

Klasifikasi Emosi Ulasan Produk *E-Commerce* Menggunakan *Support Vector Machine*

Natasha Isnaeni Raharko¹, Yuni Yamasari²

^{1,2} Teknik Informatika, Fakultas Teknik, Universitas Negeri Surabaya

¹natasha.19040@mhs.unesa.ac.id

²yuniyamasari@unesa.ac.id

Abstrak— Ulasan pada e-commerce dapat mempengaruhi keputusan pembeli dan menunjukkan emosi pembeli terhadap produk yang mereka beli. Namun, analisis terhadap ulasan dalam e-commerce tersebut tidak mudah jika dilakukan secara manual. Oleh karena itu, penelitian ini memfokuskan domain tersebut dengan Support Vector Machine (SVM). Lebih lanjut, penelitian ini mengeksplorasi kernel dari SVM, yaitu: linear, polynomial, RBF, dan sigmoid. Nilai evaluasi tertinggi yang SVM tanpa data *resampling* untuk *accuracy* diraih oleh SVM kernel *Linear* dan *Sigmoid* sebesar 66.30. *Precision*, *recall*, dan *f1-score* diraih oleh SVM kernel *Sigmoid* dengan masing-masing 66.00, 62.12, dan 62.91. Untuk data *resampling*, *accuracy* dan *precision* tertinggi dicapai oleh SVM kernel RBF dengan nilai 61.97 dan 62.04. Sedangkan, nilai *recall* dan *f1-score* dengan data *resampling* diraih oleh SVM kernel *Linear* dengan nilai 65.19 dan 64.46. Berdasarkan kinerja secara keseluruhan, *accuracy* tertinggi diraih oleh SVM kernel *Linear* dan *Sigmoid*, yaitu 66,30, *precision* tertinggi oleh SVM kernel RBF, yaitu 67,25, *recall* tertinggi oleh SVM kernel *Sigmoid*, yaitu 62,12, dan *f1-score* tertinggi oleh SVM kernel *Sigmoid*, yaitu 62,91.

Kata Kunci— Klasifikasi, Support Vector Machine, kernel, emosi, ulasan

I. PENDAHULUAN

Menurut Laudon, *E-Commerce* atau *electronic commerce* merupakan suatu proses membeli dan menjual produk-produk secara elektronik oleh konsumen dan dari perusahaan ke perusahaan dengan komputer sebagai perantara bisnis [1]. Pada umumnya, konsumen yang telah membeli produk akan diminta untuk memberi ulasan terkait produk tersebut. Konsumen akan menyatakan opini mereka terkait produk yang didapatkan. Isi ulasan tersebut dapat mempengaruhi keputusan pembelian konsumen lainnya [2]. Kepastian produk dan ulasan yang ada dapat mencerminkan level kepuasan seorang konsumen tentang produk yang mereka beli [3]. Seorang konsumen yang puas juga akan meninggalkan ulasan yang positif dan memengaruhi keputusan pembelian konsumen lainnya. Maka dari itu, penjual perlu untuk meninjau ulasan-ulasan yang ada untuk dapat memajukan usaha mereka pada e-commerce tersebut.

Salah satu hal yang dapat diperhatikan dalam meninjau ulasan adalah emosi ulasan. Ulasan yang dibuat konsumen pasti melibatkan emosi mereka terkait produk yang dibeli. Namun, berdasarkan data Statista Market Insights, jumlah pengguna e-commerce di Indonesia mencapai 178,94 juta orang pada 2022 dan diproyeksikan akan mencapai 196,47

juta pengguna hingga akhir 2023 [4]. Semakin besar pengguna suatu e-commerce, maka semakin banyak pula ulasan yang ada. Dapat disimpulkan bahwa jumlah ulasan yang ada di e-commerce tidak memungkinkan peninjauan secara manual.

Data Mining adalah sebuah proses yang bertujuan untuk mendapatkan informasi dari data dan menyajikan informasi tersebut kepada pengguna secara komprehensif [5]. Salah satu pengembangan dari *data mining* adalah *Text Mining*, sebuah pengembangan data mining yang bertujuan untuk menemukan informasi yang berguna dari data-data berupa tulisan, dokumen, atau teks dalam bentuk klasifikasi, penentuan, maupun clustering [6]. Salah satu aplikasi text mining dalam cabang klasifikasi adalah klasifikasi emosi. Klasifikasi emosi mengacu pada tugas mengenali emosi seseorang dan mengklasifikasikannya menurut reaksi dan respons mereka. Klasifikasi emosi didefinisikan sebagai pengkategorian emosi dan upaya membedakan satu emosi dengan emosi lainnya [7]. Metode tersebut dapat diaplikasikan pada ulasan pada e-commerce untuk mendapatkan informasi kepuasan konsumen secara cepat.

Untuk mendapatkan hasil klasifikasi emosi yang baik, dibutuhkan penentuan algoritma yang mampu menjalankan tugas tersebut dengan baik. Algoritma yang kerap digunakan dalam melakukan tugas klasifikasi emosi adalah *Support Vector Machine* (SVM). Dalam [8], dikatakan bahwa SVM efektif untuk melakukan beberapa tugas seperti pengkategorian artikel berita dan prediksi sentimen karena mampu mengatasi dimensionalitas tinggi dengan baik [9]serta mengenali *hyperplane* yang membagi batas diantara dua kelas. Selain itu, dikatakan juga bahwa SVM juga tidak rentan overfitting. Pada [10], SVM mencapai akurasi dengan rerata 97.01, daripada NB yang hanya memiliki rerata akurasi 90.86, menjadikan SVM lebih baik dalam tugas klasifikasi jalur minat SMA daripada NB. Pada [11], SVM dan *Logistic Regression* yang digabungkan pun memiliki nilai akurasi tertinggi diantara metode-metode lainnya. Pada [12], SVM dengan seleksi fitur *Chi-Square* dan *stemming* mencapai nilai *accuracy* 80%, *precision* 95, *recall* 95, dan *f1-score* 95,05.

Peninjauan ulasan produk pada e-commerce dalam skala besar perlu dilakukan klasifikasi emosi sebagai bagian dari proses text mining untuk mendapatkan informasi kepuasan konsumen dengan cepat. Namun, algoritma yang dipilih untuk tugas tersebut juga harus dapat melakukan tugas dengan baik. Maka dari itu, untuk menemukan algoritma

yang optimal, akan dilakukan perbandingan performa algoritma *Support Vector Machine* (SVM) dalam melakukan tugas klasifikasi emosi pada data ulasan produk pada *e-commerce*. *Support Vector Machine* dipilih karena keunggulannya dalam mengatasi data dengan dimensionalitas tinggi serta tidak rentan overfitting.

II. METODE PENELITIAN

Penelitian ini menggunakan metode *Support Vector Machine*. Tahapan-tahapan dalam penelitian digambarkan pada Gbr.1.



Gbr. 1 Tahapan Metode Penelitian

A. Perolehan Dataset

Masalah yang berusaha diselesaikan dalam penelitian ini adalah fakta bahwa ulasan pembeli lainnya dapat mempengaruhi keputusan pembeli lainnya untuk membeli suatu produk dan mencerminkan level kepuasan pembeli tersebut terhadap produk. Karena itulah, menggunakan teknik *Natural Language Processing* (NLP), bermacam pengklasifikasi dibuat untuk mencoba memisahkan ulasan produk dalam *e-commerce* dengan memperhatikan emosi dari pembeli tersebut serta mencari algoritma yang lebih unggul diantara *Support Vector Machine* (SVM). Penelitian ini juga merupakan kontribusi dalam bidang *Natural Language Processing* (NLP) bahasa Indonesia.

