. These findings provide insights for industry players and policymakers in understanding public perceptions regarding the TKDN policy on Apple products.

Keywords: Sentiment analysis, TKDN policy, Apple, machine learning, X

ABSTRAK

Kebijakan Tingkat Komponen Dalam Negeri (TKDN) terhadap produk Apple di Indonesia telah memicu berbagai opini di media sosial, terutama di platform X. Penelitian ini menganalisis sentimen publik terhadap kebijakan tersebut menggunakan metode *machine learning*. Data dikumpulkan dengan teknik *crawling*, diproses melalui *preprocessing*, dan diekstraksi menggunakan TF-IDF. Algoritma Random Forest diterapkan untuk mengklasifikasi opini menjadi kategori negatif, netral, dan positif. Hasil eksperimen menunjukkan model mencapai akurasi 91%, dengan *precision* 91%, *recall* 91% dan *f1-score* 89%. Analisis *confusion matrix* mengungkapkan bahwa model sering mengklasifikasikan sentimen sebagai netral, yang dapat disebabkan oleh ketidakseimbangan kelas dalam dataset. Untuk meningkatkan performa, penelitian selanjutnya dapat mengeksplorasi teknik penyeimbangan kelas atau model yang lebih kompleks. Temuan ini memberikan wawasan bagi industri dan pembuat kebijakan dalam memahami persepsi masyarakat terkait kebijakan TKDN pada produk Apple.

Kata Kunci: Analisis sentimen, TKDN, Apple, Machine learning, X

Riwayat Artikel :

Tanggal diterima : 19-02-2025

Tanggal revisi : 22-02-2025

Tanggal terbit : 22-02-2025

DOI :

<https://doi.org/10.31949/infotech.v11i1.13159>

INFOTECH journal by Informatika UNMA is licensed under CC BY-SA 4.0

Copyright © 2025 By Author



1. PENDAHULUAN

1.1. Latar Belakang

Apple merupakan salah satu perusahaan teknologi terbesar di dunia yang terus berinovasi dalam pengembangan perangkat keras dan lunak. Di Indonesia, kehadiran produk Apple, khususnya iPhone, sering kali dikaitkan dengan regulasi Tingkat Komponen Dalam Negeri (TKDN), yang mengatur batas minimal kandungan lokal dalam produk teknologi yang dipasarkan di dalam negeri. Kementerian Perindustrian Republik Indonesia (Kemenperin) menetapkan bahwa Apple harus memenuhi ketentuan TKDN untuk dapat menjual iPhone 16 di Indonesia. Dalam negosiasi terbaru, Apple mengajukan skema investasi inovasi, namun nilai yang diajukan masih di bawah ekspektasi teknokratis yang ditetapkan oleh Kemenperin. Meskipun Apple telah melakukan pelatihan dan pendidikan sejak 2017, kontribusi dalam riset dan pengembangan (R&D) masih dinilai belum optimal. Oleh karena itu, aspek TKDN dan investasi Apple di Indonesia menjadi topik yang menarik untuk dianalisis, khususnya perspektif opini publik (Kementerian Perindustrian Republik Indonesia, 2025).

Dalam era digital, analisis sentimen menjadi alat yang penting untuk memahami opini masyarakat terhadap kebijakan dan produk. Sentiment analysis atau opinion mining adalah cabang dari pemrosesan bahasa alami (*Natural Processing Language*, NLP) yang bertujuan untuk mengekstrak, mengidentifikasi, dan mengklasifikasi opini dalam sebuah teks (Liu, 2012). Dengan meningkatnya diskusi publik mengenai kebijakan TKDN dan strategi bisnis Apple, analisis sentimen dapat memberikan wawasan tentang bagaimana masyarakat Indonesia merespons kebijakan tersebut.

Untuk menganalisis sentimen publik terhadap kebijakan TKDN Apple, penelitian ini menggunakan data yang diperoleh dari media sosial X (Twitter). Platform ini dipilih karena menjadi salah satu sumber utama diskusi publik di Indonesia, dengan berbagai opini yang mencerminkan persepsi masyarakat terhadap kebijakan tersebut. Data dikumpulkan melalui teknik *crawling*, yang mencakup berbagai parameter seperti teks cuitan (*full text*), jumlah suka (*likes*), jumlah balasan (*replies*), dan jumlah ulang cuit (*retweets*).

Untuk melakukan analisis sentimen secara efektif, diperlukan berbasis pembelajaran mesin (*machine learning*). Salah satu algoritma yang banyak digunakan dalam klasifikasi teks adalah random forest, yang dikembangkan oleh Leo Breiman. Random forest merupakan metode *ensemble* yang terdiri dari banyak pohon keputusan (*decision trees*) untuk meningkatkan akurasi dan mengurangi *overfitting* (Breiman, 2001). Algoritma ini sangat efektif dalam menangani dataset yang kompleks atau sering digunakan dalam berbagai tugas klasifikasi, termasuk analisis sentimen.

Penelitian ini bertujuan untuk menganalisis sentimen publik terkait kebijakan TKDN Apple di Indonesia dengan menggunakan metode random forest. Dataset yang digunakan diperoleh dari media sosial X, yang kemudian diproses melalui tahapan *preprocessing*, tokenisasi dan klasifikasi. Hasil dari penelitian ini diharapkan dapat memberikan gambaran tentang persepsi masyarakat terhadap kebijakan TKDN serta efektivitas metode random forest dalam analisis sentimen.

