

Model Klasifikasi Emosi Berbasis Teks dengan Algoritma *Decision Tree* dan *Support Vector Machine*

Habib Aulia Raihan^{1*}, Herman Yuliansyah², Murinto³

¹ Program Studi Magister Informatika, Fakultas Teknologi Industri

^{2,3} Program Studi Informatika, Fakultas Teknologi Industri

Universitas Ahmad Dahlan Yogyakarta

*Email: habibaularaihan@gmail.com

Abstrak

Komunikasi berbasis teks telah menjadi salah satu sarana interaksi dalam berbagai sektor. Penelitian sebelumnya menerapkan algoritma supervised learning untuk klasifikasi emosi pada teks. Penelitian tentang klasifikasi emosi menggunakan dataset yang berbeda-beda. Namun, memunculkan potensi overfitting pada model klasifikasi emosi berbasis teks. Sehingga diperlukan penerapan cross-validation dan optimasi hyperparameter untuk memastikan kemampuan generalisasi model. Tujuan dari penelitian ini untuk membandingkan kinerja dua algoritma supervised learning yaitu *Decision Tree* (DT) dan *Support Vector Machine* (SVM) untuk klasifikasi emosi pada dataset teks berbahasa Inggris yang berisi 16.000 data teks berlabel (anger, fear, joy, love, sadness, surprise) yang bersumber dari kaggle. Dataset tersebut diproses melalui tahap cleaning, tokenizing, stopword removal, dan lemmatization, kemudian feature extraction menggunakan TF-IDF. Kedua algoritma dievaluasi menggunakan KFold dan Stratified KFold kemudian dilakukan klasifikasi untuk mengetahui hasil metrik accuracy, precision, recall, dan f1-score. Hasil evaluasi klasifikasi menunjukkan DT dengan hyperparameter memperoleh akurasi rata-rata 88% dan SVM dengan hyperparameter model memperoleh akurasi rata-rata 89%. Sedangkan hasil cross-validation menggunakan Stratified KFold menunjukkan selisih akurasi sebesar 0,02% untuk DT dan 0,15% untuk SVM. Oleh karena itu dapat disimpulkan bahwa Stratified KFold lebih baik daripada KFold untuk dataset yang tidak seimbang. Kemudian algoritma SVM dengan hyperparameter lebih unggul dibandingkan dengan DT dengan hyperparameter.

Kata kunci: *Decision Tree, Support Vector Machine, Klasifikasi Emosi, TF-IDF*

Abstract

Text-based communication has become a key means of interaction across various sectors. Previous studies have applied supervised learning algorithms to emotion classification in text. These studies used different datasets, but this diversity also introduced a risk of overfitting in text-based emotion classification models. Consequently, the use of cross-validation and hyperparameter optimization is required to ensure the model's generalization ability. The aim of this research is to compare the performance of two supervised learning algorithms—*Decision Tree* (DT) and *Support Vector Machine* (SVM)—for emotion classification on an English-language text dataset of 16,000 labeled entries (anger, fear, joy, love, sadness, surprise) sourced from Kaggle. The dataset undergoes cleaning, tokenization, stopword removal, and lemmatization, after which features are extracted using TF-IDF. Both algorithms are evaluated with K-Fold and Stratified K-Fold cross-validation, then used to compute metrics of accuracy, precision, recall, and F1-score. Classification results show that the hyperparameter-tuned DT achieved an average accuracy of 88%, while the hyperparameter-tuned SVM achieved 89%. Meanwhile, Stratified K-Fold cross-validation yielded an accuracy variance of just 0.02% for DT and 0.15% for SVM. Therefore, it can be concluded that Stratified K-Fold performs better than standard K-Fold on imbalanced datasets, and that hyperparameter-tuned SVM outperforms hyperparameter-tuned DT.

Keywords: *Decision Tree, Emotion Classification, Support Vector Machine, TF-IDF*

PENDAHULUAN

Komunikasi berbasis teks telah menjadi salah satu sarana interaksi dalam berbagai sektor, termasuk layanan pelanggan (Elinda et al., 2024; Wulan & Basri, 2024), platform hiburan (Syafia et al., 2023), dan *e-commerce* (Cahyaningtyas et al., 2021). Dalam interaksi ini, klasifikasi emosi pada teks menjadi aspek

penting untuk dipahami secara lebih mendalam guna memberikan pengalaman interaksi yang lebih personal (Arifian et al., 2024). Klasifikasi emosi berbasis teks masih perlu peningkatan, sehingga membuka ruang untuk eksplorasi penelitian lebih lanjut. Penelitian ini dilakukan agar *Natural Language Processing* (NLP) dengan algoritma *traditional machine learning*

