

ANALISIS SENTIMEN TERHADAP SKINCARE MENGUNAKAN METODE ALGORITMA *DECISION TREE*

Yustina Bete Dos Santos, Semlinda Juszandri Bulan

Sistem Informasi, STIKOM Uyelindo Kupang

Jl. Perintis Kemerdekaan 1, Kelurahan Kayu Putih, Kota Kupang, Nusa Tenggara Timur

yustinabetedossantos@gmail.com

ABSTRAK

Media sosial Instagram, telah menjadi *platform* utama bagi konsumen untuk menyampaikan opini tentang produk, termasuk dalam industri kecantikan di Indonesia. Meskipun industri kecantikan tumbuh pesat, masalah terkait klaim kualitas dan keamanan produk, seperti yang dialami oleh merk *skincare* “Daviena”, “Scora”, dan “The Originote” pada produk Serum HA Daviena, Serum Arbutin Scora, dan Serum Retinol The Originote, dapat memengaruhi kepercayaan konsumen. Tujuan penelitian untuk menganalisis sentimen konsumen terhadap ketiga produk *skincare* berdasarkan komentar di Instagram menggunakan algoritma *Decision Tree*. Dari 1006 komentar yang dikumpulkan, 468 untuk Daviena, 444 untuk Scora, dan 436 untuk The Originote yang diolah, data dibagi menjadi 80% untuk pelatihan dan 20% untuk pengujian. Hasilnya menunjukkan bahwa model *Decision Tree* mampu mengklasifikasikan sentimen dengan akurasi tinggi 94,22% untuk Daviena, 87,39% untuk Scora, dan 83,72% untuk The Originote. Presisi dan *recall* untuk Daviena adalah 89,96% (positif) dan 99,52% (negatif) dengan *recall* 99,57% (positif) dan 88,89% (negatif). Untuk Scora, presisi positif 80,74% dan negatif 97,70%, dengan *recall* 98,20% (positif) dan 76,58% (negatif). Sementara itu, The Originote memiliki presisi positif 76,73% dan negatif 95,65%, dengan *recall* 96,79% (positif) dan 70,64% (negatif). Kata-kata yang sering muncul dalam komentar termasuk “retinol” untuk Daviena, “scora” untuk Scora, dan “pake” untuk The Originote.

Kata kunci : Analisis sentimen, *Decision Tree*, Instagram, *Skincare* Daviena, Scora, The Originote.

1. PENDAHULUAN

Era digital yang berkembang pesat, terutama melalui media sosial, telah meningkatkan jumlah pengguna yang aktif memposting aktivitas sehari-hari, termasuk status yang mencerminkan suasana hati dan perilaku konsumen [1].

Instagram menjadi salah satu *platform* terpopuler di Indonesia, di mana konsumen berbagi informasi dan memberikan komentar. Opini masyarakat terbentuk dari penilaian terhadap produk dan layanan, yang dapat dianalisis untuk memahami sentimen [2].

Di Instagram, terdapat beragam opini positif dan negatif mengenai produk perawatan kulit, tubuh, dan wajah, yang mencerminkan perhatian terhadap kesehatan kulit.

Perawatan kulit (*skincare*) penting untuk melindungi kesehatan kulit dari kerusakan akibat sinar matahari, polusi, dan stres. Dengan rutinitas yang tepat, *skincare* dapat mencegah jerawat, penuaan dini, dan menjaga kelembapan kulit, yang berdampak positif pada rasa percaya diri. Di Indonesia, industri kecantikan berkembang pesat, dengan survei Populix menunjukkan bahwa 77% orang rutin membeli produk *skincare* setiap bulan, dan 93% menghabiskan sekitar Rp 250 ribu per bulan [2].

Produk *skincare* yang menonjol termasuk Scora, Daviena, dan The Originote, dengan media sosial seperti Instagram dan TikTok berperan besar dalam mempengaruhi tren *skincare*. *Skincare* “Daviena” juga menduduki peringkat pertama sebagai *brand* dengan nilai penjualan sebesar 5,1% dan “The Originote” menduduki peringkat kedua dengan nilai penjualan

sebesar 4,5% di kategori pelembab wajah pada periode Januari-Mei 2024.

Seiring meningkatnya popularitas *skincare* Daviena, perdebatan publik mengenai kualitas dan keamanan produk muncul, terutama terkait klaim kandungan *Niacinamide* pada *Body Lotion* HB Dosting Daviena yang menyatakan 9%. Namun, hasil uji laboratorium oleh “Doktif” menunjukkan kandungan *Niacinamide* hanya -1% atau tidak terdeteksi. Masalah serupa juga terjadi pada produk lain, seperti serum The Originote yang mengklaim 10% *Niacinamide* tetapi hasil uji lab menunjukkan 4,97%, dan *moisturizer gel* Scora yang mengklaim 5% tetapi hanya terdeteksi 0,8%. Ketidakpastian ini dapat mengurangi kepercayaan konsumen dan berpotensi merugikan produsen, karena klaim yang tidak sesuai dapat menyebabkan konsumen merasa dirugikan dan mengurangi minat beli. Ketidakpuasan ini dapat memicu diskusi negatif di media sosial, yang berpotensi merusak citra merk secara keseluruhan. Oleh karena itu, penting bagi perusahaan untuk memahami dan merespons perubahan sentimen konsumen.

Untuk mengatasi permasalahan ini, analisis sentimen menggunakan metode *Decision Tree* dapat diterapkan. *Decision Tree* dipilih karena memiliki interpretabilitas yang tinggi, memungkinkan konsumen untuk memahami dengan jelas bagaimana sentimen positif maupun negatif diungkapkan. Metode ini juga dapat menangani data kategorikal dan numerik dengan baik, serta fleksibel dalam berbagai konteks analisis. Selain itu, *Decision Tree* dapat mendeteksi interaksi antara sentimen, yang penting dalam

memahami bagaimana kombinasi faktor-faktor tertentu mempengaruhi persepsi konsumen. Kecepatan dan efisiensi dalam pelatihan dan prediksi juga menjadi keuntungan, terutama ketika berhadapan dengan volume data besar dari media sosial.

Penelitian yang dilakukan oleh [3] menunjukkan bahwa penggunaan *Decision Tree* berbasis SMOTE dalam analisis sentimen pada rating aplikasi Shopee mencapai akurasi sebesar 99,91% dengan nilai AUC (*Area Under Curve*) sebesar 0,999, serta *recall* dan *precision* masing-masing sebesar 99,88% dan 99,98%. Hasil ini menegaskan efektivitas *Decision Tree* dalam mengklasifikasikan sentimen pengguna dengan akurasi yang sangat tinggi.

