

Klasifikasi Usia Berdasarkan Suara Dengan Ekstraksi Ciri Mel Frequency Cepstral Coefficients Menggunakan Support Vector Machine

Mufi Oktaviani¹, Taufik Edy Sutanto², Mahmudi³

^{1,2,3}Program Studi Sains dan Teknologi Matematika, Universitas Islan Negeri Syarif Hidayatullah Jakarta, Indonesia

E-mail: ¹mufi.oktaviani17@mhs.uinjkt.ac.id, ²taufik.sutanto@uinjkt.ac.id,

³mahmud.mathlovers@gmail.com

Abstract

Each individual has a different and unique voice, this can potentially be used to determine a person's age. Classification can be done through a feature extraction process using Mel Frequency Cepstral Coefficients (MFCC). Then the extraction results can be used as features in machine learning models such as Support Vector Machines. Data was collected by recording 100 voices of informants, then the feature extraction process was carried out using MFCC, and then classified using Support Vector Machine. Age categories set are children (0-14), teenagers (15-40) and elderly (40+). Then the data samples were duplicated for variety using noise, shift and dynamic change methods so that 6000 sound samples were obtained. Our model can result in an accuracy of 75,8% and the recall obtained in the children category is a percentage of 80%. Researchers hope that these results can be a reference for human age recognition application systems for voice-based classification.

Keywords: Mel Frequency Cepstral Coefficients, Support Vector Machine, Voice

Abstrak

Setiap individu memiliki suara yang berbeda dan unik, hal ini berpotensi untuk digunakan dalam menentukan usia seseorang. Klasifikasi suara dapat dilakukan melalui proses ekstraksi ciri salah satunya dengan menggunakan Mel Frequency Cepstral Coefficients. Kemudian hasil ekstraksi dapat digunakan sebagai feature pada model machine learning seperti Support Vector Machines. Dalam penelitian ini, 100 sampel suara dengan kategori usia anak-anak (0-14 tahun), remaja (15-40 tahun) dan usia lanjut (40+ tahun) digunakan sebagai training data awal. Kemudian sampel data diduplikasi agar beragam menggunakan metode noise, shift dan dynamic change sehingga didapatkan 6000 sampel suara baru. Model kami dapat menghasilkan rata-rata akurasi sebesar 75,8% dengan tingkat keberhasilan sistem dalam menemukan kembali yang tertinggi terdapat pada kategori anak-anak dengan persentase 80%. Peneliti berharap dengan hasil ini dapat menjadi referensi untuk sistem aplikasi pengenalan manusia untuk mengklasifikasi usia berbasis suara.

Keywords: Mel Frequency Cepstral Coefficients, Suara, Support Vector Machine

1. Pendahuluan

Salah satu jenis data yang dapat dihasilkan dari tubuh manusia yang dapat mencerminkan informasi karakteristik pemiliknya adalah suara. Usia, jenis kelamin, aksen, dan aspek lain dari identitas individu dapat ditentukan dari suara yang dihasilkan [1]. Dari suara ini kita dapat mengenali jenis kelamin bahkan rentang usia si pemilik suara, karena suara dapat berfluktuasi seiring bertambahnya usia dari anak muda hingga orang tua. Automatic Speech Recognition (ASR) merupakan bidang studi yang berkembang relatif cepat dan pesat seiring dengan kemajuan teknologi [2]. Ilustrasi dasar tentang bagaimana aplikasi ASR telah memengaruhi kehidupan sehari-hari adalah perangkat lunak speech to text atau asisten virtual. Penelitian ini terinspirasi dari salah satu bidang studi yang berhubungan dengan interaksi manusia dan mesin dengan media ucapan seperti Google Assistant, Alexa, Cortana atau beberapa macam aplikasi