bisa terus dieksplorasi (Chowdhary, 2020). Perkembangan teknologi *Artificial Intelligence* (AI) memberikan peluang besar untuk mengatasi kendala tersebut (Bijaksana Putra Negara et al., 2021; Patel et al., 2022; Thomas et al., 2021; Yuliansyah et al., 2023). Permasalahan yang biasa ditemui adalah algoritma *Decision Tree* (DT) memiliki efektivitas yang rendah dalam tugas pemrosesan bahasa alami dibandingkan algoritma *traditional machine learning* yang lain seperti *Support Vector Machine*, *Naive Bayes*, dan *Random Forest* (Depari et al., 2022; Rokhman et al., 2021). DT memodelkan hubungan kompleks antar variabel dengan cara yang sederhana dan mudah dipahami. Algoritma DT membentuk model seperti pohon keputusan dengan *node* (Maruf et al., 2022; Rahayu et al., 2023). *Support Vector Machine* (SVM) memiliki keunggulan dalam menangani data berdimensi tinggi dan bersifat tidak linier, yang sering ditemukan dalam pemrosesan bahasa alami. Penelitian ini bertujuan untuk melakukan komparasi antara algoritma DT dan SVM dalam konteks klasifikasi emosi pada data teks. Penelitian serupa telah dilakukan oleh Pramesti & Pratiwi (2023), yang mengkomparasikan algoritma DT dan SVM pada dataset yang berbeda dan dengan fokus analisis yang terbatas tanpa melakukan optimasi *tuning hyperparameter* maupun teknik evaluasi *cross-validation*. Berbeda dengan penelitian tersebut, studi ini mengintegrasikan analisis fitur teks yang rinci, optimasi *hyperparameter*, serta evaluasi model dengan *Stratified K-Fold cross-validation* pada dataset besar dan tidak seimbang, sehingga diharapkan mampu menghasilkan model klasifikasi emosi yang lebih akurat dan generalisasi yang lebih kuat. Pendekatan penelitian ini mencakup ekstraksi fitur dari data teks menggunakan teknik pemrosesan bahasa alami yang kemudian diklasifikasikan menggunakan kedua algoritma tersebut (Ashraf et al., 2022; Azam et al., 2021; Susandri et al., 2023). Komparasi DT dengan SVM akan memberikan gambaran yang jelas mengapa DT bisa memiliki efektivitas yang rendah dalam pemrosesan NLP. Penelitian ini menawarkan kebaruan dalam mengintegrasikan algoritma DT dan SVM dengan *dataset* yang berbeda serta analisis fitur teks dan *tuning hyperparameter* untuk klasifikasi emosi. Hasil penelitian ini diharapkan tidak hanya memperluas pemahaman tentang pemilihan algoritma terbaik dalam klasifikasi emosi tetapi

juga memberikan kontribusi terhadap pemrosesan bahasa alami menggunakan *traditional machine learning*..

TINJAUAN PUSTAKA

2.1 Klasifikasi Emosi

Emosi merujuk pada suatu perasaan dan pikiran yang khas, suatu keadaan biologis dan psikologis dan serangkaian kecenderungan untuk bertindak (Machová et al., 2023). Klasifikasi emosi adalah proses mengenali serta memahami emosi manusia yang diekspresikan dalam berbagai bentuk, terutama melalui bahasa alami. Klasifikasi emosi mencakup ekstraksi kondisi emosional seperti *joy*, *sad*, *anger*, *fear*, *surprise* dan *love* dari teks atau ujaran manusia (Azam et al., 2021; Fernandes et al., 2024; Machová et al., 2023; Nandwani & Verma, 2021; Rohman et al., 2019). Penelitian ini penting karena emosi memiliki dampak signifikan terhadap perilaku manusia, pengambilan keputusan, interaksi sosial, serta produktivitas individu. Penggunaan bahasa alami sebagai sumber data utama dalam klasifikasi emosi menjadi tantangan tersendiri, karena bahasa merupakan fenomena yang dinamis dan sangat dipengaruhi oleh konteks, budaya, serta karakteristik personal masing-masing individu (Okta et al., 2024). Chowanda et al., (2021) menyebutkan bahwa gaya linguistik manusia sangat beragam, yang menyebabkan sulitnya menemukan representasi emosional yang universal. Lebih jauh lagi, representasi emosional sering kali bersifat implisit, samar, atau ambigu, sehingga mempersulit proses untuk menyimpulkan keadaan emosional seseorang secara konkret dan akurat.

2.2 TF-IDF

TF-IDF (*Term Frequency-Inverse Document Frequency*) merupakan salah satu metode statistik penting dalam bidang *text mining* dan *information retrieval* yang digunakan untuk mengukur signifikansi atau tingkat kepentingan suatu kata terhadap dokumen tertentu dalam suatu kumpulan dokumen (*corpus*) (Ab. Nasir et al., 2020; Acheampong et al., 2020; Chowanda et al., 2021; Kusal et al., 2022; Liu et al., 2023; Machová et al., 2023; Nandwani & Verma, 2021). Konsep dasar metode TF-IDF diperkenalkan untuk membantu mengatasi kendala dalam representasi teks yang besar, di mana pentingnya sebuah kata tidak dapat

ditentukan hanya dari seberapa sering kata tersebut muncul dalam dokumen, melainkan juga dari seberapa jarang atau seringnya kata tersebut ditemukan di dalam kumpulan dokumen secara keseluruhan (Ab. Nasir et al., 2020; Agus Setiawan & Yuliansyah, 2024).

