

ANALISIS KLASIFIKASI GENRE MUSIK MENGGUNAKAN ALGORITMA SUPPORT VECTOR MECHINE (SVM)

Gion Andrian

Program Studi Teknik Informatika, Fakultas Teknologi Informasi, Universitas Tarumanagara
Jl. Letjen S.Parmar No.1, Jakarta Barat, DKI Jakarta, Indonesia 11410

E-mail: gion.535200024@stu.untar.ac.id

ABSTRAK

Klasifikasi genre musik dengan menggunakan 14 parameter, model terbaik adalah Support Vector Machine (SVM) dengan kernel "poly". Model SVM ini berhasil mencapai akurasi tinggi sekitar 0,5571 dan memiliki nilai F1-Score yang tinggi, menjadikannya pilihan utama dalam penelitian ini. Selain itu, pengujian menunjukkan bahwa pembagian dataset dengan rasio 60:40 dan 70:30 adalah yang paling efektif dalam meningkatkan kinerja model. Dalam perbandingan dengan algoritma standar, SVM pada rasio dataset 70:30 mencapai akurasi tertinggi dengan nilai 0,5440, F1-Score 0,5396, Recall 0,5440, dan Precision 0,5450. Ini diikuti oleh K-Nearest Neighbors (K-NN) dan Decision Tree.

Hasil penelitian secara konsisten menunjukkan bahwa SVM dengan kernel "poly" memberikan hasil terbaik dalam klasifikasi genre musik. Decision Tree dengan aktivasi criterion "gini" memiliki akurasi yang kompetitif, yaitu sekitar 0,5257. Sementara itu, K-NN dengan jumlah tetangga (k) sekitar 10 memberikan akurasi sekitar 0,4638, menjadikannya pilihan yang kurang efektif dalam konteks ini. Penelitian ini memberikan wawasan penting dalam penggunaan metode klasifikasi genre musik dengan menggunakan algoritma machine learning.

Kata kunci— K-Nearest Neighbors (K-NN), Support Vector Machine (SVM), Decision Tree, Genre musik

ABSTRACT

Music genre classification using 14 parameters, the best model is Support Vector Machine (SVM) with a "poly" kernel. This SVM model managed to achieve high accuracy of around 0.5571 and has a high F1-Score value, making it the main choice in this research. In addition, testing shows that dividing the dataset in a ratio of 60:40 and 70:30 is the most effective in improving model performance. In comparison with the standard algorithm, SVM on a dataset ratio of 70:30 achieved the highest accuracy with a value of 0.5440, F1-Score 0.5396, Recall 0.5440, and Precision 0.5450. This is followed by K-Nearest Neighbors (K-NN) and Decision Tree.

Research results consistently show that SVM with the "poly" kernel provides the best results in music genre classification. Decision Tree with activation criterion "gini" has competitive accuracy, which is around 0.5257. Meanwhile, K-NN with a number of neighbors (k) of around 10 provides an accuracy of around 0.4638, making it a less effective choice in this context. This research provides important insights into the use of music genre classification methods using machine learning algorithms.

Keywords— K-Nearest Neighbors (K-NN), Support Vector Machine (SVM), Decision Tree, Genre music

1. PENDAHULUAN

Musik adalah ekspresi batin manusia yang dinyatakan melalui suara yang terorganisir dengan melodi atau ritme, serta memiliki unsur harmoni yang memikat. Musik berupa kumpulan nada yang teratur, baik dalam bentuk vokal atau nyanyian, maupun dalam bentuk dimainkan dengan alat musik, dan seringkali keduanya digabungkan [1]. Teknologi dan media massa memperluas akses kita ke beragam musik, menciptakan era globalisasi musik. Reaksi terhadap musik sangat subjektif, dipengaruhi oleh faktor individu dan situasional, menjadi fokus psikologi musik, sebuah bidang yang memahami bagaimana musik memengaruhi pengalaman dan perilaku

manusia. Ini mencakup elemen-elemen musik, produksi, pemahaman, dan dampaknya pada kehidupan kita, meskipun pengalaman ini sangat bervariasi karena berbagai faktor [2].

Berdasarkan penelitian meta-analisis yang melibatkan 40 studi dengan individu, musik telah terbukti memiliki manfaat yang signifikan dalam mengurangi tingkat stres manusia. Selain itu, penelitian ini juga mengungkapkan bahwa musik memiliki kemampuan untuk mempengaruhi berbagai aspek stres, termasuk aspek fisik, perilaku, dan psikologis, sehingga membantu mengurangi tingkat stres secara keseluruhan [3].

Banyaknya jenis musik, baik tradisional maupun modern, yang menggabungkan elemen-elemen dari keduanya, diperlukan pengelompokan musik berdasarkan karakteristik serupa. Ini membantu pendengar dalam memahami dan mengkategorikan musik berdasarkan prinsip-prinsip seperti ritme, citra artis, atau posisi dalam tangga lagu [4]. Genre-genre populer seperti pop, rock, jazz, dan metal telah menjadi pilihan mendengar yang umum. Namun, setiap genre ini memiliki penggemar sendiri, dan dengan bertambahnya musisi dan lagu-lagu baru yang terus bermunculan, klasifikasi genre menjadi semakin penting.

Saat ini, penggemar musik dapat dengan mudah menikmati berbagai lagu melalui berbagai platform, termasuk layanan streaming musik yang sangat digemari. Layanan ini memberikan kenyamanan kepada pengguna dalam menemukan lagu sesuai dengan preferensi genre yang mereka sukai dengan mengelompokkan berdasarkan genre tertentu. Meskipun klasifikasi genre musik dapat dilakukan secara manual, dan bantuan teknologi seperti machine learning sehingga lebih cepat dan akurat.

