

Analisis Sentimen Pelanggan Terhadap Layanan gerai Alfamart Di Sumba Timur Menggunakan SVM

(Customer Sentiment Analysis Towards Alfamart Outlet Services in East Sumba Using SVM)

Melvianus Nggau Behar¹, Fajar Hariadi², Alfrian Carmen Talakua³

^{1,2,3}Program Studi Teknik Informatika, Universitas Kristen Wira Wacana Sumba

E-mail: ¹melvianusnggbhr@gmail.com, ²fajar@unkriswina.ac.id, ³alfriantalakua@unkriswina.ac.id

KEYWORDS:

*Sentiment Analysis, Support Vector
Machine, Alfamart*

ABSTRACT

The growth of modern retail, such as Alfamart in East Sumba Regency, has influenced the consumption patterns of the local community. Although it is highly popular, there are still varying perceptions among customers regarding the quality of service. This study aims to analyze customer sentiment towards Alfamart's services using the Support Vector Machine (SVM) algorithm. Data was collected from 219 reviews via a Google Form questionnaire, then exported to CSV format and analyzed computationally. The analysis process includes text preprocessing (cleaning, case folding, normalization, tokenizing, stopword, stemming), TF-IDF calculation, sentiment labeling based on lexicon, and classification using SVM. The evaluation results show the model's accuracy at 76.19%, with the highest F1-score for positive sentiment (83%) and neutral sentiment (76%). Meanwhile, negative sentiment only reached 44%, indicating some weakness in classifying reviews with negative tones. The majority of reviews are neutral and positive, reflecting that Alfamart's service is generally well-received by the community. The SVM model is considered effective in mapping customer perceptions and supporting data-driven decision-making.

KATA KUNCI:

Analisis Sentimen, SVM, Alfamart

ABSTRAK

Pertumbuhan ritel modern seperti Alfamart di Kabupaten Sumba Timur memengaruhi pola konsumsi masyarakat. Meski banyak diminati, masih terdapat persepsi beragam dari pelanggan terkait kualitas layanan. Penelitian ini bertujuan untuk menganalisis sentimen pelanggan terhadap layanan Alfamart menggunakan algoritma Support Vector Machine (SVM). Data dikumpulkan dari 219 ulasan melalui kuesioner Google Form, kemudian diekspor ke format CSV dan dianalisis secara komputasional. Tahapan analisis meliputi preprocessing teks (cleaning, case folding, normalisasi, tokenisasi, stopword, stemming), perhitungan TF-IDF, pelabelan sentimen berbasis lexicon, dan klasifikasi dengan SVM. Hasil evaluasi menunjukkan akurasi model sebesar 76,19%, dengan F1-score tertinggi pada sentimen positif (83%) dan netral (76%). Sementara itu, sentimen negatif hanya mencapai 44%, menandakan masih ada kelemahan dalam klasifikasi ulasan bernada negatif. Mayoritas ulasan bersifat netral dan positif, mencerminkan layanan Alfamart cukup diterima oleh masyarakat. Model SVM dinilai efektif untuk memetakan persepsi pelanggan dan mendukung pengambilan keputusan berbasis data.

PENDAHULUAN

Perkembangan teknologi informasi dan komunikasi telah mendorong perubahan signifikan dalam cara masyarakat memenuhi kebutuhan sehari-hari [1]. Pola konsumsi masyarakat kini cenderung bergeser ke arah kemudahan, efisiensi, dan kenyamanan, yang tercermin dari meningkatnya preferensi terhadap gerai ritel modern dibandingkan pasar tradisional [2]. Minimarket seperti Alfamart dan Indomaret telah menjadi bagian penting dalam kehidupan sehari-hari masyarakat Indonesia [3].

Kehadiran Alfamart di daerah terpencil seperti Kabupaten Sumba Timur, Nusa Tenggara Timur, menjadi fenomena yang menarik untuk dikaji. Sumba Timur secara geografis merupakan wilayah kepulauan dengan tingkat aksesibilitas yang terbatas, namun seiring waktu, citra dan kenyamanan pelayanan Alfamart mulai dirasakan oleh masyarakat, yang perlahan mengubah kebiasaan belanja mereka.