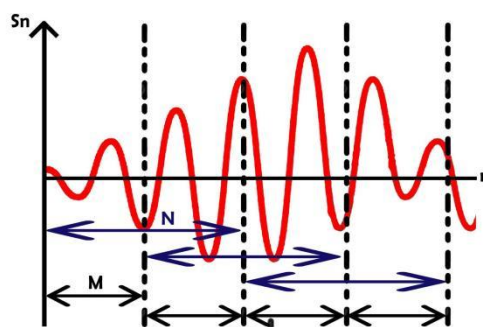
ASR lainnya [3]. Faktor-faktor yang membatasi sistem pengenalan suara seseorang mencakup ketidakmampuan mengenali fitur yang sensitif dan tidak cukup kuat untuk menjelaskan perbedaan artikulasi pembicara, variasi prosodik, dan perbedaan saluran ucapan [4]. Proses ekstraksi ciri dalam paper ini menggunakan metode Mel Frequency Cepstral Coefficients (MFCC). Metode MFCC merupakan metode ekstraksi ciri untuk mendapatkan cepstral coefficient dan frame sehingga dapat digunakan untuk proses speech recognition agar didapatkan akurasi yang lebih baik untuk pengklasifikasian [5]. Metode ini digunakan untuk melakukan ekstraksi fitur, yaitu suatu proses yang mengubah sinyal suara menjadi variabel atau peubah.

Terdapat beberapa pendekatan klasifikasi usia berdasarkan suara yang telah dilakukan sebelumnya. Dalam salah satu penelitian menggunakan MFCC dengan model K-Nearest Neighbor (KNN) untuk mengklasifikasi data suara, dengan data suara yang sudah dinormalisasikan sebanyak 40 data suara, telah dihasilkan akurasi maksimal sebesar 55,55% [6]. Pada penelitian lainnya, model Random Forest digunakan pada data dengan dua kategori usia yakni anak-anak dan dewasa yang menghasilkan akurasi terbaik sebesar 88,82% [7]. Pada paper ini, data suara yang digunakan bukanlah data suara orang asing seperti yang dilakukan dalam paper sebelumnya [6][7] dan kategori yang digunakan juga lebih banyak serta menggunakan metode yang berbeda, yaitu Support Vector Machine (SVM). Kode dan data pada paper ini dapat diakses dan divalidasi secara terbuka untuk pembaca yang berminat di <https://github.com/MufiOktaviani/age-classification-by-voice>.

2. Metodologi Penelitian

2.1. Mel Frequency Cepstrum Coefficients

Metode MFCC memiliki beberapa tahapan, yaitu pre-emphasis, frame blocking, windowing, Fast Fourier Transform, filter bank, dan DCT [9]. Pre-emphasis adalah tahapan memfilter frekuensi suara yang dimana dalam prosesnya frekuensi suara yang tinggi pada sebuah spektrum akan dipertahankan, ini bertujuan untuk menyeimbangkan spektrum sinyal suara [10]. Selanjutnya, proses frame blocking adalah proses framing pada sinyal suara dimana sinyal akan dibagi menjadi beberapa frame dengan masing-masing frame berisi N sampel sinyal yang berdekatan dan frame yang dipisahkan oleh ruang M sampel seperti yang dicontohkan pada Gambar 1. Setiap panjang frame dibagi dengan waktu 20 ms sampai 40 ms, dengan N adalah jumlah sampel dan M adalah panjang frame [11].



Gambar 1. Proses frame blocking

Proses windowing ini bertujuan untuk mengurangi efek diskontinu pada ujung-ujung frame, pada tool librosa, proses ini menggunakan fungsi windowing yang disebut dengan hamming window. Fungsi hamming window ini tidak menghasilkan sidelobe yang berlebihan dan noise yang dihasilkan tidak terlalu besar [12]. Selanjutnya, proses Fast Fourier Transform atau FFT adalah pengembangan dari algoritma Discrete Fourier Transform (DFT) yang digunakan untuk mengubah sinyal yang berawal dari domain waktu menjadi domain frekuensi. Setelah berubah menjadi domain frekuensi dilakukan filter untuk mengetahui ukuran energi dimana proses ini disebut dengan filter bank. Langkah terakhir dari proses MFCC ini yaitu Discrete Cosine Transform (DCT). Konsep dasar dari DCT adalah mengkorelasikan frekuensi sehingga menghasilkan representasi yang baik dari spectral lokal [5].

