

**TESIS**  
**REDUKSI DIMENSI DAN KLASIFIKASI MUSIC UNTUK**  
**BELAJAR MENGGUNAKAN FEATURE IMPORTANCE DAN**  
**RANDOM FOREST**



disusun oleh

**LAKSMITA DEWI SUPRABA**

**24.55.1636**

**Teknologi Media Digital**

**FAKULTAS ILMU KOMPUTER**  
**UNIVERSITAS AMIKOM YOGYAKARTA**  
**YOGYAKARTA**

**2026**

**TESIS**  
**EDUKSI DIMENSI DAN KLASIFIKASI MUSIK UNTUK**  
**BELAJAR MENGGUNAKAN FEATURE IMPORTANCE DAN**  
**RANDOM FOREST**

**DIMENSIONALITY REDUCTION AND MUSIC**  
**CLASSIFICATION FOR LEARNING USING FEATURE**  
**IMPORTANCE AND RANDOM FOREST**



disusun oleh

**LAKSMITA DEWI SUPRABA**

**24.55.1636**

**Teknologi Media Digital**

**FAKULTAS ILMU KOMPUTER**

**UNIVERSITAS AMIKOM YOGYAKARTA**

**YOGYAKARTA**

**2026**

## HALAMAN PERSETUJUAN

**REDUKSI DIMENSI DAN KLASIFIKASI MUSIK UNTUK BELAJAR  
MENGUNAKAN FEATURE IMPORTANCE DAN RANDOM FOREST**

**DIMENSIONALITY REDUCTION AND MUSIC CLASSIFICATION FOR  
LEARNING USING FEATURE IMPORTANCE AND RANDOM FOREST**

yang disusun dan diajukan oleh

**Laksmita Dewi Supraba**

**24.55.1636**

telah disetujui oleh Dosen Pembimbing Tesis  
pada tanggal <tanggal ujian>

**Dosen Pembimbing,**



**Dr. Andri Sunyoto, M.Kom**

**NIK. 190302052**

**HALAMAN PENGESAHAN**

**REDUKSI DIMENSI DAN KLASIFIKASI MUSIK UNTUK BELAJAR  
MENGUNAKAN FEATURE IMPORTANCE DAN RANDOM FOREST**

**DIMENSIONALITY REDUCTION AND MUSIC CLASSIFICATION FOR  
LEARNING USING FEATURE IMPORTANCE AND RANDOM FOREST**

yang disusun dan diajukan oleh

**Laksmi Dewi Supraba**

**24.55.1636**

Telah dipertahankan di depan Dewan Penguji  
pada tanggal 8 Januari 2026

**Susunan Dewan Penguji**

**Nama Penguji**

**Tanda Tangan**

**Prof.Dr.Ema Utami, S.Si, M.Kom.**  
**NIK. 190302037**

**Hanif Al Fatta, S.Kom, M.Kom., Ph.D.**  
**NIK. 190302096**



Tesis ini telah diterima sebagai salah satu persyaratan  
untuk memperoleh gelar Magister Komputer  
Tanggal 8 Januari 2026

**DEKAN FAKULTAS ILMU KOMPUTER**



**Prof. Dr. Kusriani, M.Kom.**  
**NIK. 190302106**

## HALAMAN PERNYATAAN KEASLIAN TESIS

Yang bertandatangan di bawah ini,

Nama mahasiswa : Laksmi Dewi Supraba  
NIM : 24.55.1636

Menyatakan bahwa Tesis dengan judul berikut:

### **REDUKSI DIMENSI DAN KLASIFIKASI MUSIK UNTUK BELAJAR MENGUNAKAN FEATURE IMPORTANCE DAN RANDOM FOREST**

Dosen Pembimbing : Dr. Andi Sunyoto, M. Kom.

1. Karya tulis ini adalah benar-benar ASLI dan BELUM PERNAH diajukan untuk mendapatkan gelar akademik, baik di Universitas AMIKOM Yogyakarta maupun di Perguruan Tinggi lainnya.
2. Karya tulis ini merupakan gagasan, rumusan dan penelitian SAYA sendiri, tanpa bantuan pihak lain kecuali arahan dari Dosen Pembimbing.
3. Dalam karya tulis ini tidak terdapat karya atau pendapat orang lain, kecuali secara tertulis dengan jelas dicantumkan sebagai acuan dalam naskah dengan disebutkan nama pengarang dan disebutkan dalam Daftar Pustaka pada karya tulis ini.
4. Perangkat lunak yang digunakan dalam penelitian ini sepenuhnya menjadi tanggung jawab SAYA, bukan tanggung jawab Universitas AMIKOM Yogyakarta.
5. Pernyataan ini SAYA buat dengan sesungguhnya, apabila di kemudian hari terdapat penyimpangan dan ketidakbenaran dalam pernyataan ini, maka SAYA bersedia menerima SANKSI AKADEMIK dengan pencabutan gelar yang sudah diperoleh, serta sanksi lainnya sesuai dengan norma yang berlaku di Perguruan Tinggi.

Yogyakarta, 8 Januari 2026

Yang Menyatakan,



Laksmi Dewi Supraba

## HALAMAN PERSEMBAHAN

Segala puji dan syukur ke hadirat Allah SWT atas limpahan rahmat, hidayah, dan keberkahan-Nya sehingga penulis dapat menyelesaikan tesis berjudul *“Reduksi Dimensi Dan Klasifikasi Musik Untuk Belajar Menggunakan Feature Importance Dan Random Forest”* Dengan penuh rasa syukur, karya ini penulis persembahkan kepada Mama dan Bapak tercinta atas doa dan kasih sayang tanpa batas serta Abah KH. Abd Nashir atas dukungan moril, kepada dr. Ismi yang selalu memberi motivasi dan energi positif, kepada diri sendiri Laksmi Dewi Supraba yang terus berjuang dan bangkit dalam setiap proses, kepada Rui atas semangat dan dukungan sehingga penulis dapat menyelesaikan studi Program Magister Teknik Informatika (PJJ-MTI) Universitas Amikom Yogyakarta pada tahun 2026 ini.

## KATA PENGANTAR

Segala puji dan syukur kehadiran Allah SWT atas berkah, rahmat dan hidayah Nya yang senantiasa dilimpahkan kepada penulis, sehingga bisa merampungkan Tesis yang berjudul "*REDUKSI DIMENSI DAN KLASIFIKASI MUSIK UNTUK BELAJAR MENGGUNAKAN FEATURE IMPORTANCE DAN RANDOM FOREST*" ini dengan baik, sebagai syarat untuk menyelesaikan Program Magister (S2) pada Program Studi Pendidikan Jarak Jauh Magister Teknik Informatika (PJJ-MTI) Universitas Amikom Yogyakarta.

Dalam penyusunan Tesis ini banyak hambatan serta rintangan yang penulis hadapi namun pada akhirnya dapat dilalui berkat adanya dukungan dan bantuan dari berbagai pihak baik secara moral maupun spiritual. Untuk itu pada kesempatan ini penulis menyampaikan ucapan terimakasih kepada:

1. Bapak Prof. Dr. M. Suyanto, M.M., selaku Rektor Universitas Amikom Yogyakarta.
2. Ibu Prof. Dr. Kusriani, M. Kom., selaku Direktur Program Pasca Sarjana Universitas Amikom Yogyakarta yang telah memberikan kesempatan dan izin untuk menempuh studi lanjut di Program Studi Pendidikan Jarak Jauh Magister Teknik Informatika (PJJ-MTI).
3. Bapak Dr. Andi Sunyoto, M. Kom., selaku pembimbing utama. Atas bimbingan dan arahan yang menjaga penulis tetap fokus dan berproses dengan baik penulis demi kesempurnaan penulisan penelitian ini.
4. Tim Penguji dari SPT, SHPT, hingga UT yang telah memberikan arahan dan wawasan lebih dalam proses penyempurnaan penulisan.
5. Seluruh Dosen Pengajar di Program Studi Pendidikan Jarak Jauh Magister Teknik Informatika (PJJ-MTI) Universitas Amikom Yogyakarta dari semester pertama hingga terakhir yang memberikan arahan, dukungan, semangat, dan sharing pengetahuan sehingga penulis mendapatkan wawasan baru yang lebih luas dalam menyelesaikan tugas disetiap studi.

6. Segenap Civitas Akademika (Pengelola dan Admisi) Program Studi Pendidikan Jarak Jauh Magister Teknik Informatika (PJJ-MTI) Universitas Amikom Yogyakarta yang telah memberikan pelayanan dan bantuan sangat baik dalam kebutuhan studi.
7. Bapak Supradoto dan Mama Kurniawati SP. d, S.E. terimakasih atas segala doa dan dukungan. Penulis percaya dalam setiap kemudahan dan kesuksesan penulis bukanlah karena jerih payah penulis, melainkan berkat secercah do'a kedua orang tua yang menembus langit.
8. Abah KH. Abdul Nashir yang memberikan nasehat moril selama penulisan tesis.
9. Pakpuh Prof. Dr. Sugiono, M.M atas support, semangat dan arahan yang diberikan.
10. Teman terdekat dan rekan penulis yang memberi support dr. Ismi Nur Aini Latifah. Terimakasih atas kesediannya berbagi keluh kesah, sedih senang, tawa duka. Penulis tidak akan sampai pada titik ini tanpa dukungan dan dampingan kalian selama ini.
11. Teman-teman Mahasiswa Program Studi Pendidikan Jarak Jauh Magister Teknik Informatika (PJJ-MTI) Universitas Amikom Yogyakarta Angkatan 2024 yang telah memberikan pengalaman, suasana dan keluarga baru.

Dengan seneng hati penulis menerima kritik dan saran yang membangun dari pembaca. Karena penulis menyadari dengan penuh bahwa dalam penyusunan penelitian ini masih banyak kekurangan. Akhir kata, semoga penelitian ini dapat memberikan manfaat bagi pembacanya.

Yogyakarta, 10 Oktober 2025

Penulis

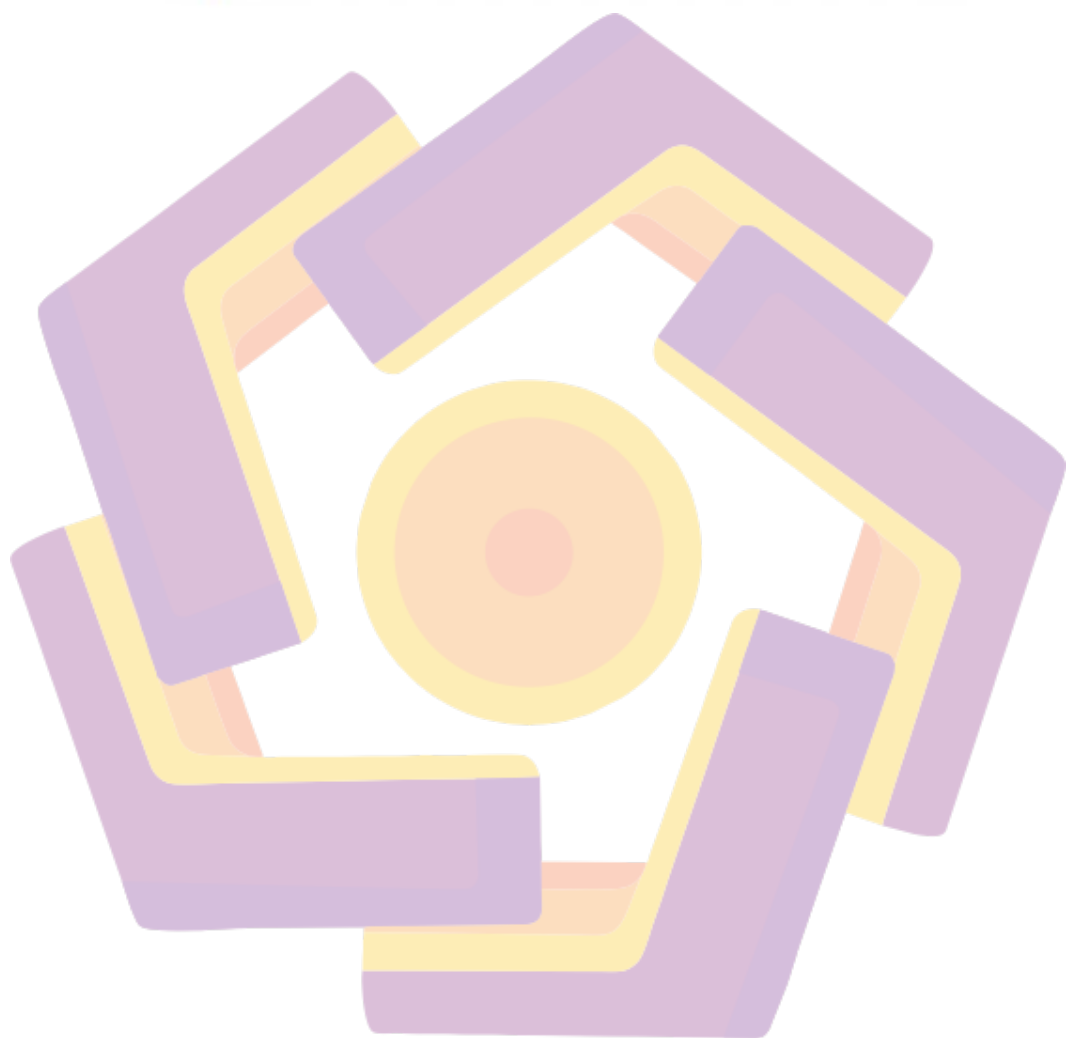
## DAFTAR ISI

|                                        |       |
|----------------------------------------|-------|
| HALAMAN JUDUL.....                     | i     |
| HALAMAN PERSETUJUAN.....               | iii   |
| HALAMAN PENGESAHAN.....                | iv    |
| HALAMAN PERNYATAAN KEASLIAN TESIS..... | v     |
| HALAMAN PERSEMBAHAN.....               | vi    |
| KATA PENGANTAR.....                    | vii   |
| DAFTAR ISI.....                        | ix    |
| DAFTAR TABEL.....                      | xiii  |
| DAFTAR GAMBAR.....                     | xiv   |
| DAFTAR ISTILAH.....                    | xvi   |
| INTISARI.....                          | xvii  |
| ABSTRACT.....                          | xviii |
| BAB I PENDAHULUAN.....                 | 1     |
| 1.1. Latar Belakang Masalah.....       | 1     |
| 1.2. Rumusan Masalah.....              | 4     |
| 1.3. Batasan Masalah.....              | 5     |
| 1.4. Tujuan Penelitian.....            | 5     |
| 1.5. Manfaat Penelitian.....           | 6     |
| BAB II TINJAUAN PUSTAKA.....           | 8     |
| 2.1. Tinjauan Pustaka.....             | 8     |
| 2.2. Keaslian Penelitian.....          | 14    |

|                                                             |           |
|-------------------------------------------------------------|-----------|
| 2.3. Landasan Teori.....                                    | 18        |
| 2.3.1 Data dan Persiapan Data .....                         | 18        |
| 2.3.1.1 Dataset.....                                        | 18        |
| 2.3.1.2 Preprocessing Data.....                             | 19        |
| 2.3.1.3 Data Split.....                                     | 19        |
| 2.3.2 Balancing Data.....                                   | 20        |
| 2.3.2.1 SMOTE (Synthetic Minority Oversampling Technique).. | 20        |
| 2.3.3 Normalisasi Data.....                                 | 21        |
| 2.3.4 <i>Dimensionaly Reduction</i> .....                   | 21        |
| 2.3.5 <i>Modeling Machine Learning Random Forest</i> .....  | 29        |
| 2.3.6 Evaluasi Model .....                                  | 32        |
| 2.3.6.1 Confusion Matrix .....                              | 32        |
| 2.3.6.2 Akurasi .....                                       | 33        |
| 2.3.6.3 Precision.....                                      | 34        |
| 2.3.6.4 Recall.....                                         | 35        |
| 2.3.6.5 F1-Score .....                                      | 36        |
| <b>BAB III METODE PENELITIAN.....</b>                       | <b>37</b> |
| 3.1. Jenis, Sifat, dan Pendekatan Penelitian.....           | 37        |
| 3.2. Metode Pengumpulan Data.....                           | 38        |
| 3.3. Metode Analisis Data.....                              | 39        |

|                                                                       |           |
|-----------------------------------------------------------------------|-----------|
| 3.4. Alur Penelitian .....                                            | 40        |
| 3.4.1 Tahapan Pengumpulan Data .....                                  | 42        |
| 3.5. Keterbatasan Metodologi dan Potensi Bias Penelitian .....        | 43        |
| <b>BAB IV HASIL PENELITIAN DAN PEMBAHASAN .....</b>                   | <b>44</b> |
| 4.1. Pengumpulan Data .....                                           | 44        |
| 4.1.1 Data Set .....                                                  | 44        |
| 4.1.2 Preprocessing Data .....                                        | 45        |
| 4.1.3 Data Split .....                                                | 46        |
| 4.2. Balancing Data .....                                             | 47        |
| 4.3. Normalisasi Data .....                                           | 48        |
| 4.4 Analisis Performa Model Sebelum dan Sesudah Reduksi Dimensi ..... | 49        |
| 4.5 <i>Dimensionality Reduction Data</i> .....                        | 50        |
| 4.6 <i>Feature Importance</i> .....                                   | 52        |
| 4.7 Hasil <i>Feature Importance</i> .....                             | 52        |
| 4.8 <i>Stepwise Feature Reduction</i> .....                           | 54        |
| 4.9 Hasil <i>Summary Stepwise Feature Reduction</i> .....             | 59        |
| 4.10 Perbandingan Stepwise Reduction Tanpa Menerapkan .....           | 65        |
| 4.11 <i>Modelling</i> .....                                           | 67        |
| 4.12 Evaluasi Model .....                                             | 67        |
| <b>BAB V PENUTUP .....</b>                                            | <b>72</b> |
| 5.1. Kesimpulan .....                                                 | 72        |
| 5.2. Saran .....                                                      | 74        |

|                      |    |
|----------------------|----|
| DAFTAR PUSTAKA ..... | 76 |
| LAMPIRAN .....       | 87 |



## DAFTAR TABEL

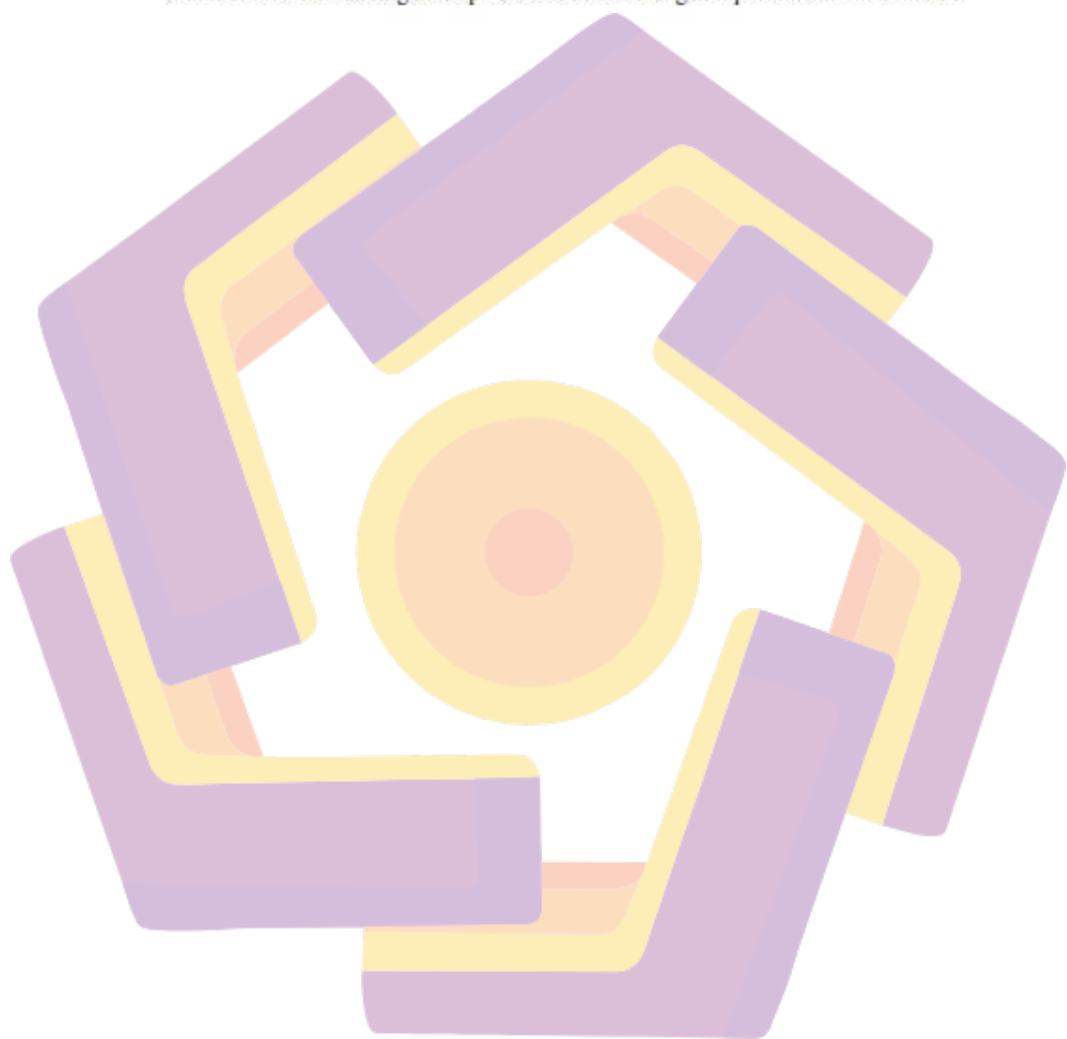
|                                                                   |    |
|-------------------------------------------------------------------|----|
| Tabel 2.1. Detail Tinjauan Pustaka dan Perbandingan dengan .....  | 12 |
| Tabel 2.2. Matriks literatur review dan posisi penelitian .....   | 14 |
| Tabel 3.1. Daftar Feature Dataset Spotify Study Music .....       | 38 |
| Tabel 4.1 Drop Data <i>Stepwise Feature Reduction</i> .....       | 61 |
| Tabel 4.2 Perbandingan Tanpa Menerapkan Feature Importance .....  | 65 |
| Tabel 4.3 Sebaran Data Setelah Normalisasi dan Sebaran Data ..... | 67 |
| Tabel 4.4 <i>Model Evaluation Summary</i> .....                   | 70 |
| Tabel 4.5 <i>Model Performance per Class</i> .....                | 70 |
| Tabel 4.6 <i>Confusion Matrix</i> .....                           | 71 |
| Tabel 5.1 <i>Selected Features and Model Accuracy</i> .....       | 73 |

## DAFTAR GAMBAR

|             |                                                               |    |
|-------------|---------------------------------------------------------------|----|
| Gambar 2.1  | <i>Classification of Dimensionally Reduction</i>              | 22 |
| Gambar 2.2  | <i>Embedded Method Random Forest using Feature Importance</i> | 27 |
| Gambar 2.3  | <i>Wrapper Method -Stepwise Selection</i>                     | 29 |
| Gambar 2.4  | <i>Random Forest</i>                                          | 30 |
| Gambar 2.5  | <i>Predicted Class</i>                                        | 33 |
| Gambar 3.1  | <i>Research FLOW Stages</i>                                   | 41 |
| Gambar 4.1  | <i>Distribution Class Before and After SMOTE</i>              | 47 |
| Gambar 4.2  | <i>Data Normalization</i>                                     | 48 |
| Gambar 4.3  | <i>Flow Dimensionally Reduction Data</i>                      | 51 |
| Gambar 4.4  | <i>Visualisasi Hasil Feature Importance</i>                   | 52 |
| Gambar 4.5  | Penghapusan Mode                                              | 56 |
| Gambar 4.6  | Penghapusan Popularity                                        | 57 |
| Gambar 4.7  | Penghapusan Duration_ms                                       | 57 |
| Gambar 4.8  | Penghapusan Danceability                                      | 57 |
| Gambar 4.9  | Penghapusan Key                                               | 58 |
| Gambar 4.10 | Penghapusan Liveness                                          | 58 |
| Gambar 4.11 | Penghapusan Instrumentalness                                  | 59 |
| Gambar 4.12 | Penghapusan Accousticness                                     | 59 |
| Gambar 4.13 | <i>Visualisasi Hasil Stepwise Feature Reduction</i>           | 60 |