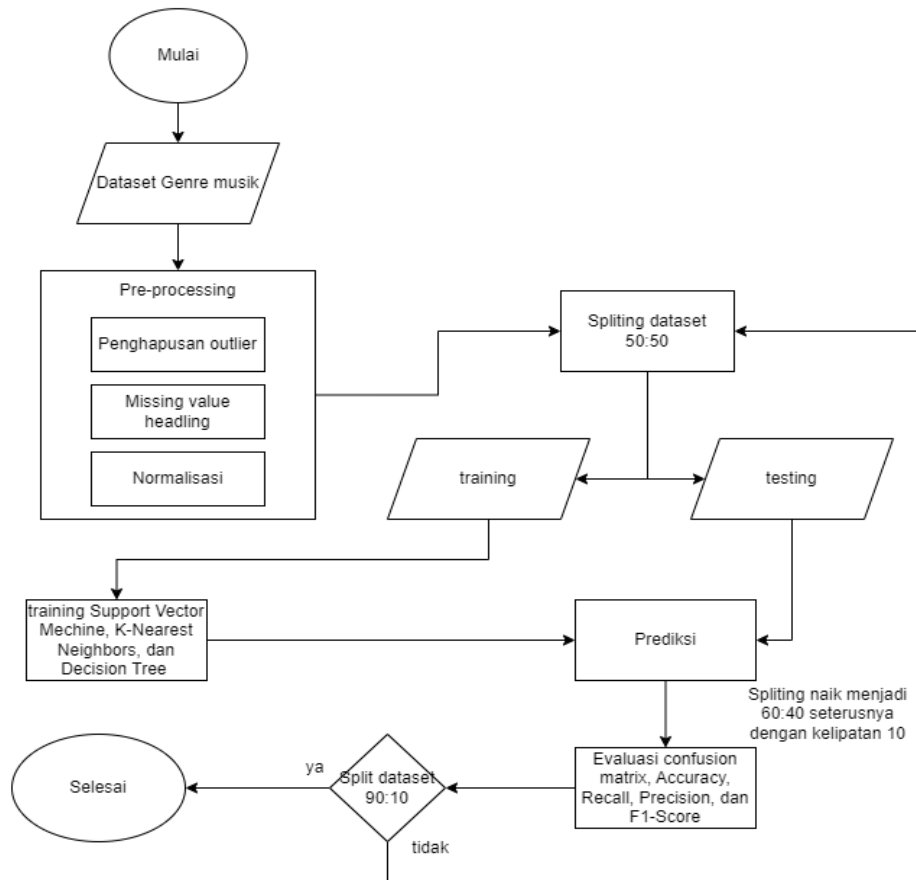
Machine learning adalah bagian dari AI, yang digunakan untuk mengetahui sebuah pola yang salah satunya digunakan untuk klasifikasi untuk pengelompokan data ke dalam kategori berdasarkan karakteristik tertentu [5]. Sebagai contoh penerapan machine learning dalam klasifikasi genre, berdasarkan penelitian sebelumnya berhasil mengklasifikasikan 1000 trek audio dengan durasi 30 detik, frekuensi 22050Hz, dan format mono 16 bit dalam file .wav, menggunakan algoritma Extra Trees, dengan tingkat akurasi tertinggi mencapai 72,5% [6].

Pada penelitian ini, bertujuan untuk melakukan klasifikasi genre dari 5005 lagu yang termasuk dalam 10 genre yang berbeda, seperti '*Electronic*', '*Anime*', '*Jazz*', '*Alternative*', '*Country*', '*Rap*', '*Blues*', '*Rock*', '*Classical*', dan '*Hip-Hop*'. Algoritma klasifikasi akan menggunakan metode *Support Vector Machine* (SVM) dan akan dibandingkan dengan algoritma K-Nearest Neighbors (KNN) serta Decision Tree untuk mendapatkan hasil terbaik yang dapat mencerminkan kemampuan setiap algoritma dalam melakukan klasifikasi genre musik. Penerapan algoritma ini akan memanfaatkan 14 fitur yang tersedia dari dataset.

2. METODE PENELITIAN

Sebelum implementasi menggunakan algoritma klasifikasi seperti Support Vector Machine, K-Nearest Neighbors, dan Decision Tree, diperlukan perancangan flowchart untuk memvisualisasikan alur program secara lebih jelas. Flowchart pada Gambar 1 menunjukkan bahwa dataset akan melalui tahapan preprocessing, termasuk penghapusan data outlier, penanganan missing value, dan normalisasi, sebelum dilakukan pembagian data menjadi dua bagian: data latih (training) dan data uji (testing).

Data latih akan digunakan untuk melatih model dengan ketiga algoritma tersebut, dan hasil pelatihan akan digunakan untuk memprediksi data uji. Setelahnya, dilakukan evaluasi menggunakan confusion matrix, serta metrik evaluasi seperti akurasi, presisi, recall, dan F1-score. Karena ada lima skenario pembagian data yaitu 50:50, 60:40, 70:30, 80:20, dan 90:10, maka akan dilakukan iterasi atau looping sebelum program berakhir untuk mengatasi semua skenario yang ada.