2.2. Ekstraksi Ciri MFCC

Pada proses ekstraksi ciri ini akan melalui tahapan pre-emphasis yang menggunakan nilai konstanta filter (α) sebesar 0,97, pada tahap pre-emphasis ini satu suara akan menghasilkan sebanyak 44100 frame. Kemudian pada tahapan frame blocking, data suara dilakukan pembungkakan dengan ukuran 25 milidetik dengan ukuran bagian tumpang tindih sebesar 10 milidetik sehingga akan menghasilkan 98. frame dengan 1102 data didalamnya. Data hasil tahapan frame blocking kemudian akan digunakan untuk tahapan windowing dengan fungsi Hamming window yang akan menghasilkan 1102 data dalam setiap frame-nya. Selanjutnya pada tahapan FFT, banyaknya data yang digunakan sebesar 512 atau $NFFT = 512$ untuk setiap frame data suara. Setelah dilakukan FFT, akan dilakukan perhitungan spektrum daya yang akan digunakan pada tahapan filter bank. Banyaknya data yang digunakan dalam perhitungan spektrum daya adalah $\frac{NFFT}{2} + 1$, sehingga menghasilkan 257 data pada setiap frame-nya. Kemudian pada tahapan filter bank, nilai yang digunakan yaitu sebesar 40 atau $Nfilt = 40$, dan hasil pada tahapan filter bank ini akan diubah kedalam satuan dB. Hasil konversi dari tahapan filter bank ini selanjutnya akan dilakukan tahapan DCT dengan banyak data yang diharapkan adalah 11 data pada setiap frame data suara.

2.3. Support Vector Machines

Cara kerja SVM dalam melakukan klasifikasi data adalah dengan menggunakan hyperplane (bidang datar) yang akan menjadi batas keputusan (decision boundary) dengan memaksimalkan margin antar kategori (class). Margin adalah jarak antara hyperplane dan data terdekat di setiap layer. Data terdekat dengan hyperplane pada setiap kelas disebut support vectors [14]. Pada dasarnya SVM merupakan model yang dapat digunakan untuk beberapa klasifikasi kelas, dalam paper ini klasifikasi dibagi menjadi tiga kelas (multi-class classification) [15]. Hyperplane terbaik dapat ditemukan dengan mengukur margin hyperlane tersebut dan mencari titik maksimalnya. Oleh karena itu, dicari hyperplane terbaik dengan nilai margin terbesar [16]. Persamaan hyperplane ditulis dalam Persamaan 1 berikut.

$$f(x) = w \cdot x + b \quad (1)$$

Dimana :