Dalam kajian perilaku konsumen, kepuasan pelanggan didefinisikan sebagai evaluasi menyeluruh atas pengalaman berbelanja, yang mencerminkan seberapa jauh harapan pelanggan terpenuhi [4]. Model *Expectation-Confirmation Theory* (ECT) menjelaskan bahwa ketika pelayanan aktual melebihi ekspektasi pelanggan, maka akan tercipta kepuasan. Sebaliknya, pelayanan yang tidak sesuai ekspektasi berpotensi menimbulkan ketidakpuasan [5]. Oleh karena itu, perusahaan ritel modern mulai mengalami penetrasi layanan modern, termasuk sektor ritel [6]. Hadirnya gerai Alfamart di berbagai titik strategis di Kabupaten Sumba Timur telah mengubah sebagian kebiasaan belanja masyarakat. Pergeseran ini mengindikasikan perubahan preferensi konsumsi yang tidak hanya dipengaruhi faktor harga dan jarak, tetapi juga perlu memahami secara menyeluruh bagaimana pelanggan menilai layanannya, tidak hanya dari sisi transaksi, tetapi juga dari perspektif kualitas interaksi, kelengkapan barang, harga, dan kenyamanan [7].

Berdasarkan penyebaran kuesioner kepada pelanggan Alfamart di Sumba Timur menghasilkan data opini yang sangat beragam. Banyak pelanggan memberikan respons positif seperti “menyenangkan”, “pelayanan ramah”, dan “sangat memudahkan”. Namun, terdapat pula komentar yang menunjukkan ketidakpuasan seperti “harga tidak sesuai”, “pelayanan kurang ramah”, atau “uang terpotong sia-sia”. Temuan ini menunjukkan bahwa terdapat persepsi beragam yang belum terstruktur secara sistematis, dan jika dibiarkan, dapat menghambat upaya peningkatan layanan oleh manajemen gerai Alfamart di wilayah tersebut.

Penelitian ini mengusulkan penggunaan analisis sentimen berbasis *Support Vector Machine* (SVM) untuk mengklasifikasikan ulasan pelanggan Alfamart di Kabupaten Sumba Timur secara akurat dan konsisten. Dengan memanfaatkan kemampuan klasifikasi SVM yang canggih, penelitian ini bertujuan untuk mengkategorikan opini pelanggan secara efektif ke dalam sentimen positif dan negatif. Pendekatan ini diharapkan dapat meningkatkan pemahaman tentang kepuasan pelanggan dan membantu Alfamart meningkatkan layanannya di wilayah tersebut.

Untuk melakukan analisis sentimen, dibutuhkan proses text mining, yaitu teknik untuk mengekstraksi informasi bermakna dari data teks tidak terstruktur. Proses ini mencakup beberapa tahapan penting, antara lain: pengumpulan data, preprocessing (seperti *cleaning*, *case folding*, *normalization*, *tokenizing*, *stopword removal*, dan *stemming*) [8]. serta ekstraksi fitur menggunakan metode *Term Frequency-Inverse Document Frequency* (TF-IDF), hingga tahap klasifikasi dengan menggunakan algoritma *Support Vector Machine* (SVM) [9]. Dalam penelitian ini, data dikumpulkan melalui kuesioner daring berbasis Google Form, yang memungkinkan responden memberikan ulasan terbuka terkait pengalaman mereka berbelanja di Alfamart. Seluruh data hasil pengisian kuesioner kemudian diunduh dan dikonversi ke dalam format CSV agar dapat dianalisis secara komputasional. Proses pengolahan data dilakukan menggunakan bahasa pemrograman *Python* di lingkungan *Google Collaboratory* (Colab) (data). Google Colab merupakan platform berbasis *cloud* yang menyediakan akses terhadap berbagai pustaka *Python* untuk data *analysis* dan machine learning, seperti *pandas*, *scikit-learn*, *nlTK*, dan *Sastrawi*, serta mendukung komputasi intensif tanpa memerlukan instalasi *local* [10].

