

KLASIFIKASI JENIS SUARA PENYANYI MENGGUNAKAN *PROBABILISTIC NEURAL NETWORK* (PNN)

Gogi Afrendo Inabuy¹, Franki Yusuf Bisilisin²
STIKOM Uyelindo Kupang
gogiinabuy23@gmail.com

ABSTRACT

Human voice is a sound produced when speaking, singing, laughing, and crying. In singing, each individual has unique voice characteristics, including variations in style, pitch, and vocal quality. Human voice types are generally divided into soprano, alto, tenor, and bass. However, identifying singers' voices still faces complex challenges. The main challenges include voice variability, overlapping characteristics between singers, and limited datasets. To overcome these, an accurate identification system is needed. This research uses 600 voice samples from 60 singers in the Talitakumi Pasir Panjang Church youth choir, with durations of 2-20 seconds in WAV format. The feature extraction method used is Mel-Frequency Cepstral Coefficients (MFCC), involving signal recording, preprocessing, signal segmentation, and Fourier transformation. For classification, Probabilistic Neural Network (PNN) is chosen due to its ability to generate probability distributions for each class. This research aims to develop a singer voice type classification system using PNN with MFCC feature extraction. It is hoped that this system can assist vocal trainers in classifying singer voice types more accurately, making a significant contribution to the field of voice and music analysis.

Keywords: Extraction, Classification, MFCC, PNN, Voice

ABSTRAK

Suara manusia adalah bunyi yang dikeluarkan saat berbicara, bernyanyi, tertawa, dan menangis. Dalam bernyanyi, setiap individu memiliki karakteristik suara unik, termasuk variasi gaya, nada, dan vokal. Jenis suara manusia umumnya terbagi menjadi sopran, alto, tenor, dan bass. Namun, identifikasi suara penyanyi masih menghadapi tantangan kompleks. Tantangan utama meliputi variabilitas suara, overlapping karakteristik antar penyanyi, dan keterbatasan dataset. Untuk mengatasi ini, dibutuhkan sistem identifikasi yang akurat. Penelitian ini menggunakan 600 sampel suara dari 60 penyanyi pada paduan suara pemuda Gereja Talitakumi Pasir Panjang, dengan durasi 2-20 detik dan format WAV. Metode ekstraksi fitur yang digunakan adalah Mel-Frequency Cepstral Coefficients (MFCC), melibatkan perekaman sinyal, preprocessing, pembagian sinyal, dan transformasi Fourier. Untuk klasifikasi, dipilih Probabilistic Neural Network (PNN) karena kemampuannya menghasilkan distribusi probabilitas untuk setiap kelas. Penelitian

ini bertujuan mengembangkan sistem klasifikasi jenis suara penyanyi menggunakan PNN dengan ekstraksi fitur MFCC. Diharapkan sistem ini dapat membantu pelatih vokal dalam mengklasifikasikan jenis suara penyanyi secara lebih akurat, memberikan kontribusi signifikan dalam bidang analisis suara dan musik.

Kata Kunci: Ekstraksi, Klasifikasi, MFCC, PNN, Suara

A. Pendahuluan

Suara manusia adalah bunyi yang dikeluarkan dari mulut manusia seperti pada saat berbicara, bernyanyi, tertawa dan menangis (Mu'tashim, 2020). Dalam hal bernyanyi, setiap manusia memiliki karakteristik jenis suara yang berbeda-beda termasuk variasi dalam gaya, nada dan vokal. Jenis suara manusia terdiri atas empat kategori yaitu sopran, alto, tenor dan bass. Namun melalui pendengaran manusia persepsi jenis suara terhadap setiap individu bisa diidentifikasi secara berbeda. Oleh karena itu dibutuhkan sebuah sistem untuk mengidentifikasi jenis suara penyanyi.