Selain itu, penelitian oleh [4] juga menunjukkan bahwa model *Decision Tree* dapat memberikan hasil yang baik dalam analisis sentimen, dengan akurasi tertinggi mencapai 81,25% pada pengujian dengan 10-fold cross validation. Ini menunjukkan bahwa metode berbasis pohon keputusan, termasuk *Decision Tree*, memiliki potensi yang besar dalam analisis sentimen. Dengan demikian, penelitian ini berjudul “Analisis Sentimen Terhadap Skincare Menggunakan Metode Algoritma *Decision Tree*” bertujuan untuk mengeksplorasi respon konsumen terhadap produk-produk tersebut, baik yang bersifat positif maupun negatif, serta memahami perubahan persepsi setelah munculnya masalah terkait klaim dan kualitas.

2. TINJAUAN PUSTAKA

2.1. Penelitian Terdahulu

[5] melakukan penelitian dengan judul “Analisis Sentimen Ulasan Produk Serum Wajah pada *Beauty Brand Somethinc* Menggunakan Metode *Naïve Bayes Classifier*”. Tujuan penelitian ini untuk mengetahui review dari produk serum *somethinc* dengan perbandingan Serum *Niacinamide* dan Serum AHA BAH. Metode yang digunakan pada penelitian ini dengan menggunakan metode *text mining* dan algoritma *Naïve Bayes*. Hasil penelitian ini menunjukkan bahwa data produk serum AHA BAH memiliki hasil lebih baik dengan jumlah data sentimen positif 675 dan nilai akurasi 80%, presisi 84%, dan *recall* 94% yang tergolong positif.

[6] melakukan penelitian dengan judul “Implementasi Algoritma *Naïve Bayes* untuk Analisis Sentimen Masyarakat pada Twitter mengenai Kepopuleran Produk *Skincare* di Indonesia”. Tujuan penelitian ini untuk mengetahui analisis sentimen masyarakat terhadap kepopuleran produk *skincare* di Indonesia. Metode yang digunakan pada penelitian ini menggunakan algoritma *Naïve Bayes* dan pengolahan data dilakukan menggunakan *software Orange*. Hasil dari penelitian ini dengan menghasilkan nilai *accuracy*, *precision*, dan *recall* diatas 80% dan dapat memberikan analisis sentimen yang positif terhadap kepopuleran produk *skincare* di Indonesia.

[7] melakukan penelitian dengan judul “Analisis Sentimen pada *Review Skincare Female Daily* Menggunakan Metode *Support Vector Machine*

(SVM)”. Tujuan penelitian ini untuk membantu perempuan agar lebih mudah mengetahui jenis produk *skincare* yang cocok pada kulit. Metode yang digunakan pada penelitian ini dengan menggunakan *support vector machine* (SVM). Hasil dari penelitian ini dengan mendapatkan akurasi sebanyak 87% dengan *recall* sebesar 90%, *precision* sebesar 84,90%, dan *f1 score* sebesar 87,37%.

[8] melakukan penelitian dengan judul “Analisis Sentimen Pengaruh Media Sosial Terhadap Minat Beli *Skincare* pada Remaja di Indonesia Menggunakan Algoritma *Naïve Bayes*”. Tujuan penelitian ini untuk menganalisis sejauh mana pengaruh media sosial terhadap minat beli *skincare* pada remaja. Metode yang digunakan pada penelitian ini dengan menggunakan algoritma *Naïve Bayes*. Hasil dari penelitian ini adalah berupa diagram sentimen positif dan negatif dengan hasil positif 0,8138 dan yang negatif 0,7526.

[9] melakukan penelitian dengan judul “Analisis Sentimen Produk *Skincare Somethinc Niacinamide* di *Female Daily* dengan *Naïve Bayes Classifier*”. Tujuan penelitian ini untuk menganalisis sentimen publik pada review produk *skincare Somethinc Niacinamide* di *Female Daily*. Metode yang digunakan pada penelitian ini dengan algoritma *Naïve Bayes Classifier*. Hasil dari penelitian ini menunjukkan bahwa model ini berhasil mengklasifikasikan sentimen dengan nilai akurasi sebesar 61%.

2.2. Analisis Sentimen

Satu teknik analisa data kualitatif yang sangat berguna adalah analisis sentimen. Tujuan dari analisis sentimen adalah untuk menafsirkan dan mengklasifikasikan emosi yang disampaikan dalam data tekstual. Salah satu jenis model analisis sentimen yaitu *emotion detection* dimana model ini sering menggunakan algoritma pembelajaran mesin yang kompleks untuk memilih berbagai emosi dari data tekstual. Kata-kata yang terkait dengan kebahagiaan, kemarahan, frustrasi, dan kegembiraan, memberi informasi tentang perasaan pelanggan saat menulis tentang produk di situs ulasan produk [10].

2.3. Decision Tree

Teknik mengambil sebuah keputusan memang banyak jenisnya, namun terdapat sebuah model pengambilan keputusan yang mungkin bisa menjadi alternatif yaitu model *Decision Tree* atau pohon keputusan. Model *Decision Tree* merupakan sebuah “pohon keputusan” yang menjulur ke bawah, diawali dengan akar pada bagian atasnya, berlanjut pada cabang dan ranting-ranting keputusan atau yang biasa disebut dengan lengan keputusan dan berakhir pada daun-daun keputusan yang bisa dipetik dan ambil sebagai kebijakan menyelesaikan masalah, menghindari masalah, atau meminimalkan risiko serta efisiensi biaya [11].at analisis *machine learning* di era *Data Science* dan *Big Data* dengan pendekatan klasifikasi terstruktur pohon yang terdiri atas node

sebagai atributnya, dan daun-daun keputusan yang berkelas-kelas di dalamnya sebagai pendukung kemungkinan hasil, biaya atas sumber daya, kebermanfaatan, serta kemungkinan risiko dari permasalahan [11].

2.4. Klasifikasi

Klasifikasi adalah proses yang bertujuan untuk menemukan sekumpulan model yang dapat menjelaskan kelas-kelas data [12].

Model ini kemudian digunakan untuk memprediksi kelas yang belum diketahui dari suatu objek. Proses klasifikasi melibatkan analisis terhadap data latih (*training set*) untuk membangun model dan data uji (*test set*) untuk mengevaluasi akurasi model tersebut. Dengan demikian, klasifikasi dapat digunakan untuk memprediksi nama atau nilai dari objek data berdasarkan karakteristik yang ada [12].

Klasifikasi adalah proses menemukan sekumpulan model atau fungsi yang mendeskripsikan dan membedakan kelas data. Tujuan dari klasifikasi adalah untuk memprediksi kelas suatu objek yang kelasnya tidak diketahui. Proses klasifikasi terdiri dari dua tahap: pembuatan model klasifikasi dari sekumpulan kelas data yang sudah ada dan didefinisikan sebelumnya (*set data latih*), dan penggunaan model tersebut untuk klasifikasi tes data (*prediksi*), dan mengukur kelas data yang ada [13].

