

ANALISIS SENTIMEN MASYARAKAT PADA MEDIA SOSIAL PLATFORM X TERHADAP PERTAMINA MENGGUNAKAN METODE DECISION TREE

Yksan Naibaho¹, Darwis Robinson Manalu², Samuel VB H. Manurung³

^{1,2,3}Fakultas Ilmu Komputer, Universitas Methodist Indonesia

¹yksannaibaho@gmail.com, ²manaludarwis@gmail.com

ABSTRACT

Platform X social media has become a platform for people to express their opinions, including on issues related to Pertamina. This study aims to analyze public sentiment on Platform X towards Pertamina using the ID3 (Iterative Dichotomiser 3) algorithm-based Decision Tree method. The data used are 2,005 tweets collected with the keyword "Pertamina Corruption". The data went through a preprocessing stage which includes case folding, tokenizing, stopword removal, and stemming. Text features were converted into binary representations (Binary Weighting) of 10 main keywords such as 'corruption', 'pertamina', and 'prosecutor'. The ID3 Decision Tree model was built recursively by selecting separator attributes based on the highest Information Gain value. The results showed that the built model had excellent performance. Evaluation on testing data (20% of the total data) produced an accuracy of 98.50%, with a precision value of 97.98%, a recall of 98.50%, and an F1-score of 98.22%. The 5-fold cross-validation results also confirmed the model's stability, with an average accuracy of 98.30% and a low standard deviation (0.0046). The contains_korupsi attribute was identified as the most informative root node in the decision tree structure. The conclusion of this study is that the Decision Tree method with the ID3 algorithm has proven effective and reliable in classifying public sentiment toward Pertamina on Platform X with high accuracy. The results of this analysis are expected to be used by Pertamina in understanding public opinion and formulating more appropriate communication strategies.

Keywords: Sentiment Analysis, Decision Tree, Pertamina, Platform X.

1. PENDAHULUAN

Di era digital yang ditandai dengan pertumbuhan eksponensial media sosial, Platform X (sebelumnya Twitter) telah menjadi barometer opini publik yang signifikan [1]. Platform ini tidak hanya berfungsi sebagai ruang komunikasi sosial, tetapi juga telah berkembang menjadi arena diskusi publik yang kritis terhadap berbagai isu strategis nasional, termasuk kinerja Badan Usaha Milik Negara (BUMN). Sebagai perusahaan energi terbesar di Indonesia, Pertamina kerap menjadi sorotan publik, terutama terkait dengan isu-isu korupsi yang muncul secara periodik.

Pada awal tahun 2025, Indonesia dikejutkan oleh pengungkapan kasus korupsi besar di lingkungan Pertamina. Kejaksan Agung menetapkan beberapa pejabat anak perusahaan Pertamina sebagai tersangka dalam kasus dugaan korupsi tata kelola minyak mentah dan produk kilang, dengan kerugian negara yang ditaksir mencapai Rp193,7 triliun [2]. Selain itu, masyarakat juga dihebohkan oleh kasus dugaan peredaran Pertamina oplosan yang tersebar di beberapa SPBU dan menjadi pembahasan hangat di berbagai media sosial [3]. Dua peristiwa ini memicu gelombang diskusi masif di Platform X, dengan ribuan tweet yang merefleksikan beragam sentimen masyarakat.

Dalam konteks komunikasi korporat di era digital, sentimen publik di media sosial memiliki pengaruh signifikan terhadap reputasi dan legitimasi organisasi. Tanggapan negatif yang berkembang di Platform X berpotensi mempengaruhi persepsi masyarakat, yang pada akhirnya dapat berdampak pada kepercayaan publik terhadap Pertamina sebagai BUMN strategis. Oleh karena itu, penting bagi manajemen

Pertamina untuk secara sistematis menganalisis dan memahami pola sentimen masyarakat guna merumuskan strategi komunikasi yang efektif dan responsif.

Metode Decision Tree dengan algoritma ID3 (Iterative Dichotomiser 3) dipilih dalam penelitian ini karena beberapa pertimbangan strategis. Pertama, metode ini telah terbukti efektif dalam penelitian sebelumnya untuk klasifikasi teks dengan akurasi tinggi [4]. Kedua, Decision Tree memberikan keunggulan dalam interpretasi hasil yang transparan melalui struktur pohon yang mudah dipahami. Ketiga, algoritma ID3 secara khusus dirancang untuk menangani atribut kategorikal, sehingga cocok untuk analisis fitur biner dalam data teks. Penelitian ini bertujuan mengisi celah literatur dengan menerapkan metode Decision Tree ID3 untuk analisis sentimen terhadap isu korupsi di BUMN sektor energi. Dengan memanfaatkan data real-time dari Platform X, penelitian ini diharapkan dapat memberikan kontribusi metodologis dalam analisis sentimen berbahasa Indonesia serta kontribusi praktis bagi manajemen komunikasi korporat Pertamina