Penelitian ini bertujuan untuk menganalisis sentimen pelanggan terhadap layanan gerai Alfamart di Kabupaten Sumba Timur, yang tercermin dalam ulasan yang diperoleh melalui kuesioner dengan diharapkan dapat memberikan wawasan yang berguna bagi manajemen gerai Alfamart di Kabupaten Sumba Timur untuk memahami persepsi pelanggan terhadap layanan mereka, sehingga dapat merancang perbaikan layanan yang

lebih tepat sasaran dan relevan dengan kebutuhan pelanggan, serta meningkatkan kepuasan dan loyalitas pelanggan.

METODE PENELITIAN

A. Alur Penelitian



Gambar 1. Alur Penelitian

1. Pengumpulan Data

Pengumpulan data dalam penelitian ini dilakukan melalui kuesioner daring berbasis Google Form. Data yang dikumpulkan berfokus pada sentimen pelanggan terhadap pengalaman pelanggan saat berbelanja di Alfamart Kabupaten Sumba Timur. Jumlah total data yang diperoleh adalah 400 entri data, yang dihimpun dalam rentang waktu Mei hingga Juni 2025. Seluruh data yang terkumpul dikonversi ke dalam format CSV untuk digunakan dalam tahap analisis lebih lanjut.

2. *Pre-processing* Data

Tahap ini bertujuan untuk membersihkan dan menormalkan teks agar dapat dianalisis secara efisien oleh algoritma klasifikasi. Langkah-langkah *preprocessing* meliputi:

- *Cleaning* dilakukan penghapusan karakter tidak relevan seperti angka, simbol, dan tanda baca berlebih.
- *Case Folding* adalah tahap mengubah seluruh huruf menjadi huruf kecil (*lowercase*). Tujuannya agar kata-kata yang dibentuk kapital yang berbeda tetap dikenali sebagai kata yang sama.
- *Normalization* Mengganti kata tidak baku atau kata slang ke dalam bentuk baku.
- *Tokenize* adalah proses memecah teks menjadi unit-unit kata (*token*).
- *Stopword* adalah proses menghapus kata-kata umum yang tidak memiliki makna penting seperti “dan”, “atau”, “adalah”.
- *Stemming* adalah proses Mengubah kata ke bentuk dasar.

3. Perhitungan TF-IDF

Setelah proses pembersihan selesai, teks diubah menjadi representasi numerik menggunakan metode TF-IDF. TF-IDF digunakan untuk menilai seberapa penting sebuah kata dalam dokumen relatif terhadap seluruh korpus data. Metode ini memperkuat kata-kata yang sering muncul dalam satu dokumen namun jarang muncul dalam dokumen lain, sehingga membantu meningkatkan performa klasifikasi.

4. Pelabelan *Lexicon*

Setiap opini teks diberi label sentimen secara semi-otomatis menggunakan pendekatan *lexicon-based*. Pendekatan ini dilakukan dengan membandingkan setiap kata dalam teks dengan daftar kata pada kamus sentimen khusus. Kamus sentimen ini berisi kata-kata yang telah diberi bobot sentimen, yakni positif, negatif, atau netral, yang disusun berdasarkan penelitian atau literatur sebelumnya.

5. Pembuatan Model

Dalam proses penelitian ini, data yang telah melalui tahap *pre-processing* kemudian dibagi menjadi dua bagian: data latih dan data uji. Data latih, yang terdiri dari 80% dari total data, digunakan untuk melatih model klasifikasi. Sementara itu, data uji, yang terdiri dari 20% dari total data, digunakan untuk menguji dan mengevaluasi kinerja model. Pembagian ini penting untuk memastikan bahwa

model tidak hanya mengingat data yang ada, tetapi juga dapat menggeneralisasi dengan baik terhadap data baru yang belum pernah dilihat sebelumnya.