1.2. Tinjauan Pustaka

Menurut Manurung & Mayatopani (2025), perkembangan teknologi smartphone, terutama dari merek global Apple, selalu menarik perhatian masyarakat dunia, termasuk Indonesia. Penelitian mereka menganalisis sentimen masyarakat Indonesia terhadap peluncuran iPhone 16 dengan menggunakan algoritma naive bayes, random forest, dan knn. Data yang digunakan berasal dari tweet masyarakat Indonesia di platform Twitter (sekarang X) antara Oktober dan November 2024. Hasil penelitian menunjukkan bahwa 51,49% tweet memiliki sentimen positif, 28,15% netral, dan 20,35% negatif. Algoritma random forest mencatat akurasi tertinggi sebesar 72,4% diikuti oleh KNN 66,9% dan naive bayes 66,3%. Penelitian ini menyimpulkan bahwa mayoritas masyarakat Indonesia memiliki sentimen positif terhadap peluncuran iPhone 16, dan random forest merupakan algoritma yang paling efektif untuk klasifikasi sentimen.

Selanjutnya Syahrohim et al. (2024), melakukan analisis sentimen terhadap opini masyarakat di Twitter setelah Pemilihan Presiden 2024 menggunakan algoritma naive bayes, support vector machine (svm), dan logistic regression. Data yang digunakan terdiri dari 4.260 tweet terkait tiga calon presiden: Ganjar Pranowo, Anies Baswedan, dan Prabowo Subianto. Hasil dari penelitian menunjukkan bahwa logistic regression memiliki performa terbaik dengan akurasi 84,39% pada dataset Ganjar Pranowo. Penelitian ini memberikan wawasan tentang perbandingan performa algoritma klasifikasi dalam analisis sentimen di media sosial, yang dapat menjadi referensi untuk penelitian serupa di masa depan.

Selanjutnya Ismail Akbar & Muhammad Faisal (2024), membandingkan analisis sentimen terhadap aplikasi PLN Mobile menggunakan algoritma *machine learning* (logistic regression, decision tree, dan random forest) dan *deep learning* (*neural network multi-layer perceptron*/MLP dan *long short-term memory*/LSTM). Data yang digunakan terdiri dari 3.000 ulasan pengguna, dengan 1.965 ulasan positif, dan 1.035 ulasan negatif. Hasil penelitian menunjukkan bahwa logistic regression dan MLP memiliki akurasi tertinggi, masing-masing sebesar 84,47%. Penelitian ini menyimpulkan bahwa kedua algoritma tersebut lebih unggul dalam mengklasifikasi sentimen dibandingkan dengan algoritma lainnya.

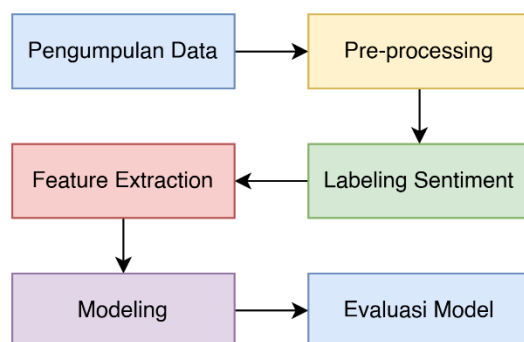
Selanjutnya Ferdian Maulana Akbar et al. (2024), melakukan analisis sentimen terkait isu rangka motor XYZ di Twitter menggunakan lima algoritma *machine learning*: naive bayes, decision tree, svm, logistic regression, dan random forest. Data diambil dari tweet periode Agustus hingga November 2023. Hasil penelitian menunjukkan bahwa random forest memiliki performa terbaik dengan f1-score sebesar 0,765. Penelitian ini merekomendasikan peningkatan kesadaran tentang pemeriksaan rangka gratis untuk memperbaiki reputasi merek XYZ di mata masyarakat.

Kemudian Fauzan et al. (2024), melakukan analisis sentimen terhadap aplikasi MyTelkomsel menggunakan data ulasan dari Google Play Store. Penelitian ini membandingkan dua metode pelabelan, yaitu *lexicon based* dan pakar bahasa Indonesia, serta menggunakan algoritma svm dan random forest. Hasil penelitian menunjukkan bahwa pelabelan oleh pakar lebih akurat 83% dibanding *lexicon based* 79%. Algoritma svm mencatat akurasi tertinggi sebesar 83% pada dataset yang dianalisis oleh pakar. Penelitian ini memberikan wawasan berharga bagi pengembang aplikasi dalam meningkatkan kualitas layanan berdasarkan analisis sentimen pengguna.

1.3. Metodologi Penelitian

Penelitian ini bertujuan untuk menganalisis sentimen masyarakat terhadap penjualan iPhone 16 di Indonesia dalam konteks kebijakan Tingkat Komponen Dalam Negeri (TKDN). Penelitian ini menggunakan pendekatan *supervised machine learning* dengan algoritma *Random Forest Classification* untuk melakukan klasifikasi sentimen ke dalam tiga kategori: positif, netral, dan negatif.