2.3 Decision Tree

Decision Tree (DT) atau pohon keputusan adalah salah satu metode klasifikasi dalam *supervised learning* yang banyak digunakan dalam berbagai bidang seperti *data mining*, pembelajaran mesin, dan analisis prediktif. Menurut Al Maruf et al., (2022), DT merupakan metode klasifikasi yang merepresentasikan proses pengambilan keputusan dalam bentuk struktur pohon. Struktur ini terdiri atas sejumlah *node* yang saling terhubung, membentuk pohon dengan satu *node* akar (*root node*) yang kemudian bercabang hingga menghasilkan *node* daun (*leaf node*) sebagai hasil akhir klasifikasi (Cahyaningtyas et al., 2021; Sinaga & Agustian, 2022; Sondakh et al., 2023).

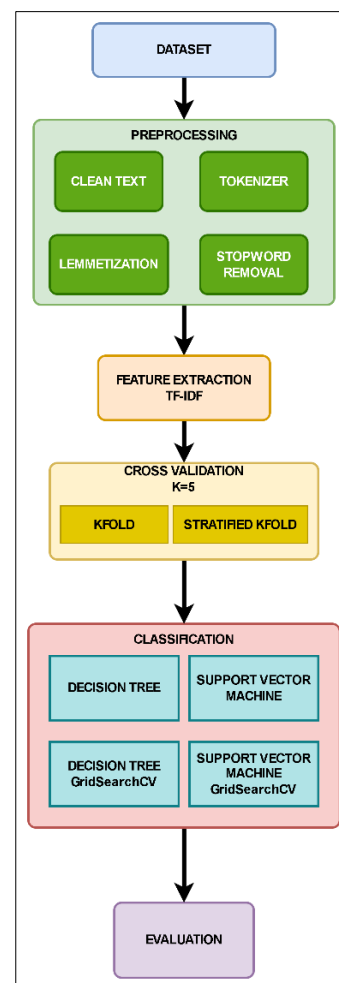
2.4 Support Vector Machine

Support Vector Machine (SVM) adalah algoritma *machine learning* yang sering digunakan dalam klasifikasi teks, termasuk untuk klasifikasi emosi (Rokhman et al., 2021). Prinsip dasar SVM adalah menemukan *hyperplane* terbaik yang memisahkan kelas-kelas data dengan margin maksimum, sehingga meningkatkan generalisasi model terhadap data yang belum pernah dilihat sebelumnya (Gunawan et al., 2023). SVM memiliki keunggulan dalam menangani data berdimensi tinggi, yang sering ditemukan dalam tugas klasifikasi teks dan emosi, melalui pemanfaatan *kernel* untuk mengatasi *non-linearity* dalam data. Beberapa *kernel* yang umum digunakan dalam SVM adalah *linear*, *polynomial*, *Radial Basis Function* (RBF), dan *sigmoid*, yang masing-masing memiliki karakteristik yang berbeda dalam mengolah data (Sontayasara et al., 2021).

METODE PENELITIAN

Metode penelitian terdiri dari enam tahapan yang ditunjukkan pada Gambar 1. Pertama, pengumpulan dataset berupa teks berbahasa Inggris berlabel enam kategori emosi (*anger*, *fear*, *joy*, *love*, *sadness*, *surprise*) dari *Kaggle*. Kedua, *preprocessing* mencakup empat sub-tahap: *clean text* untuk menghapus karakter

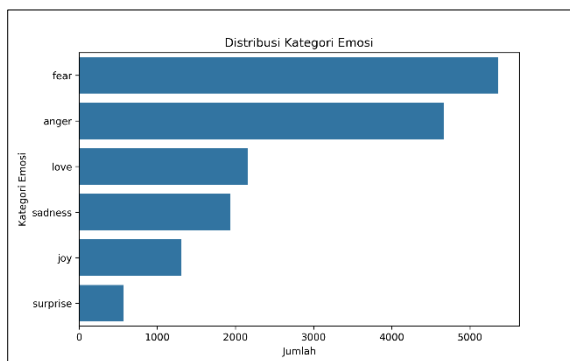
husus dan tanda baca; *tokenizing* memecah kalimat menjadi *token*, *stopword removal* menyingkirkan kata umum tanpa makna emosional, dan *lemmatization* mengembalikan *token* ke bentuk dasar. Ketiga, *feature extraction* menggunakan TF-IDF mengonversi teks menjadi vektor numerik dengan bobot kata berdasarkan frekuensi. Keempat, penerapan *K-Fold* dan *Stratified K-Fold cross-validation* ($K=5$) menjaga distribusi kelas dan mencegah *overfitting* dengan memvalidasi model pada setiap lipatan. Kelima, klasifikasi menguji DT dan SVM dalam varian tanpa penyetelan dan dengan optimasi *hyperparameter* melalui *GridSearchCV*. Keenam, evaluasi akhir mengukur kinerja model menggunakan metrik *accuracy*, *precision*, *recall*, *F1-score*, serta menilai variansi akurasi antarlipatan pada *Stratified K-Fold* untuk memastikan kemampuan generalisasi model secara menyeluruh dan *robust*.



