

Model Rekomendasi Lagu Berbasis Genre Menggunakan Metode *Random Forest* Dan *Decision Tree*

Putri Alvina¹, Yuni Yamasari²,

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

putri.21051@mhs.unesa.ac.id

yuniyamasari@unesa.ac.id

Abstrak— Penelitian ini mengembangkan model sistem rekomendasi lagu berbasis genre menggunakan algoritma *Random Forest* dan *Decision Tree*. Proses pemodelan dimulai dengan analisis *feature importance* untuk mengidentifikasi sepuluh fitur audio utama yang paling berpengaruh terhadap klasifikasi genre. Evaluasi performa dilakukan menggunakan teknik *cross-validation* guna memastikan hasil yang konsisten dan dapat digeneralisasi. Hasil pengujian menunjukkan bahwa algoritma *Random Forest* memberikan performa yang lebih unggul dibandingkan *Decision Tree*, ditunjukkan oleh nilai akurasi, presisi, *recall*, dan *F1-score* yang lebih tinggi. Secara kuantitatif, model *Random Forest* mencatat rata-rata akurasi sebesar 83%, sedangkan *Decision Tree* hanya mencapai 79%. Keunggulan ini menegaskan bahwa *Random Forest* merupakan metode yang lebih efektif dan andal untuk digunakan dalam membangun model rekomendasi lagu berdasarkan genre.

Kata Kunci— Sistem Rekomendasi, Genre Musik, *Random Forest*, *Decision Tree*, *Feature Importance*, *Cross-Validation*, *Machine Learning*.

I. PENDAHULUAN

Era Digital dimulai ketika sekelompok insinyur dari Institut Fraunhofer untuk Sirkuit Terpadu di Jerman berhasil mengembangkan teknologi kompresi musik ke dalam format MP3. Inovasi ini memungkinkan musik untuk ditransmisikan melalui kabel optik dan internet, serta menyimpan banyak lagu dalam satu *hard drive*. MP3 juga membuka jalan bagi kemudahan berbagi file musik secara online. Generasi Z menjadi kelompok konsumen terbesar dalam penggunaan berbagai layanan streaming musik, termasuk Spotify. Hingga pertengahan tahun 2024, jumlah pengguna aplikasi Spotify versi *Android* yang telah mengunduhnya tercatat lebih dari satu miliar orang.

Sistem rekomendasi merupakan salah satu cabang dari *machine learning* dengan jenis yang lebih spesifik. Sistem rekomendasi dapat ditingkatkan kualitas pengalaman pengguna dengan memanfaatkan berbagai data dan informasi dari daftar lagu serta riwayat musik pengguna. Salah satu pendekatan yang efektif untuk menganalisis data

tersebut adalah melalui *machine learning*, yaitu metode untuk mengekstraksi informasi dari data yang menggabungkan konsep dari ilmu komputer dan kecerdasan buatan (*artificial intelligence*).

Metode *Random Forest* memiliki dua kegunaan utama dalam menyelesaikan suatu permasalahan, yaitu untuk klasifikasi dan prediksi. Teknik dasar yang mendasari metode ini adalah pohon keputusan. *Random Forest* sendiri merupakan pengembangan dari metode *Classification and Regression Tree* (CART) yang menggabungkan teknik *bagging* (*bootstrap aggregating*) dan pemilihan fitur secara acak. *Bagging* adalah salah satu metode yang dapat meningkatkan kinerja algoritma klasifikasi.

Decision Tree merupakan salah satu metode klasifikasi yang paling umum digunakan. Metode ini memungkinkan pengelompokan suatu item ke dalam model pohon keputusan yang mudah dipahami. Pada dasarnya, *Decision Tree* adalah struktur pohon yang terdiri dari serangkaian atribut yang diuji dengan tujuan memprediksi suatu hasil. Setiap node internal dalam pohon menggambarkan pengujian terhadap sebuah atribut, dengan cabang yang mewakili hasil dari pengujian tersebut, dan setiap node daun menyimpan label kelas. Node yang berada di bagian paling atas pohon disebut sebagai root node atau simpul akar. Penentuan akar pohon dilakukan dengan memilih atribut yang memiliki nilai gain tertinggi atau nilai indeks entropi paling rendah.

Seiring dengan perkembangan teknologi dan semakin banyaknya lagu yang tersedia di *platform streaming* musik, pengguna sering kali kesulitan menemukan lagu yang sesuai dengan selera mereka. Untuk membantu pengguna menemukan lagu yang relevan, diperlukan sistem rekomendasi yang efektif. Metode yang digunakan dalam membangun model rekomendasi juga memiliki peran penting dalam menentukan seberapa baik sistem tersebut bekerja. Oleh karena itu, penelitian ini berfokus pada penerapan metode *Random Forest* dan *Decision Tree* untuk mengembangkan model rekomendasi lagu berbasis genre, dengan tujuan untuk mengetahui metode mana yang lebih efektif dalam memprediksi lagu yang sesuai dengan preferensi

pengguna.