Untuk akurasi klasifikasi, pertama-tama akan dinilai akurasi prediksi *classifier*. Jika *training set* digunakan untuk mengukur akurasi *classifier*, estimasi yang dihasilkan kemungkinan besar akan sangat baik karena *classifier* cenderung *overfit* pada data. Oleh karena itu, untuk mengetahui akurasi prediksi *classifier* menggunakan *test set*, yang merupakan bagian dari data tuple yang tidak digunakan dalam proses pembelajaran. Setelah membangun model untuk klasifikasi, akan muncul pertanyaan baru tentang seberapa bagus/akurat model yang dibangun. Setelah model dibangun maka akan dilakukan evaluasi terhadap kinerja *classifier*, evaluasi tersebut terdiri dari *accuracy*, *recall*, *specificity*, *precision*, *F1*, dan *Fβ* [14].

2.5. Preprocessing

Data *preprocessing* atau prapemrosesan data adalah serangkaian langkah atau tahapan yang dilakukan pada data mentah sebelum data tersebut digunakan untuk analisis lebih lanjut atau pengembangan model. Tujuan utama dari data *preprocessing* adalah untuk meningkatkan kualitas data, memastikan keakuratan hasil analisis, dan mengatasi masalah atau kekeurangan yang mungkin muncul dalam data mentah [15].

Text preprocessing merupakan tahapan dari proses awal terhadap *text* untuk mempersiapkan *text* menjadi data yang akan diolah lebih lanjut. Suatu *text* tidak dapat diproses langsung algoritma pencarian, oleh karena itu dibutuhkan *preprocessing text* untuk mengubah *text* menjadi data *numeric* [16].

2.6. Term Frequency Inverse Document Frequency (TF-IDF)

Algoritma TF/IDF digunakan dalam proses perhitungan bobot (TF-IDF) terminology kata. Algoritma ini digunakan untuk menghitung bobot setiap kata yang paling umum digunakan pada *information retrieval*. Metode ini juga dikenal efisien, mudah, dan memiliki hasil yang akurat [16].

Term Frequency Inverse Document Frequency (TF-IDF) adalah proses mengubah kata menjadi nilai numerik atau vector yang biasa disebut pembobotan [17]. Tujuan dari TF-IDF adalah untuk menghitung pentingnya sebuah kata dalam sebuah dokumen dengan memberikan bobot setiap kata dalam setiap dokumen untuk mengidentifikasi dan menghitung jumlah kemunculan kata tersebut [17].

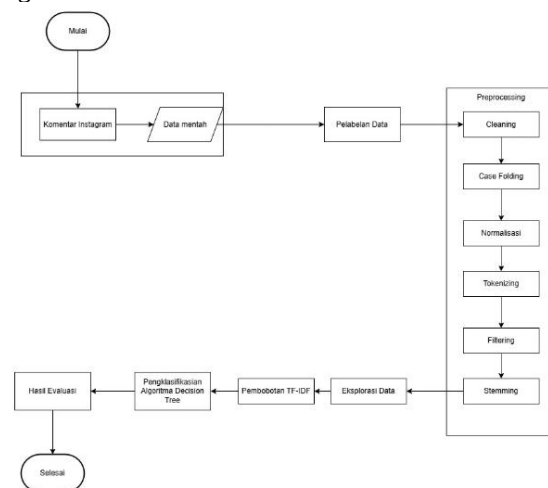
3. METODE PENELITIAN

3.1. Pengumpulan Data

Tahap pengumpulan data dari *platform* media sosial Instagram dalam bentuk teks berbahasa Indonesia. Data diambil dengan menggunakan kata kunci yang berkaitan dengan *skincare* 'Daviena', 'Scora', dan 'The Originote'. Kata kunci yang digunakan meliputi "Daviena skincare", "skincare Daviena", "skincare Scora", dan "skincare The Originote" untuk memastikan semua data relevan dengan studi kasus yang diteliti.

3.2. Alur Penelitian

Rancangan sistem dibuat untuk memberikan Gambaran mengenai alur setiap proses dari sistem yang akan dibangun. Sistem yang akan dibangun dapat digambarkan secara detail melalui Gambar 1.



Gambar 1. Alur Penelitian

Berdasarkan Gambar 1 maka perancangan alur setiap proses dari sistem yang akan dibangun dapat dijelaskan sebagai berikut.

a. Pengambilan Data

Proses dimulai dengan pengambilan data ulasan konsumen dari media sosial Instagram untuk produk *skincare* Daviena, Scora, dan The Originote.

b. Pelabelan Data

Setelah data diambil, penulis memberikan label sentimen pada setiap ulasan dengan menggunakan Python. Data diklasifikasikan menjadi dua kategori: Positif: Ulasan yang menunjukkan reaksi menyenangkan dan menciptakan kesan baik dan Negatif: Ulasan yang menunjukkan reaksi kurang menyenangkan dan menciptakan kesan buruk.

c. *Preprocessing* Data

Data yang telah dilabeli kemudian melalui beberapa tahapan *preprocessing* untuk mempersiapkan data sebelum analisis lebih lanjut:

Cleaning: Menghapus elemen yang tidak relevan seperti tanda baca, angka, dan karakter khusus.

Case Folding: Mengubah semua huruf menjadi huruf kecil untuk konsistensi.

Tokenizing: Memecah teks menjadi token (kata) berdasarkan spasi.

Filtering: Menghilangkan kata-kata yang tidak penting menggunakan daftar *stopword*.

Stemming: Mengubah kata berimbuhan menjadi bentuk dasar sesuai dengan aturan yang berlaku.

Eksplorasi Data
Setelah *preprocessing*, dilakukan eksplorasi data dengan visualisasi sentimen menggunakan *WordCloud* untuk menampilkan kata-kata yang paling sering muncul dalam setiap kategori sentimen.

e. Pembagian Data

Data yang telah diproses dibagi secara acak menjadi dua bagian: 80% untuk data *training* dan 20% untuk data *testing*, untuk memastikan model dapat diuji dengan baik.

f. Pembobotan *Term*

Pada tahap ini, setiap *term* dalam data *training* diberikan bobot menggunakan metode TF-IDF (*Term Frequency-Inverse Document Frequency*) untuk menilai pentingnya setiap kata dalam konteks dokumen.

g. Pelatihan Model

Data *training* yang telah dibobot kemudian digunakan untuk melatih model menggunakan algoritma *Decision Tree*. Proses ini bertujuan untuk menemukan pola dalam data yang dapat digunakan untuk mengklasifikasikan data testing dengan akurasi yang optimal.