Teknologi yang digunakan untuk mengidentifikasi jenis suara adalah salah satu penerapan kecerdasan buatan yang diperlukan untuk mengenali suara penyanyi berdasarkan sistem yang dibuat. Adapun penelitian oleh Svava (2017),

yang mengembangkan klasifikasi jenis suara dengan menggunakan *hidden markov models* (HMM) yang dibangun berdasarkan fitur *mel-frequency cepstral coefficients* (MFCC) menghasilkan HMM dengan fitur MFCC dengan akurasi terbaik yaitu 70% diikuti dengan akurasi 66,67%. Namun, meski kemajuan ini signifikan, identifikasi suara penyanyi masih menghadapi berbagai tantangan kompleks yang perlu diatasi. Salah satu tantangan utama adalah variabilitas dalam suara penyanyi itu sendiri. Suara seseorang dapat berubah seiring waktu karena berbagai faktor seperti usia, kondisi kesehatan, atau perubahan dalam teknik bernyanyi. Sistem identifikasi suara perlu mampu beradaptasi dengan variasi- variasi ini untuk tetap akurat. Tantangan lain muncul dari overlapping karakteristik suara antar penyanyi. Beberapa penyanyi mungkin memiliki kualitas vokal yang sangat mirip, membuat proses identifikasi menjadi lebih sulit. Hal ini

diperparah oleh keterbatasan dataset suara penyanyi yang besar dan beragam. Mengatasi semua tantangan ini dibutuhkan suatu kemajuan dalam pemrosesan sinyal, pembelajaran mesin, dan pemahaman mendalam tentang karakteristik vokal manusia. Proses ekstraksi fitur dengan menggunakan MFCC dalam analisis suara dimulai dengan merekam sinyal suara dan melakukan *preprocessing* untuk menghilangkan *noise*. Sinyal kemudian dibagi menjadi *frame-frame* kecil, diikuti oleh transformasi *fourier* untuk menganalisis spektrum frekuensi. Fitur-fitur ini dapat digunakan untuk tugas klasifikasi atau analisis suara lainnya, karena kemampuannya dalam menangkap informasi kunci terkait dengan vokal dan struktur suara.

klasifikasi adalah proses pengelompokan yang dilakukan secara sistematis pada sejumlah objek atau gagasan ke dalam kelas-kelas atau kelompok-kelompok tertentu berdasarkan berbagai ciri yang sama (Hamakonda dan Tairas, 1983). Proses klasifikasi bisa menggunakan banyak metode jaringan syaraf tiruan atau biasa

disebut *neural network*, salah satunya metode jaringan syaraf tiruan yaitu *probabilistic neural network* (PNN). PNN merupakan salah satu model pada Jaringan Syaraf Tiruan (JST) yang digunakan untuk pengklasifikasian (Adyati, et.al., 2018). PNN muncul sebagai salah satu perwujudan khusus dari *artificial*

neural network (ANN) metode ini menonjol karena integrasi konsep probabilitas dalam arsitektur ANN. Sebagai alat klasifikasi PNN tidak hanya menilai pola pada data tetapi juga menghasilkan distribusi probabilitas untuk masing-masing kelas. Hal ini memperkaya kemampuan PNN untuk menangani ketidakpastian dalam data dan memberikan tingkat keyakinan terhadap hasil klasifikasi. Penerapan PNN di dalam ANN menunjukkan *respons* terhadap kebutuhan sistem yang tidak sekedar dapat mengenali pola tetapi juga mengukur tingkat kepastian dalam pengambilan keputusan untuk klasifikasi.

Pada penelitian ini menggunakan *probabilistic neural network* (PNN) untuk mengenali jenis suara penyanyi menggunakan ekstraksi MFCC. Penelitian ini diharapkan mampu membuat sistem untuk mengidentifikasi jenis suara penyanyi berdasarkan kelas inputan suara.

B. Metode Penelitian

1. Lokasi dan Waktu Penelitian

Penelitian ini dilakukan pada paduan suara pemuda GMIT Talitakumi Pasir Panjang. Penelitian dilaksanakan mulai Januari 2024 sampai Juni 2024.

2. Bahan dan Alat Penelitian

2.1. Bahan Penelitian

Bahan penelitian berupa data suara penyanyi yang direkam dari 30 anggota paduan suara laki-laki dan perempuan pada paduan suara pemuda GMIT Talitakumi Pasir Panjang, Kota Kupang.

2.2. Alat Penelitian

Peralatan penelitian terdiri dari perangkat keras dan perangkat lunak sebagai berikut:

1. Perangkat Keras (Hardware)

- a. Laptop Asus
- b. Intel Core i7-8750H CPU @ 2.20 GHz
- c. RAM 8 GB
- d. Microphone

2. Perangkat Lunak (Software)

- a. Windows 11
- b. Microsoft Word 2013
- c. Microsoft Excel 2013
- d. Adobe Audition 2021
- e. Python
- f. Draw.io

3. Prosedur Penelitian

Penelitian dilakukan melalui beberapa tahapan, mulai dari pengumpulan data suara hingga evaluasi model klasifikasi. Adapun tahapan penelitian sebagai berikut:

3.1. Data Suara

Data suara diperoleh dari 30 anggota paduan suara pemuda GMT Talitakumi Pasir Panjang. Setiap responden menghasilkan 10 sampel suara dengan durasi 2–10 detik sehingga total data yang digunakan sebanyak 300 sampel suara.

