

ANALISIS SENTIMEN ULASAN PRODUK *BRAND SKINCARE XYZ* PADA MEDIA SOSIAL MENGGUNAKAN ALGORITMA *LOGISTIC REGRESSION*

Navisatul Marchamah ^{1*}, Siti Zulaiqoh ², Alfira Candra Inanda ³

¹⁻³ Teknik Informatika, STMIK Bina Patria Magelang, Indonesia.

Email: navisa095@gmail.com*

ABSTRAK

Perkembangan industri *skincare* mendorong meningkatnya jumlah ulasan konsumen pada *platform* media sosial. Banyaknya ulasan tersebut menyulitkan proses analisis secara manual sehingga diperlukan metode otomatis untuk mengidentifikasi sentimen pengguna. Penelitian ini bertujuan menganalisis sentimen ulasan produk *Brand Skincare XYZ* pada media sosial memakai algoritma *Logistic Regression*. Dataset terdiri dari 760 ulasan yang diperoleh dari Kaggle dan berasal dari YouTube, TikTok, Facebook, Instagram, dan Forum. Pelabelan sentimen dilakukan secara otomatis menggunakan pendekatan *Lexicon-Based Sentiment Analysis*, bukan pelabelan manual. Tahapan penelitian meliputi *preprocessing* data yang mencakup *case folding*, *tokenizing*, dan *stemming* menggunakan pustaka Sastrawi, serta klasifikasi menggunakan *Logistic Regression* dengan penanganan ketidakseimbangan kelas melalui *class weight balanced*. Validasi model dilakukan melalui *5-fold cross-validation* dan perbandingan dengan algoritma *Support Vector Machine (SVM)* untuk memastikan kestabilan dan menguji indikasi *overfitting/underfitting*. Setelah *preprocessing* diperoleh 737 dibagi menjadi 589 data latih dan 148 data uji. Hasil pelabelan menunjukkan distribusi sentimen Netral sebesar 66,08%, Positif 22,39%, dan Negatif 11,53%. Hasil evaluasi model menunjukkan nilai *Accuracy* sebesar 81,41%, *Precision* sebesar 80,81%, *Recall* sebesar 81,08%, dan *F1-Score* sebesar 80,77%.

Kata Kunci: Analisis Sentimen; *Brand Skincare XYZ*; *Logistic Regression*; Media Sosial; *TF-IDF*

ABSTRACT

The growth of the *skincare* consumer reviews across various social media platforms. Thereby requiring automated methods to identify user opinions. *Brand Skincare XYZ* product reviews on social media using the *Logistic Regression* algorithm. The dataset consisted of 760 reviews collected from Kaggle, originating from five social media platforms: YouTube, TikTok, Facebook, Instagram, and online forums. The research process included sentiment labeling using the *Lexicon-Based Sentiment Analysis* method, data preprocessing covering *case folding*, *cleaning*, *tokenizing*, *stopword removal*, and *stemming* using the Sastrawi library, feature extraction using *Term Frequency-Inverse Document Frequency (TF-IDF)*, and sentiment classification using *Logistic Regression* with class imbalance handling through *balanced class weighting*. After preprocessing, 737 reviews remained and were. The sentiment labeling results showed that neutral sentiment accounted for 66.08% of the reviews, followed by positive sentiment at 22.39% and negative sentiment at 11.53%. The evaluation results indicated that the *Logistic Regression* model achieved an *Accuracy* of 81.41%, *Precision* of 80.81%, *Recall* of 81.08%, and *F1-Score* of 80.77%. These findings demonstrate that the *Logistic Regression* algorithm performs well in classifying sentiment in *Brand Skincare X* product reviews on social media.

Keywords: *Brand Skincare XYZ*; *Logistic Regression*; *Sentiment Analysis*; *Social Media*; *TF-IDF*

1. Pendahuluan

Pertumbuhan teknologi internet menjadikan seseorang mempermudah membuat konten secara online dan menyebarkan informasi secara pesat, seperti halnya pada ulasan di sosmed. Konten yang dihasilkan seseorang bisa dipakai perusahaan sebagai sumber informasi untuk menilai tanggapan masyarakat terhadap produk maupun layanan yang ditawarkan (Daffa & Brata, 2025). Pendapat dan pengalaman pelanggan tersebut mempunyai peranan terpenting, baik bagi pihak produsen maupun calon pemakai, agar bisa mengasah paparan mengenai persepsi publik hingga membantu dalam mengambil keputusan yang lebih tepat.

Dalam beberapa tahun terakhir, industri perawatan kulit memaparkan pengembangan pesat. Diantara brand yang memikat perhatian, terutama di media sosial, diantaranya yang dalam penelitian ini disebut sebagai "Brand Skincare XYZ" (nama brand disamarkan untuk menghindari isu hak cipta). Sejak hadir di Indonesia pada tahun 2021, merek ini mempunyai respons konsumen terbaik dan semakin dikenal luas setelah produk *5X Ceramide Barrier Repair Moisturizer Gel* menjadi tren bermacam platform digital (Harnelia, 2024). Namun, banyaknya varian produk yang ditawarkan *Brand Skincare XYZ* sering kali menjadikan konsumen merasa bingung memilah dengan keadaan dan kebutuhan kulit mereka.