2. METODE PENELITIAN

2.1 Text Mining

Text mining adalah proses mengeksplorasi dan menganalisis sejumlah besar data teks tidak terstruktur yang dibantu oleh perangkat lunak yang dapat mengidentifikasi konsep, pola, topik, kata kunci, dan atribut lainnya dalam data [5]. Text Mining Merupakan metode yang digunakan untuk menangani permasalahan klasifikasi, information retrieval, information extraction serta clustering. Pada umumnya proses kerja dari Text

Mining banyak mengadopsi dari penelitian data mining namun yang menjadi perbedaan adalah pola yang digunakan oleh Text Mining diambil dari sekumpulan bahasa alami yang tidak terstruktur sedangkan dalam data mining pola yang diambil dari database yang terstruktur. Tahapan algoritma Text Mining secara umum adalah praproses teks [6].

2.1.1 Preprocessing Text

Preprocessing teks merupakan tahapan awal dalam pengolahan data teks yang bertujuan untuk membersihkan dan menyiapkan teks mentah agar dapat dianalisis lebih efektif dalam pemrosesan bahasa alami (Natural Language Processing/NLP). Proses ini dilakukan dengan menghilangkan data yang tidak relevan serta mengubah teks ke dalam format yang lebih terstruktur sehingga meningkatkan kualitas data dan mengurangi noise (Sipayung et al., 2024). Tahapan preprocessing teks meliputi:

1. Cleansing, yaitu menghapus karakter yang tidak relevan seperti simbol, angka, dan tautan.
2. Case Folding, yaitu mengubah seluruh huruf dalam dokumen menjadi huruf kecil.
3. Tokenizing, yaitu memecah teks menjadi unit kata.
4. Filtering, yaitu memilih kata-kata penting dari hasil tokenisasi.
5. Stemming, yaitu mengubah kata berimbuhan menjadi kata dasar untuk menyamakan makna kata yang sejenis.

2.2 Decision Tree

Decision Tree merupakan metode klasifikasi berbentuk struktur pohon, di mana setiap simpul internal merepresentasikan pengujian terhadap suatu atribut dan simpul daun menunjukkan hasil klasifikasi [7]. Metode ini digunakan untuk mengelompokkan data ke dalam kelas tertentu berdasarkan aturan keputusan yang terbentuk [8]. Pohon keputusan terdiri atas simpul akar, simpul internal, dan simpul daun, yang saling terhubung melalui cabang sebagai alur pengambilan keputusan.

2.2.1 Algoritma ID3

Algoritma ID3 (Iterative Dichotomiser 3) merupakan algoritma dasar dalam pembentukan pohon keputusan dengan pendekatan greedy, yaitu memilih atribut yang paling informatif sebagai pemisah data [9]. Pemilihan atribut didasarkan pada nilai information gain tertinggi yang diperoleh dari perhitungan entropy, sebagai ukuran tingkat ketidakpastian data [10].

$$Entropy(S) = \sum_{i=1}^n -p_i * \log_2 p_i \quad (1)$$

Keterangan:

$Entropy(S)$ = Tingkat ketidakpastian dalam *dataset S*

n = Jumlah total kelas dalam *dataset*

p_i = Proporsi data pada kelas i

Untuk menentukan information gain, digunakan rumus sebagai berikut :

$$Gain(S, A) = Entropy(S) - \sum_{i=1}^n \frac{|S_i|}{|S|} * Entropy(S_i) \quad (2)$$

Keterangan:

$Entropy(S)$ = Tingkat ketidakpastian dalam *dataset S*

n = Jumlah total kelas dalam *dataset*

p_i = Proporsi data pada kelas i

$Gain(S, A)$ = Pengurangan *entropy* setelah *dataset S* dibagi oleh atribut A

n = Jumlah *subset* yang dihasilkan oleh atribut A

$|S_i|$ = Jumlah data dalam *subset S_i*

$|S|$ = Jumlah total data dalam *dataset S*

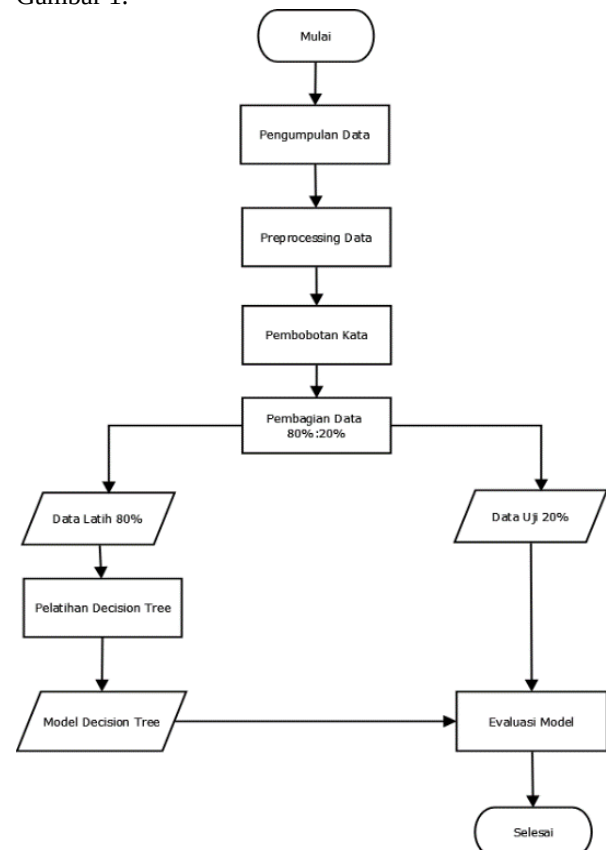
Dalam implementasinya, ID3 membangun pohon keputusan dengan memilih atribut ber-gain tertinggi sebagai simpul akar, kemudian membagi data ke dalam subset sesuai nilai atribut tersebut. Proses ini dilakukan secara berulang hingga seluruh data terklasifikasi atau tidak terdapat atribut yang dapat digunakan, sehingga dihasilkan model pohon keputusan yang dapat digunakan untuk klasifikasi data baru.

2.3 Analisis Sentimen

Analisis sentimen adalah studi komputasi dari opini-opini, sentimen, serta emosi yang diekspresikan dalam teks. Tugas dasar dalam analisis sentimen adalah mengelompokkan polaritas dari teks yang ada dalam dokumen, kalimat, atau pendapat. Analisis sentimen digunakan sebagai gambaran umum media sosial untuk mengetahui pengetahuan perasaan lebih cenderung opini positif atau opini negatif [11].

2.4 Framework Penelitian

Framework penelitian adalah kerangka kerja yang digunakan untuk membimbing jalannya sebuah penelitian agar lebih terarah dan sistematis. Adapun bentuk gambar diagram penelitian dapat dilihat pada Gambar 1.



Gambar 1. *Framework* Penelitian

2.5 Pengumpulan Data

Data yang digunakan dalam penelitian ini dikumpulkan dari platform X (sebelumnya Twitter) dengan menggunakan kata kunci terkait Pertamina. Dari proses pengumpulan awal, diperoleh data mentah sebanyak 2.005 tweet yang kemudian melalui proses cleaning untuk memastikan kualitas data. Hasil preprocessing menunjukkan distribusi sentimen sebagai berikut: negatif 96,01% (1.925 tweet), netral 3,44% (69 tweet), dan positif 0,55% (11 tweet). Distribusi ini mengindikasikan ketidakseimbangan kelas yang perlu ditangani dalam pemodelan. Seperti yang dilampirkan pada Tabel 1.

Tabel 1. Data Tweets

No.	Kalimat	Label
1.	Korupsi Lagi di Pertamina KPK Tahan Anak Eks Direktur dan 2 Tersangka Terkait Gratifikasi Pengadaan Katalis https://t.co/PqNyuG4ach	negatif
2.	18 Tersangka Korupsi Tata Kelola Minyak Pertamina Kejaksaan Periksa Pejabat https://t.co/coT5CdV3E4	negatif
...	...	...
2005	@DS_yantie Pejabat PD rakus.. Listrik masyarakat telat bayar/ga beli token listrik mati masyarakat beli bensin cash Pertamina rugi masyarakat beli tiket pesawat Garuda tp tetep rugi. Aneh negeri ini.. Korupsi semakin subur mulai dr bawah sampai kelas atas.	negatif

3. HASIL DAN PEMBAHASAN

3.1 Preprocessing Data

Melalui penekanan pada tahapan preprocessing, transformasi fitur, algoritma, dan implementasi model yang sepenuhnya disusun menggunakan pendekatan komputasional. Proses seperti case folding, tokenizing, stopword removal, dan stemming diimplementasikan melalui teknik Natural Language Processing (NLP) pada Python, yang merupakan bagian dari metodologi informatika untuk mengolah data teks tidak terstruktur.

1. Case Folding

Preprocessing dilakukan untuk membersihkan teks agar siap digunakan dalam pemodelan, terutama karena data dari media sosial sering mengandung noise. Tahap awal meliputi case folding untuk menyeragamkan huruf menjadi kecil serta cleansing menggunakan regular expression (regex) guna menghapus mention, hashtag, URL, dan karakter spesial melalui perintah `re.sub(r'@\w+|#\w+|http\S+', '', text)`. Hasil dari proses case folding dapat dilihat pada Tabel 2.