Gambar 1. Metode penelitian

3.1 Pengumpulan Dataset

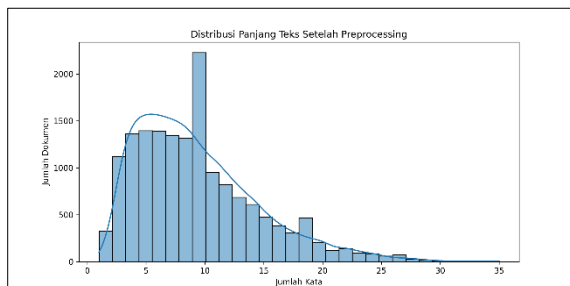
Dataset yang digunakan dalam penelitian ini bersumber dari platform *Kaggle* (<https://www.kaggle.com/datasets/parulpandey/emotion-dataset>), dan terdiri atas 16.000 data teks yang masing-masing telah diberi label berdasarkan jenis emosi yang terkandung di dalamnya. Terdapat enam kategori emosi dalam *dataset* ini, yaitu *anger*, *fear*, *joy*, *love*, *sadness*, dan *surprise*. Namun, seperti yang terlihat pada Gambar 2, jumlah data pada masing-masing kategori tidak merata. Emosi *fear* dan *anger* mendominasi jumlah data, sementara emosi seperti *surprise* memiliki representasi yang jauh lebih sedikit.



Gambar 2. Distribusi label emosi

3.2 Preprocessing

Tahap ini terdiri dari 4 langkah yaitu *clean text* untuk menghilangkan angka dan tanda baca, *tokenizer* untuk memecah teks menjadi kata agar memudahkan analisis statistik per kata, *stopword removal* untuk menghapus kata umum yang tidak informatif, dan *lemmetization* untuk mengubah kata menjadi bentuk kata dasar.



Gambar 3. Panjang teks setelah *preprocessing*

Distribusi panjang teks seperti yang ditunjukkan pada Gambar 3 setelah melalui tahap *preprocessing*. Sumbu horizontal menunjukkan jumlah kata dalam setiap teks, yang

mengindikasikan panjang teks tersebut. Sumbu vertikal menggambarkan jumlah dokumen atau teks yang memiliki panjang tertentu. Sebagai contoh, dari Gambar 3 dapat diamati bahwa terdapat sekitar 2000 dokumen yang panjang teksnya mencapai sekitar 10 kata, menunjukkan bahwa kebanyakan dokumen dalam *dataset* cenderung singkat. Selain itu, grafik juga dilengkapi dengan garis kurva biru yang menggambarkan distribusi densitas atau kepadatan data secara umum. Kurva ini memberikan informasi tambahan bahwa distribusi panjang teks dalam *dataset* cenderung miring ke kanan (*right-skewed distribution*). Artinya, sebagian besar dokumen memiliki panjang teks yang relatif pendek, sementara jumlah dokumen dengan teks yang lebih panjang jauh lebih sedikit.

3.3 Feature Extraction

Pada tahap ini dilakukan ekstraksi fitur menggunakan TF-IDF. TF-IDF dirumuskan sebagai berikut :

Term Frequency (TF):

$$TF(t, d) = \frac{ft, d}{\sum_d ft, d} \quad (1)$$

Keterangan:

ft, d adalah jumlah kemunculan kata t dalam dokumen d .

$\sum_d ft, d$ adalah total jumlah kata dalam dokumen d .

Inverse Document Frequency (IDF):

$$IDF(t) = \log\left(\frac{N}{N_t}\right) \quad (2)$$

Keterangan:

N adalah jumlah total dokumen dalam corpus.

N_t adalah jumlah dokumen di mana kata t muncul.

Setelah nilai TF dan IDF dihitung, bobot TF-IDF dari suatu kata dihitung sebagai hasil perkalian kedua nilai tersebut, dinyatakan dalam rumus berikut:

$$TF - IDF(t, d) = TF(t, d) \times IDF(t) \quad (3)$$

3.4 Cross Validation

Penelitian ini menerapkan *Cross Validation* dengan *KFold* dan *Stratified KFold* untuk masing – masing algoritma dengan $k = 5$. *K-Fold* membagi *dataset* menjadi lima bagian

yang seimbang pada setiap iterasi dengan satu bagian dijadikan data *test* dan empat bagian lainnya sebagai data *train*. Selanjutnya, *Stratified K-Fold* digunakan untuk mengatasi potensi ketidakseimbangan kelas. Teknik ini memastikan proporsi masing-masing label emosi di setiap *fold* menyerupai distribusi aslinya, sehingga setiap kelas terwakili pada data latih dan data uji.