Gambar 1 Flowchart sistem penelitian

2.1 Data

Data yang digunakan dalam penelitian ini menggunakan dataset pada situs kaggle yang diakses pada link: <https://www.kaggle.com/datasets/vicsuperman/prediction-of-music-genre/data> dengan jumlah 5005 data yang diperoleh dari platform streaming Spotify, yang mencakup musik rilisan dari tahun 2001 hingga 2018 dengan jumlah dalam arti genre musik sebanyak 10 'Electronic', 'Anime', 'Jazz', 'Alternative', 'Country', 'Rap', 'Blues', 'Rock', 'Classical', dan 'Hip- Hop' yang memiliki 18 fitur antara lain:

- instance_id: Ini adalah ID unik yang diberikan untuk setiap rekaman musik dalam dataset, memungkinkan identifikasi yang unik
- artist_name: Fitur ini mencantumkan nama artis yang menciptakan rekaman musik tersebut, memberikan informasi tentang pencipta musik.
- track_name: Ini adalah nama lagu atau judul dari rekaman musik, memberikan informasi tentang judul musik.
- popularity: Fitur ini mengindikasikan seberapa populer musik tersebut berdasarkan metrik tertentu, memberikan gambaran tentang popularitasnya di antara pendengar.
- acousticness: Mengukur sejauh mana musik bersifat akustik atau menggunakan elemen-elemen elektronik, memberikan gambaran tentang karakter akustiknya.
- danceability: Menggambarkan seberapa cocok musik untuk tarian atau gerakan tubuh, mengindikasikan tingkat kecocokan untuk aktivitas ini.
- duration_ms: Ini adalah durasi musik dalam milidetik, memberikan informasi tentang panjangnya lagu.
- energy: Mengukur tingkat energi atau kegembiraan dalam musik, memberikan gambaran tentang energi yang terpancar dari lagu tersebut.
- instrumentalness: Indikator sejauh mana musik bersifat instrumental tanpa vokal, menggambarkan karakter instrumentalnya.

- j. key: Ini adalah kunci musik yang digunakan dalam rekaman, memberikan informasi tentang nada dasar musik.
- k. liveness: Mengindikasikan seberapa "hidup" atau terasa seperti pertunjukan langsung dari musik tersebut.
- l. loudness: Mengukur kekuatan suara atau volume musik, menggambarkan tingkat volume yang dihasilkan.
- m. mode: Indikator apakah musik berada dalam mode mayor atau minor, mempengaruhi mood musik.
- n. speechiness: Mengukur tingkat wicara atau kata-kata dalam musik, memberikan gambaran tentang unsur wicara dalam lagu.
- o. tempo: Ini adalah tempo atau kecepatan musik dalam BPM (beat per menit), menggambarkan seberapa cepat atau lambat musik tersebut.
- p. obtained_date: Tanggal data diperoleh atau dicatat, memberikan informasi tentang kapan data tersebut diambil.
- q. valence: Mengukur tingkat emosi positif dalam musik, mengindikasikan seberapa bahagia atau positifnya mood musik.
- r. music_genre: Fitur kelas yang merupakan genre musik yang telah ditentukan, yang menjadi tujuan klasifikasi dalam analisis data musik ini.

2.2 Data Pre-processing

Outlier merupakan data yang memiliki karakteristik sangat berbeda dari data lain dalam dataset. Outlier dapat mempengaruhi statistik deskriptif dan analisis statistik, serta dapat memberikan informasi tentang kemungkinan adanya masalah dalam data atau proses yang diamati [7]. Oleh karena itu data outlier perlu dihapus karena dapat mempengaruhi hasil klasifikasi yang dimana tidak semua algoritma klasifikasi tahan terhadap data outlier.

Missing value merupakan salah satu masalah yang mungkin ada dalam dataset missing value merupakan kosongnya sebuah data dari beberapa fitur dalam dataset. Missing value bisa ditangani dengan menggunakan metode Listwise Deletion, menghapus kasus yang memiliki data yang hilang. Mean Substitution, menggantikan data yang hilang dengan nilai rata-rata dari variabel yang sama. Dalam penelitian ini menggantikan data yang hilang dengan modus(nilai terbanyak dalam satu fitur) [8].

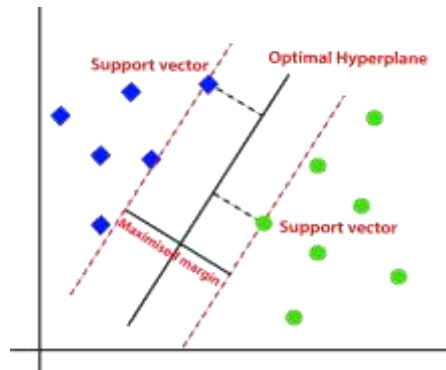
Normalisasi data sangat penting dalam klasifikasi normalisasi merupakan hasil penggabungan beberapa data sehingga data memiliki skala yang sama dan dapat dibandingkan. Salah satu teknik normalisasi yang paling terkenal adalah Min-Max scaling dengan rentang nilai 0 hingga 1 dan z-score normalization mengubah skor relevansi ke distribusi normal standar dengan mean (rata-rata) 0 dan deviasi standar (standard deviation) 1 [9]. Dalam penelitian ini menggunakan Min-Max scaling dengan persamaan 1

$$X_{norm} = \frac{X - X_{min}}{X_{max} - X_{min}} \quad (1)$$

2.3 Algoritma

2.3.1 Support Vector Machine (SVM)

Support Vector Machine (SVM) adalah sebuah algoritma klasifikasi yang dikembangkan oleh Vladimir Vapnik di tahun 1990-an. SVM sebuah metode supervised learning yang digunakan untuk melakukan klasifikasi serta regresi. Prinsip kerja SVM dengan melakukan pencarian hyperplane (bidang linear) pada gambar 1 garis berwarna hitam dalam ruang yang memiliki dimensi tinggi dengan efektif dalam memisahkan dua kelas data yang berbeda garis putus-putus dalam gambar 1. Dalam konteks klasifikasi, algoritma SVM bertujuan untuk memisahkan data ke dalam dua kelas yang berbeda dengan memaksimalkan margin (jarak) antara kedua kelas tersebut [10].