Tahapan penelitian ini mencakup beberapa langkah utama, yaitu pengumpulan data dari platform X, proses preprocessing untuk membersihkan data, labeling sentimen secara otomatis, ekstraksi fitur menggunakan metode *Term Frequency-Inverse Document Frequency* (TF-IDF), pelatihan model klasifikasi dengan Random Forest, serta evaluasi model berdasarkan metrik *accuracy*, *precision*, *recall*, dan *F1-score*. Seperti pada gambar nomor 1.



Gambar 1. Model penelitian

Berikut penjelasan setiap tahapannya:

a. Pengumpulan data

Data yang digunakan dalam penelitian ini diperoleh melalui proses *crawling* dari platform X. Data ini berisi berbagai atribut seperti teks *tweet*, tanggal unggahan, jumlah *retweet*, jumlah *like*, dan lainnya. Namun, fokus utama analisis berada pada atribut teks *tweet* (*full_text*) yang berisi opini pengguna. Proses pengumpulan dilakukan menggunakan pustaka Python, dengan kata kunci yang relevan seperti “Apple”, “iPhone 16” dan “TKDN”. Data yang berhasil dikumpulkan kemudian disimpan dalam format CSV untuk proses analisis lebih lanjut (Rizkina & Hasan, 2023).

b. Pre-processing

Preprocessing data merupakan tahap pemrosesan awal yang bertujuan untuk menyusun ulang data mentah menjadi lebih terstruktur dan siap digunakan dalam analisis (Indarwati & Februriyanti, 2023). Proses ini mencakup eliminasi informasi yang tidak relevan serta penyesuaian format data agar lebih mudah diproses. Dalam konteks analisis sentimen, terutama di media sosial, *preprocessing* menjadi tahap yang sangat penting karena teks yang digunakan seringkali bersifat tidak terstruktur, mengandung elemen tidak perlu (*noise*), serta menggunakan bahasa informal (Widia et al., 2024). Adapun tahapan *preprocessing* yang dilakukan dalam penelitian ini adalah sebagai berikut:

a). Case folding

Seluruh teks diubah menjadi huruf kecil untuk memastikan konsistensi dan menghindari perbedaan makna akibat perbedaan huruf besar dan kecil (Khairani et al., 2024).

b). Cleaning text

Proses ini meliputi penghapusan URL, *mention*, *hashtag*, angka, emoji, dan simbol yang tidak relevan. Selain itu, teks non-alfabet juga dihapus untuk meningkatkan kualitas data (Khairani et al., 2024).

c). Stopword Removal

Kata-kata umum seperti “yang”, “dan”, “di” tidak memiliki nilai informatif dalam analisis sentimen dihapus menggunakan daftar *stopword* bahasa Indonesia (Slamet et al., 2022).

d). Tokenization

Teks dipecah menjadi unit kata untuk memudahkan proses analisis lebih lanjut (Sari et al., 2023).

e). Stemming

Setiap kata dalam teks diubah ke bentuk dasar (*lemma*) menggunakan pustaka Sastrawi untuk mengurangi variasi bentuk kata yang memiliki arti sama (Rizkina & Hasan, 2023).

f). Normalisasi

Kata-kata slang atau tidak baku yang sering muncul di media sosial distandarkan menjadi kata baku untuk meningkatkan akurasi analisis.

g). *Remove Duplicates*

Baris data yang identik dihapus untuk memastikan setiap teks memiliki kontribusi unik dalam analisis.

c. Labeling Sentiment

Setelah data diproses, sentimen setiap teks diberi label berdasarkan makna teksnya. Label yang digunakan adalah positif, netral, dan negatif. Proses labeling dilakukan dengan pemberian nilai *score*. Jika nilai *score* lebih dari 0, maka labelnya positif. Jika nilai *score* sama dengan 0, maka labelnya netral, sedangkan jika nilai *score* kurang dari 0, maka labelnya negatif (Kamelia Cindy Astuti et al., 2024).

d. *Feature extraction* (TF-IDF)

Untuk mengonversi teks ke dalam bentuk numerik yang dapat digunakan oleh model klasifikasi, digunakan teknik TF-IDF (*Term Frequency-Inverse Document Frequency*). Metode ini memiliki dua faktor utama (Septiani & Isabela, 2022).

a). *Term frequency* (TF)

Mengukur seberapa sering sebuah kata muncul dalam suatu dokumen dibandingkan dengan total jumlah kata dalam dokumen tersebut.

b). *Inverse document frequency* (IDF)

Mengukur seberapa unik sebuah kata di seluruh dokumen dalam dataset. TF-IDF menghitung bobot berdasarkan kedua metrik ini, sehingga memberikan fokus lebih pada kata-kata yang penting dan informatif.

e. Klasifikasi menggunakan algoritma random forest

Random Forest adalah metode pembelajaran *ansambel* yang berbasis pada pohon keputusan (*decision tree*). Setiap pohon dalam Random Forest dibangun menggunakan subset dari data pelatihan serta subset dari fitur yang tersedia. Model ini menghasilkan prediksi akhir dengan menggabungkan atau merata-ratakan hasil prediksi dari berbagai pohon keputusan.