4. HASIL DAN PEMBAHASAN

4.1. *Crawling* Data

Crawling data dilakukan dari Januari hingga Februari 2025 di Instagram, mengumpulkan komentar terkait tiga merk produk *skincare* Serum HA Daviena, Serum Arbutin Scora, dan Serum Retinol The Originote. Total data yang diambil mencapai 1006 komentar, terdiri dari 334 komentar untuk Daviena, 332 untuk Scora, dan 340 untuk The Originote. Proses *crawling* dilakukan dengan dua metode menggunakan Apify, *platform web scraping* yang menyimpan data

dalam format *spreadsheet*, dan secara manual. *Scraping* di Apify dilakukan dengan mengambil link dari Instagram.com. Gambar 2 menunjukkan contoh hasil *crawling* menggunakan Apify.

Gambar 2. Hasil *Crawling* Menggunakan Apify

Berikut ini merupakan data mentah dari hasil *crawling* data di Apify dan secara manual disimpan dalam bentuk *file spreadsheet* yang dapat dilihat pada Gambar 3.

Index	
1	Index
2	Index
3	Index
4	Index
5	Index
6	Index
7	Index
8	Index
9	Index
10	Index
11	Index
12	Index
13	Index
14	Index
15	Index
16	Index
17	Index
18	Index
19	Index
20	Index
21	Index
22	Index
23	Index
24	Index
25	Index
26	Index
27	Index
28	Index
29	Index
30	Index
31	Index
32	Index
33	Index
34	Index
35	Index
36	Index
37	Index
38	Index
39	Index
40	Index
41	Index
42	Index
43	Index
44	Index
45	Index
46	Index
47	Index
48	Index
49	Index
50	Index
51	Index
52	Index
53	Index
54	Index
55	Index
56	Index
57	Index
58	Index
59	Index
60	Index
61	Index
62	Index
63	Index
64	Index
65	Index
66	Index
67	Index
68	Index
69	Index
70	Index
71	Index
72	Index
73	Index
74	Index
75	Index
76	Index
77	Index
78	Index
79	Index
80	Index
81	Index
82	Index
83	Index
84	Index
85	Index
86	Index
87	Index
88	Index
89	Index
90	Index
91	Index
92	Index
93	Index
94	Index
95	Index
96	Index
97	Index
98	Index
99	Index
100	Index

Gambar 3. Data Mentah *Spreadsheet*

4.2. Labeling Ulasan

Labeling data ulasan komentar Instagram dilakukan menggunakan Python di Google Colab. Kata-kata positif dan negatif diambil dari kombinasi beberapa sumber, termasuk kamus sentimen Bahasa Indonesia (INSET *Lexicon*), ulasan produk di *e-commerce*, serta istilah umum dalam konteks *skincare* dan kecantikan. Kata positif berasal dari ulasan yang menunjukkan kepuasan dan rekomendasi, sedangkan kata negatif diambil dari keluhan pengguna yang mencerminkan ketidakpuasan dan pengalaman buruk dengan produk. Gambar 4 menunjukkan proses *labelling* di Google Colab.

[illegible]

Gambar 4. Proses *Labeling* di Google Colab

Proses *labeling* yang ditunjukkan dalam gambar bertujuan untuk menganalisis sentimen ulasan produk

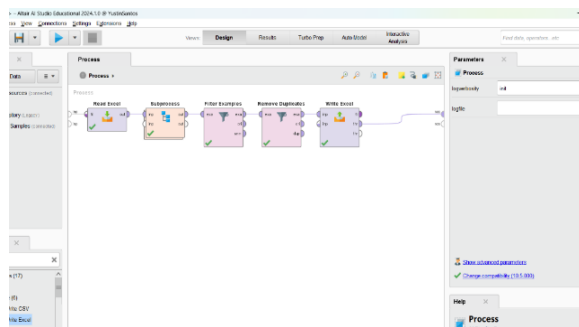
menggunakan kata-kata dan emoji. Pertama, kode menginisialisasi dua kamus untuk kata-kata positif dan negatif, serta daftar emoji yang relevan. Fungsi `get_sentiment` digunakan untuk menentukan sentimen dengan menghitung jumlah kata dan emoji positif serta negatif dalam teks. Setelah itu, *dataset* ulasan dimuat dari file CSV, di mana kode memastikan kolom yang menyimpan ulasan diidentifikasi dan menerapkan fungsi `get_sentiment` untuk menghasilkan kolom baru yang berisi hasil sentimen. Hasil akhir kemudian disimpan dalam file CSV baru, dan dapat mendownloadnya secara otomatis. Proses ini memastikan semua ulasan telah teranalisis dan diklasifikasikan sebagai positif atau negatif secara sistematis.

4.3. Pembersihan Data dan *Preprocessing Text*

Pembersihan data dan *preprocessing text* merupakan langkah penting dalam analisis sentimen untuk memastikan data yang digunakan bersih dan relevan. Proses ini meliputi penghapusan elemen yang tidak diperlukan, *case folding*, *tokenizing*, penghapusan *stopwords*, *stemming*, dan *filtering*. Tahapan ini bertujuan untuk meningkatkan kualitas data agar lebih siap dianalisis.

4.4. Proses *Cleaning*

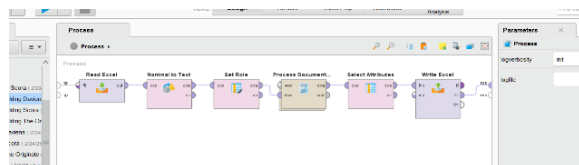
Proses ini digunakan untuk menghapus atribut seperti *hashtag*, *mention*, emoji, *whitespace*, dan lain-lain yang tidak berguna atau tidak penting. *Dataset* dibersihkan dan Gambar 5 merupakan proses *cleaning* yang dilakukan.



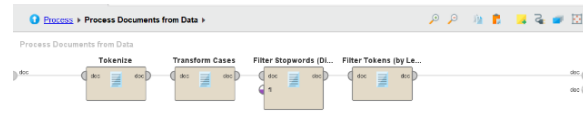
Gambar 5. Proses *Cleaning Dataset*

4.5. *Preprocessing Text*

Beberapa tahap *preprocessing text* termasuk *case folding*, *tokenizing*, *filtering*, dan *stemming*. Gambar 6 merupakan proses *preprocessing text* dan Gambar 7 merupakan tahapan *preprocessing text*.