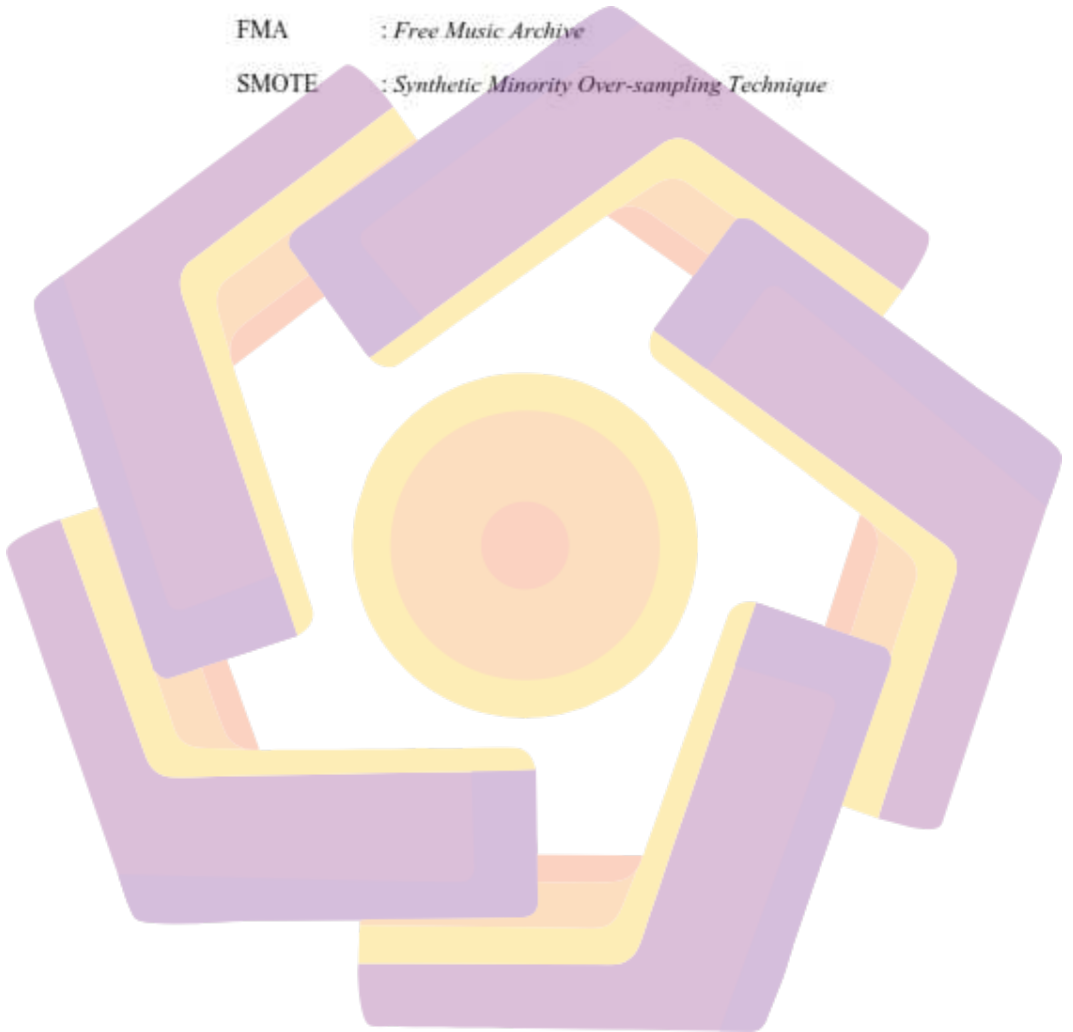
Gambar 4.14 Perbandingan Stepwise Reduction Tanpa Feature Importance..... 66

Gambar 4.15 Perbandingan Stepwise Reduction Dengan mportance..... 66



## DAFTAR ISTILAH

|       |                                                     |
|-------|-----------------------------------------------------|
| RF    | : <i>Random Forest</i>                              |
| FMA   | : <i>Free Music Archive</i>                         |
| SMOTE | : <i>Synthetic Minority Over-sampling Technique</i> |



## INTISARI

Penelitian "*Reduksi Dimensi Dan Klasifikasi Musik Untuk Belajar Menggunakan Feature Importance Dan Random Forest*" bertujuan untuk mengklasifikasikan jenis musik yang sesuai untuk aktivitas pembelajaran berdasarkan fitur audio Spotify menggunakan algoritma *Random Forest*, dengan menggunakan pendekatan analisis kepentingan fitur dan reduksi fitur. Dataset Musik Studi Spotify terdiri dari 172.819 lagu dengan 12 fitur audio, yang dikategorikan ke dalam tiga kategori utama: Lagu Pop, Soundtrack Klasik, dan Lagu Lo-fi. Metode penelitian meliputi prapemrosesan data, pemrosesan kelas menggunakan SMOTE, normalisasi data menggunakan Min-Max, analisis Kepentingan Fitur, validasi silang, dan Reduksi Fitur untuk mengidentifikasi fitur yang paling berpengaruh dalam mengklasifikasikan musik pembelajaran.

Hasil normalisasi menunjukkan bahwa semua fitur berada dalam rentang 0,0-1,0 tanpa mengubah karakteristik distribusi asli. Model RF menunjukkan kinerja yang luar biasa dengan akurasi rata-rata 99% pada validasi silang dan 99,9% pada data pelatihan, yang menunjukkan kemampuan model untuk mengenali pola musik secara efisien. Analisis Kepentingan Fitur menunjukkan bahwa energi, kenyaringan, akustik, instrumentalitas, dan keaktifan merupakan faktor paling signifikan dalam membedakan karakteristik musik untuk pembelajaran. Sementara itu, hasil Reduksi Fitur menunjukkan bahwa penghapusan fitur *mode*, *Popularitas*, *Durasi\_ms*, dan *Ketertarikan* tidak mengurangi akurasi model secara signifikan. Namun, penghapusan fitur kunci, keaktifan, instrumentalitas, dan keingintahuan menyebabkan akurasi menurun menjadi 0,973.

Penelitian ini menyimpulkan bahwa kombinasi SMOTE, normalisasi, dan *Random Forest* dengan analisis fitur berhasil mengoptimalkan klasifikasi musik pembelajaran dengan akurasi 98,77% pada data uji. Fitur akustik, instrumentalitas, dan keingintahuan direkomendasikan untuk dipertahankan karena memainkan peran penting dalam menjaga stabilitas dan akurasi model klasifikasi musik, yang mendukung proses pembelajaran.

Kata Kunci: *Spotify Study Music*, *SMOTE*, *Feature Importance*, *Feature Reduction*, *Random Forest*

## ABSTRACT

Music has a significant impact on the way a person thinks and feels in their daily activities. This study aims to categorize the types of music that are suitable for learning activities by using Spotify's audio feature, to create a more flexible and personalized music recommendation system. The dataset used comes from Spotify Study Music which consists of 172,819 songs with 12 audio features, which are grouped into three main categories, namely Pop Tracks, Classical Soundtracks, and Lo-fi Tracks. The research process includes data pre-processing, handling class imbalances using SMOTE, data normalization, feature significance Analysis, Cross Validation, and Feature Reduction. Normalization results show that all features have been in the range of 0.0-1.0 without changing the characteristics of the original distribution. The Random Forest Model performed exceptionally well with an average accuracy rate of 99% on cross-validation and 99.9% on training data, indicating the model's ability to efficiently recognize musical patterns. Important Feature Analysis shows that energy, loudness, acousticness, instrumentality, and liveness have the most significant influence in distinguishing music characteristics for learning, while mode, popularity, duration\_ms, and danceability when removed using Feature Reduction analysis show a significant decrease in accuracy. This study recommends maintaining the features of acousticness, instrumentality, and liveness because it plays an important role in maintaining the stability and accuracy of music classification models that support the learning process.

**Keyword:** Spotify Study Music, SMOTE, Feature Importance, Feature Reduction, Random Forest, Classification.

## BAB I

### PENDAHULUAN

#### 1.1. Latar Belakang Masalah

Musik telah menjadi komponen yang sangat penting dalam kehidupan manusia sehari-hari, sebagai kebutuhan dasar untuk merangsang pendengaran, yang memengaruhi aspek kognitif dan emosional [1]. Mendengarkan musik yang tepat dapat menghasilkan berbagai manfaat dan menimbulkan reaksi positif dalam aktivitas otak [2]. Ketepatan dalam memilih dan mengelompokkan musik mempunyai peranan yang sangat penting dalam meningkatkan dampak psikologis dan fisiologis, terutama musik instrumental atau musik dengan irama lembut, yang dapat membantu relaksasi dan meningkatkan konsentrasi dalam meningkatkan proses belajar [3], [4], [5]. Spotify adalah platform streaming yang menawarkan berbagai genre dengan sistem penyimpanan cloud di servernya [6], [7]. Platform ini menggunakan sistem klasifikasi yang merujuk pada dataset George Tzanetakis dari tahun 2002, yang mencakup berbagai genre seperti Blues, Klasik, Country, Disco, Hip-hop, Jazz, Metal, Pop, Reggae, Rock, serta Free Music Archive (FMA) yang mencakup genre kontemporer seperti Elektronik, Eksperimental, Folk, Hip-Hop, Instrumental, Internasional, dan Rock [8].

Riset awal Rebecca Jane Scarratt pada tahun 2019 mengeksplorasi ciri-ciri musik menggunakan data insight YouTube dengan penekanan pada perancangan analisis khusus untuk mendukung proses menuju tertidur atau musik untuk tidur [5]. Perkembangan terjadi pada tahun 2021 ketika Scarratt memperluas penelitiannya secara komprehensif untuk mengidentifikasi dan mengkarakterisasi

musik pemicu tidur yang efektif [9]. Dalam studi lanjutan pada tahun 2023, algoritma K-Means diterapkan untuk membandingkan musik yang digunakan untuk belajar dan tidur. Studi ini menunjukkan kesamaan yang signifikan dalam fitur audio, kategori genre, dan hasil pengelompokan. Hasilnya adalah lingkungan audio yang mendukung belajar dan beristirahat tanpa mengganggu konsentrasi [10]

Penelitian ini mengkaji berbagai metode untuk menyeimbangkan data, mencakup pendekatan pada level data, level algoritma, serta hybrid/ensemble, dengan menerapkan teknik seperti SMOTE, ADASYN, dan RUSBoost yang bertujuan memperbaiki distribusi kelas dan meningkatkan ketepatan klasifikasi pada data yang tidak seimbang[11]. Penelitian yang dilakukan oleh Khan dan rekan-rekan tahun 2022 meneliti dampak pemilihan fitur terhadap akurasi dalam klasifikasi popularitas musik menggunakan algoritma Random Forest. Hasil studi menunjukkan bahwa mengurangi jumlah fitur dapat mengurangi kompleksitas tanpa mengakibatkan penurunan akurasi yang signifikan[12]. Pada penelitian lain, Saragih tahun 2023 mengembangkan pendekatan itu dengan menerapkan *Random Forest Classifier*, validasi silang, dan analisis pentingnya fitur untuk meramalkan popularitas lagu berdasarkan fitur audio dari Spotify, namun belum memanfaatkan teknik penyeimbangan data seperti SMOTE[13]. Sementara itu, Maitsa dan Winarsih tahun 2025 menambahkan langkah-langkah normalisasi data, penggunaan SMOTE, dan GridSearch CV dalam klasifikasi genre lagu, tetapi belum meneliti pengaplikasian *looping Feature Importance* sebagai metode pengurangan fitur berulang yang dapat meningkatkan kinerja *Random Forest*[14]. Selain itu, Prasetya pada tahun 2024 menemukan bahwa kombinasi antara SMOTE dan *Random Forest*

dapat meningkatkan akurasi klasifikasi pada data yang tidak seimbang sampai 96,28% [15]. Dalam penelitian lain, Bernard Ginn Jr pada tahun 2024 menunjukkan kemampuan *Random Forest* dalam mengklasifikasikan suasana hati lagu dengan tingkat akurasi mencapai 95% dengan memanfaatkan fitur audio dan lirik[16]. Metode machine learning seperti *Random Forest* dianggap relevan karena kemampuannya dalam mengenali pola dari data musik yang kompleks dan bervariasi[17]. Berdasarkan hasil tersebut, terdapat kekurangan dalam penelitian yang ada, yaitu masih belum ada studi yang secara menyeluruh menggabungkan SMOTE, normalisasi data, validasi silang, analisis pentingnya fitur, serta *looping Feature Reduction* dalam satu kerangka model *Random Forest* dalam bidang musik. Penggabungan kelima elemen ini berpotensi menghasilkan model yang lebih efektif, stabil, dan akurat dalam mengklasifikasikan data musik yang sangat tidak seimbang.

Sebuah penelitian yang dilakukan pada tahun 2021 dengan judul "Menyelami Perilaku Pengguna di Spotify" mengungkapkan bahwa pendengar baru menghabiskan waktu 50% lebih lama untuk menikmati musik dibandingkan dengan pendengar veteran, disertai dengan peningkatan 0,95% pada pengikut harian untuk daftar putar nyaris dua kali lipat dari lagu-lagu yang dimiliki oleh artis di label besar serta menunjukkan bahwa platform ini juga membantu dalam meningkatkan kemampuan bahasa Inggris melalui fitur lirik, beragam genre, dan paparan terhadap bahasa asli yang melatih pengucapan dan pemahaman mendengarkan [7],[17]. Pemilihan genre musik diasumsikan berpengaruh terhadap persona pendengar secara mendalam. Analisis spesifik terhadap fitur-fitur Spotify memungkinkan

pemahaman yang lebih kompleks tentang bagaimana setiap fitur dalam genre musik berkontribusi terhadap respons pendengar. Oleh karena itu, topik ini penting untuk diteliti lebih lanjut guna memahami pola fitur Spotify dalam mengklasifikasikan jenis musik yang tepat untuk mendukung proses belajar.

Penelitian ini berfokus pada klasifikasi musik untuk tujuan pembelajaran dengan memanfaatkan fitur audio Spotify seperti tempo dan energi. Dasar penelitian ini adalah temuan Rebecca Jane Scarratt, yang telah mengidentifikasi tiga kelas utama dalam *Spotify Study Music/SSM* yaitu *Pop Tracks*, *Classical Soundtracks*, dan *Lo-fi Tracks*.

Dataset yang digunakan dalam penelitian ini terdiri atas 15 atribut dengan total 172.819 entri data. Dalam proses pengembangan model, diterapkan tahapan seleksi fitur melalui analisis feature importance untuk mengidentifikasi atribut yang memiliki kontribusi signifikan terhadap variabel target. Selanjutnya, dilakukan proses reduksi fitur (feature reduction) guna menyederhanakan representasi data serta meningkatkan efisiensi dan kinerja model. Tahap akhir melibatkan penerapan algoritma *Random Forest* untuk melakukan klasifikasi, dengan tujuan memperoleh performa yang optimal, mengenali pola data secara lebih efektif, serta memberikan dasar pengembangan sistem rekomendasi musik yang adaptif dan personal.

## **1.2. Rumusan Masalah**

1. Bagaimana proses mengolah dataset, yang diambil dari penelitian Rebecca Jane Scarratt?
2. Bagaimana mengolah dataset yang hanya menggunakan Algoritma *Random Forest* pada musik untuk belajar?

3. Bagaimana mengklasifikasikan musik belajar menggunakan Algoritma *Random Forest*?
4. Bagaimana hasil analisis data, dan hasil dari perbandingan penelitian sebelumnya menggunakan klasifikasi data *K-Means* dan penelitian menggunakan klasifikasi data *Random Forest* khusus musik untuk belajar

### 1.3. Batasan Masalah

Bagian ini memuat penjelasan tentang:

- a. Studi ini khusus menggunakan dataset yang berasal dari library spotify sesuai penelitian Rebecca Jane Scarratt tanpa memasukkan data tambahan.
- b. Untuk pemrosesan dataset, penelitian ini hanya berfokus pada penerapan Algoritma *Random Forest*.
- c. Untuk proses klasifikasi data menggunakan Algoritma *Random Forest*, *Feature Importance* dan *Feature Reduction*.
- d. Fokus penelitian khusus pada analisis musik untuk belajar, tidak termasuk kategori musik untuk kegiatan lainya.

### 1.4. Tujuan Penelitian

Bagian ini memuat penjelasan secara spesifik:

- a. Memproses data dari dataset *Spotify* khusus untuk kategori musik belajar dengan mengacu pada parameter yang digunakan dalam penelitian Rebecca Jane Scarratt. Mengolah fitur audio menggunakan Algoritma *Random Forest*

mempertahankan komponen esensial yang memengaruhi karakteristik musik belajar.

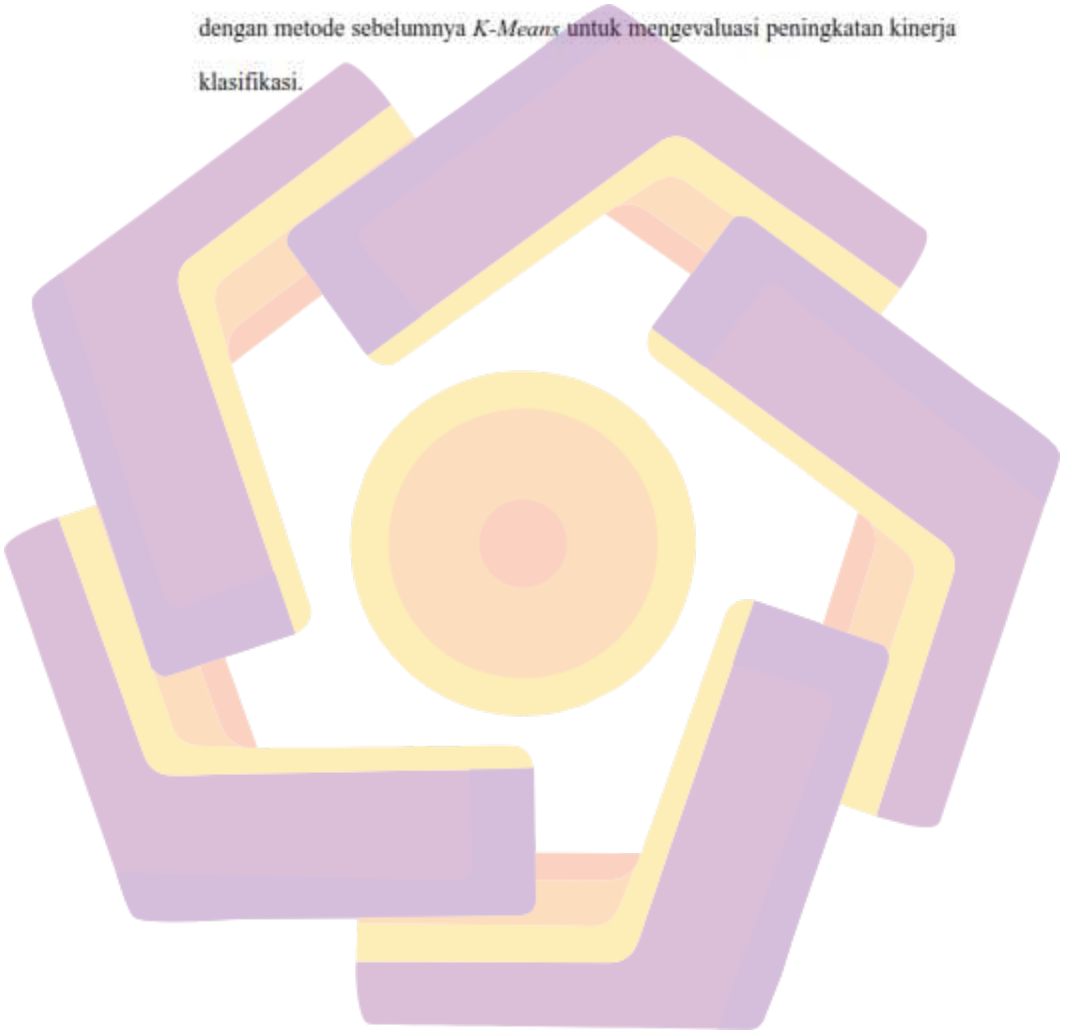
- b. Mengimplementasikan Algoritma *Random Forest* untuk mengklasifikasikan musik belajar guna mengidentifikasi pola dan kategorisasi yang lebih akurat. Dan *Feature Importance* dan *Feature Reduction* untuk mengetahui pengaruh fitur pada spotify
- c. Mengevaluasi Algoritma *Random Forest* dengan membandingkan hasilnya terhadap K-Means pada penelitian sebelumnya, sekaligus mengukur performa klasifikasi untuk musik belajar.

#### 1.5. Manfaat Penelitian

Bagian ini memuat penjelasan tentang:

- a. Manfaat secara akademis akan memberikan kontribusi dalam pengembangan metode pembelajaran mesin untuk analisis musik, dengan fokus pada klasifikasi dengan Algoritma *Random Forest* pada data audio.
- b. Manfaat praktis bagi pengguna spotify Menyediakan bantuan bagi pengguna Spotify, khususnya di kalangan pelajar dan pekerja, dalam menemukan rekomendasi musik yang sesuai untuk meningkatkan konsentrasi dalam pembelajaran, berdasarkan karakteristik audio yang terukur.
- c. Manfaat bagi pengembang aplikasi musik memberikan wawasan yang berharga dalam pengembangan sistem rekomendasi musik yang lebih personal, *Random Forest* untuk meningkatkan akurasi dalam klasifikasi musik berdasarkan tujuan penggunaannya.

- d. Manfaat metodologis menunjukkan efektivitas Algoritma *Random Forest* dalam menangani data musik yang kompleks, serta memberikan perbandingan dengan metode sebelumnya *K-Means* untuk mengevaluasi peningkatan kinerja klasifikasi.



## BAB II

### TINJAUAN PUSTAKA

#### 2.1. Tinjauan Pustaka

Berbagai penelitian telah mengeksplorasi hubungan antara musik dan aktivitas manusia dengan menggunakan berbagai pendekatan. Dalam publikasi oleh Heggli et al. pada tahun 2021, menyimpulkan bahwa musik tidur umumnya memiliki tempo yang lambat, bersifat instrumental, dan cenderung lembut, meskipun terdapat variasi yang signifikan dalam enam subkelompok [18]. Para peneliti tersebut memberikan eksplorasi yang lebih mendalam mengenai hubungan antara fitur audio dan dampak fisiologis yang ditimbulkannya. Penelitian lain yang dilakukan oleh Scarratt et al. pada tahun 2023 dalam jurnal *Scientific Reports* membandingkan musik untuk belajar dan musik untuk tidur, menemukan adanya kesamaan dalam fitur-fitur musik walaupun fungsi keduanya berbeda, sambil menekankan pentingnya konteks dalam pemilihan musik [14]. Temuan serupa juga didukung oleh penelitian Scarratt et al. pada tahun 2021 yang menganalisis daftar putar musik untuk tidur di Spotify, yang mengungkapkan bahwa preferensi terhadap musik untuk tidur bersifat individual, dengan beberapa lagu yang memiliki tempo cepat tetap digunakan untuk tujuan relaksasi. Pada penelitian tersebut, menggunakan metode KN untuk menentukan musik mana yang dapat digunakan untuk mendukung aktifitas belajar ataupun untuk memulai tidur. Data musik diambil dari database spotify dengan beberapa *keywords* khusus seperti *sleep/sleeping*, *study* atau *studying*. Setelah mendapatkan data musik yang diinginkan, lalu data direduksi secara manual dengan cara melihat deskripsi dari

data musik tersebut. Kemudian data reduksi manual tersebut di filter dan diolah menggunakan metode KN untuk mengklasifikasi data yang diperoleh.