Secara matematis, sebuah pohon keputusan membagi ruang fitur secara rekursif ke dalam beberapa wilayah, di mana setiap wilayah diberikan label kelas tertentu. Model matematika untuk sebuah pohon keputusan dapat dirumuskan sebagai berikut:

$$h(x) = \sum_{i=1}^N w_i \times I(x \in R_i) \quad (1)$$

Dimana:

- $h(x)$ adalah hasil prediksi pohon keputusan untuk input x .
- N adalah jumlah wilayah (*leaves*) dalam pohon.
- w_i adalah bobot atau probabilitas kelas yang terkait dengan wilayah R_i .
- $I(x \in R_i)$ adalah fungsi indikator yang bernilai 1 jika x termasuk dalam wilayah R_i , dan 0 jika tidak.

f. Evaluasi model

Model dievaluasi menggunakan beberapa metrik untuk memastikan performanya dalam mengklasifikasikan data sentimen:

a). *Accuracy*

Akurasi adalah rasio antara jumlah prediksi yang benar terhadap total jumlah data yang diuji. Akurasi digunakan untuk mengukur seberapa sering model membuat prediksi yang benar secara keseluruhan. Berikut merupakan rumus dari akurasi:

$$accuracy = \frac{TP + TN}{TP + TN + FP + FN} \times 100\% \quad (2)$$

b). *Precision*

Presisi adalah rasio antara jumlah prediksi positif yang benar dibandingkan dengan semua prediksi positif yang dibuat oleh model. Metrik ini penting ketika kesalahan klasifikasi positif (*False Positive*) harus diminimalkan. Berikut merupakan rumus dari presisi:

$$precision = \frac{TP}{TP + FP} \times 100\% \quad (3)$$

c). *Recall*

Recall mengukur seberapa baik model dalam menemukan semua contoh positif dalam dataset. Nilai *recall* tinggi berarti model dapat mengidentifikasi sebagian besar data positif yang ada. Berikut merupakan rumus dari presisi:

$$recall = \frac{TP}{TP + FN} \times 100\% \quad (4)$$

d). *f1-score*

F1-Score adalah rata-rata harmonik dari *Precision* dan *Recall*, yang digunakan ketika kita ingin menemukan keseimbangan antara kedua metrik tersebut. Berikut merupakan rumus dari *F1-Score*:

$$f1 - score = 2 \times \frac{precision \times recall}{precision + recall}$$

(5)

2. PEMBAHASAN

Pembahasan dalam penelitian ini bertujuan untuk menganalisis dan menginterpretasikan hasil yang diperoleh pada setiap tahapan metodologi. Berikut adalah penjabaran hasil dan analisis dari tahapan-tahapan penelitian:

a. Pengumpulan data

Data yang digunakan dalam penelitian ini berhasil diperoleh melalui *crawling* dari platform X dengan total sebanyak 4.725 tweet. Data tersebut mencakup atribut teks, jumlah *retweet*, jumlah *like*, dan informasi lainnya. Namun, hanya atribut teks (*full_text*) yang relevan digunakan untuk analisis sentimen. Proses *crawling* dilakukan dengan menggunakan kata kunci spesifik seperti "Apple", "iPhone 16" dan "TKDN", sehingga memastikan data yang diperoleh relevan dengan topik penelitian.

Dataset yang terkumpul menunjukkan keberagaman opini dari pengguna platform X, baik berupa kritik, dukungan, maupun netral terhadap kebijakan TKDN dan dampaknya pada penjualan iPhone 16 di Indonesia. Hal ini menunjukkan bahwa diskusi terkait topik ini cukup aktif di media sosial, seperti ditunjukkan pada gambar nomor 2.

conversation_id_str	created_at	favorite_count	full_text	id_str
0	1866988872540029021 Wed Dec 11 23:30:08 +0000 2024	0	Cek Apple iPhone 13 128GB Starlight dengan har...	1866988872540029021
1	1866988803346534841 Wed Dec 11 23:29:52 +0000 2024	0	Cek Apple iPhone 12 128GB Purple dengan harga ...	1866988803346534841
2	1866308492510003300 Wed Dec 11 23:28:32 +0000 2024	0	@coffeecrushhh @AzarineCosmetic Kak kalau yang...	1866308492510003300
3	1866972774822187352 Wed Dec 11 23:28:36 +0000 2024	0	@drafttext Ada aja gebrakan apple ini	1866972774822187352
4	1866972774822187352 Wed Dec 11 23:26:21 +0000 2024	0	@drafttext keren banget nih. fitur konektivitas...	1866972774822187352

Gambar 2. Dataset yang diperoleh

b. Pre-processing

Pre-processing dilakukan untuk membersihkan data teks dari elemen-elemen yang tidak relevan. Berikut adalah hasil tiap langkah *pre-processing*:

a). Case folding

Semua huruf berhasil diubah menjadi huruf kecil. Hal ini memastikan tidak ada perbedaan makna yang disebabkan oleh huruf kapital atau kecil. Contohnya, teks "APPLE adalah merek terkenal" diubah menjadi "apple adalah merek terkenal". Seperti pada gambar nomor 3.

conversation_id_str	created_at	favorite_count	full_text	id_str
0	1866988872540029021 Wed Dec 11 23:30:08 +0000 2024	0	Cek Apple iPhone 13 128GB Starlight dengan har...	1866988872540029021
1	1866988803346534841 Wed Dec 11 23:29:52 +0000 2024	0	Cek Apple iPhone 12 128GB Purple dengan harga ...	1866988803346534841
2	1866308492510003300 Wed Dec 11 23:28:32 +0000 2024	0	@coffeecrushhh @AzarineCosmetic Kak kalau yang...	1866308492510003300
3	1866972774822187352 Wed Dec 11 23:28:36 +0000 2024	0	@drafttext Ada aja gebrakan apple ini	1866972774822187352
4	1866972774822187352 Wed Dec 11 23:26:21 +0000 2024	0	@drafttext keren banget nih. fitur konektivitas...	1866972774822187352