B. Studi Literatur

Studi literatur untuk penelitian yang dilakukan adalah topik-topik yang terkait dengan *e-commerce*, emosi, data mining, text mining, langkah-langkah untuk melakukan klasifikasi, algoritma *Support Vector Machine* (SVM), *confusion matrix*, dan metrik pengukur performa model.

C. Analisis Kebutuhan

Berikut merupakan berbagai kebutuhan yang dibutuhkan untuk keberlangsung penelitian ini, yaitu:

1) Kebutuhan Perangkat Keras (*Hardware*)

Perangkat keras yang digunakan dalam penelitian guna mewujudkan tujuan penelitian yaitu perangkat laptop sebagai uji coba dengan spesifikasi berikut :

Prosesor : Intel(R) Core(TM) i7-9750H
CPU @ 2.60GHz (12 CPUs), ~2.6GHz
RAM : 16.00 GB

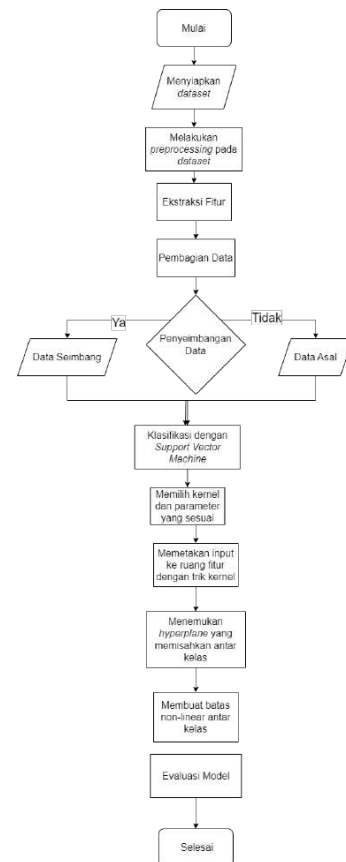
2) Kebutuhan Perangkat Lunak (*Software*)

Sistem Operasi : Windows 10 Enterprise 64-bit
Software : Python
Opera GX 104.0.4944.85

D. Perancangan Sistem

Dalam setiap penelitian, diperlukan perancangan sistem yang baik agar misi dan visi penelitian dapat tercapai dengan baik. Perencanaan sistem ini dapat disajikan dalam berbagai bentuk, seperti *flowchart*, tabel, dan sebagainya untuk mempermudah pembaca memahami cara sistem dalam penelitian bekerja.

Dalam penelitian ini, sistem yang akan dirancang digambarkan oleh Gbr. 2, yang terdiri dari perolehan dataset, pra-pemrosesan dataset, ekstraksi fitur, pembagian data menjadi data latih dan uji, penyeimbangan dataset, pelatihan model, pengujian model, dan evaluasi model.



Gbr. 2 Tahapan Pembuatan Model Klasifikasi

Berdasarkan Gbr.2, rincian perancangan sistem adalah sebagai berikut:

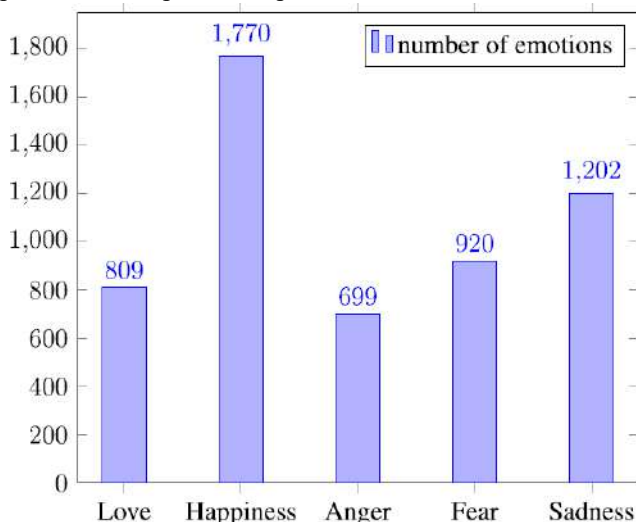
E. Perolehan Dataset

Dataset yang digunakan dalam penelitian ini adalah dataset dari artikel “PRDECT-ID Indonesian Product Reviews Dataset for Emotions Classification Tasks” oleh [13]. Dataset ini memiliki 5400 ulasan produk yang tersebar di 29 kategori dari *e-commerce* Tokopedia dan telah dilakukan anotasi emosi oleh seorang ahli psikologi klinis sehingga siap digunakan untuk tugas klasifikasi emosi. Rincian dataset digambarkan pada Gbr. 3.

Attribute	Description
Category	Product classification by category
Product Name	Name of the reviewed product
Location	City name of the shop or product seller
Price	Price in IDR of the reviewed product
Overall Rating	Overall product rating
Number Sold	Total number of products sold
Total Reviews	Total number of reviews given by the customers
Customer Rating	Product rating (range 1 to 5) from the customers
Customer Review	Product reviews given to the product by the customers
Sentiment	Sentiment labels (i.e., Positive, Negative)
Emotion	Emotion labels (i.e., Anger, Fear, Happy, Love, Sadness)

Gbr. 3 Rincian Dataset

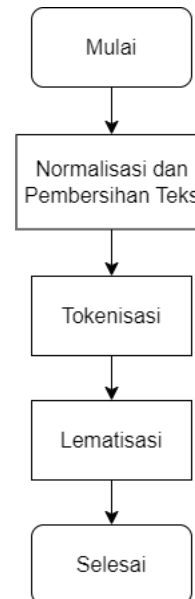
Untuk melakukan klasifikasi emosi, tidak semua atribut akan digunakan. Atribut yang digunakan dari dataset tersebut adalah atribut ‘Customer Review’ dan ‘Emotion’. ‘Customer Review’ merupakan atribut yang berisi ulasan produk yang diberikan oleh pembeli produk, sedangkan ‘Emotion’ merupakan label emosi yang terdiri dari *Anger*, *Fear*, *Happy*, *Love*, dan *Sadness*. Atribut ‘Customer Review’ akan diklasifikasikan menurut atribut ‘Emotion’. Distribusi emosi pada dataset dapat dilihat pada Gbr. 4.



Gbr. 4 Distribusi Isi Kelas dalam Set Data

Sebelum diolah lebih lanjut, serangkaian proses akan dilakukan untuk menyiapkan dataset sehingga pengembangan model dapat dilakukan.

