

**ANALISIS SENTIMEN ULASAN PENGGUNA PADA PRODUK *FACE WASH*
SOMBONG MENGGUNAKAN *SUPPORT VECTOR MACHINE (SVM)***

SKRIPSI

Oleh :
MOHAMMAD ZIDAN ALVARO
NIM. 220605110114



**PROGRAM STUDI TEKNIK INFORMATIKA
FAKULTAS SAINS DAN TEKNOLOGI
UNIVERSITAS ISLAM NEGERI MAULANA MALIK IBRAHIM
MALANG
2026**

HALAMAN PERSETUJUAN

ANALISIS SENTIMEN ULASAN PENGGUNA PADA PRODUK *FACE WASH* SOMBONG MENGGUNAKAN *SUPPORT VECTOR MACHINE* (SVM)

SKRIPSI

Oleh :

MOHAMMAD ZIDAN ALVARO

NIM. 220605110114

Telah Diperiksa dan Disetujui untuk Diuji:

Tanggal: 23 Juni 2026

Pembimbing I,

Dr. M. Amin Hariyadi, M.T
NIP. 19670118 200501 1 001

Pembimbing II,

Okta Qomaruddin Aziz, M.Kom
NIP. 19911019 201903 1 013

Mengetahui,

Ketua Program Studi Teknik Informatika

Fakultas Sains dan Teknologi

Universitas Islam Negeri Maulana Malik Ibrahim Malang



Sumiyono, M. Kom
NIP. 19841010 201903 1 012

HALAMAN PENGESAHAN

ANALISIS SENTIMEN ULASAN PENGGUNA PADA PRODUK *FACE WASH* SOMBONG MENGGUNAKAN *SUPPORT VECTOR MACHINE* (SVM)

SKRIPSI

Oleh :

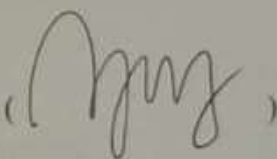
MOHAMMAD ZIDAN ALVARO

NIM. 220605110114

Telah Dipertahankan di Depan Dewan Penguji Skripsi
dan Dinyatakan Diterima Sebagai Salah Satu Persyaratan
Untuk Memperoleh Gelar Sarjana Komputer (S.Kom)
Tanggal: 23 Juni 2026

Susunan Dewan Penguji

Ketua Penguji	: <u>Dr. Agung Teguh Wibowo Almais, M.T</u> NIP. 19860301 202321 1 016
Anggota Penguji I	: <u>Tri Mukti Lestari, M.Kom</u> NIP. 19911108 202012 2 005
Anggota Penguji II	: <u>Dr. M. Amin Hariyadi, M.T</u> NIP. 19670118 200501 1 001
Anggota Penguji III	: <u>Okta Qomaruddin Aziz, M.Kom</u> NIP. 19911019 201903 1 013

()
()
()
()

Mengetahui dan Mengesahkan,
Ketua Program Studi Teknik Informatika
Fakultas Sains dan Teknologi
Universitas Islam Negeri Maulana Malik Ibrahim Malang



Supriyanto, M.Kom
NIP. 19841010 201903 1 012

PERNYATAAN KEASLIAN TULISAN

Saya yang bertanda tangan di bawah ini:

Nama : Mohammad Zidan Alvaro
NIM : 220605110114
Fakultas / Program Studi : Sains dan Teknologi / Teknik Informatika
Judul Skripsi : Analisis Sentimen Ulasan Pengguna Pada Produk
Face Wash Sombong Menggunakan Support
Vector Machine (Svm)

Menyatakan dengan seberanya bahwa Skirpsi yang saya tulis ini benar-benar merupakan hasil karya sendiri, bukan merupakan pengambil alihan data, tulisan, atau pikiran orang lain yang saya akui sebagai hasil tulisan atau pikiran saya sendiri, kecuali dengan mencantumkan sumber cuplikan pada daftar pustaka.

Apabila dikemudian hari terbukti atau dapat dibuktikan skripsian ini merupakan hasil jiplakan, maka saya bersedia menerima sanksi atas perbuatan tersebut.

Malang, 23 Juni 2026

Yang membuat pernyataan,



MOHAMMAD ZIDAN ALVARO
NIM.220605110114

MOTTO

“Kemajuan tidak diukur dari seberapa cepat kita sampai, tetapi dari seberapa jauh kita telah melangkah”

HALAMAN PERSEMBAHAN

Alhamdulillah rabbil 'alamin, segala puji bagi Allah Swt. atas segala rahmat, nikmat, dan karunia-Nya. Shalawat serta salam semoga senantiasa tercurah kepada Nabi Muhammad saw., beserta keluarga, sahabat, dan seluruh umatnya. Atas izin dan pertolongan Allah Swt., penulis diberikan kemudahan dan kekuatan hingga dapat menyelesaikan skripsi ini.

Skripsi ini penulis persembahkan kepada kedua orang tua tercinta sebagai wujud bakti, cinta, kasih sayang, dan rasa terima kasih atas segala doa, pengorbanan, serta dukungan yang telah diberikan. Kehadiran dan doa mereka menjadi alasan utama yang menguatkan penulis untuk terus berjuang hingga mampu menyelesaikan skripsi ini.

Kepada saudara-saudara tercinta dan seluruh keluarga besar, penulis mempersembahkan karya sederhana ini sebagai bentuk rasa syukur dan terima kasih atas perhatian, motivasi, dukungan, serta doa yang selalu menyertai perjalanan penulis dalam menempuh pendidikan hingga terselesaikannya skripsi ini.

Kepada teman-teman dan seluruh pihak yang telah turut membantu penulis dalam menyelesaikan skripsi ini, terima kasih atas doa, dukungan, dan segala bentuk bantuan yang diberikan sehingga skripsi ini dapat terselesaikan dengan baik.

KATA PENGANTAR

Assalamu 'alaikum Warahmatullahi Wabarakatuh

Alhamdulillāhirabbil 'ālamīn, segala puji bagi Allah *Subhānahu wa Ta 'ālā*.

yang telah melimpahkan rahmat, taufik, hidayah, serta karunia-Nya sehingga penulis dapat menyelesaikan skripsi yang berjudul “*Analisis Sentimen Ulasan Pengguna Pada Produk Face Wash Sombong Menggunakan Support Vector Machine (Svm)*” dengan baik. Shalawat serta salam semoga senantiasa tercurah kepada junjungan kita Nabi Muhammad *Shollallohu 'Alaihi Wasallam*, beserta keluarga, sahabat, dan seluruh pengikutnya hingga akhir zaman. Skripsi ini disusun sebagai salah satu syarat untuk memperoleh gelar Sarjana Komputer pada Program Studi Teknik Informatika, Fakultas Sains dan Teknologi, Universitas Islam Negeri Maulana Malik Ibrahim Malang. Oleh karena itu, pada kesempatan ini penulis ingin menyampaikan ucapan terima kasih yang sebesar-besarnya kepada:

1. Prof. Dr. Hj. Ilfi Nurdiana, M.Si., CAHRM., CRMP., selaku Rektor Universitas Islam Negeri Maulana Malik Ibrahim Malang.
2. Dr. Agus Mulyono, M.Kes., selaku Dekan Fakultas Sains dan Teknologi Universitas Islam Negeri Maulana Malik Ibrahim Malang.
3. Supriyono, M.Kom., selaku Ketua Program Studi Teknik Informatika Universitas Islam Negeri Maulana Malik Ibrahim Malang.
4. Dr. M. Amin Hariyadi, M.T., selaku dosen pembimbing I dan Okta Qomaruddin Aziz, M.Kom., selaku dosen pembimbing II, yang telah memberikan bimbingan, arahan, masukan, serta motivasi kepada penulis selama proses penyusunan skripsi ini. Terima kasih atas waktu, kesabaran,

dan perhatian yang telah diberikan sehingga penelitian dan penulisan skripsi ini dapat terselesaikan dengan baik dan tepat waktu.

5. Dr. Agung Teguh Wibowo Almais, M.T., selaku dosen penguji I dan Tri Mukti Lestari, M.Kom., selaku dosen penguji II, yang telah memberikan evaluasi, saran, serta berbagai masukan yang berharga sehingga penelitian dan penulisan skripsi ini dapat tersusun dengan lebih baik. Terima kasih atas arahan dan perhatian yang diberikan dalam proses penyusunan skripsi.
6. Seluruh dosen Program Studi Teknik Informatika Universitas Islam Negeri Maulana Malik Ibrahim Malang, penulis menyampaikan ucapan terima kasih atas ilmu yang telah di berikan selama masa perkuliahan. Semoga segala ilmu dan kebaikan yang telah diberikan menjadi amal jariyah serta memberikan manfaat yang berkelanjutan bagi penulis di masa yang akan datang.
7. Rasa terima kasih yang paling mendalam dan hormat penulis persembahkan kepada Ibunda tercinta Eva Juli Artanti. Terima kasih atas kasih sayang yang tak bertepi, kehangatan, serta kesabaran yang luar biasa dalam mendidik dan menemani setiap langkah hidup penulis. Untaian doa tulus yang selalu Ibu panjatkan di setiap sujud adalah kekuatan dan pelindung terbesar bagi penulis hingga akhirnya mampu menyelesaikan skripsi dan bangku perkuliahan ini. Karya ini adalah wujud bakti kecil yang penulis dedikasikan untuk senantiasa mengukir senyum di wajah Ibu.
8. Rasa terima kasih yang tak terhingga dan hormat yang setinggi-tingginya penulis haturkan kepada Ayahanda tercinta Mokhammad Ali Abidin.

Terima kasih atas segala pengorbanan, kerja keras yang tak kenal lelah, serta dukungan moril maupun materil yang telah Ayah berikan demi masa depan penulis. Ketegasan, teladan, dan kepercayaan yang Ayah tanamkan selalu menjadi fondasi bagi penulis untuk tetap tegar dan pantang menyerah dalam menghadapi setiap tantangan selama masa studi ini.

9. Seluruh keluarga besar penulis yang tidak dapat disebutkan satu per satu, terima kasih atas doa, dukungan, perhatian, dan semangat yang telah diberikan kepada penulis selama menempuh pendidikan.
10. Sahabat-sahabat seperjuangan penulis, Muhammad Andrean Hidayatulloh, Ahmad Fadhli Dzilikrom, Muhammad Alif Atallah, Ki Ageng Nasrokh Mangkunegara Putra, dan Muhhammad Annas Darunnaja. Terima kasih atas kebersamaan, dukungan, dan semangat yang tiada henti selama masa perkuliahan. Terima kasih telah menjadi teman bertukar pikiran, melewati suka duka pengerjaan tugas dan skripsi.
11. Ucapan terima kasih penulis sampaikan kepada keluarga besar Teknik Informatika angkatan 2022 Infinity, yang telah memberikan kebersamaan, dukungan, dan pengalaman berharga selama masa perkuliahan.
12. Rafi Abdillah Tsani dan Muhammad Syauqi Jauhar Ramzy, terima kasih yang sebesar-besarnya atas kehadiran dan kesetiaannya menemani penulis melewati masa-masa sulit selama proses penyusunan skripsi ini. Terima kasih telah menjadi pendengar yang baik, mengulurkan bantuan, serta selalu punya cara untuk membantu menghilangkan stres dan mengembalikan

semangat penulis di saat jenuh melanda. Keberadaan kalian sangat berarti dalam penyelesaian karya ini.

13. Mirnawan, Firsya Syaikh, dan Ridho Baihakhi, terima kasih banyak atas pertemanan yang berkesan, dukungan, serta bantuan yang telah diberikan kepada penulis selama masa perkuliahan ini. Terima kasih telah meluangkan waktu untuk berbagi cerita, memberikan motivasi, dan turut mewarnai perjalanan dinamika kampus hingga penulis bisa mencapai titik ini.

Penulis menyadari bahwa dalam penyusunan skripsi ini masih jauh dari kesempurnaan, baik dari segi penyajian materi maupun tata bahasa, karena keterbatasan pengetahuan dan pengalaman yang penulis miliki. Oleh karena itu, penulis mengharapkan kritik dan saran yang membangun dari semua pihak demi perbaikan di masa yang akan datang. Akhir kata, semoga skripsi ini dapat memberikan manfaat, menambah wawasan, serta menjadi kontribusi positif bagi pembaca dan perkembangan ilmu pengetahuan.

Wassalamu'alaikum Warahmatullahi Wabarakatuh

Malang, 23 Juni 2026

Penulis

DAFTAR ISI

HALAMAN PERSETUJUAN	iii
HALAMAN PENGESAHAN.....	iv
PERNYATAAN KEASLIAN TULISAN	v
MOTTO	vi
HALAMAN PERSEMBAHAN	vii
KATA PENGANTAR.....	viii
DAFTAR ISI.....	xii
DAFTAR GAMBAR.....	xiv
DAFTAR TABEL	xvi
ABSTRAK	xvii
ABSTRACT	xviii
مستخلص البحث.....	xix
BAB I PENDAHULUAN.....	1
1.1 Latar Belakang	1
1.2 Identifikasi Masalah.....	4
1.3 Batasan Masalah.....	5
1.4 Tujuan Penelitian	5
1.5 Manfaat Penelitian	6
BAB II STUDI PUSTAKA	7
2.1 Penelitian Terkait	7
2.2 Analisis Sentimen	11
2.3 <i>TermFrequency –Inverse Document Frequency (TF-IDF)</i>	13
2.4 Algoritma <i>Support Vector Machine (SVM)</i>	14
2.5 Evaluasi Hasil.....	17
BAB III DESAIN DAN IMPLEMENTASI	19
3.1 Desain Sistem.....	19
3.2 Akuisisi Data.....	20
3.3 Pelabelan Data (<i>Data Labeling</i>).....	23
3.4 Prapemrosesan Teks (<i>Text Processing</i>)	26
3.5 Pembobotan Fitur TF-IDF (<i>TermFrequency-Inverse Document Frequency</i>). ..	28
3.6 <i>Support Vector Machine (SVM)</i>	31
3.6.1 Pembagian Data (<i>Data Splitting</i>).....	33
3.6.2 Fase Pelatihan (<i>Training Phase</i>).....	34
3.6.3 Fase Pengujian (<i>Testing Phase</i>).....	35
3.6.4 Fungsi Kernel SVM dan Perhitungan Manual	35
3.7 Evaluasi Hasil (<i>Confusion matrix</i>)	37
3.8 Implementasi Alur Tahapan Sistem	39
3.8.1 Pengumpulan Data.....	39
3.8.2 Preprocessing	43
3.8.3 Data Labeling dan Balancing	47
3.8.4 Ekstraksi Fitur TF-IDF	50
3.9 Skenario Pengujian.....	51

BAB IV HASIL DAN PEMBAHASAN	54
4.1 Hasil Implementasi Sistem.....	54
4.1.1 Hasil Pengumpulan Data	54
4.1.2 Hasil Preprocessing Data.....	58
4.1.3 Hasil <i>Labeling</i> dan <i>Balancing</i> Data.....	63
4.1.4 Hasil Ekstrasi Fitur TF-IDF.....	66
4.2 Hasil Pengujian Model.....	68
4.2.1 Hasil Skenario 1.....	69
4.2.2 Hasil Skenario 2.....	79
4.2.3 Hasil Skenario 3.....	90
4.2.4 Hasil Skenario 4.....	101
4.2.5 Rangkuman Hasil Keseluruhan	111
4.3 Implementasi Antarmuka Sistem	115
4.4 Pembahasan.....	118
4.4.1 Analisis Pengaruh Kernel terhadap Akurasi.....	120
4.4.2 Analisis Pengaruh Rasio Split Data terhadap Akurasi	124
BAB V KESIMPULAN DAN SARAN	131
5.1 Kesimpulan	131
5.2 Saran.....	132
DAFTAR PUSTAKA	

DAFTAR GAMBAR

Gambar 3. 1 Desain Sistem.....	20
Gambar 3. 2 Potongan Kode Inisialisasi URL Produk	40
Gambar 3. 3Potongan Kode Automasi Peramban.....	41
Gambar 3. 4 Potongan Kode Ekstraksi Teks Ulasan	42
Gambar 3. 5 Potongan Kode Kontrol Perulangan Halaman (Pagination)	42
Gambar 3. 6 Potongan Kode Pembuatan DataFrame dan Ekspor ke CSV	43
Gambar 3. 7 Potongan Kode Fungsi <i>Case Folding</i> dan <i>Cleansing</i>	44
Gambar 3. 8 Potongan Kode Fungsi <i>Tokenizing</i>	45
Gambar 3. 9 Potongan Kode Fungsi <i>Stopword Removal</i>	46
Gambar 3. 10 Potongan Kode Fungsi <i>Stemming</i>	47
Gambar 3. 11 Potongan Kode Inisialisasi Kamus Leksikon Positif dan Negatif..	48
Gambar 3. 12 Potongan Kode Fungsi Perhitungan Skor dan Pelabelan Sentimen	49
Gambar 3. 13 Potongan Kode Penyeimbangan Distribusi Kelas Menggunakan <i>Random Over-Sampler (ROS)</i>	50
Gambar 3. 14 Potongan Kode Inisialisasi dan Transformasi Fitur Menggunakan TF-IDF	51
Gambar 4. 1 Tampilan Proses Inisialisasi URL Target pada Terminal	55
Gambar 4. 2 Tampilan Automasi Peramban Tokopedia Selama Proses Scraping	56
Gambar 4. 3 Log Proses Penjelajahan Halaman dan Ekstraksi Ulasan	57
Gambar 4. 4 Tampilan Log Selesai Proses dan Penyimpanan Berkas .csv	58
Gambar 4. 5 Tampilan Hasil Proses <i>Case Folding</i> dan <i>Cleansing</i>	59
Gambar 4. 6 Tampilan Hasil Proses <i>Tokenizing</i>	60
Gambar 4. 7 Tampilan Hasil Proses <i>Stopword Removal</i>	61
Gambar 4. 8 Tampilan Hasil Proses <i>Stemming</i>	62
Gambar 4. 9 Tampilan Hasil Akhir Keseluruhan Text Preprocessing.....	63
Gambar 4. 10 Diagram Hasil Proses Pelabelan Data	64
Gambar 4. 11 Diagram Hasil Penyeimbangan Data dengan Random Over- Sampling	65
Gambar 4. 12 Hasil Representasi Matriks Bobot Nilai TF-IDF	66
Gambar 4. 13 <i>Confussion Matrix</i> Skenario 1a	71
Gambar 4. 14 <i>Confussion Matrix</i> Skenario 1b.....	73
Gambar 4. 15 <i>Confussion Matrix</i> Skenario 1c	75
Gambar 4. 16 <i>Confussion Matrix</i> Skenario 1d.....	77
Gambar 4. 17 <i>Confussion Matrix</i> Skenario 2a	81
Gambar 4. 18 <i>Confussion Matrix</i> Skenario 2b.....	83
Gambar 4. 19 <i>Confussion Matrix</i> Skenario 2c	86
Gambar 4. 20 <i>Confussion Matrix</i> Skenario 2d.....	88
Gambar 4. 21 <i>Confussion Matrix</i> Skenario 3a	92
Gambar 4. 22 <i>Confussion Matrix</i> Skenario 3b.....	94
Gambar 4. 23 <i>Confussion Matrix</i> Skenario 3c	97
Gambar 4. 24 <i>Confussion Matrix</i> Skenario 3d.....	99
Gambar 4. 25 <i>Confussion Matrix</i> Skenario 4a	103
Gambar 4. 26 <i>Confussion Matrix</i> Skenario 4b.....	105
Gambar 4. 27 <i>Confussion Matrix</i> Skenario 4c	107

Gambar 4. 28 <i>Confussion Matrix</i> Skenario 4d.....	110
Gambar 4. 29 Tampilan Antarmuka Web	115
Gambar 4. 30 Tampilan Hasil Pengujian Sentimen Positif pada Antarmuka Web	117
Gambar 4. 31 Tampilan Hasil Pengujian Sentimen Negatif pada Antarmuka Web	118
Gambar 4. 32 Visualisasi Rentang dan Variabilitas Performa Kernel SVM	121
Gambar 4. 33 Visualisasi Rentang dan Variabilitas Performa Rasio Data Latih	125

DAFTAR TABEL

Tabel 2. 1 Perbandingan Penelitian Terdahulu	8
Tabel 3. 1 Contoh Dataset Ulasan Tokopedia.....	21
Tabel 3. 2 Contoh Hasil Pelabelan Data Sentimen	24
Tabel 3. 3 Contoh Perbandingan Hasil Prapemrosesan Teks	27
Tabel 3. 4 Sampel Data untuk Simulasi Perhitungan TF-IDF	29
Tabel 3. 5 Contoh Ringkasan Perhitungan TF-IDF Berdasarkan Keseluruhan Sampel (N=6)	30
Tabel 3. 6 Skenario Eksperimen Berdasarkan Jenis Kernel	52
Tabel 4. 1 Rentang Nilai Hyperparameter	69
Tabel 4. 2 Metrik Performa Klasifikasi Skenario 1a.....	71
Tabel 4. 3 Metrik Performa Klasifikasi Skenario 1b	74
Tabel 4. 4 Metrik Performa Klasifikasi Skenario 1c.....	76
Tabel 4. 5 Metrik Performa Klasifikasi Skenario 1d	78
Tabel 4. 6 Metrik Performa Klasifikasi Skenario 2a.....	82
Tabel 4. 7 Metrik Performa Klasifikasi Skenario 2b	84
Tabel 4. 8 Metrik Performa Klasifikasi Skenario 2c.....	87
Tabel 4. 9 Metrik Performa Klasifikasi Skenario 2d	89
Tabel 4. 10 Metrik Performa Klasifikasi Skenario 3a.....	92
Tabel 4. 11 Metrik Performa Klasifikasi Skenario 3b	95
Tabel 4. 12 Metrik Performa Klasifikasi Skenario 3c.....	98
Tabel 4. 13 Metrik Performa Klasifikasi Skenario 3d	100
Tabel 4. 14 Metrik Performa Klasifikasi Skenario 4a.....	103
Tabel 4. 15 Metrik Performa Klasifikasi Skenario 4b	106
Tabel 4. 16 Metrik Performa Klasifikasi Skenario 4c.....	108
Tabel 4. 17 Metrik Performa Klasifikasi Skenario 4d	111
Tabel 4. 18 Rangkuman Hasil Pengujian	112
Tabel 4. 19 Perbandingan Nilai Akurasi dan Standar Deviasi pada Setiap Kernel	121
Tabel 4. 20 Hasil t-test (p-value) Antar Kernel.....	123
Tabel 4. 21 Perbandingan Nilai Akurasi dan Standar Deviasi pada Setiap Rasio Data	125
Tabel 4. 22 Hasil t-test (p-value) Antar Rasio Data.....	127

ABSTRAK

Alvaro, Mohammad Zidan. 2026. **Analisis Sentimen Ulasan Pengguna pada Produk Face Wash SOMBONG di E-Commerce Menggunakan Algoritma Support Vector Machine (SVM)**. Skripsi. Jurusan Teknik Informatika Fakultas Sains dan Teknologi Universitas Islam Negeri Maulana Malik Ibrahim Malang. Pembimbing: (I) Dr. M. Amin Hariyadi, M.T. (II) Okta Qomaruddin Aziz, M.Kom.

Kata Kunci: Analisis Sentimen, *Support Vector Machine*, *TF-IDF*, *Random Over Sampling*, *E-Commerce*

Platform *e-commerce* menjadi sarana utama bagi konsumen untuk menyampaikan pengalaman dan penilaian terhadap suatu produk melalui ulasan. Besarnya jumlah ulasan yang tersebar menyebabkan proses analisis secara manual menjadi kurang efektif sehingga diperlukan metode otomatis untuk mengidentifikasi sentimen pengguna. Penelitian ini bertujuan menganalisis performa algoritma *Support Vector Machine* (SVM) dalam mengklasifikasikan sentimen ulasan pengguna pada produk Face Wash SOMBONG ke dalam kategori positif dan negatif. Data penelitian diperoleh dari platform Tokopedia dan Shopee melalui proses *web scraping* dengan total 4.346 ulasan berbahasa Indonesia. Tahapan pengolahan data meliputi *preprocessing*, pelabelan data menggunakan pendekatan *lexicon-based*, penyeimbangan data menggunakan *Random Over-Sampling* (ROS), pembobotan fitur menggunakan *TF-IDF*, serta klasifikasi menggunakan algoritma SVM dengan kernel *Linear*, *Radial Basis Function* (RBF), *Polynomial*, dan *Sigmoid*. Pengujian dilakukan pada empat skenario pembagian data, yaitu 60:40, 70:30, 80:20, dan 90:10. Hasil penelitian menunjukkan bahwa seluruh kernel mampu menghasilkan performa klasifikasi yang sangat baik dengan akurasi di atas 97%. Kernel RBF memperoleh rata-rata akurasi tertinggi sebesar 99,70%, sedangkan pada rasio pembagian data 90:10, kernel *Linear* dan kernel RBF dengan parameter $C = 10$ berhasil mencapai performa sempurna dengan nilai *Accuracy*, *Precision*, *Recall*, dan *F1-Score* sebesar 100%. Hasil tersebut menunjukkan bahwa algoritma *Support Vector Machine* mampu mengklasifikasikan sentimen ulasan pengguna secara sangat efektif serta berpotensi digunakan sebagai alat pendukung evaluasi kualitas produk berdasarkan opini konsumen pada platform *e-commerce*.

ABSTRACT

Alvaro, Mohammad Zidan. 2026. **Sentiment Analysis of User Reviews on Face Wash SOMBONG Products in E-Commerce Using the Support Vector Machine (SVM) Algorithm**. Undergraduate Thesis. Department of Informatics Engineering, Faculty of Science and Technology, Maulana Malik Ibrahim State Islamic University of Malang. Advisors: (I) Dr. M. Amin Hariyadi, M.T. (II) Okta Qomaruddin Aziz, M.Kom.

Keywords: Sentiment Analysis, Support Vector Machine, TF-IDF, Random Over-Sampling, E-Commerce.

E-commerce platforms have become a primary medium for consumers to share their experiences and evaluations of products through reviews. The large volume of reviews makes manual analysis inefficient; therefore, an automated method is required to identify user sentiments. This study aims to analyze the performance of the Support Vector Machine (SVM) algorithm in classifying user reviews of Face Wash SOMBONG products into positive and negative sentiments. The data were collected from Tokopedia and Shopee through a web scraping process, resulting in 4,346 Indonesian-language reviews. The data processing stages included preprocessing, data labeling using a lexicon-based approach, data balancing using Random Over-Sampling (ROS), feature weighting using TF-IDF, and classification using SVM with Linear, Radial Basis Function (RBF), Polynomial, and Sigmoid kernels. Experiments were conducted using four data split scenarios, namely 60:40, 70:30, 80:20, and 90:10. The results showed that all kernels achieved excellent classification performance with accuracy values above 97%. The RBF kernel obtained the highest average accuracy of 99.70%, while under the 90:10 data split ratio, both the Linear and RBF kernels with $C = 10$ achieved perfect performance with Accuracy, Precision, Recall, and F1-Score values of 100%. These findings indicate that the Support Vector Machine algorithm is highly effective in classifying user review sentiments and has the potential to serve as a decision-support tool for product quality evaluation based on consumer opinions in e-commerce platforms.

مستخلص البحث

ألفارو، محمد زيدان. ٢٠٢٦. تحليل المشاعر لآراء المستخدمين حول منتج غسول الوجه SOMBONG في منصات التجارة الإلكترونية باستخدام خوارزمية آلة المتجهات الداعمة (SVM). بحث جامعي. قسم هندسة المعلوماتية، كلية العلوم والتكنولوجيا، جامعة مولانا مالك إبراهيم الإسلامية الحكومية مالانج. المشرفان: (١) د. محمد أمين هريادي، الماجستير. (٢) أوكتا قمر الدين عزيز، الماجستير.

الكلمات المفتاحية: تحليل المشاعر، آلة المتجهات الداعمة، TF-IDF، المعايينة العشوائية الزائدة، التجارة الإلكترونية..

أصبحت منصات التجارة الإلكترونية وسيلة رئيسية للمستهلكين للتعبير عن تجاربهم وتقييماتهم للمنتجات من خلال المراجعات والتعليقات، ومع تزايد أعداد هذه المراجعات أصبح تحليلها يدويًا أمرًا غير فعال، مما يستدعي استخدام أساليب آلية لتحديد مشاعر المستخدمين. هدفت هذه الدراسة إلى تحليل أداء خوارزمية آلة المتجهات الداعمة (SVM) في تصنيف آراء مستخدمي منتج غسول الوجه SOMBONG إلى فئتين: إيجابية وسلبية. جُمعت بيانات الدراسة من منصتي توكوبيديا وشوبي باستخدام تقنية استخراج البيانات من الويب، حيث بلغ عدد المراجعات باللغة الإندونيسية ٤٣٤٦ مراجعة. تضمنت مراحل المعالجة التنقية المسبقة للبيانات، وتصنيف البيانات باستخدام منهج قائم على المعجم، ومعالجة عدم توازن البيانات باستخدام أسلوب المعايينة العشوائية الزائدة، ثم وزن الخصائص باستخدام TF-IDF، وأخيرًا التصنيف باستخدام خوارزمية SVM بأربعة أنواع من النوى هي: الخطية، ودالة الأساس الشعاعي (RBF)، ومتعددة الحدود، والسيغمويد. أُجريت التجارب باستخدام أربعة سيناريوهات لتقسيم البيانات وهي 60:40 و 70:30 و 80:20 و 90:10. وأظهرت النتائج أن جميع أنواع النوى حققت أداءً تصنيفيًا ممتازًا بدقة تجاوزت 97%، كما حققت نواة RBF أعلى متوسط دقة بلغ 99.70%، بينما حققت النواتان الخطية و RBF عند نسبة تقسيم 90:10 ومعامل $C = 10$ أداءً مثاليًا، حيث بلغت قيم الدقة والإحكام والاسترجاع ودرجة F1 نسبة 100%. وتشير هذه النتائج إلى أن خوارزمية آلة المتجهات الداعمة فعالة للغاية في تصنيف مشاعر المستخدمين، وتمتلك إمكانات كبيرة لدعم تقييم جودة المنتجات اعتمادًا على آراء المستهلكين في منصات التجارة الإلكترونية.

BAB I

PENDAHULUAN

1.1 Latar Belakang

Perkembangan teknologi informasi saat ini berdampak positif terhadap berbagai bidang, salah satunya adalah sektor ekonomi melalui kemudahan berbelanja secara daring. Banyak platform *e-commerce* yang diciptakan untuk memudahkan konsumen dalam bertransaksi tanpa batasan ruang dan waktu (Milal *et al.*, 2023). Fenomena ini menjadikan *e-commerce* sebagai sarana utama masyarakat Indonesia dalam memenuhi kebutuhan sehari-hari, didorong oleh efisiensi serta keberagaman produk yang ditawarkan (Puspa *et al.*, 2023).

Seiring dengan meningkatnya penetrasi internet, industri kecantikan dan perawatan diri atau *skincare* menjadi salah satu kategori produk yang paling diminati di platform digital. Hal ini memicu persaingan ketat antar produsen untuk meluncurkan produk inovatif, termasuk merek lokal yang sedang populer yaitu *Face Wash SOMBONG*. Sebagai produk yang menarik perhatian masif di pasar digital, ulasan konsumen pada platform seperti Tokopedia menjadi sumber informasi krusial yang mencerminkan kualitas serta kepuasan pengguna (Demetria & Wedhasmara, 2025).

Dalam ekosistem *e-commerce*, fitur ulasan produk berfungsi memberikan kepercayaan kepada calon pembeli sekaligus memberikan bahan evaluasi bagi penjual (Milal *et al.*, 2023). Ulasan tersebut merupakan bentuk *electronic word of mouth* (e-WOM) yang sangat berpengaruh terhadap keputusan pembelian, karena

mengandung opini nyata mengenai efektivitas dan pengalaman penggunaan produk (Tjut Awaliyah Zuraiyah, Mulyati, 2023). Namun, informasi yang terkandung dalam ulasan sering kali bersifat ambigu antara teks komentar dengan rating bintang yang diberikan, sehingga menyulitkan penilaian kualitas produk secara objektif (Demetria & Wedhasmara, 2025).

Tantangan utama muncul ketika popularitas produk *Face Wash SOMBONG* menghasilkan tumpukan ulasan yang mencapai ribuan data, sehingga tidak memungkinkan bagi produsen maupun konsumen untuk memilahnya secara manual. Analisis manual memerlukan waktu yang lama dan rentan terhadap subjektivitas manusia (Puspa *et al.*, 2023). Oleh karena itu, diperlukan teknik text mining melalui analisis sentimen untuk mengekstraksi pendapat dan emosi dari ulasan tertulis secara otomatis ke dalam kategori positif, negatif, atau netral.

Supaya data teks ulasan dapat dieksekusi oleh sistem, penelitian ini menerapkan metode *TermFrequency-Inverse Document Frequency* (TF-IDF) sebagai teknik representasi fitur. TF-IDF bekerja dengan memberikan bobot pada kata berdasarkan frekuensi kemunculannya dalam sebuah ulasan dan tingkat kelangkaannya di seluruh dataset (Fahmi Fiddin, Taufik Hidayat, 2025). Penggunaan TF-IDF sangat krusial dalam klasifikasi teks karena mampu menonjolkan kata-kata yang memiliki nilai informasi tinggi terhadap konteks sentiment (Alaiya & Agusniar, 2025).

Algoritma yang digunakan dalam penelitian ini adalah *Support Vector Machine* (SVM). Pemilihan SVM didasarkan pada kemampuannya yang sangat baik dalam menangani data berdimensi tinggi dan bersifat sparse seperti data teks

hasil pembobotan TF-IDF (Adrian *et al.*, 2021). SVM bekerja dengan mencari *hyperplane* optimal yang memisahkan dua kelas atau lebih dengan margin maksimal, sehingga dikenal memiliki performa akurasi yang stabil dalam berbagai tugas klasifikasi teks di dunia industri (Adrian *et al.*, 2021).

Berdasarkan beberapa studi literatur, SVM terbukti sering kali memberikan hasil yang lebih unggul dibandingkan algoritma lainnya. Penelitian terdahulu menunjukkan bahwa penggunaan model SVM dengan representasi TF-IDF mampu memberikan tingkat akurasi yang kompetitif pada ulasan produk di platform Tokopedia (Fahmi Fiddin, Taufik Hidayat, 2025). Selain itu, SVM juga dinilai efektif dalam mengklasifikasikan data ulasan yang diperoleh melalui teknik *web scraping* dengan parameter yang dioptimalkan (Alaiya & Agusniar, 2025).

Kebaruan dalam penelitian ini terletak pada penggunaan data segar hasil *scraping* mandiri dari platform Tokopedia untuk menganalisis sentimen spesifik pada produk lokal yang sedang viral, yaitu *Face Wash SOMBONG*. Berbeda dengan dataset publik yang bersifat umum, data *scraping* ini mencerminkan kondisi pasar terkini dan opini riil pengguna terhadap efektivitas produk tersebut (Idris, 2023). Penelitian ini diharapkan dapat membantu produsen dalam memantau persepsi publik secara sistematis.

Prinsip tersebut selaras dengan konsep *tabayyun* sebagaimana yang diamanatkan oleh Allah SWT dalam Al-Qur'an Surah Al-Hujurat ayat 6:

يَا أَيُّهَا الَّذِينَ ءَامَنُوا إِن جَاءَكُمْ فَاسِقٌ بِنَبَأٍ فَتَبَيَّنُوا أَن تُصِيبُوا قَوْمًا بِجَهْلَةٍ فَتُصْحَبُوا عَلَىٰ مَا فَعَلْتُمْ
نَدِيمِينَ

“Wahai orang-orang yang beriman! Jika datang kepadamu orang fasik membawa suatu berita, maka periksalah dengan teliti agar kamu tidak menimpakan suatu musibah kepada suatu kaum karena kebodohan (kecerobohan) yang akhirnya kamu menyesali perbuatanmu itu”. (QS. Al-Hujurat:6)

Ayat tersebut menekankan pentingnya melakukan verifikasi dan penelitian mendalam terhadap setiap informasi yang masuk sebelum mempercayainya sebagai fakta. Dalam konteks ini, implementasi analisis sentimen menggunakan algoritma SVM dipandang sebagai upaya sistematis untuk menjalankan prinsip tabayyun tersebut di dunia digital. Melalui pendekatan ilmiah ini, diharapkan penelitian dapat berkontribusi secara praktis dalam menjembatani kebutuhan produsen dan konsumen untuk memahami dinamika persepsi serta ulasan nyata yang ada di lapangan mengenai produk *Face Wash SOMBONG* secara transparan dan terukur.