Gambar 3. Hasil case folding

b). Cleaning text

Proses ini berhasil menghapus URL, *mention*, *hashtag*, angka, emoji, dan simbol. Selain itu, teks non-alfabet dihilangkan sehingga lebih bersih dan fokus ke isi utama. Seperti pada gambar nomor 4.

conversation_id_str	created_at	favorite_count	full_text	id_str
0	1866988872540029021 Wed Dec 11 23:30:08 +0000 2024	0	Cek Apple iPhone 13 128GB Starlight dengan har...	1866988872540029021
1	1866988803346534841 Wed Dec 11 23:29:52 +0000 2024	0	Cek Apple iPhone 12 128GB Purple dengan harga ...	1866988803346534841
2	1866308492510003300 Wed Dec 11 23:28:32 +0000 2024	0	@coffeecrushhh @AzarineCosmetic Kak kalau yang...	1866308492510003300
3	1866972774822187352 Wed Dec 11 23:28:36 +0000 2024	0	@drafttext Ada aja gebrakan apple ini	1866972774822187352
4	1866972774822187352 Wed Dec 11 23:26:21 +0000 2024	0	@drafttext keren banget nih. fitur konektivitas...	1866972774822187352

Gambar 4. Hasil cleaning text

c). Stopword Removal

Kata-kata umum yang tidak relevan dihapus untuk meningkatkan akurasi analisis sentimen. Seperti pada gambar nomor 5.

conversation_id_str	created_at	favorite_count	full_text	id_str
0	1866988872540029021 Wed Dec 11 23:30:08 +0000 2024	0	Cek Apple iPhone 13 128GB Starlight dengan har...	1866988872540029021
1	1866988803346534841 Wed Dec 11 23:29:52 +0000 2024	0	Cek Apple iPhone 12 128GB Purple dengan harga ...	1866988803346534841
2	1866308492510003300 Wed Dec 11 23:28:32 +0000 2024	0	@coffeecrushhh @AzarineCosmetic Kak kalau yang...	1866308492510003300
3	1866972774822187352 Wed Dec 11 23:28:36 +0000 2024	0	@drafttext Ada aja gebrakan apple ini	1866972774822187352
4	1866972774822187352 Wed Dec 11 23:26:21 +0000 2024	0	@drafttext keren banget nih. fitur konektivitas...	1866972774822187352

Gambar 5. Hasil stopwords removal

d). Tokenization

Teks dipecah menjadi unit kata agar lebih mudah dianalisis. Misalnya, "apple adalah merek terkenal" diubah menjadi ["apple", "adalah", "merek", "terkenal"]. Seperti pada gambar nomor 6.

conversation_id_str	created_at	favorite_count	full_text	id_str
0	1866988872540029021 Wed Dec 11 23:30:08 +0000 2024	0	Cek Apple iPhone 13 128GB Starlight dengan har...	1866988872540029021
1	1866988803346534841 Wed Dec 11 23:29:52 +0000 2024	0	Cek Apple iPhone 12 128GB Purple dengan harga ...	1866988803346534841
2	1866308492510003300 Wed Dec 11 23:28:32 +0000 2024	0	@coffeecrushhh @AzarineCosmetic Kak kalau yang...	1866308492510003300
3	1866972774822187352 Wed Dec 11 23:28:36 +0000 2024	0	@drafttext Ada aja gebrakan apple ini	1866972774822187352
4	1866972774822187352 Wed Dec 11 23:26:21 +0000 2024	0	@drafttext keren banget nih. fitur konektivitas...	1866972774822187352

Gambar 6. Hasil tokenisasi

e). Stemming

Kata-kata dikonversi ke bentuk dasar menggunakan pustaka Sastrawi. Contohnya, kata "terkenalnya" diubah menjadi "terkenal". Seperti pada gambar nomor 7.

```

Setelah Stemming:
                                stopwords_removed \
0 cek apple iphone gb starlight harga rp dapatka...
1 cek apple iphone gb purple harga rp dapatkan s...
2          kak warnanya rada pink apple ya
3          aja gebrakan apple
4 keren banget nih fitur konektivitas satelit ap...

                                stemmed
0 cek apple iphone gb starlight harga rp dapat s...
1 cek apple iphone gb purple harga rp dapat shop...
2          kak warna rada pink apple ya
3          aja gebrak apple
4 keren banget nih fitur konektivitas satelit ap...

```

Gambar 7. Hasil *stemming*

f). Normalisasi

Kata-kata slang dikonversi ke kata baku agar lebih seragam. Seperti pada gambar nomor 8.

```

Setelah Normalisasi:
                                stemmed \
0 cek apple iphone gb starlight harga rp dapat s...
1 cek apple iphone gb purple harga rp dapat shop...
2          kak warna rada pink apple ya
3          aja gebrak apple
4 keren banget nih fitur konektivitas satelit ap...

                                normalized
0 cek apple iphone gb starlight harga rp dapat s...
1 cek apple iphone gb purple harga rp dapat shop...
2          kak warna rada pink apple ya
3          saja gebrak apple
4 keren banget ini fitur konektivitas satelit ap...