F. Normalisasi dan Pembersihan Teks



Gbr. 5 Tahapan Pra-Pemrosesan Set Data

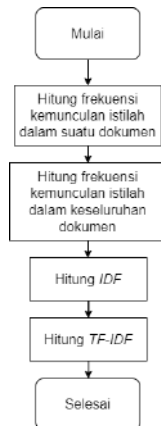
Proses-proses yang digambarkan pada Gbr. 5 memiliki rincian sebagai berikut:

- 1) *Normalisasi dan Pembersihan Teks*: Menyetarakan istilah-istilah dalam teks agar sederajat dengan mengubah istilah-istilah menjadi huruf kecil, menghapus tanda baca, angka, simbol-simbol tidak relevan, serta spasi ekstra.
- 2) *Tokenisasi*: Memilih dan memisahkan unit teks baik dalam bentuk satu kata, kumpulan kata, ataupun per kalimat.
- 3) *Lematisasi*: Menghapus imbuhan istilah-istilah dalam teks dan memastikan hasil akhirnya adalah bentuk dasar dari istilah tersebut. Teknik ini bertujuan untuk mengurangi token unik yang ada dalam dataset.

Proses-proses tersebut hanya akan diberlakukan pada atribut ‘Customer Review’. Hasil akhir dari proses ini adalah token teks yang bersih dari tanda baca, angka, simbol-simbol, dan spasi serta telah berubah ke bentuk awal token teks tersebut.

G. Ekstraksi Fitur

Setelah dilakukan pra-pemrosesan dataset, fitur-fitur dalam dataset tersebut akan direpresentasikan dalam bentuk vektor. Teknik TF-IDF akan digunakan untuk merepresentasikan atribut ‘Customer Review’ terhadap keseluruhan dokumen dengan langkah-langkah yang digambarkan oleh Gbr 6.



Gbr. 6 Tahapan Ekstraksi Fitur

Tahapan ekstraksi fitur menggunakan TF-IDF adalah sebagai berikut:

- 1) Menghitung *Term Frequency (TF)*: Menghitung frekuensi kemunculan suatu istilah atribut 'Customer Review' dalam satu dokumen

$$tf_{i,j} = \frac{f_{i,j}}{\sum_{i \in j} f_{i,j}} \quad (1)$$

Dengan keterangan bahwa i adalah kata, j adalah dokumen. Frekuensi kata dalam suatu dokumen tersebut dibagi oleh jumlah kata yang ada dokumen tersebut.

- 2) Menghitung frekuensi kemunculan suatu istilah atribut 'Customer Review' dalam keseluruhan dokumen. Hal ini disebut *Document Frequency (DF)*.
- 3) Menghitung Inverse Document Frequency suatu istilah

$$idf_{i,j} = \log \frac{|N_j|}{|DF_{i,j}|} \quad (2)$$

Dimana N adalah keseluruhan kata dalam dokumen sedangkan $DF_{i,j}$ adalah *document frequency* kata i dalam dokumen besar j .

- 4) Menghitung TF-IDF.

Hasil akhir dari proses ini adalah representasi tiap istilah yang ada pada atribut 'Customer Review' dalam bentuk vektor numerik. Nilainya dihitung dengan mengalikan TF dengan IDF.

H. Pembagian Dataset

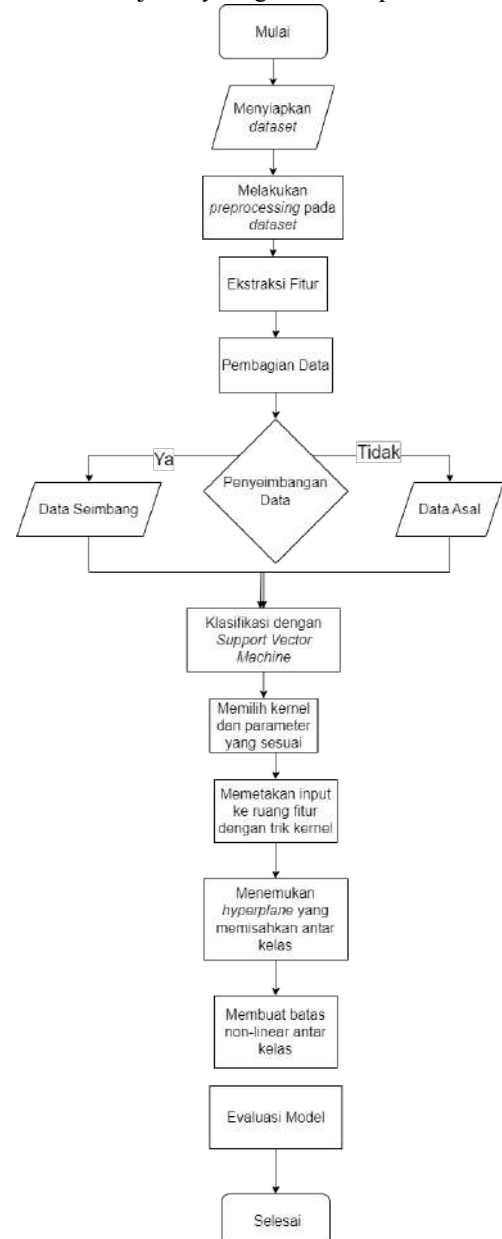
Selanjutnya, dataset akan dibagi menjadi sepuluh kombinasi data latih dan data uji untuk melatih dan menguji performa kedua algoritma menggunakan kombinasi-kombinasi data tersebut. Sepuluh kombinasi yang dimaksudkan adalah 0.05: 0.95, 0.10:0.90, 0.15:0.85, 0.20:0.80, 0.25:0.75, 0.30:0.70, 0.35:0.65, 0.40:0.60, 0.45:0.55, dan 0.50:0.50 dari data yang ada.

I. Penyeimbangan Dataset

Sebagaimana yang digambarkan oleh Gbr. 5, distribusi data dalam dataset tidak merata sehingga dapat mempengaruhi algoritma klasifikasi untuk bias terhadap label tertentu. Untuk menanggulangi hal ini, akan dilakukan penyeimbangan data pada atribut 'Customer Review' menggunakan metode resampling pada data latih sebelum melakukan pelatihan model.

J. Pelatihan Model

Langkah ini lebih jelasnya digambarkan pada Gbr. 7.



Gbr. 7 Tahapan Pembuatan Model Klasifikasi

Langkah-langkah pelatihan model adalah sebagai berikut:

- 1) Langkah pertama adalah menyiapkan dataset yang akan melalui pra-pemrosesan untuk langkah-langkah selanjutnya

- 2) Selanjutnya, dataset akan melalui berbagai proses pra-pemrosesan sebelum melalui proses ekstraksi fitur
- 3) Dataset akan melalui proses ekstraksi fitur
- 4) Selanjutnya, dataset akan dibagi menjadi data latih dan data uji. Data latih digunakan untuk melatih algoritma, sedangkan data uji digunakan untuk menguji ketepatan algoritma dalam melakukan klasifikasi
- 5) Langkah selanjutnya adalah penyeimbangan data. Akan terdapat dua data yang akan digunakan oleh algoritma, data dengan jumlah data yang seimbang diantara kelas dan data dengan jumlah yang tidak seimbang. Langkah ini dilakukan untuk mengeksplorasi performa algoritma dalam melakukan klasifikasi dataset tersebut.
- 6) Dataset kemudian akan dilakukan untuk melatih algoritma. Support Vector Machine memiliki alur kerja sebagai berikut:
 1. Langkah pertama adalah menyiapkan dataset hasil pra-pemrosesan untuk digunakan sebagai bahan pelatihan SVM.
 2. Langkah selanjutnya adalah untuk memilih kernel dan parameter yang sesuai untuk merepresentasikan dataset ke ruang fitur dan melakukan klasifikasi.
 3. Dengan bantuan trik kernel, SVM merepresentasikan fitur dari dataset ke ruang fitur.
 4. SVM berusaha menemukan ruang yang paling sesuai untuk memisahkan data pada masing-masing kelas.
 5. Setelah ruang tersebut ditemukan, SVM membuat batas non-linear pada dimensi rendah sebagai akhir dari klasifikasi.