3.2. Praproses

Pada tahap ini dilakukan pengurangan noise menggunakan teknik filtering agar suara menjadi lebih jernih dan bersih dari gangguan suara latar belakang.

3.3. Ekstraksi MFCC

Data suara diekstraksi menggunakan metode Mel Frequency Cepstral Coefficients (MFCC) untuk memperoleh karakteristik utama sinyal suara berdasarkan spektrum frekuensi dan durasi sinyal.

3.4. Pembagian Data

Data dibagi menggunakan metode K-Fold Cross Validation

dengan komposisi 80% data latih dan 20% data uji.

3.5. Probabilistic Neural Network (PNN)

Data latih diproses menggunakan algoritma Probabilistic Neural Network (PNN) yang terdiri dari input layer, pattern layer, summation layer, dan decision layer.

3.6. Pengujian

Tahap pengujian dilakukan menggunakan confusion matrix untuk mengevaluasi performa model berdasarkan nilai sensitivitas, spesifisitas, dan akurasi.

C. Hasil Penelitian dan Pembahasan Implementasi Sistem

Implementasi sistem dilakukan untuk menggambarkan bagaimana proses pengolahan data suara dilakukan mulai dari tahap pembacaan data hingga pembentukan data latih dan data uji. Pada tahap awal, sistem membaca data *audio* dari 600 *file audio* yang digunakan, yang terdiri dari enam jenis suara manusia. Setiap *file audio* diproses secara berulang untuk diubah menjadi sinyal digital sebagai dasar dalam pengolahan selanjutnya. Proses pembacaan

sinyal audio dilakukan menggunakan library Librosa, sebagaimana ditunjukkan pada potongan kode berikut.

```
signal, sr = librosa.load(path, sr=None)
```

Data audio WAV

Sinyal suara yang dihasilkan merepresentasikan amplitudo terhadap waktu dalam bentuk array, sedangkan sampling rate menunjukkan jumlah sampel yang direkam dalam setiap detik. Informasi tersebut menjadi dasar dalam menentukan panjang durasi *audio* yang diproses oleh sistem. Durasi *audio* diperoleh melalui perbandingan antara jumlah sampel terhadap nilai *sampling rate*. Dalam implementasinya, perhitungan tersebut dilakukan sebagai berikut

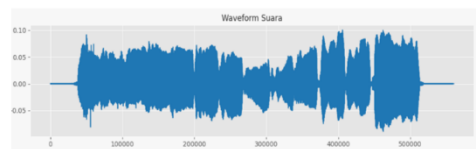
```
duration = len(signal) / sr
```

Sebagai contoh, apabila suatu *file audio* yang dipilih memiliki jumlah sampel sebanyak 576.000 dengan sampling rate sebesar 48.000 Hz, maka durasi audio dapat dihitung sebagai berikut.

$$\frac{576000}{48000} = 12 \text{ detik}$$

Hasil tersebut menunjukkan bahwa sinyal suara yang diproses memiliki panjang waktu selama 12 detik. Dengan demikian, jumlah sampel yang dihasilkan dalam bentuk *array* secara langsung merepresentasikan keseluruhan durasi sinyal suara. Untuk memberikan gambaran awal terhadap karakteristik sinyal suara, data suara divisualisasikan dalam bentuk *waveform* yang menunjukkan perubahan amplitudo terhadap waktu. Visualisasi WAV ditunjukan pada gambar 14.

```
plt.plot(signal)
```



Visualisasi tersebut menunjukkan pola gelombang suara sebagai representasi awal dari data audio yang telah dibaca oleh sistem. Berdasarkan informasi yang diperoleh, data suara tersebut selanjutnya diproses pada tahap ekstraksi fitur untuk memperoleh karakteristik yang lebih representatif.