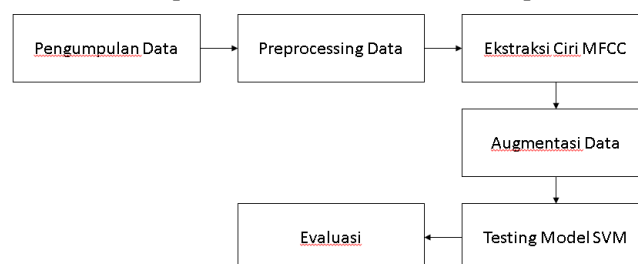
w = Parameter bobot

x = Vektor input

b = Bias

2.4. Tahapan Penelitian

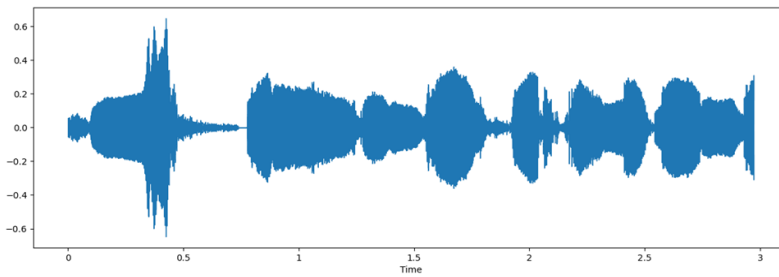
Jalannya penelitian mengenai klasifikasi suara berdasarkan usia dengan ekstraksi ciri MFCC dengan menggunakan model SVM ini diperlihatkan dengan diagram alur. Diagram alur yang dibuat pada gambar 2 ini untuk mempermudah alur informasi dari penelitian ini.



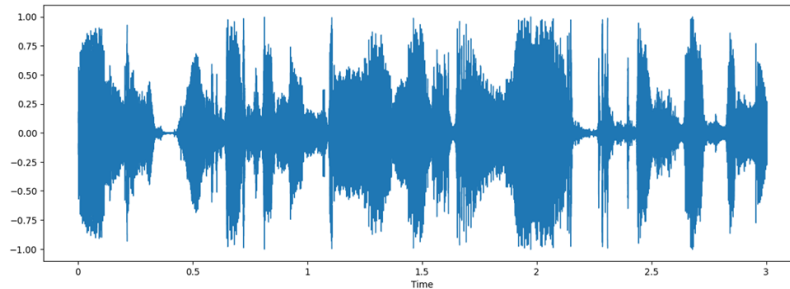
Gambar 2. Tahapan Penelitian.

2.5. Pengumpulan Data

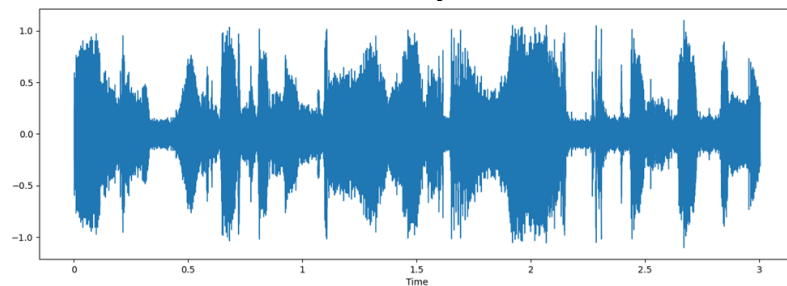
Data dalam paper ini menggunakan data primer dimana data diambil secara langsung (data primer) dari 100 orang sampel dengan pembagian 20 suara anak-anak pada usia 14 tahun ke bawah dengan contoh sinyal suara anak-anak disajikan pada Gambar 3, 40 suara usia remaja/dewasa usia 15-40 tahun dengan contoh sinyal suara remaja/dewasa disajikan pada Gambar 4 dan 40 suara usia lanjut usia 40 tahun ke atas.



Gambar 3. Contoh Sinyal Suara Anak-anak.



Gambar 4. Contoh Sinyal Suara Dewasa.



Gambar 5. Contoh Sinyal Suara Usia Lanjut.

2.6. Preprocessing Data

Dari 100 berkas rekaman suara ini akan dilakukan preprocessing dengan menghilangkan suara kosong yang ada dan dibagi menjadi 15 segmen suara. Setiap segmen suara memiliki panjang tiga detik, format *wav, sampling rate 44100 gelombang/detik, Channels mono, dan resolution 32-bit. Dengan dilakukannya preprocessing ini, total segmen suara yang awalnya hanya 100 berkas rekaman suara menjadi 1500 segmen suara dengan format berkas yang sama di setiap segmennya. Dari segmen suara yang telah terkumpul, berkas suara akan dikelompokkan sesuai kategori yang sudah ditentukan.