6. Visualisasi

Pada tahap visualisasi, hasil analisis akan disajikan dalam bentuk grafik untuk mempermudah pemahaman. Grafik batang atau *pie chart* akan digunakan untuk menunjukkan distribusi sentimen positif, negatif, dan netral. word cloud akan memvisualisasikan kata-kata penting berdasarkan perhitungan TF-IDF. Selain itu, *confusion matrix* akan digunakan untuk mengevaluasi model SVM, menunjukkan sejauh mana model berhasil dalam klasifikasi sentimen.

7. Evaluasi

Setelah melatih model *Support Vector Machine* (SVM), efektivitasnya dalam mengklasifikasikan sentimen pelanggan dinilai dengan menerapkannya pada serangkaian data uji. Evaluasi ini dilakukan menggunakan matriks kebingungan, sebuah alat yang memberikan ringkasan yang jelas tentang kinerja model dengan menampilkan jumlah klasifikasi yang benar dan salah. Matriks kebingungan membantu mengidentifikasi seberapa baik model membedakan berbagai kategori sentimen.

HASIL DAN PEMBAHASAN

PENGUMPULAN DATA

Data dikumpulkan dengan menyebarkan kuesioner daring kepada penduduk Sumba Timur, memastikan beragam tanggapan dari penduduk setempat. Pertanyaan yang ada dalam kuesioner adalah “Bagaimana pengalaman Anda saat berbelanja di Alfamart”, dengan rentang waktu dari bulan Mei hingga Juni. Data yang berhasil dikumpulkan sebanyak 219 entri data. Kemudian data tersebut disimpan dalam format CSV.



Baik
Sangat menyenangkan
Menyenangkan
Senang
Bagus
Sangat baik
Cukup puas
Cukup menyenangkan
Menyenangkan

Gambar 2. Hasil Pengumpulan Data

Gambar 2 menunjukkan beberapa ulasan terkait pengalaman masyarakat berbelanja di Alfamart. Dimana data dikumpulkan menggunakan kuesioner.

PRE-PROCESSING

Setelah mengumpulkan data, langkah berikutnya melibatkan *preprocessing* untuk menyiapkan dan membersihkan informasi untuk analisis. Proses ini sangat penting untuk memastikan data yang digunakan dalam model sudah bersih, terstruktur, dan siap untuk diproses lebih lanjut. Berikut adalah langkah-langkah *preprocessing* data yang dilakukan di Google Colab:

1. Cleaning

Data yang sudah terkumpul, langkah-langkah berikutnya adalah melakukan pembersihan (*cleaning*) data untuk menghilangkan elemen-elemen yang tidak relevan atau tidak berguna dalam analisis. Proses ini dilakukan dengan cara menghapus karakter-karakter seperti angka, simbol, atau tanda baca yang tidak memberikan informasi penting dalam teks.

Tabel 1. Hasil *Cleaning*

Sebelum	Sesudah
Menarik, Apalagi Kalo Ada Diskon ??	Menarik Apalagi Kalo Ada Diskon

2. *Case Folding*

Tahap pelipatan huruf mengubah semua huruf dalam teks menjadi huruf kecil. Proses ini membantu menstandarisasi variasi antara huruf kapital dan non-kapital, memastikan konsistensi dan meningkatkan akurasi pemrosesan teks.

Tabel 2. Hasil *Case Folding*

Sebelum	Sesudah
Menarik Apalagi Kalo Ada Diskon	menarik apalagi kalo ada diskon

3. *Normalization*

Pada tahap normalisasi, ekspresi non-standar atau kata slang diubah menjadi padanannya dalam bahasa Indonesia standar. Proses ini meningkatkan kejelasan dan keakuratan teks, memastikan konten mudah dipahami dan menyampaikan makna yang diinginkan dengan tepat kepada khalayak yang lebih luas.

Tabel 3. Hasil *Normalization*

Sebelum	Sesudah
menarik apalagi kalo ada diskon	menarik apalagi kalau ada diskon

4. *Tokenize*

Tokenisasi adalah langkah pertama dan mendasar dalam pemrosesan teks. Teks ulasan pelanggan dibagi menjadi unit-unit yang lebih kecil yang disebut token, biasanya berupa kata-kata. Proses ini membantu menyederhanakan dan mengorganisasikan teks, sehingga memudahkan analisis dan pemahaman strukturnya. Dengan memecah teks menjadi bagian-bagian yang lebih mudah dikelola, tokenisasi memungkinkan pemrosesan yang lebih efektif dalam berbagai tugas pemrosesan bahasa alami.