Gambar 6. Proses *Preprocessing Text*



Gambar 7. Tahapan *Preprocessing Text*

a. Proses *Case Folding*

Natural Language Processing (NLP) melalui proses *case folding* mengubah semua huruf dalam teks menjadi huruf kecil. Tujuannya adalah menghilangkan perbedaan antara huruf besar dan kecil, sehingga analisis teks menjadi lebih akurat dan konsisten.

b. Proses *Tokenizing*

Natural Language Processing (NLP) yang dikenal sebagai *tokenisasi* membagi teks menjadi bagian-bagian kecil yang disebut "token". Tergantung pada tujuan analisis, token ini dapat berupa kata.

c. Proses *Filtering*

Natural Language Processing (NLP) yang dikenal sebagai *filtering* memungkinkan penghapusan elemen tertentu dari teks yang dianggap tidak relevan atau tidak penting untuk analisis.

d. Proses *Stemming*

Stemming merupakan langkah dalam *Natural Language Processing* (NLP) yang bertujuan untuk mengubah kata-kata menjadi bentuk dasarnya atau akarnya. Proses ini berfungsi untuk mengelompokkan kata-kata yang memiliki arti serupa meskipun dalam bentuk yang berbeda.

4.6. *Feature Extraction Using TF-IDF*

Proses *Feature Extraction* menggunakan TF-IDF menghitung jumlah kemunculan *term* (tf), jumlah komentar yang mengandung *term* (df), dan nilai *inverse document frequency* (idf). Kemudian, tf dikalikan dengan idf untuk mendapatkan bobot *term* untuk setiap komentar. Hasil ekstraksi fitur ditampilkan pada Gambar 8, Gambar 9, dan Gambar 10.

	acid	acne	adren	adren	adren	adren	adren	adren	adren	adren	adren	adren
1	0	0	0	0	0	0	0	0	0	0	0	0
2	0	0	0	0	0	0	0	0	0	0	0	0
3	0	0	0	0	0	0	0	0	0	0	0	0
4	0	0	0	0	0	0	0	0	0	0	0	0
5	0	0	0	0	0	0	0	0	0	0	0	0
6	0	0	0	0	0	0	0	0	0	0	0	0
7	0	0	0	0	0	0	0	0	0	0	0	0
8	0	0	0	0	0	0	0	0	0	0	0	0
9	0	0	0	0	0	0	0	0	0	0	0	0
10	0	0	0	0	0	0	0	0	0	0	0	0
11	0	0	0	0	0	0	0	0	0	0	0	0
12	0	0	0	0	0	0	0	0	0	0	0	0
13	0	0	0	0	0	0	0	0	0	0	0	0
14	0	0	0	0	0	0	0	0	0	0	0	0
15	0	0	0	0	0	0	0	0	0	0	0	0
16	0	0	0	0	0	0	0	0	0	0	0	0
17	0	0	0	0	0	0	0	0	0	0	0	0
18	0	0	0	0	0	0	0	0	0	0	0	0
19	0	0	0	0	0	0	0	0	0	0	0	0
20	0	0	0	0	0	0	0	0	0	0	0	0
21	0	0	0	0	0	0	0	0	0	0	0	0
22	0	0	0	0	0	0	0	0	0	0	0	0
23	0	0	0	0	0	0	0	0	0	0	0	0
24	0	0	0	0	0	0	0	0	0	0	0	0
25	0	0	0	0	0	0	0	0	0	0	0	0
26	0	0	0	0	0	0	0	0	0	0	0	0
27	0	0	0	0	0	0	0	0	0	0	0	0
28	0	0	0	0	0	0	0	0	0	0	0	0
29	0	0	0	0	0	0	0	0	0	0	0	0
30	0	0	0	0	0	0	0	0	0	0	0	0
31	0	0	0	0	0	0	0	0	0	0	0	0
32	0	0	0	0	0	0	0	0	0	0	0	0
33	0	0	0	0	0	0	0	0	0	0	0	0
34	0	0	0	0	0	0	0	0	0	0	0	0
35	0	0	0	0	0	0	0	0	0	0	0	0
36	0	0	0	0	0	0	0	0	0	0	0	0
37	0	0	0	0	0	0	0	0	0	0	0	0
38	0	0	0	0	0	0	0	0	0	0	0	0
39	0	0	0	0	0	0	0	0	0	0	0	0
40	0	0	0	0	0	0	0	0	0	0	0	0
41	0	0	0	0	0	0	0	0	0	0	0	0
42	0	0	0	0	0	0	0	0	0	0	0	0
43	0	0	0	0	0	0	0	0	0	0	0	0
44	0	0	0	0	0	0	0	0	0	0	0	0
45	0	0	0	0	0	0	0	0	0	0	0	0
46	0	0	0	0	0	0	0	0	0	0	0	0
47	0	0	0	0	0	0	0	0	0	0	0	0
48	0	0	0	0	0	0	0	0	0	0	0	0
49	0	0	0	0	0	0	0	0	0	0	0	0
50	0	0	0	0	0	0	0	0	0	0	0	0

Gambar 8. Hasil TF-IDF *skincare* Daviena

[illegible]

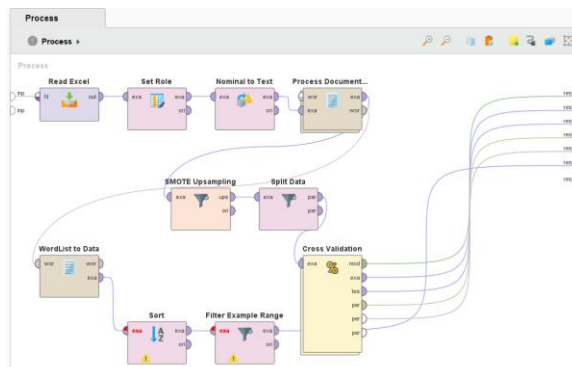
Gambar 9. Hasil TF-IDF *skincare* Scora

[illegible]

Gambar 10. Hasil TF-IDF *skincare* The Originote

4.7. Analisis Sentimen dengna Algoritma *Decision Tree*

Setelah data melalui tahap *preprocessing* dan pelabelan, tahap selanjutnya adalah melakukan pengujian analisis sentimen dengan memanfaatkan data *training* dan data *testing*. Pada Gambar 10 merupakan proses dari implementasi *Decision Tree*.



Gambar 11. Implementasi *Decision Tree*

Proses yang ditunjukkan dalam Gambar 11 menggunakan model *Decision Tree* untuk analisis data. Pertama, data diimpor dan diproses dengan menetapkan peran kolom serta mengonversi kolom nominal menjadi teks. Teknik SMOTE digunakan untuk mengatasi ketidakseimbangan kelas dengan meningkatkan jumlah data kelas minoritas. Setelah itu, *dataset* dibagi menjadi *set* pelatihan dan pengujian, diikuti oleh validasi silang untuk memastikan model tidak overfitting. Model *Decision Tree* dilatih pada data pelatihan untuk memprediksi hasil, dan akhirnya dievaluasi untuk mengukur akurasi. Proses ini memberikan kerangka kerja sistematis untuk analisis yang lebih akurat dan kuat.