Ekstarksi fitur

Setelah sinyal suara diperoleh, tahap berikutnya adalah ekstraksi fitur untuk

mengubah data suara menjadi representasi numerik yang dapat digunakan dalam proses klasifikasi. Sinyal diolah secara bertahap sehingga karakteristik utama suara dapat ditangkap dengan lebih baik. Dilakukan dengan penguatan sinyal untuk menonjolkan komponen frekuensi yang berperan penting dalam membentuk karakter suara. Hasil dari proses ini berupa sinyal yang telah diperkuat, yang kemudian digunakan sebagai masukan pada tahap berikutnya.

Proses tersebut diimplementasikan dalam program menggunakan fungsi *filter_noise* yang menerapkan metode *pre-emphasis* melalui library *Librosa*

```
signal= librosa.effects.preemphasis(signal)
```

Pada tahap ini, sinyal suara diolah secara bertahap sehingga karakteristik utama dari suara dapat ditangkap dengan lebih jelas. Proses diawali dengan penguatan sinyal (*pre-emphasis*) untuk menonjolkan komponen frekuensi yang penting. Sinyal kemudian dibagi menjadi beberapa bagian kecil (*framing*) agar perubahan karakter suara dapat diamati secara lebih detail pada setiap

segmen. Selanjutnya, setiap bagian sinyal diproses untuk mengurangi ketidakaturan pada batasnya (*windowing*), sehingga menghasilkan bentuk sinyal yang lebih stabil. Sinyal yang telah diproses kemudian diubah ke domain frekuensi melalui proses FFT, sehingga diperoleh informasi mengenai kandungan frekuensi dalam suara. Informasi tersebut selanjutnya dipetakan ke dalam skala yang menyerupai persepsi pendengaran manusia (*Mel filter bank*), sehingga energi pada setiap frekuensi dapat direpresentasikan dengan lebih baik. Nilai energi yang dihasilkan kemudian diubah ke dalam bentuk logaritmik (*log energy*) untuk menyesuaikan dinamika suara, dan selanjutnya diproses menggunakan DCT sehingga menghasilkan koefisien MFCC. Dalam implementasinya, seluruh rangkaian proses tersebut dilakukan secara otomatis menggunakan library *Librosa*, sehingga sistem dapat langsung menghasilkan nilai koefisien MFCC dari sinyal suara yang diproses.

Proses ekstraksi fitur dilakukan menggunakan fungsi *extract_mfcc*

yang memanfaatkan *library Librosa* untuk menghasilkan 13 koefisien MFCC. Hasil dari proses tersebut berupa sekumpulan nilai koefisien pada setiap frame. Untuk memperoleh satu representasi dari keseluruhan sinyal, nilai pada setiap frame diringkas dengan menghitung nilai rata-rata (*mean*) dan standar deviasi (*standard deviation*) dari masing-masing koefisien. Dengan demikian diperoleh total 26 fitur yang terdiri dari 13 nilai rata-rata dan 13 nilai standar deviasi.

```
mean = np.mean(mfcc_features.T, axis=0)
std = np.std(mfcc_features.T, axis=0)
features = np.concatenate([mean, std])
```

Sebagai ilustrasi pada salah satu data suara yang dipilih, nilai koefisien MFCC diperoleh dari hasil perhitungan pada setiap frame yang kemudian dirata-ratakan. Misalnya pada koefisien ke-13, diperoleh beberapa nilai dari tiap frame sebagai berikut:

Frame 1: -6.8

Frame 2: -7.5

Frame 3: -6.9

Frame 4: -7.1

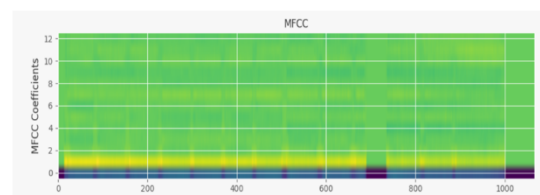
Frame 5: -7.3

Nilai-nilai tersebut diringkas untuk memperoleh satu nilai representatif:

$$(-6.8 + -7.5 + -6.9 + -7.1 + -7.3) / 5 = -7.12$$

Hasil tersebut menjadi nilai akhir untuk koefisien ke 13.

Selain nilai rata-rata, pada penelitian ini juga dihitung nilai standar deviasi dari setiap koefisien MFCC untuk menangkap variasi sinyal suara. Dengan demikian, setiap koefisien tidak hanya direpresentasikan oleh nilai rata-rata, tetapi juga oleh tingkat penyebarannya, sehingga total fitur yang digunakan dalam proses klasifikasi berjumlah 26.