Sebagai solusi dari permasalahan tersebut, penelitian ini mengusulkan penerapan dua metode pembelajaran mesin, yaitu *Random Forest* dan *Decision Tree*, untuk membangun model rekomendasi lagu berbasis fitur audio. Kedua metode ini akan menganalisis berbagai atribut audio dari lagu-lagu yang ada untuk menghasilkan rekomendasi yang lebih tepat sasaran. Dengan membandingkan kedua metode ini, diharapkan dapat ditemukan pendekatan yang paling efektif dalam membangun sistem rekomendasi lagu yang mampu menangkap selera musik pengguna dengan lebih baik.

II. METODE PENELITIAN

Subjek penelitian ini adalah algoritma pembelajaran mesin *Random Forest* dan *Decision Tree*, yang berperan sebagai teknik utama dalam membangun model rekomendasi lagu berbasis fitur audio. Penelitian ini berfokus pada penerapan kedua algoritma tersebut dalam mengolah berbagai atribut dan karakteristik audio dari lagu-lagu yang ada. Objek penelitian ini adalah dataset lagu yang mengandung berbagai fitur audio, seperti *tempo*, *danceability*, *energy*, *loudness*, serta atribut lainnya yang dapat diekstraksi dari rekaman audio. Penelitian ini berfokus pada bagaimana fitur-fitur audio tersebut diproses dan diinterpretasikan oleh kedua algoritma untuk menghasilkan prediksi atau rekomendasi lagu yang tepat. Selain itu, penelitian juga mempertimbangkan kualitas dan variasi data dalam dataset sebagai faktor yang mempengaruhi performa model rekomendasi lagu.

Penelitian ini menggunakan metode penelitian kuantitatif, yang bertujuan untuk menguji dan membandingkan performa metode *Random Forest* dan *Decision Tree* dalam membangun model rekomendasi lagu berbasis fitur audio. Penelitian ini berfokus pada peningkatan akurasi dan kualitas rekomendasi musik berdasarkan fitur-fitur audio.

A. Pengumpulan Data

Pada penelitian ini, data yang digunakan untuk membangun model rekomendasi lagu berbasis genre berasal dari dataset yang tersedia di platform Kaggle. Kaggle merupakan salah satu sumber data terbesar yang menyediakan berbagai dataset untuk keperluan analisis data dan *machine learning*. Dataset yang digunakan dalam penelitian ini terdiri dari sejumlah 32833 data lagu yang mencakup enam genre utama :

Tabel 1 Jumlah Genre dan Lagu

Genre	Jumlah Lagu
EDM	6043
Rap	5746

Pop	5507
R&B	5431
Latin	5155
Rock	4951

B. PRA PEMROSESAN DATA

Data yang telah dikumpulkan dan diekstraksi akan melalui tahap *preprocessing*. Proses ini meliputi pembersihan data (*data cleaning*), normalisasi, dan penanganan data yang hilang (*missing values*). Langkah ini penting untuk memastikan bahwa data siap digunakan oleh model *machine learning* dan menghasilkan performa yang optimal.

C. Balancing Data

Dilakukan *Random Under Sampling* (RUS) yang bertujuan menghilangkan beberapa sampel kelas mayoritas secara acak untuk mengubah distribusi data dalam dataset yang tidak seimbang dan mengubahnya menjadi lebih seimbang.

D. Feature Importance

Pentingnya fitur (*feature importance*) merupakan salah satu metode dalam tahap persiapan data (*data preparation*), khususnya termasuk dalam kategori reduksi data (*data reduction*). Pemilihan fitur bertujuan untuk menemukan nilai relevansi suatu fitur terhadap label kelas dan mengabaikan fitur yang tidak memberikan kontribusi apa pun terhadap hasil klasifikasi data

E. Eksplorasi Data

Analisis data eksploratori (EDA) untuk memeriksa data untuk mengetahui distribusi guna mengarahkan pengujian khusus terhadap hipotesis. EDA bertujuan untuk membantu pengenalan pola alami analisis. Terakhir, teknik pemilihan fitur sering kali termasuk dalam EDA.

F. Evaluasi Model

Evaluasi dilakukan dengan menggunakan metrik seperti akurasi, *precision*, *recall*, dan *F1-score* untuk mengukur seberapa baik model dalam merekomendasikan lagu. *Cross-validation* adalah penggunaan kembali data, dan kemudian data sampel yang diperoleh dibagi menjadi beberapa *training set* dan *testing set*. Model dilatih dengan *training set*, dan kualitas prakiraan model dievaluasi oleh *testing set*. *Cross-validation* sangat penting untuk membangun model pembelajaran mesin dengan membandingkan dan memilih model yang sesuai dan mengevaluasi kinerja. Dalam pembelajaran mesin, validasi silang telah diterapkan secara luas untuk mengevaluasi kinerja model. *Cross-Fold Validation* akan diukur dengan menggunakan metrik kinerja akurasi, presisi, recall, *F1-score*.