Tabel 2. Hasil Case Folding

ID.	Text_Aslai	Text_Setelah_Case_Folding
1.	Korupsi Lagi di Pertamina KPK Tahan Anak Eks Direktur dan 2 Tersangka Terkait Gratifikasi Pengadaan Katalis https://t.co/PqNyuG4ach	korupsi pertamina kpk tahan anak eks direktur sangka kait gratifikasi ada katalis
2.	18 Tersangka Korupsi Tata Kelola Minyak Pertamina Kejaksaan Periksa Pejabat https://t.co/coT5CdV3E4	sangka korupsi tata kelola minyak pertamina jaksa periksa jabat
...	...	...
2005.	@DS_yantie Pejabat PD rakus.. Listrik masyarakat telat bayar/ga beli token listrik mati masyarakat beli bensin cash Pertamina rugi masyarakat beli tiket pesawat Garuda tp tetep rugi. Aneh negeri ini.. Korupsi semakin subur mulai dr bawah sampai kelas atas.	jabat rakus listrik masyarakat telat bayar beli token listrik mati masyarakat beli bensin cash pertamina rugi masyarakat beli tiket pesawat garuda tetep rugi aneh negeri korupsi subur kelas

2. Tokenisasi

Tokenisasi merupakan proses pemecahan teks menjadi unit-unit kata atau token sehingga analisis dapat dilakukan pada level yang lebih detail. Sebelum tahap ini, diterapkan case folding dengan mengonversi seluruh teks menjadi huruf kecil menggunakan fungsi `str.lower()` untuk menstandarkan format data. Hasil dari Tokenisasi dari data tweets dapat dilihat pada Tabel 3.

Tabel 3. Hasil Tokenisasi

ID	Text_Bersih	Token_List
1	korupsi pertamina kpk tahan anak eks direktur sangka kait gratifikasi ada katalis	korupsi, pertamina, kpk, tahan, anak, eks, direktur, sangka, kait, gratifikasi, ada, katalis
2	sangka korupsi tata kelola minyak pertamina jaksa periksa jabat	sangka, korupsi, tata, kelola, minyak, pertamina, jaksa, periksa, jabat
...	...	...

ID	Text_Bersih	Token_List
2005	jabat rakus listrik masyarakat telat bayar beli token listrik mati masyarakat beli bensin cash Pertamina rugi masyarakat beli tiket pesawat Garuda tetap rugi aneh negeri korupsi subur kelas	jabat, rakus, listrik, masyarakat, telat, bayar, beli, token, listrik, mati, masyarakat, beli, bensin, cash, Pertamina, rugi, masyarakat, beli, tiket, pesawat, Garuda, tetap, rugi, aneh, negeri, korupsi, subur, kelas

3. Stopword Removal

Stopword removal dilakukan untuk menghapus kata-kata umum yang tidak memiliki makna signifikan dalam analisis sentimen, seperti “dan”, “atau”, dan “di”, sehingga analisis dapat berfokus pada kata yang lebih informatif. Proses ini dilakukan menggunakan daftar stopwords bahasa Indonesia yang telah dikustomisasi agar lebih sesuai dengan konteks data. Hasil dapat dilihat pada Tabel 4.

Tabel 4. Hasil Stopword Removal

ID	Token_Sebelum	Token_Setelah
1	korupsi, Pertamina, kpk, tahan, anak, eks, direktur, sangka, kait, gratifikasi, ada, katalis	korupsi, Pertamina, kpk, tahan, anak, eks, direktur, sangka, kait, gratifikasi, katalis
2	sangka, korupsi, tata, kelola, minyak, Pertamina, jaksa, periksa, jabat	sangka, korupsi, tata, kelola, minyak, Pertamina, jaksa, periksa, jabat
..	..	..
2005	jabat, rakus, listrik, masyarakat, telat, bayar, beli, token, listrik, mati, masyarakat, beli, bensin, cash, Pertamina, rugi, masyarakat, beli, tiket, pesawat, Garuda, tetap, rugi, aneh, negeri, korupsi, subur, kelas	jabat, rakus, listrik, masyarakat, telat, bayar, beli, token, listrik, mati, masyarakat, beli, bensin, cash, Pertamina, rugi, masyarakat, beli, tiket, pesawat, Garuda, tetap, rugi, aneh, negeri, korupsi, subur, kelas

4. Steemming

Stemming diterapkan untuk mengembalikan kata ke bentuk dasarnya menggunakan algoritma Porter Stemmer dari library Sastrawi. Tahap ini bertujuan mengurangi variasi bentuk kata sehingga representasi teks menjadi lebih sederhana dan konsisten. Hasil dari proses Steeming dapat dilihat pada Tabel 5.