Penelitian ini mengkaji berbagai metode untuk menyeimbangkan data, mencakup pendekatan pada level data, level algoritma, serta hybrid/ensemble, dengan menerapkan teknik seperti SMOTE, ADASYN, dan RUSBoost yang bertujuan memperbaiki distribusi kelas dan meningkatkan ketepatan klasifikasi pada data yang tidak seimbang [11]. Penelitian yang dilakukan oleh Khan dan rekan-rekan tahun 2022 meneliti dampak pemilihan fitur terhadap akurasi dalam klasifikasi popularitas musik menggunakan algoritma Random Forest. Hasil studi menunjukkan bahwa mengurangi jumlah fitur dapat mengurangi kompleksitas tanpa mengakibatkan penurunan akurasi yang signifikan [12]. Pada penelitian lain, Saragih tahun 2023 mengembangkan pendekatan itu dengan menerapkan *Random Forest Classifier*, validasi silang, dan analisis pentingnya fitur untuk meramalkan popularitas lagu berdasarkan fitur audio dari Spotify, namun belum memanfaatkan teknik penyeimbangan data seperti SMOTE[13]. Sementara itu, Maitsa dan Winarsih tahun 2025 menambahkan langkah-langkah normalisasi data, penggunaan SMOTE, dan GridSearch CV dalam klasifikasi genre lagu, tetapi belum meneliti pengaplikasian looping *Feature Importance* sebagai metode pengurangan fitur berulang yang dapat meningkatkan kinerja *Random Forest* [19].

*Feature Selection* merupakan proses krusial untuk meningkatkan akurasi dan mempercepat waktu komputasi dengan memilih fitur informatif dari dataset, dimana *Random Forest* menurut Yuan et al. pada tahun 2023 menggunakan mekanisme *Feature Importance* ranking berdasarkan pengurangan indeks untuk

mengukur kontribusi setiap fitur dalam meningkatkan kemurnian dataset [20]. Information Gain (IG) sebagai metode *Feature Selection* penelitian menurut Prasetyowati et al. pada tahun 2022 menghitung bobot nilai entropi maksimum untuk menentukan pentingnya fitur, dengan menetapkan threshold berdasarkan nilai standar deviasi atau median transformasi Fast Fourier Transform (FFT), sehingga hanya fitur dengan nilai IG yang lebih besar atau sama dengan threshold yang dipilih sebagai fitur informatif dalam proses klasifikasi [21].

*Feature Selection* dan *Feature Reduction* merupakan proses penting dalam machine learning untuk meningkatkan kemampuan pemahaman, waktu pemrosesan model klasifikasi, dan mengatasi masalah dimensionalitas tinggi, dimana metode *Feature Selection* yang efektif tidak hanya menghilangkan fitur yang dapat menghambat proses klasifikasi tetapi juga secara bersamaan meningkatkan akurasi klasifikasi. *Random Forest* menurut Mahmood et al. pada penelitiannya tahun 2025 menggunakan *importance (mean decrease in impurity)* sebagai metode *Feature Importance* ranking yang didefinisikan sebagai total pengurangan ketidakmurnian node yang dicapai dengan menggunakan fitur tersebut untuk membagi data, yang dirata-ratakan di semua pohon keputusan, dimana skor *importance* yang tinggi menunjukkan bahwa fitur spesifik tersebut lebih penting untuk klasifikasi dan sebaliknya [1]. Sementara itu, penelitian menurut Khabiri et al. pada tahun 2024 metode *Feature Reduction* yang dapat mengurangi jumlah dimensi dalam dataset yang terdiri dari banyak variabel saling berhubungan dengan mempertahankan sebanyak mungkin variasi dalam dataset, menghasilkan variabel baru yang tidak memiliki korelasi dan diatur berdasarkan nilai efektivitasnya, dimana dalam

penelitiannya berhasil mereduksi 46,872 fitur menjadi hanya 10 komponen utama yang mempertahankan variasi terbesar dari seluruh principal components untuk digunakan sebagai input classifier [22].

Klasifikasi merupakan salah satu tema dasar dalam deep learning yang bertujuan untuk menciptakan model atau pola dari data yang tersedia untuk mendukung proses pengambilan keputusan. Beragam algoritma untuk klasifikasi telah dibuat, seperti *Decision Tree*, *Naïve Bayes*, *KNN*, *SVM* & *Random Forest* [23].

Selain itu, Prasetya pada tahun 2024 menemukan bahwa kombinasi antara SMOTE dan *Random Forest* dapat meningkatkan akurasi klasifikasi pada data yang tidak seimbang sampai 96,28% [15]. Dalam penelitian lain, Bernard Ginn Jr pada tahun 2024 menunjukkan keampuhan *Random Forest* dalam mengklasifikasikan suasana hati lagu dengan tingkat akurasi mencapai 95% dengan memanfaatkan fitur audio dan lirik[16]. Metode machine learning seperti *Random Forest* dianggap relevan karena kemampuannya dalam mengenali pola dari data musik yang kompleks dan bervariasi[17]. Berdasarkan tinjauan literatur di atas, dapat disimpulkan bahwa berbagai studi sebelumnya telah meneliti hubungan antara musik, fitur audio, dan penerapan metode pembelajaran mesin dalam konteks klasifikasi dan prediksi. Untuk mempermudah pemahaman dan meningkatkan keterbacaan, ringkasan hasil penelitian ini disajikan dalam bentuk tabel. Tabel ini berisi informasi kunci dari setiap studi terkait, termasuk pendekatan yang digunakan dan dampaknya terhadap kinerja model terlihat seperti tabel dibawah ini:

Tabel 2.1. Detail Tinjauan Pustaka dan Perbandingan dengan Algoritma Klasifikasi Lainnya

| Peneliti & Tahun          | Fokus Penelitian                                                        | Metode & Data                                                                                                                                                                                                     | Dampak terhadap Akurasi                                                                                                                                                                  | Paper Information                                                                                                                                                                                                                                                                                                   |
|---------------------------|-------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Prasetya (2024)           | Klasifikasi data tidak seimbang                                         | SMOTE + Random Forest                                                                                                                                                                                             | Kombinasi SMOTE dan Random Forest meningkatkan akurasi hingga 96,28%                                                                                                                     | J. Prasetya and A. Abdurakhman, "Comparison Of Smote Random Forest And Smote K-Nearest Neighbors Classification Analysis On Imbalanced Data," <i>MEDIA STATISTIKA</i> , vol. 15, no. 2, pp. 198–208, Apr. 2023, doi: 10.14710/medstat.15.2.198-208.                                                                 |
| Bernard Ginn Jr (2024)    | Klasifikasi suasana hati berdasarkan lagu yang dipilih                  | Random Forest, fitur audio & lirik                                                                                                                                                                                | Random Forest mencapai akurasi hingga 95% dalam klasifikasi mood lagu                                                                                                                    | Bernard Ginn Jr and Sejoniq Johnson, "Evaluating the Efficacy of Random Forest in Mood Classification for Music Recommendation System," 2024, doi: DOI:10.13140/RG.2.2.22524.04485.                                                                                                                                 |
| Putri Ayu Firmanda (2025) | Prediksi / klasifikasi berbasis fitur dengan algoritma machine learning | Dataset numerik, Decision Tree dan Random Forest                                                                                                                                                                  | Akurasi tertinggi menggunakan Random Forest yang dilaporkan mencapai 99%, sedangkan algoritma Decision Tree mencapai 97,5%.[24]                                                          | Putri Ayu Firmanda <i>et al.</i> , "Analisis Perbandingan Decision Tree dan Random Forest dalam Klasifikasi Penjualan Produk pada Supermarket," <i>Emerging Statistics and Data Science Journal</i> , vol. 3, no. 1, 2025.                                                                                          |
| Jiulin Song (2023)        | Perbandingan Random Forest dan XGBoost pada klasifikasi emosi musik     | Turkish Music Emotion Dataset (UCL), preprocessing (duplicate & outlier removal), Z-score standardization, feature reduction berbasis Pearson correlation (pemilihan 10 fitur teratas), Random Forest dan XGBoost | Sebelum feature reduction: tidak dilaporkan secara eksplisit. Setelah feature reduction: Random Forest mencapai akurasi 80,8%, sedangkan XGBoost mencapai akurasi 75% pada data uji.[25] | J. Song, "Comparison and Analysis of Accuracy of Traditional Random Forest Machine Learning Model and XGBoost Model On Music Emotion Classification Dataset," in <i>ACM International Conference Proceeding Series</i> , Association for Computing Machinery, Oct. 2023, pp. 712–716. doi: 10.1145/3650215.3650340. |

|              |                                                                                                                 |                                                                   |                                                                                              |                                                                                                                                                                                     |
|--------------|-----------------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------|----------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Meng (2024)  | Perbandingan CNN, VGG16, dan XGBoost pada klasifikasi genre musik                                               | GTZAN Dataset; MFCC & Mel Spectrogram; XGBoost (tabular features) | XGBoost + MFCC: 97% [26]                                                                     | Y. Meng, "Music Genre Classification: A Comparative Analysis of CNN and XGBoost Approaches with Mel-frequency Cepstral Coefficients and Mel Spectrograms."                          |
| Maita (2025) | Comparison of Hyperparameter Tuning in Decision Tree and Random Forest Algorithms for Song Genre Classification | Random Forest + Decision Tree                                     | Akurasi menggunakan RF 85% namun terjadi penurunan ketika menggunakan algoritma DT 84 % [19] | Bernard Ginn Jr and Sejoniq Johnson, "Evaluating the Efficacy of Random Forest in Mood Classification for Music Recommendation System," 2024, doi: DOI:10.13140/RG.2.2.22524.04485. |

## 2.2. Keaslian Penelitian

Tabel 2.2. Matriks literatur review dan posisi penelitian  
Klasifikasi Musik Untuk Belajar Berdasarkan Fitur Audio Spotify Menggunakan Random Forest Dengan Pendekatan Analisis  
Kepentingan Fitur & Reduksi Fitur

| No | Judul                                                                      | Peneliti,<br>Media<br>Publikasi, dan<br>Tahun                                                       | Tujuan Penelitian                                                                                            | Kesimpulan                                                                                                                                                               | Saran atau Kelemahan                                                          | Perbandingan                                                                                                                                     |
|----|----------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------|
| 1  | Motivations for using music for sleep: insights from YouTube comments      | Rebecca J. Scurratt, Jan Stupacher, Peter Vuust & Kira V. Jespersen — Aarhus University, 2019       | Memahami motivasi user youtube dalam menggunakan youtube untuk tidur melalui analisis teks komentar YouTube. | Ada 4 motivasi yang ditemukan dalam penelitian ini untuk relaksasi, menghilangkan distraksi, untuk mengatur suasana hati dan kebiasaan menggunakan youtube sebelum tidur | Hanya fokus pada motivasi user menggunakan youtube tidak mengkaji fitur audio | Memberikan konteks perilaku user terhadap musik,<br><br>dapat diadaptasi untuk penelitian lanjutan yaitu penelitian tentang musik untuk belajar. |
| 2  | The music that people use to sleep: universal and subgroup characteristics | Rebecca J. Scurratt, Ole A. Heggli, Peter Vuust & Kira V. Jespersen — Frontiers in Psychology, 2021 | Identifikasi karakteristik universal musik yang digunakan untuk tidur dengan analisis big data dari Spotify. | Musik untuk tidur biasanya tenang, lembut, lambat, dan instrumental, namun ditemukan adanya perbedaan yang signifikan jika                                               | Belum mempertimbangkan faktor psikologi dan personalisasi pengguna.           | Fokus pada konteks tidur;<br><br>penelitian tesis memperluas konteks ke musik untuk belajar.                                                     |

| No | Judul                                                                                                                       | Peneliti, Media Publikasi, dan Tahun                                                                                           | Tujuan Penelitian                                                                                                           | Kesimpulan                                                                                                               | Saran atau Kelemahan                                                             | Perbandingan                                                                                                                                                                                                                                |
|----|-----------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
|    |                                                                                                                             |                                                                                                                                |                                                                                                                             | menggunakan pada musik pop dan musik yang energik.                                                                       |                                                                                  |                                                                                                                                                                                                                                             |
| 3  | Music that is used while studying and music that is used for sleep share similar musical features, genres and subgroups     | Rebecca J. Scarratt, Ole A. Heggli, Peter Vanst & Makiko Sadakata — Scientific Reports (Nature), 2023                          | Melakukan perbandingan karakteristik musik untuk belajar dan musik untuk tidur menggunakan dataset Spotify.                 | Musik belajar dan tidur memiliki fitur audio dan genre yang sangat mirip; keduanya menciptakan suasana yang menenangkan. | Tidak menilai efektivitas terhadap performa belajar atau konsentrasi.            | Fokus hanya pada klasifikasi musiknya<br><br>Penelitian tesis mengembangkan pendekatan ini dengan klasifikasi berbasis <i>Random Forest</i> dan analisis fitur importance dan kemudian fitur importance dilakukan reduksi pada setiap fitur |
| 4  | SMOTE: Synthetic Minority Over-sampling Technique                                                                           | Nitesh V. Chawla, Kevin W. Bowyer, Lawrence O. Hall & W. Philip Kegelmeyer — Journal of Artificial Intelligence Research, 2002 | Mengusulkan metode Synthetic Minority Over-sampling Technique (SMOTE) untuk menangani data klasifikasi yang tidak seimbang. | Kombinasi oversampling kelas minoritas dan undersampling kelas mayoritas meningkatkan performa klasifikasi.              | Fokus pada metodologi umum, belum diterapkan pada konteks musik atau data audio. | Fokus pada teknik SMOTE<br><br>Digunakan dalam penelitian tesis untuk mengatasi ketidakseimbangan kelas pada dataset musik belajar sehingga menerapkan SMOTE untuk balancing data.                                                          |
| 5  | Predicting song popularity based on Spotify's audio features: insights from Indonesian streaming users (Feature Importance) | Harriman Samuel Saragih — Journal of Management Analytics, 2023                                                                |                                                                                                                             | Menentukan fitur audio Spotify yang memengaruhi popularitas lagu di Indonesia menggunakan machine learning.              | Fokus pada prediksi popularitas, bukan konteks edukatif atau psikologis.         | Menjadi rujukan metodologi penggunaan <i>Feature Importance Random Forest</i> dalam klasifikasi musik belajar.                                                                                                                              |

| No | Judul                                                                                                                                              | Peneliti, Media Publikasi, dan Tahun                                                                             | Tujuan Penelitian                                                                                                                                                        | Kesimpulan                                                                                                                                                                                                                            | Saran atau Kelemahan                                                                                                                                                                                                                   | Perbandingan                                                                                                                                                        |
|----|----------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------|
|    |                                                                                                                                                    |                                                                                                                  |                                                                                                                                                                          |                                                                                                                                                                                                                                       |                                                                                                                                                                                                                                        | Penggunaan <i>Random Forest</i> dan feature importance sebagai rujukan tesis dan pembanding cara menerapkan                                                         |
| 6  | <i>Feature Importance ranking of random forest-based end-to-end learning algorithm (Feature Importance)</i>                                        | X. Yuan, S. Liu, W. Feng, and G. Dauphin-<br><i>Journal of Remote Sensing</i> , vol. 15, no. 21, pp. 5203, 2023. | Mengoptimasi data Sentinel-2 (S2) level L1C dengan metode feature expansion dan <i>Feature Importance analysis</i> untuk meningkatkan akurasi klasifikasi tanaman        | Penambahan 13 fitur vegetation indices dan seleksi fitur menggunakan Random Forest meningkatkan akurasi klasifikasi dan kemampuan prediksi dengan data input yang tidak lengkap                                                       | Kurangnya data sampel lapangan sehingga hasil tidak dapat diplotkan ke dalam peta berdasarkan koordinat lapangan<br><br>Penelitian masa depan perlu mengumpulkan lebih banyak data sampel lapangan untuk validasi dan pemetaan spasial | <i>Feature Importance</i> digunakan sebagai referensi cara penerapan fitur importance                                                                               |
| 7. | The accuracy of Random Forest performance can be improved by conducting a <i>Feature Selection</i> with a balancing strategy (Features Importance) | M. I. Prasetyowati, N. U. Maulidevi, and K. Surendro, <i>PeerJ Computer Science</i> , vol. 8, pp. e1041, 2022.   | Meningkatkan akurasi dan waktu pemrosesan Random Forest melalui <i>Feature Selection</i> dengan integrasi Information Gain (IG), Fast Fourier Transform (FFT), dan SMOTE | <i>Feature Selection</i> menggunakan IG threshold median-real meningkatkan akurasi Random Forest pada 5 dari 8 dataset (EEG Eye, Cancer, Dermatology, Urban Land Cover, dan Epilepsy) dengan peningkatan akurasi 0.0071 hingga 0.0249 | Akurasi menurun pada dataset Contraceptive Method dan Epilepsy karena jumlah fitur yang digunakan lebih sedikit (5 dari 9 fitur untuk Contraceptive Method; 97 dari 178 fitur untuk Epilepsy)                                          | Implementasi SMOTE digunakan sebagai referensi Implementasi teknik validasi kualitas synthetic data untuk memastikan SMOTE menghasilkan samples yang representative |

| No | Judul                                                                                          | Peneliti,<br>Media<br>Publikasi, dan<br>Tahun                                                                                                      | Tujuan Penelitian                                                                                   | Kesimpulan                                                                                               | Saran atau Kelemahan                                                                    | Perbandingan                                                                                                                                                      |
|----|------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| 8. | Music-Induced Emotion Recognition Based on <i>Feature Reduction</i> Using PCA From EEG Signals | Hamid Khabiri, Mohammad N. Talebi, Mehdi F. Kamran, Shadi Akbari, Farzaneh Zarrin & Fatemeh Mohandesi — Frontiers in Biomedical Technologies, 2024 | Mengenal emosi yang ditimbulkan oleh musik melalui sinyal EEG dengan reduksi fitur menggunakan PCA. | PCA efektif mereduksi fitur EEG, dan algoritma SVM mencapai akurasi tertinggi (89%) untuk emosi bahagia. | Ukuran sampel kecil dan hanya berbasis EEG; belum dikaitkan dengan fitur audio digital. | Fokus pada cara menghapus fitur menggunakan fitur reduction<br><br>Memberikan dasar teknis untuk reduksi fitur yang akan dihapus dalam klasifikasi musik belajar. |
|    |                                                                                                |                                                                                                                                                    |                                                                                                     |                                                                                                          |                                                                                         |                                                                                                                                                                   |

## 2.3. Landasan Teori

### 2.3.1 Data dan Persiapan Data

#### 2.3.1.1 Dataset

Dataset dapat dipahami sebagai kumpulan data terstruktur yang memiliki peran fundamental dalam penelitian machine learning, baik sebagai sumber daya pelatihan maupun sebagai alat koordinasi ilmiah. Dataset dalam konteks deteksi objek merupakan kumpulan gambar yang berjumlah puluhan hingga ribuan, di mana setiap gambar mengandung objek yang diberi label agar sistem dapat mempelajari dan mengenali objek tersebut [27]. Dataset membentuk tulang punggung penelitian machine learning yang terintegrasi secara mendalam dalam praktik kerja peneliti, berfungsi sebagai sumber daya untuk melatih dan menguji model [28]. Dataset benchmark menyediakan titik perbandingan yang stabil dan mengkoordinasikan para ilmuwan di sekitar masalah penelitian bersama, dimana peningkatan performa pada benchmark dianggap sebagai sinyal kunci untuk kemajuan kolektif. Dataset mencontohkan tugas machine learning melalui kumpulan pasangan input dan output, dan ketika dilembagakan sebagai benchmark, secara implisit mendukung data tersebut sebagai abstraksi bermakna dari suatu tugas atau domain masalah. Kualitas dan karakteristik dataset sangat mempengaruhi performa akurasi sistem deteksi objek, dimana data yang tidak optimal dapat menyebabkan sistem memberikan hasil yang tidak sesuai atau gagal mendeteksi objek. Sebuah dataset pada dasarnya merupakan kumpulan data yang terstruktur dan

digunakan sebagai bahan dasar dalam proses pembelajaran mesin untuk melatih serta menguji model.

#### **2.3.1.2 Preprocessing Data**

Prapemrosesan data merupakan langkah krusial dalam penambangan data yang melibatkan berbagai teknik untuk membersihkan dan mempersiapkan data mentah sebelum digunakan dalam algoritma pembelajaran mesin. Langkah ini mencakup penanganan data yang tidak sempurna seperti nilai yang hilang dan noise, transformasi data, reduksi dimensionalitas melalui pemilihan fitur dan reduksi instans, serta diskritisasi untuk menyesuaikan data dengan kebutuhan analisis[29]. Dalam konteks big data, prapemrosesan data menjadi lebih kompleks karena ukuran data yang besar membutuhkan pendekatan yang terdistribusi dan efisien, dan melibatkan proses integrasi data, pelabelan, transformasi, pembersihan, dan pemilihan fitur untuk meningkatkan kualitas dan akurasi model prediktif [30].

#### **2.3.1.3 Data Split**

Pembagian data merupakan langkah penting dalam pengembangan model pembelajaran mesin untuk memastikan kemampuan generalisasi. Dataset harus dipisahkan menjadi dua bagian utama, yaitu data pelatihan dan data pengujian. Tujuan pembagian ini adalah agar model tidak hanya menghafal data pelatihan (overfitting), tetapi juga mampu melakukan prediksi dengan baik pada data baru. Metode untuk membagi dataset dapat dilakukan dengan cara pemisahan data yang sederhana (misalnya, 70%

digunakan untuk pelatihan dan 30% untuk pengujian) atau dengan menggunakan teknik yang lebih kompleks untuk memastikan validasi model yang lebih kuat[31]. Pemilihan teknik pemisahan data yang tepat memiliki dampak besar pada akurasi, bias, dan variansi model, sehingga harus disesuaikan dengan sifat dataset dan tujuan penelitian [32].

### **2.3.2 Balancing Data**

Penyeimbangan data adalah teknik praproses yang bertujuan untuk mengatasi masalah ketidakseimbangan distribusi kelas dalam suatu dataset dengan menyesuaikan proporsi sampel antara kelas yang lebih besar dan lebih kecil untuk meningkatkan kinerja model pembelajaran mesin [33].

#### **2.3.2.1 SMOTE (Synthetic Minority Oversampling Technique)**

Teknik Oversampling Minoritas Sintetis (SMOTE) adalah metode augmentasi pengambilan sampel yang dirancang untuk mengatasi masalah ketidakseimbangan kelas dengan menciptakan sampel sintetis baru untuk kelas minoritas melalui proses interpolasi linier antara sampel dari kelas tersebut dan K tetangga terdekatnya[13]. Algoritma SMOTE dimulai dengan memilih satu sampel secara acak dari kelas minoritas, kemudian menemukan K tetangga terdekatnya, dan menghasilkan sampel sintetis baru di sepanjang garis yang menghubungkan sampel yang dipilih dengan tetangga terdekat menggunakan rumus  $z = (1 - w) x_0 + wx_1$ , di mana  $w$  adalah bilangan acak terdistribusi seragam dalam rentang (0, 1) [34]. Meskipun SMOTE telah terbukti meningkatkan efektivitas klasifikasi pada data yang tidak seimbang, metode ini memiliki beberapa keterbatasan,

seperti potensi menciptakan noise dengan meningkatkan sampel yang noise dan menghasilkan sampel sintetis yang dapat tumpang tindih dengan area kelas mayoritas, sehingga menghasilkan berbagai varian SMOTE seperti Borderline-SMOTE yang secara khusus mengambil sampel minoritas yang terlalu banyak (oversampling) yang berada di dekat batas keputusan[35], [36].