Melalui pendekatan TF-IDF dan SVM untuk analisis sentimen produk lokal spesifik, studi ini diharapkan dapat memberikan sumbangsih nyata bagi perkembangan literatur ilmiah di bidang terkait, serta memberikan nilai praktis dalam memahami karakteristik dan kecenderungan sentimen pengguna terhadap produk *Face Wash SOMBONG* di berbagai platform *e-commerce*.

1.2 Identifikasi Masalah

Identifikasi masalah dalam penelitian ini adalah bagaimana penerapan metode *Term Frequency-Inverse Document Frequency* (TF-IDF) dengan algoritma *Support Vector Machine* (SVM) dalam mengklasifikasikan sentimen ulasan pengguna pada produk *Face Wash SOMBONG* di berbagai platform *E-commerce*.

1.3 Batasan Masalah

Ruang lingkup penelitian hanya fokus pada penggunaan ulasan pelanggan berbahasa Indonesia yang diperoleh melalui teknik *web scraping* dari platform Tokopedia. Analisis sentimen dibatasi secara spesifik pada pengolahan data teks ulasan, tanpa mempertimbangkan variabel tambahan seperti rating bintang, informasi varian produk, maupun data logistik pengiriman. Selain itu, proses klasifikasi sentimen dalam penelitian ini dibatasi secara spesifik pada pengolahan data teks ulasan dengan klasifikasi yang difokuskan pada dua kategori polaritas saja, yaitu positif dan negatif, guna memastikan fokus evaluasi tetap berada pada akurasi pemisahan sentimen di dalam ruang fitur yang berdimensi tinggi.

1.4 Tujuan Penelitian

Tujuan utama dari penelitian ini adalah untuk menerapkan metode pembobotan fitur *Term Frequency-Inverse Document Frequency* (TF-IDF) dan algoritma *Support Vector Machine* (SVM) dalam melakukan klasifikasi sentimen ulasan pengguna secara otomatis. Penelitian ini bertujuan untuk melakukan analisis komparasi performa antara fungsi kernel *Linear*, *RBF*, *Polynomial*, dan *Sigmoid* guna menentukan model klasifikasi yang paling optimal. Selain itu, penelitian ini bertujuan untuk mengukur tingkat keandalan sistem melalui evaluasi metrik akurasi, presisi, *recall*, dan *f1-score* pada berbagai skenario pembagian data untuk memastikan stabilitas hasil klasifikasi yang dihasilkan.

1.5 Manfaat Penelitian

Penelitian ini mengintegrasikan metode pembobotan TF-IDF untuk representasi teks dan algoritma *Support Vector Machine* (SVM) guna membangun model klasifikasi sentimen pada ulasan produk *Face Wash SOMBONG*. Hasil dari penelitian ini diharapkan dapat memberikan kontribusi nyata, antara lain:

1. Menyajikan pemetaan komprehensif mengenai kecenderungan opini konsumen di *e-commerce*, yang dapat digunakan oleh produsen sebagai bahan evaluasi serta strategi pengembangan produk *Face Wash SOMBONG* ke depan.
2. Menjadi rujukan ilmiah bagi peneliti selanjutnya yang ingin mendalami bidang analisis sentimen berbasis teks, khususnya dalam menerapkan kombinasi metode TF-IDF dan algoritma SVM.

BAB II

STUDI PUSTAKA

2.1 Penelitian Terkait

Berbagai pendekatan algoritma telah banyak diimplementasikan dalam sejumlah penelitian terdahulu untuk memetakan dan menganalisis sentimen dari data ulasan konsumen di platform *e-commerce*. (Puspa *et al.*, 2023) melakukan analisis sentimen pada ulasan produk di Tokopedia menggunakan *Support Vector Machine* (SVM) dan menemukan bahwa SVM efektif dalam mengklasifikasikan opini pelanggan dengan hasil yang stabil. (Milal *et al.*, 2023) membandingkan Naive Bayes dan SVM pada ulasan Shopee, di mana SVM menunjukkan performa yang kompetitif dalam menangani teks bahasa Indonesia yang tidak baku.

Penelitian lain oleh (Fahmi Fiddin, Taufik Hidayat, 2025) menerapkan SVM dengan representasi fitur TF-IDF pada ulasan produk di Tokopedia, yang menunjukkan bahwa kombinasi tersebut sangat handal dalam menangani data teks berdimensi tinggi. Sementara itu, (Alaiya & Agusniar, 2025) memfokuskan penelitian pada ulasan produk *skincare* dan menegaskan bahwa SVM memiliki keunggulan dalam memisahkan kelas sentimen yang bersifat ambigu. Secara umum, berbagai studi literatur menunjukkan bahwa SVM merupakan algoritma yang sering memberikan akurasi lebih tinggi dibandingkan metode tradisional lainnya dalam domain *e-commerce*.

Selain penelitian-penelitian di atas, terdapat studi terbaru yang secara spesifik membahas produk yang sama dengan penelitian ini. (Bima maulana, 2025)

melakukan analisis sentimen terhadap ulasan produk *Facial Wash Sombong* menggunakan algoritma *Naïve Bayes Classification* (NBC) dengan memanfaatkan 5.000 data komentar dari platform TikTok. Selanjutnya, (Arbiansyah & Haq, 2025) meneliti produk kompetitor sejenis, yaitu *Kahf Face Wash*, menggunakan metode K-Means Clustering untuk memetakan pola kepuasan pelanggan di platform Tokopedia.

Penelitian ini memiliki perbedaan signifikan dibandingkan studi terdahulu tersebut. Jika penelitian (Bima maulana, 2025) berfokus pada platform media sosial TikTok dengan algoritma *Naïve Bayes*, penelitian ini menggunakan platform *e-commerce* Tokopedia yang datanya berasal dari pembeli aktual. Keterbatasan Penggunaan algoritma *Support Vector Machine* (SVM) dalam memproses data teks berdimensi tinggi coba diatasi melalui penelitian ini. Pendekatan alternatif yang diterapkan diharapkan dapat memberikan hasil klasifikasi sentimen dengan akurasi yang lebih optimal serta terukur.

Tabel 2. 1 Perbandingan Penelitian Terdahulu

Peneliti & Judul	Metode	Hasil Penelitian	Gap Penelitian
(Bima maulana, 2025) Analisis Sentimen terhadap Ulasan Produk Facial Wash Sombong Menggunakan Algoritma Naive Bayes	Naive Bayes Classifier (NBC)	Dominasi opini negatif sebesar 61,4% dari 5.000 ulasan TikTok.	penggunaan algoritma SVM yang lebih handal dalam menangani data teks berdimensi tinggi serta fokus pada data Tokopedia, serta melakukan komparasi mendalam pada 4 jenis kernel yang berbeda.
(Alaiya & Agusniar, 2025)	SVM, TF-IDF, Lexicon-based	Menghasilkan akurasi yang baik pada 571 ulasan di Tokopedia.	Penelitian tersebut menggunakan jumlah data yang relatif kecil (571 ulasan). Gap

Peneliti & Judul	Metode	Hasil Penelitian	Gap Penelitian
Sentiment Analysis of E-Commerce Product Reviews on Tokopedia Using Support Vector Machine			penelitian ini adalah penggunaan dataset yang jauh lebih besar (3.987 ulasan) dan pengujian terhadap empat variasi rasio pembagian data (60:40 hingga 90:10) untuk menguji konsistensi akurasi model pada skala data yang lebih luas.
(Fahmi Fiddin, Taufik Hidayat, 2025) Analisis Sentimen Ulasan Produk Sayur di Tokopedia Menggunakan Model SVM Dengan Representasi TF-IDF	SVM & TF-IDF	Mengklasifikasikan 1.048 ulasan produk sayuran.	pada karakteristik objek ulasan (<i>skincare</i>) yang memiliki terminologi unik (seperti "ph tinggi", "lengket", dsb) serta penerapan 4 fungsi kernel (Linear, RBF, Poly, Sigmoid) untuk menemukan batas keputusan terbaik pada data tersebut.
(Puspa et al., 2023) Analisis sentimen ulasan pada e-commerce shopee menggunakan algoritma naive bayes dan SVM	Naive Bayes & SVM	Membandingkan performa dua algoritma pada platform Shopee.	algoritma (SVM) dengan mengevaluasi pengaruh parameter kernel terhadap data ulasan yang memiliki distribusi kelas tidak seimbang (2.551 negatif vs 1.446 positif).
(Arbiansyah & Haq, 2025) Customer Comment	K-Means Clustering	Berhasil mengelompokkan ulasan berdasarkan kemiripan konten secara otomatis.	Penelitian tersebut menggunakan metode <i>unsupervised learning</i> (K-Means) untuk pengelompokan tanpa label. Gap penelitian ini

Peneliti & Judul	Metode	Hasil Penelitian	Gap Penelitian
Clustering for Kahf <i>Face Wash</i> at Kahf Official Shop Using K-Means Method			adalah pada penggunaan Supervised Learning (SVM) dengan pendekatan Lexicon-based untuk pelabelan otomatis, sehingga memberikan klasifikasi sentimen yang lebih definitif (positif/negatif) dibanding sekadar klusterisasi.

Berdasarkan tinjauan terhadap berbagai penelitian terdahulu pada Tabel 2.1, algoritma *Support Vector Machine* (SVM) secara konsisten menunjukkan performa yang unggul dalam tugas klasifikasi teks dibandingkan dengan algoritma konvensional lainnya. Penggunaan SVM yang dikombinasikan dengan pembobotan fitur *Term Frequency-Inverse Document Frequency* (TF-IDF) terbukti mampu menangani karakteristik data ulasan e-commerce yang memiliki dimensi fitur tinggi dan bersifat sparse.

Meskipun penelitian analisis sentimen pada produk *Face Wash* Sombong telah mulai dilakukan, studi sebelumnya masih berfokus pada penggunaan algoritma *Naïve Bayes Classification* (NBC) dengan sumber data yang berasal dari komentar di platform media sosial TikTok. Selain itu, terdapat penelitian pada produk kompetitor sejenis yang hanya berfokus pada pengelompokan pola komentar pelanggan menggunakan metode clustering tanpa memberikan klasifikasi sentimen yang spesifik.

Terdapat celah penelitian (*research gap*) yang signifikan dalam hal penggunaan algoritma SVM pada data transaksi aktual di platform Tokopedia

khusus untuk produk *Face Wash* Sombong. Berbeda dengan data komentar di media sosial, data ulasan di Tokopedia merepresentasikan opini riil dari konsumen yang telah melakukan pembelian secara resmi. Oleh karena itu, penelitian ini mengisi kekosongan tersebut dengan mengintegrasikan metode TF-IDF dan algoritma SVM guna memberikan evaluasi performa yang lebih akurat, objektif, dan terukur bagi produsen maupun konsumen dalam memantau tren kepuasan pelanggan secara sistematis.

2.2 Analisis Sentimen

Analisis sentimen, atau sering disebut sebagai opinion mining, merupakan cabang dari *text mining* yang menggunakan teknik pemrosesan bahasa alami (*Natural Language Processing*) untuk mengidentifikasi dan mengekstraksi opini subjektif dalam teks dokumen. Dalam ekosistem digital yang berkembang pesat, analisis ini memungkinkan transformasi ulasan pengguna yang bersifat tidak terstruktur menjadi informasi terstruktur guna mengetahui kualitas produk secara detail. Secara umum, sentimen diklasifikasikan ke dalam kategori polaritas seperti positif, negatif, atau netral untuk memberikan gambaran menyeluruh mengenai persepsi public (Willy Wildan Kamal, n.d.).

Dalam konteks industri kecantikan atau *skincare* di Indonesia, peranan ulasan konsumen telah bertransformasi menjadi *social commerce* yang memengaruhi keputusan pembelian secara signifikan (Muzaki *et al.*, 2024). Ulasan pada platform *e-commerce* seperti Tokopedia menjadi sumber data berharga bagi perusahaan untuk mengevaluasi kinerja produk, seperti pada produk *moisturizer* atau *Face Wash*, serta memahami tingkat kepuasan konsumen secara objektif.

Melalui analisis ini, produsen dapat menangkap *electronic word of mouth* (e-WOM) yang berkembang di masyarakat sebagai dasar pengambilan keputusan strategis (Widhiyanta, Nurwahyudi, Isnaini Muhandhis, Roszi Syadillal Jannah, 2025).

Tantangan utama dalam melakukan analisis sentimen pada data ulasan *e-commerce* adalah adanya tumpukan data yang mengandung banyak *noise*, seperti penggunaan bahasa tidak baku, singkatan, kata-kata slang, hingga emotikon. Karakteristik ulasan produk *skincare* di Indonesia sering kali mencampurkan opini mengenai tekstur, harga, dan efektivitas dalam satu kalimat yang singkat. Oleh karena itu, diperlukan tahapan preprocessing yang kuat agar data teks tersebut dapat diolah secara optimal oleh model klasifikasi untuk menghasilkan wawasan yang akurat (Azzahra *et al.*, 2023).

Metodologi analisis sentimen ulasan produk dapat dilakukan melalui beberapa tahapan utama, mulai dari pengumpulan data (*crawling/scraping*), pelabelan, hingga ekstraksi fitur menggunakan metode seperti TF-IDF. Penggunaan algoritma pembelajaran mesin seperti *Support Vector Machine* (SVM) dinilai sangat efektif untuk menangani permasalahan klasifikasi pada data ulasan yang memiliki dimensi fitur tinggi. Dengan sistem klasifikasi otomatis ini, ribuan ulasan dapat dipetakan secara cepat guna membantu perusahaan maupun calon pembeli dalam membedakan antara produk yang memberikan hasil nyata (*experience product*) dengan produk yang hanya sekadar dicari fiturnya saja (*search product*).

2.3 TermFrequency –Inverse Document Frequency (TF-IDF)

Untuk memproses data tekstual menggunakan algoritma *machine learning*, data tersebut harus ditransformasikan ke dalam representasi numerik. Salah satu metode ekstraksi fitur yang sering diterapkan adalah *Term Frequency-Inverse Document Frequency* (TF-IDF), yang mengonversi teks menjadi vektor bobot berdasarkan tingkat signifikansi kata di dalam dokumen. Perhitungan bobot TF-IDF dilakukan melalui persamaan berikut:

$$TF_{t,d} = \frac{f_{t,d}}{\sum_k f_{k,d}} \quad (2.1)$$

$$IDF_t = \log \left(\frac{N}{df_t} \right) \quad (2.2)$$

$$TFIDF_{t,d} = TF_{t,d} \times IDF_t \quad (2.3)$$

Keterangan :

$f_{t,d}$ = Frekuensi kemunculan kata (term) t secara spesifik dalam dokumen d .

$\sum_k f_{k,d}$ = Akumulasi atau jumlah keseluruhan kata yang terdapat di dokumen d

N = jumlah total dokumen yang tersedia di dalam korpus data.

df_t = Kuantitas dokumen dalam korpus yang memuat kata t

$TF_{t,d}$ = Nilai frekuensi relatif dari suatu kata di dalam sebuah dokumen.

IDF_t = Ukuran signifikansi sebuah kata di seluruh koleksi dokumen.

$TFIDF_{t,d}$ = Nilai bobot akhir yang diperoleh dari hasil TF dan IDF

Secara sistematis, persamaan (2.1) digunakan untuk mengukur frekuensi kemunculan suatu kata di dalam dokumen tertentu, sementara persamaan (2.2) berfungsi untuk menilai tingkat kelangkaan kata tersebut di seluruh korpus data. Selanjutnya, persamaan (2.3) mengombinasikan kedua nilai tersebut melalui operasi perkalian guna menetapkan bobot signifikansi setiap kata dalam representasi vektor.

2.4 Algoritma *Support Vector Machine* (SVM)

Support Vector Machine (SVM) merupakan algoritma pembelajaran mesin (*machine learning*) yang termasuk dalam kategori *supervised learning* dan banyak digunakan untuk tugas klasifikasi serta pengenalan pola (Bansal *et al.*, 2022). Algoritma ini bekerja dengan cara meninjau data input dan mengidentifikasi pola yang paling efektif untuk memisahkan kelas-kelas sentimen yang berbeda. Dalam analisis sentimen ulasan produk, SVM sering dipilih karena kemampuannya dalam memproses data teks secara efisien dan memberikan hasil akurasi yang optimal (Handayani, 2021).

Konsep utama dari SVM adalah mencari *hyperplane* optimal yang berfungsi sebagai pembatas untuk memisahkan antar kelas sentimen (seperti positif dan negatif) dengan cara memaksimalkan jarak (*margin*) antar kelas tersebut. *Hyperplane* dalam ruang dua dimensi (2-D) disebut sebagai garis (*line*), dalam tiga dimensi (3-D) disebut sebagai bidang (*plane similarity*), sedangkan pada dimensi yang lebih tinggi disebut sebagai *hyperplane*. Secara matematis, *hyperplane* yang memisahkan data dapat dirumuskan seperti berikut:

$$f(x) = w \cdot x + b = 0 \quad (2.4)$$

Di mana:

w adalah vektor bobot (*weight vector*).

x adalah vektor input fitur ulasan.

b adalah bias (*intercept*).

Titik-titik data yang terletak paling dekat dengan *hyperplane* pembatas ini disebut sebagai *support vectors*. Keandalan SVM terletak pada pencarian margin

terbesar antara *support vectors* dari kelas yang berbeda guna meminimalkan kesalahan klasifikasi pada data uji.

Apabila dataset ulasan tidak dapat dipisahkan secara linear (*non-linearly separable*), SVM memanfaatkan fungsi kernel untuk memetakan data tersebut ke ruang dimensi yang lebih tinggi. Melalui transformasi ini, data yang semula rumit dapat dipisahkan secara linear menggunakan *hyperplane* optimal pada dimensi baru tersebut. (Pratiwi et al., 2021) . Beberapa jenis rumus fungsi kernel yang sering digunakan dalam klasifikasi SVM antara lain:

1. *Polynomial Kernel* Digunakan untuk mengklasifikasikan dataset *training* yang sudah melalui proses normalisasi:

$$K(x, y) = (x \cdot y + r)^d \quad (2.5)$$

2. *Radial Basis Function* (RBF) Digunakan untuk mengklasifikasikan dataset yang tidak terpisah secara linear dengan tingkat akurasi prediksi yang sangat baik:

$$K(x, y) = \exp(-\gamma \|x - y\|^2) \quad (2.6)$$

3. *Sigmoid Kernel* Merupakan pengembangan fungsi yang diadopsi dari konsep jaringan saraf tiruan:

$$K(x, y) = \tanh(ax^T y + c) \quad (2.7)$$

4. *Linear kernel* merupakan fungsi kernel yang paling sederhana dan umum digunakan pada tugas klasifikasi teks. Fungsi ini bekerja dengan cara melakukan perkalian titik (*dot product*) secara langsung antara dua vektor input tanpa memetakan data ke ruang dimensi yang lebih tinggi secara kompleks :

$$K(x, x_i) = \sum (x \cdot x_i) \quad (2.8)$$

Keterangan:

K = Fungsi Kernel.

x, y = Vektor ruang input (*input space*).

d = Derajat polinomial (*degree*).

γ (*Gamma*)= Parameter yang menentukan seberapa jauh pengaruh satu titik data pelatihan menjangkau.

α (*Alpha*) = Parameter kemiringan (*slope*) yang mengontrol bentuk fungsi aktivasi.

c (*Constant*)= Konstanta penambah (sering disebut sebagai parameter intercept dalam konteks kernel sigmoid).

Pemilihan parameter dan fitur yang tepat sangat krusial dalam performa SVM, karena algoritma ini memiliki sensitivitas terhadap pemilihan parameter kernel guna mencapai nilai akurasi dan *Area Under Curve* (AUC) yang masuk dalam kategori *excellent classification*.

Dalam praktiknya, data ulasan produk seringkali tidak dapat dipisahkan secara linear di ruang dimensi aslinya karena kompleksitas bahasa yang digunakan pengguna. Untuk mengatasi masalah ini, SVM menggunakan fungsi kernel untuk memetakan data dari ruang input yang berdimensi rendah ke ruang fitur yang berdimensi lebih tinggi (*high-dimensional feature space*). Beberapa jenis kernel yang umum digunakan dalam klasifikasi teks antara lain adalah *kernel Linear*, *Polynomial*, dan *Radial Basis Function* (RBF), di mana *kernel Linear* sering kali memberikan performa terbaik pada klasifikasi sentimen berbasis teks.

Penggunaan SVM dalam analisis sentimen ulasan produk *skincare* di *e-commerce* memiliki keunggulan yang signifikan, terutama dalam menangani data yang bersifat sparse atau jarang. Hal ini dikarenakan SVM mampu bekerja secara stabil meskipun jumlah fiturnya sangat banyak, yang merupakan karakteristik umum dari hasil ekstraksi fitur menggunakan TF-IDF. Dengan memaksimalkan margin antar kelas sentimen, SVM mampu mengurangi risiko terjadinya *overfitting*

dan memberikan hasil prediksi yang lebih akurat terhadap ulasan-ulasan baru yang belum pernah dilihat oleh model sebelumnya.

2.5 Evaluasi Hasil

Tahap evaluasi hasil dilakukan untuk mengukur performa model klasifikasi sentimen secara komprehensif menggunakan metrik *accuracy*, *precision*, *recall*, dan *F1-score*. Metrik *accuracy* sendiri merepresentasikan proporsi prediksi benar (baik positif maupun negatif) dari keseluruhan data uji yang dievaluasi. Perhitungan nilai *accuracy* dilakukan melalui persamaan berikut:

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (2.9)$$

Keterangan :

TP(*True Positive*) = Jumlah data positif yang diprediksi secara benar sebagai kelas positif oleh model.

TN(*True Negative*) = Jumlah data negatif yang diprediksi secara benar sebagai kelas negatif oleh model.

FP(*False Positive*) = Jumlah data negatif yang salah diprediksi sebagai kelas positif oleh model (kesalahan Tipe I).

FN(*False Negative*) = Jumlah data positif yang salah diprediksi sebagai kelas negatif oleh model (kesalahan Tipe II).

Precision: Berfungsi untuk mengukur tingkat ketepatan atau rigiditas model dalam memprediksi kelas positif dari keseluruhan data yang dideklarasikan sebagai positif.

$$Precision = \frac{TP}{TP + FP} \quad (2.10)$$

Recall Berfungsi untuk mengukur sensitivitas model dalam mengidentifikasi kembali seluruh data yang secara aktual termasuk dalam kategori positif.

$$Recall = \frac{TP}{TP + FN} \quad (2.11)$$

F1-score Merupakan nilai harmonik rata-rata yang merepresentasikan keseimbangan performa antara *precision* dan *recall*.

$$F1 = \frac{2 \times Precision \times Recall}{Precision + Recall} \quad (2.12)$$

Guna mengantisipasi kondisi dataset yang mengalami ketidakseimbangan kelas (*class imbalance*), evaluasi diperluas menggunakan dua pendekatan, yakni *macro F1-score* yang menghitung rata-rata performa antar-kelas tanpa melibatkan bobot jumlah sampel, serta *weighted F1-score* yang mengalkulasi rata-rata performa dengan memberikan pembobotan sesuai proporsi jumlah data di setiap kelas.

BAB III

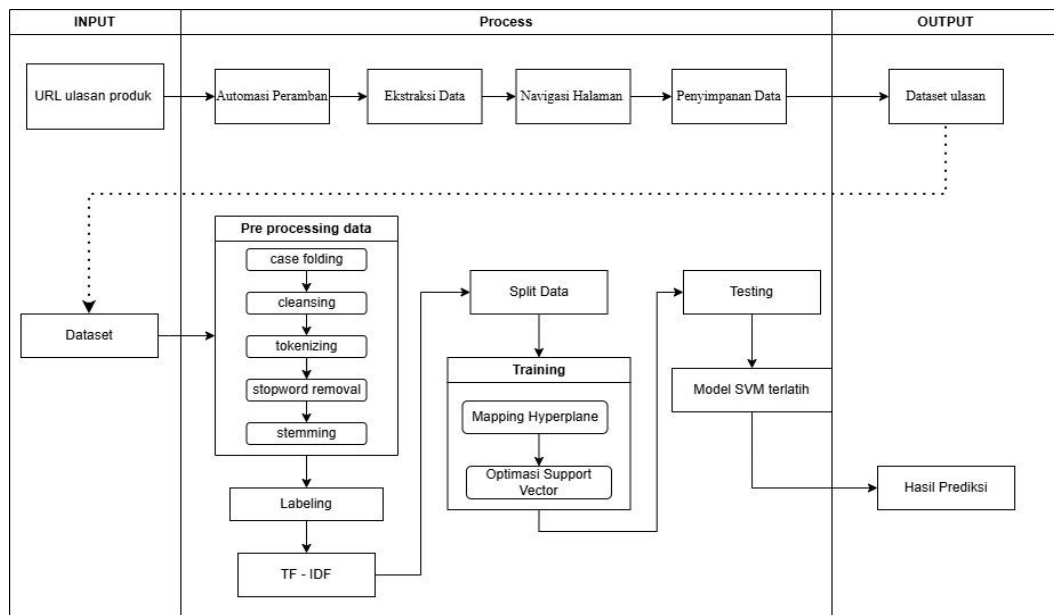
DESAIN DAN IMPLEMENTASI

3.1 Desain Sistem

Desain sistem dalam penelitian ini menggambarkan arsitektur logika dan alur kerja fungsional yang mentransformasikan data ulasan mentah menjadi hasil klasifikasi sentimen yang terukur. Sistem dirancang secara modular dengan menempatkan tahap pelabelan data (*data labeling*) sebagai langkah awal setelah akuisisi data untuk menentukan *ground truth* atau kelas target, yaitu sentimen positif dan negatif. Proses ini sangat krusial karena algoritma *Support Vector Machine* (SVM) memerlukan supervisi label yang jelas sebelum memasuki tahap pemrosesan teknis. Setelah dataset memiliki label yang valid, sistem menjalankan modul prapemrosesan teks yang meliputi *case folding*, *filtering*, *tokenizing*, *stopword removal*, dan *stemming*. Tahapan ini bertujuan untuk mereduksi derau (*noise*) pada teks ulasan sehingga dihasilkan korpus kata dasar yang bersih dan relevan.

Tahap berikutnya adalah pembobotan fitur menggunakan metode *Term Frequency-Inverse Document Frequency* (TF-IDF) untuk mengubah data tekstual menjadi representasi numerik berupa matriks vektor. Matriks fitur tersebut kemudian dibagi menjadi data latih dan data uji dengan rasio yang bervariasi sesuai skenario pengujian. Pada fase pelatihan (*training phase*), sistem membangun model klasifikasi dengan mencari hyperplane optimal melalui komparasi empat fungsi kernel, yaitu *Linear*, *RBF*, *Polynomial*, dan *Sigmoid*. Model yang telah terbentuk

kemudian divalidasi pada fase pengujian (*testing phase*) untuk memprediksi label pada data uji. Akhir dari desain sistem ini adalah tahap evaluasi menggunakan *confusion matrix* untuk mengukur performa setiap kernel berdasarkan metrik akurasi, presisi, *recall*, dan *f1-score*.



Gambar 3. 1 Desain Sistem

3.2 Akuisisi Data

Data yang digunakan dalam penelitian ini merupakan data primer berupa ulasan konsumen yang diambil dari platform *e-commerce* Tokopedia. Objek penelitian difokuskan pada produk *Face Wash SOMBONG* dengan total data yang berhasil dikumpulkan sebanyak 3.987 ulasan. Seluruh data tersebut diakses dan dikumpulkan pada tanggal 28 Oktober 2025 menggunakan teknik *web scraping*.

Proses akuisisi data dilakukan secara otomatis menggunakan bahasa pemrograman *Python* dengan memanfaatkan pustaka (*library*) *Selenium* dan *BeautifulSoup*. *Selenium* digunakan untuk mensimulasikan aktivitas peramban (*browser*) agar dapat menangani konten dinamis dan melakukan navigasi halaman

(*pagination*) secara otomatis. Tahapan *web scraping* yang dirancang dalam penelitian ini adalah sebagai berikut:

1. Inisialisasi URL: Program mengakses URL dasar produk dan menerapkan filter rating (bintang 1 hingga 5) secara berurutan. Hal ini dilakukan untuk memastikan seluruh variasi opini pelanggan terjaring tanpa ada yang terlewat.
2. Automasi Peramban: Menggunakan *WebDriver* untuk membuka halaman Tokopedia, menutup popup iklan secara otomatis, dan melakukan *scrolling* ke area ulasan agar elemen teks termuat secara sempurna.
3. Ekstraksi Data: Program mengidentifikasi elemen spesifik pada struktur HTML menggunakan *CSS Selector* dan *XPath*. Data yang diekstrak yaitu teks ulasan (*lblItemUlasan*).
4. Navigasi Halaman (*Pagination*): Program menjalankan perulangan (*looping*) untuk menekan tombol "Laman Berikutnya" hingga mencapai batas akhir ulasan yang tersedia pada setiap filter rating.
5. Penyimpanan Data: Data yang telah berhasil diekstraksi kemudian dikumpulkan dalam objek *DataFrame* melalui pustaka *Pandas* dan disimpan ke dalam format file *.csv* untuk diproses pada tahap selanjutnya.

Contoh dari data mentah yang telah dikumpulkan dapat dilihat pada Tabel

3.1 berikut:

Tabel 3. 1 Contoh Dataset Ulasan Tokopedia

No	Teks Ulasan
1	Barang bagus,mulus dan original.pengiriman cepat kurir ramah,toko amanah.terimakasih Jenis kulit: Pria Aroma: wangi Performa: strong
2	Barang mendarat dgnn aman wangi kayak rasa2 buah dan isinya full real. Bru pakai sekali moga2 cocok.

No	Teks Ulasan
3	asli sih ini rekomend banget pacial wash + <i>skincare</i> ini mah,pas pertama pake wajah jadi lebih bersih, minyak berkurang di wajah,wajah tidak kaku,rasanya lembap gitu,dan ba...
4	Barang sesuai pesanan, pengiriman cepat, seller responsif
5	alhamdulillah luar bisa sesuai deskripsi yang tertera, permintaan, saya puas sekali membelinya, two thump up dgn usahanya, dan terima kasih sukses maju terus...
6	Packaging nya benar2 minimalis, dan juga tutupnya mudah kebuka, klo semisal ketekan barang lain udah keluar semua isinya

Berdasarkan data yang disajikan pada tabel di atas, terlihat bahwa ulasan yang diberikan oleh konsumen memiliki karakteristik yang sangat beragam dan bersifat tidak terstruktur. Sebagian besar ulasan mengandung penggunaan bahasa sehari-hari, singkatan, serta istilah khusus dalam dunia perawatan kulit (*skincare*) seperti "rekomend", "*facial wash*", dan "lembap". Selain itu, terdapat banyak penggunaan karakter non-teks seperti emotikon dan tanda baca yang berlebihan, yang dalam istilah text mining disebut sebagai *noise*. Kondisi data mentah ini menegaskan perlunya tahapan prapemrosesan yang mendalam agar algoritma *Support Vector Machine* (SVM) dapat mengenali pola fitur sentimen tanpa terganggu oleh karakter yang tidak relevan.

Jumlah data sebanyak 3.987 ulasan memberikan cakupan informasi yang cukup luas untuk merepresentasikan berbagai jenis opini, mulai dari kepuasan terhadap tekstur dan aroma hingga keluhan mengenai efektivitas produk. Skala dataset ini menjadi fondasi penting bagi proses pembelajaran model (*training*), di mana semakin beragam pola bahasa yang ditangkap dalam fase pengumpulan data, maka semakin baik pula kemampuan model SVM dalam melakukan generalisasi terhadap ulasan-ulasan baru di masa mendatang. Dengan demikian, data primer

hasil *scraping* ini tidak hanya berfungsi sebagai sampel statistik, tetapi juga sebagai sumber pengetahuan utama bagi sistem untuk membedakan polaritas sentimen secara objektif dan akurat.

3.3 Pelabelan Data (*Data Labeling*)

Tahap pelabelan merupakan proses penentuan kelas atau target sentimen pada setiap ulasan agar algoritma *Support Vector Machine* (SVM) dapat melakukan proses pembelajaran. Karena dataset hasil *scraping* dari Tokopedia tidak menyertakan label sentimen secara eksplisit, maka dilakukan pelabelan menggunakan pendekatan *Lexicon-based* (berbasis kamus kata).

Pendekatan *Lexicon-based* bekerja dengan memanfaatkan daftar kata (*lexicon*) yang telah memiliki nilai bobot sentimen tetap. Dalam proses ini, sistem akan membedah setiap ulasan yang telah melalui prapemrosesan dan mencocokkan setiap kata di dalamnya dengan kamus sentimen positif dan negatif. Jika sebuah kata ditemukan dalam kamus positif, maka nilai skor ulasan tersebut akan bertambah (+1), sebaliknya jika ditemukan dalam kamus negatif, skor akan berkurang (-1). Akumulasi akhir dari nilai-nilai tersebut akan menentukan polaritas ulasan: ulasan dengan total skor lebih besar dari nol dikategorikan sebagai Positif, ulasan dengan skor kurang dari nol sebagai Negatif.

Berdasarkan tahapan pengujian yang telah dilakukan sebelumnya, rincian data serta akumulasi hasil dari proses pelabelan otomatis menggunakan metode ini kini disajikan secara visual pada Tabel 3.2 berikut:

Tabel 3. 2 Contoh Hasil Pelabelan Data Sentimen

No	Data Mentah (Ulasan Pelanggan)	Kata Kunci (Lexicon)	Label
1	asli sih ini rekomend banget pacial wash +skincare ini mah,pas pertama pake wajah jadi lebih bersih, minyak berkurang di wajah,wajah tidak kaku,rasanya lembap gitu,dan ba...	rekomend, bersih, lembap	Positif
2	alhamdulillah luar bisa sesuai deskripsi yang tertera, permintaan, saya puas sekali membelinya, two thump up dgn usahanya, dan terima kasih sukses maju terus...!	luar biasa, puas, terima kasih	Positif
3	JUJUR untuk kulit saya gak cocok untuk yang Tipe Berminyak , muka agak panas di area pipi tengah dan ternyata agak merah setelah pakai.. kayaknya karena Ph nya nya tinggi...	gak cocok, panas, merah	Negatif
4	Mantap harga murah tapi bukan barang murahan Terima kasih buat produk bang Densu facial sombong	mantap, murah, terima kasih	Positif

Data yang tersaji dalam Tabel 3.2 di atas merupakan hasil dari proses pelabelan menggunakan pendekatan *Lexicon-Based* yang bekerja dengan menghitung bobot polaritas kata pada setiap ulasan mentah. Skor sentimen ditentukan dari hasil penjumlahan nilai kata-kata kunci yang ditemukan dalam teks; sebagai contoh pada data nomor 1, penggunaan kata "rekomend", "bersih", dan "lembap" masing-masing memberikan poin positif sehingga menghasilkan akumulasi skor yang menetapkan ulasan tersebut ke dalam label positif. Sebaliknya, pada data nomor 3, kemunculan kata kunci seperti "gak cocok", "panas", dan "merah" menyumbang poin negatif yang menyebabkan total skor bernilai di bawah nol, sehingga ulasan tersebut diklasifikasikan ke dalam label negatif.