Tabel 1. *Accuracy* setiap *fold*

Fold-i	DT		SVM	
	K-5-Fold	Stratified-K-5-Fold	K-5-Fold	Stratified-K-5-Fold
Fold-1	86.25%	86.80%	84.10%	84.18%
Fold-2	86.76%	87.34%	83.52%	84.14%
Fold-3	86.52%	85.78%	83.36%	83.32%
Fold-4	85.70%	86.56%	82.89%	82.89%
Fold-5	87.11%	86.76%	83.12%	83.24%
Mean	86.47%	86.65%	83.40%	83.55%

Tabel 1 menunjukkan *accuracy* setiap *fold*, dan setiap *fold* menghasilkan skor metrik *accuracy*, *precision*, *recall*, dan *F1-score* yang kemudian dihitung rata-rata untuk menilai stabilitas model seperti yang ditunjukkan pada Tabel 2.

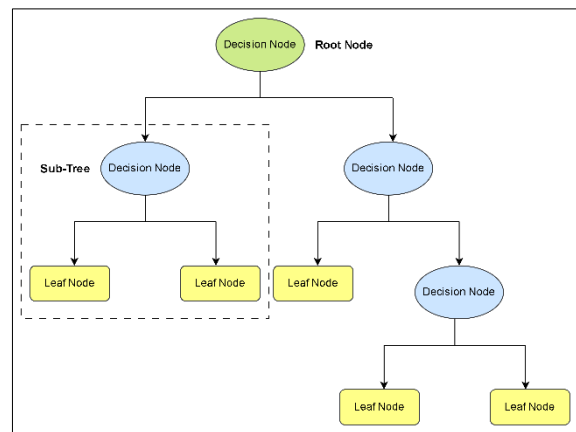
Tabel 2. *Metrics cross validation*

Metrics	DT		SVM	
	K-5-Fold	Stratified-K-5-Fold	K-5-Fold	Stratified-K-5-Fold
Accuracy	0,8647	0,8665	0,834	0,8355
Precision	0,8679	0,8677	0,8484	0,8498
Recall	0,8669	0,8667	0,834	0,8355
f1	0,867	0,8669	0,8227	0,8244

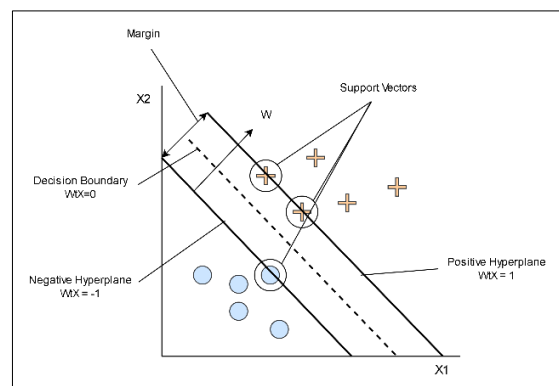
3.5 Klasifikasi Emosi

Penelitian ini menerapkan dua algoritma klasifikasi, yaitu DT dan SVM. DT membuat

setiap keputusan diambil dengan memodelkan hubungan antara berbagai variabel. DT dimulai dari pemilihan fitur paling informatif dalam penelitian ini diwakili oleh vektor TF-IDF sebagai kriteria pemisahan awal pada simpul akar. Setiap dokumen teks kemudian “ditelusuri” melalui pohon dengan membandingkan nilai fitur terhadap ambang batas yang ditentukan. Gambar 4 menunjukkan ilustrasi proses kerja algoritma DT.

Gambar 4. *Decision Tree*

Sementara itu, SVM digunakan untuk mencari *hyperplane* terbaik dengan memaksimalkan jarak antar kelas. Setiap dokumen yang telah direpresentasikan sebagai vektor TF-IDF digunakan untuk mencari *hyperplane* optimal suatu batas pemisah yang memaksimalkan margin. Gambar 5 menunjukkan ilustrasi proses kerja algoritma SVM.

Gambar 5. *Support Vector Machine*

3.6 Evaluasi

Evaluasi dilakukan untuk mengukur performa model dalam mengklasifikasikan emosi berdasarkan teks. Proses evaluasi ini untuk memberikan *metrics* evaluasi seperti *accuracy*, *precision*, *recall*, dan *F1-Score*.

Accuracy :

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (4)$$

Precision:

$$Precision = \frac{TP}{TP+FP} \quad (5)$$

Recall :

$$Recall = \frac{TP}{TP+FN} \quad (6)$$

F1:

$$F1 = 2 \times \frac{Recall \times Precision}{Recall + Precision} \quad (7)$$

Proses evaluasi dilakukan dengan 2 tahap, tahap pertama tanpa *hyperparameter* dan tahap kedua dengan *hyperparameter* menggunakan *GridSearchCV*. *GridSearchCV* adalah metode pemilihan kombinasi model dan *hyperparameter* dengan cara menguji coba satu persatu kombinasi dan melakukan validasi untuk setiap kombinasi. Tujuannya adalah menentukan kombinasi yang menghasilkan performa model terbaik yang dapat dipilih untuk dijadikan model untuk prediksi.