### 2.3.3 Normalisasi Data

Normalisasi data merupakan langkah awal yang penting dalam penambangan data untuk menyamakan rentang nilai setiap atribut tanpa kehilangan informasi asli, sehingga meningkatkan akurasi [37]. Perbedaan signifikan dalam rentang nilai antar atribut dapat menyebabkan algoritma klasifikasi yang kurang optimal [38],[39]. Metode normalisasi yang umum diterapkan meliputi Normalisasi MinMax, Normalisasi Z-Score, dan Normalisasi MaxAbs [40]. Normalisasi MinMax mengubah data menjadi rentang antara 0 dan 1 menggunakan rumus  $(x_i - \min(x)) / (\max(x) - \min(x))$ . Proses ini paling baik diterapkan pada data tanpa outlier yang signifikan dan pada algoritma yang sensitif terhadap skala seperti K-NN dan jaringan saraf tiruan. Sebaliknya, Normalisasi Z-Score lebih direkomendasikan untuk data yang mengandung outlier karena metode ini memanfaatkan rata-rata dan deviasi standar dalam perhitungannya [39]

### 2.3.4 Dimensionality Reduction

*Dimensionality Reduction* adalah konversi data dari ruang berukuran besar menjadi ruang berukuran lebih rendah sedemikian sehingga representasi berukuran rendah mempertahankan beberapa atribut yang penting dari data asli. Setiap dataset

melalui data preprocessing yang dinamakan teknik reduksi. Langkah reduksi data bertujuan untuk menghasilkan data yang lebih efisien dan siap diolah dengan baik [41]. Dibawah ini Teknik Reduksi dalam diagram agar pembaca lebih mudah memahami dengan teknik reduksi.



Gambar 2.1 *Classification of Dimensionally Reduction*

Terlihat pada teknik pada diagram diatas berikut adalah beberapa tekniknya:

#### 1. Fitur Selection:

- a. Filter Method: seleksi fitur berdasarkan statistik data / Metode seleksi fitur *filter* menggunakan perhitungan statistik untuk memberi nilai pada setiap fitur [42]. Fitur kemudian dipilih atau dibuang berdasarkan peringkat nilai tersebut. Metode ini menilai setiap fitur secara terpisah, baik secara mandiri maupun terhadap variabel target dengan beberapa teknik seperti:

##### i. Korelasi

ii. Chi-square

iii. Information

b. Wrapper Method: seleksi fitur berdasarkan performa model atau Metode seleksi fitur *wrapper* memandang pemilihan fitur sebagai proses pencarian, di mana berbagai kombinasi fitur diuji dan dibandingkan. Setiap kombinasi dievaluasi menggunakan model prediksi, dan kualitasnya dinilai berdasarkan kinerja model, seperti akurasi. beberapa teknik seperti:

i. Forward Selection

ii. Backward Elimintaion

iii. Stepwise Selection

iv. Recursive Feature Elimination

c. Embedded Method: seleksi fitur secara langsung atau Metode seleksi fitur *embedded* menentukan fitur yang paling berpengaruh terhadap akurasi model secara langsung saat proses pelatihan model berlangsung. Dengan cara ini, pemilihan fitur menjadi bagian dari proses pembentukan model itu sendiri. dengan menerapkan algoritma model beberapa teknik seperti **Random Forest (feature importance)**, Lasso /L1 Regulation, *Decision Tree*, dan XGBOOST.

Namun pada penelitian ini yang akan dibahas hanya pada Salah satu teknik reduksi fitur adalah Embedded Method, yaitu metode seleksi fitur yang dilakukan secara langsung selama proses

pembentukan model. Pada metode ini digunakan algoritma *Random Forest*, yang memiliki kemampuan untuk menentukan tingkat kepentingan fitur (*feature importance*).

#### ***Embedded Method – Random Forest - Feature Importance***

Reduksi fitur pada penelitian ini dilakukan menggunakan *Embedded Method* dengan algoritma *Random Forest*, di mana proses seleksi fitur dilakukan secara langsung selama pembentukan model[42], [43]. Dataset awal yang terdiri dari seluruh fitur dibagi menjadi data latih dan data uji, **dengan perhitungan feature Importance hanya menggunakan data latih**. *Random Forest* merupakan algoritma pembelajaran yang diawasi, menggunakan prinsip aturan *if-then-else*, dan memiliki tingkat interpretabilitas yang tinggi. *Random Forest* menerapkan **bootstrap sampling** untuk membentuk beberapa subset data latih yang berbeda, sehingga meningkatkan keanekaragaman antar pohon keputusan dan mengurangi risiko *overfitting*. Jumlah pohon yang dibangun ditentukan oleh parameter *n\_estimators* sebenarnya *n\_estimators* adalah pembentuk keanekaragaman antar pohon keputusan.

Pada setiap pohon keputusan, dilakukan **pemilihan fitur secara acak di setiap node**. Pemilihan split terbaik ditentukan berdasarkan **Gini Impurity**, yang dirumuskan sebagai:

$$Gini(t) = 1 - \sum_{k=1}^K p_k^2$$

### Penjelasan Rumus:

|           |                                                                                           |
|-----------|-------------------------------------------------------------------------------------------|
| $Gini(t)$ | : Nilai indeks Gini pada node $t$ , digunakan untuk mengukur tingkat ketidakmurnian data. |
| $T$       | : Node (simpul) pada pohon keputusan yang sedang dievaluasi.                              |
| $K$       | : Jumlah kelas pada variabel target.                                                      |
| $p_k$     | : Proporsi atau probabilitas data yang termasuk ke dalam kelas $k$ - $k$ pada node $t$ .  |
| $p_k^2$   | : Kuadrat proporsi kelas $k$ - $k$ , menunjukkan dominasi kelas tersebut dalam node.      |

di mana  $p_k$  merupakan proporsi kelas  $k$ - $k$  pada node  $t$ .

Selanjutnya dihitung penurunan Gini ( $\Delta Gini$ ) untuk setiap fitur:

$$\Delta Gini = Gini(t) - \left( \frac{N_L}{N} Gini(t_L) + \frac{N_R}{N} Gini(t_R) \right)$$

### Penjelasan Rumus

|                        |                                                                                             |
|------------------------|---------------------------------------------------------------------------------------------|
| $t$                    | : Node saat ini (sebelum di-split)                                                          |
| $t_L, t_R$             | : Node kiri dan kanan (hasil split)                                                         |
| $N$                    | : Total sampel di node $t$ (induk)                                                          |
| $N_L$                  | : Jumlah sampel di node kiri                                                                |
| $N_R$                  | : Jumlah sampel di node kanan                                                               |
| $Gini(t)$              | : Nilai impurity Gini di node induk                                                         |
| $Gini(t_L), Gini(t_R)$ | : Nilai impurity Gini di node kiri & kanan                                                  |
| $\Delta Gini$          | : Besarnya penurunan impurity karena split (semakin besar $\rightarrow$ split semakin baik) |

Nilai penurunan Gini dari setiap fitur kemudian diakumulasi dan dirata-ratakan pada seluruh pohon keputusan untuk menghasilkan **feature importance**, yang dinyatakan sebagai:

$$FI(X_j) = \frac{1}{T} \sum_{i=1}^T \Delta Gini(X_j)$$

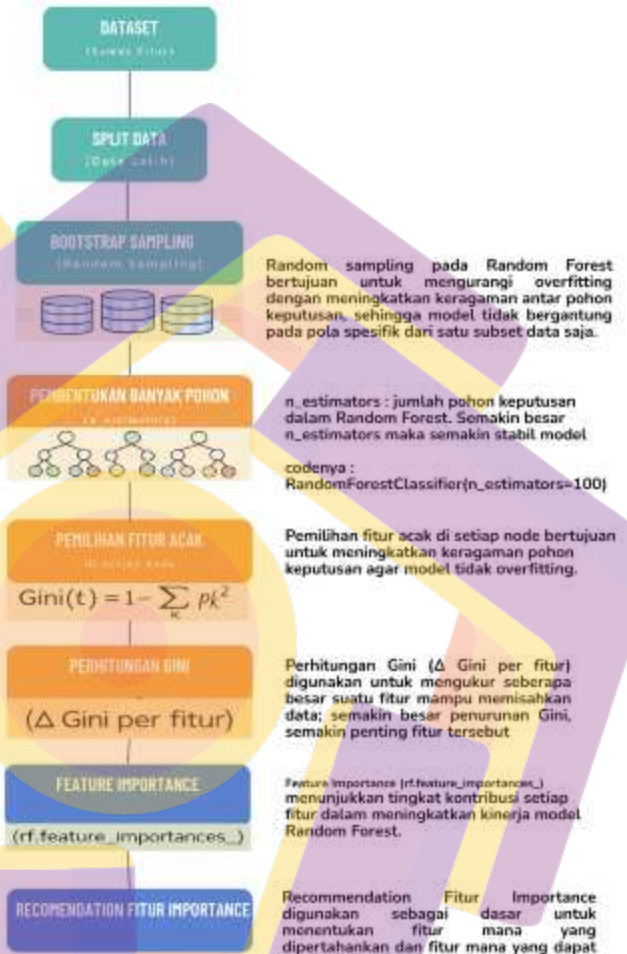
### Penjelasan Rumus

|       |                 |
|-------|-----------------|
| $X_j$ | : Fitur ke- $j$ |
|-------|-----------------|

|                      |                                                                                                  |
|----------------------|--------------------------------------------------------------------------------------------------|
| $T$                  | : Total tree (jika di <i>Random Forest</i> ). Jika hanya 1 decision tree, maka $(T=1)$           |
| $\Delta Gini_i(X_j)$ | : Penurunan impurity Gini yang dihasilkan oleh fitur $(X_j)$ pada split tertentu di tree ke- $i$ |
| $\sum_{i=1}^T$       | : Menjumlahkan semua kontribusi $\Delta Gini$ fitur $(X_j)$ di seluruh tree                      |
| $\frac{1}{T}$        | : Dirata-ratakan terhadap jumlah tree                                                            |

Pelaksanaanya pada googlecollab dapat dilakukan dengan menggunakan *RandomForestClassifier* dari pustaka *scikit-learn*, di mana peringkat pentingnya fitur diperoleh melalui *rf.feature\_importances\_[20]*.

Fitur dengan nilai feature importance yang lebih tinggi dianggap memiliki kontribusi yang lebih besar terhadap kinerja model dan dipertahankan, sedangkan fitur dengan nilai rendah dapat dieliminasi sebagai bagian dari proses reduksi fitur. Dibawah ini detail diagram yang menggambarkan bagaimana fitur importance bisa dijadikan rekomendasi.



Gambar 2.2 *Embedded Method Random Forest using Feature Importance*

### **Wrapper Method – Step Selection**

Pemilihan Bertahap adalah komponen dari *Wrapper Method* yang mempertimbangkan pemilihan fitur sebagai usaha untuk menemukan subset terbaik dengan cara penambahan dan/atau

pengurangan fitur secara iteratif melalui evaluasi model prediksi. Dalam metode *backward stepwise*, fitur dihapus secara bertahap dan pada setiap langkah, model dilatih kembali untuk menentukan pengaruh subset fitur terhadap kinerja prediksi. Kualitas pemisahan dalam model yang berdasarkan Gini lalu hingga menghasilkan fitur yang paling penting. Kemudian dari fitur yang paling penting dilakukan penerapan Evaluasi *wrapper* pada subset fitur dilakukan berdasarkan akurasi model pada data uji:

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

Keputusan eliminasi fitur diukur dari perubahan performa antar iterasi:

$$\Delta Perf = Accuracy_{baru} - Accuracy_{lama}$$

Keputusan seleksi fitur ditentukan berdasarkan aturan:

#### **Penjelasan Rumus**

$$\Delta Perf > \theta \text{ (threshold)}$$

: fitur dipertahankan dalam subset,

$$\Delta Perf \leq \theta$$

: fitur dieliminasi, dan iterasi dapat dihentikan karena tidak memberikan peningkatan performa yang berarti.

## REDUCTION TECHNIQUE: WRAPPER METHOD -STEPWISE SELECTION



Gambar 2.3 Wrapper Method -Stepwise Selection

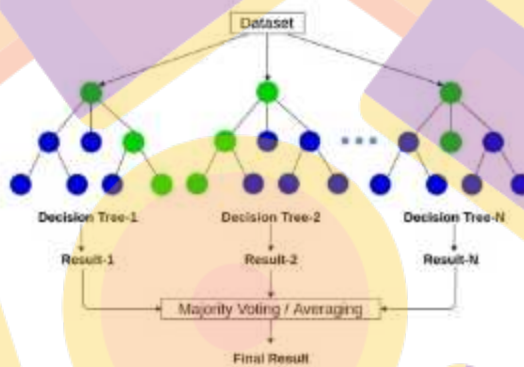
### 2.3.5 Modeling Machine Learning Random Forest

*Random Forest* adalah algoritma pembelajaran ensemble yang efektif, terdiri dari sekumpulan Pohon Keputusan yang bekerja secara kolektif. Meskipun dokumentasi yang tersedia tidak menjelaskan secara spesifik bagaimana data dipecah dalam algoritma *Random Forest* (seperti bootstrapping dan pemilihan subset fitur acak), secara umum, untuk melatih *Random Forest*, set data awal harus dibagi menjadi set pelatihan dan set pengujian. Pemisahan ini penting agar model dapat belajar dari data pelatihan dan dievaluasi secara objektif menggunakan data uji yang tidak terlibat dalam proses pelatihan. Akurasi *Random Forest* dapat ditingkatkan lebih lanjut dengan memilih kriteria pemisahan yang tepat untuk setiap pohon dalam set dan dengan mengoptimalkan metode pemisahan data awal [44].

Dalam praktiknya, untuk melatih model *Random Forest*, set data awal perlu dipecah menjadi set pelatihan dan set pengujian. Pemisahan ini memastikan bahwa model dilatih pada satu bagian data dan dievaluasi secara terpisah pada bagian lainnya, sehingga memberikan pemahaman yang akurat tentang kemampuan

generalisasi model terhadap data yang sebelumnya tidak terlihat [16], [45], [46], [47].

*Random Forest* merupakan salah satu metode dalam Machine Learning yang diterapkan untuk kegiatan seperti klasifikasi dan regresi. Khususnya, *Random Forest* terbukti mampu memberikan akurasi yang sangat baik dalam memprediksi data berdasarkan berbagai atribut [48].



Gambar 2.4 *Random Forest*

#### **Random Forest Dapat Bekerja dengan baik ketika:**

1. Dataset dengan Banyak Fitur

RF sangat ampuh pada data numerik atau kategorikal (seperti dataset Spotify yang Anda gunakan), terutama ketika jumlah fiturnya cukup besar.

2. Hubungan non-linier dan interaksi kompleks antar fitur

RF dapat menangkap pola non-linier dan interaksi antar fitur tanpa memerlukan transformasi khusus.

### 3. Data dengan Noise dan Outlier

Karena didasarkan pada beberapa pohon keputusan, RF cukup kuat terhadap noise dan kurang rentan terhadap overfitting dibandingkan pohon keputusan tunggal.

### 4. Dataset Sedang hingga Besar

RF paling efektif ketika data cukup besar sehingga setiap pohon dapat belajar dari berbagai pola.

### 5. Fitur Tidak Memerlukan Normalisasi

Tidak seperti KNN atau SVM, RF tidak memengaruhi skala data (meskipun normalisasi masih berguna dalam proses umum).

### 6. Masalah Klasifikasi dan Regresi dengan Recall Tinggi

RF seringkali menjadi referensi yang ampuh karena stabil dan memiliki tingkat akurasi yang tinggi.

### 7. Kapan Interpretasi Fitur Diperlukan

RF memberikan informasi tentang pentingnya fitur, yang sangat berguna untuk penelitian pengurangan dimensi, seperti tesis Anda.

## Kelebihan Random Forest

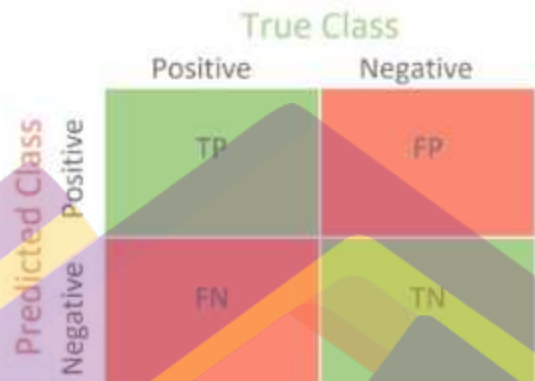
Random Forest memiliki kelebihan dalam menangani fitur musik yang berdimensi tinggi dan saling berkorelasi karena menggunakan pendekatan ensemble dengan pemilihan subset fitur secara acak pada setiap pohon, sehingga mampu mengurangi dampak korelasi antar fitur dan menangkap hubungan non-linear yang kompleks pada data audio. Selain itu, Random Forest bersifat robust terhadap noise dan menyediakan informasi feature importance yang membantu

mengidentifikasi fitur musik yang paling berkontribusi dalam proses klasifikasi. Namun, dalam konteks klasifikasi musik untuk belajar, Random Forest memiliki keterbatasan berupa interpretabilitas model yang rendah secara keseluruhan, kebutuhan komputasi yang relatif tinggi, serta potensi ketidakstabilan nilai feature importance ketika fitur sangat berkorelasi atau data berukuran kecil.

### 2.3.6 Evaluasi Model

#### 2.3.6.1 Confusion Matrix

Pengujian dilakukan untuk mendapatkan nilai akurasi berjalannya sistem deteksi penggunaan masker. Pengujian sistem dilakukan dengan bantuan *Confussion Matrix* menggunakan metrik evaluasi seperti akurasi, precision, recall, dan F1-score[49]. Evaluasi ini bertujuan untuk mengetahui sejauh mana model mampu membedakan secara efektif genre musik untuk belajar. *Confussion Matrix* adalah matriks  $N \times N$  yang digunakan untuk mengevaluasi kinerja model klasifikasi, dimana  $N$  adalah jumlah kelas target. Dengan memvisualisasikan matriks konfusi, keakuratan model dapat ditentukan dengan mengamati nilai diagonal untuk mengukur jumlah klasifikasi yang benar [50].



The diagram shows a 2x2 confusion matrix. The columns are labeled 'True Class' with 'Positive' and 'Negative'. The rows are labeled 'Predicted Class' with 'Positive' and 'Negative'. The cells are: Top-Left (TP, green), Top-Right (FP, red), Bottom-Left (FN, red), and Bottom-Right (TN, green). Diagonal lines from the top-left to bottom-right and bottom-left to top-right are drawn across the matrix.

|                          | True Class |          |
|--------------------------|------------|----------|
|                          | Positive   | Negative |
| Predicted Class Positive | TP         | FP       |
| Predicted Class Negative | FN         | TN       |

Gambar 2.5 *Predicted Class*

Matriks konfusi adalah matriks persegi yang kolom-kolomnya mewakili nilai sebenarnya dan baris-barisnya mewakili nilai prediksi model, begitu pula sebaliknya.

- TP - True Positif: nilai sebenarnya positif dan model memprediksi nilai positif
- FP - False Positif: prediksi Anda positif dan hasilnya salah. (Juga dikenal sebagai kesalahan tipe 1)
- FN - False Negatif: Prediksi Anda negatif dan hasilnya juga salah. (juga dikenal sebagai kesalahan tipe 2)
- TN - True Negatif: Nilai sebenarnya adalah negatif dan model memprediksi nilai negatif.

#### 2.3.6.2 Akurasi

Akurasi mencerminkan kedekatan hasil analisis dengan nilai sebenarnya, yang menggambarkan ketepatan data serta berkaitan dengan kesalahan sistematis atau bias.

Akurasi =

$$\frac{(TP+TN)}{(TP+FP+FN+TN)} \times 100\%$$

- TP - True Positif: nilai sebenarnya positif dan model memprediksi nilai positif
- FP - False Positif: prediksi Anda positif dan hasilnya salah. (Juga dikenal sebagai kesalahan tipe 1)
- FN - False Negatif: Prediksi Anda negatif dan hasilnya juga salah. (juga dikenal sebagai kesalahan tipe 2)
- TN - True Negatif: Nilai sebenarnya adalah negatif dan model memprediksi nilai negatif.

### 2.3.6.3 Precision

Presisi merujuk pada variabilitas yang terjadi dari beberapa kali pengukuran atau pengujian, yang mencerminkan akurasi data dan berhubungan dengan kesalahan acak [51].

Precision =

$$\frac{TP}{(TP+FP)} \times 100\%$$

- TP - True Positif: nilai sebenarnya positif dan model memprediksi nilai positif
- FP - False Positif palsu: prediksi Anda positif dan hasilnya salah. (Juga dikenal sebagai kesalahan tipe 1)

- FN - False Negatif: Prediksi Anda negatif dan hasilnya juga salah. (juga dikenal sebagai kesalahan tipe 2)
- TN - True Negatif: Nilai sebenarnya adalah negatif dan model memprediksi nilai negatif.

#### 2.3.6.4 Recall

*Recall* merujuk pada proporsi kasus positif yang sebenarnya yang berhasil diprediksi dengan akurat. Konsep ini mengukur sejauh mana kasus positif yang sesungguhnya dapat dicakup oleh aturan +P (Prediksi Positif). Tujuan fungsional dari metrik ini adalah untuk mencerminkan jumlah kasus relevan yang dapat diidentifikasi oleh aturan +P. Fungsi ini cenderung meremehkan pemulihan informasi, terutama dengan asumsi bahwa terdapat banyak dokumen relevan, bahwa subset yang berhasil ditemukan tidak terlalu signifikan, dan bahwa kita tidak memiliki informasi yang mencukupi mengenai pentingnya dokumen yang tidak dikembalikan. Dalam konteks pembelajaran mesin dan linguistik komputasi, faktor memori sering kali diabaikan atau dimediasi, karena fokus utama terletak pada keyakinan kita terhadap aturan atau pengklasifikasi [52]. Namun, dalam area linguistik komputasi dan terjemahan mesin, keberadaan memori telah terbukti memiliki pengaruh yang signifikan dalam memprediksi keberhasilan penyelarasan kata. Di samping itu, *recall* atau sensitivitas dapat dideskripsikan sebagai rasio antara jumlah prediksi positif yang benar dengan keseluruhan data yang benar positif.

Recall) =

$$\frac{TP}{(TP+FN)} \times 100\%$$

### 2.3.6.5 F1-Score

*F1-Score* adalah metrik evaluasi pembelajaran mesin alternatif yang menilai kemampuan prediktif model dengan mengelompokkan performa berdasarkan kelas bukan performa keseluruhan, seperti yang dilakukan dengan presisi. Skor *F1* menggabungkan dua metrik yang bersaing: skor presisi dan perolehan model, yang menyebabkan penggunaannya secara luas dalam literatur terkini. *F1 Score* merupakan perbandingan rata-rata *presisi* dan *recall* yang dibobotkan [53].

F1 Score =

$$\frac{2 \times \text{Recall} \times \text{Precision}}{\text{Recall} + \text{Precision}}$$

## BAB III

### METODE PENELITIAN

#### 3.1. Jenis, Sifat, dan Pendekatan Penelitian

Penelitian ini merupakan studi eksperimental yang bertujuan untuk meningkatkan presisi algoritma *Random Forest* dalam mengklasifikasikan audio dari Spotify. Kumpulan data /dataset yang dipakai mencakup beberapa atribut audio musik yang mencerminkan ciri khas lagu, seperti kecepatan, panjang, kemampuan menari, dan tingkat instrumen. Dataset yang tidak seimbang ditangani menggunakan metode SMOTE (*Synthetic Minority Over-sampling Technique*) untuk menyeimbangkan distribusi kelas. Hal ini dilakukan dengan menerapkan teknik-teknik penting fitur untuk mengidentifikasi fitur-fitur yang paling signifikan, diikuti dengan reduksi fitur untuk menghilangkan fitur-fitur yang kurang relevan. Dari segi karakteristik, penelitian ini bersifat deskriptif, memberikan gambaran sistematis tentang eksperimen dan membandingkannya dengan penelitian sebelumnya. Pendekatan kuantitatif digunakan karena hasil penelitian bersifat objektif dan dalam bentuk numerik, untuk memastikan validitas dan integritas temuan.