Gambar 2. Grafik MFCC

Visualisasi tersebut menunjukkan distribusi nilai koefisien MFCC terhadap waktu dalam bentuk *heatmap*. Perbedaan warna menunjukkan besar kecilnya nilai koefisien, di mana warna yang lebih terang merepresentasikan nilai yang

lebih tinggi, sedangkan warna yang lebih gelap menunjukkan nilai yang lebih rendah. Hasil akhir dari proses ini berupa vektor fitur yang telah diringkas menjadi 26 fitur yang terdiri dari nilai rata-rata dan standar deviasi dari 13 koefisien MFCC. Nilai tersebut digabungkan menjadi satu vektor fitur yang digunakan sebagai input pada proses klasifikasi.

Pembagian data

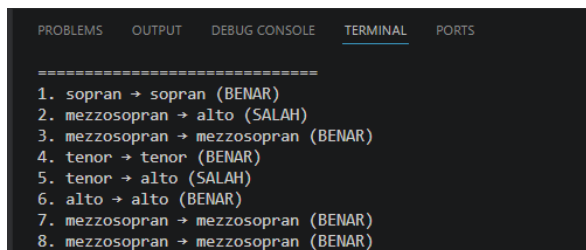
Dataset yang digunakan terdiri dari 600 data suara yang berhasil diproses, Setelah proses ekstraksi fitur selesai dilakukan, seluruh data yang telah direpresentasikan dalam bentuk fitur selanjutnya disiapkan untuk proses pelatihan dan pengujian model. Pada tahap awal, data dibagi menjadi dua bagian, yaitu data latih dan data uji. Pembagian dilakukan dengan perbandingan 80% sebagai data latih dan 20% sebagai data uji. Dari proses tersebut diperoleh sebanyak 480 data latih dan 120 data uji. Data latih digunakan untuk membentuk model, sedangkan data uji digunakan untuk melihat kemampuan model dalam

mengklasifikasikan data yang belum pernah diproses sebelumnya. Untuk memperoleh gambaran performa model yang lebih menyeluruh, proses pembagian data tidak hanya dilakukan satu kali, tetapi diulang menggunakan metode *StratifiedGroupKFold* sebanyak 5 fold. Pada setiap iterasi, sebagian data digunakan sebagai data latih dan sebagian lainnya sebagai data uji, kemudian dilakukan pengujian secara bergantian hingga seluruh data mendapatkan kesempatan sebagai data uji. Pembagian data pada setiap fold menghasilkan jumlah data uji yang relatif sama, yaitu 120 data, sedangkan sisanya digunakan sebagai data latih. Proses ini dilakukan secara konsisten sehingga hasil yang diperoleh pada setiap percobaan tetap sama ketika dijalankan ulang.

Pemodelan

Pada tahap pemodelan, klasifikasi dilakukan menggunakan metode Probabilistic Neural Network (PNN). Data uji yang telah direpresentasikan dalam bentuk fitur MFCC dibandingkan dengan seluruh

data latih pada setiap kelas. Proses ini diawali dengan menghitung jarak antara data uji dan data latih menggunakan jarak *Euclidean*, kemudian nilai jarak tersebut diubah menjadi nilai probabilitas melalui fungsi *Gaussian*. Nilai probabilitas pada setiap kelas selanjutnya dirata-ratakan, dan kelas dengan nilai tertinggi dipilih sebagai hasil klasifikasi. Berdasarkan hasil proses pada program, salah satu data uji menghasilkan prediksi sebagai berikut pada gambar dibawah ini.



```
=====
1. sopran → sopran (BENAR)
2. mezzosopran → alto (SALAH)
3. mezzosopran → mezzosopran (BENAR)
4. tenor → tenor (BENAR)
5. tenor → alto (SALAH)
6. alto → alto (BENAR)
7. mezzosopran → mezzosopran (BENAR)
8. mezzosopran → mezzosopran (BENAR)
```