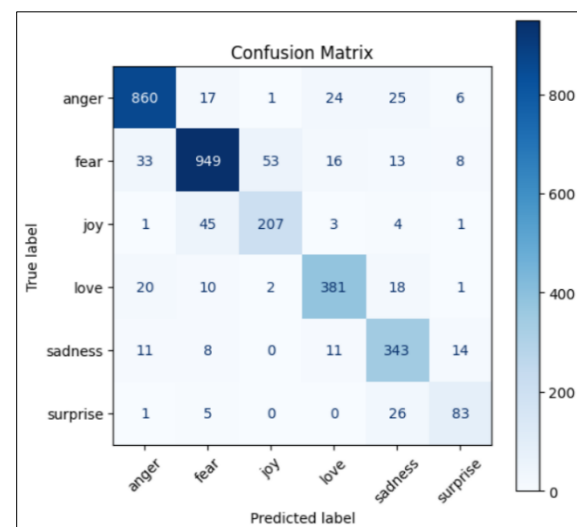
HASIL DAN PEMBAHASAN

GridSearchCV digunakan untuk menghasilkan kombinasi parameter terbaik. Parameter yang dihasilkan *GridSearchCV* untuk TF-IDF Vectorizer adalah *max_features*=10000 dan *ngram_range*=(1,2). Parameter yang digunakan untuk algoritma DT yaitu *max_depth*=None, *min_sample_split*=10, *criterion*= 'gini', dan *random_split*=42. Parameter yang digunakan untuk SVM yaitu C = 1, *gamma*= 'scale', *kernel*= 'linear' dan *random_state*=42. Cross validation untuk data train yang digunakan adalah *StratifiedKFold* dengan *n_splits*=5.

Tabel 3. *Metrics* evaluasi

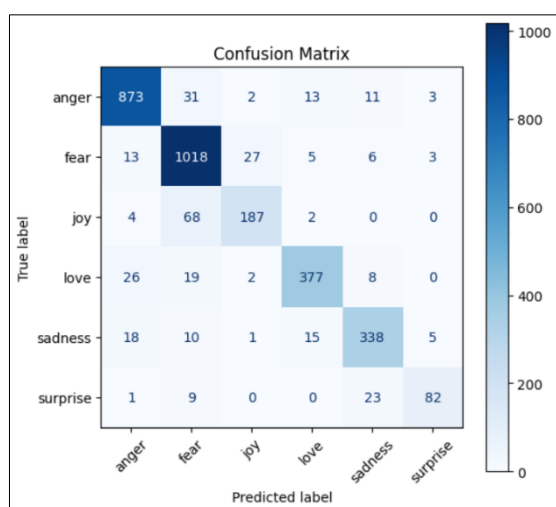
Metric s	DT		SVM	
	Baseli ne model	Hyperpara meter Tuning	Baseli ne model	Hyperpara meter Tuning
Accur acy	0,8797	0,8822	0,8703	0,8984
Precisi on	0,839	0,8406	0,896	0,8902
Recall	0,8499	0,8484	0,77	0,8435
f1	0,8441	0,844	0,8149	0,8641

Metric evaluasi setiap algoritma ditunjukkan pada Tabel 3. Algoritma DT menghasilkan rata – rata akurasi sebesar 87,97% tanpa *hyperparameter* dan akurasi 88,22% dengan *hyperparameter*. *Confusion Matrix* algoritma DT yang menggunakan *hyperparameter* ditunjukkan pada Gambar 6.



Gambar 6. *Confusion Matrix* DT dengan *hyperparameter*

Sedangkan algoritma SVM menghasilkan rata – rata akurasi sebesar 87,03% tanpa *hyperparameter* dan 89,84% dengan *hyperparameter*. *Confusion Matrix* algoritma SVM yang menggunakan *hyperparameter* ditunjukkan pada Gambar 7.



Gambar 7. Confusion Matrix SVM dengan hyperparameter

SIMPULAN

Penelitian ini berhasil membandingkan dua algoritma *machine learning*, DT dan SVM untuk klasifikasi emosi pada teks. Dengan memanfaatkan TF-IDF untuk ekstraksi fitur dan *stratified 5-fold cross-validation* untuk evaluasi, SVM dengan *hyperparameter* menunjukkan performa yang lebih baik dibandingkan DT dengan *hyperparameter*. Keunggulan ini diperoleh baik pada data latih maupun pada data uji.

Penelitian selanjutnya bisa mengeksplorasi teknik *traditional machine learning* untuk klasifikasi emosi pada teks *live chat*.

DAFTAR PUSTAKA

- Ab Nasir, A. F., Seok Nee, E., Sern Choong, C., Shahrizan Abdul Ghani, A., Abdul Majeed, A. P. P., Adam, A., & Furqan, M. (2020). Text-based emotion prediction system using machine learning approach. *IOP Conference Series: Materials Science and Engineering*, 769(1). <https://doi.org/10.1088/1757-899X/769/1/012022>
- Acheampong, F. A., Wenyu, C., & Nunoo-Mensah, H. (2020). Text-based emotion detection: Advances, challenges, and opportunities. *Engineering Reports*, 2(7). <https://doi.org/10.1002/eng2.12189>
- Agus Setiawan, H., & Yuliansyah, H. (2024). Aspect-Based Sentiment Analysis of User Reviews on the Game "Honkai: Star Rail"