Alhamdulillah berfungsi dengan baik. Packaging aman. Respon cepat dan ramah. Seller dan kurir amanah	alhamdulillah berfungsi dengan baik packaging aman respon cepat dan ramah seller dan kurir amanah
Ini pembelian ke 5 krupuknya mantap, pengirimannya cepat dan tidak ada yang rusak. Terima kasih kepada penjual dan kurir sicepat	ini pembelian ke krupuknya mantap pengirimannya cepat dan tidak ada yang rusak terima kasih kepada penjual dan kurir sicepat

TABEL II
HASIL TOKENISASI TEKS

Data Normalisasi	Hasil Tokenisasi
alhamdulillah berfungsi dengan baik packaging aman respon cepat dan ramah seller dan kurir amanah	['alhamdulillah', 'berfungsi', 'dengan', 'baik', 'packaging', 'aman', 'respon', 'cepat', 'dan', 'ramah', 'seller', 'dan', 'kurir', 'amanah']
ini pembelian ke krupuknya mantap pengirimannya cepat dan tidak ada yang rusak terima kasih kepada penjual dan kurir sicepat	['ini', 'pembelian', 'ke', 'krupuknya', 'mantap', 'pengirimannya', 'cepat', 'dan', 'tidak', 'ada', 'yang', 'rusak', 'terima', 'kasih', 'kepada', 'penjual', 'dan', 'kurir', 'sicepat']

TABEL III
UKURAN HASIL LEMATISASI TEKS

Data Tokenisasi	Hasil Lematisasi
['alhamdulillah', 'berfungsi', 'dengan', 'baik', 'packaging', 'aman', 'respon', 'cepat', 'dan', 'ramah', 'seller', 'dan', 'kurir', 'amanah']	['alhamdulillah', 'fungsi', 'dengan', 'baik', 'packaging', 'aman', 'respon', 'cepat', 'dan', 'ramah', 'seller', 'dan', 'kurir', 'amanah']
['ini', 'pembelian', 'ke', 'krupuknya', 'mantap', 'pengirimannya', 'cepat', 'dan', 'tidak', 'ada', 'yang', 'rusak', 'terima', 'kasih', 'kepada', 'penjual', 'dan', 'kurir', 'sicepat']	['ini', 'beli', 'ke', 'krupuknya', 'mantap', 'kurir', 'cepat', 'dan', 'tidak', 'ada', 'yang', 'rusak', 'terima', 'kasih', 'kepada', 'jual', 'dan', 'kurir', 'sicepat']

- 3) Ekstraksi Fitur: mengubah kata menjadi vektor di ruang dimensi tinggi. Teknik ini diperlihatkan prosesnya oleh Tabel IV, V, VI, VII, dan VIII.

TABEL IV
CONTOH DATA PRA-PEMROSESAN

	Data Mentah	Data Olah
D1	barang bagus, berfungsi dengan baik, seller ramah, pengiriman cepat	['barang', 'bagus', 'fungsi', 'dengan', 'baik', 'seller', 'ramah', 'kurir', 'cepat']
D2	bagus sesuai harapan penjual nya juga ramah. trimakasih pelapak ??	['bagus', 'suai', 'harap', 'jual', 'nya', 'juga', 'ramah', 'trimakasih', 'pelapak']

TABEL I
UKURAN FONT UNTUK MAKALAH

Data Mentah	Hasil Konversi
-------------	----------------

K. Evaluasi Model

Setelah melakukan pelatihan kedua algoritma, performa tiap algoritma akan dievaluasi. Performa model akan dievaluasi untuk melihat nilai-nilai dalam confusion matrix, yaitu accuracy, precision, recall, dan f1-score. Evaluasi model akan dilakukan dengan membagi dataset menjadi sepuluh bagian dan mengevaluasi model masing-masing untuk menemukan model dengan performa yang terbaik.

III. HASIL DAN PEMBAHASAN

A. Implementasi

- 1) Pengambilan *dataset*: dari keseluruhan *dataset*, diambil kolom 'Customer Review' dan 'Emotion'.
- 2) Pra-pemrosesan *Dataset*: Dilakukan normalisasi dan pembersihan teks, tokenisasi, serta lematisasi pada *dataset*. Pra-pemrosesan ditunjukkan oleh Tabel I, II, dan III.

D3	barang bagus dan respon cepat, harga bersaing dengan yg lain.	['barang', 'bagus', 'dan', 'respon', 'cepat', 'harga', 'saing', 'dengan', 'yg', 'lain']
----	---	---

TABEL V
HASIL PERHITUNGAN TERM-FREQUENCY 1

Kata	Term Frequency (TF)		
	D1	D2	D3
Barang	1	0	1
Bagus	1	1	1
Fungsi	1	0	0
Dengan	1	0	1
Baik	1	0	0
Seler	1	0	0
Ramah	1	1	0
Kirim	1	0	0
Cepat	1	0	1
Suai	0	1	0
Harap	0	1	0
Jual	0	1	0
Nya	0	1	0
Juga	0	1	0
Trimakasih	0	1	0
Pelapak	0	1	0
Dan	0	0	1
Respon	0	0	1
Harga	0	0	1
Saing	0	0	1
Yg	0	0	1
Lain	0	0	1

TABEL VI
HASIL PERHITUNGAN TERM-FREQUENCY 2

Kata	Term Frequency (TF)		
	D1	D2	D3
Barang	0,11	0	0,1
Bagus	0,11	0,11	0,1
Fungsi	0,11	0	0
Dengan	0,11	0	0,1
Baik	0,11	0	0
Seler	0,11	0	0
Ramah	0,11	0,11	0
Kirim	0,11	0	0
Cepat	0,11	0	0,1
Suai	0	0,11	0
Harap	0	0,11	0
Jual	0	0,11	0
Nya	0	0,11	0
Juga	0	0,11	0
Trimakasih	0	0,11	0
Pelapak	0	0,11	0
Dan	0	0	0,1
Respon	0	0	0,1
Harga	0	0	0,1

Saing	0	0	0,1
Yg	0	0	0,1
Lain	0	0	0,1

$$tf_{ij} = \frac{f_{ij}}{\sum_{i \in j} f_{ij}} \quad (3)$$

Dengan keterangan bahwa i adalah kata, j adalah dokumen. Frekuensi kata dalam suatu dokumen tersebut dibagi oleh jumlah kata yang ada dokumen tersebut

TABEL VII
HASIL PERHITUNGAN DOCUMENT-FREQUENCY

Kata	Term Frequency (TF)			Document Frequency (DF)
	D1	D2	D3	
Barang	0,11	0	0,1	2
Bagus	0,11	0,11	0,1	3
Fungsi	0,11	0	0	1
Dengan	0,11	0	0,1	2
Baik	0,11	0	0	1
Seler	0,11	0	0	1
Ramah	0,11	0,11	0	2
Kirim	0,11	0	0	1
Cepat	0,11	0	0,1	2
Suai	0	0,11	0	1
Harap	0	0,11	0	1
Jual	0	0,11	0	1
Nya	0	0,11	0	1
Juga	0	0,11	0	1
Trimakasih	0	0,11	0	1
Pelapak	0	0,11	0	1
Dan	0	0	0,1	1
Respon	0	0	0,1	1
Harga	0	0	0,1	1
Saing	0	0	0,1	1
Yg	0	0	0,1	1
Lain	0	0	0,1	1

$$IDF_{ij} = \log \frac{|N_i|}{|DF_{ij}|} \quad (2)$$

Dimana N adalah keseluruhan kata dalam dokumen sedangkan DF_{ij} adalah *document frequency* kata i dalam dokumen besar j.