Hasil tersebut menunjukkan bahwa

```
cm = confusion_matrix(y_ts, y_pred, labels=model_k.classes)
akurasi = np.trace(cm) / np.sum(cm)

tp = np.diag(cm)

sensitivitas = np.mean(tp / np.sum(cm, axis=1))

spesifisitas = np.mean(
    (np.sum(cm) - np.sum(cm, axis=0) - np.sum(cm, axis=1) + tp) /
    (np.sum(cm) - np.sum(cm, axis=1))
)
```

data uji memiliki tingkat kemiripan paling tinggi terhadap kelas mezzo sopran dibandingkan kelas lainnya, Sebagai ilustrasi dilakukan perhitungan probabilitas pada masing-masing kelas berdasarkan

rata-rata hasil fungsi Gaussian terhadap data latih pada setiap kelas

Pengujian Sistem

Pengujian sistem dilakukan untuk mengetahui kemampuan model dalam mengklasifikasikan jenis suara berdasarkan fitur MFCC yang telah diperoleh pada tahap sebelumnya. Data hasil pembagian fold pada proses sebelumnya digunakan sebagai data latih dan data uji untuk menghasilkan confusion matrix dan parameter evaluasi model. Proses evaluasi dilakukan menggunakan metode 5-fold cross validation, dimana pada setiap fold digunakan 480 data sebagai data latih dan 120 data sebagai data uji.

Confusion matrix digunakan untuk mengetahui jumlah prediksi benar dan salah pada setiap kelas suara. Nilai diagonal utama confusion matrix menunjukkan jumlah data yang berhasil diklasifikasikan dengan benar pada masing-masing kelas. Nilai tersebut kemudian digunakan dalam perhitungan akurasi, sedangkan sensitivitas dan spesifisitas digunakan untuk melihat kemampuan model dalam mengenali dan membedakan

setiap kelas suara pada masing-masing *fold* pengujian. Hasil pengujian pada setiap *fold* ditunjukkan melalui *confusion matrix* yang diperoleh dari proses klasifikasi data uji pada masing-masing *fold*. *Fold* pertama ditunjukkan pada

Analisis Kelebihan dan Kekurangan Sistem

Setelah dilakukan proses pengujian sistem menggunakan metode MFCC untuk ekstraksi fitur dan PNN untuk proses klasifikasi, dapat diketahui beberapa kelebihan dan keterbatasan dari sistem yang dikembangkan. Analisis ini dilakukan berdasarkan hasil pengujian sistem serta evaluasi performa model yang telah dilakukan pada dataset penelitian.

Kelebihan sistem

Sistem yang dibangun mampu melakukan klasifikasi jenis suara manusia menggunakan fitur MFCC dan metode PNN. Berdasarkan hasil pengujian menggunakan metode 5-fold cross validation, sistem memperoleh rata-rata akurasi sebesar 70.83%,

sehingga model mampu mengenali pola suara pada masing-masing kelas dengan cukup baik. Sistem juga mampu melakukan proses klasifikasi secara otomatis melalui antarmuka GUI. Pengguna hanya perlu memilih file audio, kemudian sistem akan melakukan proses ekstraksi fitur MFCC dan klasifikasi menggunakan metode PNN hingga menghasilkan jenis suara yang dikenali. Hal tersebut memudahkan pengguna dalam melakukan proses pengenalan jenis suara manusia tanpa melakukan perhitungan secara manual. Selain itu, sistem mampu mengenali karakteristik kemiripan antar jenis suara.

Beberapa kesalahan klasifikasi yang terjadi masih berada pada kelas suara yang memiliki karakteristik berdekatan, seperti alto dengan mezzosopran atau bass dengan baritone. Hal tersebut menunjukkan bahwa metode PNN masih mampu mengenali pola kemiripan karakteristik suara berdasarkan fitur MFCC yang digunakan.

Kekurangan sistem

Sistem yang dibangun masih memiliki beberapa keterbatasan dalam proses klasifikasi jenis suara. Kesalahan klasifikasi masih dapat terjadi pada jenis suara yang memiliki karakteristik frekuensi dan pola suara yang berdekatan. Kondisi tersebut menyebabkan beberapa data suara dapat diklasifikasikan ke kelas lain yang memiliki kemiripan karakteristik suara. Hasil klasifikasi juga dipengaruhi oleh kualitas rekaman audio yang digunakan. Perbedaan pengucapan vokal, noise pada rekaman, serta kondisi mikrofon dapat mempengaruhi hasil ekstraksi fitur MFCC sehingga berdampak pada hasil klasifikasi sistem. Selain itu, metode PNN cukup sensitif terhadap perubahan fitur suara yang diperoleh dari data audio. Perubahan kecil pada karakteristik suara dapat mempengaruhi jarak antar fitur sehingga hasil klasifikasi dapat berubah pada beberapa kondisi pengujian tertentu.