Tabel 5. Hasil Steeming

ID	Token Sebelum Steemming	Token Setelah Steemming
1	korupsi, Pertamina, kpk, tahan, anak, eks, direktur, sangka, kait, gratifikasi, katalis	korupsi, Pertamina, kpk, tahan, anak, eks, direktur, sangka, kait, gratifikasi, katalis
2	sangka, korupsi, tata, kelola, minyak, Pertamina, jaksa, periksa, jabat	sangka, korupsi, tata, kelola, minyak, Pertamina, jaksa, periksa, jabat
...	...	...
2005	jabat, rakus, listrik, masyarakat, telat, bayar, beli, token, listrik, mati, masyarakat, beli, bensin, cash, Pertamina, rugi, masyarakat, beli, tiket, pesawat, Garuda, tetap, rugi, aneh, negeri, korupsi, subur, kelas	jabat, rakus, listrik, masyarakat, telat, bayar, beli, token, listrik, mati, masyarakat, beli, bensin, cash, Pertamina, rugi, masyarakat, beli, tiket, pesawat, Garuda, tetap, rugi, aneh, negeri, korupsi, subur, kelas

3.1.1 Pelabelan Data

Pelabelan sentimen dilakukan secara manual dengan meninjau konteks setiap tweet menggunakan pendekatan context-based annotation. Proses ini dibantu aplikasi Label Studio untuk memastikan konsistensi dan ketepatan anotasi. Sentimen negatif ditentukan berdasarkan kata bernuansa kritik seperti “korupsi”, “rugi”, “mahal”, sentimen positif berdasarkan kata apresiatif seperti “baik” atau “bagus”, sedangkan sentimen netral ditandai oleh kata informatif tanpa muatan opini. Hasil pelabelan ini digunakan sebagai dasar perhitungan entropy, information gain, dan pembentukan pohon keputusan ID3.

3.1.2 Pembobotan Kata

Setelah preprocessing, dilakukan pembobotan kata menggunakan binary weighting untuk mengidentifikasi fitur yang relevan. Analisis dilakukan terhadap 2.005 tweet, menghasilkan sepuluh kata kunci dengan frekuensi tertinggi yang ditampilkan pada Tabel 6.

Tabel 6. Frekuensi Kemunculan 10 Kata Kunci Penting

Kata_Kunci	Frekuensi	Persentase
korupsi	1938	96,66
Pertamina	1909	95,21
sangka	423	21,1
riza	400	19,95
minyak	408	20,35
cholid	397	19,8
jagung	315	15,71
rugi	284	14,16
negara	256	12,77
jaksa	286	14,26

Kesepuluh kata kunci tersebut dijadikan atribut biner, dengan nilai 1 jika kata muncul dalam tweet dan 0 jika tidak. Hasil pembobotan membentuk dataset terstruktur yang siap digunakan dalam algoritma ID3, sebagaimana ditunjukkan pada Tabel 7.

Tabel 7. Hasil Pembobotan Kata

Kalimat	Contains Korupsi	Contains Pertamina	...	Contains Jaksa	Sentimen
1	1	1	...	0	Negatif
2	1	1	...	1	Negatif
3	1	1	...	0	Negatif
...	...	...	...	...	...
4	1	1	...	0	Negatif
5	1	1	...	0	Negatif

3.1.2 Menghitung Entropy Awal (Total Entropy)

Entropy awal dihitung untuk mengukur tingkat ketidakpastian (*impurity*) dalam dataset sebelum dilakukan pemisahan berdasarkan fitur. Perhitungan ini menggunakan distribusi kelas sentimen yang ada dalam dataset, dimana komposisi data sebagai berikut: sentimen negatif sebanyak 1.925 tweet (96,01%), sentimen netral 69 tweet (3,44%), dan sentimen positif 11 tweet (0,55%).

Sebelum melakukan perhitungan information gain, dataset dibagi menjadi data training dan testing dengan rasio 80:20. Hasil pembagian dataset menunjukkan bahwa dari total 2.005 data, sebanyak 1.604 sampel (80%) dialokasikan sebagai data training dan 401 sampel (20%) sebagai data testing. Pembagian ini dilakukan dengan metode stratifikasi untuk mempertahankan proporsi distribusi sentimen pada kedua subset, sehingga memastikan model dapat mempelajari pola klasifikasi secara efektif dan diuji pada data yang representatif.