Menguatnya ulasan di internet menjadikan konsumen semakin kesulitan dalam menemukan informasi yang benar-benar relevan untuk mendukung keputusannya. Komentar yang diperoleh dari bermacam *platform* media sosial umumnya masih berupa data teks yang belum tersusun secara sistematis, sehingga diperlukan progres pengolahan awal sebelum data tersebut dapat dianalisis lebih memadai (Rifaldi et al., 2023). Pengolahan awal tersebut biasanya mencakup normalisasi huruf (*case folding*), pemecahan teks menjadi token (*tokenizing*), penghapusan kata-kata yang tidak memiliki makna penting (*stopword removal*), serta proses *stemming* untuk mengubah kata menjadi bentuk dasarnya (Rifaldi et al., 2023).

Dengan bermacam opini pada ulasan, dibutuhkan suatu pendekatan yang bisa mengolah secara otomatis. Diantara metode yang banyak dipakai yakni analisis sentimen, sebagai mengelompokkan kecenderungan, maupun netral (Harnelia, 2024).

Dalam bidang analisis sentimen, sejumlah algoritma machine learning telah banyak diterapkan untuk mengelompokkan opini yang terdapat dalam data teks, di antaranya *Naïve Bayes*, *Support Vector Machine* (SVM), dan *Logistic Regression*. *Logistic Regression* cara yang dipakai untuk menjalin bermacam variabel terikat yang mempunyai dua atau lebih kategori (Daffa & Brata, 2025). Iskandar dan Nataliani (2021) yang membandingkan performa algoritma *Naïve Bayes*, SVM, dan k-NN dalam analisis sentimen berbasis aspek pada ulasan gadget menemukan bahwa performa masing-masing algoritma bervariasi bergantung pada karakteristik data dan aspek yang dianalisis (Iskandar & Nataliani, 2021). Penelitian yang dilaksanakan oleh Daffa, Brata, dan Pangestu (2025) mengenai klasifikasi ulasan pelanggan produk es krim membandingkan kinerja *Naive Bayes* dan *Logistic Regression*. Hasil penelitian tersebut memaparkan bahwasanya *Logistic Regression* mengasah tingkatan akurasi yang dibandingkan *Naive Bayes*, baik pada data sebelum maupun sesudah dilaksanakannya penyeimbangan terhadap *Synthetic Minority Over-sampling Technique* (SMOTE). Model terbaik yang diperoleh mampu mencapai akurasi sebesar 94,56%.

Walaupun demikian, pemakaian *Logistic Regression* tidak selalu menghasilkan performa yang tinggi pada setiap kasus. Mulyono dan Saprudin (2025) melaksanakan penelitian terhadap komentar berbahasa Indonesia pada *platform* YouTube yang membahas isu ketenagakerjaan. Dengan menerapkan kombinasi *Logistic Regression*, TF-IDF, dan SMOTE, model yang dibangun hanya memperoleh akurasi sebesar 44%. Selain itu, kemampuan model dalam mengidentifikasi kelas positif relatif rendah dengan nilai F1-score sebesar 0,29, sedangkan pada kelas negatif

mencapai 0,55. Temuan tersebut memaparkan bahwasanya ketidakseimbangan jumlah data antar kelas serta karakteristik komentar yang cenderung mengandung makna tersirat dapat menjadi hambatan dalam proses klasifikasi sentimen berbahasa Indonesia (Mulyono & Saprudin, 2025).

Penelitian terkait penerapan dan perbandingan algoritma *Logistic Regression* maupun *Support Vector Machine* (SVM) turut banyak dilakukan pada berbagai domain data teks berbahasa Indonesia maupun mancanegara. (Styawati et al., 2021) memakai SVM untuk menganalisis sentimen masyarakat terhadap Program Kartu Prakerja pada *platform* Twitter. Selanjutnya, (Lidinillah et al., 2023) membandingkan performa *Logistic Regression* dan SVM pada ulasan *platform Steam* di Twitter, sedangkan (Prasetyo, 2025) menyandingkan *Logistic Regression*, SVM, dan *Random Forest* pada analisis sentimen aplikasi Gopay. Di sisi lain, penelitian oleh (Wati et al., 2023) menyoroti pengaruh pembobotan TF-IDF dan penerapan *stopword removal* terhadap peningkatan performa klasifikasi sentimen pada teks berbahasa Indonesia, yang relevan dengan tahapan *preprocessing* pada penelitian ini.

Meskipun kajian mengenai analisis sentimen telah terjalin pesat, penelitian yang berfokus pada ulasan produk *Brand Skincare XYZ* di media sosial dengan memakai algoritma *Logistic Regression* dan melibatkan tiga kategori sentimen, yaitu positif, negatif, serta netral, masih belum banyak ditemukan (Sikana et al., 2025). Penelitian sebelumnya terkait produk *Brand Skincare XYZ* lebih banyak memakai algoritma SVM dengan klasifikasi dua sentimen saja, yakni positif dan negatif (Harnelia, 2024). Dengan demikian, efektivitas *Logistic Regression* dalam menangani kasus yang sama, terutama pada skema klasifikasi tiga kelas sentimen, masih perlu diteliti lebih lanjut.