Gambar 2 Visualisasi SVM (sumber : pemrogramanmatlab.com)

$$f(x) = \sum_{i=1}^N a_i y_i K(x_i, x) + b \quad (2)$$

SVM memiliki rumus pada persamaan 2 yang dimana $K(x_i, x)$ sebagai karnel. Karnel yang digunakan dalam klasifikasi karnel linear merupakan karnel yang paling sederhana yang memebentuk garis lurus pemisah antara kelas dan sangat cocok dengan data yang dapat dipisahkan oleh garis lurus seperti data dengan 2 kelas untuk rumus dapat dilihat pada persamaan 3 [11].

$$K(x_i, x) = x_i \cdot x \quad (3)$$

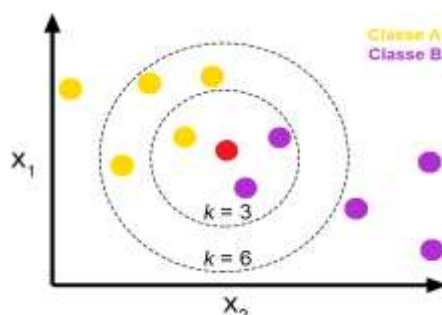
Selain karnel linear SVM memiliki karnel Polinomial merupakan hasil klasifikasi berdasarkan arah dan magnitudo vektor dalam ruang dimensi rendah dengan menggunakan fungsi polinomial untuk transformasi data dengan persamaan 4 [11].

$$K(x_i, x) = (x_i \cdot x + c)^d \quad (4)$$

Karnel lainnya adalah *Kernel Gaussian* (RBF) dalam SVM memungkinkan pemisahan data non- linear dan memilik parameter σ untuk mengontrol sejauh pengaruh setiap titik data pada pemisahan dengan persamaan 5 [12].

$$K(x_i, x_j) = \frac{\|x_i - x_j\|^2}{2\sigma^2} \quad (5)$$

2.3.2 K-Nearest Neighbors (K-NN)



Gambar 3 Visualisasi K-NN (Sumber: medium.com)

K-Nearest Neighbors (KNN) adalah algoritma yang digunakan untuk klasifikasi maupun regresi dengan cara kerja mencari k (jumlah tetangga terdekat). Sebagai contoh, jika $k=3$, algoritma akan mencari tiga anggota terdekat dari kluster data yang ada. Namun, jika $k=6$, cakupan pencarian akan lebih luas, mencari jarak terdekat dari enam anggota kluster data. Dengan demikian, nilai k mempengaruhi seberapa banyak tetangga terdekat yang akan digunakan dalam pengambilan keputusan. Sebagai contoh pada gambar 3 jika $k=3$ maka akan mencari jarak terdekat dari 3 anggota kluster dan jika 6 maka cakupan akan lebih luas mencari jarak terdekat dari 6 anggota kluster. Rumus perhitungan data di lihat pada persamaan 6 [13].

$$d(x_i, x_j) = \sqrt{\sum_{k=1}^n (x_{i,k} - x_{j,k})^2} \quad (6)$$

Perhitungan jarak antar data dalam kluster biasanya menggunakan metrik jarak seperti *Euclidean distance* (jarak Euclidean), *Manhattan distance* (jarak Manhattan).

2.3.3 Decision Tree

Decision tree adalah metode hirarkis yang digunakan untuk membagi ruang fitur dataset ke dalam subspace yang berbeda untuk klasifikasi, di mana sampel baru diklasifikasikan dengan mengikuti aturan keputusan di sepanjang jalur pohon tersebut [14]. Cara kerjanya adalah dengan membagi dataset menjadi beberapa kelompok berdasarkan aturan keputusan yang mengukur kualitas pemisahan sampel. Pemisahan ini dilakukan dengan memilih atribut terbaik dan nilai ambang yang paling cocok, yang kemudian membentuk cabang-cabang pohon. Proses *Decision Tree* adalah melibatkan pemilihan atribut terbaik, pembagian data, dan pengambilan keputusan hingga mencapai node daun yang mewakili prediksi akhir. Rumus *Information Gain* digunakan untuk memilih atribut terbaik dengan rumus again pada persamaan 7 [15].

$$Gain(S, A) = Entropy(S) - \sum_{v \in Values(A)} \frac{|s_v|}{|S|} \cdot Entropy(s_v) \quad (6)$$

2.4 Evaluasi

2.4.1 Confusion Matrix

Confusion Matrix adalah matriks berupa bilangan bulat dan non-negatif menghitung berapa banyak data yang diklasifikasikan dengan benar dan berapa banyak yang salah [16]. Contoh pada gambar 3 adalah matrix 2 kali 2 yang dimana merupakan confusion matrix dari hasil klasifikasi 2 kelas (0 dan 1). True Positive (TP) merupakan objek positif atau kelas 1 yang berhasil diprediksi benar, True Negative (TN) merupakan objek negatif atau kelas 0 yang berhasil diprediksi benar, False Positive (FP) merupakan objek positif atau kelas 1 yang salah diprediksi, dan False Negative (FN) merupakan objek negatif atau kelas 0 yang salah di prediksi.

		Actual Values	
		Positive (1)	Negative (0)
Predicted Values	Positive (1)	TP	FP
	Negative (0)	FN	TN

Gambar4 Confusion Matrix

2.4.2 Akurasi (*Accuracy*)

Akurasi merupakan pendekatan untuk menilai sejauh mana metode klasifikasi dapat mencapai hasil yang konsisten dengan karakteristik sebenarnya dari tes yang digunakan [17]. Akurasi adalah suatu metrik yang mengukur sejauh mana metode klasifikasi berhasil dalam memprediksi dengan benar kasus positif dan negatif dalam kasus klasifikasi dua kelas dengan persamaan 1.