TABEL VIII
HASIL PERHITUNGAN INVERSE DOCUMENT-FREQUENCY

Kata	TF			DF	IDF
	D1	D2	D3		
Barang	0,11	0	0,1	2	0,176
Bagus	0,11	0,11	0,1	3	0
Fungsi	0,11	0	0	1	0,477
Dengan	0,11	0	0,1	2	0,176
Baik	0,11	0	0	1	0,477
Seler	0,11	0	0	1	0,477
Ramah	0,11	0,11	0	2	0,176
Kirim	0,11	0	0	1	0,477
Cepat	0,11	0	0,1	2	0,176

Suai	0	0,11	0	1	0,477
Harap	0	0,11	0	1	0,477
Jual	0	0,11	0	1	0,477
Nya	0	0,11	0	1	0,477
Juga	0	0,11	0	1	0,477
Trimakasih	0	0,11	0	1	0,477
Pelapak	0	0,11	0	1	0,477
Dan	0	0	0,1	1	0,477
Respon	0	0	0,1	1	0,477
Harga	0	0	0,1	1	0,477
Saing	0	0	0,1	1	0,477
Yg	0	0	0,1	1	0,477
Lain	0	0	0,1	1	0,477

Setelah nilai *Inverse Document Frequency* (IDF) telah ditemukan, nilai tersebut dikalikan dengan *Term Frequency* (TF) untuk membuat nilai *Term Frequency-Inverse Document Frequency* (TF-IDF). Hasilnya dapat dilihat pada Tabel IX.

TABEL IX
HASIL PERHITUNGAN TERM FREQUENCY-INVERSE DOCUMENT FREQUENCY

Kata	TF-IDF		
	D1	D2	D3
Barang	0,01936	0	0,0176
Bagus	0	0	0
Fungsi	0,05247	0	0
Dengan	0,01936	0	0,0176
Baik	0,05247	0	0
Seler	0,05247	0	0
Ramah	0,01936	0,01936	0
Kirim	0,05247	0	0
Cepat	0,01936	0	0,0176
Suai	0	0,05247	0
Harap	0	0,05247	0
Jual	0	0,05247	0
Nya	0	0,05247	0
Juga	0	0,05247	0
Trimakasih	0	0,05247	0
Pelapak	0	0,05247	0
Dan	0	0	0,0477
Respon	0	0	0,0477
Harga	0	0	0,0477
Saing	0	0	0,0477
Yg	0	0	0,0477
Lain	0	0	0,0477

- 2) Pembagian *Dataset*: melakukan pembagian pasangan data untuk menemukan dimanakah keoptimalan algoritma untuk melakukan klasifikasi dengan menggunakan pasangan data tersebut. Pasangan data yang dimaksudkan adalah 0.05: 0.95, 0.10:0.90, 0.15:0.85, 0.20:0.80, 0.25:0.75, 0.30:0.70, 0.35:0.65, 0.40:0.60, 0.45:0.55, dan 0.50:0.50 dari data yang ada.
- 5) Melakukan *resampling*: *dataset* yang tidak seimbang dapat menyebabkan bias pada algoritma yang dipilih.

Untuk mengatasi hal ini, digunakan teknik *resampling* yaitu *oversampling*. *Oversampling* yang digunakan adalah *Random Oversampling*. Teknik ini mengisi kelas minoritas lainnya secara acak.

- 6) Pelatihan Model SVM tanpa *resampling*: keempat trik kernel SVM digunakan untuk membandingkan performa, yaitu *linear*, *poly*, *rbf*, dan *sigmoid*. Selain *linear*, ketiga kernel lainnya menggunakan trik kernel untuk memetakan vektor dataset ke ruang fitur. Sebelum pelatihan, tidak dilakukan *resampling dataset*.
- 7) Pelatihan Model SVM dengan *resampling*: keempat trik kernel SVM digunakan untuk membandingkan performa, yaitu *linear*, *poly*, *rbf*, dan *sigmoid*. Sebelum pelatihan, dilakukan *resampling dataset*. Selain *linear*, ketiga kernel lainnya menggunakan trik kernel untuk memetakan vektor dataset ke ruang fitur. Sebelum pelatihan, dilakukan *resampling dataset*.

B. Evaluasi Model

- 1) *Support Vector Machine* tanpa *resampling*

TABEL X
SVM LINEAR TANPA RESAMPLING

Uji:Latih	Accuracy	Precision	Recall	F1-Score
95:5	66.30	65.62	61.88	62.72
90:10	64.81	64.20	59.38	60.37
85:15	62.35	60.99	56.96	57.84
80:20	62.22	60.33	56.30	57.00
75:25	62.37	60.60	56.71	57.47
70:30	62.35	60.65	56.88	57.65
65:35	62.54	60.75	56.88	57.65
60:40	62.18	60.24	56.21	57.02
55:45	60.66	58.18	54.91	55.60
50:50	60.30	58.06	54.09	54.82

Tabel X menggambarkan hasil evaluasi SVM kernel linear. Pada baris pertama, 95 data latih digunakan untuk melatih algoritma dan 5 data digunakan untuk menguji performa model. Menurut tabel tersebut, nilai accuracy, precision, recall, dan f1-score konsisten menurun seiring menurunnya data latih yang digunakan untuk melatih algoritma dan meningkatnya data uji yang digunakan untuk mengevaluasi model. Dari tabel ini, model terbaik dicapai oleh baris pertama, model dengan 95 data latih dan 5 data uji dengan nilai accuracy 66,30, nilai precision 65,62, nilai recall 61,88, dan nilai f1-score 62,72.

TABEL XI
SVM POLYNOMIAL TANPA RESAMPLING

Uji:Latih	Accuracy	Precision	Recall	F1-Score
95:5	54.07	59.93	42.64	37.51
90:10	53.33	58.51	41.78	35.02
85:15	51.98	55.98	40.37	33.83
80:20	52.22	55.83	40.16	33.09
75:25	51.04	56.74	39.41	31.67
70:30	50.99	56.63	39.34	31.86
65:35	50.16	55.67	37.98	30.70
60:40	50.88	54.83	38.32	31.04

55:45	49.75	54.55	37.08	30.79
50:50	47.19	55.20	34.54	28.68

Tabel XI menggambarkan hasil evaluasi *SVM* kernel *polynomial*. Menurut tabel tersebut, nilai *accuracy*, *precision*, *recall*, dan *f1-score* konsisten menurun seiring menurunnya data latih yang digunakan untuk melatih algoritma dan meningkatnya data uji yang digunakan untuk mengevaluasi model. Dari tabel ini, model terbaik dicapai oleh baris pertama, model dengan 95 data latih dan 5 data uji dengan nilai *accuracy* 54,07, nilai *precision* 59,9, nilai *recall* 42,64, dan nilai *f-score* 37,51.