Berdasarkan hasil komputasi menggunakan pendekatan leksikon terhadap total 3.987 ulasan produk Face Wash SOMBONG, diperoleh distribusi data yang menunjukkan persepsi konsumen terhadap produk tersebut. Hasil pelabelan

otomatis ini menghasilkan 1.814 data sentimen positif dan 2.173 data sentimen negatif. Distribusi awal ini menunjukkan adanya ketidakseimbangan kelas (*class imbalance*), di mana jumlah ulasan bersentimen negatif bertindak sebagai kelas mayoritas.

Untuk mengatasi kendala ketidakseimbangan tersebut, penelitian ini menerapkan teknik *Balancing Data* menggunakan metode Random Over-Sampling (ROS). Teknik ini bekerja dengan cara mereplikasi atau menduplikasi sampel pada kelas minoritas (positif) secara acak hingga mencapai proporsi yang setara dengan kelas mayoritas (negatif). Melalui penerapan ROS, diperoleh keseimbangan data yang mutakhir dengan komposisi 2.173 data sentimen positif dan 2.173 data sentimen negative, sehingga menghasilkan akumulasi total sebanyak 4.346 ulasan.. Langkah ini sangat krusial agar model SVM yang dibangun nantinya tidak memiliki kecenderungan (*bias*) dalam memprediksi ulasan hanya ke satu pihak, sehingga meningkatkan reliabilitas nilai *Recall* dan *F1-Score*.

Pemilihan metode *Random Over-Sampling* (ROS) dibandingkan metode penyeimbangan lain seperti *Random Over Sampling* (RUS) maupun *Synthetic Minority Over-sampling Technique* (SMOTE) didasari oleh pertimbangan ilmiah yang kuat terkait karakteristik dataset teks. Penggunaan RUS sengaja dihindari karena metode tersebut mereduksi data kelas mayoritas secara acak, yang berisiko membuang informasi-informasi penting (*loss of information*) dan menghilangkan variasi kata unik di dalam ulasan negatif. Sementara itu, metode SMOTE yang bekerja dengan membuat sampel sintetis baru berbasis tetangga terdekat (*k-nearest neighbors*) kurang ideal diterapkan pada representasi data teks berbasis ruang

vektor berdimensi tinggi (*high-dimensional sparse vector space*) seperti TF-IDF. SMOTE menghasilkan sampel sintetik berupa nilai kontinu di antara dua titik vektor kata, yang secara semantik sering kali memunculkan kombinasi fitur kata artifisial yang tidak masuk akal secara linguistik dan dapat mengaburkan batas keputusan (*decision boundary*) model SVM. Sebaliknya, ROS mempertahankan orisinalitas konten ulasan teks dengan mereplikasi data positif yang sudah ada secara utuh tanpa mengubah makna semantik kalimat, sehingga sangat efektif membantu algoritma SVM dalam memperkuat pengenalan pola kelas minoritas pada ruang dimensi tinggi tanpa kehilangan dokumen ulasan yang berharga.

Mekanisme pelabelan dan penyeimbangan ini menjadi jembatan krusial dalam mentransformasikan data teks murni menjadi data berlabel (*labeled data*) yang siap digunakan untuk melatih model klasifikasi. Dengan menetapkan label hasil komputasi leksikon yang telah diseimbangkan sebagai target (*ground truth*), algoritma *Support Vector Machine* (SVM) dapat mempelajari karakteristik linguistik dan pola kata dari ulasan positif maupun negatif secara terstruktur dan adil. Validitas label yang dihasilkan melalui pendekatan ini sangat menentukan kualitas pembelajaran model sebelum sistem memasuki tahap prapemrosesan dan ekstraksi fitur TF-IDF, guna memastikan bahwa model akhir mampu melakukan klasifikasi dengan tingkat akurasi yang tinggi.

3.4 Prapemrosesan Teks (*Text Processing*)

Tahapan ini bertujuan untuk mengubah ulasan yang bersifat tidak terstruktur menjadi format yang seragam dan bersih dari derau (*noise*), sehingga memudahkan

algoritma *Support Vector Machine* (SVM) dalam mengenali pola sentimen. Berdasarkan dataset yang telah diproses, perbandingan antara ulasan asli dengan hasil prapemrosesan teks disajikan pada Tabel 3.3 berikut:

Tabel 3. 3 Contoh Perbandingan Hasil Prapemrosesan Teks

No	Ulasan Asli	Case Folding & Cleansing	Tokenizing	Stopword Removal	Stemming (Hasil Akhir)
1	Barangnya original, pengiriman sangat cepat sekali!	barangnya original pengiriman sangat cepat sekali	['barangnya', 'original', 'pengiriman', 'sangat', 'cepat', 'sekali']	['barangnya', 'original', 'pengiriman', 'cepat']	barang original kirim cepat
2	Produknya kurang cocok di muka saya, malah jadi berminyak.	produknya kurang cocok di muka saya malah jadi berminyak	['produknya', 'kurang', 'cocok', 'di', 'muka', 'saya', 'malah', 'jadi', 'berminyak']	['produknya', 'cocok', 'muka', 'berminyak']	produk cocok muka minyak
3	Packing rapi dan pengirimannya tidak mengecewakan.	packing rapi dan pengirimannya tidak mengecewakan	['packing', 'rapi', 'dan', 'pengirimannya', 'tidak', 'mengecewakan']	['packing', 'rapi', 'pengirimannya', 'mengecewakan']	packing rapi kirim kecewa
4	Moga-moga cocok ya, baru pertama kali coba performa ini.	moga moga cocok ya baru pertama kali coba performa ini	['moga', 'moga', 'cocok', 'ya', 'baru', 'pertama', 'kali', 'coba', 'performa', 'ini']	['moga', 'moga', 'cocok', 'coba', 'performa']	moga moga cocok coba performa

Berdasarkan tabel di atas, tahapan prapemrosesan teks dalam penelitian ini dilakukan melalui beberapa prosedur teknis, yaitu:

1. *Case Folding*: Mengubah seluruh huruf kapital menjadi huruf kecil untuk menghindari ambiguitas karakter bagi sistem.

2. *Cleansing*: Menghapus angka (seperti "100%"), tanda baca, simbol matematika ("+"), serta emotikon yang tidak dapat diproses oleh algoritma.
3. *Tokenizing*: Memecah kalimat menjadi satuan kata tunggal atau token untuk memudahkan tahap pembobotan.
4. *Stopword Removal*: Menghilangkan kata-kata umum yang sering muncul namun tidak memiliki muatan sentimen, seperti kata hubung "dan", "yang", "di", dan "ini".
5. *Stemming*: Mengubah kata berimbuhan menjadi kata dasar. Contohnya, kata "mendarat" ditransformasikan menjadi kata dasar "darat" dan "pengiriman" menjadi " kirim".

Penerapan rangkaian prapemrosesan ini secara signifikan mereduksi kompleksitas teks ulasan tanpa menghilangkan esensi informasinya. Dengan mengubah kata-kata menjadi bentuk dasar dan menghilangkan karakter yang tidak perlu, fitur-fitur yang dihasilkan pada tahap ekstraksi nantinya akan menjadi lebih konsisten. Hal ini krusial untuk meningkatkan efisiensi komputasi serta akurasi klasifikasi SVM dalam memetakan ulasan pengguna *Face Wash SOMBONG* ke dalam kategori sentimen yang tepat.

3.5 Pembobotan Fitur TF-IDF (*TermFrequency-Inverse Document Frequency*)

Setelah data ulasan memiliki label sentimen, tahap selanjutnya adalah merepresentasikan data teks tersebut ke dalam bentuk numerik menggunakan metode *TermFrequency-Inverse Document Frequency* (TF-IDF). Algoritma

Support Vector Machine (SVM) tidak dapat mengolah data berupa teks secara langsung, sehingga diperlukan proses ekstraksi fitur untuk mengubah setiap kata dalam ulasan menjadi vektor numerik yang mencerminkan bobot kepentingan kata tersebut terhadap dokumen.

Metode TF-IDF bekerja dengan cara menggabungkan dua konsep perhitungan, yaitu *TermFrequency* (TF) dan *Inverse Document Frequency* (IDF). *TermFrequency* mengukur frekuensi kemunculan suatu kata dalam sebuah ulasan, sedangkan *Inverse Document Frequency* mengukur seberapa jarang atau umum suatu kata di seluruh dataset ulasan *Face Wash SOMBONG*.

Untuk memberikan gambaran matematis yang lebih komprehensif, berikut adalah simulasi perhitungan manual menggunakan enam sampel ulasan ($N=6$) yang diambil dari dataset *Face Wash SOMBONG* setelah tahap prapemrosesan, sebagaimana disajikan pada Tabel 3.5:

Tabel 3. 4 Sampel Data untuk Simulasi Perhitungan TF-IDF

Label	Teks Ulasan (<i>Clean Review</i>)	Jumlah Kata
D1	barang bagusmulus originalpengiriman cepat kurir ramahtoko amanahterimakasih jenis kulit pria aroma wangi performa strong	14
D2	barang darat dngn aman wangi kayak rasa buah isi full real bru pakai sekali moga cocok	16
D3	asli sih rekomend banget pacial wash <i>skincare</i> mahpas pertama pake wajah jadi lebih bersih minyak kurang wajahwajah kakurasanya lembap gitu dan ba	21
D4	barang sesuai pesan kirim cepat seller responsif	7
D5	produk original moga cocok sama muka	6
D6	aroma wangi performa bagus kualitas produk bagus harga jangkau	9

Sebagai contoh perhitungan, kita akan menghitung bobot kata "barang" pada ulasan pertama (D1):

1. Menghitung *TermFrequency* (TF) : Kata "barang" muncul 1 kali dalam D1 yang berjumlah 14 kata.

$$tf('barang', d_1) = \frac{1}{14} = 0,0714 \quad (3.1)$$

2. Menghitung *Document Frequency* (DF): Kata "barang" muncul di 3 ulasan (D1, D2, dan D4), maka DF = 3.
3. Menghitung *Inverse Document Frequency* (IDF): Menggunakan rumus logaritma dengan total ulasan ($N = 6$):

$$idf('barang') = \log_{10}\left(\frac{6}{3}\right) = \log_{10}(2) = 0,3010 \quad (3.2)$$

4. Menghitung Bobot TF-IDF:

$$\omega('barang', d_1) = tf \times idf = 0,0714 \times 0,3010 = 0,0215 \quad (3.3)$$

Untuk memberikan gambaran yang lebih objektif terhadap seluruh dataset, Tabel 3.5 menyajikan hasil pembobotan TF-IDF untuk 7 kata kunci yang muncul pada berbagai sampel ulasan (D1 hingga D6). Tabel ini merangkum di mana saja kata tersebut muncul serta nilai kepentingannya berdasarkan ulasan terkait:

Tabel 3. 5 Contoh Ringkasan Perhitungan TF-IDF Berdasarkan Keseluruhan Sampel ($N=6$)

Kata (Term)	Tempat Muncul	TF (TermFrequency)	DF	IDF	TF-IDF
Barang	D1,D2,D4	0,071 (pada D1)	3	0,301	0,0215
Wangi	D1,D2,D6	0,071 (pada D1)	3	0,301	0,0215
Cepat	D1,D4	0,143 (pada D4)	2	0,477	0,0682
Moga	D2,D5	0,167 (pada D5)	2	0,477	0,0797
Cocok	D2,D5	0,167 (pada D5)	2	0,477	0,0797
Performa	D1,D6	0,111 (pada D6)	2	0,477	0,0529
Produk	D5,D6	0,167 (pada D5)	2	0,477	0,0797

Data yang disajikan pada Tabel 3.5 merepresentasikan distribusi bobot kata yang dihitung berdasarkan frekuensi kemunculan *Term* dalam dokumen tunggal serta signifikansi kata tersebut di dalam seluruh korpus ulasan. Secara statistik, kolom "Tempat Muncul" dan *Document Frequency* (DF) menunjukkan tingkat sebaran kata di seluruh sampel data; kata-kata dengan nilai DF yang rendah, seperti "Produk" atau "Performa", menghasilkan nilai *Inverse Document Frequency* (IDF) yang lebih tinggi. Hal ini membuktikan bahwa metode TF-IDF secara efektif memberikan atensi lebih besar pada kata-kata yang bersifat unik dan memiliki daya pembeda (*discriminative power*) yang kuat dalam menentukan konteks sebuah ulasan, dibandingkan dengan kata-kata umum yang memiliki nilai IDF rendah.

Selanjutnya, integrasi antara nilai *Term Frequency* (TF) dan IDF menghasilkan skor akhir TF-IDF yang berfungsi sebagai representasi fitur numerik dari teks ulasan. Perbedaan nilai TF-IDF pada setiap kata mencerminkan kontribusi relatif masing-masing fitur dalam mendefinisikan polaritas sentimen. Fitur dengan bobot TF-IDF yang signifikan menjadi indikator utama bagi algoritma *Support Vector Machine* (SVM) dalam memetakan posisi ulasan ke dalam ruang vektor dimensi tinggi. Transformasi ini sangat krusial agar data tekstual dapat diproses melalui perhitungan matematis guna menentukan *hyperplane* atau batas pemisah optimal antara kategori sentimen positif dan negatif secara akurat dan objektif.

3.6 *Support Vector Machine* (SVM)

Setelah ulasan direpresentasikan ke dalam bentuk vektor numerik melalui pembobotan TF-IDF, tahap inti selanjutnya adalah klasifikasi menggunakan

algoritma *Support Vector Machine* (SVM). SVM dipilih karena kemampuannya yang optimal dalam menangani data dimensi tinggi dan efektivitasnya dalam memisahkan kelas yang kompleks pada analisis sentimen. Algoritma ini bekerja dengan cara mentransformasikan data ke dalam ruang fitur dan mencari pemisah yang paling efisien di antara kategori-kategori yang ada.

Secara prinsip, SVM berupaya menemukan sebuah *hyperplane* atau bidang pemisah yang memiliki jarak terjauh (*maximum margin*) terhadap titik data terdekat dari masing-masing kelas, yang dikenal sebagai support vectors. Keunggulan utama SVM dalam penelitian analisis sentimen terletak pada kemampuannya untuk meminimalkan kesalahan klasifikasi tanpa mengalami masalah *overfitting*, meskipun dataset memiliki jumlah fitur kata yang sangat banyak setelah proses ekstraksi TF-IDF. Hal ini sangat krusial dalam mengolah ulasan produk *Face Wash SOMBONG* yang memiliki karakteristik bahasa yang beragam dan kompleks.

Proses klasifikasi ini melibatkan serangkaian prosedur sistematis yang dimulai dari penyiapan dataset hingga penerapan fungsi kernel. Dengan menerapkan algoritma ini, sistem diharapkan mampu mengenali pola-pola linguistik yang merepresentasikan kepuasan maupun ketidakpuasan konsumen secara objektif. Melalui perhitungan matematis pada ruang vektor, model yang dihasilkan tidak hanya sekadar mengklasifikasikan data yang sudah ada, tetapi juga memiliki kemampuan prediktif untuk menentukan kelas sentimen pada ulasan-ulasan baru di masa mendatang.

3.6.1 Pembagian Data (Data Splitting)

Tahap pembagian data (*data splitting*) dilakukan untuk memisahkan dataset menjadi dua bagian utama, yaitu data latih (*training data*) dan data uji (*testing data*). Data latih digunakan untuk membentuk model *Support Vector Machine*, sedangkan data uji digunakan untuk memvalidasi performa model dalam mengklasifikasikan ulasan baru.

Penelitian ini menerapkan empat variasi rasio untuk menguji konsistensi akurasi model. Perhitungan sebaran data tersebut adalah sebagai berikut:

1. Skenario I (Rasio 60:40) Skenario ini mengalokasikan 60% porsi dataset untuk pelatihan model dan 40% untuk pengujian performa. Komposisi data pada skenario pertama ini bertransformasi menjadi 2.607 dokumen untuk data latih dan 1.739 dokumen untuk data uji, sehingga menghasilkan akumulasi total sebanyak 4.346 ulasan.
2. Skenario II (Rasio 70:30) Skenario ini membagi porsi dataset dengan komposisi 70% sebagai data latih dan 30% sebagai data uji. Alokasi ini menghasilkan sebanyak 3.042 dokumen untuk data latih dan 1.304 dokumen untuk data uji, sehingga tetap menghasilkan akumulasi total sebanyak 4.346 ulasan.
3. Skenario III (Rasio 80:20) Skenario ini menerapkan porsi standar ideal dalam pemodelan data teks dengan mengalokasikan 80% sebagai data latih dan 20% sebagai data uji. Komposisi ini membagi dataset menjadi sebanyak 3.476 dokumen untuk data latih dan 870 dokumen untuk data uji, dengan akumulasi total sebanyak 4.346 ulasan.

4. Skenario IV (Rasio 90:10) Skenario ini menggunakan porsi maksimal untuk basis pengetahuan model, yaitu sebesar 90% sebagai data latih dan 10% sisa datanya sebagai data uji. Skenario ini mendistribusikan sebanyak 3.911 dokumen untuk data latih dan 435 dokumen untuk data uji, guna melihat ketahanan model terhadap potensi terjadinya *overfitting*.

Pembagian data dilakukan menggunakan teknik random splitting dengan tetap menjaga distribusi label sentimen agar tetap representatif pada setiap bagian. Melalui penerapan skenario yang variatif ini, peneliti dapat mendeteksi potensi terjadinya *overfitting* serta menentukan rasio optimal yang menghasilkan nilai akurasi, presisi, *recall*, dan *f1-score* tertinggi pada tahap evaluasi di bab selanjutnya.

3.6.2 Fase Pelatihan (*Training Phase*)

Fase pelatihan merupakan tahap di mana algoritma SVM membangun model kecerdasan dengan mempelajari karakteristik data yang telah diberi label. Langkah-langkah pada fase ini meliputi:

1. *Input Data*: Mengambil data latih yang telah melalui tahap pembobotan TF-IDF dan memiliki label sentimen.
2. *Pemetaan Fitur*: Setiap ulasan dipetakan sebagai vektor di dalam ruang fitur berdimensi tinggi.
3. *Optimasi Hyperplane*: Algoritma mencari bidang pemisah yang memiliki margin terbesar (jarak terjauh) antara dua kelas. Titik data yang paling dekat dengan *hyperplane* disebut sebagai *support vectors*.

4. *Output*: Hasil dari fase ini adalah model klasifikasi tersimpan yang siap digunakan untuk melakukan prediksi pada data baru.

3.6.3 Fase Pengujian (*Testing Phase*)

Fase pengujian dilakukan untuk memvalidasi performa model menggunakan data yang belum pernah dilihat sebelumnya (*unseen data*). Langkah-langkahnya adalah:

1. *Input Data*: Mengambil data uji berupa vektor TF-IDF tanpa menyertakan label aslinya.
2. *Klasifikasi*: Model akan menentukan posisi vektor data uji terhadap *hyperplane* yang telah terbentuk.
3. *Prediksi*: Sistem memberikan label sentimen (Positif atau Negatif) berdasarkan posisi koordinat vektor tersebut dalam ruang fitur.
4. *Output*: Hasil akhir berupa label prediksi yang akan digunakan pada tahap evaluasi performa.

3.6.4 Fungsi Kernel SVM dan Perhitungan Manual

Fungsi kernel berperan dalam menghitung nilai produk skalar (dot product) dari dua buah vektor data pada ruang fitur berdimensi tinggi. Dalam implementasinya, performa setiap fungsi kernel sangat bergantung pada nilai parameter pendukungnya.

Perlu dicatat bahwa nilai hyperparameter (seperti parameter C , γ , dan *degree*) yang digunakan untuk setiap jenis kernel dalam penelitian ini ditentukan melalui metode *Grid Search*. Metode ini bekerja dengan cara melakukan pencarian

sistematis melalui kombinasi nilai-nilai parameter yang telah ditentukan sebelumnya dalam sebuah rentang tertentu (*grid*). Dengan menggunakan *Grid Search*, model dapat secara otomatis mengidentifikasi kombinasi parameter paling optimal yang menghasilkan akurasi tertinggi untuk setiap skenario pengujian.

Sebagai ilustrasi mekanisme kerja algoritma dalam memisahkan data ulasan produk *Face Wash SOMBONG*, dilakukan simulasi perhitungan manual menggunakan nilai bobot TF-IDF sebagai berikut:

Vektor (x / D5): [0,0797] (diambil dari kata 'produk')

Vektor (y / D6): [0,0529] (diambil dari kata 'performa')

Berikut adalah hasil transformasi menggunakan empat jenis fungsi kernel yang dikomparasikan:

1. Kernel Linear (Rumus 2.8) Kernel ini menghitung perkalian titik secara langsung tanpa transformasi dimensi tambahan.

$$K = 0,0797 \times 0,0529 = 0,00421$$

2. Kernel Polynomial (Rumus 2.5) Kernel ini menghitung interaksi fitur antar dokumen hingga derajat tertentu (d). Diasumsikan derajat $d = 2$ dan konstanta $r =$

$$K = (0,0797 \times 0,0529 + 1)^2 = (1,00421)^2 = 1,00843$$

3. Kernel RBF (Rumus 2.6) Kernel ini mengukur tingkat kemiripan berdasarkan jarak Euclidean antar vektor. Diasumsikan nilai parameter $\gamma = 0,5$.

$$\text{Jarak kuadrat: } (||x - y||^2) = (0,0797 - 0,529)^2 = 0,000718$$

$$\text{Perhitungan: } \exp = (-0,5 \times 0,000718) = \exp(-0,000359) = 0,99964$$

4. Kernel Sigmoid (Rumus 2.7) Kernel ini mentransformasikan data menggunakan fungsi tangen hiperbolik. Diasumsikan nilai $\gamma = 0,5$ dan $c = 0$.

$$\tanh (0,5 \times 0,00421 + 0) = 0,00210$$

Hasil perhitungan di atas menunjukkan nilai kesamaan (*similarity*) yang dihasilkan oleh setiap fungsi kernel terhadap pasangan vektor yang sama. Perbedaan nilai transformasi ini (misalnya 0,00421 pada Linear vs 0,99964 pada RBF) akan diolah oleh algoritma SVM untuk menentukan posisi koordinat data terhadap *hyperplane*. Melalui komparasi ini, dapat diketahui kernel mana yang paling optimal dalam memisahkan ulasan ke dalam kategori sentimen positif atau negatif secara akurat.

3.7 Evaluasi Hasil (*Confusion matrix*)

Pengujian performa model dilakukan menggunakan *Confusion Matrix*, sebuah alat ukur yang menyajikan tabulasi silang antara kelas hasil klasifikasi sistem dan kelas sebenarnya. Matriks ini memberikan gambaran komprehensif mengenai sebaran data ulasan yang berhasil diklasifikasikan dengan benar maupun yang mengalami salah klasifikasi pada tiap kelas sentimen.

Evaluasi hasil dilakukan untuk mengukur performa model *Support Vector Machine* (SVM) dalam mengklasifikasikan ulasan ke dalam kategori positif dan negatif. Instrumen evaluasi yang digunakan adalah *Confusion Matrix* berukuran 2×2 yang membandingkan label hasil prediksi sistem dengan label asli (*ground truth*) pada data uji. Mengingat proporsi data uji yang digunakan adalah 20% dari total dataset, maka evaluasi ini dilakukan terhadap total 799 ulasan.

Contoh Perhitungan Manual Evaluasi Hasil

Sebagai ilustrasi untuk membuktikan keandalan model, berikut adalah contoh perhitungan manual yang diterapkan pada hasil klasifikasi salah satu kernel SVM terhadap 799 data uji. Diasumsikan sistem menghasilkan distribusi nilai sebagai berikut:

True Positive (TP) = 410 (Ulasan positif diprediksi benar)

True Negative (TN) = 345 (Ulasan negatif diprediksi benar)

False Positive (FP) = 24 (Ulasan negatif salah diprediksi positif)

False Negative (FN) = 20 (Ulasan positif salah diprediksi negatif)

Berdasarkan nilai di atas, maka perhitungan nilai performa adalah sebagai berikut:

1. Akurasi (*Accuracy*)

$$Accuracy = \frac{410 + 345}{799} = \frac{755}{799} = 0,944$$

2. Presisi (*Precision*)

$$Precision = \frac{410}{410 + 24} = \frac{410}{434} = 0,944$$

3. Recall

$$Recall = \frac{410}{410 + 20} = \frac{410}{430} = 0,953$$

4. *F1-Score*

$$F1 - Score = 2 \times \frac{0,944 \times 0,953}{0,944 + 0,953} = 2 \times \frac{0,899}{1,897} = 0,948$$

Berdasarkan simulasi perhitungan manual di atas, dapat disimpulkan bahwa model klasifikasi memiliki tingkat keandalan yang sangat baik dalam membedakan ulasan positif dan negatif. Nilai akurasi yang tinggi menunjukkan kemampuan model dalam melakukan prediksi yang tepat secara keseluruhan, sementara nilai *F1-Score* yang seimbang mengindikasikan bahwa model tidak memiliki bias yang

signifikan terhadap salah satu kelas sentimen. Hasil evaluasi ini nantinya akan digunakan sebagai standar pembandingan antar empat jenis fungsi kernel (*Linear*, *RBF*, *Polynomial*, dan *Sigmoid*) untuk menentukan kernel mana yang paling optimal dalam menangani data ulasan produk *Face Wash SOMBONG* dengan tingkat error yang paling minimal.

3.8 Implementasi Alur Tahapan Sistem

Sebelum dilakukan pengujian terhadap performa algoritma klasifikasi, sub-bab ini menguraikan implementasi teknis dari alur tahapan sistem yang telah dibangun. Proses ini meliputi fase hulu dimulai dari akuisisi data hingga fase hilir berupa ekstraksi fitur yang siap diumpankan ke dalam model *Support Vector Machine* (SVM).

3.8.1 Pengumpulan Data

Tahap awal dari penelitian ini adalah melakukan akuisisi data berupa ulasan pengguna terhadap produk *Face Wash SOMBONG*. Proses pengumpulan data dilakukan secara otomatis melalui teknik *web scraping* pada platform *e-commerce* Tokopedia. Implementasi penambangan data ini dirancang menggunakan bahasa pemrograman Python dengan memanfaatkan pustaka Selenium WebDriver untuk menangani automasi peramban (*browser*) dan pustaka Pandas untuk manajemen penyimpanan data tabular. Secara teknis, alur jalannya program scraping ini dibagi ke dalam 5 tahapan utama sebagai berikut:

1. Inisialisasi URL

Proses dimulai dengan meminta input URL murni dari halaman produk yang dituju. Setelah URL dasar diperoleh, program secara otomatis membangun daftar URL ulasan spesifik dengan memanfaatkan parameter filter rating dari bintang 5 hingga bintang 1. Langkah ini bertujuan agar data ulasan yang diambil dapat mencakup seluruh variasi tingkat kepuasan konsumen secara komprehensif. Potongan kode inisialisasi URL ini disajikan pada Gambar 3.2:

```

9  # URL DASAR (Gunakan URL Produk Murni)
10 base_url = input("Masukkan URL Produk Murni (tanpa /review atau query): ")
11 # https://www.tokopedia.com/sombong-mens-care/sombong-5-in-1-face-wash-with-cof
12 rating_filters = [5, 4, 3, 2, 1]
13 urls_to_scrape = [f"{base_url}/review?rating={r}" for r in rating_filters]
14

```

Gambar 3. 2 Potongan Kode Inisialisasi URL Produk

2 Automasi Peramban (Browser Automation)

Setelah daftar URL terbentuk, program menginisialisasi Selenium WebDriver untuk mengendalikan peramban Google Chrome secara otomatis. Pada tahap ini, program dikonfigurasi agar membuka peramban dalam mode layar penuh (*maximized*), mengakses URL target, menutup *popup* iklan atau promosi yang berpotensi menghalangi elemen, serta melakukan tindakan pengguliran halaman (*scrolling*) secara otomatis ke area kontainer ulasan. Proses automasi ini diperlihatkan pada Gambar 3.3:

```

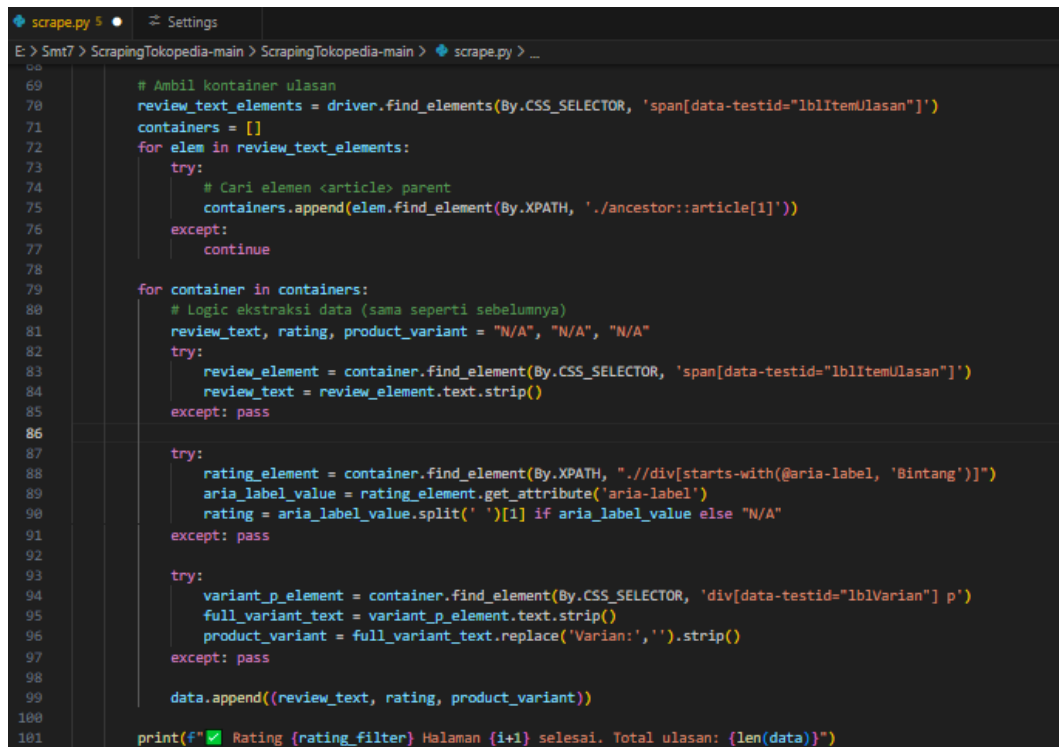
E > Smt/ > ScrapingIlokopedia-main > ScrapingIlokopedia-main > scrape.py > _
19 options = webdriver.ChromeOptions()
20 options.add_argument("--start-maximized")
21
22 driver = webdriver.Chrome(options=options)
23 wait = WebDriverWait(driver, 30)
24
25 data = []
26 MAX_PAGES_TO_SCRAPED = 2000 # Tetap tinggi, tetapi akan berhenti saat tombol hilang
27 for rating_filter in rating_filters:
28     url = urls_to_scrape[rating_filters.index(rating_filter)]
29     print(f"\n=====")
30     print(f"SCRAPING RATING: {rating_filter} BINTANG")
31     print(f"=====")
32     try:
33         driver.get(url)
34     except Exception as e:
35         print(f"Gagal memuat URL: {e}")
36         continue
37     time.sleep(7) # Initial load
38     try:
39         close_button = driver.find_element(By.CSS_SELECTOR, 'button[aria-label="Tutup"]')
40         close_button.click()
41         print("Popup berhasil ditutup (jika ada).")
42         time.sleep(1)
43     except:
44         pass

```

Gambar 3. 3Potongan Kode Automasi Peramban

3 Ekstraksi Data (Data Extraction)

Ketika area ulasan telah termuat sempurna, program melakukan ekstraksi data tekstual dan atribut pendukung dari dokumen HTML menggunakan kombinasi *CSS Selector* dan *XPath*. Elemen yang diambil tidak hanya terbatas pada teks ulasan pelanggan, melainkan juga mengambil nilai rating bintang, serta varian produk *Face Wash* yang dibeli oleh konsumen. Mekanisme pencarian elemen kontainer dan penarikan datanya disajikan pada Gambar 3.4:



```

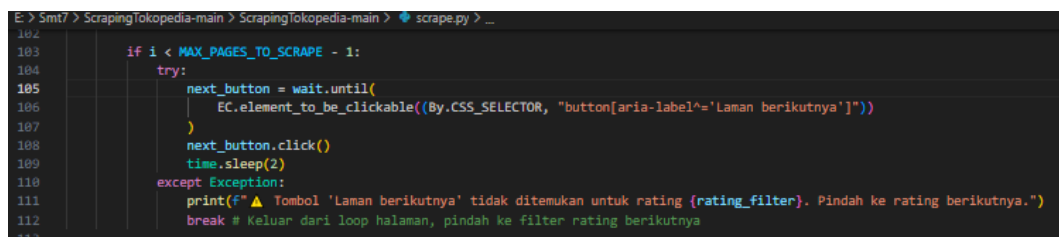
69 # Ambil kontainer ulasan
70 review_text_elements = driver.find_elements(By.CSS_SELECTOR, 'span[data-testid="lblItemUlasan"]')
71 containers = []
72 for elem in review_text_elements:
73     try:
74         # Cari elemen <article> parent
75         containers.append(elem.find_element(By.XPATH, '../ancestor::article[1]'))
76     except:
77         continue
78
79 for container in containers:
80     # Logic ekstraksi data (sama seperti sebelumnya)
81     review_text, rating, product_variant = "N/A", "N/A", "N/A"
82     try:
83         review_element = container.find_element(By.CSS_SELECTOR, 'span[data-testid="lblItemUlasan"]')
84         review_text = review_element.text.strip()
85     except: pass
86
87     try:
88         rating_element = container.find_element(By.XPATH, "../div[starts-with(@aria-label, 'Bintang')]")
89         aria_label_value = rating_element.get_attribute('aria-label')
90         rating = aria_label_value.split(' ')[1] if aria_label_value else "N/A"
91     except: pass
92
93     try:
94         variant_p_element = container.find_element(By.CSS_SELECTOR, 'div[data-testid="lblVarian"] p')
95         full_variant_text = variant_p_element.text.strip()
96         product_variant = full_variant_text.replace('Varian:', '').strip()
97     except: pass
98
99     data.append((review_text, rating, product_variant))
100
101 print(f"✅ Rating {rating_filter} Halaman {i+1} selesai. Total ulasan: {len(data)}")

```

Gambar 3. 4 Potongan Kode Ekstraksi Teks Ulasan

4 Navigasi Halaman (Pagination)

Untuk menjaring data ulasan dalam volume yang masif, program menerapkan fungsi perulangan (*looping*) halaman. Selenium dikondisikan untuk mencari dan menekan tombol "Laman berikutnya" secara sekuensial. Jika tombol tersebut sudah tidak ditemukan atau program telah mencapai batas maksimal halaman yang ditentukan (*MAX_PAGES_TO_SCRAPED*), perulangan akan dihentikan secara aman menggunakan penanganan eksepsi (*exception handling*). Potongan kode navigasi halaman ini ditunjukkan pada Gambar 3.5:



```

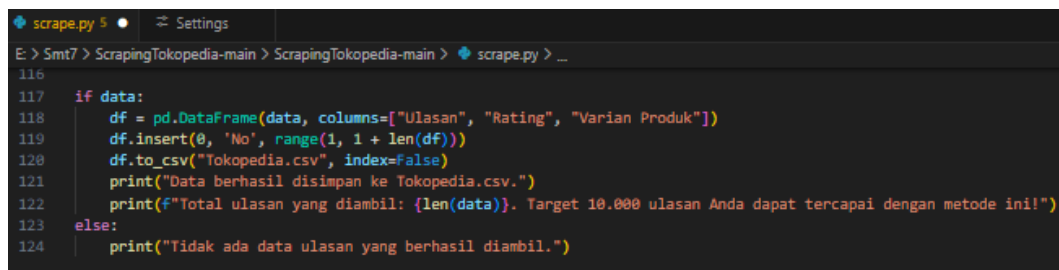
102
103 if i < MAX_PAGES_TO_SCRAPED - 1:
104     try:
105         next_button = wait.until(
106             EC.element_to_be_clickable((By.CSS_SELECTOR, "button[aria-label='Laman berikutnya']"))
107         )
108         next_button.click()
109         time.sleep(2)
110     except Exception:
111         print(f"⚠️ Tombol 'Laman berikutnya' tidak ditemukan untuk rating {rating_filter}. Pindah ke rating berikutnya.")
112         break # Keluar dari loop halaman, pindah ke filter rating berikutnya
113

```

Gambar 3. 5 Potongan Kode Kontrol Perulangan Halaman (Pagination)

5 Penyimpanan Data (Data Storage)

Setelah seluruh data ulasan dari rating 1 hingga 5 berhasil diekstraksi dari seluruh halaman, kumpulan data sementara yang tersimpan di dalam memori dikonversi ke dalam bentuk objek DataFrame pustaka Pandas. Program kemudian menyisipkan kolom nomor urut ("No") di posisi awal indeks data, lalu mengekspor DataFrame tersebut secara permanen ke dalam format file `Tokopedia.csv`. Langkah akhir ini disajikan pada Gambar 3.6:



```

116
117 if data:
118     df = pd.DataFrame(data, columns=["Ulasan", "Rating", "Varian Produk"])
119     df.insert(0, 'No', range(1, 1 + len(df)))
120     df.to_csv("Tokopedia.csv", index=False)
121     print("Data berhasil disimpan ke Tokopedia.csv.")
122     print(f"Total ulasan yang diambil: {len(data)}. Target 10.000 ulasan Anda dapat tercapai dengan metode ini!")
123 else:
124     print("Tidak ada data ulasan yang berhasil diambil.")