Penelitian ini tidak menerapkan *k-fold cross-validation*. Sebagai gantinya, pengendalian overfitting dilakukan melalui pemisahan data secara tegas antara data latih dan data uji yang bersifat independen. Model *Random Forest* dilatih sepenuhnya menggunakan data latih, sementara evaluasi kinerja dilakukan menggunakan data uji yang tidak terlibat dalam proses pelatihan. Pendekatan ini

bertujuan untuk memastikan bahwa model tidak hanya menghafal pola pada data latih, tetapi mampu melakukan generalisasi terhadap data baru.

### 3.2. Metode Pengumpulan Data

Penelitian ini menggunakan data sekunder yang dikumpulkan melalui metode analisis data sekunder. Data ini merupakan catatan yang telah dikumpulkan sebelumnya yang dapat diakses publik dan memiliki izin untuk digunakan. Sumber data ini biasanya dapat ditemukan di repositori daring seperti Kaggle, GitHub, Drive, dan sejenisnya. Dalam melakukan penelitian kali ini peneliti mengambil dataset yang tersedia secara publik. Tahap praproses dilakukan untuk memastikan kualitas data sebelum digunakan pada proses seleksi fitur dan klasifikasi. Tabel 3.1 menampilkan dataset yang diambil dari source tersebut

Tabel 3.1. Daftar Feature Dataset Spotify Study Music

| No     | Danceability | Energy | ... | Class | Label                        |
|--------|--------------|--------|-----|-------|------------------------------|
| 0      | 0.645        | 0.114  | ... | 2     | <i>Classical Soundtracks</i> |
| 1      | 0.613        | 0.526  | ... | 1     | <i>Pop Tracks</i>            |
| 2      | 0.788        | 0.446  | ... | 1     | <i>Pop Tracks</i>            |
| 4      | 0.717        | 0.193  | ... | 3     | <i>Lo-fi Tracks</i>          |
| 5      | 0.521        | 0.289  | ... | 1     | <i>Pop Tracks</i>            |
| ...    | ...          | ...    | ... | ...   | ...                          |
| 172814 | 0.435        | 0.183  | ... | 3     | <i>Lo-fi Tracks</i>          |
| 172815 | 0.221        | 0.32   | ... | 1     | <i>Pop Tracks</i>            |
| 172816 | 0.419        | 0.0592 | ... | 3     | <i>Lo-fi Tracks</i>          |
| 172817 | 0.479        | 0.0899 | ... | 3     | <i>Lo-fi Tracks</i>          |
| 172818 | 0.205        | 0.468  | ... | 1     | <i>Pop Tracks</i>            |

Tabel 3.1 menjelaskan tabel musik dengan label utama yang digunakan pada Spotify yang diperoleh dari Github. Penelitian ini merujuk pada penelitian yang

berkesinambungan yang dilakukan oleh Rebecca Jane Scarrat dari tahun 2019 hingga 2023. Dalam dataset ini terdapat 172818 data musik, dengan dataset ini penelitian bertujuan untuk mengembangkan model klasifikasi musik untuk belajar yang lebih akurat.

### 3.3. Metode Analisis Data

Agar penelitian dapat mencapai tujuan dan hasil dapat dibuktikan dengan cara empiris, serangkaian langkah analisis data dilakukan dengan sistematis. Analisis difokuskan pada penentuan pentingnya fitur (*feature important*) dan reduksi fitur (*feature recognition*) untuk memilih variabel audio yang paling mempengaruhi kebiasaan mendengarkan musik sebagai alat belajar, serta untuk mengidentifikasi fitur musik yang paling relevan selama proses belajar.

Tahap awal dalam pra-pemrosesan data adalah menangani masalah ke integrasi kelas pada pelatihan data. Untuk mencapai hal itu, diterapkan teknik Synthetic Minority Over-sampling Technique (SMOTE) untuk menyeimbangkan distribusi kelas. Setelah proses penyeimbangan, dilakukan normalisasi data. Langkah ini penting untuk memastikan keseragaman dalam skala data, sehingga meningkatkan stabilitas model pembelajaran dan memperkuat validitas perbandingan antar fitur, meskipun algoritma *Random Forest* cukup tangguh terhadap perbedaan skala data.

Setelah pra-pemrosesan, analisis kepentingan fitur dilakukan dengan menggunakan algoritma *Random Forest* untuk menilai kontribusi relatif setiap fitur audio terhadap variabel target. Berdasarkan hasil analisis ini, prosedur seleksi dan

pengurangan dimensi fitur diterapkan untuk meminimalkan redundansi serta mempertahankan fitur-fitur yang berperan penting dalam proses pembelajaran.

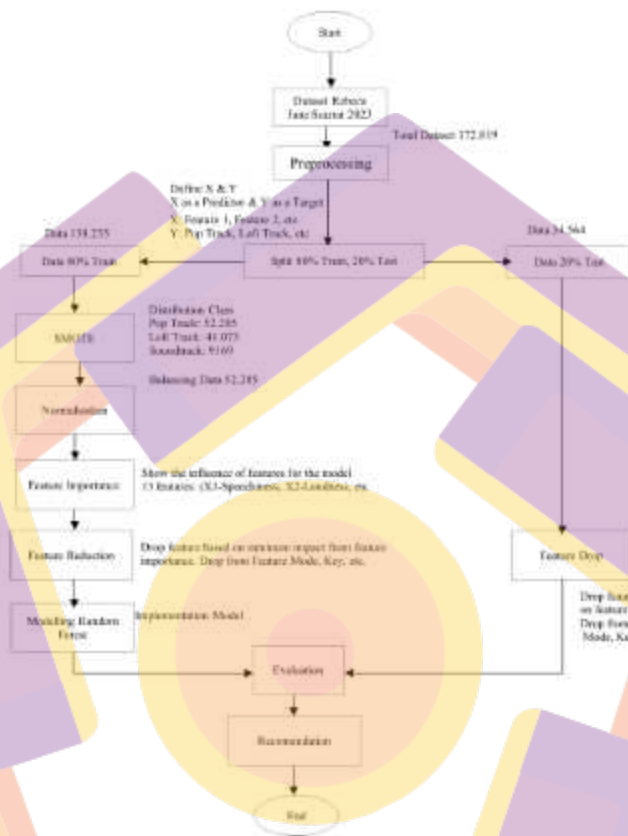
Relevansi fitur musik dinilai dengan Merujuk pada kinerja model klasifikasi yang diperoleh, baik sebelum maupun setelah dimensi dikurangi. Selanjutnya model *Random Forest* diterapkan pada pengujian data untuk kinerja dan stabilitasnya. Pengujian dilakukan dengan berbagai konfigurasi pelatihan untuk mendapatkan kinerja yang optimal. Hasil dari tahap ini dijadikan dasar untuk menilai efektivitas model dalam memanfaatkan representasi fitur musik yang sudah dipilih.

Tahap terakhir meliputi evaluasi model kinerja dengan menggunakan metrik akurasi dan analisis konfusi matriks. Evaluasi dilakukan pada uji data yang tidak terlibat dalam proses pelatihan untuk mengukur kemampuan generalisasi model terhadap data yang baru. Confusion matriks digunakan untuk menganalisis distribusi prediksi yang benar dan salah di setiap kelas, sehingga memberikan gambaran yang komprehensif mengenai kekuatan dan kelemahan model klasifikasi.

### 3.4. Alur Penelitian

Rangkaian proses penelitian yang akan dilakukan dapat dilihat pada Gambar

3.1 *Research FLOW Stages* sebagai berikut:



Gambar 3.1 Research FLOW Stages

Proses penelitian ini melibatkan beberapa langkah utama, dimana proses dimulai dengan pemrosesan data awal untuk memastikannya siap digunakan. Selanjutnya, data *Spotify Music Study /SSM* dipisahkan menjadi dua bagian: data (train\_x) 80% = 138.255 dan data (test\_y) 20% = 34.564 dari total data 172.819. Selanjutnya, penyeimbangan data dilakukan menggunakan metode SMOTE untuk

memastikan distribusi kelas yang seimbang. Dilakukan normalisasi data agar data lebih proporsional dan tidak menimbulkan bias pada model pembelajaran mesin.

Langkah selanjutnya adalah analisis fitur importance untuk memahami pengaruh setiap fitur terhadap model, dilanjutkan dengan pengembangan model menggunakan algoritma *Random Forest* setelah model selesai, validasi silang dilakukan untuk menilai konsistensi dan kualitasnya.

Dan juga reduksi dimensi dilakukan secara bertahap berdasarkan feature importance, dan evaluasi performa dilakukan pada setiap tahap menggunakan data uji. Yaitu langkah terakhir reduction feature dilakukan, yaitu pengujian ulang model dengan menghilangkan fitur-fitur tertentu satu per satu untuk mengamati dampaknya terhadap akurasi. Hal ini menghasilkan model dengan kinerja terbaik dan fitur-fitur yang paling relevan.

#### **3.4.1 Tahapan Pengumpulan Data**

Sebagaimana dijelaskan pada bagian Metode Pengumpulan Data, penelitian ini memanfaatkan data publik yang bersumber dari penelitian sebelumnya. Pendekatan ini dipilih untuk meningkatkan efisiensi waktu penelitian dengan menghilangkan kebutuhan pelabelan data manual yang melibatkan para ahli, sekaligus berfungsi untuk memvalidasi keabsahan hasil yang diperoleh. Tahap awal penelitian adalah Pra-pemrosesan Data yang dilakukan pada data latih, dengan tujuan agar data yang digunakan komprehensif, tidak redundan, dan berdistribusi seimbang. Proses ini diawali dengan mendefinisikan variabel independen (X) dan dependen (Y), dilanjutkan dengan memilah data menjadi data latih (80%) dan data uji (20%). Karena data latih menunjukkan bias kelas, maka diterapkan teknik

SMOTE (Synthetic Minority Over-sampling Technique) untuk menyeimbangkan distribusi data, dilanjutkan dengan normalisasi data untuk memastikan keseragaman skala fitur.

Langkah selanjutnya berfokus pada Pemilihan Fitur, dimulai dengan mengukur tingkat kepentingan fitur menggunakan model *Random Forest* (RF) pada data latih SMOTE. Analisis ini mengidentifikasi fitur audio Spotify yang paling menonjol. Berdasarkan hasil ini, reduksi fitur dilakukan untuk menentukan subset fitur audio terbaik yang memiliki dampak terbesar pada variabel target, terutama ketika fitur dengan tingkat kepentingan rendah dihilangkan. Setelah fitur optimal dipilih, model *Random Forest* kinerjanya dievaluasi menggunakan data uji. Hasil evaluasi disajikan melalui metrik seperti Confusion Matrix, Precision, Recall, dan F1-Score. Data dan kinerja klasifikasi kemudian disajikan dalam visualisasi grafis untuk memudahkan interpretasi dan menunjukkan perbedaan hasil model *Random Forest* dibandingkan dengan penelitian sebelumnya yang menggunakan metode K-Means.

### **3.5. Keterbatasan Metodologi dan Potensi Bias Penelitian**

Studi ini mungkin dipengaruhi oleh bias metodologis yang berasal dari pilihan algoritma klasifikasi. Algoritma yang digunakan terbatas dengan menerapkan pada *Random Forest* sehingga hasil yang diperoleh mungkin dipengaruhi oleh karakteristik spesifik dari algoritma-algoritma tersebut. Lebih lanjut, penggunaan akurasi sebagai tolok ukur utama untuk menilai kinerja model memperkenalkan potensi bias ke dalam evaluasi, terutama jika distribusi kelas dalam dataset tidak merata.

## BAB IV

### HASIL PENELITIAN DAN PEMBAHASAN

#### 4.1. Pengumpulan Data

##### 4.1.1 Data Set

Data yang digunakan dalam studi ini berasal dari data publik yang dikumpulkan oleh Rebecca Jane Scarratt (2023). Data yang digunakan diperoleh dari github

[https://github.com/RebeccaJaneScarratt/Study-Sleep-](https://github.com/RebeccaJaneScarratt/Study-Sleep-Analyses/blob/main/SSM_all_withClass.csv)

[Analyses/blob/main/SSM\\_all\\_withClass.csv](https://github.com/RebeccaJaneScarratt/Study-Sleep-Analyses/blob/main/SSM_all_withClass.csv)). Studi ini berfokus pada analisis elemen musik dalam mendukung proses pembelajaran, dengan mengklasifikasikan data musik untuk aktivitas pembelajaran ke dalam dua kategori: musik untuk tidur dan musik untuk belajar. Analisis dilakukan terhadap lima belas fitur utama Spotify yang berpotensi memengaruhi efektivitas pembelajaran, yaitu:

1. Danceability mengukur tingkat ketergerakan atau keasikan lagu.
2. Energy menilai intensitas emosional lagu,
3. Loudness untuk mengidentifikasi tingkat kebisingan audio,
4. Speechiness menganalisis kejelasan intonasi vocal,
5. Acousticness sebagai indikator karakteristik akustik,
6. Instrumentalness, membedakan musik instrumental dan vocal,
7. Liveness, mengukur kesan kehidupan dalam audio,
8. Valence, mengevaluasi pesan emosional lagu,
9. Tempo, pengukur kecepatan musik,
10. Key, merepresentasikan kunci musik,

11. Mode
12. Duration (ms), durasi dalam milidetik,
13. Popularity, mengindikasikan tingkat popularitas.
14. Class
15. Class\_label

Klasifikasi utama dalam studi ini dibagi menjadi tiga kategori sasaran, yaitu *Pop Tracks*, *Classical Soundtracks*, dan *Lo-fi Tracks*. Dataset yang digunakan terdiri dari 15 kolom dari total data mencapai 172.819 entri.

#### 4.1.2 Preprocessing Data

Pada tahap ini dilakukan untuk memastikan kualitas dan kesiapan dataset sebelum implementasi algoritma pembelajaran mesin [54]. Proses ini mencakup beberapa prosedur sistematis, yaitu:

- 1 Tahap pertama membuat link ke googledrive agar data bisa dibaca
- 2 Kemudian menghapus informasi yang kosong atau redundant [55].
- 3 Dengan menghapus kolom pada dataset yang tidak digunakan untuk penelitian diantaranya kolom ['Unnamed: 0', 'X.1', 'X', 'playlistID', 'playlistName', 'nTracks', 'type', 'owner', 'description', 'url', 'nFoll\_x', 'nFoll\_y', 'TrackName', 'TrackID', 'SampleURL', 'ReleaseYear', 'Genres']
- 4 Menampilkan dataset setelah kolom dihapus
- 5 Penanganan pada kolom class yang memiliki nilai 4.0 Ambient Track karena yang akan diteliti hanya pada class 1.0 - *Pop Tracks*, 2.0 - *Classical Soundtracks Tracks*, dan 3.0 - *Lo-fi Tracks*
- 6 Menampilkan sebaran dataset setelah menghapus nilai 4.0 Ambient Track

- 7 Penanganan missing values, duplikasi dan data kosong, atau data tidak lengkap
- 8 Kemudian mengubah tipe data harus menggunakan string, karena sebelumnya menampilkan data class label yang berisi bukan data text yang berisi huruf seperti *Pop Tracks*, *Classical Soundtracks* dan *Lo-fi Tracks*.
- 9 Lalu pada kolom class dan class\_label jika ada yang kosong /missing value lakukan hapus data
- 10 Lalu melihat lagi dataset apakah sudah menampilkan anpa data yang hanya string .
- 11 Baru melakukan split data.

#### 4.1.3 Data Split

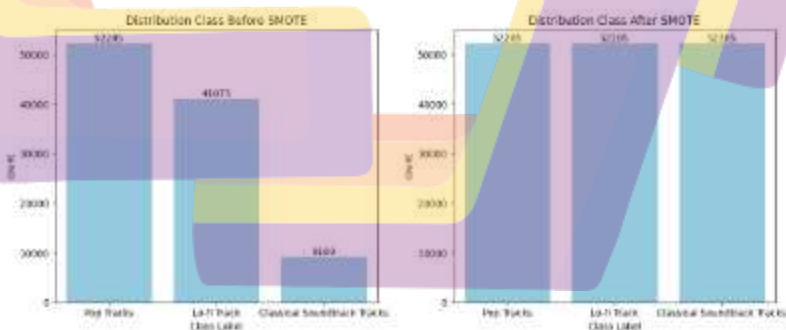
Tahapan pembagian data dalam proses ini dilakukan untuk mempersiapkan data sebelum pelatihan model sehingga dapat digunakan secara optimal[32], [56] . berikut langkah- langkah split data:

1. Mendefinisikan X dan Y, X as a predictor and Y as a target.
2. Pilih kolom X (*danceability, energy, loudness, speechiness, acousticness, instrumentalness, liveness, valence, tempo, key, mode, duration\_ms, Popularity*) dan Y (*class*). Berisi 13 fitur dan 1 target yang bersumber dari dataset SSM.
3. Ratio pembagian X dan Y, Pembagian data memisahkan data menjadi dataset awal dibagi menjadi 80% data train dan 20% data tes. pembagian data ini menggunakan library py yaitu *Scikit-learn* yang merupakan library

yang sangat populer untuk *machine learning* dan analisis data untuk memastikan evaluasi model yang objektif dan representatif

#### 4.2. Balancing Data

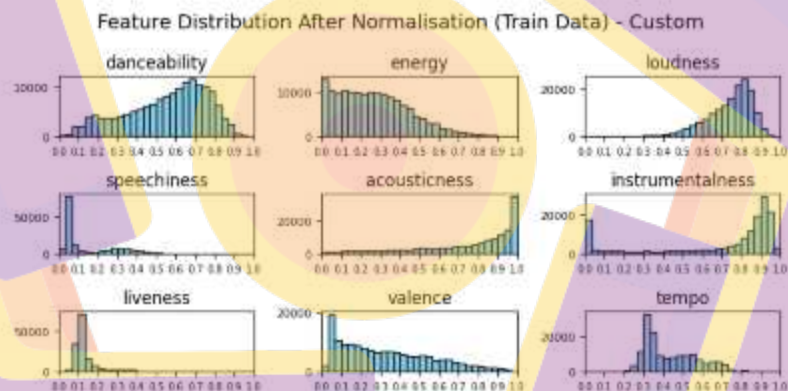
Setelah data di-*split*, sistem mengimplementasikan SMOTE dengan parameter  $k\_neighbors = 5$  untuk menangani ketidakseimbangan kelas dalam dataset dengan menciptakan data sintetis berdasarkan rata-rata dari lima tetangga terdekat, sehingga diperoleh keseimbangan kelas yang optimal tanpa menimbulkan risiko *overfitting* [23], [24]. Seperti yang ditunjukkan pada Gambar 3, data sebelum penerapan SMOTE menunjukkan adanya ketidakseimbangan jumlah *Pop Tracks* 52.285, *Lo-fi Tracks* 41.073, dan *Classical Soundtracks* 9.169 data kelas. Namun setelah SMOTE diterapkan, ketiga kelas kini menjadi seimbang dengan total 52.285 entri untuk masing-masing kelas, yang menunjukkan bahwa metode ini berhasil dalam menyeimbangkan distribusi data serta meningkatkan kualitas klasifikasi yang dilakukan terlihat pada Gambar 4.1 dibawah ini.



Gambar 4.1 *Distribution Class Before and After SMOTE*

### 4.3. Normalisasi Data

Normalisasi sebaiknya diterapkan pada setiap fitur secara terpisah, karena setiap fitur memiliki ukuran dan satuan yang berbeda. Hal ini membuat hasil normalisasi menjadi lebih adil dan tidak menimbulkan bias dalam model mesin pembelajaran. Proses normalisasi data biasanya dilakukan untuk masing-masing fitur, bukan untuk keseluruhan kumpulan data menggunakan min-max normalization. Nilai minimum dan maksimum dihitung untuk setiap fitur secara terpisah [57],[58]. Dibawah ini menunjukkan normalisasi data terlihat pada Gambar 4.2. Data Normalization.



Gambar 4.2 Data Normalization

Tujuan grafik ini agar tiap fitur ada pada skala yang seimbang yaitu skala 0-1, skala angka memengaruhi cara model menghitung dan belajar. Dari fenomena yang tampak, fitur danceability memiliki sebaran yang merata di tengah dengan puncak sekitar 0.6, yang menunjukkan bahwa kebanyakan lagu memiliki tingkat ketergerakan yang sedang. Fitur energy dan loudness cenderung berkumpul di nilai

tinggi 0.8–1.0, yang mencerminkan bahwa sebagian besar lagu memiliki intensitas dan volume suara yang tinggi. Sebaliknya, speechiness dan liveness menunjukkan puncak tinggi pada nilai rendah sekitar 0.0–0.2, yang berarti kebanyakan lagu tidak banyak memiliki unsur vokal percakapan dan jarang direkam langsung. Fitur acousticness dan instrumentalness menunjukkan peningkatan menuju nilai 1.0, menandakan banyaknya lagu dengan karakter akustik dan musik instrumental yang kuat. Di sisi lain, valence cenderung menurun dari 0.0 ke 1.0, menunjukkan lebih banyak lagu dengan suasana netral hingga sedih, dan tempo menunjukkan sebaran yang cukup bervariasi di tengah, menandakan adanya variasi dalam kecepatan lagu. Secara keseluruhan, grafik ini menunjukkan bahwa proses normalisasi berhasil membuat semua fitur memiliki skala yang seragam tanpa mengubah bentuk distribusi aslinya, sehingga data sudah siap untuk pelatihan model klasifikasi.

#### 4.4 Analisis Performa Model Sebelum dan Sesudah Reduksi Dimensi

Analisis sebelum dan sesudah reduksi dilakukan untuk mengetahui perubahan kinerja model yang terjadi, yang akan dijabarkan seperti dibawah ini, diantaranya:

1. Sebelum pengurangan dimensi, bias data berasal dari fitur asli yang berdiri sendiri. Setelah reduksi, fitur-fitur tersebut diubah menjadi komponen baru yang dihasilkan dari kombinasi banyak fitur yang merepresentasikan pola utama data yang dapat dilihat pada data normalisasi.
2. Sebelum pengurangan dimensi, model sering kali bias terhadap fitur yang dominan atau berlebihan. Setelah pengurangan dimensi, bias bergeser ke pola dan informasi data global.

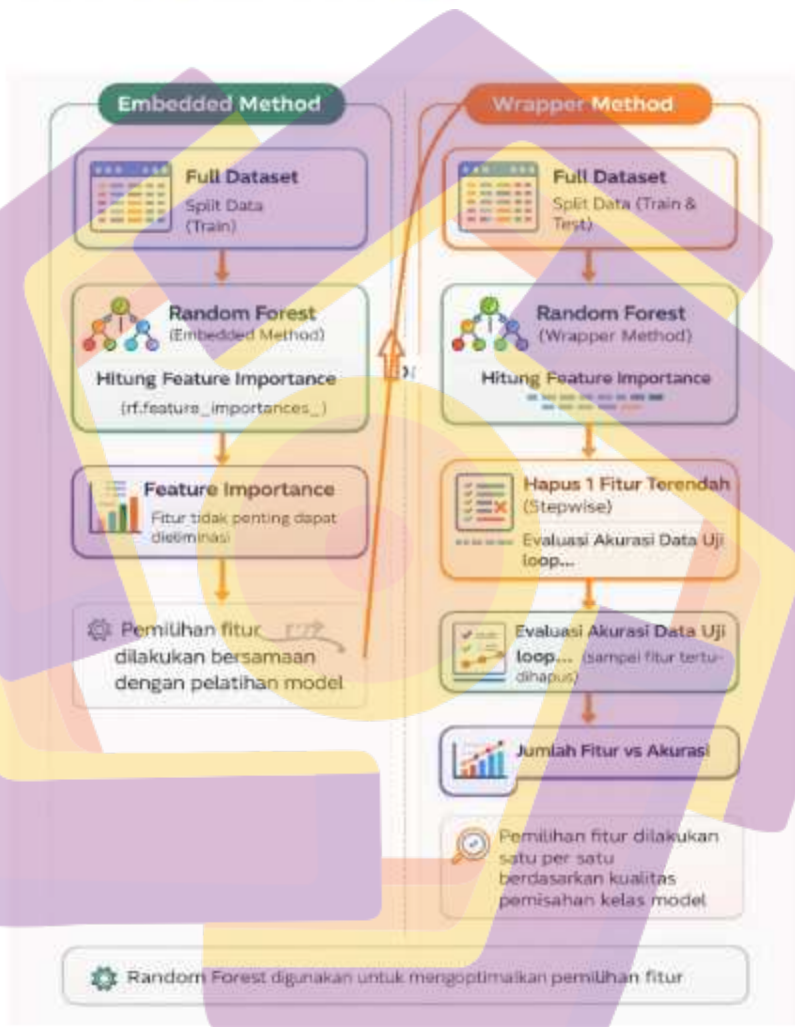
3. Total jumlah data sebelum di lakukan reduksi dimensi dan setelah melakukan SMOTE 52. 285 sample dari ketiga kelas utama.
4. Setelah itu baru menerapkan reduksi dimensi berupa feature importance untuk melihat seberapa penting sebuah fitur.
5. Feature importance dilakukan untuk memilih fitur-fitur yang paling relevan dan informatif, sehingga model menjadi lebih efisien, mudah diinterpretasikan, dan tetap mempertahankan performa klasifikasi.