TABEL XII
SVM RBF TANPA RESAMPLING

Uji:Latih	Accuracy	Precision	Recall	F1-Score
95:5	65.19	67.25	59.36	59.61
90:10	62.96	63.32	56.02	56.44
85:15	59.63	59.56	52.10	52.37
80:20	61.11	61.40	52.99	53.16
75:25	60.00	61.55	52.16	52.38
70:30	60.19	61.59	52.31	52.42
65:35	60.48	61.42	52.32	52.48
60:40	60.05	60.42	51.38	51.31
55:45	59.22	59.50	50.54	50.65
50:50	58.22	58.44	48.95	48.53

Tabel XII menggambarkan hasil evaluasi *SVM* kernel *RBF*. Menurut tabel tersebut, nilai *accuracy*, *precision*, *recall*, dan *f1-score* konsisten menurun seiring menurunnya data latih yang digunakan untuk melatih algoritma dan meningkatnya data uji yang digunakan untuk mengevaluasi model. Dari tabel ini, model terbaik dicapai oleh baris pertama, model dengan 95 data latih dan 5 data uji dengan nilai *accuracy* 65,19, nilai *precision* 67,25, nilai *recall* 59,36%, dan nilai *f-score* 69,61

TABEL XIII
SVM SIGMOID TANPA RESAMPLING

Uji:Latih	Accuracy	Precision	Recall	F1-Score
95:5	66.30	66.00	62.12	62.91
90:10	64.07	63.20	58.37	59.20
85:15	61.23	59.58	55.42	56.18
80:20	61.67	59.54	55.48	56.11
75:25	61.63	60.08	56.16	56.94
70:30	62.10	60.24	56.50	57.18
65:35	62.28	60.76	56.53	57.34
60:40	61.76	59.75	55.70	56.42
55:45	60.86	58.32	54.90	55.56
50:50	60.41	58.16	53.95	54.62

Tabel XIII menggambarkan hasil evaluasi *SVM* kernel *Sigmoid*. Menurut tabel tersebut, nilai *accuracy*, *precision*, *recall*, dan *f1-score* konsisten menurun seiring menurunnya data latih yang digunakan untuk melatih algoritma dan meningkatnya data uji yang digunakan untuk mengevaluasi model. Dari tabel ini, model terbaik dicapai oleh baris pertama, model dengan 95 data latih dan 5 data uji dengan nilai *accuracy* 66,30, nilai *precision* 66, nilai *recall* 62,12, dan nilai *f-score* 62,91

2) Support Vector Machine dengan resampling

TABEL XIV
SVM LINEAR DENGAN RESAMPLING

Uji:Latih	Accuracy	Precision	Recall	F1-Score
5:95	64.44	62.51	61.97	62.04
10:90	63.33	60.76	60.31	60.39
15:85	61.73	59.16	58.89	58.80
20:80	60.56	57.59	57.44	57.11
25:75	62.96	59.91	59.86	59.73
30:70	62.35	59.35	58.89	58.90
35:65	62.28	58.96	59.16	58.97
40:60	61.76	58.83	58.95	58.66
45:55	60.58	57.42	57.18	57.14
50:50	61.00	57.97	57.49	57.59

Tabel XIV menggambarkan hasil evaluasi *SVM* kernel *linear*. Menurut tabel tersebut, nilai *accuracy*, *precision*, *recall*, dan *f1-score* konsisten menurun seiring menurunnya data latih yang digunakan untuk melatih algoritma dan meningkatnya data uji yang digunakan untuk mengevaluasi model. Dari tabel ini, model terbaik dicapai oleh baris pertama, model dengan 9 data latih dan data uji dengan nilai *accuracy* 64,44 , nilai *precision* 62,51, nilai *recall* 61,97 , dan nilai *f-score* 62,04

TABEL XV
SVM POLYNOMIAL DENGAN RESAMPLING

Uji:Latih	Accuracy	Precision	Recall	F1-Score
5:95	51.85	57.88	42.58	37.78
10:90	52.41	56.10	42.15	38.52
15:85	52.59	54.24	40.77	36.25
20:80	53.52	52.62	40.52	36.76
25:75	54.30	55.70	41.63	38.28
30:70	53.09	54.00	40.50	36.37
35:65	53.92	57.46	40.99	37.26
40:60	53.15	56.16	40.23	36.87
45:55	51.73	54.99	38.76	34.86
50:50	50.93	53.30	37.70	33.66

Tabel XV menggambarkan hasil evaluasi *SVM* kernel *polynomial*. Menurut tabel tersebut, nilai tertinggi *accuracy*, *precision*, *recall*, dan *f1-score* ada pada model yang berbeda-beda. *Accuracy* tertinggi pada 25:75, *precision* pada 95:5, *recall* pada 95:5, dan *f1-score* pada 10:90

TABEL XVI
SVM RBF DENGAN RESAMPLING

Uji:Latih	Accuracy	Precision	Recall	F1-Score
5:95	65.19	64.46	61.16	61.39
10:90	64.07	62.62	59.45	59.66
15:85	62.35	59.94	57.10	57.35
20:80	63.52	60.69	57.66	57.81
25:75	63.11	60.58	57.11	57.74

30:70	62.35	59.97	56.51	56.90
35:65	63.60	60.87	57.71	58.39
40:60	63.29	60.84	57.73	58.48
45:55	63.33	61.23	57.43	58.15
50:50	62.70	61.16	56.55	57.37

Tabel XVI menggambarkan hasil evaluasi *SVM* kernel *rbf*. Menurut tabel tersebut, nilai *accuracy*, *precision*, *recall*, dan *f1-score* sempat menurun dan stabil naik turun dari data latih 85 dan 50. Hasil evaluasi tertinggi dicapai oleh data latih 95 dan data uji 5 dengan nilai *accuracy* 65,19 %, nilai *precision* 64,46, nilai *recall* 61,16, dan nilai *f-score* 61,39

TABEL XVII
SVM SIGMOID DENGAN RESAMPLING

Uji:Latih	Accuracy	Precision	Recall	F1-Score
5:95	61.85	60.10	59.76	59.62
10:90	62.22	59.32	59.31	59.23
15:85	61.60	59.32	59.21	58.99
20:80	59.54	56.96	57.16	56.44
25:75	62.22	59.22	60.06	59.36
30:70	61.79	58.85	58.97	58.71
35:65	61.32	58.43	59.15	58.47
40:60	60.93	58.34	58.86	58.13
45:55	60.62	57.73	58.22	57.70
50:50	60.37	57.42	57.52	57.26

Tabel XVII menggambarkan hasil evaluasi *SVM* kernel *sigmoid*. Menurut tabel tersebut, nilai *accuracy* sempat naik turun di angka 61-59. *Accuracy* tertinggi dicapai oleh model 10:90 dan 25:75. *Precision* stabil menurun dari model dengan data latih tertinggi. Nilai *recall* tertinggi ada pada model 25:65. *F1-score* tidak stabil di semua jenis model dan mencapai nilai tertinggi pada 5:95.