Proses penerapan *Decision Tree* dalam penelitian ini akan menggunakan data *testing* dan data *training* untuk mencapai nilai akurasi yang paling optimal. *Dataset* dalam penelitian ini dibagi menjadi dua bagian, yaitu 20% untuk data uji dan 80% untuk data latih. Berikut Gambar 12, Gambar 13, dan Gambar 14 merupakan hasil implementasi *Decision Tree* pada ketiga produk *skincare* Daviena, Scora dan The Originote.

Estimation Results												
Case #	Task type	Activity	Time (min)	Success rate (%)	Confidence (%)	WIS	MSD	RMSE	MAE	MAPE	MAPE _{95%}	MAPE _{5%}
1	Fixed	Fixed	0.075	0.981	0	0	0	0	0	0	0	0
2	Fixed	Fixed	0.075	0.981	0	0	0	0	0	0	0	0
3	Fixed	Fixed	0.075	0.981	0	0	0	0	0	0	0	0
4	Fixed	Fixed	0.075	0.981	0	0	0	0	0	0	0	0
5	Neural	Neural	1	0	0	0	0	0	0	0	0	0
6	Fixed	Fixed	0.075	0.981	0	0	0	0	0	0	0	0
7	Fixed	Fixed	0.075	0.981	0	0	0	0	0	0	0	0
8	Fixed	Fixed	0.075	0.981	0	0	0	0	0	0	0	0
9	Fixed	Fixed	0.075	0.981	0	0	0	0	0	0	0	0
10	Fixed	Fixed	0.075	0.981	0	0	0	0	0	0	0	0
11	Fixed	Fixed	0.075	0.981	0	0	0	0	0	0	0	0
12	Fixed	Fixed	0.075	0.981	0	0	0	0	0	0	0	0
13	Fixed	Fixed	0.075	0.981	0	0	0	0	0	0	0	0
14	Fixed	Fixed	0.075	0.981	0	0	0	0	0	0	0	0
15	Fixed	Fixed	0.075	0.981	0	0	0	0	0	0	0	0
16	Fixed	Fixed	0.075	0.981	0	0	0	0	0	0	0	0
17	Fixed	Fixed	0.075	0.981	0	0	0	0	0	0	0	0
18	Fixed	Fixed	0.075	0.981	0	0	0	0	0	0	0	0
19	Fixed	Fixed	0.075	0.981	0	0	0	0	0	0	0	0
20	Fixed	Fixed	0.075	0.981	0	0	0	0	0	0	0	0
21	Fixed	Fixed	0.075	0.981	0	0	0	0	0	0	0	0
22	Fixed	Fixed	0.075	0.981	0	0	0	0	0	0	0	0
23	Fixed	Fixed	0.075	0.981	0	0	0	0	0	0	0	0
24	Fixed	Fixed	0.075	0.981	0	0	0	0	0	0	0	0

Gambar 12. Hasil Implementasi *Decision Tree* skincare Daviena

[illegible]

Gambar 13. Hasil Implementasi *Decision Tree* skincare Scora

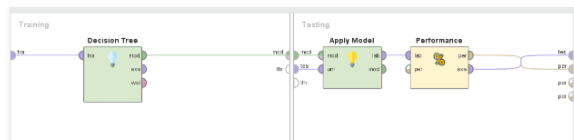
The screenshot shows the RStudio environment with a script editor at the top and a console at the bottom. The main window displays a data frame named 'data' with 19 rows and 10 columns. The columns are: Name, Vorname, geburtsjahr, adressstrasse, wohnort, sex, alter, arbeit, einkommen, and abgabe. The data represents survey results for various professions and their associated income and taxes.

Name	Vorname	geburtsjahr	adressstrasse	wohnort	sex	alter	arbeit	einkommen	abgabe	abgabe
1	Paul	1940	11017	10001	M	2	0	0	0	0
2	Paul	1940	11017	10001	M	2	0	0	0	0
3	Paul	1940	11017	10001	M	2	0	0	0	0
4	Paul	1940	11017	10001	M	2	0	0	0	0
5	Paul	1940	11017	10001	M	2	0	0	0	0
6	Paul	1940	11017	10001	M	2	0	0	0	0
7	Paul	1940	11017	10001	M	2	0	0	0	0
8	Paul	1940	11017	10001	M	2	0	0	0	0
9	Paul	1940	11017	10001	M	2	0	0	0	0
10	Paul	1940	11017	10001	M	2	0	0	0	0
11	Paul	1940	11017	10001	M	2	0	0	0	0
12	Paul	1940	11017	10001	M	2	0	0	0	0
13	Paul	1940	11017	10001	M	2	0	0	0	0
14	Paul	1940	11017	10001	M	2	0	0	0	0
15	Paul	1940	11017	10001	M	2	0	0	0	0
16	Paul	1940	11017	10001	M	2	0	0	0	0
17	Paul	1940	11017	10001	M	2	0	0	0	0
18	Paul	1940	11017	10001	M	2	0	0	0	0
19	Paul	1940	11017	10001	M	2	0	0	0	0

Gambar 14. Hasil Implementasi *Decision Tree* skincare The Originote

4.8. Evaluasi

Hasil evaluasi untuk kedua metode *Decision Tree* berupa *confusion matrix* yang mencakup nilai akurasi, presisi, dan *recall* yang diperoleh dari data *test*. Untuk menentukan nilai *confusion matrix*, penelitian ini akan menerapkan *k-fold cross validation* dengan $k=10$ agar hasil yang diperoleh menjadi maksimal. Pada Gambar 14 merupakan tahap dari *Cross Validation*.



Gambar 15. Tahap Cross Validation

Berikut ini merupakan hasil dari *Confusion Matrix* pada *skincare* Daviena, Scora, dan The Originote yang dilakukan pada tahap *cross validation*.

accuracy: 94.22% +/- 2.51% (micro average: 94.22%)

	true Positif	true Negatif	class precision
pred Positif	233	26	89.96%
pred Negatif	1	238	99.52%
class recall	99.57%	88.89%	

Gambar 16. Hasil *Confusion Matrix* *skincare* Daviena

Gambar 16 menunjukkan *skincare* Daviena menunjukkan performa yang sangat baik dengan akurasi 94,22%, yang berarti model *Decision Tree* berhasil mengklasifikasikan data dengan benar sebanyak 94,23%. *Margin error* yang diperoleh adalah $\pm 2,51\%$, menunjukkan bahwa kesalahan dalam pengambilan sampel cukup kecil. Dalam analisis sentimen, terdapat 234 sentimen positif dan 234 sentimen negatif. Presisi model untuk prediksi positif mencapai 89,96% dan untuk prediksi negatif 99,52%. Selain itu, nilai *recall* menunjukkan bahwa model berhasil menemukan kembali informasi dengan rasio 99,57% untuk data positif dan 88,89% untuk data negatif.