Tabel 4. Hasil *Tokenize*

Sebelum	Sesudah
menarik apalagi kalau ada diskon	['menarik', 'apalagi', 'kalau', 'ada', 'diskon']

5. *Stopword*

Setelah proses *tokenize*, langkah selanjutnya adalah penghapusan *stopword*. Kata henti adalah kata-kata umum seperti "dan", "yang", dan "atau" yang sering muncul dalam teks tetapi seringkali hanya mengandung sedikit informasi bermakna. Kata-kata ini biasanya disaring selama pemrosesan bahasa alami untuk meningkatkan akurasi analisis.

Tabel 5. Hasil *Stopword*

Sebelum	Sesudah
['menarik', 'apalagi,kalau,ada,diskon]	[menarik, diskon]

6. *Stemming*

Setelah penghapusan *stopword*, *Stemming* adalah proses penyederhanaan kata menjadi bentuk dasar atau akar katanya untuk menyederhanakan analisis teks. Misalnya, kata-kata seperti “berjalan”, “berlari”, dan “melompat” dapat diubah menjadi bentuk dasarnya, yaitu “berjalan”, “berlari”, dan “melompat”. Teknik ini membantu meningkatkan akurasi pencarian dan efisiensi pemrosesan teks.

Tabel 6. Hasil *Stemming*

Sebelum	Sesudah
[menarik, diskon]	[tarik, diskon]

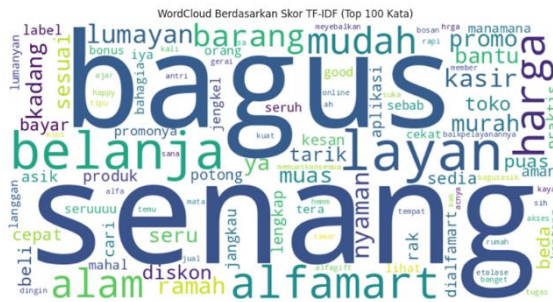
PERHITUNGAN TF-IDF

Setelah pra-pemrosesan awal teks, metode *Term Frequency-Inverse Document Frequency* (TF-IDF) diterapkan untuk memberikan bobot pada kata. Teknik ini mempertimbangkan seberapa sering sebuah kata muncul dalam satu dokumen, serta seberapa jarang kata tersebut muncul di seluruh koleksi dokumen. Dengan demikian, TF-IDF secara efektif menyoroti istilah-istilah yang sangat penting atau khas dalam tinjauan tertentu, sekaligus mengabaikan kata-kata umum yang sering muncul di banyak dokumen. Hasil perhitungan TF-IDF dapat dilihat pada tabel ini:

Tabel 7. Hasil Perhitungan TF-IDF

kata	Hasil TF-IDF
tarik	0.731
diskon	0.467
nyaman	1
aman	1

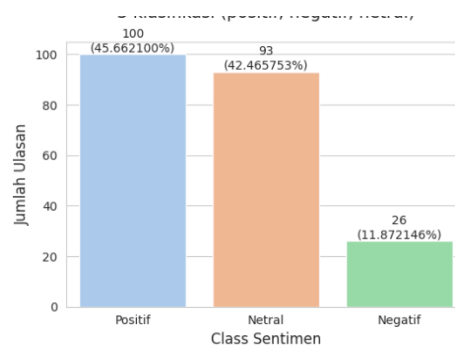
Adapun *wordcloud* dari hasil perhitungan TF-IDF, dapat dilihat pada gambar ini:

Gambar 3. *Wordcloud* Hasil Perhitungan TF-IDF

Pada gambar 3 menunjukkan kata-kata yang muncul dengan ukuran yang lebih besar memiliki nilai TF-IDF yang lebih tinggi. Dalam hal ini, kata “bagus” dan “senang” tampak lebih dominan, yang menunjukkan bahwa kata-kata tersebut memiliki bobot atau relevansi yang lebih besar berdasarkan perhitungan TF-IDF. *Wordcloud* ini memberikan visualisasi dari hasil perhitungan TF-IDF, di mana kata-kata yang lebih jarang muncul dalam keseluruhan dokumen akan memiliki nilai TF-IDF yang lebih tinggi dan tampil lebih besar di dalam *cloud*.