Dari kompilasi hasil evaluasi model-model di atas,Tabel XVIII menunjukkan nilai-nilai terbaik dari masing-masing jenis model yang dilatih dengan data yang tidak di *resample*.

TABEL XVIII
PERBANDINGAN PERFORMA MODEL-MODEL TANPA DATA RESAMPLING

Jenis Algoritma		Accuracy	Precision	Recall	F1-Score	Rasio Data
SVM	Linear	66.30	65.62	61.88	62.72	5:95
	Polynomial	54.07	59.93	42.64	37.51	5:95
	RBF	65.19	67.25	59.36	59.61	5:95
	Sigmoid	66.30	66.00	62.12	62.91	5:95

Melihat Tabel XVIII, model yang melakukan performa klasifikasi terbaik secara keseluruhan dalam kasus ini adalah *SVM sigmoid* dengan nilai *accuracy*, *precision*, *recall*, dan *f1-score*, 66.30, 66.00, 62.12, dan 62.91 pada rasio data latih dan uji 95:5.

Nilai terdekat dengan model *SVM sigmoid* adalah model *SVM linear*, dengan nilai *accuracy* yang sama, selisih nilai *precision* 0,39 dengan *SVM sigmoid* bernilai 66 dan nilai *SVM linear* 65,62, selisih nilai *recall* 0,24 dengan *SVM sigmoid* bernilai 62,12 dan *SVM linear* bernilai 61,99, serta selisih nilai *f1-score* 0,2 dengan *SVM sigmoid* bernilai 62,19 dan *SVM linear* 62,72.

TABEL XIX
PERBANDINGAN PERFORMA MODEL-MODEL DENGAN DATA RESAMPLING

Jenis Algoritma		Accuracy	Precision	Recall	F1-Score	Rasio Data
SVM	Linear	64.44	62.51	61.97	62.04	5:95
	Polynomial	54.30	55.70	41.63	38.28	25:75
	RBF	65.19	64.46	61.16	61.39	5:95
	Sigmoid	62.22	59.22	60.06	59.36	25:75

Tabel XIX membandingkan nilai hasil evaluasi klasifikasi model dilatih oleh data yang telah di-*resample*. Dari keseluruhan model, tidak ada model yang memiliki nilai evaluasi tertinggi secara keseluruhan. *SVM rbf* memiliki nilai tertinggi dalam *accuracy* dan *precision*, yaitu 65,19 dan 64,46, namun *SVM linear* memiliki nilai tertinggi dalam *recall* dan *f1-score*, yaitu 61,97 dan 62,04. Kedua model ini memiliki nilai yang saling menyaingi. Terdapat selisih *accuracy* 0,75, dengan *SVM rbf* bernilai 65,19 dan *SVM linear* bernilai 64,44, selisih *precision* 1,95 dengan *SVM rbf* bernilai 64,46 dan *SVM linear* bernilai 62,51, selisih *recall* 0,81 dengan *SVM linear* bernilai 62,04 dan *SVM rbf* bernilai 61,16, serta selisih *f1-score* 0,65 dengan *SVM linear* bernilai 62,04 dan *SVM rbf* bernilai 61,39.

Setelah menampilkan nilai hasil evaluasi tiap algoritma dan jenisnya secara terpisah, yaitu baik dengan pelatihan model yang menggunakan *dataset* tanpa *resampling* atau pelatihan model menggunakan *dataset* dengan *resampling*, Tabel XXVI menampilkan kompilasi nilai hasil evaluasi model dari seluruh jenis algoritma dan jenis-jenisnya untuk melakukan klasifikasi dalam kedua skenario.

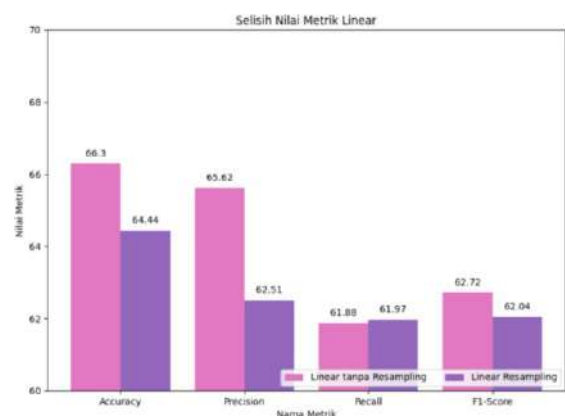
TABEL XX
PERBANDINGAN PERFORMA KESELURUHAN JENIS MODEL

Jenis Algoritma		Accuracy	Precision	Recall	F1-Score	Rasio Data
SVM	Linear	66,30	65.62	61.88	62.72	5:95
	Linear Resampling	64.44	62.51	61.97	62.04	5:95

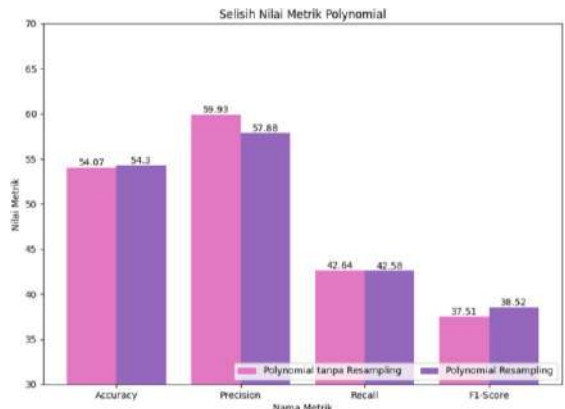
	Polynomi al	54.07	59.93	42.64	37.51	5:95
	Polynomi al Resamplin g	54.30	57.88	42.58	38.52	-
	RBF	65.19	67.25	59.36	59.61	5:95
	RBF Resamplin g	65.19	64.46	61.16	61.39	5:95
	Sigmoid	66.30	66.00	62.12	62.91	5:95
	Sigmoid Resamplin g	62.22	60.10	60.06	59.62	-

Menurut Tabel XX, dari keseluruhan model yang ada, nilai tertinggi untuk *accuracy* adalah SVM kernel *Sigmoid* dan *Linear*, dengan nilai 66,3. Nilai terbesar *precision* adalah oleh SVM kernel *RBF*, dengan nilai 67,25. Nilai terbesar *recall* adalah oleh SVM kernel *Sigmoi* dengan nilai 62,12. Dan nilai terbesar untuk *F1-Score* adalah SVM kernel *Sigmoid*, dengan nilai 62,1. SVM memiliki performa baik dalam penelitian ini, kecuali untuk model SVM *Polynomial*. Model SVM *Polynomial* baik dengan data *resampling*, memiliki nilai yang menduduki posisi terendah ketiga. SVM *Polynomial* tanpa data *resampling* menduduki posisi terendah ketiga dalam *accuracy* dan *f1-score*, sedangkan model SVM *Polynomial* dengan data *resampling* menduduki posisi terendah ketiga pada *precision* dan *recall*.