Sebelum proses klasifikasi, data suara melalui proses ekstraksi ciri menggunakan metode MFCC. Seluruh segmen suara yang ada diubah ke dalam bentuk digital dengan menggunakan tool librosa [8]. Setelah proses pre-processing, hanya dua detik pertama dari transformasi data suara digunakan pada proses ekstraksi ciri suara menggunakan MFCC. Metode MFCC memiliki beberapa tahapan, yaitu pre-emphasis, frame blocking, windowing, Fast Fourier Transform, filter bank, dan DCT [9].

2.7. Augmentasi Data

Setelah dilakukan ekstraksi ciri, setiap data suara akan diduplikasi dengan teknik augmentasi. Teknik augmentasi yang dipilih yaitu dengan menambahkan suara noise, metode shift (pergeseran) suara, dan juga dengan metode penambahan dinamik suara [13]. Tujuan dari dilakukannya audio data augmentasi adalah memperkaya dataset untuk proses pelatihan model dengan dataset yang bervariasi dan memiliki data yang cukup banyak. Artinya, variasi data suara yang sudah dilakukan augmentasi ini haruslah sebuah kemungkinan yang akan dilihat oleh model pada saat dilakukannya proses uji model. Dengan demikian, jelas bahwa pilihan teknik data augmentasi yang digunakan untuk melatih data suara ini harus dipilih dengan hati-hati.

3. Hasil dan Pembahasan

3.1. Pengolahan Data

Data yang akan digunakan dalam penelitian ini adalah data primer. Cara pengambilan data dalam penelitian ini yaitu dilakukan secara langsung. Dengan merekam suara narasumber sebanyak 100 orang, dengan setiap narasumber membaca sebuah kalimat yang sama yang sudah disiapkan oleh peneliti. Perbandingan suara laki-laki dan suara perempuan sama yakni 1:1, data suara juga dibagi menjadi 20:40:40 untuk masing-masing kategori. Tabel 1 berikut merupakan pembagian data suara yang dilakukan.

Tabel 1. Pembagian Data Suara

Kategori	Narasumber
Anak-Anak	20
Dewasa	40
Lanjut Usia	40

Dari 100 berkas rekaman suara ini akan dilakukan *preprocessing* dengan menghilangkan suara kosong yang ada dan dibagi menjadi 15 segmen suara. Setiap segmen suara memiliki panjang tiga detik, format **wav*, *sampling rate* 44100 gelombang/detik, Channels mono, dan *resolution* 32-bit. Dengan dilakukannya *preprocessing* ini, total segmen suara yang awalnya hanya 100 berkas rekaman suara menjadi 1500 segmen suara dengan format berkas yang sama di setiap segmennya. Dari 1500 data suara yang telah *dipreprocessing*, selanjutnya data di ekstraksi ciri MFCC.

Seluruh segmen suara yang ada diubah ke dalam bentuk digital dengan menggunakan *tool librosa* [8]. Setelahnya hanya dua detik pertama dari transformasi data suara digunakan pada proses ekstraksi ciri suara menggunakan MFCC. Selanjutnya setelah dilakukan ekstraksi ciri, setiap data suara akan diduplikasi dengan teknik augmentasi. Teknik augmentasi yang dipilih yaitu dengan menambahkan suara *noise*, metode *shift* (pergeseran) suara, dan juga dengan metode penambahan dinamik suara [13].

Maka hasil dari proses augmentasi data ini data suara bertambah empat kali lipat dari data suara awal hanya 1500 data menjadi 6000 data suara. Pada Tabel 2 berikut merupakan contoh empat keluaran dari data yang sudah di ekstraksi ciri MFCC dan dilakukan duplikasi.