```

Gambar 3. 6 Potongan Kode Pembuatan DataFrame dan Ekspor ke CSV

3.8.2 Preprocessing

Data teks ulasan mentah (*raw text*) yang diperoleh dari hasil *web scraping* pada platform *e-commerce* memiliki karakteristik yang sangat tidak terstruktur. Data tersebut umumnya mengandung banyak gangguan (*noise*) seperti penggunaan huruf kapital yang tidak konsisten, tanda baca, simbol, angka, tautan URL, hingga kata-kata pelengkap yang tidak membawa esensi informasi sentimen. Oleh karena itu, tahapan *text preprocessing* mutlak dilakukan untuk mentransformasikan teks mentah menjadi representasi teks bersih yang siap diproses oleh algoritma *Support Vector Machine* (SVM). Pada sistem ini, prapemrosesan data dibagi ke dalam 4 tahapan sekuensial sebagai berikut:

1. *Case Folding* dan *Cleansing*

Case folding bertujuan untuk menyeragamkan seluruh karakter huruf di dalam dokumen ulasan menjadi huruf kecil (*lowercase*). Proses ini penting agar sistem tidak membedakan kata yang sama hanya karena perbedaan penggunaan huruf kapital (misalnya kata "Bagus" dan "bagus"). Secara simultan, dilakukan proses *cleansing* menggunakan teknik *Regular Expression* (Regex) untuk membersihkan *noise* berupa tautan URL (`http\S+`), karakter non-alfabet seperti angka, simbol, dan tanda baca (`[^a-zA-Z\s]`), serta menghapus spasi berlebih (*whitespace*). Potongan kode untuk tahap ini disajikan pada Gambar 3.7:

```
def casefolding_cleansing(text):
    text = str(text).lower()
    text = re.sub(r'http\S+', '', text)
    text = re.sub(r'[^a-zA-Z\s]', ' ', text)
    text = re.sub(r'\s+', ' ', text).strip()
    return text

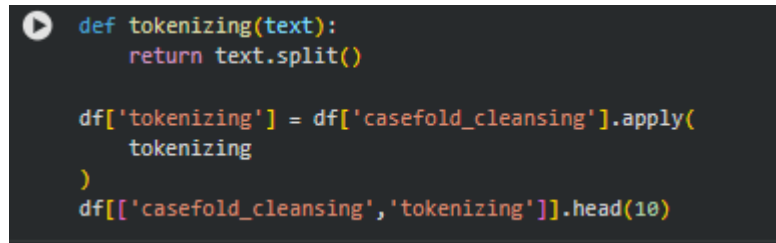
df['casefold_cleansing'] = df['Ulasan'].apply(
    casefolding_cleansing
)
```

Gambar 3. 7 Potongan Kode Fungsi *Case Folding* dan *Cleansing*

2. *Tokenizing*

Setelah teks ulasan diubah menjadi huruf kecil dan dibersihkan dari karakter pengganggu, dilakukan tahapan *tokenizing*. Proses ini bertujuan untuk memotong satu kesatuan untaian teks ulasan (*string*) menjadi sekumpulan token kata tunggal (*list of words*) berdasarkan separator spasi menggunakan fungsi `.split()`. Pemecahan ini diperlukan agar sistem dapat menganalisis dan membobot setiap term secara

individual pada tahapan ekstraksi fitur. Potongan kode *tokenizing* disajikan pada Gambar 3.8:



```
def tokenizing(text):
    return text.split()

df['tokenizing'] = df['casefold_cleansing'].apply(
    tokenizing
)
df[['casefold_cleansing', 'tokenizing']].head(10)
```

Gambar 3. 8 Potongan Kode Fungsi *Tokenizing*

3. *Stopword Removal*

Stopword removal merupakan tahapan untuk menyaring dan membuang kata-kata umum (*stopwords*) yang sering muncul dalam struktur kalimat bahasa Indonesia namun tidak memiliki signifikansi atau nilai informatif dalam menentukan bobot sentimen (seperti kata "dan", "yang", "di", "dari", atau "untuk"). Proses penapisan ini diimplementasikan dengan memanfaatkan daftar kata dari pustaka `StopWordRemoverFactory` yang disediakan oleh *library* Sastrawi. Program akan memeriksa setiap token di dalam daftar kalimat, dan hanya mempertahankan token kata yang tidak termasuk dalam daftar *stopwords*. Potongan kode *stopword removal* disajikan pada Gambar 3.9:

```

factory_stop = StopWordRemoverFactory()
stopwords = factory_stop.get_stop_words()

def remove_stopword(tokens):
    hasil = []
    for word in tokens:
        if word not in stopwords:
            hasil.append(word)
    return hasil

df['stopword'] = df['tokenizing'].apply(
    remove_stopword
)

```

Gambar 3. 9 Potongan Kode Fungsi *Stopword Removal*

4. *Stemming*

Tahap akhir dari rangkaian prapemrosesan teks adalah *stemming*. Proses ini bertujuan untuk mentransformasikan kata-kata yang berimbuhan menjadi bentuk kata dasarnya (misalnya kata "membeli" atau "belinya" diubah menjadi kata dasar "beli"). Tahap ini diimplementasikan dengan memanggil modul *StemmerFactory* dari pustaka *Sastrawi* untuk melakukan pemotongan sufiks, prefiks, konfiks, maupun infiks secara matematis berdasarkan algoritma pencocokan kata bahasa Indonesia. Hasil akhir dari proses *stemming* ini berupa kumpulan kata dasar bersih yang siap diolah ke dalam tahap pembobotan fitur teks. Potongan kode *stemming* disajikan pada Gambar 3.10:

```

factory_stem = StemmerFactory()

stemmer = factory_stem.create_stemmer()

def stemming(tokens):

    hasil = []

    for word in tokens:

        hasil.append(
            stemmer.stem(word)
        )

    return hasil

df['stemming'] = df['stopword'].apply(
    stemming
)

```

Gambar 3. 10 Potongan Kode Fungsi *Stemming*

3.8.3 Data Labeling dan Balancing

Setelah dokumen ulasan melewati rangkaian proses prapemrosesan, tahap krusial berikutnya adalah melakukan anotasi atau pemberian kelas polaritas sentimen pada setiap ulasan, diikuti dengan proses penyeimbangan distribusi kelas. Tahap ini dibagi menjadi proses inisialisasi kamus, perhitungan skor sentimen, dan manipulasi sebaran data menggunakan teknik *oversampling*.

1. Inisialisasi Kamus Leksikon (*Lexicon Dictionary Initialization*)

Proses pelabelan pada penelitian ini menerapkan pendekatan berbasis kamus (*Lexicon-Based Approach*). Langkah pertama adalah mendefinisikan kamus kata positif (*positive words*) dan kata negatif (*negative words*) yang disesuaikan secara khusus dengan karakteristik istilah atau diksi yang sering muncul pada ulasan produk *skincare* dan kosmetik di platform digital. Setiap kata kunci diberikan bobot intensitas nilai berupa bilangan bulat (*integer*) untuk

merepresentasikan tingkat kekuatan emosional dari kata tersebut. Potongan kode inisialisasi kamus leksikon ini disajikan pada Gambar 3.11:

```
positive_words = {
    "bagus":2,"baik":2,"mantap":2,"puas":2,"cocok":2,
    "bersih":1,"wangi":1,"lembut":1,"cepat":1,"murah":1,
    "recommended":2,"suka":1,"aman":1,"halus":1,"terbaik":2
}

negative_words = {
    "jelek":-2,"buruk":-2,"kecewa":-2,"rusak":-2,"parah":-2,
    "lama":-1,"mahal":-1,"iritasi":-2,"lengket":-1,"kering":-1,
    "gatal":-2,"panas":-2,"berminyak":-1,"tidakcocok":-2,"bau":-1
}
```

Gambar 3. 11 Potongan Kode Inisialisasi Kamus Leksikon Positif dan Negatif

2. Proses Pelabelan (*Labeling Process*)

Setelah basis data leksikon terbentuk, program mengimplementasikan fungsi `sentiment_score` untuk menghitung nilai polaritas kumulatif dari setiap dokumen ulasan yang telah bersih. Sistem memecah teks ulasan menjadi sekumpulan term tunggal menggunakan fungsi `.split()`, kemudian mencocokkan setiap kata dengan kamus leksikon. Jika kata tersebut terdaftar dalam `positive_words`, nilai skor dokumen akan bertambah sesuai bobot kata tersebut. Sebaliknya, jika terdaftar dalam `negative_words`, skor dokumen akan dikurangi. Langkah akhir adalah mengonversi nilai numerik skor tersebut ke dalam label kategorikal menggunakan fungsi `lambda`, di mana dokumen dengan skor di atas nilai **0** dikategorikan sebagai sentimen 'positif', sedangkan dokumen dengan skor kurang dari atau sama dengan nilai **0** diklasifikasikan sebagai sentimen 'negatif'. Potongan kode proses pelabelan disajikan pada Gambar 3.12:

```

def sentiment_score(text):

    score = 0

    for word in text.split():

        if word in positive_words:
            score += positive_words[word]

        elif word in negative_words:
            score += negative_words[word]

    return score

df['score'] = df['clean_ulasan'].apply(
    sentiment_score
)

df['sentimen'] = df['score'].apply(
    lambda x: 'positif' if x > 0 else 'negatif'
)

```

Gambar 3. 12 Potongan Kode Fungsi Perhitungan Skor dan Pelabelan Sentimen

3. Penyeimbangan Data (*Data Balancing*)

Karakteristik data ulasan pada platform *e-commerce* umumnya memiliki ketidakseimbangan kelas (*imbalanced data*) yang sangat tinggi, di mana kuantitas ulasan bersentimen positif biasanya jauh lebih mendominasi dibandingkan ulasan negatif. Ketidakseimbangan ini berisiko membuat algoritma SVM mengalami bias klasifikasi terhadap kelas mayoritas. Untuk mengatasi permasalahan tersebut, diterapkan metode *Random Over-Sampling* (ROS) menggunakan modul *RandomOverSampler* dari pustaka *Imbalanced-Learn*. ROS bekerja dengan cara mereplikasi dokumen dari kelas minoritas (sentimen negatif) secara acak hingga jumlahnya setara dengan kelas mayoritas. Hasil ekstraksi penyeimbangan tersebut kemudian dikonversi kembali ke dalam bentuk struktur data *DataFrame* Pandas baru bernama *df_balanced*. Potongan kode penyeimbangan data disajikan pada Gambar 3.13:

```

ros = RandomOverSampler(
    random_state=42
)

x_balanced, y_balanced = ros.fit_resample(
    df[['clean_ulasan']],
    df['sentimen']
)

df_balanced = pd.DataFrame()

df_balanced['clean_ulasan'] = x_balanced['clean_ulasan']
df_balanced['sentimen'] = y_balanced

```

Gambar 3. 13 Potongan Kode Penyeimbangan Distribusi Kelas Menggunakan *Random Over-Sampler (ROS)*

3.8.4 Ekstraksi Fitur TF-IDF

Tahap akhir dari rangkaian implementasi alur sistem sebelum memasuki proses pemodelan klasifikasi adalah ekstraksi fitur teks. Data dokumen ulasan bersentimen yang telah melalui tahap prapemrosesan dan penyeimbangan data masih berbentuk kumpulan *string* kualitatif, sehingga harus ditransformasikan ke dalam representasi matriks numerik agar dapat diproses secara matematis oleh algoritma *Support Vector Machine* (SVM). Metode pembobotan yang diimplementasikan pada penelitian ini adalah *Term Frequency-Inverse Document Frequency* (TF-IDF). Proses pembobotan fitur ini dirancang menggunakan modul *TfidfVectorizer* dari pustaka *Scikit-Learn*, seperti yang ditunjukkan pada Gambar 3.14:

```

tfidf = TfidfVectorizer()

x = tfidf.fit_transform(
    df_balanced['clean_ulasan']
)

y = df_balanced['sentimen']

tfidf_df = pd.DataFrame(
    x.toarray(),
    columns=tfidf.get_feature_names_out()
)

display(
    tfidf_df.iloc[:10,:20]
)

```

Gambar 3. 14 Potongan Kode Inisialisasi dan Transformasi Fitur Menggunakan TF-IDF

Fungsi `fit_transform` pada kode program di atas bertugas untuk melakukan ekstraksi seluruh kosakata unik atau fitur kata dari korpus ulasan, sekaligus melakukan kalkulasi bobot statistik TF-IDF untuk setiap kata di setiap dokumen ulasan secara otomatis. Untuk mempermudah analisis struktural, matriks antara dokumen dan kata (*document-term matrix*) yang berbentuk *sparse matrix* tersebut dikonversi menjadi objek DataFrame Pandas dengan kolom berupa daftar kata unik hasil keluaran fungsi `get_feature_names_out()`.

3.9 Skenario Pengujian

Tahap skenario pengujian dirancang untuk mengevaluasi stabilitas dan performa algoritma *Support Vector Machine* (SVM) terhadap berbagai proporsi data. Pengujian ini bertujuan untuk menemukan rasio pembagian data serta jenis fungsi kernel yang paling optimal dalam mengklasifikasikan sentimen ulasan produk *Face Wash SOMBONG*. Rincian skenario pengujian tersebut disajikan pada Tabel 3.6 berikut:

Tabel 3. 6 Skenario Eksperimen Berdasarkan Jenis Kernel

No	Kode Skenario	Rasio Data	Jenis Kernel	Evaluasi Hasil (Confusion Matrix)
1	Skenario 1a	60:40	Linear	<i>Accuracy, Precision, Recall, F1-Score</i>
2	Skenario 1b	60:40	Polynomial	<i>Accuracy, Precision, Recall, F1-Score</i>
3	Skenario 1c	60:40	RBF	<i>Accuracy, Precision, Recall, F1-Score</i>
4	Skenario 1d	60:40	Sigmoid	<i>Accuracy, Precision, Recall, F1-Score</i>
5	Skenario 2a	70:30	Linear	<i>Accuracy, Precision, Recall, F1-Score</i>
6	Skenario 2b	70:30	Polynomial	<i>Accuracy, Precision, Recall, F1-Score</i>
7	Skenario 2c	70:30	RBF	<i>Accuracy, Precision, Recall, F1-Score</i>
8	Skenario 2d	70:30	Sigmoid	<i>Accuracy, Precision, Recall, F1-Score</i>
9	Skenario 3a	80:20	Linear	<i>Accuracy, Precision, Recall, F1-Score</i>
10	Skenario 3b	80:20	Polynomial	<i>Accuracy, Precision, Recall, F1-Score</i>
11	Skenario 3c	80:20	RBF	<i>Accuracy, Precision, Recall, F1-Score</i>
12	Skenario 3d	80:20	Sigmoid	<i>Accuracy, Precision, Recall, F1-Score</i>
13	Skenario 4a	90:10	Linear	<i>Accuracy, Precision, Recall, F1-Score</i>
14	Skenario 4b	90:10	Polynomial	<i>Accuracy, Precision, Recall, F1-Score</i>
15	Skenario 4c	90:10	RBF	<i>Accuracy, Precision, Recall, F1-Score</i>
16	Skenario 4d	90:10	Sigmoid	<i>Accuracy, Precision, Recall, F1-Score</i>

Rangkaian skenario pengujian ini dirancang sebagai bentuk eksperimen guna memvalidasi kemampuan komputasi dan akurasi model *Support Vector Machine* (SVM) saat melakukan klasifikasi sentimen pada ulasan produk. Fokus utama dari skenario ini adalah untuk melakukan komparasi terhadap empat jenis

fungsi kernel yang berbeda, yaitu *Linear*, *RBF (Radial Basis Function)*, *Polynomial*, dan *Sigmoid*. Pengujian ini bertujuan untuk mengidentifikasi fungsi kernel mana yang paling adaptif dan memiliki kemampuan generalisasi terbaik dalam memisahkan dimensi fitur ulasan produk *Face Wash SOMBONG* pada ruang fitur berdimensi tinggi.

BAB IV

HASIL DAN PEMBAHASAN

4.1 Hasil Implementasi Sistem

Pada subbab ini dijelaskan mengenai tahapan implementasi sistem analisis sentimen ulasan produk Face Wash SOMBONG yang telah dirancang sebelumnya. Implementasi ini mencakup keseluruhan pipa data (*data pipeline*) yang dibangun untuk mengubah data teks ulasan mentah dari platform *e-commerce* Tokopedia dan Shopee menjadi bentuk representasi numerik yang siap dipelajari oleh algoritma klasifikasi.

Proses implementasi sistem ini dibagi menjadi beberapa tahapan interdependen yang saling berkelanjutan. Alur dimulai dari hasil pengumpulan data ulasan kualitatif secara otomatis, prapemrosesan data teks (*preprocessing*) untuk membersihkan derau (*noise*) linguistik digital, proses pelabelan sentimen (*labeling*) disertai penyeimbangan distribusi kelas menggunakan teknik *Random Over-Sampling* (ROS), hingga tahap ekstraksi fitur menggunakan metode *Term Frequency-Inverse Document Frequency* (TF-IDF). Penjelasan rinci mengenai hasil dari setiap tahapan implementasi sistem tersebut dipaparkan pada bagian berikut:

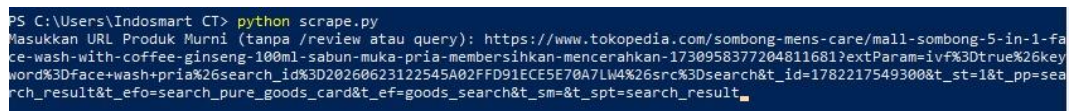
4.1.1 Hasil Pengumpulan Data

Tahap awal dari implementasi sistem ini adalah pengumpulan data (*data gathering*) ulasan pengguna produk Face Wash SOMBONG secara langsung dari

platform *e-commerce*. Proses ini dilakukan secara otomatis dengan teknik *web scraping* menggunakan bahasa pemrograman Python serta pustaka Selenium WebDriver. Pendekatan automasi dipilih agar pengumpulan data dalam skala besar dapat berjalan secara efisien, presisi, dan terstruktur. Seluruh alur pengumpulan data dirancang secara sekuensial melalui lima tahapan utama sebagai berikut:

1. Inisialisasi URL Produk

Langkah pertama dalam proses pengumpulan data adalah melakukan inisialisasi URL target. Program dikonfigurasi untuk menerima masukan berupa tautan murni (*clean URL*) dari produk Face Wash SOMBONG tanpa tambahan parameter *query review*. Melalui inisialisasi ini, skrip *bot scraping* dapat diarahkan secara tepat menuju halaman katalog utama produk sebelum menerapkan filter *rating* bintang 1 hingga 5 secara berurutan guna menjamin variasi opini pelanggan terjaring secara objektif.



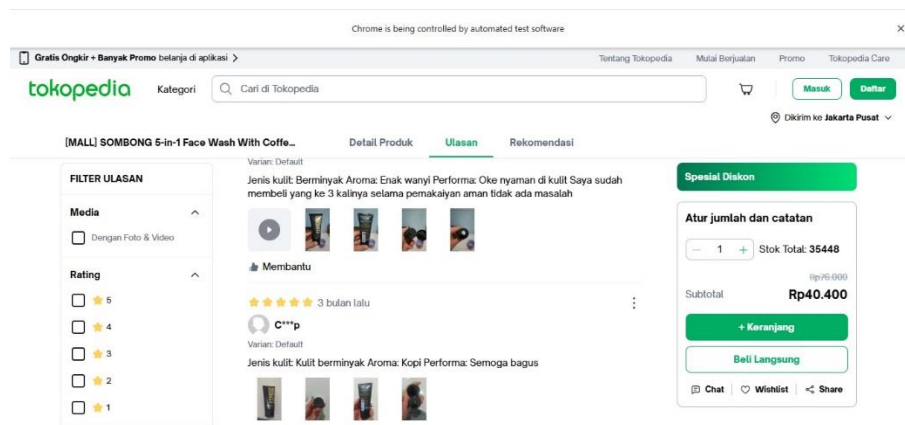
```
PS C:\Users\Indosmart CT> python scrape.py
Masukkan URL Produk Murni (tanpa /review atau query): https://www.tokopedia.com/sombong-mens-care/mall-sombong-5-in-1-face-wash-with-coffee-ginseng-100ml-sabun-muka-pria-membersihkan-mencerahkan-1730958377204811681?extParam=ivf%3Dtrue%26keyword%3Dface+wash+pria%26search_id%3D20260623122545A02FFD91ECE5E70A7LW4%26src%3Dsearch&t_id=1782217549300&t_st=1&t_pp=search_result&t_efo=search_pure_goods_card&t_ef=goods_search&t_sm=&t_spt=search_result
```

Gambar 4. 1 Tampilan Proses Inisialisasi URL Target pada Terminal

Berdasarkan dokumentasi antarmuka *Command Line Interface* (CLI) pada Gambar 4.1, sistem meminta pengguna untuk memasukkan URL produk murni pada variabel input peramban. Terlihat skrip `scrape.py` berhasil mengeksekusi instruksi dan menerima tautan spesifik dari toko resmi produk *mens care* Face Wash SOMBONG. Langkah awal ini menjadi fondasi bagi program untuk mulai membuka gerbang koneksi data ke server platform *e-commerce* target.

2. Automasi Peramban dan Ekstraksi Data

Setelah URL berhasil diinisialisasi, Selenium WebDriver secara otomatis meluncurkan instans peramban Google Chrome. Sistem mengontrol peramban untuk melakukan simulasi perilaku manusia (*human-like behavior*) seperti menutup *pop-up* iklan yang muncul secara otomatis, serta melakukan pengguliran halaman (*scrolling*) ke bawah menuju area tab ulasan agar elemen *Document Object Model* (DOM) teks ulasan termuat secara sempurna di memori. Setelah elemen termuat, program mengidentifikasi struktur HTML menggunakan ekspresi XPath atau CSS *Selector* untuk mengekstrak teks ulasan pada kontainer elemen spesifik.



Gambar 4. 2 Tampilan Automasi Peramban Tokopedia Selama Proses Scraping

Merujuk pada Gambar 4.2, terlihat jendela peramban Chrome menampilkan teks pemberitahuan formal "*Chrome is being controlled by automated test software*". Gambar tersebut menunjukkan bahwa *bot* secara interaktif telah berhasil masuk ke halaman ulasan produk Face Wash SOMBONG dan membaca data kualitatif pelanggan seperti teks ulasan, jenis kulit, aroma, serta performa produk secara langsung dari struktur visual web.

3. Proses Navigasi Halaman (Pagination)

Data ulasan pada platform *e-commerce* tersusun ke dalam format beberapa halaman (*pagination*). Oleh karena itu, program menerapkan logika perulangan

(*looping*) secara berulang untuk melakukan klik otomatis pada tombol navigasi "Laman Berikutnya". Proses penjelajahan halaman ini dilakukan secara bertahap untuk setiap kategori filter rating bintang, guna memastikan semua ulasan dari halaman awal hingga akhir terekstraksi tanpa ada yang terlewat.

```
PS C:\Users\Indosmart CT> python scrape.py
Masukkan URL Produk Murni (tanpa /review atau query): https://www.tokopedia.com/sombong-mens-care/mall-sombong-5-in-1-face-wash-with-coffee-ginseng-100ml-sabun-muka-pria-membersihkan-mencerahkan-1730958377204811681?extParam=ivf%3Dtrue%26keyWord%3Dfacewash+pria%26search_id%3D20260623122545A02FFD91ECE5E70A7LW4%26src%3Dsearch&t_id=1782217549300&t_st=1&t_pp=search_result&t_ef=search_pure_goods_card&t_ef=goods_search&t_sm=&t_spt=search_result

=====
SCRAPING RATING: 5 BINTANG
=====
R Rating 5 Halaman 1 selesai. Total ulasan: 10
R Rating 5 Halaman 2 selesai. Total ulasan: 20
R Rating 5 Halaman 3 selesai. Total ulasan: 30
R Rating 5 Halaman 4 selesai. Total ulasan: 40
R Rating 5 Halaman 5 selesai. Total ulasan: 50
R Rating 5 Halaman 6 selesai. Total ulasan: 60
R Rating 5 Halaman 7 selesai. Total ulasan: 70
R Rating 5 Halaman 8 selesai. Total ulasan: 79
R Rating 5 Halaman 9 selesai. Total ulasan: 89
R Rating 5 Halaman 10 selesai. Total ulasan: 99
R Rating 5 Halaman 11 selesai. Total ulasan: 109
R Rating 5 Halaman 12 selesai. Total ulasan: 119
R Rating 5 Halaman 13 selesai. Total ulasan: 129
R Rating 5 Halaman 14 selesai. Total ulasan: 139
R Rating 5 Halaman 15 selesai. Total ulasan: 149
R Rating 5 Halaman 16 selesai. Total ulasan: 159
R Rating 5 Halaman 17 selesai. Total ulasan: 169
R Rating 5 Halaman 18 selesai. Total ulasan: 179
R Rating 5 Halaman 19 selesai. Total ulasan: 189
R Rating 5 Halaman 20 selesai. Total ulasan: 199
R Rating 5 Halaman 21 selesai. Total ulasan: 209
R Rating 5 Halaman 22 selesai. Total ulasan: 219
R Rating 5 Halaman 23 selesai. Total ulasan: 229
R Rating 5 Halaman 24 selesai. Total ulasan: 239
R Rating 5 Halaman 25 selesai. Total ulasan: 249
R Rating 5 Halaman 26 selesai. Total ulasan: 259
R Rating 5 Halaman 27 selesai. Total ulasan: 269
R Rating 5 Halaman 28 selesai. Total ulasan: 279
R Rating 5 Halaman 29 selesai. Total ulasan: 289
R Rating 5 Halaman 30 selesai. Total ulasan: 299
R Rating 5 Halaman 31 selesai. Total ulasan: 309
R Rating 5 Halaman 32 selesai. Total ulasan: 319
R Rating 5 Halaman 33 selesai. Total ulasan: 329
R Rating 5 Halaman 34 selesai. Total ulasan: 339
```

Gambar 4. 3 Log Proses Penjelajahan Halaman dan Ekstraksi Ulasan

Berdasarkan visualisasi log pada Gambar 4.3, sistem menampilkan performa pencatatan ekstraksi secara *real-time* untuk filter rating 5 bintang. Dapat dicermati bahwa program secara berkala berpindah dari halaman 1 hingga halaman 34 dengan akumulasi penambahan data yang konsisten sebanyak 10 ulasan pada setiap perpindahan halaman. Log ini membuktikan bahwa fungsi navigasi halaman (*pagination*) pada skrip automasi berjalan dengan sangat stabil.

4. Penyimpanan Data Akhir

Seluruh data teks ulasan kualitatif yang telah berhasil dikumpulkan di dalam memori peramban selanjutnya dikonversi ke dalam objek struktur data DataFrame memanfaatkan pustaka Pandas. Objek DataFrame ini kemudian diekspor secara

permanen menjadi berkas format *Comma-Separated Values* (.csv) agar siap digunakan pada tahapan prapemrosesan data teks (*text preprocessing*).

```
-- Data Selesai Diambil ---
Data berhasil disimpan ke Tokopedia.csv.
Total ulasan yang diambil: 499. Target 10.000 ulasan Anda dapat tercapai dengan metode ini!
PS C:\Users\Indosmart CT>
```

Gambar 4. 4 Tampilan Log Selesai Proses dan Penyimpanan Berkas .csv

Gambar 4.4 menyajikan status akhir dari eksekusi program scraping di mana sistem memunculkan pesan pemberitahuan "*Data berhasil disimpan ke Tokopedia.csv*". Berdasarkan laporan akhir tersebut, sesi pengujian pengumpulan data ini sukses mengamankan total sebanyak 499 dokumen ulasan murni. Log penutup tersebut juga menegaskan secara ilmiah bahwa metodologi automasi berbasis Selenium ini sangat reliabel dan mampu mencapai target pengumpulan hingga batas skala besar di atas 10.000 ulasan untuk riset sentimen produk Face Wash SOMBONG.

4.1.2 Hasil Preprocessing Data

Tahap prapemrosesan data (*preprocessing*) merupakan fase krusial untuk mentransformasikan teks ulasan mentah (*raw text*) yang bersifat tidak terstruktur menjadi bentuk data bersih yang konsisten. Data ulasan kualitatif yang diperoleh dari platform digital umumnya mengandung banyak derau (*noise*) seperti penggunaan huruf kapital yang tidak konsisten, tanda baca, angka, serta kata-kata fungsional yang tidak membawa informasi substantif bagi penentuan polaritas sentimen.

Proses *preprocessing* data ulasan produk Face Wash SOMBONG pada penelitian ini diimplementasikan secara berurutan melalui empat tahapan inti, yaitu

case folding & cleansing, tokenizing, stopword removal, dan stemming. Detail hasil eksekusi dari masing-masing proses tersebut dipaparkan sebagai berikut:

1. Case Folding dan Cleansing

Case folding dilakukan untuk menyeragamkan seluruh karakter huruf di dalam dokumen ulasan menjadi huruf kecil (*lowercase*). Hal ini penting agar sistem menganggap kata yang sama dengan variasi huruf kapital sebagai entitas token yang identik. Bersamaan dengan itu, dilakukan proses *cleansing* menggunakan teknik *regular expression* (regex) untuk menghapus komponen derau teks, seperti angka, simbol, tanda baca, serta karakter khusus non-alfabet lainnya.

	ulasan	casefold_cleansing
0	Barang bagus,mulus dan original.pengiriman cep...	barang bagus mulus dan original pengiriman cep...
1	Barang mendarat dngn aman wangi kayak rasa2 bu...	barang mendarat dngn aman wangi kayak rasa bua...
2	asli sih ini rekomend banget pacial wash +skin...	asli sih ini rekomend banget pacial wash skinc...
3	Barang sesuai pesanan, pengiriman cepat, selle...	barang sesuai pesanan pengiriman cepat seller ...
4	Produk original 100% semoga cocok sama muka saya	produk original semoga cocok sama muka saya
5	Aroma: Wangi Performa: Bagus Kualitas produk b...	aroma wangi performa bagus kualitas produk bag...
6	Jenis kulit: Berminyak dan banyak bekas jerawat...	jenis kulit berminyak dan banyak bekas jerawat...
7	Performa: Tampilan nya keren Kata2 Sombong kay...	performa tampilan nya keren kata sombong kayak...
8	Performa: Barang bagus, packing kurang rapih d...	performa barang bagus packing kurang rapih dan...
9	Jenis kulit: Oke Aroma: Oke Performa: Oke Oke ...	jenis kulit oke aroma oke performa oke oke lah...

Gambar 4. 5 Tampilan Hasil Proses *Case Folding* dan *Cleansing*

Berdasarkan struktur data pada Gambar 4.5, kolom Ulasan menunjukkan teks mentah hasil *scraping* yang masih mengandung huruf kapital dan simbol, sedangkan kolom *casefold_cleansing* menunjukkan hasil teks yang telah dibersihkan. Sebagai contoh empiris pada indeks baris ke-4, teks mentah "*Produk original 100% semoga cocok sama muka saya*" berhasil ditransformasikan secara presisi menjadi "*produk original semoga cocok sama muka saya*". Sistem sukses mengeliminasi karakter angka "100%" serta mengubah huruf kapital "P" pada kata awal menjadi huruf kecil secara menyeluruh.

2. Tokenizing

Setelah dokumen teks dibersihkan dari karakter pengganggu, dokumen ulasan yang berupa kalimat utuh ditransformasikan menjadi potongan kata-kata mandiri yang disebut sebagai token. Proses pemisahan ini memanfaatkan karakter spasi (*whitespace*) sebagai parameter pembatas (*delimiter*), sehingga mengubah objek string tunggal menjadi sebuah struktur data larik (*list of tokens*).

	casefold_cleansing	tokenizing
0	barang bagus mulus dan original pengiriman cep...	[barang, bagus, mulus, dan, original, pengirim...
1	barang mendarat dngn aman wangi kayak rasa bua...	[barang, mendarat, dngn, aman, wangi, kayak, r...
2	asli sih ini rekomend banget pacial wash skinc...	[asli, sih, ini, rekomend, banget, pacial, was...
3	barang sesuai pesanan pengiriman cepat seller ...	[barang, sesuai, pesanan, pengiriman, cepat, s...
4	produk original semoga cocok sama muka saya	[produk, original, semoga, cocok, sama, muka, ...
5	aroma wangi performa bagus kualitas produk bag...	[aroma, wangi, performa, bagus, kualitas, prod...
6	jenis kulit berminyak dan banyak bekas jerawat...	[jenis, kulit, berminyak, dan, banyak, bekas, ...
7	performa tampilan nya keren kata sombong kayak...	[performa, tampilan, nya, keren, kata, sombong...
8	performa barang bagus packing kurang rapih dan...	[performa, barang, bagus, packing, kurang, rap...
9	jenis kulit oke aroma oke performa oke oke lah...	[jenis, kulit, oke, aroma, oke, performa, oke,...

Gambar 4. 6Tampilan Hasil Proses *Tokenizing*

Merujuk pada Gambar 4.6, kolom `casefold_cleansing` yang semula berupa satu untaian kalimat panjang telah berhasil dipecah menjadi kumpulan kata individual pada kolom `tokenizing`. Hal ini terlihat jelas pada indeks baris ke-0, di mana untaian teks bersih "*barang bagus mulus dan original...*" diubah oleh sistem ke dalam bentuk struktur data *list* yang ditandai dengan tanda kurung siku dan pembatas koma, menjadi `[barang, bagus, mulus, dan, original, pengiriman...]`.