- *Accuracy*

Proporsi sampel yang diidentifikasi dengan tepat oleh model terhadap jumlah total sampel dikenal sebagai akurasi prediksi model.

$$\text{Accuracy} = \frac{TP+TN}{TP+TN+FP+FN}$$

- *Presisi*

Keakuratan model diukur berdasarkan rasio nilai positif yang berhasil dikategorikan terhadap semua sampel positif yang diantisipasi.

$$\text{Presisi} = \frac{TP}{TP+FP}$$

- *Recall*

Penarikan kembali suatu model didefinisikan sebagai proporsi sampel positif yang diprediksi dengan benar terhadap semua sampel positif.

$$\text{Recall} = \frac{TP}{TP+FN}$$

- *F1-Score*

F-Score dari model menentukan rata-rata dari presisi dan recall.

$$\text{F1-Score} = \frac{\text{Presisi} \cdot \text{Recall}}{\text{Presisi} + \text{Recall}}$$

III. HASIL DAN PEMBAHASAN

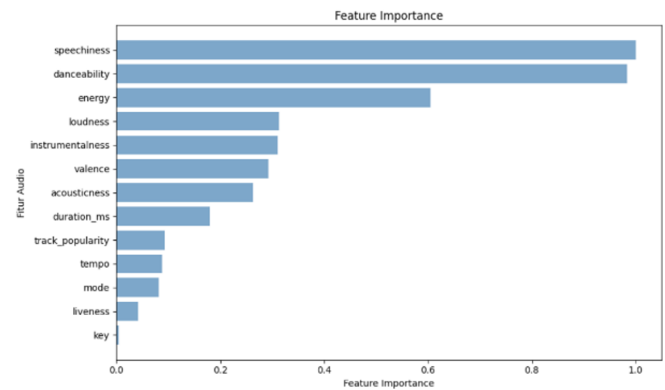
Penelitian ini menunjuk kepada pemahaman tentang fitur audio apa saja yang mempengaruhi penentuan genre. Klasifikasi *genre playlist* lagu menggunakan model *Random Forest* dan *Decision Tree*. Dataset dibersihkan dan distandarisasi, serta hanya fitur-fitur numerik terpenting yang digunakan.

A. *Feature importance*

Pemilihan Fitur adalah salah satu teknik *Data Preparation* terutama dalam lingkup *data reduction*. Pemilihan fitur bertujuan untuk menemukan nilai relevansi suatu fitur terhadap label kelas dan mengabaikan fitur yang tidak memberikan kontribusi apa pun terhadap hasil klasifikasi data. Untuk menemukan fitur yang lebih berhubungan dengan *preferensi music* setiap generasi penelitian ini mengadopsi metode umum untuk pemilihan fitur, uji F. Uji F merupakan Teknik penting untuk analisis *varians* (ANOVA). Menghitung nilai F ANOVA untuk memperkirakan Tingkat ketergantungan linear antar variable,

dalam penelitian ini masing-masing fitur. Rumus untuk menghitung uji F sebagai berikut :

$$F = \frac{MSB}{MSW}$$

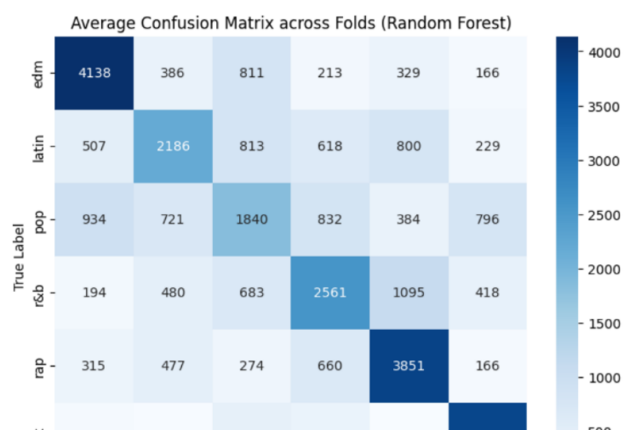


Gambar 1. *Feature Importance*

Gambar 1 menunjukkan grafik *Feature Importance* dari berbagai fitur audio dalam suatu model klasifikasi, kemungkinan untuk genre lagu. Terlihat bahwa *speechiness*, *danceability*, dan *energy* merupakan tiga fitur yang paling berpengaruh dalam model, dengan nilai penting mendekati 1, menandakan kontribusi dominan terhadap prediksi. Fitur seperti *loudness*, *instrumentalness*, dan *valence* juga memiliki pengaruh sedang, sementara *track_popularity*, *tempo*, *mode*, *liveness*, dan terutama *key* memiliki pengaruh yang sangat kecil terhadap hasil model. Ini menunjukkan bahwa karakteristik vokal dan irama lebih penting dibandingkan atribut teknis seperti kunci atau tempo dalam menentukan genre atau klasifikasi lagu dalam model ini.

B. *Random Forest* dan *Decision Tree* tanpa *Undersampling*

Penerapan model *Random Forest* untuk mengklasifikasikan genre lagu berdasarkan fitur audio, menggunakan teknik validasi silang *Stratified K-Fold* sebanyak 10 lipatan. Dataset dibersihkan (tanpa di *balancing*), fitur numerik dipilih dan distandarisasi, serta label target diencode.