```

Gambar 8. Hasil *normalisasi*

g). Remove Duplicates

Data duplikat dihapus, sehingga dataset yang digunakan lebih unik dan tidak bias. Dari 4.725 data awal, sebanyak 1.089 data duplikat dihapus, menyisakan 3.636 data unik. Seperti pada gambar nomor 9.

```

Setelah Remove Duplicates:
                                normalized
0 cek apple iphone gb starlight harga rp dapat s...
1 cek apple iphone gb purple harga rp dapat shop...
2          kak warna rada pink apple ya
3          saja gebrak apple
4 keren banget ini fitur konektivitas satelit ap...

```

Gambar 9. Hapus duplikat

c. Labeling Sentiment

Proses pelabelan sentimen menghasilkan tiga kategori: positif, negatif, dan netral. Ini dilakukan dengan otomatis menggunakan Sebanyak 3.045 *tweet* data tergolong netral, 398 *tweet* negatif, dan 193 *tweet* positif. Seperti pada gambar nomor 10.

```

Hasil Labeling Sentimen:
                                cleaned_text sentiment
0 cek apple iphone gb starlight harga rp dapat s... netral
1 cek apple iphone gb purple harga rp dapat shop... netral
2          kak warna rada pink apple ya              negatif
3          saja gebrak apple                          netral
4 keren banget ini fitur konektivitas satelit ap... netral

```

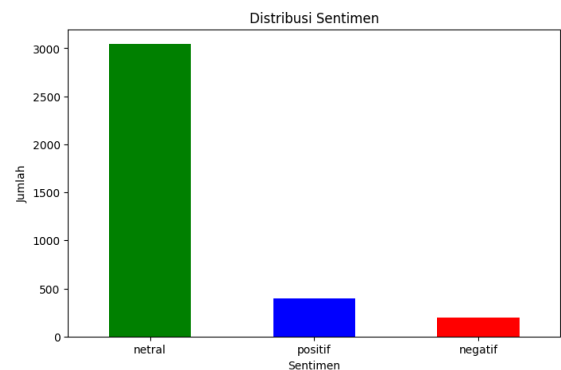
Gambar 10. Labeling sentimen

Distribusi ini menunjukkan bahwa sebagian besar opini pengguna terhadap kebijakan TKDN bersifat netral. Hal ini dapat mengindikasikan kurangnya ketertarikan atau pemahaman yang cukup mengenai kebijakan ini. Seperti pada gambar nomor 11.

Gambar 11. Distribusi *labeling sentimen*

d. Feature extraction (TF-IDF)

Pada tahap ini, metode *Term Frequency - Inverse Document Frequency* (TF-IDF) digunakan untuk



mengubah teks menjadi bentuk numerik yang dapat diproses oleh model klasifikasi. Parameter yang digunakan dalam TF-IDF meliputi *max_features=5000*, yang hanya mempertimbangkan 5.000 kata atau frasa yang paling sering muncul dalam dataset, serta *ngram_range=(1,2)* yang mempertimbangkan unigram dan bigram. Selain itu, *stopword* dalam bahasa Inggris dihapus untuk meningkatkan relevansi fitur yang digunakan dalam analisis.

Hasil ekstraksi menghasilkan matriks berukuran (3636, 5000), yang berarti terdapat 3.636 *tweet* yang direpresentasikan oleh 5.000 fitur numerik. Matriks ini kemudian digunakan sebagai input untuk model klasifikasi dalam menentukan sentimen dari teks. Seperti pada gambar nomor 12.

TF-IDF Shape: (3636, 5000)

Gambar 12. Hasil TF-IDF

e. Klasifikasi menggunakan algoritma random forest

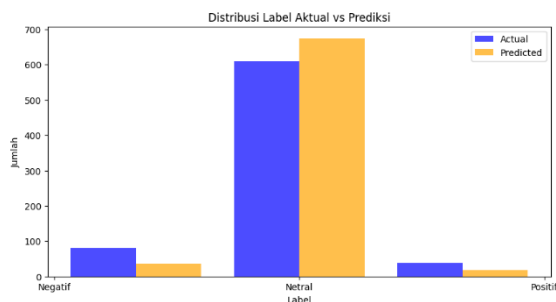
Model klasifikasi dalam penelitian ini menggunakan algoritma random forest dengan pembagian data sebesar 80% untuk *training* (2.908 data) dan 20% untuk *testing* (728 data). Model dilatih menggunakan parameter *default*, dan hasil pelatihan menunjukkan bahwa random forest memiliki kemampuan yang baik dalam menangkap pola data dalam analisis sentimen ini. Seperti pada gambar 13.

Hasil Evaluasi Model:				
Accuracy: 0.9052197802197802				
Classification Report:				
	precision	recall	f1-score	support
negatif	0.88	0.39	0.55	38
netral	0.90	1.00	0.95	610
positif	0.95	0.44	0.60	80
accuracy			0.91	728
macro avg	0.91	0.61	0.70	728
weighted avg	0.91	0.91	0.89	728

Gambar 13. Hasil klasifikasi random forest

Berdasarkan hasil pengujian, model mencapai akurasi sebesar 91%, yang menunjukkan bahwa algoritma ini mampu mengklasifikasikan sebagian besar data dengan benar. Namun, ketika melihat distribusi prediksi pada masing-masing kelas sentimen, terdapat kecenderungan model untuk lebih sering mengklasifikasikan sentimen sebagai netral, sedangkan kategori negatif dan positif cenderung

kurang terprediksi. Distribusi hasil prediksi dibandingkan data aktual dapat dilihat pada gambar 14.