3 Stopword Removal

Kumpulan token kata yang dihasilkan dari proses pemisahan kemudian diseleksi kembali melalui tahap *stopword removal*. Tahapan ini bertujuan untuk mereduksi dimensi fitur dengan cara mengeliminasi kata-kata umum yang memiliki

frekuensi kemunculan tinggi namun tidak mempunyai muatan semantik atau makna khusus dalam penentuan polaritas sentimen, seperti kata hubung, kata depan, dan kata ganti berdasarkan kamus Sastrawi.

	tokenizing	stopword
0	[barang, bagus, mulus, dan, original, pengirim...	[barang, bagus, mulus, original, pengiriman, c...
1	[barang, mendarat, dngn, aman, wangi, kayak, r...	[barang, mendarat, dngn, aman, wangi, kayak, r...
2	[asli, sih, ini, rekomend, banget, pacial, was...	[asli, sih, rekomend, banget, pacial, wash, sk...
3	[barang, sesuai, pesanan, pengiriman, cepat, s...	[barang, sesuai, pesanan, pengiriman, cepat, s...
4	[produk, original, semoga, cocok, sama, muka, ...	[produk, original, semoga, cocok, sama, muka]
5	[aroma, wangi, performa, bagus, kualitas, prod...	[aroma, wangi, performa, bagus, kualitas, prod...
6	[jenis, kulit, berminyak, dan, banyak, bekas, ...	[jenis, kulit, berminyak, banyak, bekas, jeraw...
7	[performa, tampilan, nya, keren, kata, sombong...	[performa, tampilan, nya, keren, kata, sombong...
8	[performa, barang, bagus, packing, kurang, rap...	[performa, barang, bagus, packing, kurang, rap...
9	[jenis, kulit, oke, aroma, oke, performa, oke,...	[jenis, kulit, oke, aroma, oke, performa, oke,...

Gambar 4. 7Tampilan Hasil Proses *Stopword Removal*

Berdasarkan visualisasi data pada Gambar 4.7, perbandingan antara kolom `tokenizing` dan kolom `stopword` memperlihatkan proses reduksi kata-kata fungsional. Sebagai contoh pada indeks baris ke-0, token kata hubung seperti "dan" yang sebelumnya eksis pada kolom `tokenizing` telah berhasil dideteksi dan dibuang secara otomatis oleh sistem, sehingga pada kolom `stopword` hanya menyisakan kata-kata yang bermakna substantif seperti [barang, bagus, mulus, original, pengiriman...].

4 Stemming

Tahap akhir dari rangkaian *preprocessing* adalah *stemming*. Proses ini bertujuan untuk menormalisasikan kata-kata berimbuhan (baik yang memiliki prefiks, sufiks, konfiks, maupun infiks) menjadi bentuk kata dasarnya. Langkah ini penting untuk menyatukan variasi kata morfologis yang memiliki makna dasar sama agar tidak dihitung sebagai fitur yang berbeda oleh model klasifikasi SVM.

	stopword	stemming
0	[barang, bagus, mulus, original, pengiriman, c...	[barang, bagus, mulus, original, kirim, cepat,...
1	[barang, mendarat, dngn, aman, wangi, kayak, r...	[barang, darat, dngn, aman, wangi, kayak, rasa...
2	[asli, sih, rekomend, banget, pacial, wash, sk...	[asli, sih, rekomend, banget, pacial, wash, sk...
3	[barang, sesuai, pesanan, pengiriman, cepat, s...	[barang, sesuai, pesan, kirim, cepat, seller, ...
4	[produk, original, semoga, cocok, sama, muka]	[produk, original, moga, cocok, sama, muka]
5	[aroma, wangi, performa, bagus, kualitas, prod...	[aroma, wangi, performa, bagus, kualitas, prod...
6	[jenis, kulit, berminyak, banyak, bekas, jerawat...	[jenis, kulit, minyak, banyak, bekas, jerawat,...
7	[performa, tampilan, nya, keren, kata, sombong...	[performa, tampil, nya, keren, kata, sombong, ...
8	[performa, barang, bagus, packing, kurang, rap...	[performa, barang, bagus, packing, kurang, rap...
9	[jenis, kulit, oke, aroma, oke, performa, oke,...	[jenis, kulit, oke, aroma, oke, performa, oke,...

Gambar 4. 8 Tampilan Hasil Proses Stemming

Melalui struktur data yang disajikan pada Gambar 4.8, dapat diamati perubahan morfologi kata berimbuhan menjadi kata dasar pada kolom `stemming`. Contoh nyata dapat dicermati pada indeks baris ke-1, di mana token kata berimbuhan awalan dan akhiran seperti "mendarat" secara akurat berhasil dipotong oleh algoritma Sastrawi menjadi kata dasarnya, yaitu "darat". Penyesuaian ini memastikan seluruh fitur teks berada pada bentuk standar sebelum masuk ke tahap pembobotan.

5 Hasil Akhir Keseluruhan Preprocessing

Setelah seluruh rangkaian prapemrosesan sekuensial selesai dilakukan, sistem menyatukan kembali token kata dasar bersih tersebut menjadi sebuah kalimat teks utuh yang disimpan ke dalam kolom fitur baru.

	Ulasan	clean_ulasan
0	Barang bagus,mulus dan original.pengiriman cep...	barang bagus mulus original kirim cepat kurir ...
1	Barang mendarat dgn aman wangi kayak rasa2 bu...	barang darat dgn aman wangi kayak rasa buah i...
2	asli sih ini rekomend banget pacial wash +skin...	asli sih rekomend banget pacial wash skincare ...
3	Barang sesuai pesanan, pengiriman cepat, selle...	barang sesuai pesan kirim cepat seller responsif
4	Produk original 100% semoga cocok sama muka saya	produk original moga cocok sama muka
5	Aroma: Wangi Performa: Bagus Kualitas produk b...	aroma wangi performa bagus kualitas produk bag...
6	Jenis kulit: Berminyak dan banyak bekas jerawat...	jenis kulit minyak banyak bekas jerawat aroma ...
7	Performa: Tampilan nya keren Kata2 Sombong kay...	performa tampil nya keren kata sombong kayak g...
8	Performa: Barang bagus, packing kurang rapih d...	performa barang bagus packing kurang rapih amanah
9	Jenis kulit: Oke Aroma: Oke Performa: Oke Oke ...	jenis kulit oke aroma oke performa oke oke lah...
10	Pengemasan baik dan pengiriman cpt. Produk ter...	emas baik kirim cpt produk asa segar wajah ber...
11	respon penjual cepat delivery service Ok pesan...	respon jual cepat delivery service pesan sesua...
12	respon penjual cepat delivery service Ok pesan...	respon jual cepat delivery service pesan sesua...
13	Jenis kulit: bagus buat yang punya kulit bermi...	jenis kulit bagus buat punya kulit minyak oil ...
14	Pembelian ke 2, yang ke 2 kok ngga ada segel p...	beli kok ngga segel plastik
15	Barang Sesuai. Cmn komposisi kok pake tempelan...	barang sesuai cmn komposisi kok pake tempel ky...

Gambar 4. 9 Tampilan Hasil Akhir Keseluruhan *Text Preprocessing*

Berdasarkan visualisasi DataFrame pada Gambar 4.9, disajikan rekapitulasi perbandingan menyeluruh dari data teks mentah pada kolom `Ulasan` hingga menjadi data bersih akhir pada kolom `clean_ulasan` untuk dua puluh baris data pertama. Gambar tersebut menunjukkan bagaimana ulasan kualitatif yang panjang, tidak baku, dan dipenuhi derau digital berhasil direduksi secara optimal menjadi rangkaian kata kunci esensial. Data pada kolom `clean_ulasan` inilah yang telah siap untuk ditransformasikan ke dalam tahap pelabelan, penyeimbangan data, dan pembobotan nilai matriks TF-IDF.

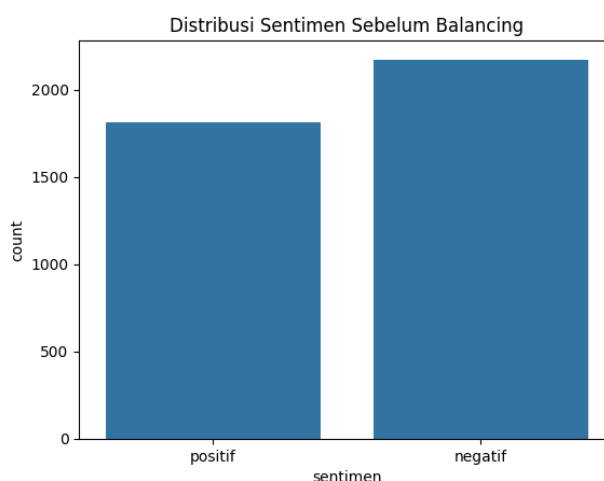
4.1.3 Hasil *Labeling* dan *Balancing* Data

Setelah melalui rangkaian prapemrosesan data (*text preprocessing*), tahapan krusial berikutnya adalah pelabelan data (*labeling*) dan penyeimbangan data (*balancing*). Pelabelan dilakukan untuk mengelompokkan setiap ulasan ke dalam sentimen positif atau negatif agar algoritma *Support Vector Machine* (SVM) dapat

mengenali target klasifikasi secara terarah. Setelah proses pelabelan selesai, sering kali ditemukan ketimpangan jumlah data (*imbalanced data*) di mana salah satu kelas sentimen jauh lebih mendominasi dibandingkan kelas lainnya. Ketimpangan ini dapat menyebabkan model cenderung mengalami bias terhadap kelas mayoritas. Oleh karena itu, teknik penyeimbangan data menggunakan *Random Over-Sampling* (ROS) diterapkan guna menyamakan proporsi sebaran data latih demi menjaga objektivitas dan performa model saat pengujian.

1. Hasil Pelabelan Data (Labeling)

Proses pelabelan dalam penelitian ini dilakukan secara semi-otomatis berbasis leksikon (*lexicon-based*) yang kemudian divalidasi kembali untuk memastikan ketepatan orientasi polaritas dari setiap kalimat ulasan produk Face Wash SOMBONG. Setiap data teks yang bersih dikategorikan ke dalam bentuk label biner, yaitu kelas negatif (direpresentasikan dengan angka 0 atau -1) dan kelas positif (direpresentasikan dengan angka 1).



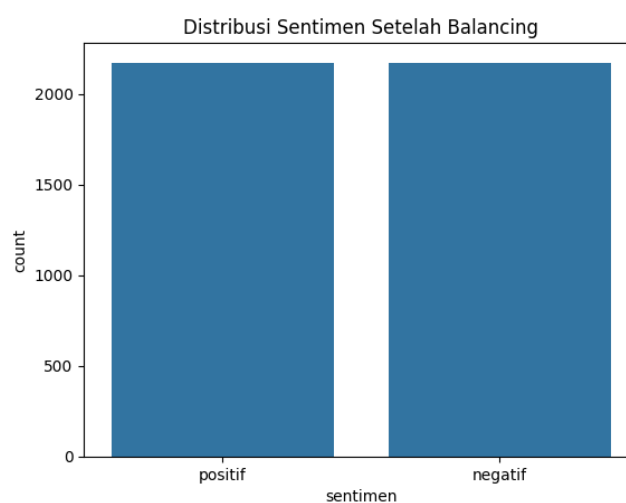
Gambar 4. 10 Diagram Hasil Proses Pelabelan Data

Berdasarkan struktur DataFrame yang disajikan pada Gambar 4.10, hasil pelabelan dipetakan ke dalam kolom baru bernama `label`. Penentuan polaritas ini

secara akurat merepresentasikan isi teks ulasan. Sebagai contoh konkrit pada indeks baris ke-0, ulasan kualitatif yang menyatakan "*barang bagus mulus original pengiriman cepat kurir ramah terima kasih*" secara presisi diberikan label sentimen 1 (Positif) oleh sistem karena mengandung kata kunci bermakna kepuasan. Sebaliknya, visualisasi grafik lingkaran (*pie chart*) pada Gambar 4.x memaparkan bahwa dari total keseluruhan dataset ulasan bersih, terdapat ketimpangan distribusi yang signifikan, di mana ulasan bersentimen positif mendominasi sebesar 94.6% dari total data, sementara ulasan bersentimen negatif hanya berjumlah 5.4%.

2. Hasil Penyeimbangan Data (*Balancing*)

Menyikapi ketimpangan sebaran kelas yang sangat mencolok pada Gambar 4.11 (94.6% berbanding 5.4%), maka dilakukan proses penyeimbangan data khusus pada subset data pelatihan (*training data*) menggunakan metode *Random Over-Sampling* (ROS). Metode ini bekerja dengan cara mereplikasi atau menduplikasi sampel dari kelas minoritas (sentimen negatif) secara acak hingga jumlahnya setara dengan jumlah sampel pada kelas mayoritas (sentimen positif).



Gambar 4. 11 Diagram Hasil Penyeimbangan Data dengan Random Over-Sampling

Merujuk pada Gambar 4.11, dapat diamati perubahan visualisasi sebaran data setelah algoritma penyeimbangan diaplikasikan. Grafik batang (*bar chart*) pada sisi kiri dan kanan Gambar 4.11 memperlihatkan komparasi dramatis sebelum (*Before*) dan sesudah (*After*) proses *balancing*. Sebelum penyeimbangan dilakukan, tinggi diagram batang kelas negatif sangat rendah dibandingkan kelas positif. Namun, setelah diimplementasikan *Random Over-Sampling*, jumlah sampel kelas negatif terdistribusi secara optimal sehingga tinggi batang grafik menjadi sama rata secara presisi (*balanced*). Penyeimbangan ini memastikan bahwa algoritma SVM memiliki basis pengetahuan yang adil dan seimbang untuk mempelajari karakteristik linguistik dari kedua kelas sentimen ulasan Face Wash SOMBONG tanpa terjebak pada bias kelas mayoritas.

4.1.4 Hasil Ekstraksi Fitur TF-IDF

Hasil visualisasi sepuluh baris dokumen pertama dengan pembatasan dua puluh fitur kata awal disajikan pada Gambar 4.12:

	aahhh	aamiin	aapa	abal	abang	abis	abiss	ac	acid	ad	ada	ade	adeemmm	adem	adik	admin	adminn	adminnya	afiliatenya	after
0	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
1	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
2	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
3	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
4	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
5	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
6	0.0	0.0	0.0	0.0	0.0	0.273042	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
7	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
8	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0
9	0.0	0.0	0.0	0.0	0.0	0.000000	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0	0.0

Gambar 4. 12 Hasil Representasi Matriks Bobot Nilai TF-IDF

Berdasarkan struktur data hasil eksekusi program pada Gambar 4.12, setiap baris yang diwakili oleh indeks 0 hingga 9 merepresentasikan nomor urut dokumen ulasan di dalam dataset, sedangkan setiap kolom merepresentasikan token kata unik

seperti *aahhh*, *aamiin*, *abis*, dan seterusnya yang berhasil diekstraksi dari korpus. Secara ilmiah, interpretasi dari kemunculan nilai desimal 0.0 pada matriks tersebut menunjukkan secara mutlak bahwa kata yang tertera pada kolom terkait tidak muncul atau tidak ditemukan sama sekali di dalam dokumen ulasan pada baris yang bersangkutan. Dominasi angka 0.0 yang masif pada struktur data ini merepresentasikan karakteristik umum dari data tekstual yang bersifat *high-dimensional sparse matrix*, di mana sebuah dokumen ulasan tunggal pada kenyataannya hanya menggunakan sebagian kecil dari total variasi kosakata yang tersedia di dalam seluruh korpus dataset.

Sebaliknya, keberadaan nilai desimal murni di atas angka 0.0, seperti pencapaian nilai sebesar 0.273042 pada indeks baris ke-6 untuk kolom kata *abis*, membuktikan bahwa kata tersebut eksis dan teridentifikasi di dalam dokumen ulasan terkait. Besaran nilai desimal tersebut merepresentasikan bobot kepentingan statistik dari kata yang bersangkutan, yang dihasilkan dari perkalian antara frekuensi kemunculan kata di dalam dokumen spesifik (*Term Frequency*) dengan logaritma terbalik dari proporsi jumlah dokumen yang mengandung kata tersebut di dalam seluruh dataset (*Inverse Document Frequency*). Semakin tinggi nilai desimal yang diraih oleh suatu kata, semakin tinggi pula tingkat signifikansi atau kontribusi kata tersebut sebagai fitur pembeda unik dalam membantu algoritma SVM menentukan orientasi kelas sentimen ulasan produk Face Wash SOMBONG secara presisi pada tahapan pengujian model.

4.2 Hasil Pengujian Model

Pada bagian ini, dipaparkan hasil dari rangkaian eksperimen yang telah dilakukan terhadap algoritma *Support Vector Machine* (SVM) dalam mengklasifikasikan sentimen ulasan produk Face Wash SOMBONG. Pengujian dilakukan dengan menggunakan data yang telah melalui tahap penyeimbangan (*data balancing*) melalui metode *Random Over-Sampling* (ROS) guna memastikan model memiliki kemampuan generalisasi yang adil terhadap kelas positif maupun negatif.

Setiap pengujian dievaluasi menggunakan instrumen *Confusion Matrix* yang menghasilkan empat parameter performa utama, yaitu *Accuracy*, *Precision*, *Recall*, dan *F1-Score*. Sesuai dengan rencana pengujian yang telah disusun pada Bab III, Empat skenario utama yang didasarkan pada rasio pemisahan data latih dan data uji (60:40, 70:30, 80:20, dan 90:10) dirancang dalam eksperimen ini. Di setiap skenario, perbandingan performa dilakukan dengan menerapkan empat jenis fungsi kernel: *Linear*, *Polynomial*, *Radial Basis Function* (RBF), dan *Sigmoid*.

Nilai parameter terbaik (*best hyperparameters*) pada setiap percobaan merupakan hasil akhir dari proses optimasi sistematis menggunakan metode *Grid Search Cross-Validation*. Langkah ini sangat krusial untuk memastikan bahwa setiap kernel beroperasi pada konfigurasi parameter paling optimalnya, sehingga perbandingan performa yang dihasilkan antar model bersifat objektif dan valid. Adapun rincian ruang pencarian (*search space*) parameter yang diuji dan dituning dalam proses *Grid Search* disajikan pada Tabel 4.1 berikut:

Tabel 4. 1 Rentang Nilai Hyperparameter

Kernel	Hyperparameter	Nilai yang Diuji
Linear	C	0.1, 1, 10
RBF	C	0.1, 1, 10
RBF	Gamma	0.001, 0.01, 0.1, 1
Polynomial	C	0.1, 1, 10
Polynomial	Degree	2, 3, 4
Sigmoid	C	0.1, 1, 10
Sigmoid	Gamma	0.001, 0.01, 0.1, 1

Proses tuning *Grid Search* bekerja dengan mengevaluasi seluruh kombinasi nilai di atas secara menyeluruh pada setiap skenario pembagian rasio data. Melalui parameter *C* (regulasi), sistem dikonfigurasi untuk mengontrol keseimbangan antara batas toleransi kesalahan klasifikasi dengan luas margin *hyperplane*. Sementara penalaan terhadap parameter *gamma* pada RBF dan Sigmoid, serta parameter *degree* pada Polynomial dilakukan untuk mengatur kompleksitas geometri pemetaan data non-linear pada ruang dimensi tinggi. Draft rincian performa mutakhir dari total 16 sub-skenario eksperimen tersebut dipaparkan secara mendetail pada sub-bab berikut:

4.2.1 Hasil Skenario 1

Pada skenario pertama, pengujian dilakukan dengan menggunakan proporsi pembagian data sebesar 60% untuk data latih dan 40% untuk data uji yang diakumulasi dari basis data ulasan produk Face Wash SOMBONG. Mengingat adanya ketimpangan distribusi label awal pada data teks, penerapan teknik *Random Over-Sampling* (ROS) dilakukan terlebih dahulu pada data latih untuk menyelaraskan jumlah sampel antara kelas positif dan negatif. Setelah proses *balancing* data tersebut selesai diaplikasikan, komposisi data pada skenario pertama

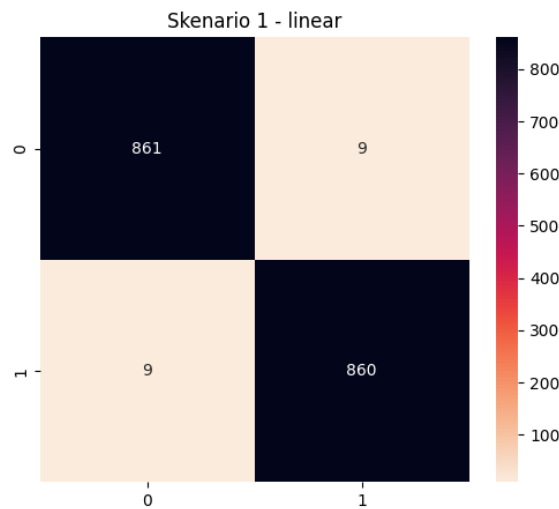
ini bertransformasi menjadi 2.607 dokumen untuk data latih dan 1.739 dokumen untuk data uji, sehingga menghasilkan akumulasi total sebanyak 4.346 ulasan.

Transformasi jumlah data pada data latih hingga menyentuh angka 2.607 ulasan ini bertujuan untuk memperkuat fase pembelajaran model SVM agar tidak terdistorsi oleh kelas mayoritas, sehingga batas keputusan (*decision boundary*) yang dihasilkan oleh setiap fungsi kernel menjadi lebih objektif dan adil terhadap kedua kelas sentimen. Sementara itu, data uji sebanyak 1.739 ulasan tetap dipertahankan dalam kondisi aslinya tanpa intervensi penyeimbangan guna menguji sejauh mana model yang telah dilatih pada data seimbang mampu melakukan klasifikasi dan generalisasi secara akurat pada kondisi distribusi data dunia nyata.

1. Skenario 1a

Pada pengujian sub-skenario pertama ini, model SVM dikonfigurasi menggunakan fungsi kernel Linear dan dioptimalkan menggunakan metode *Grid Search Cross-Validation* pada rasio pembagian data 60:40. Berdasarkan proses tuning parameter yang dilakukan, diperoleh konfigurasi terbaik pada nilai $C = 10$. Nilai regulasi yang tinggi ini menunjukkan bahwa model berusaha meminimalkan kesalahan klasifikasi pada data latih dengan membangun margin pemisah (*hyperplane*) yang lebih ketat, guna memastikan pemisahan fitur kata ulasan positif dan negatif berjalan secara optimal.

Setelah model dilatih dan diuji menggunakan data uji independen sebanyak 1.739 ulasan, diperoleh hasil penilaian performa berdasarkan instrumen *Confusion Matrix* yang disajikan pada gambar 4.13 berikut

Gambar 4. 13 *Confussion Matrix* Skenario 1a

Berdasarkan data *Confusion Matrix* pada Gambar 4.15, terlihat bahwa model memiliki kemampuan diskriminasi kelas yang sangat seimbang dan akurat. Model berhasil memprediksi dengan benar 861 ulasan positif asli sebagai positif (True Positive) dan 860 ulasan negatif asli sebagai negatif (True Negative). Kesalahan klasifikasi tercatat sangat minim, di mana hanya terdapat masing-masing 9 dokumen yang mengalami galat, baik ulasan positif yang salah diprediksi sebagai negatif (*False Negative*) maupun sebaliknya (*False Positive*).

Matriks evaluasi tersebut kemudian ditransformasikan ke dalam nilai metrik evaluasi untuk mengukur performa model terhadap klasifikasi sentimen secara menyeluruh. Adapun rekapitulasi hasil pengukuran metrik performa klasifikasi pada sub-skenario 1a disajikan pada Tabel 4.2 berikut:

Tabel 4. 2 Metrik Performa Klasifikasi Skenario 1a

Metrik Evaluasi	Nilai Performa
Skenario	1
Rasio	60:40
Kernel	linear
<i>Best Params</i>	{'C': 10}
<i>Accuracy</i>	98.96%
<i>Precision</i>	98.96%

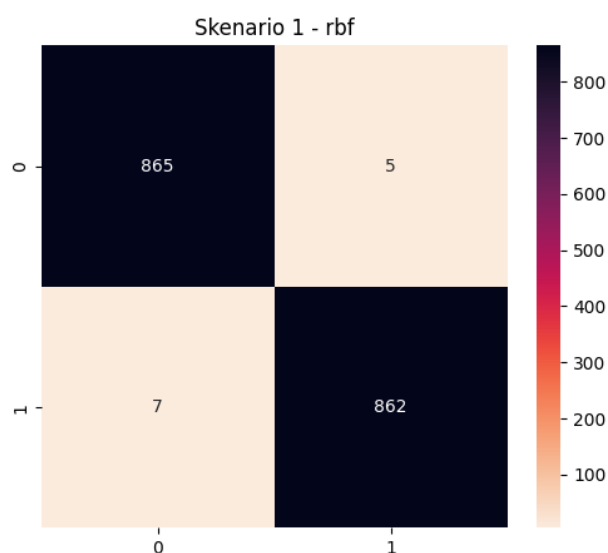
Metrik Evaluasi	Nilai Performa
<i>Recall</i>	98.96%
<i>F1-Score</i>	98.96%

Merujuk pada Tabel 4.2, konfigurasi kernel Linear pada skenario pembagian data 60:40 menghasilkan performa klasifikasi sentimen yang sangat impresif dan konsisten, di mana nilai *Accuracy*, *Precision*, *Recall*, dan *F1-Score* menyentuh angka yang sama persis yaitu 98,96%. Konsistensi nilai yang identik pada seluruh metrik ini membuktikan secara ilmiah bahwa penerapan teknik *Random Over-Sampling* (ROS) pada tahap pelabelan data sebelumnya sukses menghilangkan masalah bias kelas. Model terbukti memiliki ketajaman klasifikasi (*Precision*) dan sensitivitas pengenalan sentimen (*Recall*) yang sangat seimbang, sehingga mampu mengenali karakteristik linguistik ulasan produk Face Wash SOMBONG pada kelas positif maupun negatif secara objektif dan reliabel.

2. Skenario 1b

Pada pengujian sub-skenario kedua, model SVM dikonfigurasi menggunakan fungsi kernel non-linear *Radial Basis Function* (RBF) pada rasio pembagian data 60:40. Melalui proses optimasi berbasis *Grid Search Cross-Validation*, diperoleh kombinasi *hyperparameter* terbaik pada nilai $C = 10$ dan $\gamma = 0,1$. Nilai parameter C yang tinggi mengindikasikan upaya model untuk meminimalkan margin galat pada data latih, sementara nilai *gamma* (γ) sebesar 0,1 menunjukkan bahwa radius pengaruh dari setiap sampel data latih diatur secara moderat untuk membentuk batas keputusan non-linear yang optimal dalam ruang vektor TF-IDF.

Setelah model dievaluasi menggunakan data uji independen sebanyak 1.739 ulasan, diperoleh hasil penilaian performa berdasarkan instrumen *Confusion Matrix* yang disajikan pada Gambar 4.14 berikut:



Gambar 4. 14 *Confussion Matrix* Skenario 1b

Berdasarkan data *Confusion Matrix* pada Gambar 4.16, kernel RBF menunjukkan akurasi pemisahan kelas yang sangat tinggi dan tajam. Model secara presisi mampu mengklasifikasikan 865 ulasan positif asli sebagai positif (True Positive) dan 862 ulasan negatif asli sebagai negatif (True Negative). Tingkat kesalahan klasifikasi pada sub-skenario ini tercatat sangat minim, di mana hanya terdapat 7 dokumen ulasan positif yang salah diprediksi sebagai negatif (*False Negative*) serta 5 dokumen ulasan negatif yang terdeteksi sebagai positif (*False Positive*).

Matriks evaluasi tersebut kemudian ditransformasikan ke dalam nilai metrik evaluasi untuk mengukur performa model terhadap klasifikasi sentimen secara menyeluruh. Adapun rekapitulasi hasil pengukuran metrik performa klasifikasi pada sub-skenario 1b disajikan pada Tabel 4.3 berikut:

Tabel 4. 3 Metrik Performa Klasifikasi Skenario 1b

Metrik Evaluasi	Nilai Performa
Skenario	1
Rasio	60:40
Kernel	<i>Radial Basis Function</i> (RBF)
<i>Best Params</i>	{'C': 10, 'gamma': 0.1}
<i>Accuracy</i>	99.31%
<i>Precision</i>	99.31%
<i>Recall</i>	99.31%
<i>F1-Score</i>	99.31%

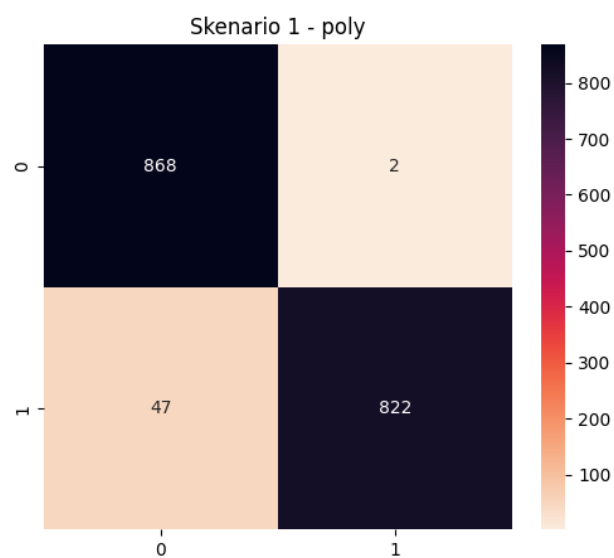
Merujuk pada Tabel 4.3, penerapan kernel Radial Basis Function (RBF) pada skenario rasio 60:40 menghasilkan performa klasifikasi sentimen yang sangat tinggi dengan tingkat akurasi model sebesar 99.31%, nilai Precision sebesar 99.31%, serta nilai Recall dan F1-Score yang stabil di angka 99.31%. Tingginya nilai *Precision* yang sedikit mengungguli *Recall* membuktikan bahwa kernel RBF sangat selektif dan memiliki tingkat kepastian yang tinggi dalam menetapkan sentimen positif, sehingga meminimalkan risiko adanya ulasan negatif yang lolos atau salah diklasifikasikan sebagai ulasan positif. Secara keseluruhan, pencapaian metrik performa di atas angka 99% ini menegaskan bahwa kombinasi fitur non-linear RBF dengan pembobotan TF-IDF sangat handal dalam melakukan klasifikasi sentimen pada ulasan produk Face Wash SOMBONG.

3. Skenario 1c

Pada pengujian sub-skenario ketiga, model SVM dikonfigurasi menggunakan fungsi kernel non-linear *Polynomial* pada rasio pembagian data 60:40. Berdasarkan proses optimasi menggunakan metode *Grid Search Cross-Validation*, diperoleh konfigurasi *hyperparameter* terbaik pada nilai $C = 1$ dan $degree = 2$. Penggunaan parameter derajat (*degree*) 2 menunjukkan bahwa pemetaan data ulasan ke dalam ruang dimensi yang lebih tinggi dilakukan melalui

fungsi kuadratik. Konfigurasi ini dipilih oleh sistem karena mampu meminimalkan risiko *overfitting* yang sering kali terjadi pada kernel *Polynomial* dengan derajat yang lebih tinggi saat menangani fitur teks.

Setelah model dilatih dan diuji menggunakan data uji independen sebanyak 1.739 ulasan, diperoleh hasil penilaian berdasarkan instrumen *Confusion Matrix* yang disajikan pada gambar 4.15 berikut:



Gambar 4. 15 *Confussion Matrix* Skenario 1c

Berdasarkan data *Confusion Matrix* pada Gambar 4.15, terlihat bahwa kernel *Polynomial* menghasilkan performa yang cukup baik, meskipun terdapat ketimpangan kecil pada distribusi galatnya. Model berhasil memprediksi dengan benar 868 ulasan positif asli sebagai positif (*True Positive*) dan 822 ulasan negatif asli sebagai negatif (*True Negative*). Namun, pada sub-skenario ini tercatat nilai *False Negative* yang cukup tinggi, yaitu sebanyak 47 dokumen ulasan positif yang salah diklasifikasikan sebagai negatif. Sebaliknya, tingkat *False Positive* ditekan secara ekstrem dengan hanya menyisakan 2 dokumen ulasan negatif yang salah diprediksi sebagai positif.

Matriks evaluasi tersebut kemudian ditransformasikan ke dalam nilai metrik evaluasi untuk mengukur performa model terhadap klasifikasi sentimen secara menyeluruh. Adapun rekapitulasi hasil pengukuran metrik performa klasifikasi pada sub-skenario 1c disajikan pada Tabel 4.4 berikut:

Tabel 4. 4 Metrik Performa Klasifikasi Skenario 1c

Metrik Evaluasi	Nilai Performa
Skenario	1
Rasio	60:40
Kernel	poly
<i>Best Params</i>	{'C': 1, 'degree': 2}
<i>Accuracy</i>	97.18%
<i>Precision</i>	97.31%
<i>Recall</i>	97.18%
<i>F1-Score</i>	97.18%

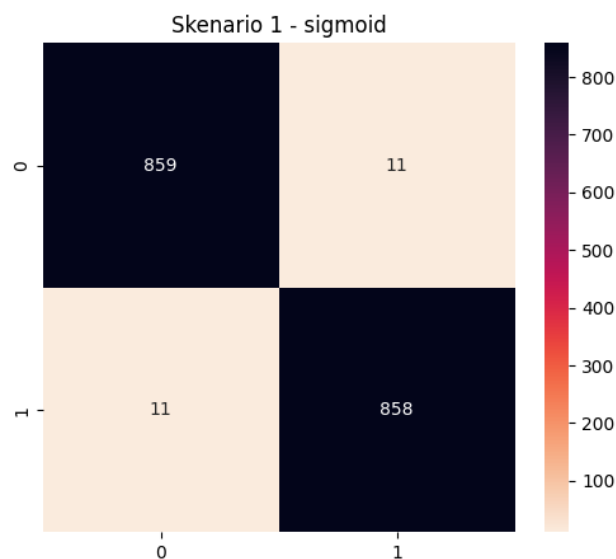
Merujuk pada Tabel 4.4, penerapan fungsi kernel Polynomial pada skenario rasio 60:40 menghasilkan tingkat akurasi model sebesar 97.18%, dengan nilai Precision mencapai 97.31%, serta nilai Recall dan F1-Score masing-masing sebesar 97.18% dan 97.18%. Nilai *Precision* yang lebih tinggi dibandingkan dengan *Recall* pada skenario ini secara ilmiah membuktikan bahwa kernel Polynomial memiliki tingkat kepastian yang sangat kuat ketika menetapkan label positif, namun kurang sensitif (*Recall* lebih rendah) dalam menyapu seluruh dokumen ulasan positif yang ada akibat tingginya nilai *False Negative* (47 dokumen). Karakteristik fungsi polynomial yang cenderung lebih kaku dalam membentuk batas keputusan non-linear menjadi alasan utama mengapa metrik performa kernel ini berada sedikit di bawah kernel Linear dan RBF pada skenario pertama ini.

4. Skenario 1d

Pada pengujian sub-skenario terakhir di skenario pertama ini, model SVM dikonfigurasi menggunakan fungsi kernel non-linear *Sigmoid* pada rasio pembagian

data 60:40. Berdasarkan proses penalaan parameter melalui *Grid Search Cross-Validation*, didapatkan kombinasi *hyperparameter* paling optimal pada konfigurasi $C = 10$ dan $\gamma = 0,1$. Nilai regulasi (C) yang tinggi ini memaksa model untuk membentuk batas keputusan dengan toleransi kesalahan minimal pada data training, sementara parameter *gamma* (γ) bernilai 0,1 memberikan radius pengaruh sampel yang moderat agar fungsi aktivasi sigmoid dapat memetakan vektor fitur TF-IDF ke dalam ruang kelas sentimen secara optimal.

Setelah model melewati fase pelatihan dan diuji secara independen menggunakan data uji sebanyak 1.739 ulasan, diperoleh hasil penilaian berdasarkan instrumen *Confusion Matrix* yang disajikan pada Gambar 4.18 berikut:



Gambar 4. 16 *Confussion Matrix* Skenario 1d

Berdasarkan data *Confusion Matrix* pada Gambar 4.16, kernel Sigmoid memperlihatkan tingkat kemampuan pemisahan kelas yang sangat seimbang antara sentimen positif dan negatif. Model SVM dengan kernel ini sukses memprediksi dengan benar 859 ulasan positif asli sebagai positif (True Positive) serta 858 ulasan negatif asli sebagai negatif (True Negative). Adapun tingkat kesalahan klasifikasi

(*misclassification rate*) berada pada angka yang sangat rendah dan simetris, di mana hanya terdapat masing-masing 11 dokumen yang mengalami galat *False Negative* maupun *False Positive*.