Using Naïve Bayes Classifier. *SISTEMASI*, 13(5), 1956. <https://doi.org/10.32520/stmsi.v13i5.4343>

- Arifian, A., Astuti, R., & Muhamad Basysyar, F. (2024). Analisis Sentimen Opini Supporter Pengguna Youtube terhadap Sistem Pembelian Tiket Pertandingan Persib menggunakan Metode Naïve Bayes. *Jurnal Informatika Dan Rekayasa Perangkat Lunak*, 6(1), 250–257. <https://doi.org/10.36499/jinrpl.v6i1.10310>
- Ashraf, N., Khan, L., Butt, S., Chang, H.-T., Sidorov, G., & Gelbukh, A. (2022). Multi-label emotion classification of Urdu tweets. *PeerJ Computer Science*, 8, e896. <https://doi.org/10.7717/peerj-cs.896>
- Azam, N., Ahmad, T., & Ul Haq, N. (2021). Automatic emotion recognition in healthcare data using supervised machine learning. *PeerJ Computer Science*, 7, e751. <https://doi.org/10.7717/peerj-cs.751>
- Bijaksana Putra Negara, A., Muhandi, H., Sajid, F., & DrHHadari Nawawi, J. (2021). Perbandingan Algoritma Klasifikasi terhadap Emosi Tweet Berbahasa Indonesia. *JEPIN (Jurnal Edukasi Dan Penelitian Informatika)*, 7(2). <https://doi.org/10.26418/jp.v7i2>
- Cahyaningtyas, C., Nataliani, Y., & Widiastari, I. R. (2021). Analisis Sentimen Pada Rating Aplikasi Shopee Menggunakan Metode Decision Tree Berbasis SMOTE. *AITI: Jurnal Teknologi Informasi*, 18(2), 173–184. <https://doi.org/10.24246/aiti.v18i2.173-184>
- Chowanda, A., Sutoyo, R., Meiliana, & Tanachutiwat, S. (2021). Exploring Text-based Emotions Recognition Machine Learning Techniques on Social Media Conversation. *Procedia Computer Science*, 179, 821–828. <https://doi.org/10.1016/j.procs.2021.01.099>
- Chowdhary, K. R. (2020). Natural Language Processing. In *Fundamentals of Artificial Intelligence* (pp. 603–649). Springer India. https://doi.org/10.1007/978-81-322-3972-7_19
- Depari, D. H., Widiastiwi, Y., & Santoni, M. M. (2022). Perbandingan Model Decision Tree, Naive Bayes dan Random Forest untuk Prediksi Klasifikasi Penyakit Jantung. *Informatik : Jurnal Ilmu*