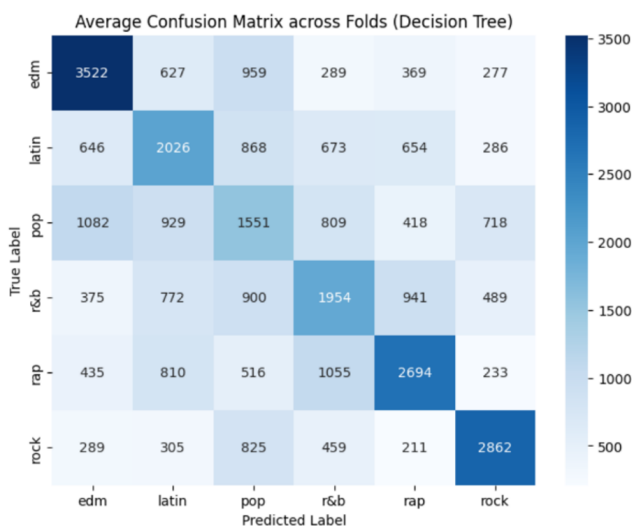
Gambar 3. *Confusion Matrix Decision Tree Tanpa Undersampling*

Gambar 3 menunjukkan *confusion matrix* rata-rata model *Decision Tree* dalam mengklasifikasikan enam genre lagu. Model cukup baik mengenali genre seperti *edm*, *rap*, dan *rock*, yang terlihat dari nilai diagonal yang tinggi. Namun, masih banyak terjadi kesalahan klasifikasi pada genre seperti *pop*, *latin*, dan *r&b*, yang sering tertukar satu sama lain. Hal ini menunjukkan bahwa model masih kesulitan membedakan genre dengan karakteristik audio yang mirip, sehingga perlu perbaikan, misalnya melalui pemilihan model yang lebih kompleks atau fitur yang lebih representatif.

Gambar 2. *Confusion Matrix Random Forest Tanpa Undersampling*

Gambar 2 menunjukkan *confusion matrix* dari model *Random Forest* dalam mengklasifikasikan genre lagu. Nilai diagonal menunjukkan jumlah prediksi yang benar untuk setiap genre, dengan hasil tertinggi pada genre *edm*, *rap*, dan *rock*, yang berarti model cukup akurat mengenali ketiganya. Sebaliknya, genre seperti *pop*, *latin*, dan *r&b* lebih sering salah diklasifikasikan ke genre lain, menunjukkan model mengalami kesulitan membedakan genre-genre tersebut. Visualisasi ini membantu mengevaluasi kekuatan dan kelemahan model dalam mengenali pola audio tiap genre.

Klasifikasi genre *playlist* lagu menggunakan model *Decision Tree*. Dataset dibersihkan dan distandarisasi, serta hanya fitur-fitur numerik terpenting yang digunakan. Target (genre) dikodekan secara numerik menggunakan *LabelEncoder*.



EVALUASI		Random Forest	Decision Tree
F1	Accuracy	0.5669	0.4372
	Presisi	0.5591	0.4362
	Recall	0.5669	0.4372
	F1-Score	0.5616	0.4360
F2	Accuracy	0.5528	0.4233
	Presisi	0.5465	0.4236
	Recall	0.5528	0.4233
	F1-Score	0.5486	0.4234
F3	Accuracy	0.5614	0.4261
	Presisi	0.5501	0.4256
	Recall	0.5614	0.4261
	F1-Score	0.5534	0.4253
F4	Accuracy	0.5556	0.4411
	Presisi	0.5458	0.4383
	Recall	0.5556	0.4411
	F1-Score	0.5488	0.4393
F5	Accuracy	0.5653	0.4339
	Presisi	0.5676	0.4335
	Recall	0.5653	0.4339
	F1-Score	0.5599	0.4332
F6	Accuracy	0.5583	0.4206
	Presisi	0.5587	0.4263
	Recall	0.5583	0.4206
	F1-Score	0.5508	0.4228
F7	Accuracy	0.5592	0.4339
	Presisi	0.5537	0.4302
	Recall	0.5592	0.4339
	F1-Score	0.5554	0.4316
F8	Accuracy	0.5532	0.4156

	<i>Presiasi</i>	0.5416	0.4150
	<i>Recall</i>	0.5532	0.4156
	<i>F1-Score</i>	0.5457	0.4151
F9	<i>Accuracy</i>	0.5628	0.4189
	<i>Presiasi</i>	0.5521	0.4215
	<i>Recall</i>	0.5628	0.4189
	<i>F1-Score</i>	0.5554	0.4198
F10	<i>Accuracy</i>	0.5603	0.4278
	<i>Presiasi</i>	0.5508	0.4302
	<i>Recall</i>	0.5603	0.4278
	<i>F1-Score</i>	0.5537	0.4285
Hasil Rata-Rata	<i>Accuracy</i>	0.5596	0.4278
	<i>Presiasi</i>	0.5506	0.4280
	<i>Recall</i>	0.5596	0.4278
	<i>F1-Score</i>	0.5533	0.4275

Tabel 2. Evaluasi terbaik dari *Random Forest* dan *Decision Tree* tanpa *Undersampling*

Tabel 2 menunjukkan perbandingan kinerja model *Random Forest* dan *Decision Tree* tanpa *Undersampling* dalam 10 lipatan (F1–F10) menggunakan metrik *Accuracy*, *Precision*, *Recall*, dan *F1-Score*. Secara konsisten, *Random Forest* menghasilkan skor yang lebih tinggi di seluruh metrik dibandingkan dengan *Decision Tree*, menunjukkan bahwa model ini lebih mampu menangkap pola dalam data. Rata-rata akurasi *Random Forest* mencapai 0.5603, sedangkan *Decision Tree* hanya 0.4278. Performa yang lebih baik ini juga tercermin pada metrik lainnya seperti *Precision* (0.5506 vs 0.4280) dan *F1-Score* (0.5533 vs 0.4275), mengindikasikan bahwa *Random Forest* lebih andal dalam mengklasifikasikan genre lagu secara keseluruhan.