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (7)$$

2.4.3 Presisi (*Precision*)

Presisi adalah metode untuk mengukur tingkat keakuratan dalam mengidentifikasi objek sebagai kelas positif (biasanya kelas 1 dalam klasifikasi dua kelas). Presisi di hitung berdasarkan dengan membagi TP dibagi dengan jumlah prediksi positif TP dan FP dengan persamaan 1. Hasil presisi akan memberikan gambaran sejauh mana hasil prediksi positif di klasifikasikan dengan benar [18].

$$Precision = \frac{TP}{TP + FP} \quad (8)$$

2.4.4 Recall (*Sensitivity* atau *True Positive Rate*)

Recall merupakan metode pengukuran untuk mengidentifikasi semua hasil positif yang benar dengan membagi TP dengan jumlah positif yang sebenarnya ada TP dan FN dengan persamaan 1. Recall menggambarkan sejauh mana sistem "mengingat" hasil positif yang harus diidentifikasi [18].

$$Recall = \frac{TP}{TP + FN} \quad (9)$$

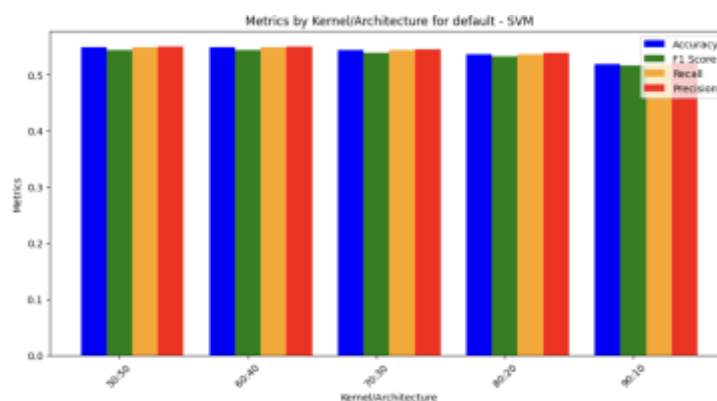
2.4.5 F1-Score

F1 score adalah metrik evaluasi yang menggabungkan baik presisi dan untuk memberikan gambaran kinerja yang seimbang antara mengidentifikasi hasil positif yang benar dan mengidentifikasi semua hasil positif yang relevan [18].

$$f1 = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (10)$$

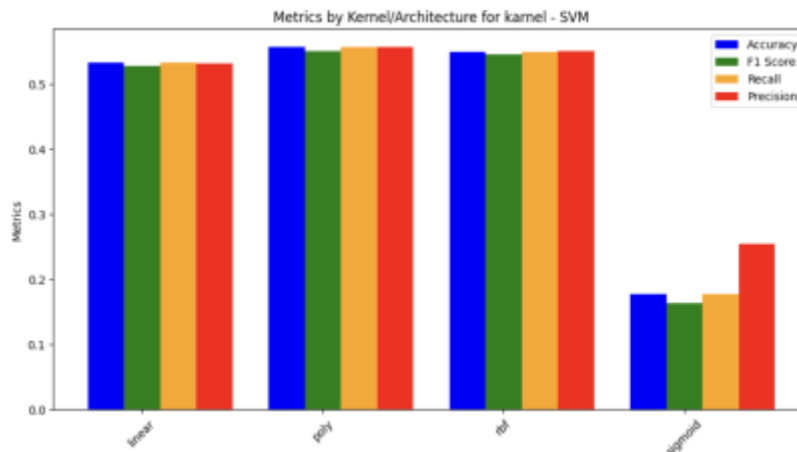
3 HASIL DAN PEMBAHASAN

3.4 Support Vector Machine (SVM)



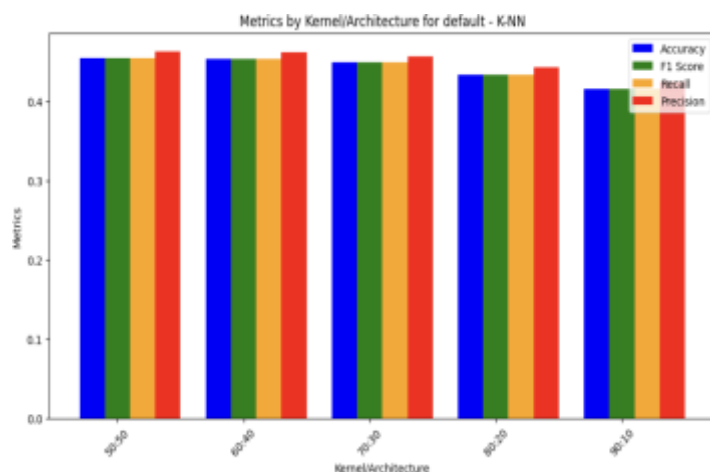
Gambar 5 grafik perbandingan splitting SVM

Hasil evaluasi kinerja model Support Vector Machine (SVM dengan scenario menggunakan pembagian dataset memiliki pola pola yang cukup konsisten dapat dilihat pada gambar 5. Dalam eksperimen ini, dataset dibagi dengan rasio data latih dan uji (50:50 hingga 90:10). Terlihat penurunan yang nyata dalam semua metrik evaluasi, termasuk akurasi, presisi, recall, dan F1-Score. Akurasi mulai dari 0,54932 pada rasio 50:50 hingga turun menjadi 0,51987 pada rasio 90:10, sementara presisi, recall, dan F1-Score juga mengikuti tren penurunan yang sama. Dengan kata lain, semakin sedikit data data uji makan akan cenderung mengalami penurunan yang mungkin terjadi karena overftting. Algoritma SVM cenderung menurun dengan akurasi paling tinggi pada perbandingan 50:50 dan cenderung mirip hingga perbandingan 70:30 dan terendah di 90:10.