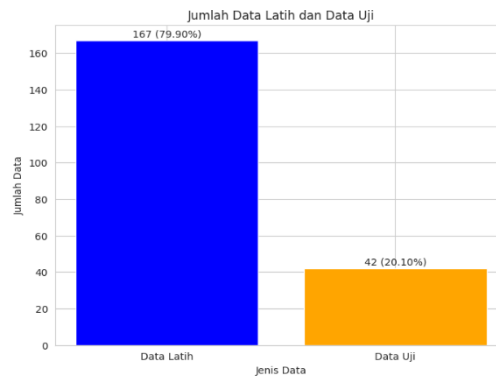
PELABELAN *LEXICON*

Pada tahap pelabelan leksikon, setiap kata dalam ulasan diberi label sentimen berdasarkan kamus yang telah ditentukan. Proses ini memungkinkan identifikasi sentimen keseluruhan yang diungkapkan dalam teks. Analisis data berlabel menunjukkan bahwa mayoritas ulasan bersifat positif, yaitu sebesar 45,66%, netral, yaitu sebesar 42,46% dan negatif sebesar 11,87%. Temuan ini menunjukkan bahwa sebagian besar pengguna cenderung memiliki opini yang seimbang atau positif, yang menunjukkan penerimaan yang umumnya positif terhadap subjek yang diulas.

Gambar 4. Hasil Pelabelan *Lexicon*

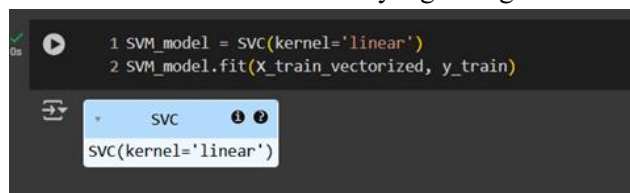
PEMBUATAN MODEL SVM

Setelah data melalui tahap pelabelan sentimen, langkah selanjutnya adalah pembuatan model klasifikasi menggunakan algoritma *Support Vector Machine* (SVM). Pada tahap ini, data yang telah diproses dan diberi label sentimen dibagi menjadi dua bagian, yaitu data latih sebesar 80% dan data uji sebesar 20% dari keseluruhan data. Pembagian ini bertujuan untuk memastikan bahwa model dapat belajar dari sebagian besar data yang tersedia (data latih), serta dapat diuji performanya terhadap data yang belum pernah dilihat sebelumnya (data uji), guna mengevaluasi kemampuan generalisasi model.



Gambar 5. Pembagian Data Latih Dan Data Uji

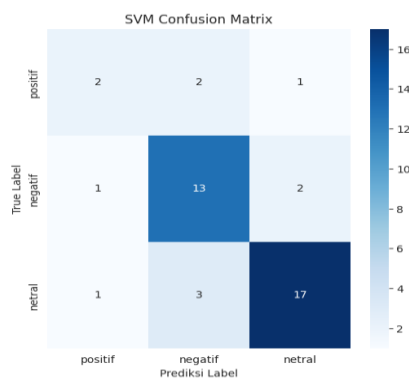
Setelah data diubah kedalam bentuk vektor, tahap berikutnya adalah melatih model SVM. Model ini dibuat menggunakan SVC (*Support Vector Classification*) dari pustaka *sklearn.svm* dengan parameter *kernel* = “linear”. *Kernel* linear dipilih karena sesuai untuk data teks yang sering kali bersifat linear *separable*.