EVALUASI		Random Forest	Decision Tree
F1	<i>Accuracy</i>	0.5571	0.4468
	<i>Presiasi</i>	0.5518	0.4499
	<i>Recall</i>	0.5571	0.4468
	<i>F1-Score</i>	0.5537	0.4476
F2	<i>Accuracy</i>	0.5449	0.4453
	<i>Presiasi</i>	0.5387	0.4475
	<i>Recall</i>	0.5449	0.4453
	<i>F1-Score</i>	0.5407	0.4456
F3	<i>Accuracy</i>	0.5568	0.4587
	<i>Presiasi</i>	0.5478	0.4601

	<i>Recall</i>	0.5568	0.4587
	<i>F1-Score</i>	0.5507	0.4592
	<i>Accuracy</i>	0.5532	0.4423
F4	<i>Presiasi</i>	0.5462	0.4448
	<i>Recall</i>	0.5532	0.4423
	<i>F1-Score</i>	0.5479	0.4430
	<i>Accuracy</i>	0.5589	0.4456
F5	<i>Presiasi</i>	0.5508	0.4483
	<i>Recall</i>	0.5589	0.4456
	<i>F1-Score</i>	0.5533	0.4466
	<i>Accuracy</i>	0.5504	0.4462
F6	<i>Presiasi</i>	0.5423	0.4436
	<i>Recall</i>	0.5504	0.4462
	<i>F1-Score</i>	0.5442	0.4446
	<i>Accuracy</i>	0.5489	0.4310
F7	<i>Presiasi</i>	0.5441	0.4346
	<i>Recall</i>	0.5489	0.4310
	<i>F1-Score</i>	0.5452	0.4318
	<i>Accuracy</i>	0.5553	0.4435
F8	<i>Presiasi</i>	0.5458	0.4428
	<i>Recall</i>	0.5553	0.4435
	<i>F1-Score</i>	0.5488	0.4428
	<i>Accuracy</i>	0.5628	0.4573
F9	<i>Presiasi</i>	0.5541	0.4571
	<i>Recall</i>	0.5628	0.4573
	<i>F1-Score</i>	0.5570	0.4571
	<i>Accuracy</i>	0.5585	0.4476
F10	<i>Presiasi</i>	0.5505	0.4474
	<i>Recall</i>	0.5585	0.4476
	<i>F1-Score</i>	0.5530	0.4471
	<i>Accuracy</i>	0.5547	0.4464
Hasil Rata-Rata	<i>Presiasi</i>	0.5472	0.4476
	<i>Recall</i>	0.5547	0.4464
	<i>F1-Score</i>	0.5495	0.4465

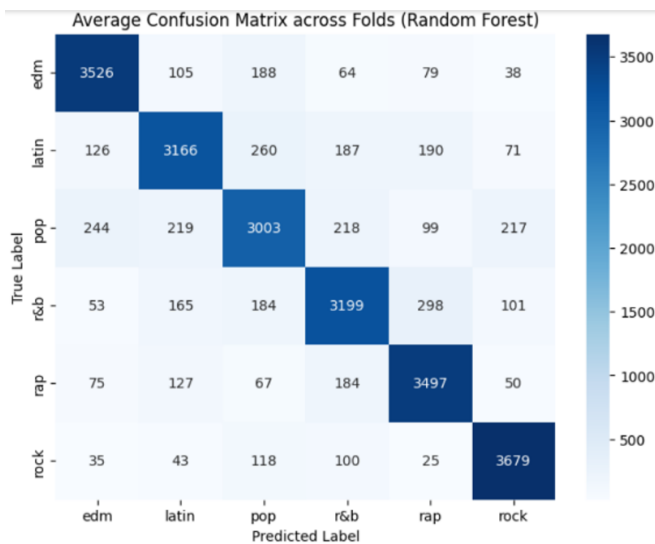
Tabel 3. Evaluasi terburuk dari *Random Forest* dan *Decision Tree* tanpa *Undersampling*

Tabel 3 memperlihatkan hasil evaluasi paling buruk dari model *Random Forest* dan *Decision Tree* tanpa *Undersampling* pada 10 fold menggunakan empat metrik utama: *Accuracy*, *Precision*, *Recall*, dan *F1-Score*. Dari data tersebut, terlihat bahwa *Random Forest* secara konsisten mengungguli *Decision Tree* pada setiap lipatan dan seluruh

metrik evaluasi. Misalnya, rata-rata *Accuracy Random Forest* berada di angka 0.5545, sedangkan *Decision Tree* hanya 0.4468. Begitu juga dengan *Precision* (0.5472 vs 0.4462), *Recall* (0.5547 vs 0.4464), dan *F1-Score* (0.5495 vs 0.4465), yang menunjukkan bahwa *Random Forest* lebih efektif dalam menangani klasifikasi genre musik dibandingkan *Decision Tree*. Secara keseluruhan, hasil ini mengindikasikan bahwa pendekatan ensemble seperti *Random Forest* lebih mampu menangkap kompleksitas data dibandingkan pohon keputusan tunggal.