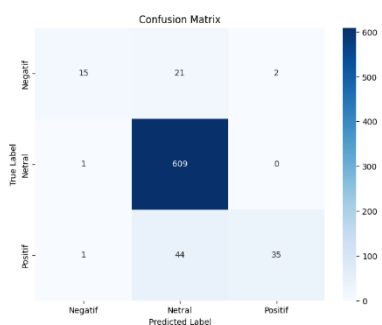


Gambar 14. Distribusi label aktual dan prediksi

Dominasi sentimen netral dalam hasil prediksi menunjukkan bahwa model cenderung memilih kategori mayoritas. Sebagai contoh, dari total 3.636 *tweet*, model memprediksi lebih dari 80% data sebagai netral, meskipun pada data aktual jumlah kategori netral lebih rendah dari prediksi tersebut. Hal ini mengindikasikan adanya ketidakseimbangan kelas, yang dapat mempengaruhi kinerja model dalam mengenali sentimen negatif dan positif. Untuk memahami performa model secara lebih mendalam, evaluasi lebih lanjut dilakukan menggunakan *confusion matrix* serta metrik *akurasi*, *precision*, *recall*, dan *f1-score*.

f. Evaluasi model

Evaluasi model dilakukan dengan menggunakan *confusion matrix* dan metrik performa lainnya. Berdasarkan hasil evaluasi, model random forest menunjukkan kemampuan tinggi dalam mengklasifikasikan data netral dengan benar, tetapi masih memiliki kesalahan prediksi pada kategori negatif dan positif. Dari total 728 data uji, model berhasil mengklasifikasikan 15 data negatif, 609 data netral, dan 35 data positif dengan benar. Namun, model juga melakukan kesalahan klasifikasi, di mana 21 data negatif dan 44 data positif diprediksi sebagai netral, serta terdapat 2 data negatif yang diprediksi sebagai positif. Selain itu, terdapat 1 data netral yang diklasifikasikan sebagai negatif, serta 1 data positif yang salah diklasifikasikan sebagai negatif. Seperti pada gambar nomor 15.



Gambar 15. Tabel *confusion matrix*

Metrik evaluasi model menunjukkan akurasi sebesar 91%, yang berarti model memiliki tingkat kesalahan relatif rendah dalam klasifikasi keseluruhan. Namun, akurasi yang tinggi ini perlu diperjelas dengan metrik lain untuk memastikan keseimbangan antar kelas sentimen. *Precision* yang diperoleh juga mencapai 91%, yang menunjukkan bahwa sebagian besar prediksi model sesuai dengan label aslinya. *Recall* sebesar 91% menunjukkan bahwa model cukup baik dalam mengenali data dari setiap kategori sentimen. Meskipun demikian, *F1-score* sebesar 89% mengindikasikan adanya ketidakseimbangan dalam prediksi model terhadap kelas-kelas yang lebih kecil, seperti sentimen negatif dan positif. Seperti pada gambar 16.

```

Metrics Evaluasi Model:
Accuracy: 0.91
Precision: 0.91
Recall: 0.91
F1 Score: 0.89

```

Gambar 16. Hasil metrik evaluasi

Untuk meningkatkan kinerja model, khususnya dalam mengenali sentimen negatif dan positif, dapat diterapkan beberapa teknik seperti *oversampling* pada kelas minoritas (negatif dan positif), *undersampling* pada kelas mayoritas (netral), atau menggunakan metode *cost-sensitive learning*. Dengan pendekatan ini, model diharapkan dapat lebih sensitif terhadap variasi sentimen dalam data, sehingga hasil klasifikasi menjadi lebih seimbang dan akurat.

3. KESIMPULAN

Penelitian ini telah berhasil menerapkan algoritma random forest untuk analisis sentimen terkait penjualan iPhone 16 dan kebijakan TKDN. Hasil evaluasi menunjukkan bahwa model memiliki akurasi sebesar 91%, dengan *precision* 91%, *recall* 91%, dan *f1-score* 89%. Metrik ini menunjukkan bahwa model memiliki performa yang baik dalam mengklasifikasikan sentimen. Namun, berdasarkan *confusion matrix*, model cenderung lebih sering mengklasifikasikan sentimen sebagai netral. Hal ini mengindikasikan adanya ketidakseimbangan kelas dalam dataset, yang berdampak pada rendahnya prediksi untuk kategori negatif dan positif.

Meskipun model menunjukkan kinerja yang cukup baik, masih terdapat kesalahan klasifikasi, terutama dalam mendeteksi sentimen negatif dan positif. Untuk meningkatkan performa, penelitian selanjutnya dapat menerapkan teknik penyeimbangan kelas, seperti *oversampling* atau *cost-sensitive learning*, guna mengatasi ketidakseimbangan data. Selain itu, eksplorasi model yang lebih kompleks, seperti *deep learning* atau *ensemble learning*, berpotensi meningkatkan akurasi prediksi. Pendekatan lain seperti pemanfaatan *word embeddings* dan sentimen *lexicon*

juga dapat dieksplorasi untuk meningkatkan kualitas representasi teks, sehingga menghasilkan klasifikasi yang lebih akurat.