- Total instance:
 $1540 + 55 + 9 = 1604$
- Distribusi kelas:
 - Negatif: 1540
 - Netral: 55
 - Positif: 9

Menghitung Entropi Awal (*Entropi Dataset*)

$$\begin{aligned}
 Entropi(S) &= -\frac{1540}{1604} \log_2 \left(\frac{1540}{1604} \right) \\
 &\quad - \frac{55}{1604} \log_2 \left(\frac{55}{1604} \right) \\
 &\quad - \frac{9}{1604} \log_2 \left(\frac{9}{1604} \right) \\
 Entropi(S) &= -0.9601 \times \log_2(0.9601) - 0.0343 \\
 &\quad \times \log_2(0.0343) - 0.00561 \\
 &\quad \times \log_2(0.00561) \\
 Entropi(S) &= 0.0560 + 0.1671 + 0.0447 \approx 0.2678
 \end{aligned}$$

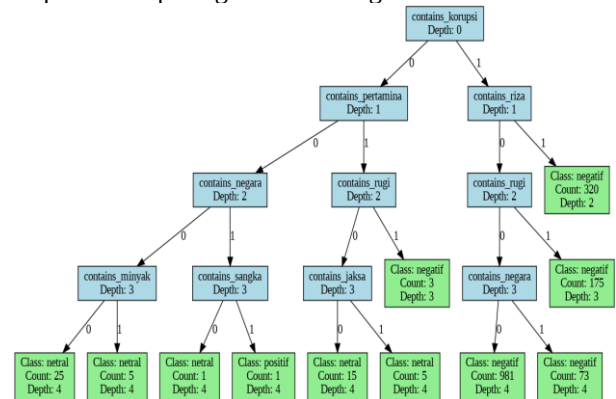
Proses ini dilakukan hingga node akhir. Setelah menghitung information gain, maka semua hasil yang didapatkan dapat dilihat pada Tabel 8.

Tabel 7. Hasil Perhitungan Entropy dan Information Gain pada Node 7 (contains_korupsi=0, contains_pertamina=0)

3.1.3 Hasil Algoritma ID3

Pembentukan pohon keputusan pada penelitian ini dilakukan melalui mekanisme pemilihan atribut

berbasis information gain untuk memaksimalkan pemisahan kelas sentimen. Hasil perhitungan menunjukkan bahwa contains_korupsi memiliki gain tertinggi sehingga ditetapkan sebagai node akar, karena atribut ini paling efektif menurunkan tingkat ketidakpastian kelas. Dataset kemudian dipartisi menjadi dua subset berdasarkan nilai atribut tersebut, dan pada masing-masing subset dilakukan kembali perhitungan entropy dan gain secara independen untuk memilih atribut pemisah berikutnya. Proses ini berlangsung secara rekursif hingga setiap cabang mencapai kondisi homogen atau tidak ada atribut tersisa. Dapat dilihat pada gambar 2 sebagai berikut.



Gambar 2. Pohon Keputusan

Berdasarkan struktur pohon keputusan yang terbentuk dari proses sebelumnya, dapat diekstraksi rule-rule prediksi dapat dilihat pada tabel 8.

Tabel 8. Rule-Rule Prediksi Algoritma ID3

No Rule	Kondisi	Keputusan
1	IF contains_korupsi = 0 AND contains_pertamina = 1 AND contains_negara = 1 AND contains_riza = 1	Positif
2	IF contains_korupsi = 0 AND contains_pertamina = 1 AND contains_negara = 1 AND contains_riza = 0	Netral
3	IF contains_korupsi = 0 AND contains_pertamina = 1 AND contains_negara = 0 AND contains_riza = 1	Negatif
4	IF contains_korupsi = 0 AND contains_pertamina = 1 AND contains_negara = 0 AND contains_riza = 0 AND contains_sangka = 1	Netral
5	IF contains_korupsi = 0 AND contains_pertamina = 1 AND contains_negara = 0 AND contains_riza = 0 AND contains_sangka = 0	Negatif
6	IF contains_korupsi = 0 AND contains_pertamina = 0 AND contains_riza = 1	Netral
7	IF contains_korupsi = 0 AND contains_pertamina = 0 AND contains_riza = 0	Positif

No Rule	Kondisi	Keputusan
8	IF contains_korupsi = 1 AND contains_jaksa = 1 AND contains_jagung = 1	Negatif
9	IF contains_korupsi = 1 AND contains_jaksa = 1 AND contains_jagung = 0 AND contains_rugi = 1	Negatif
10	IF contains_korupsi = 1 AND contains_jaksa = 1 AND contains_jagung = 0 AND contains_rugi = 0	Negatif
11	IF contains_korupsi = 1 AND contains_jaksa = 0 AND contains_chalid = 1	Negatif
12	IF contains_korupsi = 1 AND contains_jaksa = 0 AND contains_chalid = 0	Negatif

3.2 Evaluasi Model dengan Confusion Matrix

Untuk mendapatkan pemahaman yang lebih mendalam mengenai pola kesalahan klasifikasi yang dilakukan oleh model, analisis confusion matrix diperlukan. Confusion matrix merupakan tabel yang menampilkan distribusi prediksi model dibandingkan dengan label sebenarnya untuk setiap kelas sentimen. Analisis ini penting untuk mengidentifikasi apakah model memiliki bias tertentu atau kesulitan dalam membedakan kelas-kelas spesifik. Hasil confusion matrix dari pengujian model terhadap 401 data testing disajikan secara detail pada Tabel 9.