C. *Random Forest* dan *Decision Tree* dengan *Undersampling*

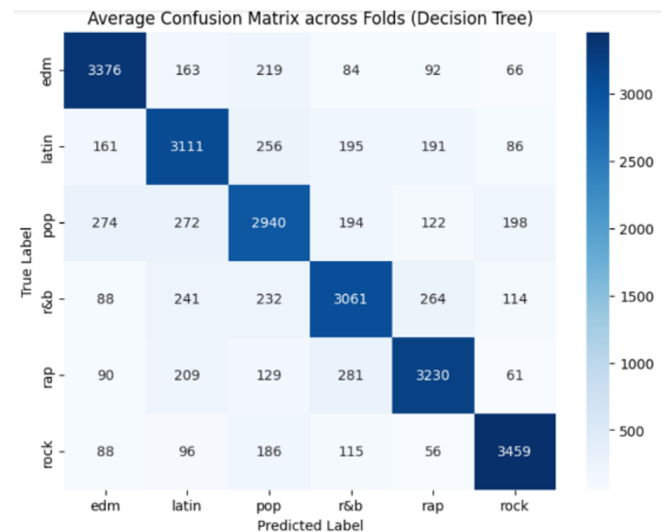
Model *Random Forest* dikembangkan dan diuji untuk set data. Metrik kinerja adalah akurasi, presisi, *recall*, dan skor F1 untuk menganalisis efektivitas model.



Gambar 4. *Confusion Matrix Random Forest* Dengan *Undersampling*

Gambar 4 adalah implementasi sistem rekomendasi genre musik menggunakan algoritma *Random Forest* yang dievaluasi dengan *Stratified K-Fold Cross Validation* sebanyak 10 lipatan. Dataset yang digunakan telah seimbang dan mencakup fitur-fitur audio penting seperti *speechiness*, *valence*, *tempo*, dll. Setiap fitur dinormalisasi menggunakan *StandardScaler*, dan label genre dikodekan menjadi angka dengan *LabelEncoder*. Model dilatih dan diuji pada setiap lipatan untuk menghitung metrik evaluasi seperti *accuracy*, *precision*, *recall*, dan *F1-score*. Hasil dari setiap lipatan dirata-rata untuk menilai performa keseluruhan model. Akhirnya, ditampilkan visualisasi *confusion matrix* rata-rata

untuk melihat seberapa akurat model dalam mengklasifikasikan genre musik.



Gambar 5. *Confusion Matrix Decision Tree* Dengan *Undersampling*

Gambar 5 adalah implementasi model rekomendasi genre lagu menggunakan *Decision Tree* yang dievaluasi dengan *Stratified K-Fold Cross-Validation* sebanyak 10 lipatan. Data diformat ulang dengan standarisasi dan target label diubah menjadi angka. Pada setiap lipatan, model *Decision Tree* dilatih dan diuji, lalu dihitung metrik evaluasi seperti akurasi, presisi, *recall*, dan *F1-score*. Setelah semua lipatan selesai, hasil dari setiap metrik dirata-ratakan untuk menilai performa umum model. Terakhir, *confusion matrix* gabungan divisualisasikan untuk menunjukkan seberapa baik model mengenali genre.

EVALUASI		<i>Random Forest</i>	<i>Decision Tree</i>
F1	<i>Accuracy</i>	0.8246	0.7846
	<i>Presisi</i>	0.8236	0.7851
	<i>Recall</i>	0.8246	0.7846
	<i>F1-Score</i>	0.8235	0.7847
F2	<i>Accuracy</i>	0.8337	0.8004
	<i>Presisi</i>	0.8335	0.8008
	<i>Recall</i>	0.8337	0.8004
	<i>F1-Score</i>	0.8331	0.8004
F3	<i>Accuracy</i>	0.8363	0.7917
	<i>Presisi</i>	0.8356	0.7917
	<i>Recall</i>	0.8363	0.7917
	<i>F1-Score</i>	0.8357	0.7915

F4	Accuracy	0.8475	0.8046
	Presisi	0.8467	0.8045
	Recall	0.8475	0.8046
	F1-Score	0.8466	0.8043
F5	Accuracy	0.8367	0.8063
	Presisi	0.8358	0.8058
	Recall	0.8367	0.8063
	F1-Score	0.8357	0.8058
F6	Accuracy	0.8450	0.8075
	Presisi	0.8448	0.8077
	Recall	0.8450	0.8075
	F1-Score	0.8448	0.8074
F7	Accuracy	0.8204	0.7900
	Presisi	0.8195	0.7897
	Recall	0.8204	0.7900
	F1-Score	0.8190	0.7898
F8	Accuracy	0.8333	0.7996
	Presisi	0.8337	0.8009
	Recall	0.8333	0.7996
	F1-Score	0.8335	0.8000
F9	Accuracy	0.8375	0.7996
	Presisi	0.8369	0.8003
	Recall	0.8375	0.7996
	F1-Score	0.8366	0.7998
F10	Accuracy	0.8475	0.8063
	Presisi	0.8470	0.8059
	Recall	0.8475	0.8063
	F1-Score	0.8467	0.8060
Hasil Rata-Rata	Accuracy	0.8363	0.7990
	Presisi	0.8357	0.7992
	Recall	0.8363	0.7990
	F1-Score	0.8355	0.7990