D. Kesimpulan

Berdasarkan hasil penelitian yang telah dilakukan, sistem

klasifikasi jenis suara manusia menggunakan metode MFCC dan PNN berhasil diimplementasikan. Sistem mampu melakukan proses klasifikasi jenis suara secara otomatis berdasarkan *file audio* yang dipilih oleh pengguna. Proses pengujian menggunakan metode *5-fold cross validation* menghasilkan akurasi sebesar 0.8000, sensitivitas sebesar 0.8000, dan spesifisitas sebesar 0.9600 dari *fold* terbesar yaitu *fold* empat dan keseluruhan nya menghasilkan rata-rata akurasi 0.7063, sensitivitas sebesar 0.7083 dan sensitivitas sebesar 0.9417. Hasil tersebut menunjukkan bahwa metode MFCC dan PNN mampu mengenali pola karakteristik suara manusia dengan cukup baik. Berdasarkan hasil pengujian yang dilakukan, beberapa kesalahan klasifikasi masih terjadi pada jenis suara yang memiliki karakteristik berdekatan, seperti alto dengan mezzosopran atau *bass* dengan *baritone*. Meskipun demikian, sistem tetap mampu mengenali kemiripan pola suara antar kelas berdasarkan fitur MFCC yang digunakan. Selain itu, GUI yang

dibangun juga mampu mempermudah pengguna dalam melakukan proses pengujian dan klasifikasi file audio secara langsung.

DAFTAR PUSTAKA

- Adyati, R. D., Nasution, Y. N. dan Wahyuningsih, S. 2019. Klasifikasi *Probabilistic Neural Network* (PNN) pada data diagnosa penyakit demam berdarah dengue (DBD) tahun 2018. Di dalam: Seminar Nasional Matematika dan Statistika. (15 mei 2019): Samarinda (ID): 2657-232X. **[internet]. [di akses 6 January 2014]**. Tersedia pada: <https://jurnal.fmipa.unmul.ac.id/index.php/SNMSA/article/view/521>
- Azizah, P.D.I. 2022. Penerapan Probabilistic Neural Network Pada Klasifikasi Berat Bayi Baru Lahir. Jurnal Riset Statistika **[internet]. [diakses 6 Januari 2024]**. 1(2):152-159. Tersedia pada: <https://doi.org/10.29313/jrs.v1i2.524>
- Bueno, A. 2009. Representasi niali HOS dan model MFCC sebagai ekstraksi ciri pada system identifikasi pembicara dilingkungan bernoise menggunakan HMM **[Disertasi]**. Depok (ID): Universitas Indonesia. **[internet] [diakses 10 Januari 2024]** Tersedia pada: <https://lib.ui.ac.id/file?file=digit>
- al/2016- 6/20426547-D958-Agus%20Buono.pdf
- Fathurohman, A. 2021. Machine Learning Untuk Pendidikan: Mengapa Dan Bagaimana. Jurnal Informatika dan Teknologi Komputer **[Internet]. [diakses 6 Januari 2024]**. 1(3):2809-9230. Tersedia pada: <https://journal.amikveteran.ac.id/index.php/jitek/article/view/306/238>
- Hamakonda, T.P., dan Tairas, J.N.P. 1983. Pengantar Klasifikasi persepuluhan deway. Jakarta (ID): Gunung Mulia. Tersedia pada : <https://balaiyanpus.jogjaprovg.go.id/opac/detail-opac?id=273016>
- Heriyanto, H., Hartati, S. dan Putra, A.E. 2018. Ekstraksi Ciri Mel Frequency Cepstral Coefficient (MFCC) dan Rerata Coefficient Untuk Pengecekan Bacaan Al-qur'an. Jurnal TELEMATIKA **[internet]. [diakses 6 Januari 2024]**. 15(2):99-108. Tersedia pada: <http://jurnal.upnyk.ac.id/index.php/telematika/article/view/3123/2455>
- Husna, F. B. 2021. Implementasi Data Mining Menggunakan Algoritma C4.5 Pada Klasifikasi Penjualan Hijab. **[Skripsi]**. Malang (ID): Matematika, Universitas Islam Negeri Maulana Malik Ibrahim. **[internet] [diakses 10 Januari 2024]** Tersedia pada: <http://etheses.uin-malang.ac.id/32882/1/17610069.pdf>