Berdasarkan latar belakang yang telah dipaparkan, penelitian ini memfokuskan suatu penerapan algoritma *Logistic Regression* dalam mengelompokkan sentimen ulasan produk *Brand Skincare XYZ* yang diperoleh dari berbagai media sosial ke dalam tiga kategori, yaitu positif, negatif, dan netral. Dengan itu, penelitian ini juga bertujuan untuk mengetahui tingkat kinerja algoritma tersebut dalam proses klasifikasi sentimen. Melalui penelitian ini berharap bisa mengasah informan mengenai pandangan konsumen terhadap produk *Brand Skincare XYZ*, sekaligus menjadi masukan bagi perusahaan dalam upaya meningkatkan mutu produk. Di samping itu, hasil penelitian diharapkan mampu memperkaya kajian ilmiah mengenai analisis sentimen pada teks berbahasa Indonesia.

2. Metoda

Penelitian ini menerapkan pendekatan kuantitatif dengan memanfaatkan teknik analisis sentimen untuk mengidentifikasi dan mengelompokkan opini pemakai produk *Brand Skincare XYZ* yang tersebar di media sosial. Tahapan penelitian meliputi pengumpulan data, pra proses data, pemberian label sentimen, pembagian dataset, pembangunan model *Logistic Regression*, pengujian performa model, serta interpretasi hasil penelitian. Data yang dipakai berasal dari dataset Kaggle yang memuat ulasan mengenai produk *Brand Skincare XYZ* dari berbagai *platform* digital, seperti Instagram, YouTube, TikTok, Facebook, dan forum diskusi daring (dataset diperoleh dari repositori Kaggle, lihat rujukan (Harnelia, 2023), pada Daftar Pustaka). dalam penelitian ini sebanyak 760 ulasan yang ditulis dalam bahasa Indonesia. Penentuan label sentimen dilakukan menggunakan pendekatan Lexicon-Based Sentiment Analysis. Metode ini bekerja dengan menghitung selisih antara jumlah kata yang mempunyai makna positif dan jumlah kata yang bermakna negatif pada setiap ulasan. Perhitungan dilakukan menggunakan rumus berikut:

$$Score = \sum KataPositif - \sum KataNegatif \quad (1)$$

Apabila nilai skor lebih besar dari nol, ulasan dikategorikan sebagai sentimen positif. Sementara itu, nilai skor sama dengan nol memperlihatkan sentimen netral karena tidak terdapat kata yang mengandung sentimen atau jumlah kata positif dan negatif berada dalam kondisi seimbang. Perlu ditegaskan bahwa seluruh proses pelabelan sentimen pada penelitian ini dilakukan secara otomatis menggunakan pendekatan *Lexicon-Based Sentiment Analysis* berdasarkan kamus kata positif dan negatif yang telah disusun, bukan melalui pelabelan manual oleh annotator manusia. Sebelum memasuki tahap analisis, data terlebih dahulu melalui proses preprocessing yang bertujuan untuk membersihkan teks hingga lebih mudah diolah oleh sistem (Khairunnisa et al., 2021). Adapun tahapan preprocessing yang dilakukan adalah sebagai berikut:

1. *Case Folding*

Case folding menjadi progres normalisasi teks dengan mengubah seluruh huruf menjadi huruf kecil (*lowercase*). Langkah ini dilakukan agar kata yang sama tidak dianggap berbeda hanya karena perbedaan penggunaan huruf kapital.

2. *Cleaning*

Pada tahap ini pembersihan berbagai elemen yang tidak relevan terhadap analisis sentimen, seperti tautan situs web, mention akun, hashtag, emoji, angka, tanda baca, serta karakter khusus lainnya. Proses ini bertujuan menghasilkan teks yang lebih terfokus pada isi ulasan.

3. *Tokenizing*

Setelah teks dibersihkan, proses selanjutnya adalah memecah kalimat menjadi unit-unit kata yang disebut token. Pemisahan dilakukan berdasarkan spasi sehingga setiap kata dapat dianalisis secara terpisah.

4. *Stopword Removal*

Tahap *stopword removal* dilakukan untuk menghilangkan kata-kata umum yang frekuensi kemunculannya tinggi tetapi tidak memberikan informasi penting terhadap penentuan sentimen.

5. *Stemming*

Proses stemming bertujuan mengubah kata yang mengandung imbuhan menjadi bentuk dasarnya dengan memanfaatkan pustaka Sastrawi (Albab et al., 2023). Sebagai contoh, kata "membantu" diubah menjadi "bantu", "membersihkan" menjadi "bersih", dan "memakai" menjadi "guna". Tahapan ini dilakukan untuk mengurangi variasi bentuk kata agar lebih konsisten.

Ekstraksi fitur pada penelitian ini menggunakan metode *Term Frequency-Inverse Document Frequency* (TF-IDF), yang terbukti efektif dalam merepresentasikan bobot kata pada dokumen teks berdasarkan frekuensi kemunculannya (Septiani & Isabela, 2022). Parameter yang digunakan adalah untuk membatasi jumlah fitur maksimal sebanyak 5.000 kata, kombinasi unigram dan bigram untuk menangkap pola satu kata maupun dua kata sekaligus, serta pengabaian kata yang hanya muncul pada kurang dari dua dokumen.