Matriks evaluasi tersebut selanjutnya ditransformasikan ke dalam nilai metrik evaluasi untuk mengukur performa model terhadap klasifikasi sentimen secara menyeluruh. Rekapitulasi hasil pengukuran metrik performa klasifikasi pada sub-skenario 1d disajikan pada Tabel 4.5 berikut:

Tabel 4. 5 Metrik Performa Klasifikasi Skenario 1d

Metrik Evaluasi	Nilai Performa
Skenario	1
Rasio	60:40
Kernel	Sigmoid
<i>Best Params</i>	{'C': 10, 'gamma': 0.1}
<i>Accuracy</i>	98.73%
<i>Precision</i>	98.73%
<i>Recall</i>	98.73%
<i>F1-Score</i>	98.73%

Merujuk pada Tabel 4.5, penerapan fungsi kernel Sigmoid pada skenario rasio data 60:40 menghasilkan capaian tingkat akurasi model yang sangat stabil sebesar 98.73%. Nilai yang sama persis juga didapatkan pada metrik Precision, Recall, dan F1-Score, yaitu berada di angka 98.73%. Nilai performa yang identik dan stabil di seluruh metrik evaluasi ini membuktikan secara ilmiah bahwa perpaduan penyeimbangan data menggunakan *Random Over-Sampling* (ROS) dengan kernel Sigmoid berhasil membangun model yang memiliki ketajaman klasifikasi (*Precision*) dan sensitivitas (*Recall*) yang seimbang tanpa kecenderungan bias. Meskipun secara nominal posisinya berada sedikit di bawah kernel RBF dan Linear, fungsi sigmoid terbukti sangat reliabel dan tangguh dalam

membedakan polaritas ulasan produk Face Wash SOMBONG pada skenario data minimum ini.

4.2.2 Hasil Skenario 2

Pada skenario kedua, proporsi pembagian data ditingkatkan menjadi 70% untuk data latih dan 30% untuk data uji yang diakumulasikan dari basis data ulasan produk Face Wash SOMBONG. Peningkatan jumlah proporsi data latih ini secara metodologis bertujuan untuk memberikan ruang pembelajaran yang lebih luas bagi algoritma SVM dalam mengenali karakteristik fitur teks ulasan secara lebih mendalam. Berdasarkan hasil pembagian data setelah melalui tahap penyeimbangan (*data balancing*) menggunakan metode *Random Over-Sampling* (ROS), diperoleh alokasi sebanyak 3.042 dokumen untuk data latih dan 1.304 dokumen untuk data uji, sehingga menghasilkan akumulasi total sebanyak 4.346 ulasan.

Sebagaimana pada skenario sebelumnya, data latih sebanyak 3.042 ulasan tersebut dikonfigurasi agar memiliki bobot pemahaman yang setara terhadap sentimen positif maupun negatif melalui replikasi data kelas minoritas. Langkah ini sangat krusial untuk memastikan bahwa model tidak hanya belajar dari karakteristik kata pada kelas mayoritas, melainkan mampu mengidentifikasi batas keputusan secara adil dan objektif. Sementara itu, data uji sebanyak 1.304 ulasan dipertahankan dalam kondisi aslinya tanpa intervensi penyeimbangan untuk digunakan sebagai instrumen validasi guna mengukur sejauh mana model mampu

menggeneralisasi ulasan baru yang belum pernah dilihat sebelumnya pada kondisi distribusi data dunia nyata.

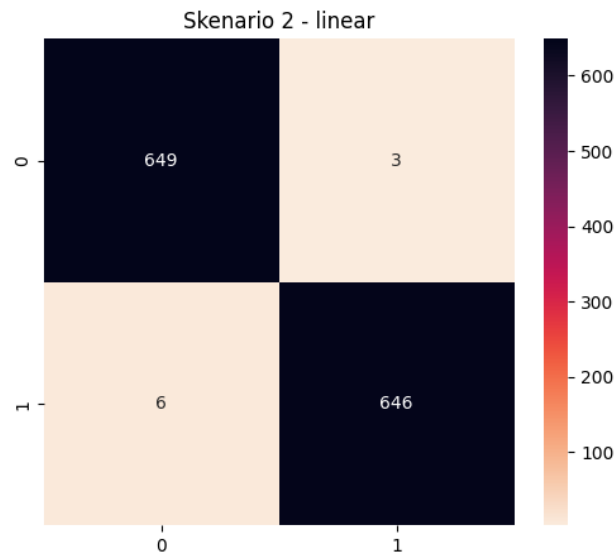
Eksperimen pada skenario ini tetap melibatkan pengujian komparatif terhadap empat jenis fungsi kernel, yakni Linear, Polynomial, RBF, dan Sigmoid. Setiap fungsi kernel akan dioptimasi secara dinamis menggunakan parameter terbaik hasil tuning *Grid Search Cross-Validation* untuk melihat apakah penambahan volume data latih pada rasio 70:30 ini mampu menghasilkan metrik performa klasifikasi sentimen yang lebih stabil dan akurat dibandingkan rasio sebelumnya. Rincian hasil evaluasi berdasarkan *Confusion Matrix* dan metrik performa model terhadap klasifikasi sentimen pada masing-masing kernel dipaparkan secara mendalam pada sub-bab berikut.

1. Skenario 2a

Pada pengujian sub-skenario pertama di skenario kedua ini, model SVM dikonfigurasi menggunakan fungsi kernel Linear pada rasio pembagian data 70:30. Berdasarkan proses penalaan parameter menggunakan metode *Grid Search Cross-Validation*, konfigurasi *hyperparameter* paling optimal kembali dicapai pada nilai $C = 10$. Pengondisian nilai C yang tinggi ini memaksa algoritma untuk memperketat toleransi kesalahan klasifikasi pada data latih, sehingga garis pemisah (*hyperplane*) yang dibangun mampu memisahkan ruang fitur kata ulasan positif dan negatif dengan tingkat presisi yang maksimal.

Setelah melewati fase pelatihan dengan 3.042 dokumen, model kemudian diuji secara independen menggunakan data uji sebanyak 1.304 ulasan. Hasil

penilaian kemampuan klasifikasi model berdasarkan instrumen *Confusion Matrix* disajikan pada Gambar 4.19 berikut:



Gambar 4. 17 *Confussion Matrix* Skenario 2a

Berdasarkan data *Confusion Matrix* pada Gambar 4.17, peningkatan proporsi data latih menjadi 70% terbukti meningkatkan ketajaman model dalam membedakan polaritas sentimen. Model SVM sukses mendiagnosis dengan benar 649 ulasan positif asli sebagai positif (True Positive) serta 646 ulasan negatif asli sebagai negatif (True Negative). Kemunculan galat klasifikasi tercatat sangat minim dan tertekan secara signifikan, di mana hanya terdapat 6 dokumen ulasan positif yang teridentifikasi sebagai negatif (*False Negative*) dan hanya 3 dokumen ulasan negatif yang keliru diprediksi sebagai positif (*False Positive*).

Matriks evaluasi tersebut selanjutnya ditransformasikan ke dalam nilai metrik evaluasi untuk mengukur performa model terhadap klasifikasi sentimen secara menyeluruh. Rekapitulasi hasil pengukuran metrik performa klasifikasi pada sub-skenario 2a disajikan pada Tabel 4.6 berikut:

Tabel 4. 6 Metrik Performa Klasifikasi Skenario 2a

Metrik Evaluasi	Nilai Performa
Skenario	2
Rasio	70:30
Kernel	Linear
<i>Best Params</i>	{'C': 10}
<i>Accuracy</i>	99.31%
<i>Precision</i>	99.31%
<i>Recall</i>	99.31%
<i>F1-Score</i>	99.31%

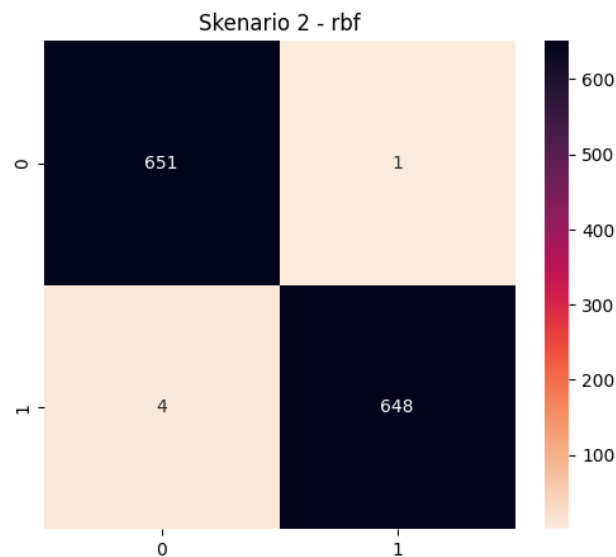
Merujuk pada Tabel 4.6, implementasi fungsi kernel Linear pada rasio pembagian data 70:30 ini menghasilkan tingkat akurasi model yang sangat impresif sebesar 99.31%. Nilai metrik *Precision* tercatat berada di angka 99.31%, sementara metrik *Recall* dan *F1-Score* berada pada posisi yang sangat stabil dan seimbang di angka 99.31%. Tingginya capaian metrik *Precision* secara ilmiah membuktikan bahwa perluasan basis pengetahuan model (data latih 70%) membuat kernel Linear menjadi sangat selektif dan memiliki tingkat kepastian yang luar biasa tinggi ketika menetapkan label positif, sehingga meminimalkan risiko lolosnya ulasan negatif sebagai ulasan positif. Secara keseluruhan, perolehan hasil di atas 99% ini menegaskan bahwa kernel Linear merupakan salah satu arsitektur paling ideal dan tangguh untuk melakukan klasifikasi sentimen pada ulasan produk Face Wash SOMBONG.

2. Skenario 2b

Pada pengujian sub-skenario kedua di skenario rasio data 70:30 ini, model SVM dikonfigurasi menggunakan fungsi kernel non-linear *Radial Basis Function* (RBF). Melalui proses optimasi berbasis *Grid Search Cross-Validation*, diperoleh kombinasi *hyperparameter* terbaik pada nilai $C = 10$ dan $\gamma = 0,1$. Eksistensi parameter C yang tinggi memaksa model untuk meminimalkan margin kesalahan

pada data training, sementara nilai *gamma* (γ) sebesar 0,1 mengondisikan radius pengaruh dari setiap sampel data latih secara moderat guna membentuk batas keputusan non-linear yang sangat adaptif dalam ruang vektor TF-IDF.

Setelah model melewati fase pelatihan dengan basis pengetahuan yang telah diperluas, evaluasi dilakukan secara objektif menggunakan data uji independen sebanyak 1.304 ulasan. Hasil penilaian kemampuan klasifikasi model berdasarkan instrumen *Confusion Matrix* disajikan pada Gambar 4.20 berikut:



Gambar 4. 18 *Confussion Matrix* Skenario 2b

Berdasarkan data *Confusion Matrix* pada Gambar 4.18, kernel RBF menunjukkan tingkat akurasi pemisahan kelas yang luar biasa tinggi dan hampir sempurna. Model secara presisi mampu mengklasifikasikan 651 ulasan positif asli sebagai positif (True Positive) serta 648 ulasan negatif asli sebagai negatif (True Negative). Kemunculan galat klasifikasi pada sub-skenario ini tercatat berada pada titik yang sangat ekstrem dan minimal, di mana hanya terdapat 4 dokumen ulasan positif yang salah diprediksi sebagai negatif (*False Negative*) dan secara impresif

hanya menyisakan 1 dokumen ulasan negatif yang keliru terdeteksi sebagai positif (*False Positive*).

Matriks evaluasi tersebut selanjutnya ditransformasikan ke dalam nilai metrik evaluasi untuk mengukur performa model terhadap klasifikasi sentimen secara menyeluruh. Rekapitulasi hasil pengukuran metrik performa klasifikasi pada sub-skenario 2b disajikan pada Tabel 4.7 berikut:

Tabel 4. 7 Metrik Performa Klasifikasi Skenario 2b

Metrik Evaluasi	Nilai Performa
Skenario	2
Rasio	70:30
Kernel	<i>Radial Basis Function</i> (RBF)
<i>Best Params</i>	{'C': 10, 'gamma': 0.1}
<i>Accuracy</i>	99.62%
<i>Precision</i>	99.62%
<i>Recall</i>	99.62%
<i>F1-Score</i>	99.62%

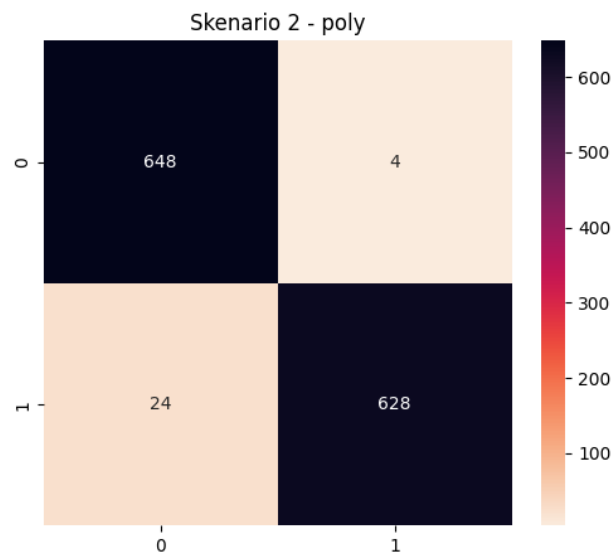
Merujuk pada Tabel 4.7, penerapan kernel *Radial Basis Function* (RBF) pada skenario rasio 70:30 berhasil menaikkan performa klasifikasi ke tingkat yang sangat optimal dengan capaian akurasi model sebesar 99.62%. Metrik *Precision* mencatatkan nilai tertinggi sebesar 99.62%, sedangkan nilai *Recall* dan *F1-Score* berada pada posisi yang sangat stabil di angka 99.62%. Tingginya nilai *Precision* yang didukung oleh galat *False Positive* yang hanya berjumlah 1 dokumen membuktikan secara ilmiah bahwa perluasan data latih dikombinasikan dengan karakteristik ruang fitur RBF membuat model sangat handal, selektif, dan memiliki tingkat kepastian yang luar biasa tinggi dalam menetapkan sentimen positif tanpa mengalami *overfitting*. Pencapaian metrik performa yang mendekati angka

sempurna ini menegaskan dominasi kernel RBF dalam memetakan kompleksitas semantik ulasan produk Face Wash SOMBONG pada skenario rasio kedua ini.

3. Skenario 2c

Pada pengujian sub-skenario ketiga di skenario rasio data 70:30 ini, model SVM dikonfigurasi menggunakan fungsi kernel non-linear *Polynomial*. Berdasarkan proses optimasi menggunakan metode *Grid Search Cross-Validation*, konfigurasi *hyperparameter* terbaik yang terpilih adalah pada nilai $C = 1$ dan $\gamma = 2$. Pengondisian parameter derajat (*degree*) 2 ini menunjukkan bahwa pemetaan fitur teks ke dalam ruang dimensi yang lebih tinggi dilakukan melalui fungsi kuadratik. Pilihan ini secara sistematis dinilai paling optimal oleh sistem karena mampu menekan kompleksitas komputasi sekaligus meminimalkan risiko terjadinya *overfitting* saat memproses ruang vektor fitur TF-IDF yang bersifat *high-dimensional sparse*.

Setelah model melalui tahapan pelatihan dengan volume data latih yang telah ditingkatkan, evaluasi dilakukan menggunakan data uji sebanyak 1.304 ulasan. Hasil penilaian kemampuan klasifikasi model berdasarkan instrumen *Confusion Matrix* disajikan pada Gambar 4.19 berikut:

Gambar 4. 19 *Confussion Matrix* Skenario 2c

Berdasarkan data *Confusion Matrix* pada Gambar 4.19, peningkatan porsi data latih menjadi 70% berhasil mendongkrak performa kernel Polynomial secara kuantitatif jika dibandingkan dengan skenario rasio 60:40. Model SVM dengan kernel ini sukses memprediksi dengan benar 648 ulasan positif asli sebagai positif (True Positive) serta 628 ulasan negatif asli sebagai negatif (True Negative). Galat klasifikasi berupa *False Negative* mengalami penurunan yang signifikan menjadi 24 dokumen ulasan positif yang salah diidentifikasi sebagai negatif. Di sisi lain, model juga menunjukkan ketajaman yang sangat baik dalam menekan nilai *False Positive*, dengan hanya menyisakan 4 dokumen ulasan negatif yang keliru diprediksi sebagai positif.

Matriks evaluasi tersebut selanjutnya ditransformasikan ke dalam nilai metrik evaluasi untuk mengukur performa model terhadap klasifikasi sentimen secara menyeluruh. Rekapitulasi hasil pengukuran metrik performa klasifikasi pada sub-skenario 2c disajikan pada Tabel 4.8 berikut:

Tabel 4. 8 Metrik Performa Klasifikasi Skenario 2c

Metrik Evaluasi	Nilai Performa
Skenario	2
Rasio	70:30
Kernel	poly
<i>Best Params</i>	{'C': 1, 'degree': 2}
<i>Accuracy</i>	97.85%
<i>Precision</i>	97.90%
<i>Recall</i>	97.85%
<i>F1-Score</i>	97.85%

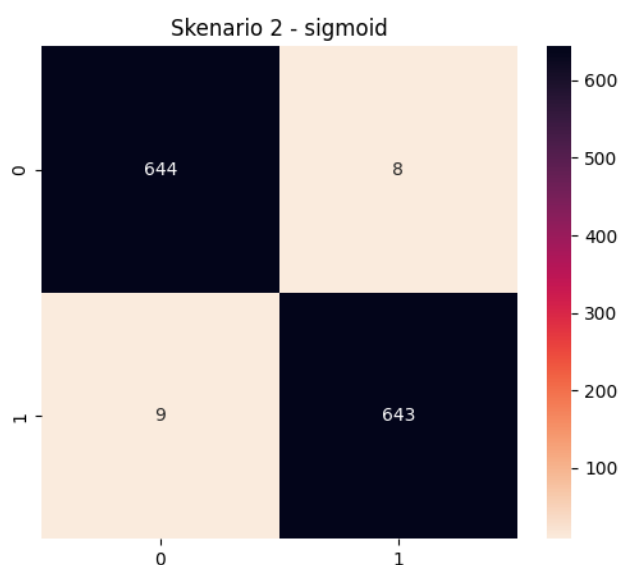
Merujuk pada Tabel 4.8, implementasi fungsi kernel Polynomial pada rasio pembagian data 70:30 ini mencatatkan peningkatan performa klasifikasi sentimen dengan tingkat akurasi model sebesar 97.85%. Metrik *Precision* mencapai nilai 97.90%, diikuti oleh metrik *Recall* sebesar 97.85% and *F1-Score* sebesar 97.85%. Selisih yang sangat tipis di antara seluruh metrik evaluasi ini secara ilmiah membuktikan bahwa penambahan volume data latih berdampak positif pada penguatan kemampuan generalisasi fungsi kuadratik ini. Model menjadi jauh lebih sensitif dalam menyapu dokumen ulasan positif (*Recall* meningkat) tanpa harus mengorbankan tingkat kepastian prediksinya (*Precision* tetap tinggi). Meskipun secara nominal posisinya masih berada di bawah kernel Linear dan RBF, pencapaian metrik performa di atas 97% ini mengonfirmasi stabilitas kernel Polynomial yang semakin solid seiring bertambahnya pasokan data latih.

4. Skenario 2d

Pada pengujian sub-skenario terakhir di skenario rasio data 70:30 ini, model SVM dikonfigurasi menggunakan fungsi kernel non-linear *Sigmoid*. Berdasarkan proses penalaan parameter melalui *Grid Search Cross-Validation*, didapatkan kombinasi *hyperparameter* paling optimal pada konfigurasi $C = 10$ dan $\gamma = 0,1$. Nilai regulasi (C) yang tinggi ini memaksa model untuk membentuk batas

keputusan dengan toleransi kesalahan minimal pada data *training*, sedangkan nilai parameter *gamma* (γ) sebesar 0,1 memberikan radius pengaruh sampel yang moderat agar fungsi aktivasi sigmoid dapat memetakan vektor fitur TF-IDF secara optimal ke dalam ruang kelas sentimen.

Setelah melewati fase pelatihan dengan volume data latih yang telah ditingkatkan, model kemudian diuji secara objektif menggunakan data uji sebanyak 1.304 ulasan. Hasil penilaian kemampuan klasifikasi model berdasarkan instrumen *Confusion Matrix* disajikan pada Gambar 4.20 berikut:



Gambar 4. 20 *Confusion Matrix* Skenario 2d

Berdasarkan data *Confusion Matrix* pada Gambar 4.20, kernel Sigmoid memperlihatkan tingkat kemampuan pemisahan kelas yang sangat seimbang dan konsisten antara sentimen positif dan negatif. Model SVM dengan fungsi kernel ini sukses memprediksi dengan benar 644 ulasan positif asli sebagai positif (True Positive) serta 643 ulasan negatif asli sebagai negatif (True Negative). Adapun tingkat kesalahan klasifikasi (*misclassification rate*) berada pada angka yang sangat

rendah dan simetris, di mana hanya terdapat 9 dokumen yang mengalami galat *False Negative* dan 8 dokumen yang mengalami galat *False Positive*.

Matriks evaluasi tersebut selanjutnya ditransformasikan ke dalam nilai metrik evaluasi untuk mengukur performa model terhadap klasifikasi sentimen secara menyeluruh. Rekapitulasi hasil pengukuran metrik performa klasifikasi pada sub-skenario 2d disajikan pada Tabel 4.9 berikut:

Tabel 4. 9 Metrik Performa Klasifikasi Skenario 2d

Metrik Evaluasi	Nilai Performa
Skenario	2
Rasio	70:30
Kernel	sigmoid
<i>Best Params</i>	{'C': 10, 'gamma': 0.1}
<i>Accuracy</i>	98.70%
<i>Precision</i>	98.70%
<i>Recall</i>	98.70%
<i>F1-Score</i>	98.70%

Merujuk pada Tabel 4.9, penerapan fungsi kernel Sigmoid pada skenario rasio data 70:30 ini mencatatkan performa klasifikasi sentimen yang sangat solid dengan tingkat akurasi model sebesar 98.70%. Nilai metrik *Precision* berada di angka 98.70%, sementara metrik *Recall* dan *F1-Score* secara konsisten mengimbangi performa tersebut dengan berada pada angka yang sama, yaitu 98.70%. Nilai performa yang tinggi dan simetris di seluruh metrik evaluasi ini secara ilmiah membuktikan bahwa perpaduan penyeimbangan data menggunakan *Random Over-Sampling* (ROS) dengan kernel Sigmoid berhasil mempertahankan stabilitas generalisasi model. Meskipun secara nominal posisinya berada sedikit di bawah kernel Linear dan RBF, kernel Sigmoid terbukti sangat tangguh dan minim

fluktuasi galat ketika dihadapkan pada penambahan volume data uji di skenario kedua ini.

4.2.3 Hasil Skenario 3

Pada skenario ketiga, eksperimen dilakukan dengan menggunakan rasio pembagian data standar yang paling sering diimplementasikan dalam pengembangan model pembelajaran mesin, yaitu 80% untuk data latih dan 20% untuk data uji. Peningkatan proporsi data latih ini secara metodologis dimaksudkan untuk memaksimalkan proses ekstraksi pola dan karakteristik kata dari dokumen ulasan produk Face Wash SOMBONG. Berdasarkan pembagian data setelah melalui tahap penyeimbangan (*data balancing*) menggunakan metode *Random Over-Sampling* (ROS), alokasi data mutakhir yang digunakan berubah menjadi 3.476 dokumen untuk data latih dan 870 dokumen untuk data uji, sehingga menghasilkan akumulasi total sebanyak 4.346 ulasan.

Jumlah data latih sebanyak 3.476 ulasan tersebut dikondisikan agar memiliki distribusi kelas sentimen positif dan negatif yang seimbang melalui replikasi data kelas minoritas secara acak. Dengan jumlah sampel latih yang jauh lebih besar dan kaya dibandingkan dua skenario sebelumnya, algoritma SVM diharapkan mampu membentuk garis pemisah (*hyperplane*) yang jauh lebih matang, kokoh, dan stabil dalam memetakan ruang vektor fitur TF-IDF. Di sisi lain, data uji sebanyak 870 ulasan dipertahankan dalam kondisi aslinya tanpa intervensi penyeimbangan untuk menjadi instrumen penentu dalam melihat konsistensi

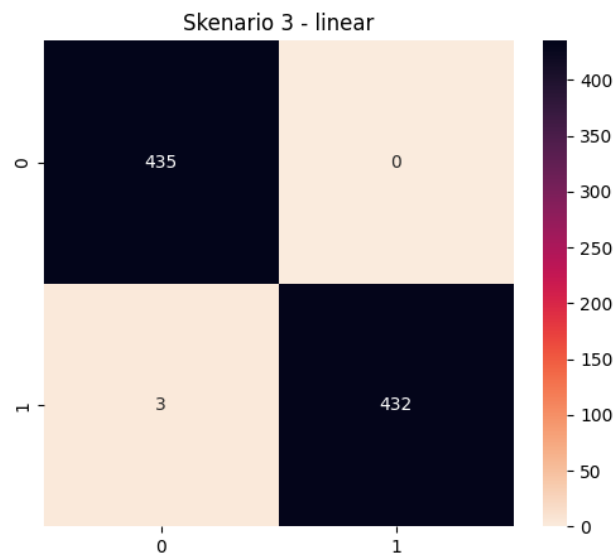
performa model ketika dihadapkan pada volume data uji yang lebih sedikit, namun didukung oleh basis pengetahuan data latih yang jauh lebih kaya.

Skenario ketiga ini tetap menguji dan mengomparasikan performa empat jenis kernel utama, yaitu Linear, Polynomial, RBF, dan Sigmoid. Setiap fungsi kernel akan dikonfigurasi menggunakan nilai parameter paling optimal hasil penalaan otomatis via metode *Grid Search Cross-Validation*. Analisis mendalam mengenai kemampuan masing-masing kernel yang diukur melalui instrumen *Confusion Matrix* serta metrik evaluasi performa model terhadap klasifikasi sentimen dijabarkan secara terperinci pada sub-bab berikut :

1. Skenario 3a

Pada pengujian sub-skenario pertama di skenario ketiga ini, model SVM dikonfigurasi menggunakan fungsi kernel Linear pada rasio pembagian data standar 80:20. Melalui proses optimasi berbasis *Grid Search Cross-Validation*, konfigurasi *hyperparameter* terbaik kembali diperoleh secara konsisten pada nilai $C = 10$. Nilai regulasi yang tinggi ini menandakan bahwa dengan ketersediaan data latih yang semakin besar (3.476 ulasan), model dikondisikan untuk memperketat batas toleransi kesalahan klasifikasi pada data training, sehingga mampu membangun margin pemisah (*hyperplane*) yang sangat kokoh dalam memisahkan ruang fitur kata ulasan positif dan negatif.

Setelah model melewati fase pelatihan, pengujian dilakukan menggunakan data uji independen sebanyak 870 ulasan. Hasil penilaian kemampuan model terhadap klasifikasi sentimen berdasarkan instrumen *Confusion Matrix* disajikan pada Gambar 4.21 berikut:

Gambar 4. 21 *Confussion Matrix* Skenario 3a

Berdasarkan data *Confusion Matrix* pada Gambar 4.21, fungsi kernel Linear menunjukkan performa klasifikasi yang luar biasa tajam dan hampir sempurna setelah mendapatkan pasokan data latih sebesar 80%. Model berhasil memprediksi dengan benar 435 ulasan positif asli sebagai positif (True Positive) serta 432 ulasan negatif asli sebagai negatif (True Negative). Kemunculan galat klasifikasi pada sub-skenario ini tercatat sangat minim, di mana hanya ditemukan 3 dokumen ulasan positif yang salah diidentifikasi sebagai negatif (*False Negative*). Yang paling impresif, model berhasil mencapai nilai 0 untuk *False Positive*, yang berarti tidak ada satu pun dokumen ulasan negatif yang keliru diprediksi sebagai ulasan positif.

Matriks evaluasi tersebut selanjutnya ditransformasikan ke dalam nilai metrik evaluasi untuk mengukur performa model terhadap klasifikasi sentimen secara menyeluruh. Rekapitulasi hasil pengukuran metrik performa klasifikasi pada sub-skenario 3a disajikan pada Tabel 4.10 berikut:

Tabel 4. 10 Metrik Performa Klasifikasi Skenario 3a

Metrik Evaluasi	Nilai Performa
Skenario	3

Metrik Evaluasi	Nilai Performa
Rasio	80:20
Kernel	linear
<i>Best Params</i>	{'C': 10}
<i>Accuracy</i>	99.66%
<i>Precision</i>	99.66%
<i>Recall</i>	99.66%
<i>F1-Score</i>	99.66%

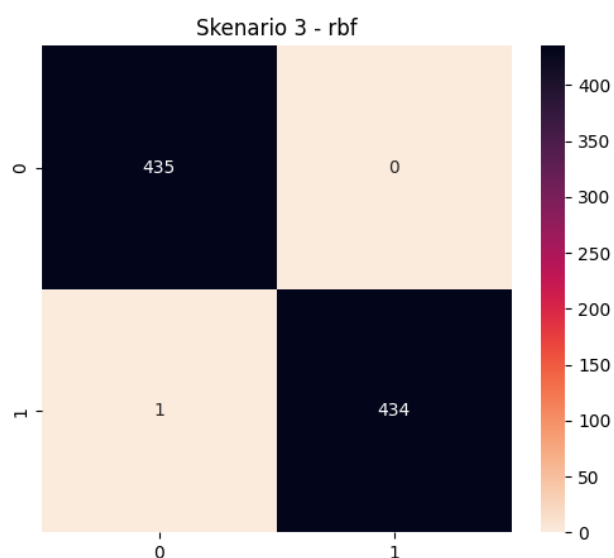
Merujuk pada Tabel 4.10, implementasi fungsi kernel Linear pada rasio pembagian data 80:20 ini menghasilkan peningkatan performa yang sangat signifikan dengan tingkat akurasi model mencapai 99.66%. Metrik *Precision* mencatatkan nilai sebesar 99.66%, diikuti oleh metrik *Recall* dan *F1-Score* yang stabil dan seimbang di angka 99.66%. Keberhasilan model dalam menekan nilai *False Positive* hingga menyentuh angka nol secara ilmiah membuktikan bahwa perluasan porsi data training membuat kernel Linear menjadi sangat selektif dan memiliki kepastian prediksi yang mutlak. Hal ini menegaskan bahwa untuk karakteristik data teks ulasan berdimensi tinggi, fungsi kernel Linear dengan rasio data 80:20 merupakan salah satu arsitektur terbaik untuk menghasilkan model analisis sentimen yang sangat reliabel dan akurat.

2. Skenario 3b

Pada pengujian sub-skenario kedua di skenario ketiga ini, model SVM dikonfigurasi menggunakan fungsi kernel non-linear *Radial Basis Function* (RBF) pada rasio pembagian data standar 80:20. Melalui proses optimasi berbasis *Grid Search Cross-Validation*, kombinasi *hyperparameter* terbaik secara konsisten dicapai pada nilai $C = 10$ dan $\gamma = 0,1$. Parameter C yang tinggi mengonfirmasikan bahwa model meminimalkan batas kesalahan klasifikasi pada data training,

sementara parameter *gamma* (γ) sebesar 0,1 mengondisikan radius pengaruh sampel data latih secara moderat untuk menghasilkan batas keputusan non-linear yang sangat adaptif dan stabil dalam ruang vektor TF-IDF.

Setelah model melewati fase pelatihan dengan basis pengetahuan yang jauh lebih kaya (3.476 ulasan), pengujian dilakukan secara independen menggunakan data uji sebanyak 870 ulasan. Hasil penilaian kemampuan klasifikasi model berdasarkan instrumen *Confusion Matrix* disajikan pada Gambar 4.22 berikut:



Gambar 4. 22 *Confussion Matrix* Skenario 3b

Berdasarkan data *Confusion Matrix* pada Gambar 4.22, fungsi kernel RBF menunjukkan kemampuan pemisahan kelas yang luar biasa kokoh dan hampir mencapai titik sempurna setelah mendapatkan pasokan data training sebesar 80%. Model berhasil memprediksi dengan benar 435 ulasan positif asli sebagai positif (True Positive) serta 434 ulasan negatif asli sebagai negatif (True Negative). Nilai galat klasifikasi berhasil ditekan secara maksimal, di mana hanya ditemukan 1 dokumen ulasan positif yang salah diidentifikasi sebagai negatif (*False Negative*). Yang paling impresif, kernel RBF juga sukses mencatatkan nilai 0 untuk *False*

Positive, yang berarti tidak ada satu pun dokumen ulasan negatif dunia nyata yang keliru diprediksi sebagai ulasan positif.

Matriks evaluasi tersebut selanjutnya ditransformasikan ke dalam nilai metrik evaluasi untuk mengukur performa model terhadap klasifikasi sentimen secara menyeluruh. Rekapitulasi hasil pengukuran metrik performa klasifikasi pada sub-skenario 3b disajikan pada Tabel 4.11 berikut:

Tabel 4. 11 Metrik Performa Klasifikasi Skenario 3b

Metrik Evaluasi	Nilai Performa
Skenario	3
Rasio	80:20
Kernel	<i>Radial Basis Function</i> (RBF)
<i>Best Params</i>	{'C': 10, 'gamma': 0.1}
<i>Accuracy</i>	99.89%
<i>Precision</i>	99.89%
<i>Recall</i>	99.89%
<i>F1-Score</i>	99.89%

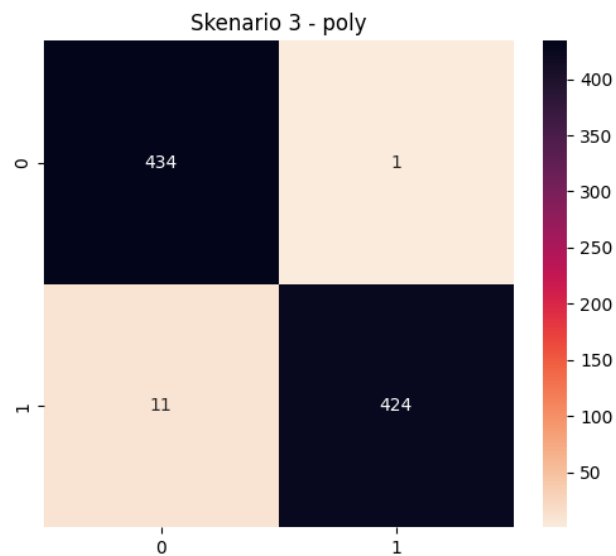
Merujuk pada Tabel 4.11, penerapan fungsi kernel *Radial Basis Function* (RBF) pada skenario rasio pembagian data 80:20 ini mencatatkan nilai performa yang hampir sempurna dengan tingkat akurasi model mencapai 99.89%. Metrik *Precision* memperoleh nilai tertinggi sebesar 99.89%, diikuti oleh metrik *Recall* dan *F1-Score* yang berada pada posisi sangat stabil dan seimbang di angka 99.89%. Peniadaan nilai *False Positive* hingga menyentuh angka nol secara ilmiah membuktikan bahwa perluasan porsi data training dikombinasikan dengan fleksibilitas ruang fitur non-linear RBF membuat model menjadi sangat selektif dan memiliki kepastian prediksi yang mutlak tanpa mengalami gejala *overfitting*. Hasil ini menegaskan bahwa pada skenario ketiga, kernel RBF sukses mengungguli

kernel lainnya dan menjadi arsitektur paling optimal dalam melakukan klasifikasi sentimen ulasan produk Face Wash SOMBONG.