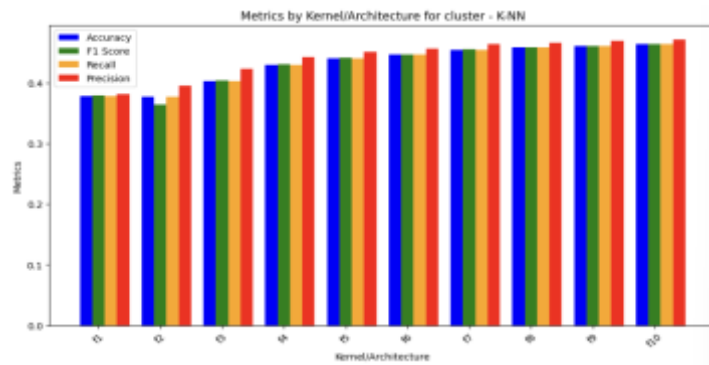


Gambar 6 grafik perbandingan karnel SVM

Hasil pengujian dengan berbagai jenis kernel, seperti linear, polynomial (poly), radial basis function (rbf) atau gaussian , dan sigmoid mendapatkan hasil perbandingan hasil pada gambar 6. Kernel polynomial memiliki akurasi tinggi (0,5571) dan nilai presisi (0,5518), recall (0,5571), serta F1-Score (0,5573) yang tinggi sedangkan kernel linear juga menunjukkan kinerja yang baik dengan keseimbangan antara metrik evaluasi, menunjukkan akurasi (0,5503), presisi (0,5459), recall (0,5503), serta F1-Score (0,5512) yang kompetitif. Sementara itu, kernel RBF juga menghasilkan hasil yang bersaing dengan akurasi yang baik. Di sisi lain, kernel sigmoid berbadning jauh dengan ketiga karnel dengan akurasi sebesar 0,177933333. Kesimpulan utama adalah bahwa dalam pemilihan kernel SVM, Kernel polynomial mungkin menjadi pilihan terbaik dalam konteks ini. *K-Nearest Neighbors (K-NN)*



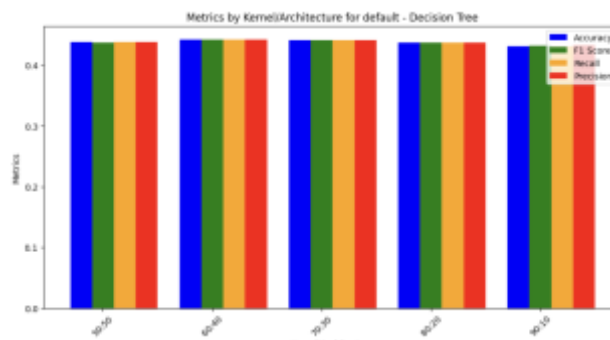
Gambar 7 grafik perbandinga rasio data K-NN



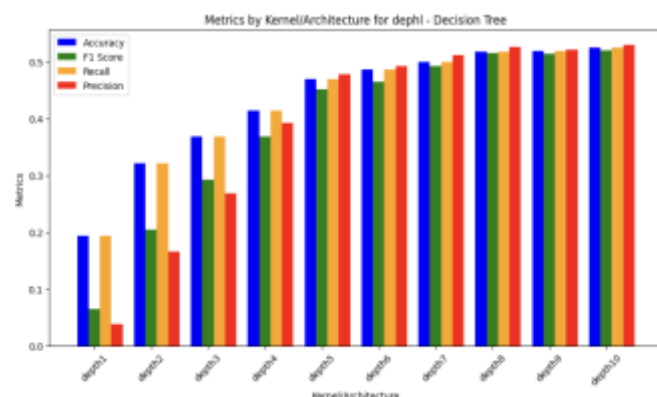
Gambar 8 grafik pebandingan jumlah K

Hasil evaluasi algoritma K-Nearest Neighbors (K-NN) dengan berbagai rasio pembagian dataset (50:50, 60:40, 70:30, 80:20, dan 90:10) menunjukkan dampak signifikan dari rasio tersebut terhadap kinerja model dapat dilihat pada gambar 7. Secara keseluruhan, terlihat bahwa semakin besar rasio data latih dan uji cenderung meningkat pada rasio 50:50, akurasi mencapai sekitar 0,45496, tetapi menurun seiring dengan peningkatan rasio hingga mencapai sekitar 0,41589 pada rasio 90:10. Dalam hal keseimbangan antara data pelatihan dan pengujian, rasio 60:40 dan 70:30 memberikan hasil yang mirip dan konsisten dari rasio 50:50 dengan akurasi sekitar 0,4537 dan 0,4494, masing-masing.

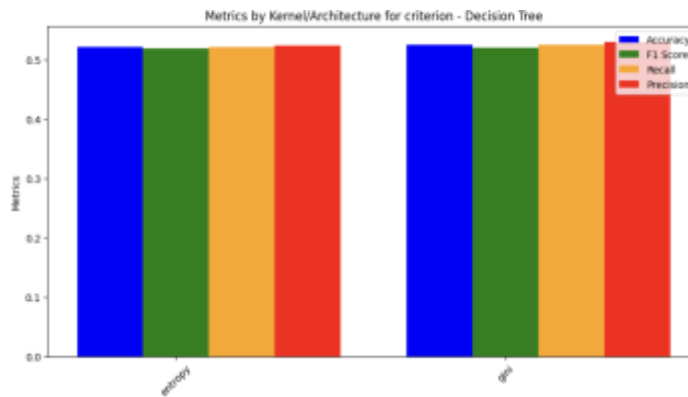
Hasil evaluasi K-NN dengan berbagai nilai k (1 hingga 10) menunjukkan bahwa, secara umum peningkatan jumlah neighbors atau K membawa dampak positif terhadap kinerja model perbandingan dapat dilihat pada gambar 8. Terlihat bahwa F1-Score, yang mencakup keseimbangan antara presisi dan recall, meningkat seiring dengan peningkatan jumlah k. Pada k1, F1-Score sekitar 0,3781, sementara pada k10, F1-Score mencapai sekitar 0,4639. Dengan kata lain, model K-NN menjadi lebih baik dalam tugas klasifikasi dengan penambahan jumlah cluster.