#### **4.5 Dimensional Reduction Data**

Tahap akhir dalam pra-pengolahan data adalah reduksi data. Meskipun ilmuwan data umumnya bekerja dengan kumpulan data berukuran besar, jumlah data yang terlalu banyak justru dapat menimbulkan permasalahan, seperti meningkatnya kompleksitas analisis dan beban komputasi. Dalam konteks predictive analytics, data umumnya direpresentasikan dalam bentuk tabel dua dimensi yang terdiri atas variabel (kolom) dan kasus atau record (baris). Oleh karena itu, data berukuran besar perlu dianalisis dan disederhanakan melalui proses reduksi data agar diperoleh dataset yang lebih ringkas, relevan, dan efisien tanpa menghilangkan informasi penting[59].

Teknik dalam reduksi data yang digunakan dalam penelitian ini yang pertama adalah teknik *embedded method* untuk menyeleksi menggunakan RF - *Feature Importance* untuk mengetahui fitur mana paling penting dan teknik *wrapper method -Step Wise Selection* untuk mengetahui bahwa penghapusan fitur berdasarkan *feature importance* bisa menekan akurasi sehingga dapat memberikan

rekomendasi[42], [43]. Diagram penjelasan dapat dilihat pada Gambar 4.3 *Flow Dimensionality Reduction Data* di bawah ini.



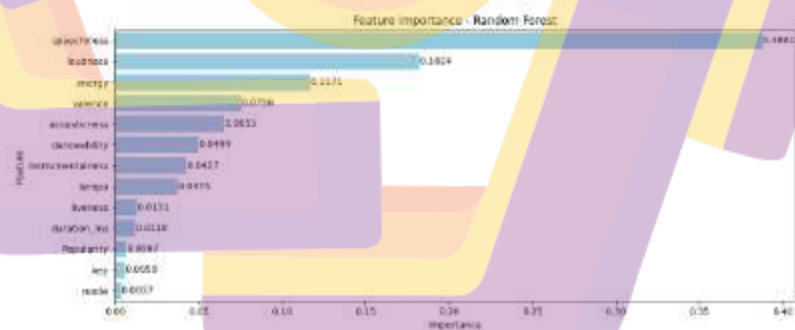
Gambar 4.3 *Flow Dimensionality Reduction Data*

#### 4.6 Feature Importance

Data yang digunakan dalam melakukan pemilihan fitur importance yaitu data train hasil normalisasi data. Pada proses ini *Random Forest* digunakan untuk melatih model, bukan untuk melakukan prediksi namun untuk mengukur seberapa pentingnya tiap fitur. Setelah model selesai dibuat dan diuji, sistem dapat memberikan rekomendasi musik kepada pengguna.

#### 4.7 Hasil Feature Importance

Berdasarkan hasil data, sistem akan menyarankan jenis musik dengan fitur yang paling sesuai untuk mendukung aktivitas belajar. Rekomendasi ini membantu pengguna memilih musik yang tepat agar proses belajar menjadi lebih efektif dan fokus. Hasil feature importance menunjukkan bahwa beberapa fitur memiliki kontribusi yang relatif rendah terhadap proses klasifikasi. Hasil analisis *Feature Importance* dari model *Random Forest* dapat dilihat pada Gambar 4.4 *Visualisation of Feature Importance* di bawah ini.



Gambar 4.4 Visualisasi Hasil Feature Importance

Dapat disimpulkan dari visualisasi *feature importance* bahwa fitur audio Spotify yang paling berperan dalam mendukung proses pembelajaran meliputi

1. **Speechiness** (tingkat dominasi unsur ucapan/verbal dalam audio),
2. **Loudness** (tingkat kekerasan suara),
3. **Energy** (intensitas energi audio),
4. **Valence** (nuansa emosional),
5. **Acousticness** (tingkat unsur akustik/alami),
6. **Danceability** (keteraturan ritme audio),
7. **Instrument** (karakter instrumental/non-verbal),
8. dan **Tempo** (kecepatan ritme musik/BPM), sedangkan fitur lain seperti **Liveness**, **Duration\_ms**, **Popularity**, **Key**, dan **Mode** pada tahap ini masih menunjukkan akurasi yang lebih rendah dalam pemodelan awal, dan pola temuan ini sejalan dengan dan diperkuat oleh penelitian Franziska Goltz dan Makiko Sadakata (2021) yang berjudul “**Do You Listen To Music While Studying? A Portrait Of How People Use Music To Optimize Their Cognitive Performance**” yang menegaskan bahwa memilih musik berdasarkan karakter audio seperti (**Genre lofi, instrumental, plano lembut, ambient, classical slow, dan background music yang minimalis.**) mampu meminimalkan gangguan lebih efektif dalam mengoptimalkan kinerja kognitif selama proses pembelajaran[60].

#### 4.8 Stepwise Feature Reduction.

Setelah model diuji langkah selanjutnya adalah melakukan pengurangan fitur untuk mengetahui seberapa besar kontribusi tiap fitur terhadap kinerja model. Dapat dijabarkan alasan menggunakan stepwise feature reduction diantaranya:

1. Random Forest mengevaluasi fitur secara independen, sementara kinerja optimal sering muncul dari interaksi antar fitur, yang hanya terungkap ketika fitur diuji dan dihilangkan secara bertahap.
2. Penghapusan dilakukan pada fitur yang paling rendah yaitu fitur mode.
3. Akurasi awal bukanlah titik akhir, melainkan garis dasar; pengurangan bertahap dilakukan untuk menentukan apakah penghapusan fitur tertentu benar-benar meningkatkan atau menurunkan kinerja, bukan hanya untuk menunjukkan pentingnya fitur tersebut.
4. Oleh karena itu, pengurangan bertahap adalah proses validasi empiris lebih lanjut untuk menemukan kombinasi fitur yang paling efisien dan stabil, bukan hanya yang memiliki skor tinggi.
5. Metode ini berguna untuk menemukan kombinasi fitur yang paling efektif yang tetap memberikan akurasi tinggi dengan jumlah fitur yang lebih sedikit, sehingga membuat model lebih efisien dan sederhana.

*Stepwise Feature Reduction* berarti proses menyederhanakan banyak fitur musik menjadi beberapa komponen utama agar analisis pada musik lebih efisien dan akurat [61],[62], [63]. Berikut adalah proses stepwise feature reduction

# 1. Melakukan Import library

```
from sklearn.ensemble import RandomForestClassifier
from sklearn.metrics import accuracy_score
import pandas as pd
import matplotlib.pyplot as plt
```

# 2. Menghapus 8 fitur based on feature importance yaitu yang paling bawah ke atas

```
('mode', 'Popularity', 'duration_ms', 'danceability', 'key', 'liveness', 'instrumentalness', 'acousticness')
```

# 3. Menyiapkan variabel penyimpanan

```
Results, dropped = [], []
```

# 4. Looping fitur yang dihapus, 8 fitur dihapus satu persatu setelah 1 proses penghapusan selesai lalu akan mulai dari fitur pertama lagi untuk proses penghapusan

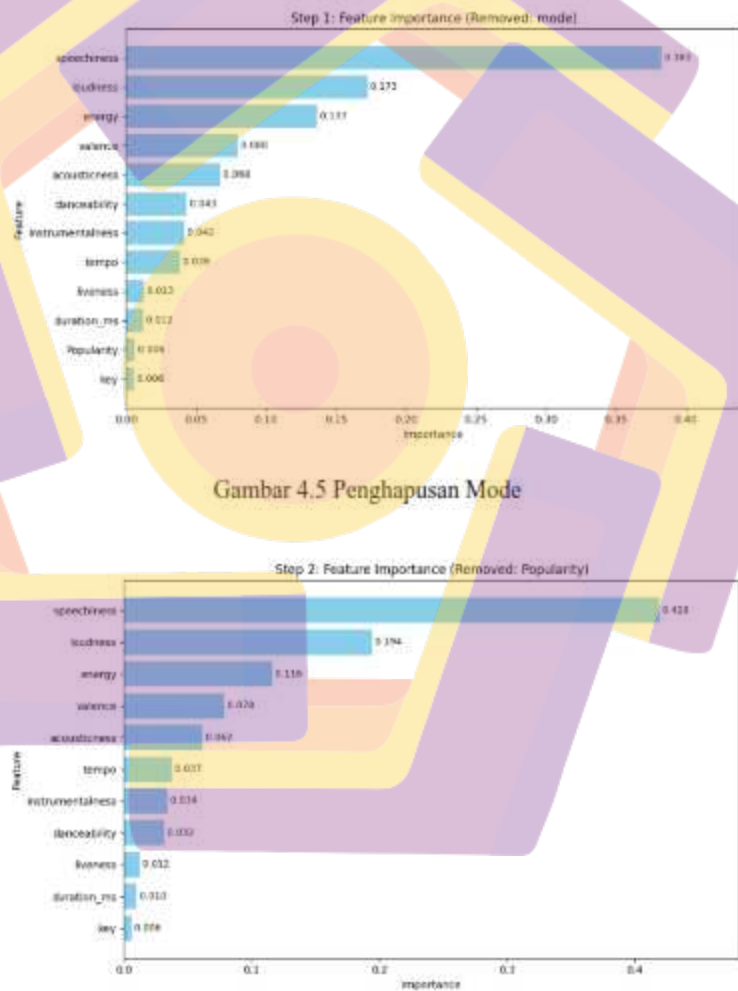
```
for step, f in enumerate(drop_order, start=1)
```

# 5. Setelah dihapus fitur yang di daftar hapus akan di tambahkan ke daftar hapus seperti

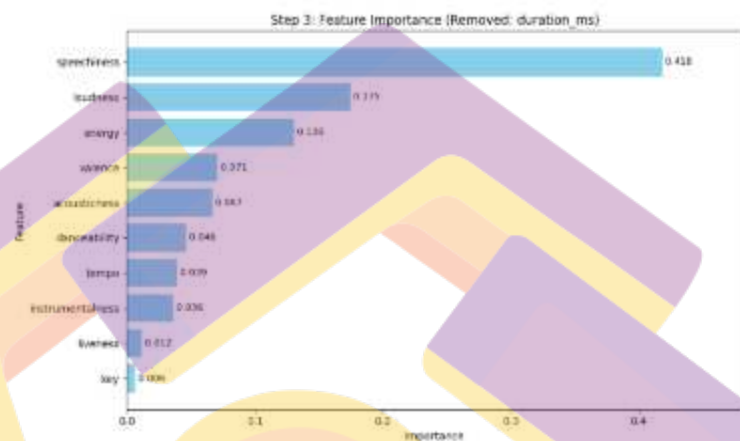
- a. Penghapusan pertama mode
- b. Penghapusan fitur kedua mode, popularity, dst

# 6. Melatih ulang model Random Forest

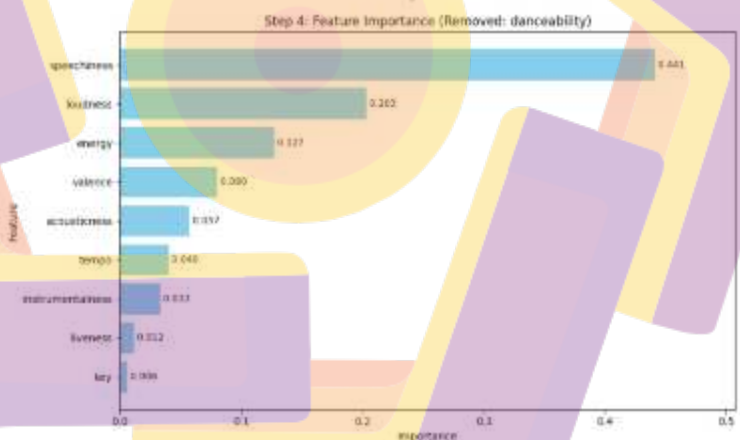
7. Mengukur akurasi model di data test
8. Menyimpan hasil evaluasi tiap step
9. Menampilkan grafik importance (opsional untuk analisis lanjutan)
10. Menampilkan tabel hasil akhir. Untuk kelengkapan coding bisa merujuk pada lampiran. Dan hasil gambar grafik yang dihasilkan seperti dibawah ini:



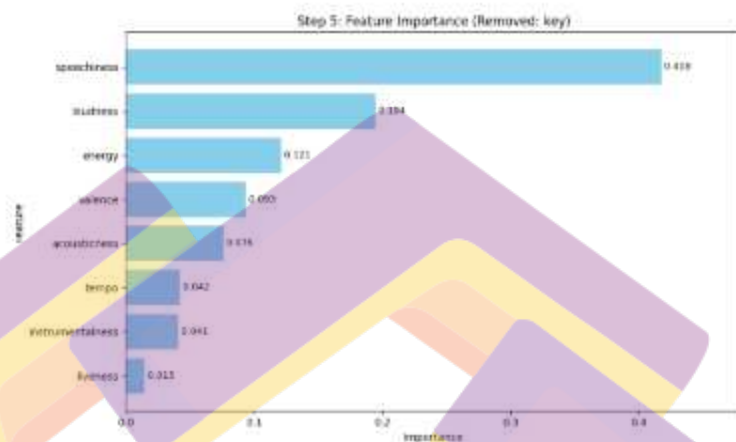
Gambar 4.6 Penghapusan Popularity



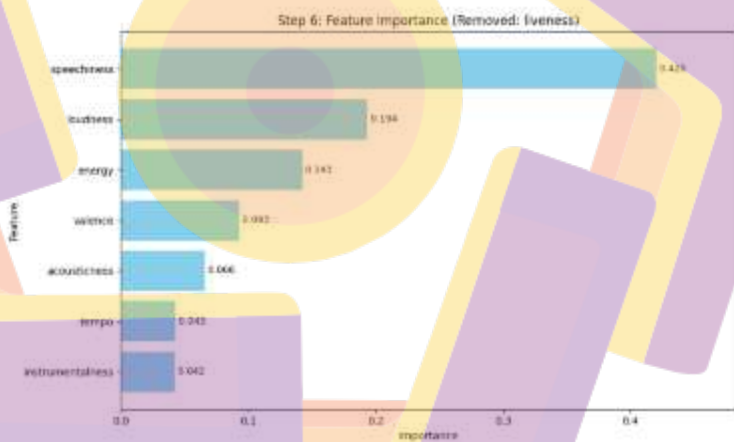
Gambar 4.7 Penghapusan Duration\_ms



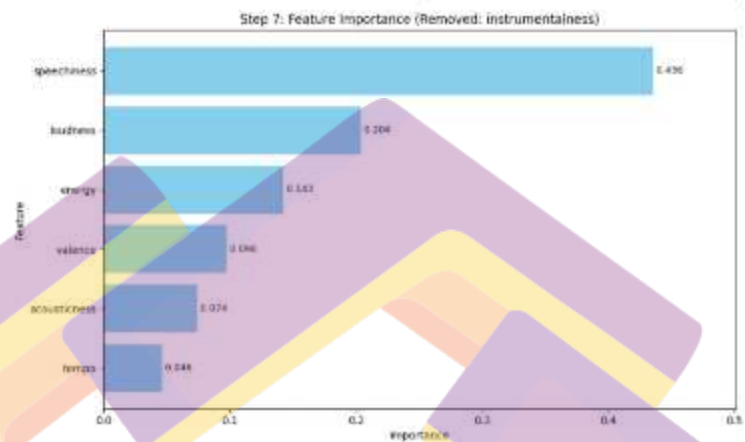
Gambar 4.8 Penghapusan Danceability



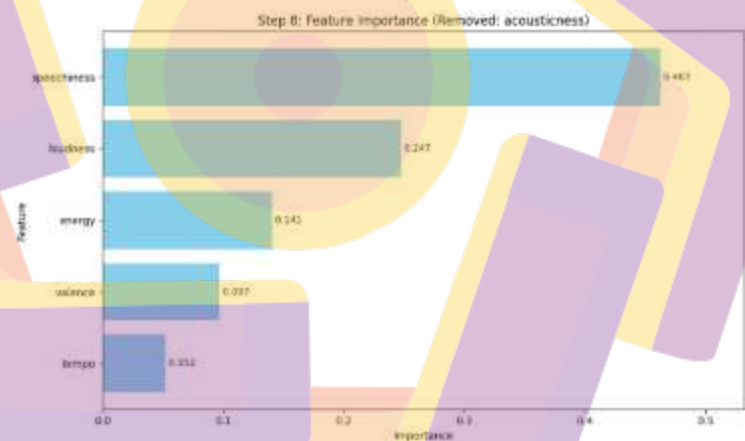
Gambar 4.9 Penghapusan Key



Gambar 4.10 Penghapusan Liveness



Gambar 4.11 Penghapusan Instrumentality



Gambar 4.12 Penghapusan Acousticness

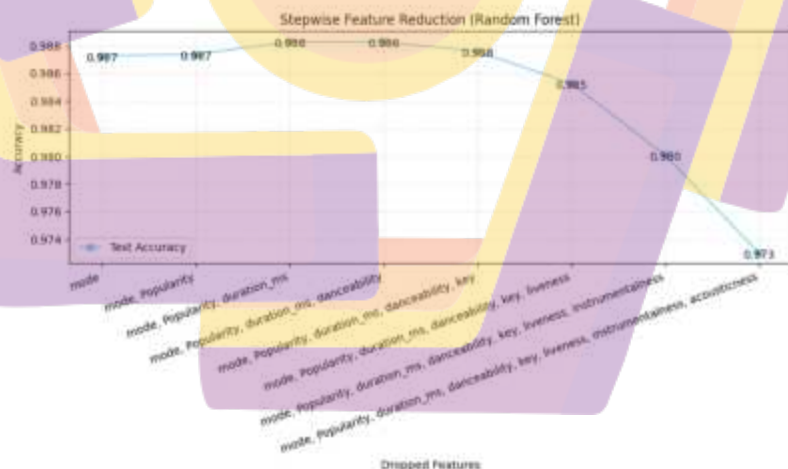
#### 4.9 Hasil Summary Stepwise Feature Reduction

Dengan menerapkan stepwise reduction Peningkatan grafik ditunjukkan oleh

1. nilai tertinggi dengan keakuratan model RF tetap konsisten di antara 0.987–0.988
2. peningkatan pada data test ketika fitur *mode*, *popularity*, *duration\_ms*, dan *danceability* dihilangkan, yang berarti penghapusan fitur-fitur tersebut tidak mempengaruhi kinerja model.

Namun, tingkat akurasi mulai terjadi penurunan ditunjukkan oleh

1. nilai keakuratan RF turun menjadi 0.973
2. penurunan pada data test setelah fitur *key*, *liveness*, *instrumentalness*, dan *acousticness* dihilangkan, menunjukkan bahwa fitur-fitur tersebut memiliki dampak signifikan terhadap kinerja model. Oleh karena itu, penelitian ini merekomendasikan agar fitur *acousticness*, *instrumentalness*, dan *liveness* dipertahankan karena fitur-fitur tersebut berperan besar dalam menjaga stabilitas dan ketepatan prediksi model.



Gambar 4.13 Visualisasi Hasil Stepwise Feature Reduction

Untuk mempermudah membaca hasil reduksi fitur maka dibuatkan tabel untuk menjelaskan step fitur yang di hapus satu persatu. Terlihat pada tabel 4.1 dibawah ini.

Tabel 4.1 Drop Data *Stepwise Feature Reduction*.

| Step Looping Delete Feature | Fitur Yang Dihapus                                                                         | Fitur Yang Aman Di Hapus Lebih Awal | Fitur Yang Dipertahankan | Fitur Pendukung Tidak Boleh Dihapus Bersamaan Dengan Fitur Awal Maupun Fitur Yang Dipertahankan | Test Accuracy |
|-----------------------------|--------------------------------------------------------------------------------------------|-------------------------------------|--------------------------|-------------------------------------------------------------------------------------------------|---------------|
| Step 1                      | Mode                                                                                       | Mode                                |                          |                                                                                                 | 0.987         |
| Step 2                      | mode, popularity                                                                           | popularity                          |                          |                                                                                                 | 0.987         |
| Step 3                      | mode, popularity, duration_ms                                                              |                                     | duration_ms              |                                                                                                 | 0.988         |
| Step 4                      | mode, popularity, duration_ms, danceability                                                |                                     | danceability             |                                                                                                 | 0.988         |
| Step 5                      | mode, popularity, duration_ms, danceability, key                                           |                                     | key                      |                                                                                                 | 0.988         |
| Step 6                      | mode, popularity, duration_ms, danceability, key, liveness                                 |                                     |                          | liveness,                                                                                       | 0,985         |
| Step 7                      | mode, popularity, duration_ms, danceability, key, liveness, instrumentalness               |                                     |                          | instrumentalness                                                                                | 0,980         |
| Step 8                      | mode, popularity, duration_ms, danceability, key, liveness, instrumentalness, acousticness |                                     |                          | acousticness                                                                                    | 0,973         |

Dapat ditarik kesimpulan dalam penerapan stepwisw feature reduction memberi dampak:

### 1. Tahap Awal Penghapusan Fitur

Penghapusan fitur mode, mode-popularity, hingga mode-popularity-duration\_ms menunjukkan peningkatan akurasi dari 0,987 menjadi 0,988. Artinya, fitur mode dan popularity tidak memberikan kontribusi signifikan terhadap performa klasifikasi.

### 2. Tahap Tengah (Akurasi Optional)

Pada penghapusan fitur hingga danceability dan key, akurasi tetap bertahan di nilai tertinggi 0,988, menunjukkan bahwa:

- a) Model stabil
- b) Informasi utama masih terjaga
- c) Tidak terjadi degradasi performa

### 3. Tahap Lanjut (Penurunan Performa)

Penghapusan mulai liveness, instrumentalness, dan acoudticness. Terjadi penurunan akurasi bertahap dari 0,985  $\rightarrow$  0,980  $\rightarrow$  0,973. artinya fitur-fitur tersebut masih memiliki kontribusi pendukung, dan penghapusannya menyebabkan hilangnya informasi penting.

4. Akurasi tertinggi sebesar 0,988 diperoleh pada komposisi fitur setelah penghapusan mode dan popularity, serta masih mempertahankan fitur duration\_ms, danceability, dan key.