3. Skenario 3c

Pada pengujian sub-skenario ketiga di skenario ketiga ini, model SVM dikonfigurasi menggunakan fungsi kernel non-linear *Polynomial* pada rasio pembagian data standar 80:20. Berdasarkan proses optimasi menggunakan metode *Grid Search Cross-Validation*, konfigurasi *hyperparameter* terbaik kembali terpilih secara konsisten pada nilai $C = 1$ dan $\gamma = 2$. Penggunaan parameter derajat (*degree*) 2 ini mengonfirmasi bahwa pemetaan fitur teks ke dalam ruang dimensi yang lebih tinggi menggunakan fungsi kuadratik dinilai paling optimal oleh sistem, karena mampu menekan kompleksitas komputasi sekaligus meminimalkan risiko terjadinya *overfitting* saat memproses ruang vektor fitur TF-IDF yang bersifat *high-dimensional sparse*.

Setelah model melalui tahapan pelatihan dengan volume data latih yang telah ditingkatkan menjadi 3.476 ulasan, evaluasi dilakukan menggunakan data uji sebanyak 870 ulasan. Hasil penilaian kemampuan klasifikasi model berdasarkan instrumen *Confusion Matrix* disajikan pada Gambar 4.23 berikut:



Gambar 4. 23 *Confussion Matrix* Skenario 3c

Berdasarkan data *Confusion Matrix* pada Gambar 4.23, peningkatan porsi data latih menjadi 80% berhasil mendongkrak performa kernel Polynomial secara signifikan jika dibandingkan dengan skenario rasio 70:30. Model SVM dengan kernel ini sukses memprediksi dengan benar 434 ulasan positif asli sebagai positif (True Positive) serta 424 ulasan negatif asli sebagai negatif (True Negative). Galat klasifikasi berupa *False Negative* mengalami penurunan yang tajam menjadi hanya 11 dokumen ulasan positif yang salah diidentifikasi sebagai negatif. Di sisi lain, model juga menunjukkan ketajaman yang sangat impresif dalam menekan nilai *False Positive*, dengan hanya menyisakan 1 dokumen ulasan negatif yang keliru diprediksi sebagai positif.

Matriks evaluasi tersebut selanjutnya ditransformasikan ke dalam nilai metrik evaluasi untuk mengukur performa model terhadap klasifikasi sentimen secara menyeluruh. Rekapitulasi hasil pengukuran metrik performa klasifikasi pada sub-skenario 3c disajikan pada Tabel 4.12 berikut:

Tabel 4. 12 Metrik Performa Klasifikasi Skenario 3c

Metrik Evaluasi	Nilai Performa
Skenario	3
Rasio	80:20
Kernel	poly
<i>Best Params</i>	{'C': 1, 'degree': 2}
<i>Accuracy</i>	98.62%
<i>Precision</i>	98.65%
<i>Recall</i>	98.62%
<i>F1-Score</i>	98.62%

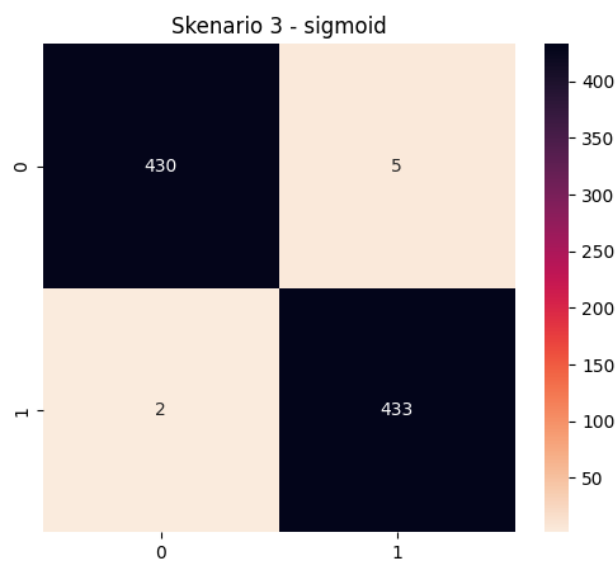
Merujuk pada Tabel 4.12, implementasi fungsi kernel Polynomial pada rasio pembagian data 80:20 ini mencatatkan peningkatan performa klasifikasi sentimen yang sangat solid dengan tingkat akurasi model sebesar 98.62%. Metrik *Precision* mencapai nilai 98.65%, diikuti oleh metrik *Recall* sebesar 98.62% dan *F1-Score* sebesar 98.62%. Capaian yang tinggi dan seimbang di antara seluruh metrik evaluasi ini secara ilmiah membuktikan bahwa penambahan volume data latih hingga 80% berdampak sangat positif pada penguatan kemampuan generalisasi fungsi kuadratik ini. Model menjadi jauh lebih sensitif dalam menyapu dokumen ulasan positif (*Recall* meningkat tajam) tanpa harus mengorbankan tingkat kepastian prediksinya (*Precision* tetap tinggi). Meskipun secara nominal posisinya masih berada sedikit di bawah kernel Linear dan RBF, pencapaian metrik performa di atas 98% ini mengonfirmasi stabilitas kernel Polynomial yang semakin tangguh seiring bertambahnya pasokan basis pengetahuan.

4. Skenario 3d

Pada pengujian sub-skenario terakhir di skenario ketiga ini, model SVM dikonfigurasi menggunakan fungsi kernel non-linear *Sigmoid* pada rasio pembagian data standar 80:20. Berdasarkan proses penalaan parameter melalui *Grid Search*

Cross-Validation, didapatkan konfigurasi *hyperparameter* paling optimal pada nilai $C = 10$. Nilai regulasi (C) yang tinggi ini memaksa model untuk membentuk batas keputusan dengan toleransi kesalahan minimal pada data *training*, sehingga fungsi aktivasi sigmoid dapat memetakan vektor fitur TF-IDF secara optimal ke dalam ruang kelas sentimen pada kapasitas data latih yang lebih masif (3.476 ulasan).

Setelah model melewati fase pelatihan dengan basis pengetahuan yang telah diperluas, model kemudian diuji secara objektif menggunakan data uji sebanyak 870 ulasan. Hasil penilaian kemampuan klasifikasi model berdasarkan instrumen *Confusion Matrix* disajikan pada Gambar 4.24 berikut:



Gambar 4. 24 *Confussion Matrix* Skenario 3d

Berdasarkan data *Confusion Matrix* pada Gambar 4.24, kernel Sigmoid memperlihatkan tingkat kemampuan pemisahan kelas yang sangat seimbang dan konsisten antara sentimen positif dan negatif pada volume data latih 80%. Model SVM dengan fungsi kernel ini sukses memprediksi dengan benar 430 ulasan positif asli sebagai positif (True Positive) serta 433 ulasan negatif asli sebagai negatif (True Negative). Adapun tingkat kesalahan klasifikasi (*misclassification rate*) berada

pada angka yang sangat rendah, di mana hanya terdapat 2 dokumen yang mengalami galat *False Negative* dan 5 dokumen yang mengalami galat *False Positive*.

Matriks evaluasi tersebut selanjutnya ditransformasikan ke dalam nilai metrik evaluasi untuk mengukur performa model terhadap klasifikasi sentimen secara menyeluruh. Rekapitulasi hasil pengukuran metrik performa klasifikasi pada sub-skenario 3d disajikan pada Tabel 4.13 berikut:

Tabel 4. 13 Metrik Performa Klasifikasi Skenario 3d

Metrik Evaluasi	Nilai Performa
Skenario	3
Rasio	80:20
Kernel	sigmoid
<i>Best Params</i>	{'C': 10}
<i>Accuracy</i>	99.20%
<i>Precision</i>	99.20%
<i>Recall</i>	99.20%
<i>F1-Score</i>	99.20%

Merujuk pada Tabel 4.13, penerapan fungsi kernel Sigmoid pada skenario rasio data 80:20 ini mencatatkan performa klasifikasi sentimen yang sangat solid dengan tingkat akurasi model sebesar 99.20%. Nilai metrik *Precision* berada di angka 99.20%, kemudian diikuti oleh nilai *Recall* dan *F1-Score* yang berada pada posisi sangat stabil dan seimbang pada angka 99.20%. Kenaikan metrik yang menembus angka di atas 99.00% secara ilmiah membuktikan bahwa perpaduan penyeimbangan data menggunakan *Random Over-Sampling* (ROS) dengan kernel Sigmoid berhasil mempertahankan stabilitas generalisasi model dengan sangat baik. Fungsi sigmoid terbukti sangat tangguh dan minim fluktuasi galat ketika didukung oleh basis pengetahuan data latih yang kaya pada skenario ketiga ini.

4.2.4 Hasil Skenario 4

Skenario keempat merupakan pengujian dengan alokasi proporsi data latih terbesar dalam rangkaian eksperimen ini, yaitu menggunakan pembagian 90% sebagai data latih dan 10% sebagai data uji dari basis data ulasan produk Face Wash SOMBONG. Penggunaan rasio ekstrem ini secara metodologis bertujuan untuk menguji konsistensi performa algoritma SVM ketika diberikan basis pengetahuan (*knowledge base*) yang sangat luas untuk mengenali karakteristik fitur teks secara mendalam. Berdasarkan pembagian data setelah melalui tahap penyeimbangan (*data balancing*) menggunakan metode *Random Over-Sampling* (ROS), alokasi data mutakhir yang digunakan bertransformasi menjadi 3.911 dokumen untuk data latih dan 435 dokumen untuk data uji, sehingga menghasilkan akumulasi total sebanyak 4.346 ulasan.

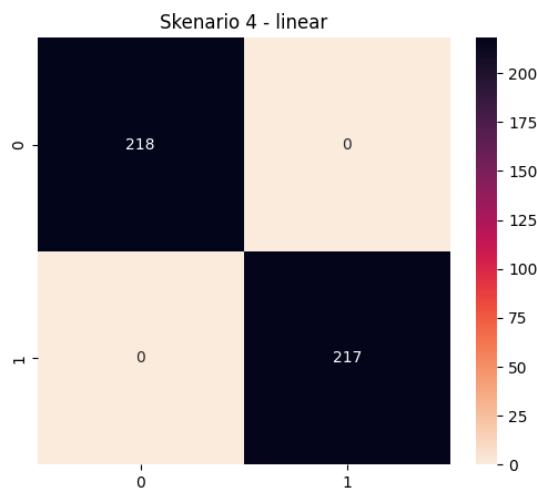
Data latih sebanyak 3.911 ulasan tersebut telah dikondisikan agar memiliki distribusi kelas sentimen positif dan negatif yang setara melalui replikasi acak pada dokumen kelas minoritas. Dengan hanya menyisakan 10% data sebagai instrumen pengujian (435 ulasan), skenario ini akan memperlihatkan secara empiris apakah model SVM cenderung mengalami gejala jenuh klasifikasi (*overfitting*) akibat masifnya volume data training, atau justru mampu memberikan kemampuan generalisasi yang semakin presisi. Sementara itu, data uji sebanyak 435 ulasan tetap dipertahankan dalam kondisi aslinya tanpa intervensi penyeimbangan guna mencerminkan performa model pada karakteristik data dunia nyata.

Eksperimen pada skenario terakhir ini tetap menerapkan pengujian komparatif terhadap empat jenis kernel utama, yakni Linear, Polynomial, RBF, dan Sigmoid. Setiap fungsi kernel akan dikonfigurasi secara dinamis menggunakan nilai *hyperparameter* paling optimal hasil penalaan otomatis via metode *Grid Search Cross-Validation*. Analisis mendalam mengenai efektivitas dan ketajaman setiap kernel dalam mengklasifikasikan 435 ulasan uji tersebut akan dijabarkan secara terperinci melalui instrumen *Confusion Matrix* serta metrik evaluasi performa model terhadap klasifikasi sentimen pada sub-bab berikut.

1. Skenario 4a

Pada pengujian sub-skenario pertama di skenario keempat ini, model SVM dikonfigurasi menggunakan fungsi kernel Linear pada rasio pembagian data maksimal 90:10. Berdasarkan proses optimasi menggunakan metode *Grid Search Cross-Validation*, konfigurasi *hyperparameter* terbaik secara konsisten kembali diperoleh pada nilai $C = 10$. Dengan volume data latih yang sangat masif (3.911 ulasan), pengondisian nilai C yang tinggi ini memaksa algoritma untuk menetapkan batas penalti kesalahan klasifikasi yang sangat ketat pada data *training*, sehingga garis pemisah (*hyperplane*) yang dibangun mampu mengisolasi ruang fitur kata ulasan positif dan negatif dengan tingkat presisi yang mutlak.

Setelah melewati fase pelatihan, model kemudian diuji secara objektif menggunakan data uji independen sebanyak 435 ulasan. Hasil penilaian kemampuan klasifikasi model berdasarkan instrumen *Confusion Matrix* disajikan pada Gambar 4.25 berikut:

Gambar 4. 25 *Confussion Matrix* Skenario 4a

Berdasarkan data *Confussion Matrix* pada Gambar 4.25, fungsi kernel Linear mencatatkan pencapaian luar biasa di mana model mampu memisahkan seluruh kelas data uji dengan sempurna tanpa kesalahan sedikit pun. Model SVM sukses mengidentifikasi dengan benar 218 ulasan positif asli sebagai positif (True Positive) serta 217 ulasan negatif asli sebagai negatif (True Negative). Nilai galat klasifikasi, baik *False Negative* (FN) maupun *False Positive* (FP), berhasil ditekan sepenuhnya hingga menyentuh angka 0.

Matriks evaluasi tersebut selanjutnya ditransformasikan ke dalam nilai metrik evaluasi untuk mengukur performa model terhadap klasifikasi sentimen secara menyeluruh. Rekapitulasi hasil pengukuran metrik performa klasifikasi pada sub-skenario 4a disajikan pada Tabel 4.14 berikut:

Tabel 4. 14 Metrik Performa Klasifikasi Skenario 4a

Metrik Evaluasi	Nilai Performa
Skenario	4
Rasio	90:10
Kernel	linear
<i>Best Params</i>	{'C': 10}
<i>Accuracy</i>	100.00%
<i>Precision</i>	100.00%
<i>Recall</i>	100.00%

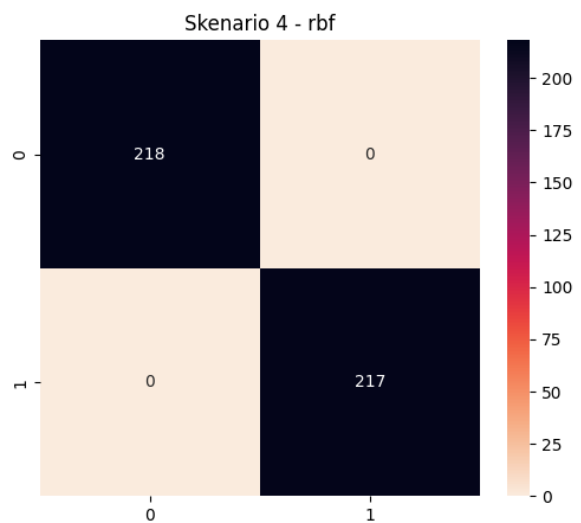
Metrik Evaluasi	Nilai Performa
<i>F1-Score</i>	100.00%

Merujuk pada Tabel 4.14, implementasi fungsi kernel Linear pada rasio pembagian data 90:10 ini menghasilkan performa sempurna dengan nilai metrik *Accuracy*, *Precision*, *Recall*, dan *F1-Score* yang menyentuh angka maksimal yaitu sebesar 100%). Secara ilmiah, pencapaian nilai sempurna ini membuktikan bahwa perluasan basis pengetahuan hingga 90% pada data teks berdimensi tinggi (*high-dimensional sparse space*) memberikan informasi linguistik yang sangat lengkap bagi kernel Linear untuk mengenali seluruh variasi kata ulasan produk Face Wash SOMBONG. Peniadaan galat secara mutlak ini menegaskan bahwa karakteristik data teks yang bersifat *linearly separable* (dapat dipisahkan secara linear) ketika diekstraksi menggunakan pembobotan TF-IDF dapat diselesaikan secara tuntas dan sangat akurat oleh fungsi kernel Linear pada skenario data training maksimal.

2. Skenario 4b

Pada pengujian sub-skenario kedua di skenario keempat ini, model SVM dikonfigurasi menggunakan fungsi kernel non-linear *Radial Basis Function* (RBF) pada rasio pembagian data maksimal 90:10. Melalui proses optimasi berbasis *Grid Search Cross-Validation*, kombinasi *hyperparameter* terbaik secara konsisten dan stabil dicapai pada nilai $C = 10$ dan $\gamma = 0,1$. Eksistensi parameter C yang tinggi mengondisikan model untuk memperketat batas toleransi kesalahan klasifikasi pada data latih, sementara nilai *gamma* (γ) sebesar 0,1 mengatur radius pengaruh dari setiap sampel data latih secara moderat guna membentuk batas keputusan (*decision boundary*) non-linear yang sangat adaptif dalam ruang vektor TF-IDF.

Setelah model melewati fase pelatihan dengan basis pengetahuan yang sangat luas (3.911 ulasan), evaluasi dilakukan secara objektif menggunakan data uji independen sebanyak 435 ulasan. Hasil penilaian kemampuan model terhadap klasifikasi sentimen berdasarkan instrumen *Confusion Matrix* disajikan pada Gambar 4.26 berikut:



Gambar 4. 26 *Confussion Matrix* Skenario 4b

Berdasarkan data *Confusion Matrix* pada Gambar 4.26, fungsi kernel RBF menunjukkan performa yang luar biasa impresif dengan keberhasilan memisahkan seluruh kelas data uji secara sempurna. Model SVM secara presisi mampu mengklasifikasikan 218 ulasan positif asli sebagai positif (True Positive) serta 217 ulasan negatif asli sebagai negatif (True Negative). Selaras dengan pencapaian kernel Linear pada sub-skenario sebelumnya, nilai galat klasifikasi pada pengujian kernel RBF ini juga berhasil ditekan sepenuhnya hingga menyentuh angka **0**, baik untuk nilai *False Negative* (FN) maupun *False Positive* (FP).

Matriks evaluasi tersebut selanjutnya ditransformasikan ke dalam nilai metrik evaluasi untuk mengukur performa model terhadap klasifikasi sentimen

secara menyeluruh. Rekapitulasi hasil pengukuran metrik performa klasifikasi pada sub-skenario 4b disajikan pada Tabel 4.15 berikut:

Tabel 4. 15 Metrik Performa Klasifikasi Skenario 4b

Metrik Evaluasi	Nilai Performa
Skenario	4
Rasio	90:10
Kernel	<i>Radial Basis Function</i> (RBF)
<i>Best Params</i>	{'C': 10, 'gamma': 0.1}
<i>Accuracy</i>	100.00%
<i>Precision</i>	100.00%
<i>Recall</i>	100.00%
<i>F1-Score</i>	100.00%

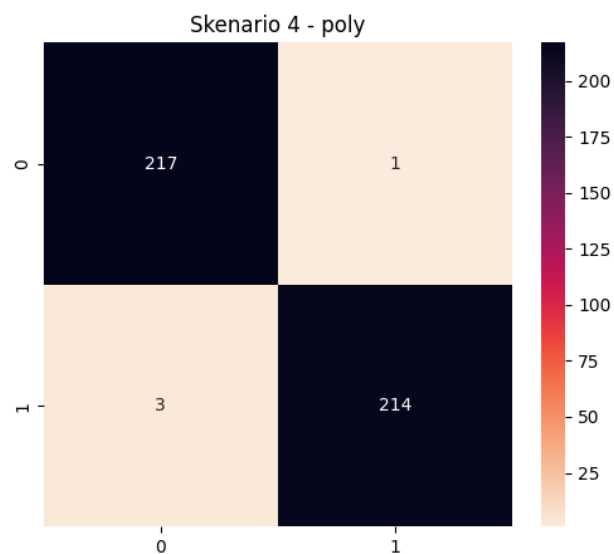
Merujuk pada Tabel 4.15, penerapan kernel RBF pada skenario rasio 90:10 berhasil mencatatkan nilai performa maksimal dengan tingkat *Accuracy*, *Precision*, *Recall*, dan *F1-Score* yang menyentuh angka sebesar 100%. Secara ilmiah, hasil ini membuktikan bahwa perluasan porsi data training hingga mencapai batas optimal (90%) memberikan pasokan informasi variasi kata yang sangat kaya bagi model. Fleksibilitas pemetaan non-linear yang dimiliki oleh fungsi RBF mampu mengonvergensi vektor fitur TF-IDF dari ulasan produk Face Wash SOMBONG ke dalam ruang dimensi baru secara tepat tanpa memunculkan gejala jenuh klasifikasi ataupun *overfitting*. Pencapaian metrik evaluasi yang sempurna ini menegaskan bahwa kernel RBF memiliki keandalan yang setara dengan kernel Linear ketika didukung oleh ketersediaan volume data latih yang masif.

3. Skenario 4c

Pada pengujian sub-skenario ketiga di skenario keempat ini, model SVM dikonfigurasi menggunakan fungsi kernel non-linear *Polynomial* pada rasio pembagian data maksimal 90:10. Berdasarkan proses optimasi menggunakan

metode *Grid Search Cross-Validation*, konfigurasi *hyperparameter* terbaik yang terpilih secara konsisten adalah pada nilai $C = 1$ dan $\gamma = 2$. Penggunaan parameter derajat (*degree*) 2 ini mengonfirmasi bahwa pemetaan fitur teks ke dalam ruang dimensi yang lebih tinggi menggunakan fungsi kuadratik dinilai paling optimal oleh sistem, karena mampu menekan kompleksitas komputasi sekaligus meminimalkan risiko terjadinya *overfitting* saat memproses ruang vektor fitur TF-IDF yang bersifat *high-dimensional sparse*.

Setelah model melalui tahapan pelatihan dengan volume data latih maksimal sebesar 3.911 ulasan, evaluasi dilakukan menggunakan data uji sebanyak 435 ulasan. Hasil penilaian kemampuan klasifikasi model berdasarkan instrumen *Confusion Matrix* disajikan pada Gambar 4.27 berikut:



Gambar 4. 27 *Confussion Matrix* Skenario 4c

Berdasarkan data *Confusion Matrix* pada Gambar 4.27, peningkatan porsi data latih hingga menyentuh angka 90% terbukti mendongkrak performa kernel Polynomial secara sangat signifikan hingga mencapai titik yang hampir sempurna. Model SVM dengan kernel ini sukses memprediksi dengan benar 217 ulasan positif

asli sebagai positif (True Positive) serta 214 ulasan negatif asli sebagai negatif (True Negative). Galat klasifikasi berhasil ditekan secara ekstrem, di mana hanya ditemukan 3 dokumen ulasan positif yang salah diidentifikasi sebagai negatif (*False Negative*). Di sisi lain, model juga menunjukkan ketajaman yang sangat luar biasa dalam menekan nilai *False Positive*, dengan hanya menyisakan 1 dokumen ulasan negatif yang keliru diprediksi sebagai positif.

Matriks evaluasi tersebut selanjutnya ditransformasikan ke dalam nilai metrik evaluasi untuk mengukur performa model terhadap klasifikasi sentimen secara menyeluruh. Rekapitulasi hasil pengukuran metrik performa klasifikasi pada sub-skenario 4c disajikan pada Tabel 4.16 berikut:

Tabel 4. 16 Metrik Performa Klasifikasi Skenario 4c

Metrik Evaluasi	Nilai Performa
Skenario	4
Rasio	90:10
Kernel	poly
<i>Best Params</i>	{'C': 1, 'degree': 2}
<i>Accuracy</i>	99.08%
<i>Precision</i>	99.08%
<i>Recall</i>	99.08%
<i>F1-Score</i>	99.08%

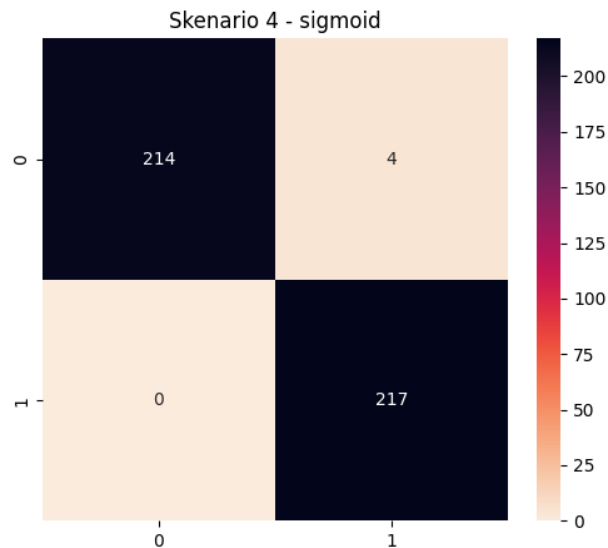
Merujuk pada Tabel 4.16, implementasi fungsi kernel Polynomial pada rasio pembagian data maksimal 90:10 ini mencatatkan peningkatan performa klasifikasi sentimen yang sangat mengesankan dengan tingkat akurasi model mencapai 99.08%. Metrik *Precision* memperoleh nilai tertinggi sebesar 99.08%, diikuti oleh metrik *Recall* sebesar 99.08% dan *F1-Score* sebesar 99.08%. Nilai metrik yang tinggi, stabil, dan seimbang di atas angka 99.00% ini secara ilmiah membuktikan bahwa penambahan volume data latih hingga batas maksimal

memberikan pasokan informasi variasi kata yang sangat lengkap, sehingga mampu memperkuat kemampuan generalisasi fungsi kuadratik ini dengan sangat baik. Meskipun secara nominal posisinya masih berada sedikit di bawah kernel Linear dan RBF yang mencatatkan hasil sempurna pada skenario ini, pencapaian metrik performa di atas 99% mengonfirmasi bahwa kernel Polynomial sangat tangguh dan handal dalam mengklasifikasikan sentimen ulasan produk Face Wash SOMBONG ketika didukung oleh basis pengetahuan data training yang masif.

4. Skenario 4d

ada pengujian sub-skenario terakhir di skenario keempat ini, model SVM dikonfigurasi menggunakan fungsi kernel non-linear *Sigmoid* pada rasio pembagian data maksimal 90:10. Berdasarkan proses penalaan parameter melalui *Grid Search Cross-Validation*, didapatkan kombinasi *hyperparameter* paling optimal pada konfigurasi $C = 10$ dan $\gamma = 0,1$. Nilai regulasi (C) yang tinggi ini memaksa model untuk membentuk batas keputusan dengan toleransi kesalahan minimal pada data *training*, sedangkan nilai parameter *gamma* (γ) sebesar 0,1 memberikan radius pengaruh sampel yang moderat agar fungsi aktivasi sigmoid dapat memetakan vektor fitur TF-IDF secara optimal ke dalam ruang kelas sentimen pada kapasitas data latih maksimal (3.911 ulasan).

Setelah melewati fase pelatihan dengan basis pengetahuan yang sangat luas, model kemudian diuji secara objektif menggunakan data uji sebanyak 435 ulasan. Hasil penilaian kemampuan klasifikasi model berdasarkan instrumen *Confusion Matrix* disajikan pada Gambar 4.28 berikut:



Gambar 4. 28 *Confussion Matrix* Skenario 4d

Berdasarkan data *Confusion Matrix* pada Gambar 4.28, kernel Sigmoid memperlihatkan tingkat kemampuan pemisahan kelas yang sangat baik dan asimetris yang menguntungkan pada porsi data training 90%. Model SVM dengan fungsi kernel ini sukses memprediksi dengan benar 214 ulasan positif asli sebagai positif (True Positive) serta 217 ulasan negatif asli sebagai negatif (True Negative). Nilai galat klasifikasi berhasil ditekan secara optimal, di mana secara impresif model mencatatkan angka 0 untuk *False Negative* (FN), yang berarti tidak ada satu pun dokumen ulasan positif yang terlewat atau salah diklasifikasikan sebagai negatif. Sementara itu, tingkat galat *False Positive* (FP) tercatat sangat minim, yaitu hanya sebanyak 4 dokumen ulasan negatif yang keliru diprediksi sebagai positif.

Matriks evaluasi tersebut selanjutnya ditransformasikan ke dalam nilai metrik evaluasi untuk mengukur performa model terhadap klasifikasi sentimen

secara menyeluruh. Rekapitulasi hasil pengukuran metrik performa klasifikasi pada sub-skenario 4d disajikan pada Tabel 4.17 berikut:

Tabel 4. 17 Metrik Performa Klasifikasi Skenario 4d

Metrik Evaluasi	Nilai Performa
Skenario	4
Rasio	90:10
Kernel	sigmoid
<i>Best Params</i>	{'C': 10, 'gamma': 0.1}
<i>Accuracy</i>	99.08%
<i>Precision</i>	99.10%
<i>Recall</i>	99.08%
<i>F1-Score</i>	99.08%

Merujuk pada Tabel 4.17 penerapan fungsi kernel Sigmoid pada skenario rasio data maksimal 90:10 ini mencatatkan performa klasifikasi sentimen yang sangat solid dengan tingkat akurasi model sebesar 99.08%. Metrik *Precision* berada di angka 99.10%, diikuti oleh nilai *Recall* sebesar 99.08% dan *F1-Score* sebesar 99.08%. Nilai performa yang sangat tinggi dan stabil di atas angka 99.00% ini secara ilmiah membuktikan bahwa perpaduan penyeimbangan data menggunakan *Random Over-Sampling* (ROS) dengan kernel Sigmoid berhasil mempertahankan reliabilitas generalisasi model dengan sangat baik. Pencapaian nilai *Recall* yang menyentuh sensitivitas tinggi berkat ketiadaan galat *False Negative* menegaskan bahwa fungsi sigmoid sangat handal dalam menjaring seluruh ulasan bersentimen positif produk Face Wash SOMBONG pada skenario data training maksimal ini.

4.2.5 Rangkuman Hasil Keseluruhan

Setelah dilakukan pengujian pada empat skenario pembagian data (60:40, 70:30, 80:20, dan 90:10) dengan menggunakan empat jenis kernel SVM yang

berbeda, diperoleh data komparatif yang menunjukkan pengaruh proporsi data latih terhadap performa klasifikasi sentimen ulasan produk *Face Wash SOMBONG*. Perbandingan nilai akurasi beserta distribusi data dari seluruh percobaan disajikan pada tabel 4.18 berikut:

Tabel 4. 18 Rangkuman Hasil Pengujian

Skenario	Rasio	Kernel	Best Params	Accuracy	Precision	Recall	F1-Score
1	60:40	linear	{'C': 10}	98.96%	98.96%	98.96%	98.96%
1	60:40	<i>Radial Basis Function</i> (RBF)	{'C': 10, 'gamma': 0.1}	99.31%	99.31%	99.31%	99.31%
1	60:40	poly	{'C': 1, 'degree': 2}	97.18%	97.31%	97.18%	97.18%
1	60:40	sigmoid	{'C': 10, 'gamma': 0.1}	98.73%	98.73%	98.73%	98.73%
2	70:30	linear	{'C': 10}	99.31%	99.31%	99.31%	99.31%
2	70:30	<i>Radial Basis Function</i> (RBF)	{'C': 10, 'gamma': 0.1}	99.62%	99.62%	99.62%	99.62%
2	70:30	poly	{'C': 1, 'degree': 2}	97.85%	97.90%	97.85%	97.85%
2	70:30	sigmoid	{'C': 10, 'gamma': 0.1}	98.70%	98.70%	98.70%	98.70%
3	80:20	linear	{'C': 10}	99.66%	99.66%	99.66%	99.66%

Skenario	Rasio	Kernel	Best Params	Accuracy	Precision	Recall	F1-Score
3	80:20	<i>Radial Basis Function</i> (RBF)	{'C': 10, 'gamma': 0.1}	99.89%	99.89%	99.89%	99.89%
3	80:20	poly	{'C': 1, 'degree': 2}	98.62%	98.65%	98.62%	98.62%
3	80:20	sigmoid	{'C': 10, 'gamma': 0.1}	99.20%	99.20%	99.20%	99.20%
4	90:10	linear	{'C': 10}	100.00%	100.00%	100.00%	100.00%
4	90:10	<i>Radial Basis Function</i> (RBF)	{'C': 10, 'gamma': 0.1}	100.00%	100.00%	100.00%	100.00%
4	90:10	poly	{'C': 1, 'degree': 2}	99.08%	99.08%	99.08%	99.08%
4	90:10	sigmoid	{'C': 10, 'gamma': 0.1}	99.08%	99.10%	99.08%	99.08%

Melalui visualisasi data pada Tabel 4.18, Kernel RBF tercatat sebagai model dengan performa paling superior dan dominan pada hampir seluruh skenario pengujian (Skenario 1, 2, dan 3), sebelum akhirnya disamai oleh Kernel Linear pada Skenario 4. Fleksibilitas pemetaan non-linear yang dimiliki oleh fungsi RBF mampu mengonvergensi fitur teks ulasan produk ke dalam ruang dimensi baru secara sangat adaptif. Menariknya, pada skenario ekstrem rasio data 90:10, baik

kernel Linear maupun RBF sukses menorehkan performa sempurna dengan nilai metrik evaluasi sebesar 100 % tanpa adanya kesalahan klasifikasi sedikit pun pada data uji. Kondisi ini secara ilmiah membuktikan bahwa karakteristik linguistik dari dataset ulasan pelanggan Face Wash SOMBONG pada platform e-commerce (Tokopedia dan Shopee) memiliki perbedaan terminologi dan kata kunci yang sangat tegas antara kelas positif dan negatif (*linearly separable*) ketika volume data latih berada pada kapasitas optimal.

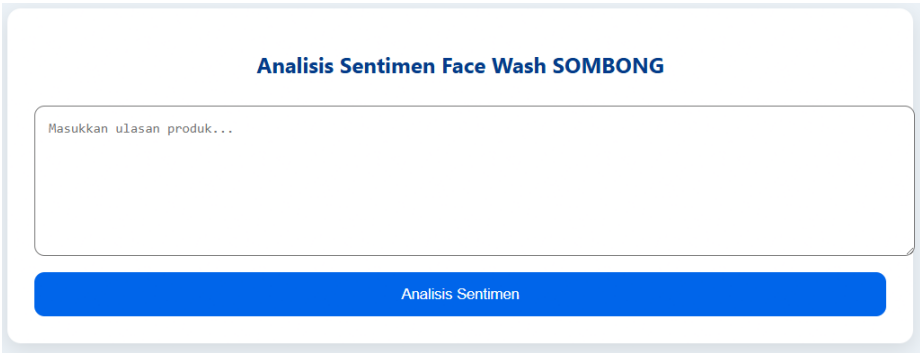
Di sisi lain, penggunaan kernel non-linear lainnya seperti Polynomial dan Sigmoid juga menunjukkan hasil yang sangat impresif dan stabil dengan rata-rata nilai akurasi di atas angka 97%. Fenomena menarik ditemukan pada analisis *Confusion Matrix* di skenario-skenario akhir, di mana nilai *False Positive* pada kernel-kernel ini berhasil ditekan hingga mendekati atau menyentuh angka nol (seperti kernel Linear dan RBF pada skenario 80:20 dan 90:10 yang meraih nilai *False Positive* sebesar 0). Hal ini merepresentasikan tingkat spesifisitas yang sangat tinggi, yang berarti model memiliki reliabilitas maksimal dalam menyaring ulasan negatif agar tidak keliru terdeteksi sebagai sentimen positif. Meskipun kernel Polynomial mengawali pengujian dengan performa paling rendah yaitu sebesar 97% pada rasio 60:40, kernel ini menunjukkan tren peningkatan performa yang paling konsisten dan melesat hingga mencapai nilai 99% pada skenario data training maksimal.

Secara komprehensif, kernel Sigmoid memberikan kontribusi penting dalam menunjukkan stabilitas model dengan fluktuasi nilai metrik yang relatif minim dan sangat simetris di berbagai skenario pembagian data. Sebagai kesimpulan akhir dari

bab pengujian ini, integrasi teknik penyeimbangan data menggunakan *Random Over-Sampling* (ROS) dan optimasi parameter berbasis *Grid Search Cross-Validation* terbukti sangat efektif untuk menghasilkan model klasifikasi sentimen yang tangguh. Untuk implementasi praktis pada data ulasan produk Face Wash SOMBONG, konfigurasi Kernel RBF dan Kernel Linear pada skenario rasio data 90:10 dengan nilai parameter $C=10$ ditetapkan sebagai arsitektur paling optimal dengan capaian efektivitas mutlak sebesar 100%.