Setelah seluruh tahapan praproses selesai dilaksanakan, data kemudian dipisahkan menjadi dua bagian, yaitu data pelatihan dan data pengujian. Pembagian dataset dilakukan dengan perbandingan 80:20, sehingga diperoleh 589 data untuk proses pelatihan model (*training set*) dan 148 data untuk proses evaluasi (*testing set*). Dalam proses pembagian tersebut diterapkan metode *stratified sampling* agar distribusi setiap kategori sentimen, baik positif, negatif, maupun netral, tetap proporsional pada kedua kelompok data.

Model klasifikasi dibangun menggunakan algoritma *Logistic Regression* dari *library Scikit-learn* dengan konfigurasi sebagai berikut: *solver* L-BFGS sebagai algoritma optimasi yang efisien untuk dataset berukuran sedang, pendekatan multinomial untuk menangani klasifikasi tiga kelas (Positif, Negatif, Netral) secara serentak menggunakan fungsi *softmax*, batas iterasi maksimal 1.000 untuk memastikan konvergensi model, serta *class weight balanced* untuk menangani ketidakseimbangan distribusi kelas secara otomatis dengan memberikan bobot lebih besar pada kelas minoritas.

Kinerja model dinilai memakai empat indikator yang umum dipakai dalam evaluasi klasifikasi teks. *Accuracy* dipakai untuk mengetahui proporsi prediksi yang sesuai dengan kondisi sebenarnya dari seluruh data pengujian. *Precision* memaparkan tingkat ketepatan model dalam memberikan label pada suatu kategori sehingga meminimalkan kesalahan klasifikasi. *Recall* dipakai untuk mengukur keahlian cara dalam mengidentifikasi seluruh data yang termasuk ke dalam masing-masing kelas. Sementara itu, *F1-Score* merupakan ukuran yang menggabungkan nilai *precision* dan *recall* secara seimbang melalui perhitungan rata-rata harmonik.

Selain evaluasi menggunakan skema pembagian data latih dan data uji (*train-test split*), penelitian ini juga menerapkan metode *5-Fold Stratified Cross-Validation* untuk menilai kestabilan performa model secara lebih menyeluruh. Pada pendekatan ini, keseluruhan data (737 data) dibagi menjadi lima subset (*fold*) dengan proporsi kelas sentimen yang tetap seimbang pada setiap *fold*. Model dilatih dan diuji sebanyak lima kali, di mana pada setiap iterasi satu *fold* digunakan sebagai data uji dan empat *fold* lainnya digunakan sebagai data latih. Rata-rata dari kelima hasil pengujian tersebut kemudian dihitung sebagai ukuran akhir kestabilan dan konsistensi performa model.

3. Hasil dan Pembahasan

3.1 Distribusi Sentimen

Proses pemberian label sentimen pada penelitian ini memakai pendekatan *Lexicon-Based Sentiment Analysis*, yang mengkategorikan ulasan ke dalam sentimen positif, negatif, dan netral. Tabel 1.

Tabel 1. Distribusi Sentimen Ulasan Produk Brand Skincare XYZ

Sentimen	Jumlah Data	Presentase
Netral	499	65,66%
Positif	176	23,16%
Negatif	85	11,18%
Total	760	100

Berdasarkan Tabel 1, sentimen netral mendominasi dataset dengan jumlah 499 ulasan atau 65,66% dari keseluruhan data. Sentimen positif berjumlah 176 ulasan atau 23,16%, sedangkan sentimen negatif sebanyak 85 ulasan atau 11,18%. Berdasarkan hasil tersebut, mayoritas ulasan yang diberikan pengguna cenderung berupa penyampaian informasi atau pengalaman tanpa menunjukkan penilaian yang secara jelas mengarah pada sentimen positif maupun negatif terhadap produk *Brand Skincare XYZ*.

3.2 Hasil *Preprocessing* Data

Sebelum memasuki tahap klasifikasi, data terlebih dahulu melalui progres *preprocessing* untuk memperbaiki kualitas dan konsistensi data. Tahapan yang dilakukan mencakup *cleaning*,

case folding, tokenisasi, penghapusan *stopword*, serta *stemming*. Setelah seluruh proses tersebut selesai diterapkan, diperoleh 737 data yang siap digunakan untuk analisis lebih lanjut. Ulasan dari total 760 data awal. Data yang tidak memenuhi kriteria analisis, seperti data kosong atau tidak mengandung teks yang dapat diproses, pada Gambar 1

```

Data setelah preprocessing: 737

Comment \
0 Heran sama skintific 😊, gue pake fw ini itu ga...
1 Abis moisturizer yg paket d atas abis itu taha...
2 @zuer.id aku oily msih stay d biru pen pindah ...
3 @zuer.id kak sbelum pake skintific jetwatku b...
4 uuuu udhh mauu abis 2 botol yg facial wash nya...

Comment_clean
0 heran skintific gue cocok malah nambah jerawat...
1 abis moisturizer paket atas abis tahap boleh s...
2 oily msih stay biru pen pindah panthenol tpi t...
3 sbelum skintific jetwatku brsih skrg mlh jerw...
4 uuuu udhh mauu abis botol facial wash nyaa nyo...