### 5. Rekomendasi fitur yang di sesuaikan dengan stepwise

Fitur yang direkomendasikan untuk sistem rekomendasi musik belajar:

- a) duration\_ms
- b) danceability
- c) key. Ketiga fitur ini terbukti menjaga akurasi tertinggi secara konsisten.
- 6. **Fitur yang tidak direkomendasikan/** tidak meningkatkan akurasi, bahkan menurunkan performa ketika dihapus terlalu jauh
  - a) mode
  - b) popularity
  - c) liveness
  - d) instrumentalness
  - e) acousticness
- 7. Berdasarkan hasil penerapan pengurangan fitur bertahap dengan Random Forest, tampaknya penghapusan fitur mode dan popularitas justru meningkatkan akurasi model. Akurasi tertinggi, yaitu 0,988, dicapai ketika fitur-fitur kunci seperti duration\_ms, danceability, dan key dipertahankan. Sebaliknya, ketika fitur-fitur pendukung seperti liveness, instrumentalness, dan acousticness dihapus, kinerja model mulai menurun. Oleh karena itu, dapat disimpulkan bahwa duration\_ms, danceability, dan key adalah fitur-fitur terpenting dan juga ditambah dengan fitur pendukung (liveness, instrumentalness, dan acousticness) dalam membangun sistem rekomendasi musik untuk mendukung aktivitas pembelajaran.

Penelitian yang dilakukan angel dengan judul “**Background Music And Cognitive Performance 1,2**” pada tahun 2010 menunjukan bahwa musik instrumental dengan ritme yang konsisten yang sering dijumpai dalam genre seperti

ambient atau klasik instrument lembut lebih cocok untuk mendukung kegiatan belajar.[64] dan juga dipertegas dalam penelitian lain yang dilakukan oleh Lucas Kiss dalam penelitian berjudul **“The effect of preferred background music on task-focus in sustained attention”** pada tahun 2020 menunjukkan bahwa musik yang sesuai untuk belajar bukan ditentukan oleh genre, melainkan oleh karakter audio yang stabil, tanpa unsur verbal/lirik, serta selaras dengan preferensi pengguna, karena mampu mendukung peningkatan fokus kognitif terutama pada tugas belajar dengan beban kognitif yang rendah.[65]. Ditambah pendapat secara personal oleh expert judgement Berliana Putri, S.Pd., Gr., seorang guru seni dengan kompetensi di bidang musik, komposisi musik yang tersedia pada platform Spotify secara umum telah menunjukkan kesesuaian dalam perancangan musik yang ditujukan untuk aktivitas belajar.

#### **Validasi Kontribusi Fitur**

Validasi kontribusi fitur dengan nilai feature importance tinggi dilakukan melalui analisis hasil stepwise feature reduction. Pendekatan ini mengevaluasi perubahan performa model dengan menghapus fitur secara bertahap berdasarkan nilai feature importance. Jika penghapusan fitur dengan nilai feature importance rendah tidak menyebabkan penurunan performa, sedangkan penghapusan fitur dengan nilai feature importance tinggi menurunkan kinerja klasifikasi secara signifikan, maka dapat disimpulkan bahwa fitur-fitur tersebut memiliki kontribusi nyata terhadap performa model dan bukan merupakan hubungan kebetulan (*spurious correlation*).

#### 4.10 Perbandingan Stepwise Reduction Tanpa Menerapkan Feature Importance

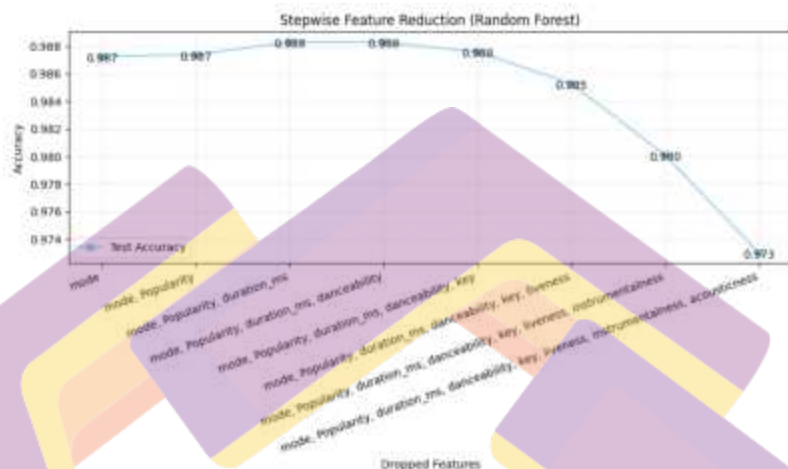
Hasil percobaan menunjukkan bahwa Reduksi fitur tanpa feature importance berpotensi menghilangkan fitur-fitur yang paling relevan, sehingga mampu menyebabkan penurunan akurasi dan tidak bisa digunakan sebagai dasar pengambilan keputusan. Dengan demikian, penggunaan feature importance terbukti lebih efektif dalam menjaga keseimbangan performa klasifikasi. Agar lebih mempermudah dalam membaca konteks tersebut maka, akan menjelaskan dalam tabel dibawah ini:

Tabel 4.2 Perbandingan Tanpa Menerapkan Feature Importance.

| Aspek                      | Dengan Feature Importance | Tanpa Feature Importance |
|----------------------------|---------------------------|--------------------------|
| Dasar pemilihan fitur      | Kontribusi ke model       | Asumsi / acak / korelasi |
| Risiko hapus fitur penting | Rendah                    | Tinggi                   |
| Akurasi                    | Stabil / meningkat        | Cenderung turun          |
| Overfitting                | Berkurang                 | Tidak terkontrol         |
| Interpretabilitas          | Tinggi                    | Rendah                   |
| Konsistensi hasil          | Tinggi                    | Rendah                   |

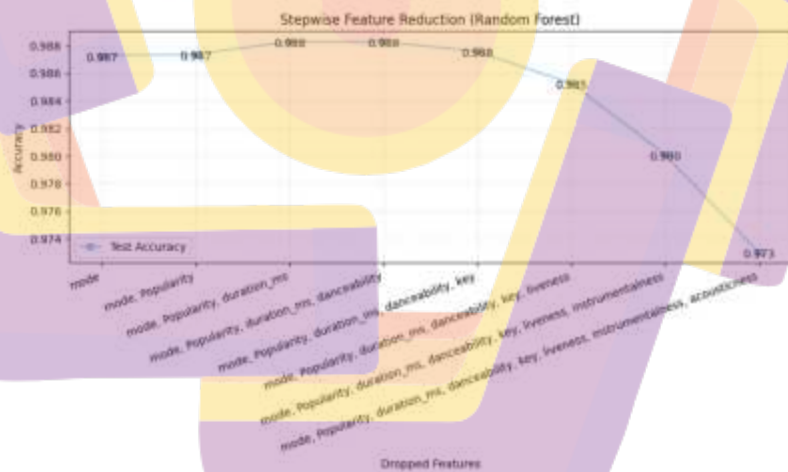
Hasil penerapan feature reduction menunjukkan bahwa tidak terdapat perubahan performa yang signifikan antara model dengan dan tanpa reduksi fitur. Namun demikian, penggunaan feature importance memberikan nilai tambah yang jelas, yaitu kemampuan untuk mengidentifikasi tingkat kontribusi masing-masing fitur terhadap kinerja model. Informasi ini divisualisasikan melalui grafik feature importance, sehingga memungkinkan peneliti memperoleh gambaran yang lebih objektif mengenai fitur-fitur yang paling berpengaruh. Temuan tersebut selanjutnya dapat dijadikan pedoman dalam penerapan stepwise feature reduction, sehingga proses penghapusan fitur dapat dilakukan secara lebih terarah dan akurat tanpa mengorbankan kinerja klasifikasi. Dibawah ini akurasi

#### Tanpa Feature Importance



Gambar 4.14 Perbandingan Stepwise Reduction Tanpa Feature Importance

Dengan Feature Importance



Gambar 4.15 Perbandingan Stepwise Reduction Dengan Feature Importance

#### 4.11 Modelling

Setelah itu, kita akan melatih model akhir menggunakan algoritma *Random Forest* pada seluruh data pelatihan ( $x_{train\_bal}$  dan  $y_{train\_bal}$ ) yang sudah diproses dengan SMOTE dan sebelumnya. Setelah model selesai dilatih, prediksi akan dilakukan pada data uji yang belum pernah dikenali oleh model sebelumnya. Terakhir, kita akan menilai kinerja model dengan membandingkan hasil prediksi dengan label asli menggunakan berbagai metrik evaluasi. Memilih metode evaluasi yang tepat sangat penting untuk memastikan kemampuan generalisasi model [47]. Sementara algoritma *Random Forest*, yang bekerja dengan membagi data menjadi cabang-cabang berdasarkan atribut untuk membentuk struktur pohon keputusan, telah terbukti efektif dalam mengidentifikasi pola-pola kompleks secara akurat dan mudah ditafsirkan [19], [45], [66].

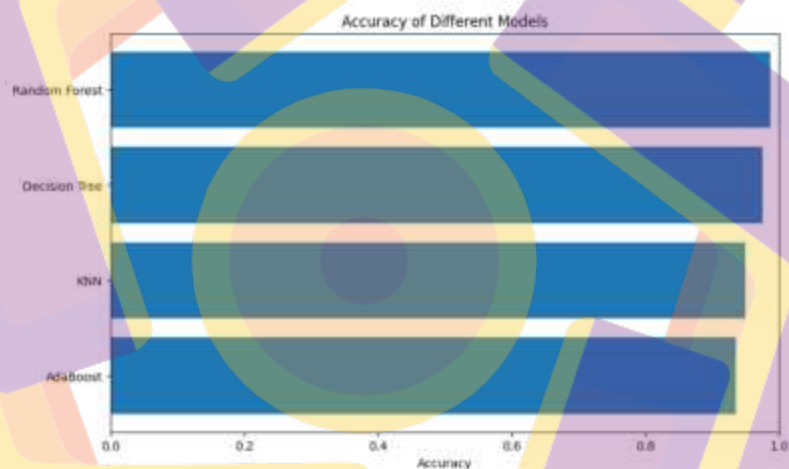
#### 4.12 Evaluasi Model

Hasil evaluasi menunjukkan bahwa reduksi dimensi berbasis feature importance mampu mempertahankan bahkan meningkatkan performa klasifikasi *Random Forest* selama fitur yang dieliminasi memiliki kontribusi rendah.

Tabel 4.3 Sebaran Data Setelah Normalisasi dan Sebaran Data Setelah Reduksi Dimensi

| Kondisi           |                            | Sebaran data setelah di SMOTE    |
|-------------------|----------------------------|----------------------------------|
| Normalisasi       |                            | 52, 285 dari ketiga sample kelas |
| Feature Reduction | Feature Importance         | 52, 285 dari ketiga sample kelas |
|                   | Stepwise Feature Reduction | 52, 285 dari ketiga sample kelas |

Berdasarkan hasil pengujian yang ditunjukkan pada gambar 4.14 dibawah ini *Random Forest* memiliki tingkat akurasi tertinggi dibandingkan dengan algoritma lain. Keunggulan ini dapat dikaitkan dengan karakteristik *Random Forest* sebagai metode ensemble yang dapat mengurangi varians model dengan menggabungkan beberapa pohon keputusan. Namun, hasil ini tidak dapat dianggap berlaku secara universal, karena kinerja algoritma sangat dipengaruhi oleh sifat dataset yang digunakan.



Gambar 4.16 Perbandingan Model

Berdasarkan hasil evaluasi, model RF menunjukkan kinerja yang sangat tinggi dengan akurasi 98,77% dan nilai presisi, recall, serta skor F1 sebesar 0,99, yang membuktikan efektivitas metode yang digunakan. Pada tingkat kelas, model secara konsisten memberikan skor 0,99 di semua metrik untuk kelas ketiga, dengan jumlah prediksi benar yang sangat tinggi untuk setiap kelas. *Matriks Confusion Matrix* menunjukkan bahwa dari 25.632 data uji, hanya terdapat 147 kesalahan, yang mengonfirmasi kemampuan model dalam mengklasifikasikan musik secara

akurat untuk data pembelajaran, sebagaimana dapat dilihat pada gambar 4.15 dibawah ini

```

Akurasi test : 0.9877186741571834

Classification Report:
              precision    recall  f1-score   support

     1.0         0.99      0.99      0.99       13196
     2.0         0.99      0.99      0.99       13348
     3.0         0.99      0.99      0.99       10009

 accuracy          0.99      0.99      0.99      25632
 macro avg         0.99      0.99      0.99      25632
 weighted avg      0.99      0.99      0.99      25632

Confusion Matrix:
[[13047   25   124]
 [   17  2328    1]
 [   141    7  9942]]

```

Gambar 4.17 Evaluasi Performa Model

Dari gambar 4.15 Evaluasi Performa Model bisa didetailkan kembali melalui beberapa tabel dibawah ini untuk mendapatkan informasi yang lebih *clear* dan detail. Pada tabel 4.3 *Model Evaluation Summary* menjelaskan bahwa Model *Random Forest* menunjukkan kinerja yang sangat baik pada tahap pengujian, dengan tingkat akurasi sebesar 0,9877 atau 98,77%. Hal ini mengindikasikan bahwa model mampu mengklasifikasikan sebagian besar data uji secara tepat. Selain itu, nilai rata-rata presisi, recall, dan F1-score masing-masing diperoleh sebesar 0,99 yang mencerminkan kemampuan model dalam menghasilkan prediksi yang akurat, mendeteksi kelas dengan tepat, serta mempertahankan keseimbangan antara kesalahan positif dan negatif. Dengan total 25.632 data uji yang digunakan, hasil evaluasi ini menunjukkan bahwa model *Random Forest* memiliki performa yang sangat stabil, konsisten, dan efektif dalam melakukan klasifikasi data musik berdasarkan fitur yang digunakan. Dengan demikian, model ini layak diterapkan sebagai metode klasifikasi pada penelitian ini

serta memiliki potensi untuk digunakan pada studi lanjutan dengan skala data yang lebih besar.

Tabel 4.4 *Model Evaluation Summary*

| Metric              | Value                |
|---------------------|----------------------|
| Model               | <i>Random Forest</i> |
| Test Accuracy       | 0.9877 (98.77%)      |
| Macro Avg Precision | 0.99                 |
| Macro Avg Recall    | 0.99                 |
| Macro Avg F1-Score  | 0.99                 |
| Total Test Samples  | 25,632               |

Kemudian untuk penjelasan model *performace per class* Evaluasi kinerja model pada masing-masing kelas menunjukkan hasil yang sangat baik, dengan nilai presisi, recall, dan F1-score yang mencapai 0.99 untuk seluruh kelas. Hal ini menandakan bahwa model dapat mengklasifikasikan data secara konsisten dengan tingkat kesalahan yang minimal pada setiap kategori. Dari keseluruhan dataset yang terdiri dari sampel kelas 1: 13.196, sampel kelas 2: 2.346, dan sampel kelas 2: 10.090, model mampu mengidentifikasi dengan tepat sebanyak 13.047, 2.328, dan 9.942 sampel untuk masing-masing kelas. Pencapaian ini membuktikan bahwa model memiliki kemampuan yang kuat dalam membedakan ciri-ciri distinktif antar kelas secara akurat terlihat pada tabel 4.4 *Model Performace per Class*.

Tabel 4.5 *Model Performance per Class*

| Class   | Precision | Recall | F1-Score | Support | Correct Predictions |
|---------|-----------|--------|----------|---------|---------------------|
| Class 1 | 0.99      | 0.99   | 0.99     | 13,196  | 13,047              |
| Class 2 | 0.99      | 0.99   | 0.99     | 2,346   | 2,328               |
| Class 3 | 0.99      | 0.99   | 0.99     | 10,090  | 9,942               |

Untuk penjelasan *Conduction Matrix* pada kelas 1 (Baris Pertama) menjelaskan dari 13.196 titik data yang sebenarnya Kelas 1, model memprediksi 13.047 sampel dengan tepat sebagai Kelas 1. Namun, terdapat kesalahan klasifikasi:

25 sampel diprediksi sebagai Kelas 2 padahal sebenarnya Kelas 1, dan 124 sampel diprediksi sebagai Kelas 3 padahal sebenarnya Kelas 1. Oleh karena itu, dari total titik data Kelas 1, model membuat kesalahan pada 149 sampel.

Lalu pada **kelas 2 (Baris Kedua)** dari 2.346 titik data yang sebenarnya Kelas 2, model memprediksi 2.328 sampel dengan tepat sebagai Kelas 2. Terdapat sangat sedikit kesalahan klasifikasi: 17 sampel yang sebenarnya Kelas 2 diprediksi sebagai Kelas 1, dan hanya 1 sampel yang diprediksi sebagai Kelas 3. Hal ini menunjukkan bahwa Kelas 2 memiliki karakteristik yang paling mudah dimodifikasi, dengan total kesalahan hanya 18 sampel.

Dan **kelas 3 (Baris Ketiga)** Dari 10.090 titik data yang sebenarnya termasuk dalam Kelas 3, model memprediksi 9.942 sampel dengan tepat sebagai Kelas 3. Kesalahan terbesar terjadi pada 141 sampel yang sebenarnya termasuk Kelas 3 tetapi diprediksi sebagai Kelas 1, sementara 7 sampel diprediksi sebagai Kelas 2. Total kesalahan di Kelas 3 adalah 148 sampel, yang menunjukkan adanya tumpang tindih karakteristik yang signifikan, terutama dengan Kelas 1. terlihat pada tabel 4.5 *Confusion Matrix* dibawah ini.

Tabel 4.6 *Confusion Matrix*

| Actual \ Predicted | Class 1 | Class 2 | Class 3 | Total Actual |
|--------------------|---------|---------|---------|--------------|
| Class 1            | 13,047  | 25      | 124     | 13,196       |
| Class 2            | 17      | 2,328   | 1       | 2,346        |
| Class 3            | 141     | 7       | 9,942   | 10,090       |
| Total Predicted    | 13,205  | 2,360   | 10,067  | 25,632       |

## BAB V

### PENUTUP

#### 5.1. Kesimpulan

Penelitian ini memiliki batasan terkait kemungkinan adanya bias algoritmik, di mana metode *Random Forest* seringkali lebih efektif pada kumpulan data yang memiliki fitur yang rumit. Untuk itu, disarankan agar penelitian selanjutnya melakukan uji coba kemampuan model pada kumpulan data yang lain serta menerapkan metrik evaluasi tambahan seperti presisi, recall, dan F1-score untuk mengurangi risiko bias dalam hasil.

Studi ini juga menunjukkan bahwa kombinasi SMOTE, reduksi fitur, dan normalisasi Min-Max berhasil mengoptimalkan dataset dan meningkatkan stabilitas model pembelajaran. Proses normalisasi secara efektif menyamakan skala fitur ke rentang 0,0–1,0 dengan tetap mempertahankan karakteristik distribusi asli, sehingga mendukung pelatihan model yang lebih efisien dan konsisten.

Pengklasifikasi *Random Forest* mencapai kinerja yang luar biasa, mencapai akurasi 98,77% pada dataset uji dan presisi, recall, serta skor F1 sebesar 0,99 di semua kelas, dengan akurasi 99% selama validasi silang. Dengan hanya 147 sampel yang salah diklasifikasikan dari 25.632 sampel dalam dataset uji, model ini terbukti sangat andal dalam mengidentifikasi musik yang cocok untuk pembelajaran.

Penelitian ini membuktikan bahwa penerapan reduksi dimensi menggunakan metode feature importance pada *Random Forest* mampu mengurangi jumlah fitur tanpa menurunkan performa klasifikasi, serta meningkatkan efisiensi model. Hasil evaluasi fitur menunjukkan bahwa fitur

1. Akustik (Acoustic),
2. Instrumentalitas (Instrumentalness),
3. Keaktifan (liveness),
4. Energi (Energy), dan
5. Kenyaringan (Loudness) merupakan faktor yang paling berpengaruh terhadap kinerja klasifikasi. Selain itu, fitur
6. Key,
7. Valensi (valence), dan
8. Tempo juga berperan dalam mendukung diferensiasi genre musik secara efektif.

Secara keseluruhan, alur kerja yang diusulkan berhasil mengklasifikasikan genre musik yang berpotensi meningkatkan proses pembelajaran, khususnya genre Pop, Classical Soundtracks, dan Lo-fi, yang ditandai oleh karakteristik fitur akustik, instrumentalitas, keaktifan, energi, dan kenyaringan pada data Spotify.

Tabel 5.1 *Selected Features and Model Accuracy*

| Category                  | Feature          | Drop Feature | Accuracy Model |
|---------------------------|------------------|--------------|----------------|
| Primary Critical Features | Acousticness     | No           | 0.988          |
|                           | Instrumentalness | No           | 0.988          |
|                           | Liveness         | No           | 0.988          |
| Core Supporting Features  | Energy           | No           | 0.988          |
|                           | Loudness         | No           | 0.988          |
|                           | Speechiness      | No           | 0.988          |
|                           | Valence          | No           | 0.988          |
|                           | Tempo            | No           | 0.988          |
| Removable / Less Impact   | Mode             | Yes          | 0.987-0.988    |
|                           | Popularity       | Yes          | 0.987-0.988    |
|                           | Duration_ms      | Yes          | 0.987-0.988    |
|                           | Danceability     | Yes          | 0.987-0.988    |
|                           | Key              | Optional     | 0.988          |

## 5.2. Saran

Berdasarkan temuan ini, terdapat beberapa rekomendasi untuk meningkatkan efisiensi dan pengembangan lebih lanjut. Pertama, fitur-fitur yang kurang berdampak seperti

1. *Mode*,
2. *Popularity*,
3. *Duration*
4. *Danceability*, dan
5. *Key* sebaiknya dihapus atau tidak digunakan dalam model untuk menghasilkan arsitektur yang lebih sederhana dan efisien dengan tetap mempertahankan tingkat akurasi yang tinggi karena Pengurangan fitur yang berlebihan berpotensi menghilangkan informasi penting, sehingga mengurangi kinerja model.

Kedua, studi serupa pada dataset *Spotify Sleep Music (SSM)* direkomendasikan untuk mengembangkan sistem rekomendasi musik yang lebih spesifik.

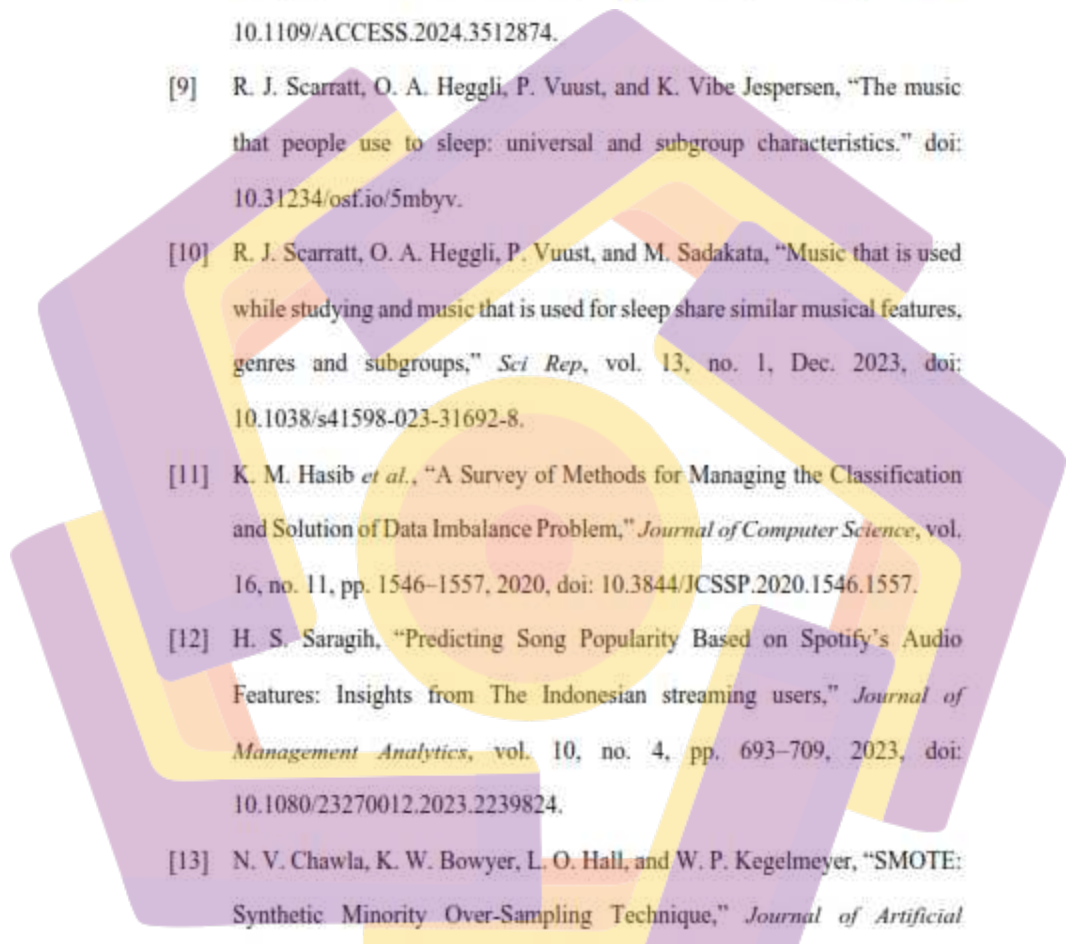
Kemudian penelitian ini masih memiliki keterbatasan dalam penerapan pada sistem rekomendasi musik pembelajaran secara real-time. Oleh karena itu, diperlukan pengembangan lebih lanjut untuk mengatasi tantangan efisiensi komputasi dan latensi sistem. Pengembangan sistem klasifikasi musik menjadi aplikasi rekomendasi pembelajaran secara real-time menghadapi tantangan utama berupa keterbatasan waktu komputasi dan kebutuhan respons yang cepat. Untuk menjaga keseimbangan antara akurasi dan efisiensi, reduksi dimensi melalui feature

selection berbasis feature importance dan stepwise feature reduction dapat digunakan untuk mempertahankan fitur yang paling relevan sehingga mempercepat proses inferensi. Selain itu, model klasifikasi seperti Random Forest perlu dioptimalkan, misalnya dengan mengurangi jumlah atau kedalaman pohon, agar sistem tetap akurat dan efisien dalam lingkungan real-time.