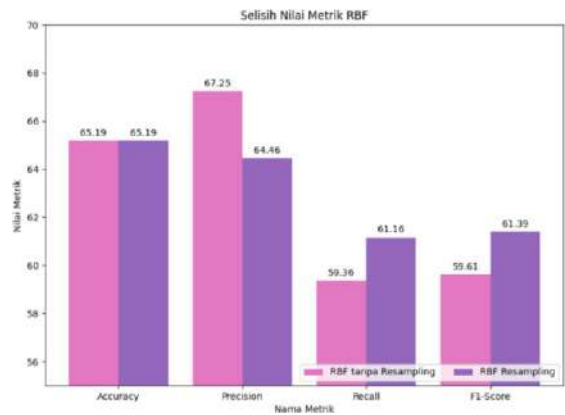
Menurut Tabel XX, melakukan data *resample* tidak memiliki efek yang signifikan pada hampir semua model. Pada model-model SVM *Linear*, *accuracy*-nya menurun 1,86, *precision* turun 3,11, nilai *recall* naik 0,09, dan nilai *f1-score* turun 0,74. Pada model SVM *Polynomial*, nilai *accuracy* menurun 0,23, nilai *precision* menurun 2,05, nilai *recall* menurun 0,06, dan nilai *f1-score* naik 0,74. Pada model SVM *RBF*, nilai *accuracy* tetap, *precision* menurun 2,79, nilai *recall* naik 1,8, dan nilai *f1-score* naik. Terakhir, pada SVM *Sigmoid*, nilai *accuracy* menurun 4,08, nilai *precision* naik 0,1, nilai *recall* menurun 2,06, dan nilai *f1-score* menurun 2,57.



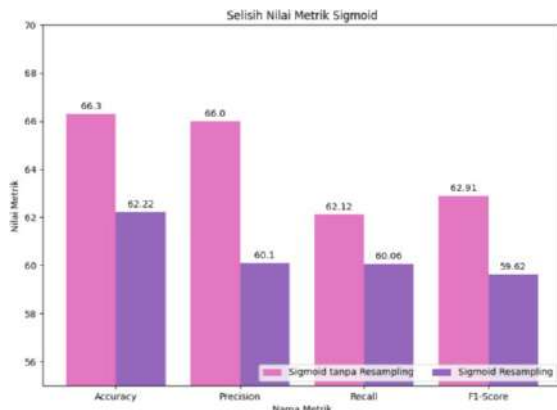
Gbr. 8 Perbandingan Nilai Evaluasi Model SVM *Linear*



Gbr. 9 Perbandingan Nilai Evaluasi Model SVM *Polynomial*



Gbr. 10 Perbandingan Nilai Evaluasi Model SVM *RBF*



Gbr. 11 Perbandingan Nilai Evaluasi Model SVM Sigmoid

IV. KESIMPULAN

Model klasifikasi emosi berbasis teks dibangun melalui berbagai macam metode. Langkah pertama adalah untuk melakukan normalisasi dan pembersihan teks, yaitu penghapusan tanda baca, simbol, angka, dan mengecilkan huruf. Selanjutnya dilakukan tokenisasi untuk memecah kalimat menjadi kata. Terakhir adalah melakukan lematisasi yaitu pengembalian kata menjadi bentuk bakunya. Langkah selanjutnya adalah untuk mengekstrak fitur dari kata-kata tersebut dengan mengubah kata-kata menjadi vektor di ruang fitur menggunakan TF-IDF. Lalu, dataset dibagi menjadi sepuluh pasang data latih dan uji untuk mengevaluasi performa algoritma.

Sepuluh pasang data ini dibagi kembali dengan melakukan resampling di salah satu set. Dengan ini, 2 jenis set data dan 20 hasil klasifikasi dari satu jenis algoritma. Setelah pelatihan, kedua algoritma dievaluasi menggunakan precision_recall_fscore_support dan accuracy_score.

Secara keseluruhan, model SVM sigmoid mencapai nilai terbesar di keempat nilai evaluasi model dengan accuracy 66,3%, precision 66%, recall 62,12%, dan f1-score 62,92%.

REFERENSI

- [1] M. Pradana, "KLASIFIKASI JENIS-JENIS BISNIS E-COMMERCE DI INDONESIA," *Neo-bis*, vol. 9, no. 2, 2015.
- [2] M. S. Ullal, C. Spulbar, I. T. Hawaldar, V. Popescu, and R. Birau, "The Impact of Online Reviews on E-commerce Sales in India: a Case Study," *Economic Research-Ekonomska Istrazivanja*, vol. 34, no. 1, pp. 2408–2422, 2021, doi: 10.1080/1331677X.2020.1865179.
- [3] C. Changchit and T. Klaus, "Determinants and Impact of Online Reviews on Product Satisfaction," *Journal of Internet Commerce*, vol. 19, no. 1, pp. 82–102, Jan. 2020, doi: 10.1080/15332861.2019.1672135.
- [4] R. Mustajab, "Pengguna E-Commerce RI Diproyeksi Capai 196,47 Juta pada 2023 - DataIndonesia," *DataIndonesia.id*.
- [5] G. Schuh *et al.*, "Data mining definitions and applications for the management of production complexity," in *52nd CIRP Conference on Manufacturing Systems*, Elsevier B.V., 2019, pp. 874–879. doi: 10.1016/j.procir.2019.03.217.

- [6] C. E. Joergensen Munthe, N. Astuti Hasibuan, and H. Hutabarat, "Penerapan Algoritma Text Mining Dan TF-RF Dalam Menentukan Promo Produk Pada Marketplace," *RESOLUSI: Rekayasa Teknik Informatika dan Informasi*, vol. 2, no. 3, pp. 110–115, 2022, [Online]. Available: <https://djournals.com/resolusi>
- [7] L. J. Zheng, J. Mountstephens, and J. Teo, "Four-class emotion classification in virtual reality using pupillometry," *J Big Data*, vol. 7, no. 1, Dec. 2020, doi: 10.1186/s40537-020-00322-9.
- [8] J. Hartmann, J. Huppertz, C. Schamp, and M. Heitmann, "Comparing automated text classification methods," *International Journal of Research in Marketing*, vol. 36, no. 1, pp. 20–38, Mar. 2019, doi: 10.1016/j.ijresmar.2018.09.009.
- [9] E. Leopold, "Text Categorization with Support Vector Machines. How to Represent Texts in Input Space?," 2002.
- [10] O. Arifin and T. B. Sasongko, "ANALISA PERBANDINGAN TINGKAT PERFORMANSI METODE SUPPORT VECTOR MACHINE DAN NAÏVE BAYES CLASSIFIER UNTUK KLASIFIKASI JALUR MINAT SMA," in *Seminar Nasional Teknologi Informasi dan Multimedia 2018*, 2018, pp. 67–72.
- [11] N. P. Damayanti, D. E. Prameswari, W. Puspita, and P. S. Sundari, "Classification of Hate Comments on Twitter Using a Combination of Logistic Regression and Support Vector Machine Algorithm," *Journal of Information System Exploration and Research*, vol. 2, no. 1, Jan. 2024, doi: 10.52465/joiser.v2i1.229.
- [12] A. Wibowo Haryanto and E. Kholid Mawardi, "Influence of Word Normalization and Chi-squared Feature Selection on Support Vector Machine (SVM) Text Classification."
- [13] R. Sutoyo, S. Achmad, A. Chowanda, E. Widhi Andangsari, and S. M. Isa, "PRDECT-ID Indonesian Product Reviews Dataset for Emotions Classification Tasks," *Data Brief*, vol. 44, 2022, doi: 10.17632/574v66hf2v.1.