PUSTAKA

- Breiman, L. (2001). Random Forest. *Machine Learning*, 45(1), 5–32. <https://doi.org/10.1023/A:1010933404324>
- Fauzan, F. J., M Afdal, Rice Novita, & Mustakim. (2024). PENERAPAN MACHINE LEARNING PADA ANALISIS SENTIMEN APLIKASI MYTELKOMSEL MENGGUNAKAN DATA ULASAN GOOGLE PLAYSTORE. *Indonesian Journal of Computer Science*, 13(3). <https://doi.org/10.33022/ijcs.v13i3.4024>
- Ferdian Maulana Akbar, Robby Hermansyah, Sofian Lusa, Dana Indra Sensuse, Nadya Safitri, & Damayanti Elisabeth. (2024). Analisis Sentimen untuk Evaluasi Reputasi Merek Motor XYZ Berkaitan dengan Isu Rangka Motor di Twitter Menggunakan Pendekatan <i>Machine Learning</i>. *Jurnal Teknologi Informasi Dan Ilmu Komputer*, 11(3), 647–654. <https://doi.org/10.25126/jtiik.938663>
- Indarwati, K. D., & Februariyanti, H. (2023). Analisis Sentimen Terhadap Kualitas Pelayanan Aplikasi Gojek Menggunakan Metode Naive Bayes Classifier. *JATISI (Jurnal Teknik Informatika Dan Sistem Informasi)*, 10(1). <https://doi.org/10.35957/jatisi.v10i1.2643>
- Ismail Akbar, & Muhammad Faisal. (2024). Perbandingan Analisis Sentimen PLN Mobile: Machine Learning vs. Deep Learning. *JOINTECS (Journal of Information Technology and Computer Science)*, 8(1), 1–10.
- Kamelia Cindy Astuti, Andri Firmansyah, & Agus Riyadi. (2024). Implementasi Text Mining Untuk Analisis Sentimen Masyarakat Terhadap Ulasan Aplikasi Digital Korlantas Polri pada Google Play Store. *Remik : Riset Dan E-Jurnal Manajemen Informatika Komputer*, 8(1), 383–394.
- Kementerian Perindustrian Republik Indonesia. (2025, January 9). *Pernyataan Menperin terkait TKDN Apple*. Kementerian Perindustrian Republik Indonesia.
- Khairani, U., Mutiawani, V., & Ahmadian, H. (2024). Pengaruh Tahapan Preprocessing Terhadap Model Indobert Dan Indobertweet Untuk Mendeteksi Emosi Pada Komentar Akun Berita Instagram. *Jurnal Teknologi Informasi Dan Ilmu Komputer*, 11(4), 887–894. <https://doi.org/10.25126/jtiik.1148315>
- Liu, B. (2012). *Sentiment Analysis and Opinion Mining*. Springer International Publishing. <https://doi.org/10.1007/978-3-031-02145-9>
- Manurung, C. E., & Mayatopani, H. (2025). Sentiment Analysis of Indonesian Society Toward the Launch of iPhone 16 Using Naive Bayes, Random Forest, and KNN Algorithms. *Jurnal Komputer, Informasi Dan Teknologi*, 5(1).
- Rizkina, N. Q., & Hasan, F. N. (2023). Analisis Sentimen Komentar Netizen Terhadap Pembubaran Konser NCT 127 Menggunakan Metode Naive Bayes. *Journal of Information System Research (JOSH)*, 4(4), 1136–1144.
- Sari, T. A., Sinduningrum, E., & Hasan, F. N. (2023). Analisis Sentimen Ulasan Pelanggan Pada Aplikasi Fore Coffee Menggunakan Metode Naive Bayes. *KLIK Kaji. Ilm. Inform. Dan Komput*, 3(6), 773–779.
- Septiani, D., & Isabela, I. (2022). Analisis term frequency inverse document frequency (tf-idf) dalam temu kembali informasi pada dokumen teks. *SINTESIA J. Sist. Dan Teknol. Inf. Indones*, 1(2), 81–88.
- Slamet, R., Gata, W., Novtariany, A., Hilyati, K., & Jariyah, F. A. (2022). Analisis sentimen Twitter terhadap penggunaan artis Korea Selatan sebagai brand ambassador produk kecantikan lokal. *INTECOMS J. Inf. Technol. Comput. Sci*, 5(1), 145–153.
- Syahrohim, I., Saputra, S. D., Saputra, R. W., Pranatawijaya, V. H., & Priskila, R. (2024). PERBANDINGAN ANALISIS SENTIMEN SETELAH PILPRES 2024 DI TWITTER MENGGUNAKAN ALGORITMA MACHINE LEARNING. *Jurnal Informatika Dan Teknik Elektro Terapan*, 12(2). <https://doi.org/10.23960/jitet.v12i2.4249>
- Widia, Zalfa Yunda Aqsalia, Septiana Sari, Nabila Umi Khoirunisa, & Fandi Kurniawan. (2024). Optimasi Algoritma Naive Bayes Untuk Menganalisis Sentimen Pada Konten Pemindahan Ibu Kota di Youtube. *Journal of Computer and Information Systems Ampera*, 5(2), 68–83.