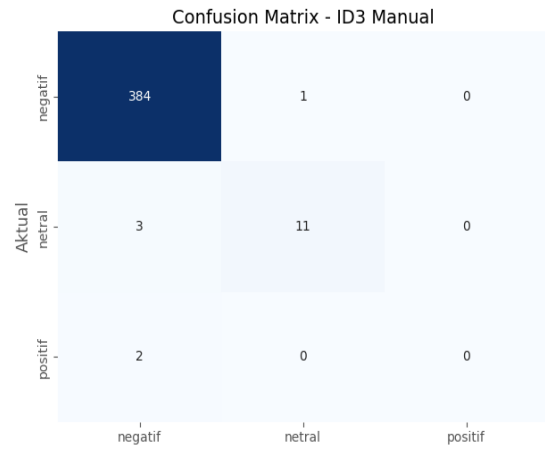
Tabel 9. Confusion Matrix Data Testing

	Prediksi Negatif	Prediksi Netral	Prediksi Positif	Class Precision
Aktual Negatif	384	1	0	99.74%
Aktual Netral	3	11	0	91.67%
Aktual Positif	2	0	0	0%
Class Recall	98.72%	78.57%	0%	

Berdasarkan pada tabel 9, analisis Confusion Matrix:

1. Kelas Negatif: Model menunjukkan performa sangat baik dengan *precision* 99.74% dan *recall* 98.72%. Hanya 4 dari 389 instance negatif yang salah diklasifikasikan.
2. Kelas Netral: *Precision* mencapai 91.67% namun *recall* hanya 78.57%, menunjukkan bahwa model agak konservatif dalam memprediksi kelas netral.
3. Kelas Positif: Model tidak berhasil mengidentifikasi instance positif karena jumlah data *training* yang sangat terbatas (hanya 9 instance).

Visualisasi *confusion matrix* yang dihasilkan oleh program dapat dilihat pada Gambar 3, yang memberikan representasi grafis dari pola kesalahan klasifikasi.



Gambar 3. Confusion Matrix Hasil Klasifikasi

Selain itu, keluaran program yang menampilkan nilai *accuracy*, *precision*, *recall*, dan *F1-score* secara numerik dapat dilihat pada Gambar 4.

```

*** === HASIL EVALUASI MODEL ID3 ===
Akurasi Model: 0.9850
Precision: 0.9798
Recall: 0.9850
F1-Score: 0.9822

Classification Report:
              precision    recall  f1-score   support

negatif      0.99         1.00         0.99         385
netral       0.92         0.79         0.85          14
positif      0.00         0.00         0.00           2

accuracy      0.63         0.59         0.61         401
macro avg     0.63         0.59         0.61         401
weighted avg  0.98         0.99         0.98         401

Confusion Matrix:
[[384  1  0]
 [ 3 11  0]
 [ 2  0  0]]

```

Model ID3 manual berhasil dilatih dan dievaluasi!

Gambar 4. Hasil Akurasi, Precision, Recall, dan F1-Score Model ID3

3.3 Kontribusi Penelitian

Penelitian ini memberikan kontribusi penting dalam bidang analisis sentimen berbasis *machine learning*, khususnya pada kajian opini publik terhadap Pertamina di media sosial seperti platform X, dengan menunjukkan bahwa algoritma Decision Tree menggunakan pendekatan ID3 mampu menghasilkan model klasifikasi yang sangat akurat meskipun hanya memanfaatkan fitur biner sederhana berbasis kata kunci, sehingga membuktikan bahwa metode yang relatif ringan secara komputasi tetap efektif apabila pemilihan fitur dilakukan secara optimal menggunakan information gain, selain itu, penelitian ini menghasilkan seperangkat rule klasifikasi yang bersifat interpretable dan transparan sehingga pola sentimen dapat dipahami secara logis melalui struktur pohon keputusan tanpa bergantung pada model black-box, serta memberikan kontribusi praktis berupa gambaran nyata dari kecenderungan sentimen masyarakat yang dapat dimanfaatkan sebagai bahan pertimbangan dalam evaluasi komunikasi publik untuk Pertamina, pemantauan reputasi digital, dan pengambilan keputusan strategis berbasis data.