Gambar 9 grafik perbadigan rasio Decision Tree



Gambar 10 Grafik kedalaman Decision Tree



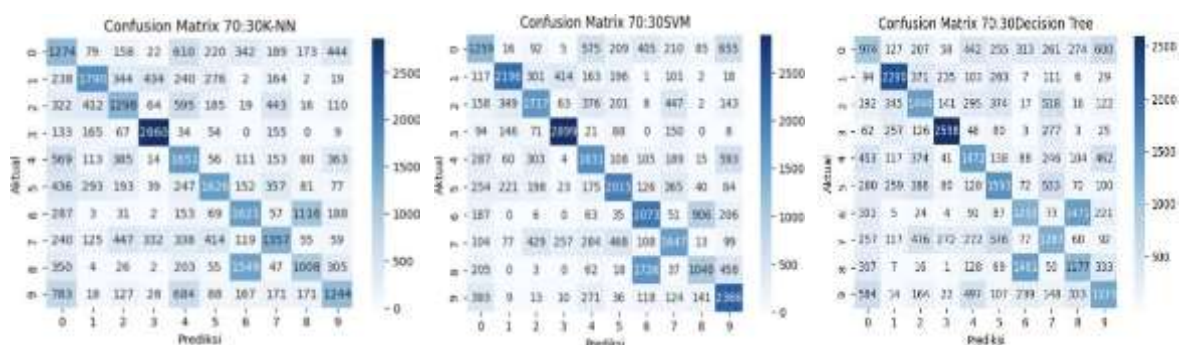
Gambar 11 grafik perbandingan entropy

Hasil evaluasi algoritma *Decision Tree* dengan berbagai rasio pembagian dataset (50:50, 60:40, 70:30, 80:20, dan 90:10) menunjukkan bahwa rasio pembagian dataset memengaruhi kinerja model secara signifikan dapat dilihat pada gambar 9. Secara umum, rasio data pelatihan terhadap data pengujian, akurasi model cenderung meningkat. Dalam kasus ini, model yang dilatih dengan rasio 60:40 dan 70:30 memberikan akurasi sekitar 0,442, sementara model dengan rasio 50:50 hanya mencapai akurasi sekitar 0,438. Namun, saat rasio data pelatihan sangat rendah (90:10), akurasi model mengalami penurunan signifikan menjadi sekitar 0,431.

Hasil *Decision Tree* dengan berbagai kedalaman (*depth*) menunjukkan perkembangan yang menarik pada gambar 10. Secara umum, terlihat bahwa seiring dengan peningkatan kedalaman pohon keputusan, kinerja model juga meningkat. Pada kedalaman 1, akurasi model hanya sekitar 0,1942, namun, dengan peningkatan kedalaman hingga 10, akurasi mencapai sekitar 0,5255. Selain itu, presisi, recall, dan F1-Score juga menunjukkan tren peningkatan seiring dengan peningkatan kedalaman pohon keputusan. Ini menunjukkan bahwa dalam kasus ini, menggunakan pohon keputusan dengan kedalaman yang lebih dalam dapat menghasilkan model yang lebih baik dalam tugas klasifikasi yang diberikan.

Hasil evaluasi kinerja model *Decision Tree* dengan dua kriteria pemilihan node, yaitu "entropy" dan "gini," menunjukkan perbedaan yang cukup menarik padagambar 11. Meskipun perbedaan antara kriteria "entropy" dan "gini" tidak signifikan, terlihat bahwa model yang menggunakan kriteria "gini" sedikit lebih unggul dalam hal akurasi dengan nilai sekitar 0,5257 dibandingkan dengan model "entropy" yang memiliki akurasi sekitar 0,5221. Namun, baik "entropy" maupun "gini" menghasilkan model dengan presisi, recall, dan F1-Score yang relatif serupa, meskipun "gini" mungkin sedikit lebih unggul dalam beberapa metrik. Dalam kasus ini, pemilihan kriteria "gini" tampaknya memberikan hasil yang sedikit lebih baik, tetapi perbedaannya tidak signifikan.

3.5 Confusion Matrix



Gambar 12 Confusion Matrix

Dapat dilihat dari hasil confusion matrix terdapat banyak kesalahan pada kelas 6 dan 8 pada semua algoritma hal tersebut terjadi dikarenakan genre musik yang memiliki karakteristik yang hampir mirip sehingga ke-3 algoritma gagal untuk mengidentifikasi genre musik pada kelas 6 dan 8.

4 KESIMPULAN

Berdasarkan hasil analisis klasifikasi genre musik dengan menggunakan 14 parameter, dapat disimpulkan bahwa model terbaik untuk tugas ini adalah model Support Vector Machine (SVM) dengan kernel "poly." Model SVM dengan kernel "poly" menunjukkan akurasi yang tinggi sekitar 0,5571 dan nilai F1-Score yang tinggi, sehingga menjadi pilihan terbaik dalam konteks ini. Selain itu, pemilihan rasio pembagian dataset juga mempengaruhi kinerja model. Berdasarkan pengujian didapatkan scenario pembagian dataset 60:40 dan 70:30 merupakan scenario paling baik. Berdasarkan perbandingan dengan algoritma standar pada rasio dataset 70:30 disimpulkan SVM memiliki akurasi tertinggi dengan 0,5440, F1-score 0,5396, Recall 0,5440, dan Precision 0,5450 diikuti oleh K-NN dan Decision tree.