```

Gambar 1. Hasil *Preprocessing* Data

3.3 Pembagian Data

Dataset melampaui dengan menggunakan perbandingan 80:20 dengan teknik *stratified sampling*. Hasil pembagian data ditunjukkan pada Tabel 2.

Tabel 2. Pembagian Dataset

Dataset	Jumlah Data
Data Latih	589
Data Uji	148
Total	737

3.4 Hasil Klasifikasi *Logistic Regression*

Proses klasifikasi dilakukan memakai algoritma *Logistic Regression* dengan *solver* L-BFGS, pendekatan multinomial, batas iterasi maksimal 1.000, dan *class weight balanced* pada Tabel 3.

Tabel 3. Hasil Evaluasi Model *Logistic Regression*

Matriks	Nilai
Accuracy	81,41%
Precision	80,81%
Recall	81,08%
F1-Score	80,77%

Berdasarkan hasil yang disajikan pada Tabel 3, algoritma *Logistic Regression* memperoleh tingkat akurasi sebesar 81,41%. Hasil tersebut memaparkan bahwasanya cara tersebut memberikan klasifikasi yang tepat pada sebagian besar data pengujian. Nilai *precision* yang mencapai 80,81% mengindikasikan bahwasanya prediksi sentimen yang dihasilkan mempunyai tingkat ketelitian yang cukup tinggi.

Adapun nilai *F1-Score* sebesar 80,77% mencerminkan keseimbangan yang baik antara *precision* dan *recall*. Secara keseluruhan, hasil evaluasi tersebut menunjukkan bahwa model *Logistic Regression* memiliki kemampuan yang cukup andal dalam mengelompokkan sentimen pada ulasan produk *Brand Skincare XYZ*.

3.5 Validasi Model dan Perbandingan Algoritma

Dilakukan perbandingan antara akurasi pada data latih dan data uji untuk mengevaluasi kestabilan model. Hasil pengujian menunjukkan bahwa model *Logistic Regression* memperoleh akurasi sebesar 95,07% pada data latih dan 80,41% pada data uji, sehingga terdapat selisih (*gap*)

sebesar 14,66 poin persentase. Selisih ini mengindikasikan adanya kecenderungan *overfitting* ringan, yang kemungkinan besar disebabkan oleh jumlah data latih yang relatif terbatas (589 data) untuk menangani ruang fitur TF-IDF yang cukup besar, serta ketidakseimbangan jumlah data antar kelas sentimen.

Untuk memastikan bahwa performa model tidak hanya bergantung pada satu skema pembagian data tertentu, dilakukan validasi silang (*cross-validation*) dengan metode *5-Fold Stratified Cross-Validation* pada seluruh data (737 data). Hasil validasi silang ditunjukkan pada Tabel 4.

Tabel 4. Hasil *5-Fold Stratified Cross-Validation* (Logistic Regression)

Matrik	Fold 1	Fold 2	Fold 3	Fold 4	Fold 5
Akurasi	83.11 %	78.38 %	80.95 %	82.31 %	78.91 %

Penggunaan metode *k-fold cross-validation* dipilih karena mampu memberikan estimasi performa model yang lebih stabil dibandingkan pengujian tunggal, sebagaimana ditunjukkan oleh Misriati dkk. (2024) yang menerapkan pendekatan *k-fold cross-validation* untuk mengevaluasi kestabilan model klasifikasi pada studi kasus berbeda (Misriati et al., 2024). Selain itu, dilakukan pengujian tambahan menggunakan algoritma *Support Vector Machine* (SVM) dengan kernel linear pada data dan skema pembagian yang identik dengan Logistic Regression sebagai pembandingan. Perbandingan hasil kedua algoritma ditunjukkan pada Tabel 5.

Tabel 5. Perbandingan Validasi Model *Logistic Regression* dan SVM

Metrik	Logistic Regression	SVM (Linear Kernel)
Akurasi Data Latih (Train)	95,07%	96,77%
Akurasi Data Uji (Test)	80,41%	82,43%
Precision	80,08%	82,01%
Recall	80,41%	82,43%
F1-Score	80,13%	82,07%
Rata-rata 5-Fold CV	80,30%	80,43%
Standar Deviasi CV	1,93%	1,12%

Berdasarkan Tabel 5, algoritma SVM dengan kernel linear menghasilkan performa yang sedikit lebih baik dibandingkan *Logistic Regression* pada seluruh metrik evaluasi, dengan selisih akurasi data uji sebesar 2,02 poin persentase (82,43% berbanding 80,41%). Pada data latih, SVM juga mencatatkan akurasi yang lebih tinggi (96,77% berbanding 95,07%), namun dengan *gap train-test* yang relatif sebanding dengan *Logistic Regression* (14,34 poin persentase). Hasil validasi silang menunjukkan rata-rata akurasi yang hampir setara antara kedua algoritma, yaitu 80,30% ($\pm 1,93\%$) untuk *Logistic Regression* dan 80,43% ($\pm 1,12\%$) untuk SVM, dengan standar deviasi SVM yang lebih kecil menunjukkan performa yang sedikit lebih stabil di berbagai subset data. Dengan demikian, kedua algoritma menunjukkan performa yang relatif sebanding dalam mengklasifikasikan sentimen ulasan produk Brand Skincare XYZ.