Tabel 4. Evaluasi terbaik dari *Random Forest* dan *Decision Tree* dengan *Undersampling*

Berdasarkan hasil perbandingan pada tabel 4 evaluasi 10-fold cross-validation antara model *Random Forest* dan *Decision Tree*, terlihat bahwa *Random Forest* secara konsisten memberikan performa yang lebih baik dalam semua metrik evaluasi : *Accuracy*, *Precision*, *Recall*, dan *F1-Score*. Rata-rata akurasi *Random Forest* sebesar 0.8363, sedangkan *Decision Tree* hanya 0.7900. Hal yang sama berlaku untuk presisi (0.8357 vs 0.7992), recall (0.8363 vs 0.7990), dan *F1-*

score (0.8355 vs 0.7990). Selisih ini menunjukkan bahwa *Random Forest* mampu melakukan generalisasi yang lebih baik terhadap data dibandingkan *Decision Tree*, karena sifat ensemble-nya yang menggabungkan banyak pohon keputusan dan mengurangi risiko *overfitting*. Dengan demikian, *Random Forest* merupakan pilihan model yang lebih unggul untuk sistem rekomendasi genre lagu dalam penelitian ini.

EVALUASI		<i>Random Forest</i>	<i>Decision Tree</i>
F1	Accuracy	0.8150	0.7754
	Presisi	0.8141	0.7774
	Recall	0.8150	0.7754
	F1-Score	0.8143	0.7761
F2	Accuracy	0.8367	0.7933
	Presisi	0.8361	0.7949
	Recall	0.8367	0.7933
	F1-Score	0.8360	0.7936
F3	Accuracy	0.8254	0.7883
	Presisi	0.8245	0.7879
	Recall	0.8254	0.7883
	F1-Score	0.8246	0.7880
F4	Accuracy	0.8408	0.8037
	Presisi	0.8395	0.8040
	Recall	0.8408	0.8037
	F1-Score	0.8396	0.8035
F5	Accuracy	0.8313	0.7983
	Presisi	0.8307	0.7979
	Recall	0.8313	0.7983
	F1-Score	0.8305	0.7978
F6	Accuracy	0.8371	0.8058
	Presisi	0.8369	0.8064
	Recall	0.8371	0.8058
	F1-Score	0.8369	0.8059
F7	Accuracy	0.8192	0.7992
	Presisi	0.8181	0.7989
	Recall	0.8192	0.7992
	F1-Score	0.8177	0.7990
F8	Accuracy	0.8283	0.7992
	Presisi	0.8288	0.7997
	Recall	0.8283	0.7992
	F1-Score	0.8285	0.7993
F9	Accuracy	0.8296	0.7913
	Presisi	0.8289	0.7916

	Recall	0.8296	0.7913
	F1-Score	0.8288	0.7911
F10	Accuracy	0.8375	0.8104
	Presisi	0.8373	0.8104
	Recall	0.8375	0.8104
	F1-Score	0.8371	0.8103
Hasil Rata-Rata	Accuracy	0.8301	0.7965
	Presisi	0.8295	0.7969
	Recall	0.8301	0.7965
	F1-Score	0.8294	0.7965

Tabel 5. Evaluasi terburuk dari *Random Forest* dan *Decision Tree* dengan *Undersampling*

Pada tabel 5 merupakan hasil perbandingan evaluasi 10-fold cross-validation hasil terburuk dari seluruh percobaan, model *Random Forest* tetap menunjukkan performa yang lebih baik dibandingkan *Decision Tree* dalam setiap metrik evaluasi utama. Meskipun performanya menurun, *Random Forest* mencatat rata-rata akurasi sebesar 0.8301, sementara *Decision Tree* hanya 0.7965. Begitu pula dengan rata-rata presisi (0.8295 vs 0.7965), recall (0.8301 vs 0.7965), dan F1-score (0.8294 vs 0.7965). Penurunan kinerja ini mungkin disebabkan oleh variasi data pelatihan yang mempengaruhi hasil tiap fold, namun *Random Forest* masih unggul secara keseluruhan karena kemampuannya dalam menangani kompleksitas data dengan lebih stabil. Kesimpulannya, bahkan dalam kondisi performa terendah, *Random Forest* tetap lebih andal dibandingkan *Decision Tree* untuk tugas klasifikasi genre lagu.