Hasil penelitian disimpulkan bahwa SVM memiliki performa dan akurasi yang tinggi dan konsisten diikuti oleh Decision Tree dengan aktivasi criterion "gini" memiliki akurasi 0,5257 dan K-NN dengan perbedaan yang tidak jauh dengan jumlah $k=10$ memiliki akurasi 0,4638. Hasil tersebut berbeda dari penelitian sebelumnya K-NN dengan $k=3$ dan pengukuran jarak Euclidean mencapai akurasi 85,38%, sementara SVM memiliki akurasi 66,9% untuk klasifikasi genre musik yang berkisar dari 1 hingga 13, mewakili genre musik gamelan Bali [19]. Hasil penelitian lain memberikan hasil algoritma Decision Tree dan Extreme Gradient Boosting memiliki akurasi sekitar 51% dan 72% dengan 10 jenis genre dan 1000 dataset [20].

Hasil tersebut menunjukkan bahwa mencapai akurasi yang tinggi dalam pengklasifikasian genre musik masih menjadi tantangan, terutama karena kompleksitas parameter yang memengaruhi hasil klasifikasi dan variasi dalam penelitian yang berbeda. Dalam kesimpulan, perbaikan lebih lanjut diperlukan untuk mencapai akurasi yang lebih baik dalam klasifikasi genre musik.

DAFTAR PUSTAKA

- [1] S. Widhyatama, *Sejarah musik dan apresiasi seni*. PT Balai Pustaka (Persero), 2012.
- [2] A. Gabriellson, *Strong Experiences with Music: Music is much more than just music*. OUP Oxford, 2011. [Daring]. Tersedia pada: <https://books.google.co.id/books?id=qnGqpqoimc4C>
- [3] M. P. Dewi, "Studi Metaanalisis: Musik untuk menurunkan stres," *Jurnal Psikologi*, vol. 36, no. 2, hlm. 106–115, 2009.
- [4] P. G. Christenson dan J. B. Peterson, "Genre and gender in the structure of music preferences," *Communic Res*, vol. 15, no. 3, hlm. 282–301, 1988.
- [5] I. D. Id, *Machine Learning: Teori, Studi Kasus dan Implementasi Menggunakan Python*, vol. 1. Unri Press, 2021.
- [6] E. A. Laksana dan F. Sulianta, "Analisis dan studi komparatif algoritma klasifikasi genre musik," *Semnasteknomedia Online*, vol. 5, no. 1, hlm. 1–2, 2017.
- [7] S. Walfish, "A review of statistical outlier methods," *Pharmaceutical technology*, vol. 30, no. 11, hlm. 82, 2006.
- [8] H. Kang, "The prevention and handling of the missing data," *kja*, vol. 64, no. 5, hlm. 402–406, Mei 2013, doi: 10.4097/kjae.2013.64.5.402.
- [9] M. Montague dan J. A. Aslam, "Relevance score normalization for metasearch," dalam *Proceedings of the tenth international conference on Information and knowledge management*, 2001, hlm. 427–433.
- [10] M. A. Hearst, S. T. Dumais, E. Osuna, J. Platt, dan B. Scholkopf, "Support vector machines," *IEEE Intelligent Systems and their Applications*, vol. 13, no. 4, hlm. 18–28, 1998, doi: 10.1109/5254.708428.
- [11] A. Patle dan D. S. Chouhan, "SVM kernel functions for classification," dalam *2013 International Conference on Advances in Technology and Engineering (ICATE)*, 2013, hlm. 1–9. doi: 10.1109/ICAAdTE.2013.6524743.

- [12] S. S. Keerthi dan C. Lin, "Asymptotic Behaviors of Support Vector Machines with Gaussian Kernel," *Neural Comput*, vol. 15, no. 7, hlm. 1667–1689, 2003, doi: 10.1162/089976603321891855.
- [13] B. Rajagopalan dan U. Lall, "A k-nearest-neighbor simulator for daily precipitation and other weather variables," *Water Resour Res*, vol. 35, no. 10, hlm. 3089–3101, 1999.
- [14] A. J. Myles, R. N. Feudale, Y. Liu, N. A. Woody, dan S. D. Brown, "An introduction to decision tree modeling," *Journal of Chemometrics: A Journal of the Chemometrics Society*, vol. 18, no. 6, hlm. 275–285, 2004.
- [15] Y.-Y. Song dan L. U. Ying, "Decision tree methods: applications for classification and prediction," *Shanghai Arch Psychiatry*, vol. 27, no. 2, hlm. 130, 2015.
- [16] R. Susmaga, "Confusion Matrix Visualization," dalam *Intelligent Information Processing and Web Mining*, M. A. Kłopotek, S. T. Wierzchoń, dan K. Trojanowski, Ed., Berlin, Heidelberg: Springer Berlin Heidelberg, 2004, hlm. 107–116.
- [17] S. A. Livingston dan C. Lewis, "Estimating the consistency and accuracy of classifications based on test scores," *J Educ Meas*, vol. 32, no. 2, hlm. 179–197, 1995.
- [18] C. Goutte dan E. Gaussier, "A Probabilistic Interpretation of Precision, Recall and F- Score, with Implication for Evaluation," dalam *Advances in Information Retrieval*, D. E. Losada dan J. M. Fernández-Luna, Ed., Berlin, Heidelberg: Springer Berlin Heidelberg, 2005, hlm. 345– 359.
- [19] I. G. Harsemadi, "Perbandingan Kinerja Algoritma K-NN dan SVM dalam Sistem Klasifikasi Genre Musik Gamelan Bali," *INFORMATICS FOR EDUCATORS AND PROFESSIONAL: Journal of Informatics*, vol. 8, no. 1, hlm. 1–10, 2023.
- [20] S. Lutfiani, T. H. Saragih, F. Abadi, M. R. Faisal, dan D. Kartini, "Perbandingan Metode Extreme Gradient Boosting Dan Metode Decision Tree Untuk Klasifikasi Genre Musik," *Jurnal Informatika Polinema*, vol. 9, no. 4, hlm. 373–382, 2023.