Sebagai pelengkap analisis distribusi sentimen, visualisasi *wordcloud* digunakan untuk menampilkan kata-kata yang paling sering muncul pada keseluruhan ulasan produk *Brand Skincare XYZ*, sebagaimana ditunjukkan pada Gambar 2. Visualisasi *wordcloud* juga digunakan

untuk memperkuat interpretasi hasil analisis teks, sejalan dengan pendekatan (Supriyanto & Rosalin, 2023) dalam menganalisis sentimen program Merdeka Belajar melalui *text analysis wordcloud* dan *word frequency*.

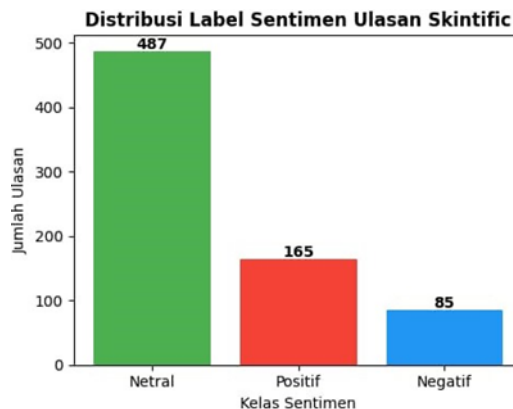


Gambar 2. Wordcloud Keseluruhan Ulasan Produk *Brand Skincare XYZ*

Berdasarkan Gambar 2, kata-kata yang paling dominan muncul antara lain “cocok”, “kulit”, “jerawat”, “skincare”, dan “moisturizer”, yang mencerminkan bahwa sebagian besar ulasan membahas kecocokan produk dengan jenis kulit serta efeknya terhadap kondisi kulit dan jerawat pengguna.

3.6 Distribusi Label Sentimen Ulasan Produk *Brand Skincare XYZ*

Proses pelabelan sentimen pada ulasan produk *Brand Skincare XYZ* menghasilkan tiga kelompok sentimen, yaitu positif, negatif, dan netral. Jumlah data yang termasuk ke dalam masing-masing kategori tersebut disajikan pada Gambar 3.



Gambar 3. Distribusi Label Sentimen Ulasan Produk *Brand Skincare XYZ*

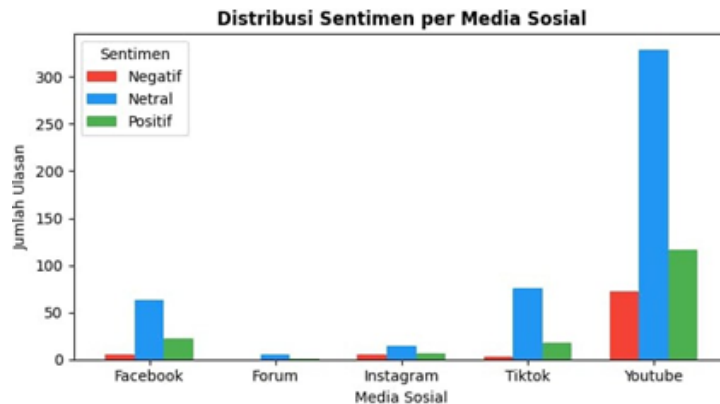
Berdasarkan Gambar 3, distribusi sentimen ulasan produk *Brand Skincare XYZ* didominasi oleh sentimen netral dengan jumlah 487 ulasan. Selanjutnya, sentimen positif berjumlah 165 ulasan, dominasi sentimen netral mengindikasikan bahwa pengguna cenderung menyampaikan pengalaman penggunaan produk secara objektif, seperti memberikan informasi mengenai kandungan, tekstur, cara penggunaan, maupun hasil yang diperoleh setelah menggunakan produk. Sementara itu, sentimen positif menunjukkan adanya kepuasan pengguna terhadap kualitas dan efektivitas produk *Brand Skincare XYZ*. Di sisi lain, sentimen negatif mencerminkan adanya keluhan atau ketidakpuasan pengguna yang dapat berkaitan dengan efektivitas produk, kecocokan pada jenis kulit tertentu, maupun pengalaman penggunaan lainnya.

Distribusi data yang tidak seimbang antara kelas netral, positif, dan negatif berpotensi mempengaruhi kinerja model klasifikasi. Kelas netral yang memiliki jumlah data paling banyak memberikan lebih banyak informasi bagi model selama proses pelatihan. Dengan itu, pada

penelitian ini digunakan parameter *class_weight = balanced* pada algoritma *Logistic Regression* untuk membantu mengurangi pengaruh ketidakseimbangan data selama proses klasifikasi.

3.7 Distribusi Sentimen Ulasan Produk *Brand Skincare XYZ* berdasarkan Media Sosial

Selain menganalisis distribusi sentimen secara keseluruhan, penelitian ini juga melaksanakan analisis berdasarkan sumber media sosial untuk mengetahui persebaran sentimen pada masing-masing *platform*, sebagaimana ditampilkan pada Gambar 4.