- Komputer*, 18(3), 239. <https://doi.org/10.52958/iftk.v18i3.4694>
- Elinda, E., Yuliansyah, H., & Latiffi, M. I. A. (2024). Sentiment Analysis of the Sheikh Zayed Grand Mosque's Visitor Reviews on Google Maps Using the VADER Method. *International Journal of Advances in Data and Information Systems*, 5(1), 71–84. <https://doi.org/10.59395/ijadis.v5i1.1320>
- Fernandes, J. V. M. R., Alexandria, A. R. de, Marques, J. A. L., Assis, D. F. de, Motta, P. C., & Silva, B. R. dos S. (2024). Emotion Detection from EEG Signals Using Machine Deep Learning Models. *Bioengineering*, 11(8). <https://doi.org/10.3390/bioengineering11080782>
- Gunawan, L., Anggreainy, M. S., Wiha, L., Santy, Lesmana, G. Y., & Yusuf, S. (2023). Support vector machine based emotional analysis of restaurant reviews. *Procedia Computer Science*, 216, 479–484. <https://doi.org/10.1016/j.procs.2022.12.160>
- Kusal, S., Patil, S., Choudrie, J., Kotecha, K., Vora, D., & Pappas, I. (2022). A Review on Text-Based Emotion Detection -- Techniques, Applications, Datasets, and Future Directions. *ArXiv*, 2205. <https://doi.org/10.48550/arXiv.2205.03235>
- Liu, X., Shi, T., Zhou, G., Liu, M., Yin, Z., Yin, L., & Zheng, W. (2023). Emotion classification for short texts: an improved multi-label method. *Humanities and Social Sciences Communications*, 10(1), 306. <https://doi.org/10.1057/s41599-023-01816-6>
- Machová, K., Szabóová, M., Paralič, J., & Mičko, J. (2023). Detection of emotion by text analysis using machine learning. *Frontiers in Psychology*, 14. <https://doi.org/10.3389/fpsyg.2023.1190326>
- Maruf, A. Al, Ziyad, Z. M., Haque, Md. M., & Khanam, F. (2022). Emotion Detection from Text and Sentiment Analysis of Ukraine Russia War using Machine Learning Technique. *International Journal of Advanced Computer Science and Applications*, 13(12), 2022. <https://doi.org/10.14569/IJACSA.2022.01312101>
- Nandwani, P., & Verma, R. (2021). A review on sentiment analysis and emotion detection from text. *Social Network Analysis and Mining*, 11(1), 81. <https://doi.org/10.1007/s13278-021-00776-6>
- Okta, B., Miranda, S., Yuliansyah, H., & Biddinika, M. K. (2024). Machine Translation Indonesian Bengkulu Malay Using Neural Machine Translation-LSTM. *IJCCS (Indonesian Journal of Computing and Cybernetics Systems)*, 18(3), 1–5. <https://doi.org/https://doi.org/10.22146/ijccs.98384>
- Patel, N., Patel, F., & Kumar Bharti, S. (2022). Live Emotion Verifier for Chat Applications Using Emotional Intelligence. In *Smart Innovation, Systems and Technologies*, 267, 11–19. Springer Science and Business Media Deutschland GmbH. https://doi.org/10.1007/978-981-16-6616-2_2
- Pramesti, L. A., & Pratiwi, N. (2023). Analisis Sentimen Twitter Terhadap Program MBKM Menggunakan Decision Tree dan Support Vector Machine. *Journal of Information System Research (JOSH)*, 4(4), 1145–1154. <https://doi.org/10.47065/josh.v4i4.3807>
- Rahayu, K., Fitria, V., Septhya, D., Rahmaddeni, R., & Efrizoni, L. (2023). Klasifikasi Teks untuk Mendeteksi Depresi dan Kecemasan pada Pengguna Twitter Berbasis Machine Learning. *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, 3(2), 108–114. <https://doi.org/10.57152/malcom.v3i2.780>
- Rohman, A. N., Utami, E., & Raharjo, S. (2019). Deteksi Kondisi Emosi pada Media Sosial Menggunakan Pendekatan Leksikon dan Natural Language Processing. *Eksplora Informatika*, 9(1), 70–76. <https://doi.org/10.30864/eksplora.v9i1.277>
- Rokhman, K. A., Berlilana, B., & Arsi, P. (2021). Perbandingan Metode Support Vector Machine Dan Decision Tree Untuk Analisis Sentimen Review Komentar Pada Aplikasi Transportasi Online. *Journal of Information System Management (JOISM)*, 3(1), 1–7. <https://doi.org/10.24076/JOISM.2021v3i1.341>
- Sinaga, H. H., & Agustian, S. (2022). Pebandingan Metode Decision Tree dan XGBoost untuk Klasifikasi Sentimen Vaksin Covid-19 di Twitter. *Jurnal Nasional Teknologi Dan Sistem Informasi*, 8(3), 107–114. <https://doi.org/10.25077/TEKNOSI.v8i3.2022.107-114>

- Sondakh, D. E., Maringka, R. C., Ayorbaba, F. P., Mangi, J. S. C. B. T., & Pungus, S. R. (2023). Emotion Mining User Review of the BRImo Mobile Banking Application Using the Decision Tree Algorithm. *Jurnal Sisfokom (Sistem Informasi Dan Komputer)*, 12(3), 350–355. <https://doi.org/10.32736/sisfokom.v12i3.1721>
- Sontayasara, T., Jariyapongpaiboon, S., Promjun, A., Seelpipat, N., Saengtabtim, K., Tang, J., & Leelawat, N. (2021). Twitter Sentiment Analysis of Bangkok Tourism During COVID-19 Pandemic Using Support Vector Machine Algorithm. *Journal of Disaster Research*, 16(1), 24–30. <https://doi.org/10.20965/jdr.2021.p0024>
- Susandri, S., Defit, S., & Tajuddin, M. (2023). Sentiment Labeling And Text Classification Machine Learning For Whatsapp Group. *JITK (Jurnal Ilmu Pengetahuan Dan Teknologi Komputer)*, 9(1), 119–125. <https://doi.org/10.33480/jitk.v9i1.4201>
- Syafia, A. N., Hidayattullah, M. F., & Suteddy, W. (2023). Studi Komparasi Algoritma SVM Dan Random Forest Pada Analisis Sentimen Komentar Youtube BTS. *Jurnal Informatika: Jurnal Pengembangan IT*, 8(3), 207–212. <https://doi.org/10.30591/jpit.v8i3.5064>
- Thomas, S., Yuliana, & Noviyanti. P. (2021). Study Analisis Metode Analisis Sentimen pada YouTube. *Journal of Information Technology*, 1(1), 1–7. <https://doi.org/10.46229/jifotech.v1i1.201>
- Wulan, P. P., & Basri, H. (2024). Analisis Sentimen Terhadap Layanan Nasabah Bank Menggunakan Teknik Klasifikasi Naive Bayes. *Jurnal Kecerdasan Buatan Dan Teknologi Informasi*, 3(2), 68–74. <https://doi.org/10.69916/jkbt.v3i2.131>
- Yuliansyah, H., Wahyuni Sukei, T., Asti Mulasari, S., & Nur Syamilah Wan Ali, W. (2023). Bulletin of Social Informatics Theory and Application Artificial intelligence in malnutrition research: a bibliometric analysis. *Bulletin of Social Informatics Theory and Application*, 7(1), 32–42. <https://doi.org/10.31763/businta.73i1.605>