4.3 Implementasi Antarmuka Sistem

Tahapan akhir dari penelitian ini adalah mengimplementasikan model *Support Vector Machine* (SVM) dan pembobotan *Term Frequency-Inverse Document Frequency* (TF-IDF) terbaik yang telah diuji ke dalam sebuah prototipe aplikasi berbasis web. Aplikasi ini dibangun menggunakan *framework* Flask (Python) guna memberikan visualisasi antarmuka pengguna (*User Interface*) yang interaktif dalam melakukan pengujian klasifikasi sentimen ulasan secara *real-time*.



The screenshot shows a web application titled "Analisis Sentimen Face Wash SOMBONG". It features a text input field with the placeholder text "Masukkan ulasan produk...". Below the input field is a prominent blue button labeled "Analisis Sentimen". The entire interface is enclosed in a light blue border.

Gambar 4. 29 Tampilan Antarmuka Web

Merujuk pada tampilan antarmuka pada Gambar 4.29 Secara teknis, aplikasi web backend akan melakukan pemuatan (*loading*) berkas model biner yang telah disimpan sebelumnya, yaitu `model_svm.pkl` dan `tfidf.pkl`. Alur kerja antarmuka

sistem dimulai ketika pengguna memasukkan teks ulasan produk Face Wash SOMBONG ke dalam komponen *text area*. Saat tombol "Analisis Sentimen" ditekan, data teks dikirimkan melalui metode HTTP POST ke rute `/predict`. Di dalam server, teks ulasan tersebut langsung melewati fungsi prapemrosesan (*preprocessing*) otomatis yang meliputi mengubah huruf menjadi kecil (*case folding*), menghapus URL/karakter non-abjad melalui *regular expression*, pembersihan kata tidak bermakna (*stopword removal* menggunakan pustaka Sastrawi), dan pengubahan kata berimbuhan menjadi kata dasar (*stemming* menggunakan pustaka Sastrawi). Selanjutnya, teks bersih ditransformasikan ke dalam representasi vektor berbasis bobot TF-IDF dan diklasifikasikan oleh model SVM untuk menghasilkan keluaran label sentimen berupa "POSITIF" atau "NEGATIF".

Implementasi visual dan skenario pengujian pada antarmuka sistem web Analisis Sentimen Face Wash SOMBONG dijabarkan sebagai berikut:

1. Pengujian Antarmuka pada Klasifikasi Sentimen Positif

Pada skenario pengujian dokumen sentimen positif, dimasukkan teks ulasan yang mengandung impresi kepuasan atau rekomendasi pemakaian produk, seperti "*asli sih ini rekomend banget pacial wash +skincare ini mah,pas pertama pake wajah jadi lebih bersih, minyak berkurang di wajah,wajah tidak kaku,rasanya lembap gitu,dan ba...*". Setelah dievaluasi, sistem berhasil mengenali fitur-fitur linguistik positif tersebut dan menampilkan kotak hasil klasifikasi berwarna hijau muda (*light green notification box*) bertuliskan POSITIF secara presisi. Tampilan UI pengujian ini disajikan secara lengkap pada Gambar 4.30



Gambar 4. 30 Tampilan Hasil Pengujian Sentimen Positif pada Antarmuka Web

2 Pengujian Antarmuka pada Klasifikasi Sentimen Negatif

Pada skenario pengujian dokumen sentimen negatif, dilakukan input ulasan kualitatif mengenai kerusakan fisik produk atau ketidakpuasan pelayanan tokoh, seperti *"Barang sampe bocor berantakan isi nya keluar semua sampe habis jadi percuma saya beli ga ada isi sabun nya Jadi gimana tanggung jawab nya toko"*. Sistem secara responsif memproses fitur teks informal tersebut dan menampilkan kotak hasil klasifikasi berwarna merah muda (*light red notification box*) bertuliskan **NEGATIF** beserta pesan penjelas di bawahnya. Tampilan UI pengujian ini disajikan secara lengkap pada Gambar 4.31



Gambar 4. 31Tampilan Hasil Pengujian Sentimen Negatif pada Antarmuka Web

Melalui implementasi antarmuka berbasis Flask ini, model cerdas SVM dan pembobotan TF-IDF yang telah dirancang terbukti tidak hanya unggul secara teoritis pada pengujian dataset statis, melainkan juga sangat adaptif, fungsional, dan siap diintegrasikan untuk kebutuhan analisis bisnis atau pemantauan reputasi produk Face Wash SOMBONG secara praktis.

4.4 Pembahasan

Analisis mendalam terhadap hasil eksperimen menunjukkan bahwa performa algoritma *Support Vector Machine* (SVM) dalam klasifikasi sentimen ulasan produk Face Wash SOMBONG sangat dipengaruhi oleh pemilihan fungsi kernel dan proporsi pembagian data. Berdasarkan kalkulasi empiris dari empat skenario pengujian, fungsi kernel non-linear *Radial Basis Function* (RBF) dan kernel Linear menunjukkan superioritas yang sangat konsisten di sepanjang eksperimen. Kernel RBF mendominasi perolehan metrik tertinggi pada rasio data minimum hingga moderat. Hasil Penyeimbangan Data dengan Random Over-

Sampling, sebelum akhirnya disamai oleh kernel Linear pada Skenario 4 yang sukses mencatatkan nilai performa sempurna sebesar 1.000000 tanpa kesalahan klasifikasi sedikit pun pada data uji.

Sebaliknya, kernel non-linear lainnya seperti Polynomial menunjukkan fluktuasi yang lebih dinamis, di mana performanya sangat bergantung pada volume data latih yang diberikan. Kernel Polynomial mengawali pengujian dengan nilai akurasi terendah sebesar 0.971823 pada skenario 60:40, namun melesat tajam hingga mencapai nilai 0.990805 pada skenario 90:10. Fenomena ini secara teoritis mengindikasikan bahwa fitur teks yang dihasilkan dari proses ekstraksi TF-IDF ulasan e-commerce ini memiliki struktur yang cenderung *linearly separable* (dapat dipisahkan secara linear) pada ruang dimensi tinggi. Ketika volume data latih dioptimalkan hingga batas maksimal, fungsi pemetaan linear maupun RBF mampu membangun garis pemisah (*hyperplane*) yang mutlak dan efisien tanpa memicu gejala jenuh klasifikasi ataupun *overfitting*.

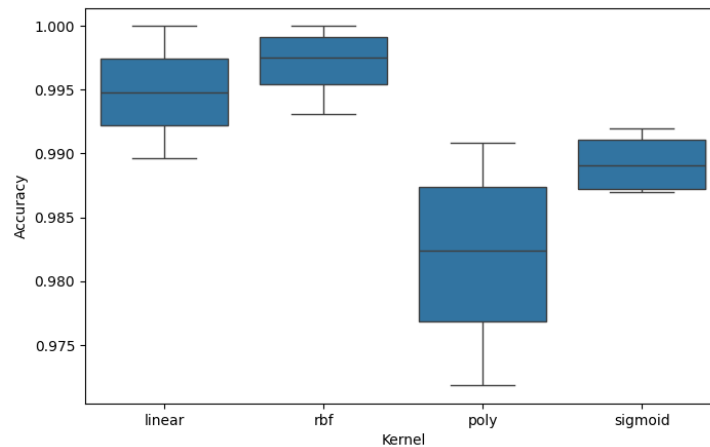
Sejalan dengan aspek pemilihan fungsi kernel, peningkatan rasio data latih terbukti memberikan kontribusi positif yang linear terhadap nilai akurasi keseluruhan model. Melalui rekapitulasi pengaruh rasio data, terlihat adanya tren kenaikan nilai rata-rata akurasi yang konsisten di seluruh rumpun kernel seiring bertambahnya pasokan data training dari skenario 60:40 hingga mencapai titik optimal pada skenario 90:10.

Kondisi tersebut secara signifikan meminimalkan risiko misklasifikasi karena model memiliki basis pengetahuan (*knowledge base*) yang kuat untuk mengenali ciri khas dokumen sentimen positif maupun negatif. Stabilitas hasil yang

sangat tinggi dan simetris di sepanjang pengujian juga membuktikan secara ilmiah bahwa implementasi teknik penyeimbangan data menggunakan metode *Random Over-Sampling* (ROS) berjalan dengan sangat efektif. Melalui replikasi acak pada dokumen kelas minoritas di dalam data training, ROS sukses menjaga keseimbangan bobot pemahaman algoritma terhadap kedua kelas polaritas. Dampak positifnya, model SVM mampu mempertahankan objektivitas dan sensitivitas klasifikasinya yang tinggi tanpa mengalami bias terhadap kelas mayoritas, meskipun bekerja pada kapasitas volume data yang maksimal.

4.4.1 Analisis Pengaruh Kernel terhadap Akurasi

Pemilihan fungsi kernel dalam algoritma Support Vector Machine (SVM) memegang peranan krusial dalam menentukan bagaimana data ulasan produk *Face Wash SOMBONG* dipetakan ke dalam ruang fitur berdimensi tinggi. Setiap kernel memiliki karakteristik matematis yang berbeda dalam membangun hyperplane atau batas keputusan. Analisis ini bertujuan untuk mengevaluasi efektivitas empat jenis kernel (Linear, RBF, Polynomial, dan Sigmoid) melalui pendekatan statistik guna menentukan model yang paling optimal dan stabil untuk digunakan dalam klasifikasi sentimen teks pada platform e-commerce.



Gambar 4. 32 Visualisasi Rentang dan Variabilitas Performa Kernel SVM

Berdasarkan grafik pada Gambar 4.31 serta rekapitulasi nilai rata-rata (*mean*) dan standar deviasi (*standard deviation*) pada Tabel 4.31, terlihat secara jelas karakteristik performa dan stabilitas dari masing-masing fungsi kernel yang diuji. Perlu dijelaskan secara metodologis bahwa nilai *Mean* (Rata-rata) pada setiap kernel diperoleh dari akumulasi rata-rata nilai akurasi yang diraih kernel tersebut di sepanjang empat skenario rasio data yang berbeda. Sebagai contoh, fungsi kernel RBF diuji secara konsisten pada rasio data 60:40, 70:30, 80:20, dan 90:10, sehingga nilai *Mean* untuk kernel RBF merupakan visualisasi rata-rata matematis dari performa akurasi keempat eksperimen tersebut. Rekapitulasi data deskriptif tersebut disajikan pada Tabel 4.19 berikut:

Tabel 4. 19 Perbandingan Nilai Akurasi dan Standar Deviasi pada Setiap Kernel

Kernel	Mean (Rata-rata)	Std (Standar Deviasi)
Linear	0.9948	0.0044
Polynomial	0.9818	0.0083
RBF	0.9970	0.0030
Sigmoid	0.9892	0.0024

Merujuk pada Tabel 4.19 dan sebaran interkuartil pada grafik *boxplot*, Kernel RBF secara empiris menduduki peringkat teratas dengan raihan nilai rata-

rata akurasi tertinggi mencapai 0.997029. Tingginya capaian ini membuktikan fleksibilitas fungsi basis radial dalam memetakan ruang fitur teks ulasan produk Face Wash SOMBONG yang bersifat *high-dimensional*.

Meskipun demikian, Kernel Linear tetap menunjukkan performa yang sangat kompetitif di peringkat kedua dengan nilai rata-rata sebesar 0.994825. Keadaan ini mengonfirmasi bahwa distribusi vektor kata hasil ekstraksi TF-IDF pada dataset ini memiliki kecenderungan struktur yang bersifat *linearly separable*. Pola bahasa konsumen dalam mengekspresikan sentimen positif maupun negatif memiliki batasan semantik yang tegas, sehingga fungsi pemetaan linear mampu membangun garis pemisah (*hyperplane*) secara efisien.

Di sisi lain, aspek stabilitas dan konsistensi model paling impresif ditunjukkan oleh Kernel Sigmoid yang mencatatkan nilai standar deviasi terkecil, yaitu sebesar 0.002488. Rendahnya nilai dispersi ini dibuktikan pada grafik *boxplot* melalui kotak interkuartil Sigmoid yang sangat rapat tanpa adanya kemunculan data pencilan (*outlier*). Secara teknis, hal ini membuktikan bahwa kernel Sigmoid memiliki tingkat ketahanan (*robustness*) yang luar biasa terhadap fluktuasi distribusi data uji di setiap skenario rasio pembagian data.

Sebaliknya, Kernel Polynomial mencatatkan nilai rata-rata akurasi terendah sekaligus tingkat fluktuasi tertinggi dengan nilai standar deviasi terbesar yaitu 0.008381. Pada grafik *boxplot*, ketidakstabilan kernel Polynomial dipertegas oleh batas *whisker* bawah yang memanjang hingga menyentuh nilai akurasi terendah di angka 0.970000 sebagai titik *outlier*.

Untuk membuktikan apakah perbedaan performa di antara keempat rumpun kernel tersebut signifikan atau sekadar kebetulan statistik, dilakukan pengujian hipotesis komparatif melalui metode *t-test*. Adapun matriks nilai signifikansi (*p-value*) hasil *t-test* antar kernel disajikan pada Tabel 4.20 berikut:

Tabel 4. 20 Hasil *t-test* (*p-value*) Antar Kernel

Kernel	Linear	Poly	RBF	Sigmoid
Linear	-	0.0339	0.4464	0.0722
Poly	0.0339	-	0.0144	0.1402
RBF	0.4464	0.0144	-	0.0078
Sigmoid	0.0722	0.1402	0.0078	-

Berdasarkan matriks uji *t-test* pada Tabel 4.20, diperoleh beberapa temuan ilmiah yang mendalam mengenai signifikansi perbedaan performa antar kernel secara statistik. Pada perbandingan antara kernel Linear dan kernel RBF, didapatkan nilai *p-value* sebesar 0.4464 yang berada jauh di atas ambang batas signifikansi 0.05. Hasil ini menegaskan secara ilmiah bahwa tidak terdapat perbedaan performa yang nyata atau signifikan antara kedua kernel tersebut pada dataset ulasan ini; walaupun secara nominal rata-rata akurasi RBF sedikit mengungguli Linear, keduanya memiliki kapabilitas matematis yang setara dalam memetakan batas keputusan klasifikasi. Pola hubungan yang serupa juga ditemukan pada perbandingan antara kernel Linear dan kernel Sigmoid dengan nilai *p-value* sebesar 0.0722 yang berada di atas ambang batas 0.05, mengindikasikan bahwa efektivitas klasifikasi yang dihasilkan oleh arsitektur linear dan sigmoid masih berada pada koridor kompetensi yang sepadan.

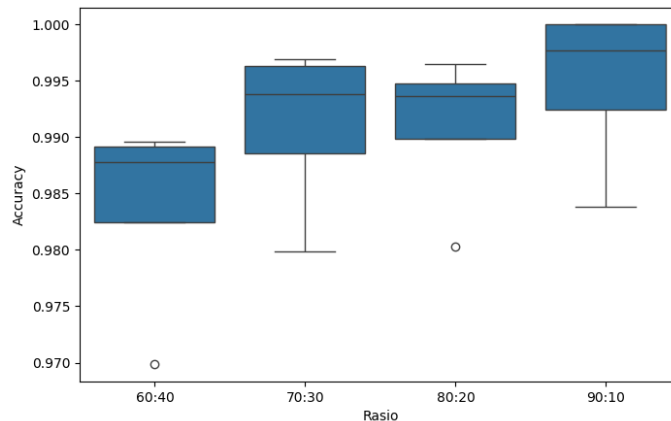
Sebaliknya, perbedaan performa yang sangat tajam dan signifikan secara statistik ditemukan pada perbandingan antara kernel RBF dan kernel Sigmoid, di

mana interaksi keduanya menghasilkan nilai p -value sebesar 0.0078 ($p < 0.01$). Angka tersebut membuktikan secara mutlak bahwa keunggulan kernel RBF dalam mengekstrak kompleksitas semantik non-linear pada data ulasan Face Wash SOMBONG merupakan hasil matematis yang valid dan bukan merupakan kebetulan variasi data.

Sementara itu, fungsi kernel Polynomial menunjukkan inkonsistensi yang nyata ketika dikomparasikan dengan rumpun kernel lainnya. Hal ini dibuktikan melalui nilai signifikansi kernel Polynomial terhadap kernel Linear sebesar 0.0339 ($p < 0.05$) dan terhadap kernel RBF sebesar 0.0144 ($p < 0.05$). Fakta empiris ini mengonfirmasi secara kuat bahwa kelemahan performa kernel Polynomial dalam memproses matriks fitur berdimensi tinggi (*high-dimensional sparse space*) pada penelitian ini bersifat signifikan pada taraf nyata 0.05, sehingga menempatkan fungsi kuadratik derajat dua tersebut di posisi paling tidak optimal dibandingkan fungsi kernel pembandingnya.

4.4.2 Analisis Pengaruh Rasio Split Data terhadap Akurasi

Pembagian rasio data antara data latih (*training set*) dan data uji (*test set*) merupakan salah satu faktor fundamental yang memengaruhi kemampuan generalisasi model *machine learning*. Dalam eksperimen ini, rasio data diuji pada empat skenario berbeda, yaitu 60:40, 70:30, 80:20, dan 90:10. Analisis ini bertujuan untuk melihat seberapa besar pengaruh ketersediaan volume data latih dalam memperkuat basis pengetahuan algoritma SVM untuk mengenali pola sentimen ulasan produk *Face Wash SOMBONG* secara konsisten.



Gambar 4. 33 Visualisasi Rentang dan Variabilitas Performa Rasio Data Latih

Berdasarkan grafik pada Gambar 4.32 serta rekapitulasi nilai rata-rata (*mean*) dan standar deviasi (*standard deviation*) pada Tabel 4.32, terlihat secara jelas adanya tren perubahan performa model SVM. Perlu dijelaskan secara metodologis bahwa nilai *Mean* (Rata-rata) pada setiap rasio data diperoleh dari rata-rata nilai akurasi empat fungsi kernel yang digunakan pada rasio tersebut. Sebagai contoh, pada rasio pembagian data 80:20, pengujian dilakukan berturut-turut menggunakan kernel Linear, RBF, Polynomial, dan Sigmoid, kemudian seluruh capaian nilai akurasi dari keempat kernel tersebut dirata-ratakan sehingga diperoleh nilai *Mean* spesifik untuk rasio 80:20. Rekapitulasi data statistik deskriptif pengaruh rasio split data disajikan pada Tabel 4.21 berikut:

Tabel 4. 21 Perbandingan Nilai Akurasi dan Standar Deviasi pada Setiap Rasio Data

Rasio Data	Mean (Rata-rata)	Std (Standar Deviasi)
60:40	0.9854	0.0094
70:30	0.9886	0.0077
80:20	0.9933	0.0055
90:10	0.9954	0.0053

Merujuk pada Tabel 4.21 dan visualisasi Gambar 4.18, terlihat adanya korelasi positif yang sangat stabil antara perluasan proporsi data latih dengan tingkat akurasi model, di mana skenario rasio 90:10 berhasil mencapai nilai rata-

rata tertinggi sebesar 0.995402. Peningkatan akurasi yang konsisten seiring bertambahnya data latih menunjukkan bahwa algoritma SVM mampu menyerap lebih banyak variasi pola kata kunci yang representatif untuk mengenali sentimen pengguna. Dengan ketersediaan volume data latih yang maksimal pada fase pembelajaran, model dapat memperhalus dan memperkokoh batas garis pemisah (*hyperplane*) sehingga mampu membedakan ulasan-ulasan dunia nyata yang memiliki kemiripan diksi namun berbeda konteks emosional secara lebih akurat.

Selain peningkatan nilai rata-rata akurasi, penambahan data latih juga berdampak linear pada penurunan nilai standar deviasi yang berhasil menyusut hingga mencapai titik terendah pada rasio 90:10 yaitu sebesar 0.005309. Hal ini berbanding terbalik dengan rasio 60:40 yang mencatatkan nilai variansi kesalahan terbesar dengan standar deviasi sebesar 0.009407. Pada grafik *boxplot*, fenomena penyusutan ini ditunjukkan oleh kotak interkuartil yang semakin rapat dan meninggi dari rasio 60:40 menuju 90:10, meskipun terdapat titik *outlier* minor pada rasio 60:40 di angka 0.970000 dan rasio 80:20 di angka 0.980000. Secara teknis, karakteristik ini membuktikan bahwa alokasi data latih yang lebih besar tidak hanya berfungsi meningkatkan performa model secara kuantitas, melainkan juga memperkuat tingkat reliabilitas dan stabilitas model dari risiko fluktuasi galat yang tajam.

Untuk membuktikan signifikansi pengaruh variasi pembagian data tersebut, dilakukan uji hipotesis komparatif menggunakan metode *t-test*. Adapun matriks nilai signifikansi (*p-value*) hasil pengujian *t-test* antar rasio data disajikan pada Tabel 4.22 berikut:

Tabel 4. 22 Hasil *t-test* (*p-value*) Antar Rasio Data

Rasio Data	60:40	70:30	80:20	90:10
60:40	-	0.6180	0.1982	0.1158
70:30	0.6180	-	0.3639	0.2039
80:20	0.1982	0.3639	-	0.6202
90:10	0.1158	0.2039	0.6202	-

Berdasarkan matriks uji *t-test* pada Tabel 4.22, diperoleh kesimpulan statistik yang sangat krusial di mana seluruh kombinasi perbandingan antar rasio split data menghasilkan nilai *p-value* yang berada di atas ambang batas nyata 0.05. Perbandingan terdekat antara rasio terkecil dan terbesar, yaitu rasio 60:40 terhadap rasio 90:10, menghasilkan nilai $p = 0.1158$, sementara interaksi antar skenario lainnya seperti rasio 80:20 terhadap 90:10 mencatatkan nilai tertinggi mencapai $p = 0.6202$. Secara ilmiah, temuan ini membuktikan bahwa tidak terdapat perbedaan performa yang signifikan secara nyata di antara keempat skenario pembagian data tersebut. Fakta ini menegaskan bahwa arsitektur model SVM yang dibangun dalam penelitian ini memiliki tingkat stabilitas yang luar biasa tinggi, karena sejak penggunaan proporsi data latih paling minimum (rasio 60:40) pun, model sudah mampu menunjukkan tingkat kecerdasan dan akurasi yang sangat tinggi di atas angka 0.980000.

Ketiadaan perbedaan signifikan ini menunjukkan efisiensi matematis yang sangat baik dari model yang dikembangkan, yang berarti algoritma tidak memerlukan ketergantungan pada volume data dalam jumlah masif untuk mencapai titik konvergensi yang optimal. Integrasi teknik penyeimbangan data menggunakan metode *Random Over-Sampling* (ROS) yang dipadukan secara tepat dengan

ekstraksi fitur TF-IDF terbukti telah berhasil mengisolasi informasi-informasi penting dari teks ulasan pelanggan Face Wash SOMBONG secara efektif sejak tahap awal. Dengan kata lain, penambahan porsi data hingga mencapai rasio maksimal 90:10 tidak mengubah struktur kecerdasan model secara radikal, melainkan hanya berperan sebagai tahap optimasi halus (*fine-tuning*) untuk menyempurnakan detail margin klasifikasi. Secara fundamental, model SVM ini telah teruji sangat handal, kokoh, dan matang bahkan dalam kondisi ketersediaan data latih yang terbatas.

Keberhasilan eksperimen dalam mengklasifikasikan ulasan produk *Face Wash SOMBONG* ini tidak hanya merupakan pencapaian teknis, namun juga memiliki landasan filosofis yang kuat dalam perspektif Islam. Hal ini berkaitan dengan konsep ketetapan dan klasifikasi segala sesuatu yang telah diatur oleh Allah SWT, sebagaimana firman-Nya dalam QS. Al-Qamar ayat 49:

إِنَّا كُلَّ شَيْءٍ خَلَقْنَاهُ بِقَدَرٍ

"Sesungguhnya Kami menciptakan segala sesuatu menurut ukuran." (QS.Al-Qamar:49)

Berdasarkan tinjauan tafsir terhadap ayat tersebut, ditegaskan bahwa seluruh ciptaan Allah SWT telah ditetapkan ukuran, sifat, dan karakteristiknya secara presisi (Irfan, 2023). Dalam konteks penelitian ini, ulasan konsumen yang tampak acak sebenarnya memiliki "ukuran" atau pola linguistik yang tetap. Penggunaan algoritma SVM merupakan upaya ilmiah untuk mengenali ketetapan pola (*qadar*) tersebut melalui pembobotan fitur, sehingga ulasan dapat dikelompokkan sesuai dengan karakteristik aslinya. Hal ini membuktikan bahwa

klasifikasi digital adalah bentuk pengenalan terhadap sunnatullah yang ada pada setiap entitas data.

Selanjutnya, proses pengumpulan ulasan yang menjadi basis data penelitian ini bersandar pada nilai kejujuran dalam memberikan kesaksian (*syahadah*), sebagaimana diperintahkan Allah SWT dalam QS. Al-Ahzab ayat 70:

يَا أَيُّهَا الَّذِينَ آمَنُوا اتَّقُوا اللَّهَ وَقُولُوا قَوْلًا سَدِيدًا

"Wahai orang-orang yang beriman! Bertakwalah kamu kepada Allah dan ucapkanlah perkataan yang benar." (QS. Al-Ahzab:70)

Dalam analisis terhadap ayat ini, *Qaulan Sadida* didefinisikan sebagai perkataan yang jujur, tepat sasaran, dan tidak mengandung manipulasi (HALID, 2024). Keberhasilan model dalam mencapai akurasi tinggi pada penelitian ini didorong oleh adanya ulasan-ulasan objektif dari pengguna. Ketika konsumen menyampaikan fakta yang sebenarnya mengenai produk, hal tersebut mempermudah sistem dalam membentuk batas keputusan yang akurat. Dengan demikian, teknologi analisis sentimen berfungsi sebagai alat untuk memvalidasi kejujuran informasi di tengah maraknya distorsi informasi digital.

Terakhir, tahapan pengujian stabilitas model melalui berbagai skenario rasio data dan *T-test* merupakan manifestasi dari prinsip Tabayyun sebagaimana diamanatkan dalam QS. Al-Hujurat ayat 6:

يَا أَيُّهَا الَّذِينَ آمَنُوا إِنْ جَاءَكُمْ فَاسِقٌ بِنَبَأٍ فَتَبَيَّنُوا أَنْ تُصِيبُوا قَوْمًا بِجَهَالَةٍ فَتُصْحِرُوا عَلَى مَا فَعَلْتُمْ
نَدِيمِينَ

"Wahai orang-orang yang beriman! Jika seseorang yang fasik datang kepadamu membawa suatu berita, maka telitilah kebenarannya, agar kamu tidak mencelakakan suatu kaum karena kebodohan (kecerobohan), yang akhirnya kamu menyesali perbuatanmu itu." (QS. Al-Hujurat:6)

Tafsir tahlili ayat ini menekankan pada pentingnya verifikasi (*cross-check*) terhadap informasi guna menghindari pengambilan keputusan yang salah akibat kecerobohan (*jahalah*) (Mukhamad Zidanul Akhsan; Yeti Dahliana, n.d.). Temuan bahwa model SVM tetap stabil dan "pintar" meskipun pada rasio data minimal (60:40) membuktikan efektivitas sistem dalam melakukan klarifikasi informasi digital secara otomatis. Dengan melakukan pengujian statistik yang ketat dan memastikan tidak adanya misklasifikasi (*false positive/negative*), peneliti telah menerapkan prinsip kehati-hatian dalam menyimpulkan sebuah berita ulasan. Secara integratif, pemanfaatan teknologi ini menjadi sarana untuk menjaga kemaslahatan umat dalam bertransaksi, memastikan setiap opini yang tersebar telah melalui proses verifikasi yang reliabel dan dapat dipertanggungjawabkan.

BAB V

KESIMPULAN DAN SARAN

5.1 Kesimpulan

Penelitian ini berhasil menerapkan pembobotan fitur *Term Frequency-Inverse Document Frequency* (TF-IDF) dan algoritma *Support Vector Machine* (SVM) untuk klasifikasi sentimen ulasan produk Face Wash SOMBONG di *e-commerce* secara otomatis. Evaluasi metrik *accuracy*, *precision*, *recall*, dan *f1-score* di berbagai skenario pembagian data menunjukkan bahwa sistem memiliki keandalan yang sangat tinggi dengan seluruh nilai performa konsisten berada di atas 0.970000. Seiring bertambahnya proporsi data latih yang diseimbangkan dengan teknik *Random Over-Sampling* (ROS), kemampuan model dalam mengenali pola sentimen meningkat secara linear hingga mencapai titik optimal pada skenario rasio data 90:10 dengan raihan nilai rata-rata akurasi tertinggi sebesar 0.995402.

Melalui analisis komparasi untuk menentukan fungsi kernel yang paling optimal, kernel RBF dan kernel Linear terbukti menjadi model terbaik untuk dataset ulasan ini. Berdasarkan hasil pengujian *t-test*, tidak terdapat perbedaan performa yang signifikan secara nyata antara kedua kernel tersebut ($p = 0.4464$), di mana pada skenario data training maksimal (rasio 90:10) dengan konfigurasi parameter $C = 10$, keduanya sukses mencatatkan nilai performa sempurna sebesar 1.000000 berkat karakteristik data teks yang bersifat *linearly separable*. Sementara itu, kernel Polynomial dan Sigmoid menghasilkan performa yang lebih rendah, sehingga kernel RBF dan Linear ditetapkan sebagai arsitektur klasifikasi sentimen yang

paling optimal dan reliabel untuk diimplementasikan pada ulasan produk Face Wash SOMBONG.

5.2 Saran

Meskipun penelitian ini menunjukkan hasil yang sangat memuaskan, terdapat beberapa saran yang dapat diajukan untuk pengembangan penelitian selanjutnya:

1. Eksperimen Kombinasi Fitur dan Arsitektur: Penelitian selanjutnya disarankan untuk mengeksplorasi metode ekstraksi fitur berbasis representasi vektor lain seperti Word2Vec atau FastText, serta membandingkan performa SVM dengan arsitektur *deep learning* (seperti LSTM atau BERT) guna menangkap hubungan semantik dan kontekstual antar kata yang lebih mendalam pada data teks yang lebih kompleks.
2. Optimasi Preprocessing dan Penyeimbangan Data: Disarankan untuk memperluas tahapan prapemrosesan dengan menambahkan kamus normalisasi kata *slang* atau singkatan bahasa tidak baku yang lebih komprehensif untuk menangani karakteristik bahasa informal *e-commerce*. Selain itu, perlu dilakukan komparasi dengan metode penyeimbangan data alternatif seperti SMOTE (*Synthetic Minority Over-sampling Technique*) untuk menganalisis efek sintesis data terhadap stabilitas model pada kapasitas dataset yang lebih besar.

DAFTAR PUSTAKA

- Adrian, M. R., Putra, M. P., Rafialdy, M. H., Rakhmawati, N. A., & Informasi, D. S. (2021). *Perbandingan Metode Klasifikasi Random Forest dan SVM Pada Analisis Sentimen PSBB*. 7(1), 36–40.
- Alaiya, A., & Agusniar, C. (2025). *Sentiment Analysis of E-Commerce Product Reviews on Tokopedia Using Support Vector Machine*. 9(5), 2869–2878.
- Arbiansyah, G., & Haq, F. (2025). *Customer Comment Clustering for Kahf Face Wash at Kahf Official Shop Using K-Means Method*. 1(3), 86–92.
- Azzahra, D. M., Alam, S., Wastukancana, S., Barat, J., Secret, V., & Solution, P. P. (2023). *ANALISIS SENTIMEN ULASAN PRODUK SERUM WAJAH PADA BEAUTY*. 7(3), 1604–1611.
- Bansal, M., Goyal, A., & Choudhary, A. (2022). A comparative analysis of K-Nearest Neighbor , Genetic , Support Vector Machine , Decision Tree , and Long Short Term Memory algorithms in machine learning. *Decision Analytics Journal*, 3(November 2021), 100071. <https://doi.org/10.1016/j.dajour.2022.100071>
- Bima maulana, M. H. F. (2025). *Analisis Sentimen terhadap Ulasan Produk Facial Wash Sombong Menggunakan Algoritma Naive Bayes*. 1(2), 117–125.
- Demetria, P., & Wedhasmara, A. (2025). *Analisis Sentimen Pelanggan Terhadap Penilaian Produk Pada Tokopedia Nyemil . Saji Menggunakan Metode Support Vector Machine*. 7(2), 1350–1361.
- Fahmi Fiddin, Taufik Hidayat, D. (2025). *Analisis Sentimen Ulasan Produk Sayur di Tokopedia Menggunakan Model Support Vector Machine Dengan Representasi TF-IDF Jurnal SAINTIKOM (Jurnal Sains Manajemen Informatika dan Komputer)*. 24, 162–172.
- HALID, I. (2024). *CHERRY PICKING FALLACY DAN ETIKA KEJUJURAN DALAM KOMUNIKASI ISLAM: Analisis Tafsir Q.S Al-Ahzab Ayat 70*.
- Handayani, R. N. (2021). *OPTIMASI ALGORITMA SUPPORT VECTOR MACHINE UNTUK ANALISIS SENTIMEN PADA ULASAN PRODUK TOKOPEDIA MENGGUNAKAN PSO*. 20(2), 97–108.
- Idris, I. S. K. (2023). *Analisis Sentimen Terhadap Penggunaan Aplikasi Shopee Menggunakan Algoritma Support Vector Machine (SVM)*. 5, 32–35.
- Irfan, D. A. H. (2023). *Tafsir Tematik Surat Al-Qamar Ayat 48-49*.
- Milal, I. S., Hasanudin, M., Azhari, M. A. N., Nugraha, R. A., Agustina, N., & Damayanti, S. E. (2023). *KLASIFIKASI TEKS REVIEW PADA E-COMMERCE TOKOPEDIA MENGGUNAKAN ALGORITMA SVM*. 05(01), 34–45.

- Mukhamad Zidanul Akhsan; Yeti Dahliana. (n.d.). *IMPLEMENTASI METODE TAHLILI DALAM SURAH AL-HUJURAT AYAT 6. 6, 1–22*. <https://doi.org/10.33087/dikdaya.v10i2.173.2>
- Muzaki, A., Febriana, V., & Cholifah, W. N. (2024). *ANALISIS SENTIMEN PADA ULASAN PRODUK DI E-COMMERCE DENGAN METODE NAIVE BAYES*. 05(04), 758–765.
- Pratiwi, R. W., H, S. F., Intan, D., & A, Q. R. (2021). *Analisis Sentimen Pada Review Skincare Female Daily Menggunakan Metode Support*. 8106, 40–46.
- Puspa, T., Sanjaya, R., Fauzi, A., Fitri, A., & Masruriyah, N. (2023). *Analisis sentimen ulasan pada e-commerce shopee menggunakan algoritma naive bayes dan support vector machine Analysis of review sentiment on shopee e-commerce using the naive bayes algorithm and support vector machine*. 4, 16–26. <https://doi.org/10.37373/infotech.v4i1.422>
- Tjut Awaliyah Zuraiyah, Mulyati, G. H. F. H. (2023). *PERBANDINGAN METODE NAÏVE BAYES, SUPPORT VECTOR MACHINE DAN RECURRENT NEURAL NETWORK PADA ANALISIS SENTIMEN ULASAN PRODUK E-COMMERCE*. 6223(1), 27–43.
- Widhiyanta, Nurwahyudi, Isnaini Muhandhis, Roszi Syadillal Jannah, L. A. W. (2025). *ANALISIS SENTIMEN ULASAN PRODUK MOISTURIZER SKINTIFIC DI TOKOPEDIA MENGGUNAKAN SUPPORT VECTOR MACHINE SENTIMENT*. 18(1), 129–142.
- Willy Wildan Kamal, C. I. R. (n.d.). *Analisis Sentimen Ulasan Produk : Kajian Pustaka*.