Gambar 4. Distribusi Sentimen Ulasan Produk *Brand Skincare XYZ* Berdasarkan Media Sosial

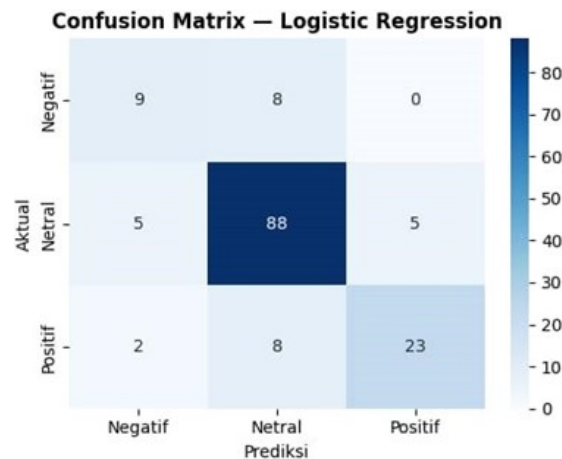
Berdasarkan Gambar 4, jumlah ulasan produk *Brand Skincare XYZ* paling banyak berasal dari *platform* YouTube dibandingkan media sosial lainnya. Pada *platform* tersebut, sentimen netral mendominasi jumlah ulasan, diikuti oleh sentimen positif dan negatif. Kondisi yang sama juga terlihat pada Facebook, TikTok, Instagram, dan Forum, di mana sentimen netral memiliki jumlah yang lebih tinggi dibandingkan sentimen positif maupun negatif.

Dominasi sentimen netral pada seluruh platform memaparkan bahwasanya pemakaian tersebut cenderung mengasih komentar yang bersifat informatif tanpa memaparkan dorongan opini yang kuat. Sementara itu, sentimen positif menunjukkan adanya kepuasan pengguna terhadap produk *Brand Skincare XYZ*, sedangkan sentimen negatif menggambarkan adanya keluhan atau pengalaman penggunaan yang kurang memuaskan.

Hasil ini menunjukkan bahwa persepsi pengguna terhadap produk *Brand Skincare XYZ* di berbagai media sosial cenderung netral hingga positif. Selain itu, tingginya jumlah ulasan pada YouTube mengindikasikan bahwa platform tersebut menjadi salah satu sumber utama diskusi dan pertukaran informasi mengenai produk *Brand Skincare XYZ* dibandingkan media sosial lainnya.

3.8 Confusion Matrix Hasil Klasifikasi *Logistic Regression*

Untuk menganalisis kemampuan model dalam mengklasifikasikan sentimen positif, negatif, dan netral secara lebih rinci, digunakan confusion matrix. Adapun hasil confusion matrix model *Logistic Regression* ditunjukkan pada Gambar 5.



Gambar 5. Confusion Matrix Hasil Klasifikasi Logistic Regression

Berdasarkan hasil visualisasi pada Gambar 5, algoritma *Logistic Regression* dapat melaksanakan prediksi sentimen dengan cukup baik pada sebagian besar data pengujian. Performa terbaik terlihat pada kategori netral, di mana model berhasil memberikan prediksi yang tepat terhadap 88 dari total 98 data. Pada kategori positif, sebanyak 23 data dari 33 data uji berhasil diklasifikasikan secara benar, sedangkan pada kategori negatif model mampu mengenali 9 dari 17 data dengan tepat. Meskipun demikian, masih ditemukan beberapa kesalahan prediksi, khususnya ketika data dengan sentimen positif maupun negatif dikategorikan sebagai sentimen netral. Hal tersebut menunjukkan bahwa model lebih optimal dalam mendeteksi ulasan netral dibandingkan dengan dua kategori lainnya. Kondisi ini kemungkinan dipengaruhi oleh jumlah data netral yang lebih dominan dalam dataset. Secara keseluruhan, hasil *confusion matrix* memperlihatkan bahwa model *Logistic Regression* memiliki kemampuan klasifikasi yang cukup efektif dengan nilai akurasi mencapai 81,08%.

4. Simpulan dan Saran

Berdasarkan penelitian yang telah dilakukan, algoritma *Logistic Regression* dapat diterapkan untuk mengklasifikasikan sentimen ulasan produk *Brand Skincare XYZ* dari media sosial ke dalam tiga kategori, yaitu positif, negatif, dan netral. Hasil evaluasi menunjukkan bahwa model mampu memberikan performa yang cukup baik dengan nilai *accuracy* sebesar 81,41%, *precision* 80,81%, *recall* 81,08%, dan *F1-score* 80,77%. Hal ini menunjukkan bahwa *Logistic Regression* dapat mengenali pola sentimen pada ulasan pengguna secara efektif.

Hasil analisis memperlihatkan bahwa sentimen netral memiliki jumlah paling dominan dibandingkan sentimen positif dan negatif. Sebagian besar pengguna cenderung memberikan ulasan berupa informasi atau pengalaman penggunaan produk tanpa menunjukkan opini yang kuat. Penelitian selanjutnya disarankan menggunakan dataset yang lebih luas, metode pelabelan yang lebih akurat, serta melakukan perbandingan dengan algoritma lain untuk memperoleh hasil klasifikasi yang lebih optimal.