- Irmansyah, D., Lokatara, B. A., Saputra, I. M., Faradillah, S., Wulansari, A. 2023. Analisis Perkembangan Artificial Intelligence Dalam Bidang Bisnis: Systematic Literature Review. *Jurnal Teknologi Informasi* [internet]. [diakses 6 Januari 2024]. 4(2):2748–8546. Tersedia pada: <https://jurnal.dharmawangsa.ac.id/index.php/djtechno/article/view/3404/pdf>
- Irwansyah, R. 2021. Deteksi Kepribadian Seseorang Berdasarkan Pola Gambar Pohon Dengan Metode *Probabilistic Neural Network* [Skripsi]. Bandung (ID): Teknik Informatika, Universitas Komputer Indonesia. [internet] [diakses 4 Januari 2024] Tersedia pada: <https://elibrary.unikom.ac.id/id/eprint/4956/>
- Karsito, K. dan Susanti, S. 2021. Klasifikasi Kelayakan Peserta Pengajuan Kredit Rumah Dengan Algoritma Naïve Bayes Di Perumahan Azzura Residencia. *Journal Information Thecnology* [Internet]. [diakses 6 Januari 2024]. 9(3):2407–3903. Tersedia pada: <https://www.jurnal.pelitabangsa.ac.id/index.php/sigma/article/view/509>
- Mardiana, L. Kusnandar, D. dan Satyahadewi, N. 2022. Analisis Diskriminan Dengan K Fold Cross Validation Untuk Klasifikasi Kualitas Air Di Kota Pontianak. *Jurnal Ilmiah Berkala* [Internet]. [diakses 6 Januari 2024]. 11(1):2302-9854. Tersedia pada: <https://jurnal.untan.ac.id/index.php/jbmstr/article/view/51608>
- Mu'tashim, Naufal. 2020. Identifikasi Jenis Suara Berdasarkan Jangkauan Vokal Menggunakan Algoritma *Linear Predictive Coding* dan *K-Nearest Neighbor* [Skripsi]. Bandung (ID): Teknik Informatika, Institut Teknologi Nasional Bandung. [internet] [diakses 4 Januari 2024] Tersedia pada: <https://eprints.itenas.ac.id/1331/>
- Pratama, K. B. 2021. Klasifikasi Jenis Vokal Manusia menggunakan MFCC dan Convolutional Neural Network. [Skripsi]. Jakarta (ID): Teknik Informatika, Universitas Telkom. [internet] [diakses 10 Januari 2024] Tersedia pada: <https://repository.telkomuniversity.ac.id/pustaka/172877/klasifikasi-jenis-vokal-manusia-menggunakan-mfcc-dan-convolutional-neural-network.html>
- Putri, V.K.M. dan Nailufar, N.N. 2022. Jenis Suara Manusia: Sopran, Alto, Tenor, Baritone dan Bass. Kompas.com. [internet]. [diakses 6 Januari 2024]. Tersedia pada: <https://www.kompas.com/skola/read/2021/01/22/170114269/jenis-suara-manusia-sopran-alto-tenor-baritone-dan-bass?page=all>
- Svara, A. A. W. 2017. Implementasi Pengenalan Pola untuk Klasifikasi Jenis Suara

Penyanyi dengan Tinjauan Audio **[Skripsi]**. Yogyakarta (ID): Ilmu Komputer, Universitas Gadjah Mada. **[internet] [diakses 4 Januari 2024]** Tersedia pada: <https://etd.repository.ugm.ac.id/penelitian/detail/112333>

Umam, A. K. 2023. Augmentasi data pada model klasifikasi jenis vokal menggunakan CNN dengan fitur ekstraksi mel spectrogram **[Skripsi]**. Jakarta (ID): Teknik Informatika, UIN Syarif Hidayatullah Jakarta. **[internet] [diakses 10 Januari 2024]** Tersedia pada: <https://repository.uinjkt.ac.id/dspace/bitstream/123456789/75599/1/AHMA-D%20KHAIRUL%20UMAM-FST.pdf>

Umar, R., Riadi, I. dan Hanif, A. 2018. Analisis Bentuk Pola Menggunakan Ekstraksi Ciri Mel-Frequency Cepstral Coefficients (MFCC). *Journal Cogito Smart* **[internet]**. **[diakses 6 Januari 2024]**. 4(2): 294–304. Tersedia pada: <https://cogito.unklab.ac.id/index.php/cogito/article/view/130/93>