IV. KESIMPULAN

Penelitian ini berhasil membangun model rekomendasi lagu berbasis genre dengan memanfaatkan algoritma *Random Forest* dan *Decision Tree*. Berdasarkan analisis *feature importance*, diperoleh 10 fitur utama yang paling mempengaruhi klasifikasi genre lagu. Hasil evaluasi menunjukkan bahwa model *Random Forest* memiliki performa yang lebih baik dibandingkan *Decision Tree*, dengan akurasi yang lebih tinggi dalam mengklasifikasikan genre lagu. Secara kuantitatif, model *Random Forest* mencatat rata-rata akurasi sebesar 83%, sementara *Decision Tree* hanya mencapai 79%. Selain itu, evaluasi model dilakukan menggunakan teknik *cross-validation* untuk memastikan bahwa hasil yang diperoleh tidak hanya akurat pada satu subset data, tetapi juga konsisten dan dapat digeneralisasi ke data lain. Teknik ini membantu mengurangi risiko *overfitting* dan meningkatkan keandalan performa model secara

keseluruhan. Oleh karena itu, *Random Forest* dapat dianggap sebagai metode yang lebih efektif untuk digunakan dalam sistem rekomendasi lagu berbasis genre.

V. REFERENSI

- [1] Ruddin, D. R. E. Indrajit, dan H. Santoso membahas dampak teknologi informasi terhadap perkembangan industri musik di Indonesia dalam Jurnal Pendidikan Sains dan Komputer, Vol. 2, No. 1, hlm. 124–136, tahun 2022.
- [2] R. C. Maringka dan tim melakukan kajian perkembangan musik melalui platform Spotify dengan memanfaatkan Structured Query Language (SQL), diterbitkan pada tahun 2021.
- [3] M. V. Anggoro dan M. Izzatillah merancang sistem rekomendasi musik berbasis Android menggunakan metode collaborative filtering, yang diterbitkan di jurnal String, Vol. 7, No. 1, hlm. 1, tahun 2022, DOI: 10.30998/String.V7i1.10300.
- [4] M. A. Madani bersama rekan-rekannya menerapkan teknik machine learning dalam sistem rekomendasi musik, dipublikasikan di Vol. 1, No. 1, hlm. 40–49, tahun 2024.
- [5] S. Mahmuda menerapkan algoritma Random Forest untuk mengklasifikasikan konten pada kanal YouTube, dimuat dalam Jurnal Jendela Matematika, Vol. 2, No. 01, hlm. 21–31, tahun 2024, DOI: 10.57008/Jjm.V2i01.633.
- [6] R. H. Rachmat Hidayat membandingkan metode Naïve Bayes dan Decision Tree C4.5 untuk analisis sentimen pada produk Es Teh Indonesia di Twitter, diterbitkan dalam Jurnal Siskom-KB, Vol. 7, No. 2, hlm. 88–98, tahun 2024, DOI: 10.47970/Siskom-Kb.V7i2.607.
- [7] Y. Cahyaningtyas, Y. Nataliani, dan I. R. Widiarsari meneliti sentimen pengguna terhadap rating aplikasi Shopee dengan metode Decision Tree berbasis SMOTE, diterbitkan di Aiti, Vol. 18, No. 2, hlm. 173–184, tahun 2021, DOI: 10.24246/Aiti.V18i2.173-184.
- [8] S. Abdulmalikov dan kolaborator mengevaluasi performa machine learning dalam klasifikasi emosi berbasis EEG dengan beragam metode seleksi fitur, dimuat di *Applied Sciences (Switzerland)*, tahun 2024, DOI: 10.3390/App142210511.
- [9] W. Chen dan rekan-rekan menyusun tinjauan komprehensif mengenai pembelajaran dengan data tidak seimbang serta aplikasinya di berbagai bidang, dalam *Artificial Intelligence Review*, Vol. 57, No. 6, tahun 2024, DOI: 10.1007/S10462-024-10759-6.
- [10] J. Arvidsson menyediakan dataset berisi 30.000 lagu dari Spotify yang tersedia secara daring di Kaggle, diakses pada Juli 2025 melalui: <https://www.kaggle.com/datasets/joebeachcapital/30000-spotify-songs>.

- [11] S. Farhani dan A. Qoiriah melakukan klasterisasi data musik menggunakan algoritma K-Means, diterbitkan di Vol. 6, hlm. 566–572, tahun 2024.
- [12] A. Elsoud dan rekan-rekannya membandingkan teknik undersampling dalam penanganan data tidak seimbang dengan berbagai tingkat ketimpangan, dipublikasikan dalam *International Journal of Advanced Computer Science and Applications*, Vol. 15, No. 8, tahun 2024, DOI: 10.14569/Ijacs.2024.01508124.
- [13] K. Al-Tamimi, M. Salem, dan A. Al-Alami meneliti pemanfaatan seleksi fitur dalam klasifikasi genre musik, dipresentasikan dalam prosiding ITT 2020, Vol. 9, No. 4, DOI: 10.1109/Itt51279.2020.9320778.
- [14] Tim dari M. Computing mengulas preferensi musik daring antar generasi dari perspektif akustik dalam publikasi bulan Agustus tahun 2020.
- [15] J. Qiu menganalisis teknik evaluasi model dengan cross-validation, termasuk penerapannya dan kemajuan terbaru, dipresentasikan dalam prosiding ICFTBA 2024, hlm. 69–72, DOI: 10.54254/2754-1169/99/2024ox0213.
- [16] J. Kaliappan dan tim menyoroti peran cross-validation dalam model machine learning untuk deteksi dini kematian janin dalam kandungan, dimuat di *Diagnostics*, Vol. 13, No. 10, tahun 2023, DOI: 10.3390/Diagnostics13101692.