accuracy: 87.39% +/- 2.42% (micro average: 87.39%)

	true Positif	true Negatif	class precision
pred Positif	218	52	80.74%
pred Negatif	4	170	97.70%
class recall	98.20%	76.58%	

Gambar 17. Hasil *Confusion Matrix* *skincare* Daviena

Gambar 17 menunjukkan *skincare* Scora memiliki akurasi yang lebih rendah, yaitu 87,39%, yang menunjukkan bahwa model *Decision Tree* dapat mengklasifikasikan data dengan benar sebanyak 87,39%. *Margin error* yang diperoleh adalah $\pm 2,42\%$, yang juga mencerminkan kesalahan dalam pengambilan sampel. Terdapat 222 sentimen positif dan 222 sentimen negatif dalam dataset ini. Presisi untuk prediksi positif adalah 80,74% dan untuk prediksi negatif 97,70%. Nilai *recall* menunjukkan keberhasilan model dalam menemukan informasi kembali, dengan hasil 98,20% untuk data positif dan 76,58% untuk data negatif.

accuracy: 83.72% +/- 2.76% (micro average: 83.72%)

	true Positif	true Negatif	class precision
pred Positif	211	64	76.73%
pred Negatif	7	154	95.60%
class recall	96.75%	70.64%	

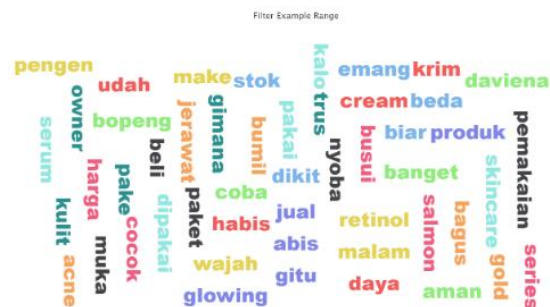
Gambar 18. Hasil *Confusion Matrix* *skincare* Daviena

Gambar 18 menunjukkan *skincare* The Originote memiliki akurasi terendah di antara ketiga produk,

yaitu 83,72%. *Margin error* yang diperoleh adalah $\pm 2,76\%$, menunjukkan adanya kesalahan dalam pengambilan sampel. *Dataset* ini terdiri dari 218 sentimen positif dan 218 sentimen negatif. Presisi untuk prediksi positif adalah 76,73% dan untuk prediksi negatif 95,65%. Nilai *recall* menunjukkan bahwa model berhasil menemukan kembali informasi dengan rasio 96,79% untuk data positif dan 70,64% untuk data negatif. Meskipun akurasinya lebih rendah, model ini masih menunjukkan kemampuan yang baik dalam klasifikasi sentimen

4.9. Wordcloud

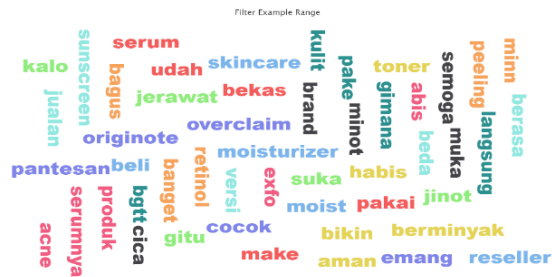
Wordcloud menampilkan *dataset* yang diambil dari kata-kata yang sering digunakan dalam komentar. Berikut ini merupakan hasil kata yang sering muncul atau yang sering digunakan dalam komentar Instagram.

Gambar 19. Wordcloud *skincare* Daviena

Gambar 19 menampilkan kata yang sering muncul pada *skincare* Daviena dengan jumlah muncul paling banyak pada kata “retinol” dengan jumlah muncul sebanyak 70, “pake” sebanyak 62, “beli” sebanyak 44, “serum” sebanyak 38, dan “daviena” sebanyak 35 kata yang sering muncul.

Gambar 20. Wordcloud *skincare* Scora

Gambar 20 menampilkan kata yang sering muncul pada *skincare* Scora dengan jumlah muncul paling banyak pada kata “scora” dengan jumlah muncul sebanyak 63, “serum” sebanyak 48, “pake” sebanyak 35, “beli” sebanyak 30, dan “kalo” sebanyak 20 kata yang sering muncul.



Gambar 21. Wordcloud skincare The Originote

Gambar 21 menampilkan kata yang sering muncul pada *skincare* The Originote dengan jumlah muncul paling banyak pada kata “pake” dengan jumlah muncul sebanyak 49, “serum” sebanyak 43, “beli” sebanyak 30, “produk” sebanyak 29, dan “udah” sebanyak 27 kata yang sering muncul.

4.10. Pembahasan

Daviena menunjukkan performa terbaik dalam analisis sentimen dengan akurasi model mencapai 94,22%. Model ini efektif dalam mengidentifikasi sentimen pengguna, didukung oleh presisi positif 89,96% dan *recall* positif 99,57%. Kata-kata yang sering muncul, seperti “retinol” dan “serum”, menunjukkan bahwa konsumen menghargai bahan aktif dalam produk. Keseimbangan antara sentimen positif dan negatif mencerminkan beragam pengalaman pengguna, penting untuk memahami penerimaan produk di pasar.

Scora memiliki akurasi model sebesar 87,39%, meskipun lebih rendah dibandingkan Daviena. Presisi positifnya mencapai 80,74% dan *recall* positif 98,20%, menunjukkan efektivitas dalam mengidentifikasi sentimen positif. Kata-kata dominan seperti “scora” dan “serum” menunjukkan kesadaran pengguna terhadap merk ini. Namun, Scora perlu meningkatkan beberapa aspek, terutama dalam klasifikasi sentimen, untuk memperbaiki penerimaan di kalangan konsumen.

The Originote mencatat akurasi terendah di antara ketiga produk, yaitu 83,72%, dengan presisi positif 76,73% dan *recall* positif 96,79%. Meskipun mampu menangkap sebagian besar sentimen positif, model ini sering salah mengklasifikasikan komentar positif, menunjukkan tantangan dalam analisis sentimen. Kata-kata umum seperti “produk” dan “pake” menunjukkan kurangnya pengenalan terhadap fitur spesifik. Untuk meningkatkan penerimaan, The Originote perlu fokus pada pengembangan fitur yang lebih dikenal dan meningkatkan akurasi klasifikasi sentimen.

Secara keseluruhan, analisis menunjukkan bahwa Serum HA Daviena merupakan produk dengan akurasi dan pemahaman sentimen pengguna terbaik, sementara Scora dan The Originote menghadapi tantangan yang perlu diatasi untuk meningkatkan penerimaan di kalangan konsumen.