Ketiga, untuk memperkaya karakteristik emosional dan kenyamanan dalam klasifikasi musik, disarankan untuk menambahkan fitur baru seperti tempo, valensi, dan ucapan guna melengkapi fitur audio yang sudah ada dan meningkatkan personalisasi dalam sistem rekomendasi musik yang dirancang untuk dukungan pembelajaran.

## DAFTAR PUSTAKA

- [1] D. Wulan *et al.*, "Penggunaan Seni Musik dalam Mendukung Perkembangan Kognitif dan Emosional Siswa SD," 2023. doi: 10.69688/jpip.v1i2.15.
- [2] K. B. Batt-Rawden, "The benefits of self-selected music on health and well-being," *Arts in Psychotherapy*, vol. 37, no. 4, pp. 301–310, Sep. 2010, doi: 10.1016/j.aip.2010.05.005.
- [3] A. G. Jondya and B. H. Iswanto, "Indonesian's Traditional Music Clustering Based on Audio Features," in *Procedia Computer Science*, Elsevier B.V., 2017, pp. 174–181. doi: 10.1016/j.procs.2017.10.019.
- [4] M. J. Pachali and H. Datta, "What Drives Demand for Playlists on Spotify?," *Marketing Science*, vol. 44, no. 1, pp. 54–64, Jan. 2025, doi: 10.1287/mksc.2022.0273.
- [5] R. Jane Scarratt, J. Stupacher, P. Vuust, and K. Vibe Jespersen, "Motivations for using music for sleep: insights from YouTube comments," 2024. doi: 10.31234/osf.io/28um6.
- [6] K. Jacobson, V. Murali, E. Newett, B. Whitman, and R. Yon, "Music Personalization at Spotify," Association for Computing Machinery (ACM), Sep. 2016, pp. 373–373. doi: 10.1145/2959100.2959120.
- [7] B. Zhang *et al.*, "Understanding user behavior in Spotify," in *Proceedings - IEEE INFOCOM*, 2013, pp. 220–224. doi: 10.1109/INFOCOM.2013.6566767.

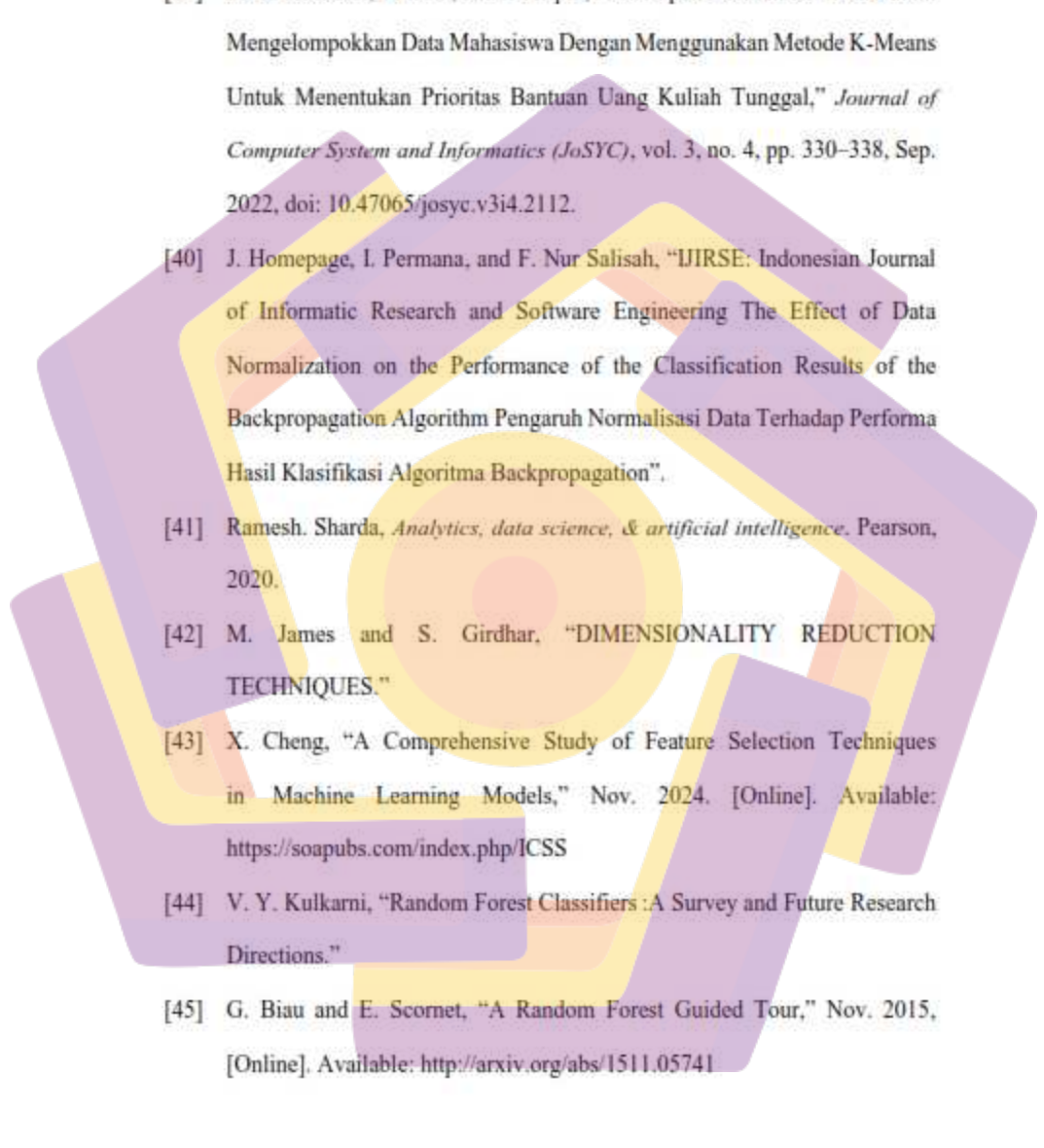
- 
- [8] J. Shen and G. Xiao, "Music Genre Classification Based on Functional Data Analysis," *IEEE Access*, pp. 1–1, 2024, doi: 10.1109/ACCESS.2024.3512874.
- [9] R. J. Scarratt, O. A. Heggli, P. Vuust, and K. Vibe Jespersen, "The music that people use to sleep: universal and subgroup characteristics." doi: 10.31234/osf.io/5mbyv.
- [10] R. J. Scarratt, O. A. Heggli, P. Vuust, and M. Sadakata, "Music that is used while studying and music that is used for sleep share similar musical features, genres and subgroups," *Sci Rep*, vol. 13, no. 1, Dec. 2023, doi: 10.1038/s41598-023-31692-8.
- [11] K. M. Hasib *et al.*, "A Survey of Methods for Managing the Classification and Solution of Data Imbalance Problem," *Journal of Computer Science*, vol. 16, no. 11, pp. 1546–1557, 2020, doi: 10.3844/JCSSP.2020.1546.1557.
- [12] H. S. Saragih, "Predicting Song Popularity Based on Spotify's Audio Features: Insights from The Indonesian streaming users," *Journal of Management Analytics*, vol. 10, no. 4, pp. 693–709, 2023, doi: 10.1080/23270012.2023.2239824.
- [13] N. V. Chawla, K. W. Bowyer, L. O. Hall, and W. P. Kegelmeyer, "SMOTE: Synthetic Minority Over-Sampling Technique," *Journal of Artificial Intelligence Research*, vol. 16, pp. 321–357, 2002, doi: 10.1613/jair.953.
- [14] A. A. Maita and N. A. S. Winarsih, "Comparison of Hyperparameter Tuning in Decision Tree and Random Forest Algorithms for Song Genre

- Classification," 2025. [Online]. Available: <http://jurnal.polibatam.ac.id/index.php/JAIC>
- [15] J. Prasetya and A. Abdurakhman, "Comparison Of Smote Random Forest And Smote K-Nearest Neighbors Classification Analysis On Imbalanced Data," *MEDIA STATISTIKA*, vol. 15, no. 2, pp. 198–208, Apr. 2023, doi: 10.14710/medstat.15.2.198-208.
- [16] Bernard Ginn Jr and Sejoniq Johnson, "Evaluating the Efficacy of Random Forest in Mood Classification for Music Recommendation System," 2024, doi: DOI:10.13140/RG.2.2.22524.04485.
- [17] K. A. Nida Qurrota and R. Wahyuni, "The Use Of Songs On The Spotify App to Improve Students Listening Ability," *English Journal Antartika*, vol. 2, no. 2, 2024, [Online]. Available: <https://ejournal.mediaantartika.id/index.php/eja/index>
- [18] O. A. Heggli, J. Stupacher, and P. Vuust, "Diurnal fluctuations in musical preference," *R Soc Open Sci*, vol. 8, no. 11, 2021, doi: 10.1098/rsos.210885.
- [19] A. A. Maita and N. A. S. Winarsih, "Comparison of Hyperparameter Tuning in Decision Tree and Random Forest Algorithms for Song Genre Classification," Aug. 2025, doi: <https://doi.org/10.30871/jaic.v9i4.10142>.
- [20] X. Yuan, S. Liu, W. Feng, and G. Dauphin, "Feature Importance Ranking of Random Forest-Based End-to-End Learning Algorithm," *Remote Sens (Basel)*, vol. 15, no. 21, Nov. 2023, doi: 10.3390/rs15215203.
- [21] M. I. Prasetyowati, N. U. Maulidevi, and K. Surendro, "The accuracy of Random Forest performance can be improved by conducting a feature

- selection with a balancing strategy," *PeerJ Comput Sci*, vol. 8, 2022, doi: 10.7717/PEERJ-CS.1041.
- [22] T. Mahmood, M. Usman, and C. Conrad, "Selecting Relevant Features for Random Forest-Based Crop Type Classifications by Spatial Assessments of Backward Feature Reduction," *PFG - Journal of Photogrammetry, Remote Sensing and Geoinformation Science*, vol. 93, no. 2, pp. 173–196, Apr. 2025, doi: 10.1007/s41064-024-00329-4.
- [23] S. N. Bardab, T. M. Ahmed, and T. A. A. Mohammed, "Data mining classification algorithms: An overview," *International Journal of Advanced and Applied Sciences*, vol. 8, no. 2, pp. 1–5, Feb. 2021, doi: 10.21833/ijaas.2021.02.001.
- [24] P. Ayu Firnanda *et al.*, "Analisis Perbandingan Decision Tree dan Random Forest dalam Klasifikasi Penjualan Produk pada Supermarket," *Emerging Statistics and Data Science Journal*, vol. 3, no. 1, 2025.
- [25] J. Song, "Comparison and Analysis of Accuracy of Traditional Random Forest Machine Learning Model and XGBoost Model On Music Emotion Classification Dataset," in *ACM International Conference Proceeding Series*, Association for Computing Machinery, Oct. 2023, pp. 712–716. doi: 10.1145/3650215.3650340.
- [26] Y. Meng, "Music Genre Classification: A Comparative Analysis of CNN and XGBoost Approaches with Mel-frequency Cepstral Coefficients and Mel Spectrograms."

- [27] R. Imantiyar, ; DThomas, and H. Fudholi, "Kajian Pengaruh Dataset dan Bias Dataset terhadap Performa Akurasi Deteksi Objek," vol. 14, no. 2, 2021, doi: 10.33322/petir.v14i2.1150.
- [28] B. Koch, E. Denton, A. Hanna, G. Research, S. Francisco, and J. G. Foster, "Reduced, Reused and Recycled: The Life of a Dataset in Machine Learning Research." [Online]. Available: <https://paperswithcode.com>
- [29] S. García, S. Ramírez-Gallego, J. Luengo, J. M. Benítez, and F. Herrera, "Big data preprocessing: methods and prospects," *Big Data Anal.*, vol. 1, no. 1, Dec. 2016, doi: 10.1186/s41044-016-0014-0.
- [30] T. Gori, A. Sunyoto, and H. Al Fatta, "Preprocessing Data dan Klasifikasi untuk Prediksi Kinerja Akademik Siswa," *Jurnal Teknologi Informasi dan Ilmu Komputer*, vol. 11, no. 1, pp. 215–224, Feb. 2024, doi: 10.25126/jtiik.20241118074.
- [31] Joaquin. Quiñero-Candela, *Dataset shift in machine learning*. MIT Press, 2009.
- [32] A. Nurhopipah and U. Hasanah, "Dataset Splitting Techniques Comparison For Face Classification on CCTV Images," *IJCCS (Indonesian Journal of Computing and Cybernetics Systems)*, vol. 14, no. 4, p. 341, Oct. 2020, doi: 10.22146/ijccs.58092.
- [33] I. Alabdulmohsin, X. Wang, A. Steiner, P. Goyal, A. D'Amour, and X. Zhai, "CLIP the Bias: How Useful is Balancing Data in Multimodal Learning?," Mar. 2024, [Online]. Available: <http://arxiv.org/abs/2403.04547>

- [34] D. Elreedy, A. F. Atiya, and F. Kamalov, "A theoretical distribution analysis of synthetic minority oversampling technique (SMOTE) for imbalanced learning," *Mach Learn*, vol. 113, no. 7, pp. 4903–4923, Jul. 2024, doi: 10.1007/s10994-022-06296-4.
- [35] G. Ryan, P. Katarina, and D. Suhartono, "MBTI Personality Prediction Using Machine Learning and SMOTE for Balancing Data Based on Statement Sentences," *Information (Switzerland)*, vol. 14, no. 4, Apr. 2023, doi: 10.3390/info14040217.
- [36] I. Komang Dharmendra, I. Made, A. W. Putra, and Y. P. Atmojo, "Evaluasi Efektivitas SMOTE dan Random Under Sampling pada Klasifikasi Emosi Tweet," *Informatics for Educators And Professionals: Journal of Informatics*, vol. 9, no. 2, pp. 192–193, 2024.
- [37] M. Sholeh, D. Andayati, R. Yuliana Rachmawati, P. Studi Informatika, and F. Teknologi Informasi dan Bisnis, "DATA MINING MODEL KLASIFIKASI MENGGUNAKAN ALGORITMA K-NEAREST NEIGHBOR DENGAN NORMALISASI UNTUK PREDIKSI PENYAKIT DIABETES DATA MINING MODEL CLASSIFICATION USING ALGORITHM K-NEAREST NEIGHBOR WITH NORMALIZATION FOR DIABETES PREDICTION."
- [38] R. G. Whendasmoro and J. Joseph, "Analisis Penerapan Normalisasi Data Dengan Menggunakan Z-Score Pada Kinerja Algoritma K-NN," *JURIKOM (Jurnal Riset Komputer)*, vol. 9, no. 4, p. 872, Aug. 2022, doi: 10.30865/jurikom.v9i4.4526.

- 
- [39] M. R. Kusnaldi, T. Gulo, and S. Aripin, "Penerapan Normalisasi Data Dalam Mengelompokkan Data Mahasiswa Dengan Menggunakan Metode K-Means Untuk Menentukan Prioritas Bantuan Uang Kuliah Tunggal," *Journal of Computer System and Informatics (JoSYC)*, vol. 3, no. 4, pp. 330–338, Sep. 2022, doi: 10.47065/josyc.v3i4.2112.
- [40] J. Homepage, I. Permana, and F. Nur Salisah, "IJIRSE: Indonesian Journal of Informatic Research and Software Engineering The Effect of Data Normalization on the Performance of the Classification Results of the Backpropagation Algorithm Pengaruh Normalisasi Data Terhadap Performa Hasil Klasifikasi Algoritma Backpropagation".
- [41] Ramesh. Sharda, *Analytics, data science, & artificial intelligence*, Pearson, 2020.
- [42] M. James and S. Girdhar, "DIMENSIONALITY REDUCTION TECHNIQUES."
- [43] X. Cheng, "A Comprehensive Study of Feature Selection Techniques in Machine Learning Models," Nov. 2024. [Online]. Available: <https://soapubs.com/index.php/ICSS>
- [44] V. Y. Kulkarni, "Random Forest Classifiers :A Survey and Future Research Directions."
- [45] G. Biau and E. Scornet, "A Random Forest Guided Tour," Nov. 2015, [Online]. Available: <http://arxiv.org/abs/1511.05741>

- [46] J. Elektronik, I. K. Udayana, I. L. Simarmata, W. Supriana, and S. Kuta, "Music Genre Classification Using Random Forest Model," vol. 12, no. 1, pp. 2654–5101, 2023.
- [47] J. Elektronik, I. K. Udayana, I. L. Simarmata, W. Supriana, and S. Kuta, "Music Genre Classification Using Random Forest Model," vol. 12, no. 1, pp. 2654–5101, 2023.
- [48] T. Gori, A. Sunyoto, and H. Al Fatta, "Preprocessing Data dan Klasifikasi untuk Prediksi Kinerja Akademik Siswa," *Jurnal Teknologi Informasi dan Ilmu Komputer*, vol. 11, no. 1, pp. 215–224, Feb. 2024, doi: 10.25126/jtiik.20241118074.
- [49] C. Yunpeng, D. U. Zhihui, and S. Lifeng, "A Simulation System and Algorithm Evaluation for Clustered Streaming Media Server \*," 2010.
- [50] *2nd ISRITI 2019 proceeding : the 2nd International Seminar on Research of Information Technology and Intelligent Systems 2019 : "The future & challenges of extended intelligence" : Yogyakarta, Indonesia, 05-06 December 2019*, IEEE, 2019.
- [51] R. Reznikov and S. Turlakova, "IMPORTANCE OF MACHINE LEARNING AND DATA SCIENCE IN MODERN BUSINESS," *Efektivna ekonomika*, no. 5, May 2024, doi: 10.32702/2307-2105.2024.5.13.
- [52] S. M. Aslam, A. K. Jilani, J. Sultana, and L. Almutairi, "Feature Evaluation of Emerging E-Learning Systems Using Machine Learning: An Extensive Survey," 2021, *Institute of Electrical and Electronics Engineers Inc.* doi: 10.1109/ACCESS.2021.3077663.

- [53] F. Azuaje, "Witten IH, Frank E: Data Mining: Practical Machine Learning Tools and Techniques 2nd edition," *Biomed Eng Online*, vol. 5, no. 1, Dec. 2006, doi: 10.1186/1475-925x-5-51.
- [54] L. Alzubaidi *et al.*, "A Survey On Deep Learning Tools Dealing With Data Scarcity: Definitions, Challenges, Solutions, Tips, And Applications," *J Big Data*, vol. 10, no. 1, Dec. 2023, doi: 10.1186/s40537-023-00727-2.
- [55] Agung Teguh Wibowo Amais and Adi Susilo, "Principal Component Analysis-Based Data Clustering for Labeling of Level Damage Sector in Post-Natural Disasters," vol. XX, 2017, doi: 10.1109/ACCESS.2023.3275852.
- [56] B. Vrigazova, "The Proportion for Splitting Data into Training and Test Set for the Bootstrap in Classification Problems," *Business Systems Research*, vol. 12, no. 1, pp. 228–242, May 2021, doi: 10.2478/bsrj-2021-0015.
- [57] M. Shantal, Z. Othman, and A. A. Bakar, "Impact of Missing Data on Correlation Coefficient Values: Deletion and Imputation Methods for Data Preparation," *Malaysian Journal of Fundamental and Applied Sciences*, vol. 19, no. 6, pp. 1052–1067, Nov. 2023, doi: 10.11113/mjfas.v19n6.3098.
- [58] G. A. B. Suryanegara, Adiwijaya, and M. D. Purbolaksono, "Improved Classification Results in the Random Forest Algorithm for Detection of Diabetes Patients Using the Normalization Method," *Jurnal RESTI*, vol. 5, no. 1, pp. 114–122, Mar. 2021, doi: 10.29207/resti.v5i1.2880.
- [59] Y. I. Sulistya and C. Danuputri, "Analisis perbandingan Reduction Technique dengan metode Dimentional Reduction dan Cross Validation

- pada dataset breast cancer," *Indonesian Journal of Data and Science (IJODAS)*, vol. 3, no. 2, pp. 82–88, 2022.
- [60] F. Goltz and M. Sadakata, "Do you listen to music while studying? A portrait of how people use music to optimize their cognitive performance," *Acta Psychol (Amst)*, vol. 220, Oct. 2021, doi: 10.1016/j.actpsy.2021.103417.
- [61] H. Khabiri, M. N. Talebi, M. F. Kamran, S. Akbari, F. Zarrin, and F. Mohandesi, "Music-Induced Emotion Recognition Based on Feature Reduction Using PCA From EEG Signals," *Frontiers in Biomedical Technologies*, vol. 11, no. 1, pp. 59–68, Dec. 2024, doi: 10.18502/fbt.v11i1.14512.
- [62] Babu Kaji Baniya and Joonwhoan Lee, *Audio Feature Reduction and Analysis for Automatic Music Genre Classification*. San Diego: [Institute of Electrical and Electronics Engineers], 2014.
- [63] G. Mostafa, H. Mahmoud, T. Abd El-Hafeez, and M. E. ElAraby, "Feature Reduction For Hepatocellular Carcinoma Prediction Using Machine Learning Algorithms," *J Big Data*, vol. 11, no. 1, Dec. 2024, doi: 10.1186/s40537-024-00944-3.
- [64] L. A. ANGEL, D. J. POLZELLA, and G. C. ELVERS, "BACKGROUND MUSIC AND COGNITIVE PERFORMANCE 1,2," *Percept Mot Skills*, vol. 110, no. 3C, pp. 1059–1064, Jun. 2010, doi: 10.2466/04.11.22.pms.110.c.1059-1064.

- [65] L. Kiss and K. J. Linnell, "The effect of preferred background music on task-focus in sustained attention," *Psychol Res*, vol. 85, no. 6, pp. 2313–2325, Sep. 2021, doi: 10.1007/s00426-020-01400-6.
- [66] H. A. Salman, A. Kalakech, and A. Steiti, "Random Forest Algorithm Overview," *Babylonian Journal of Machine Learning*, vol. 2024, pp. 69–79, Jun. 2024, doi: 10.58496/bjml/2024/007.

## LAMPIRAN