Gambar 6. Model SVM

VISUALISASI

Setelah model *Support Vector Machine* (SVM) dibangun dan diuji, langkah selanjutnya adalah visualisasi hasil analisis untuk mempermudah pemahaman dan interpretasi. Hasil analisis diubah menjadi *confusion matriks*, untuk menunjukkan distribusi sentimen positif, negatif, dan netral dalam ulasan pelanggan.

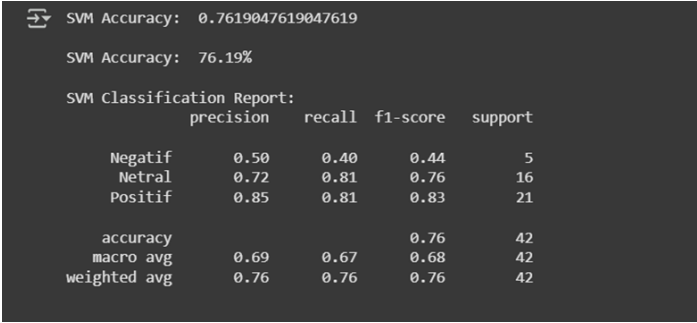


Gambar 7. Confusion Matriks

EVALUASI

Setelah melatih model *Support Vector Machine* (SVM), efektivitasnya dalam mengklasifikasikan sentimen pelanggan dievaluasi menggunakan data uji terpisah. Hasil evaluasi menunjukkan bahwa model memiliki *Accuracy* sebesar 76,19%, yang berarti model ini berhasil memprediksi dengan tepat 76,19% dari keseluruhan data uji. Untuk metrik *Precision*, kategori sentimen positif memiliki nilai 0.85, yang menunjukkan bahwa 85% dari ulasan yang diprediksi sebagai positif benar-benar bersifat positif. Sementara itu, *Recall* untuk sentimen positif tercatat 0.81, yang berarti model berhasil mengidentifikasi 81% dari semua ulasan positif yang

ada. *F1-score* untuk sentimen positif, yang menggabungkan *Precision* dan *Recall*, mencapai 0.83, menunjukkan keseimbangan yang baik antara kedua metrik tersebut. Di sisi lain, untuk sentimen netral, model memiliki *Precision* sebesar 0.72%, *Recall* sebesar 0.81%, dan *F1-score* sebesar 0.76%, yang menunjukkan kinerja yang cukup baik. Namun, untuk sentimen negatif, model menunjukkan kinerja yang lebih rendah dengan *Precision* sebesar 0.50%, *Recall* sebesar 0.40%, dan *F1-score* sebesar 0.44%, menandakan kesulitan dalam mengklasifikasikan ulasan dengan sentimen negatif. Meskipun model efektif dalam mengidentifikasi sentimen positif dan netral, perlu ada perbaikan pada klasifikasi sentimen negatif untuk meningkatkan kinerjanya. Matriks kebingungan (*confusion matrix*) yang disertakan juga memberikan wawasan lebih lanjut tentang distribusi dan kesalahan klasifikasi pada setiap kategori sentimen.



```

SVM Accuracy: 0.7619047619047619
SVM Accuracy: 76.19%
SVM Classification Report:
      precision    recall  f1-score   support

Negatif      0.50      0.40      0.44         5
Netral       0.72      0.81      0.76        16
Positif      0.85      0.81      0.83        21

accuracy          0.76          0.76          0.76          42
macro avg         0.69          0.67          0.68          42
weighted avg      0.76          0.76          0.76          42
  
```

Gambar 8. Hasil *Accuracy*, *Precision*, *Recall*, Dan *F1-score*

KESIMPULAN DAN SARAN

Studi ini menganalisis 219 ulasan pelanggan gerai Alfamart di Kabupaten Sumba Timur dengan menggunakan algoritma *Support Vector Machine* (SVM). Analisis ini bertujuan untuk mengklasifikasikan dan memahami sentimen pelanggan secara akurat. Model ini mencapai tingkat akurasi 76,19%, yang menunjukkan efektivitasnya dalam memproses dan menginterpretasikan umpan balik pelanggan dalam konteks regional ini. Namun, performa model untuk sentimen negatif masih rendah, yaitu 44%. Hal ini menunjukkan bahwa model kesulitan dalam mendeteksi opini negatif. Dari 219 ulasan yang dianalisis, 100 ulasan (45,66%) bersifat positif, 93 ulasan (42,46%) bersifat netral, dan 26 ulasan (11,87%) bersifat negatif. Secara keseluruhan, dapat disimpulkan bahwa persepsi pelanggan terhadap layanan Alfamart di Sumba Timur cenderung netral hingga positif, dan meskipun model SVM efektif, perlu adanya peningkatan pada klasifikasi sentimen negatif.