5. KESIMPULAN DAN SARAN

Penelitian ini berhasil menerapkan algoritma *Decision Tree* untuk analisis sentimen di media sosial Instagram terhadap produk serum *skincare* Daviena, Scora, dan The Originote, menunjukkan efektivitas dalam mengidentifikasi pandangan konsumen. Hasil analisis menunjukkan keseimbangan sentimen di antara ketiga produk, dengan masing-masing produk memiliki jumlah sentimen positif dan negatif yang seimbang. Model analisis sentimen menunjukkan kinerja yang baik, dengan akurasi tertinggi pada Daviena sebesar 94,22%, diikuti oleh Scora sebesar 87,39%, dan The Originote sebesar 83,72%. Untuk Daviena, presisi positif mencapai 89,96% dan *recall* positif 99,57%, sedangkan presisi negatif adalah 99,52% dan *recall* negatif 88,89%. Untuk Scora, presisi positif adalah 80,74% dan *recall* positif 98,20%, dengan presisi negatif 97,70% dan *recall* negatif 76,58%.

Sementara itu, The Originote memiliki presisi positif 76,73% dan *recall* positif 96,79%, serta presisi negatif 95,65% dan *recall* negatif 70,64%. Kata-kata yang sering muncul dalam komentar mencerminkan pengalaman konsumen, seperti “retinol” untuk Daviena, “scora” untuk Scora, dan “pake” untuk The Originote. Untuk penelitian selanjutnya, disarankan untuk membandingkan beberapa metode atau algoritma, memperluas klasifikasi sentimen menjadi multikelas, serta memperpanjang rentang waktu pengumpulan ulasan. Selain itu, eksplorasi sumber data tambahan dari platform lain seperti TikTok dan X juga dianjurkan untuk mendapatkan gambaran yang lebih komprehensif mengenai sentimen pengguna terhadap produk *skincare*.

DAFTAR PUSTAKA

- [1] M. I. Putri and I. Kharisudin, “Analisis Sentimen Pengguna Aplikasi Marketplace Tokopedia Pada Situs Google Play Menggunakan Metode Support Vector Machine (SVM), Naïve Bayes , dan Logistic Regression,” vol. 5, pp. 759–766, 2022.
- [2] I. P. D. W. Darmawan, G. A. Pradnyana, and I. B. N. Pascima, “Optimasi Parameter Support Vector Machine Dengan Algoritma Genetika Untuk Analisis Sentimen Pada Media Sosial Instagram,” *SINTECH (Science Inf. Technol. J.*, vol. 6, no. 1, pp. 58–67, 2023, doi: 10.31598/sintechjournal.v6i1.1245.
- [3] C. Cahyaningtyas, Y. Nataliani, and I. R. Widiarsari, “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.
- [4] S. I. R. Adi, B. Bakkara, K. A. Zega, F. N. Vielita, and N. A. Rakhmawati, “Analisis Sentimen Masyarakat Terhadap Progress Ikn Menggunakan Model Decision Tree,” *JIKA (Jurnal Inform.*, vol. 8, no. 1, p. 57, 2024, doi:

- 10.31000/jika.v8i1.9803.
- [5] D. Azzahra, Meisa; Totohendarto and S. Moch, Hafid; Alam, “Analisis Sentimen Ulasan Produk Serum Wajah Pada Beauty Brand Somethinc Menggunakan Metode Naïve Bayes Classifier,” *JATI (Jurnal Mhs. Tek. Inform.*, vol. 7, no. 3, pp. 1604–1611, 2023, doi: 10.36040/jati.v7i3.6929.
 - [6] A. F. Setyaningsih, D. Septiyani, and S. R. Widiyari, “Implementasi Algoritma Naïve Bayes untuk Analisis Sentimen Masyarakat pada Twitter mengenai Kepopuleran Produk Skincare di Indonesia,” *J. Teknol. Inform. dan Komput.*, vol. 9, no. 1, pp. 224–235, 2023, doi: 10.37012/jtik.v9i1.1409.
 - [7] L. A. Pramesti and N. Pratiwi, “Analisis Sentimen Twitter Terhadap Program MBKM Menggunakan Decision Tree dan Support Vector Machine,” *J. Inf. Syst. Res.*, vol. 4, no. 4, pp. 1145–1154, 2023, doi: 10.47065/josh.v4i4.3807.
 - [8] S. Lutfiani, R. Astuti, and Fadhil M Basysyar, M. Kom, “Analisis Sentimen Pengaruh Media Sosial Terhadap Minat Beli Skincare Pada Remaja Di Indonesia Menggunakan Algoritma Naïve Bayes,” *JATI (Jurnal Mhs. Tek. Inform.*, vol. 8, no. 3, pp. 2957–2961, 2024, doi: 10.36040/jati.v8i3.9614.
 - [9] D. A. WP, J. D. Firizqi, and Z. A. Amalia, “Analisis Sentimen Produk Skincare Somethinc Niacinamide di Female Daily dengan Naïve Bayes Classifier,” *J. Media Inform. Budidarma*, vol. 8, no. 2, p. 946, 2024, doi: 10.30865/mib.v8i2.7571.
 - [10] Z. Munawar *et al.*, *Big Data Analytics: Konsep, Implementasi, dan Aplikasi Terkini*. Bandung (ID): Kaizen Media Publishing, 2023.
 - [11] J. . Nursiyono, *Machine Learning dengan R Teori dan Praktikum*. Malang (ID): Media Nusa Creative, 2023.
 - [12] R. . Sari, D. Pratama, M. Kurniawan, I. . Ginting, A. Fantika, and Y. Abhipraya, *Klasifikasi Data Mining*. Payakumbuh (ID): Serasi Media, 2024.
 - [13] T. G. Santoso, “Analisis Sentimen Pada Tweet Dengan Tagar #Bpjsrasarentenir Menggunakan Metode Support Vectore Machine (Svm),” (*Doctoral Diss. Univ. Islam Riau*), pp. 12–13, 2021.
 - [14] A. Nugroho, *Data Science Menggunakan Bahasa R Analisis Data, Visualisasi serta Pemodelan*. Yogyakarta (ID): Andi, 2022.
 - [15] P. . Rahayu *et al.*, *Buku Ajar Data Mining*. Jambi (ID): PT Sonpedia Publishing Indonesia, 2024.
 - [16] Y. Findawati and M. . Rosid, *Buku Ajar Text Mining*. Sidoarjo (ID): Umsida Press, 2020.
 - [17] T. . Rani and E. Sahputra, *Penerapan Machine Learning Terhadap Analisis Sentimen Masyarakat dalam Pemilihan Presiden 2024*. Pekalongan (ID): NEM, 2024.