4. KESIMPULAN

Berdasarkan hasil implementasi dan pengujian yang telah dilakukan, penelitian ini berhasil menerapkan metode Decision Tree dengan algoritma Iterative Dichotomiser 3 (ID3) untuk mengelompokkan sentimen masyarakat terhadap Pertamina di platform X. Model dibangun melalui tahapan pengumpulan 2.005 tweet dengan kata kunci “Pertamina Korupsi”, preprocessing teks, ekstraksi 10 fitur biner berdasarkan kata kunci dominan, serta pembentukan pohon keputusan secara rekursif menggunakan information gain tertinggi, dengan atribut contains_korupsi sebagai simpul akar. Hasil pengujian menunjukkan bahwa model mencapai tingkat akurasi sebesar 98,50% dengan nilai precision 97,98%, recall 98,50%, dan F1-score 98,22%. Validasi menggunakan 5-fold cross-validation mengonfirmasi stabilitas model dengan rata-rata akurasi 98,30% dan standar deviasi rendah (0,0046), serta menunjukkan performa yang lebih unggul dibandingkan penelitian sejenis, sehingga membuktikan efektivitas dan keandalan metode yang diterapkan.

REFERENSI

- [1] D. R. Manalu, M. C. L. Tobing, and M. Yohanna, “Analisis Sentimen Twitter Terhadap Wacana Penundaan Pemilu Dengan Metode Support Vector Machine,” vol. 6, no. 2, pp. 149–156, 2024.
- [2] Kompas.com, “Kejagung Periksa Pejabat Pertamina dalam Kasus Korupsi Minyak Mentah,” *Kompas.com*, 2025. .
- [3] CNN Indonesia, “Kejagung buka suara isu Pertamina oplosan,” *cnn indonesia*, 2025. .
- [4] C. Cahyaningtyas, Y. Nataliani, and I. R. Widiyari, “Analisis Sentimen Pada Rating Aplikasi Shopee Menggunakan Metode Decision Tree Berbasis SMOTE,” *Aiti*, vol. 18, no. 2, pp. 173–184, 2021, doi: 10.24246/aiti.v18i2.173-184.
- [5] F. Fathonah and A. Herliana, “Penerapan Text Mining Analisis Sentimen Mengenai Vaksin Covid - 19 Menggunakan Metode Naïve Bayes,” *J. Sains dan Inform.*, vol. 7, no. 2, pp. 155–164, 2021, doi: 10.34128/jsi.v7i2.331.
- [6] T. Ridwansyah, “Implementasi Text Mining Terhadap Analisis Sentimen Masyarakat Dunia Di Twitter Terhadap Kota Medan Menggunakan K-Fold Cross Validation Dan Naïve Bayes Classifier,” *KLIK Kaji. Ilm. Inform. dan Komput.*, vol. 2, no. 5, pp. 178–185, 2022, doi: 10.30865/klik.v2i5.362.
- [7] R. Antika, A. Rifa’i, F. Dikananda, D. Indriya Efendi, and R. Narasati, “Penerapan Algoritma Decision Tree Berbasis Pohon Keputusan Dalam Klasifikasi Penyakit Jantung,” *JATI (Jurnal Mhs. Tek. Inform.)*, vol. 7, no. 6, pp. 3688–3692, 2024, doi: 10.36040/jati.v7i6.8264.
- [8] L. Qadrini, A. Sepperwali, and A. Aina, “Decision Tree Dan Adaboost Pada Klasifikasi Penerima Program Bantuan Sosial,” *J. Inov. Penelit.*, vol. 2, no. 7, pp. 1959–1966, 2021.
- [9] Nengah Widya Utami and I Wayan Budi Suryawan, “Implementasi Data Mining Untuk Mengklasifikasikan Produk Pada Sebuah Supermarket Menggunakan Algoritma Id3 Pada Orange,” *Smart Techno (Smart Technol. Informatics Technopreneurship)*, vol. 3, no. 1, pp. 33–36, 2021, doi: 10.59356/smart-techno.v3i1.33.
- [10] A. Sifaunajah and R. D. Wahyuningtyas, “Penggunaan Algoritma ID3 Untuk Klasifikasi Data Calon Peserta Didik,” *CSRID (Computer Sci. Res. Its Dev. Journal)*, vol. 14, no. 2, p. 103, 2022, doi: 10.22303/csrid.14.2.2022.103-112.
- [11] A. Nurian, “Analisis Sentimen Ulasan Pengguna Aplikasi Google Play Menggunakan Naïve Bayes,” *J. Inform. dan Tek. Elektro Terap.*, vol. 11, no. 3s1, pp. 829–835, 2023, doi: 10.23960/jitet.v11i3s1.3348.