Penelitian selanjutnya disarankan untuk menambah data, guna memperbaiki distribusi sentimen dan meningkatkan akurasi model. Penggunaan metode klasifikasi lain seperti *Random Forest* atau *Naive Bayes* juga dapat diperbandingkan. Selain itu, peningkatan kualitas kamus *lexicon* dan pemanfaatan data dari media sosial dapat memperkaya hasil analisis.

DAFTAR PUSTAKA

- [1] Wiryany D, Natasha S, Kurniawan R. 2022. Perkembangan Teknologi Informasi dan Komunikasi Terhadap Perubahan Sistem Komunikasi Indonesia. *Jurnal Nomosleca*. 8(2): 242–252, doi: 10.26905/nomosleca.v8i2.8821.
- [2] Saputra Y, Rosihan R. I, Spalanzani W, Kumalasari R, Riyanti H. 2022. Analisis Perilaku Konsumen Dalam Memutuskan Minimarket Sebagai Tempat Berbelanja. *Jurnal Rekavasi*. 10(1): 45–55, doi: 10.34151/rekavasi.v10i1.3880.
- [3] Palilu A. 2022. Analysis of the Impact of and Indomaret Minimarket Hilirization for Community Economy and Traditional Markets in Sorong City. *JURNAL JENDELA ILMU*. 3(2): 46–51.
- [4] Sari T. K, Akhmad K. A, Rahmawati E. D. 2024. Pengaruh Kualitas Pelayanan, Kelengkapan Produk Dan Harga Produk Terhadap Kepuasan Konsumen Pada ‘Mitra Swalayan’ Kartasura Swalayan sebagai

- Tujuan Perbelanjaan pada Bisnis Ritel. *Journal of Management and Creative Business*.. 2(3): 324-340.
- [5] AlSokkar A. A. M, Law E. L, AlMajali D. A, AlGasawneh J. A, AlShinwan M. 2024. An Indexed Approach for Expectation-Confirmation Theory: A Trust-Based Model. *Electronic Markets*. 34(12): 1-17, doi: 10.1007/s12525-024-00694-3.
- [6] Meliana D, Riswati J, Astuti D. 2025. Analisis Perkembangan Bisnis Ritel di Indonesia. *Journal of Business Economics and Management*. 1(3): 235–243.
- [7] Sitepu M. B, Munthe I. R, Harahap S. Z. 2022. Implementation of Support Vector Machine Algorithm for Shopee Customer Sentiment Analysis Sink. *Jurnal dan Penelitian Teknik Informatika*. 7(2): 619–627.
- [8] Bijaksana A. P. N. 2023. The Influence of Applying Stopword Removal and SMOTE on Indonesian Sentiment Classification. *LONTAR KOMPUTER*. 14(3): 1–5, doi: 10.24843/LKJITI.2023.v14.i03.p05.
- [9] Rifaldy F, Sibaroni Y, Prasetyowati S. S. 2025. Effectiveness of Word2Vec and TF-IDF in Sentiment Classification on Online Investment Platforms Using Support Vector Machine. *JUPI (Jurnal Ilmiah Penelitian dan Pembelajaran Informatika)*. 10(2): 863–874.
- [10] Agstriningtyas A. S, Pratama A. Y, Roso K, Krismahardi A. 2024. Visualisasi Data Kesadaran dan Penerapan Ekonomi Sirkular di Kota Malang Menggunakan Python dan Google Colab. *Jurnal Penelitian Inovatif (JUPIN)*. 5(1): 285–294, doi: 10.54082/jupin.963.