Daftar Pustaka

- Albab, M. U., P., Y. K., & Fawaiq, M. N. (2023). Optimization of the Stemming Technique on Text Preprocessing President 3 Periods Topic. *Jurnal Transformatika*, 20(2), 1–12. <https://doi.org/10.26623/transformatika.v2.0i2.5374>
- Daffa, F. I., & Brata, A. D. (2025). Analisis Sentimen Ulasan Pelanggan menggunakan Algoritma Naive Bayes dan Logistic Regression. <https://doi.org/10.22441/jitkom.v9i2.003>

- Harnelia, H. (2024). ANALISIS SENTIMEN REVIEW SKINCARE SKINTIFIC DENGAN ALGORITMA SUPPORT VECTOR MACHINE (SVM). *Jurnal Informatika dan Teknik Elektro Terapan*, 12(2). <https://doi.org/10.23960/jitet.v12i2.4095>
- Hernelia. (2023). Analisis Sentimen Review Skintific-SVM [Dataset]. <https://www.kaggle.com/code/harnelia/analisis-sentimen-review-skincare-skintific-svm>
- Iskandar, J. W., & Nataliani, Y. (2021). Perbandingan Naïve Bayes, SVM, dan k-NN untuk Analisis Sentimen Gadget Berbasis Aspek. *Jurnal RESTI (Rekayasa Sistem Dan Teknologi Informasi)*, 5(6), 1120–1126. <https://doi.org/10.29207/resti.v5i6.3588>
- Khairunnisa, S., Adiwijaya, A., & Faraby, S. A. (2021). Pengaruh Text Preprocessing terhadap Analisis Sentimen Komentar Masyarakat pada Media Sosial Twitter (Studi Kasus Pandemi COVID-19). *JURNAL MEDIA INFORMATIKA BUDIDARMA*, 5(2), 406. <https://doi.org/10.30865/mib.v5i2.2835>
- Lidinillah, E. R., Rohana, T., & Juwita, A. R. (2023). Analisis sentimen twitter terhadap steam menggunakan algoritma logistic regression dan support vector machine. *TEKNOSAINS : Jurnal Sains, Teknologi dan Informatika*, 10(2), 154–164. <https://doi.org/10.37373/tekno.v10i2.440>
- Misriati, T., Aryanti, R., Sagiyanto, A., Fachri, M., & Ramadhani, A. (2024). Klasifikasi Multi Label untuk Deteksi Keseimbangan Emosi Pengguna Media Sosial Menggunakan K-Fold Cross Validation. *Journal of Information System Research (JOSH)*, 6(1), 697–704. <https://doi.org/10.47065/josh.v6i1.6033>
- Mulyono, H. S., & Saprudin, U. (2025). Efektivitas Logistic Regression dalam Analisis Sentimen Berbahasa Indonesia pada Komentar YouTube tentang Isu Ketenagakerjaan. <https://doi.org/jimik.v6i3.1481>
- Prasetyo, R. A. (2025). Perbandingan Algoritma Logistic Regression, SVM, dan Random Forest pada Analisis Sentimen Aplikasi Gopay. *Jurnal Informatika: Jurnal Pengembangan IT*, 10(4), 1176–1188. <https://doi.org/10.30591/jpit.v10i4.8796>
- Rifaldi, D., Abdul Fadlil, & Herman. (2023). Teknik Preprocessing Pada Text Mining Menggunakan Data Tweet “Mental Health.” *Decode: Jurnal Pendidikan Teknologi Informasi*, 3(2), 161–171. <https://doi.org/10.51454/decode.v3i2.131>
- Septiani, D., & Isabela, I. (2022). ANALISIS TERM FREQUENCY INVERSE DOCUMENT FREQUENCY (TF-IDF) DALAM TEMU KEMBALI INFORMASI PADA DOKUMEN TEKS. <https://journal.unj.ac.id/unj/index.php/SINTESIA/article/view/39364>
- Sikana, N., Winardi, S., -, G., Situmorang, G. F., & Lubis, R. (2025). Analisis Sentimen untuk Ulasan Produk E-Commerce Shopee Menggunakan BERT. *Jurnal Sifo Mikroskil*, 26(2). <https://doi.org/10.55601/jsm.v26i2.1796>
- Styawati, S., Hendrastuty, N., & Isnain, A. R. (2021). Analisis Sentimen Masyarakat Terhadap Program Kartu Prakerja Pada Twitter Dengan Metode Support Vector Machine. *Jurnal Informatika: Jurnal Pengembangan IT*, 6(3), 150–155. <https://doi.org/10.30591/jpit.v6i3.2870>
- Supriyanto, B. F. S., & Rosalin, S. (2023). Analisis Sentimen Program Merdeka Belajar dengan Text Analysis Wordcloud & Word Frequency. *Jurnal Minfo Polgan*, 12(1), 25–32. <https://doi.org/10.33395/jmp.v12i1.12312>
- Wati, R., Ernawati, S., & Rachmi, H. (2023). Pembobotan TF-IDF Menggunakan Naïve Bayes pada Sentimen Masyarakat Mengenai Isu Kenaikan BIPIH. *Jurnal Manajemen Informatika (JAMIKA)*, 13(1), 84–93. <https://doi.org/10.34010/jamika.v13i1.9424>