Tabel 2. Contoh keluaran data

	Label	0	1	2	...	214	215
0	Elderly_Age	-12.407	-11.985	-11.655	...	-7.657	-10.716
1	Elderly_Age	-8.066	-8.411	-9.476	...	0.000	0.000
2	Elderly_Age	8.074	4.949	1.153	...	-17.112	-12.740
3	Elderly_Age	-13.836	-13.284	-11.197	...	-9.384	-5.843
4	Elderly_Age	-20.120	-20.273	-21.291	...	0.000	0.000

Setelah diperoleh dataset dengan fitur dan label hasil dari proses ekstraksi ciri menggunakan MFCC serta menduplikasi data suara, proses selanjutnya yaitu membagi data suara menjadi data latih dan data uji. Dalam paper ini, data latih yang digunakan yaitu sebesar 75% dari dataset yaitu sebanyak 4500 data suara, dan untuk data uji yang akan digunakan yaitu sebesar 25% dari dataset yaitu sebanyak 1500 data suara.

3.2. Optimasi Parameter SVM

Dalam paper ini, parameter terbaik didapatkan dengan menggunakan metode *Grid Search Cross Validation* [17], dengan memasukkan parameter-parameter dari model SVM *kernel* = [*rbf*, *poly*, *sigmoid*], nilai konstanta *C* = [0.1, 1, 10, 100], dan nilai *gamma* = [1, 0.1, 0.01, 0.001]. Untuk mengetahui keakuratan yang terbaik dalam pencarian parameter ini, banyaknya pengulangan (*k*) yang digunakan pada *Cross Validation* ini adalah 10. Dari pencarian optimasi parameter ini, mendapatkan hasil klasifikasi terbaik dengan *kernel* = *poly*, nilai konstanta *C* = 0.1 dan *gamma* = 1.

Pada pencarian parameter sebelumnya, belum terdapat parameter derajat maka dilakukan kembali pencarian akurasi terbaik secara manual menggunakan nilai parameter hasil optimasi

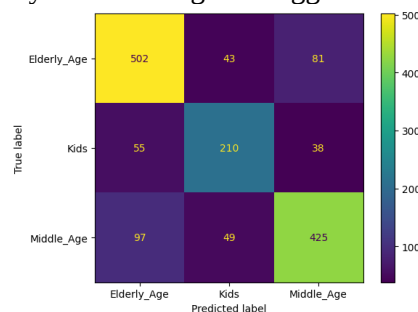
sebelumnya yaitu *kernel* = *poly*, nilai konstanta $C = 0.1$, $\gamma = 1$ dan beberapa nilai *derajat*. Keluaran akurasi terbaik dihasilkan dari masukan nilai *derajat* = 2 dengan akurasi data latih sebesar 75% seperti yang terlampir pada Tabel 3.

Tabel 3. Hasil Akurasi SVM.

Nilai Derajat	Akurasi Data Latih	Akurasi Data Uji
0	0,406	0,380
1	0,587	0,517
2	1	0,750
3	1	0,676
4	1	0,633

3.3. Hasil Klasifikasi Dan Evaluasi Model

Berdasarkan hasil akurasi terbaik yang telah didapatkan sebelumnya, selanjutnya dilakukan evaluasi model. Evaluasi model dilakukan untuk mengetahui efektifitas model dalam mengklasifikasi kelas yang berbeda. Salah satu hasil evaluasi model optimal SVM yang didapatkan pada tahap sebelumnya adalah dengan menggunakan *confusion matrix*.



Gambar 6. Hasil Confusion Matrix.

Berdasarkan *confusion matrix* yang diberikan pada Gambar 6 diperoleh hasil akurasi dari metode SVM sebesar $\frac{1137}{1500} = 0,758$. Lebih lanjut, nilai *precision* dan *recall*-nya disajikan pada Tabel 4. Pada Tabel 4 dapat dilihat tingkat ketepatan dari hasil prediksi yang diberikan metode SVM adalah 77% pada kategori usia anak-anak, 70% pada kategori dewasa dan 78% pada kategori usia lanjut. Serta terlihat tingkat keberhasilan sistem metode ini dalam menemukan kembali kategori pada usia anak-anak sebesar 80%, pada kategori dewasa 69% dan pada kategori usia lanjut 74%.

Tabel 4. Evaluasi Model.

Kelas Usia	<i>Precision</i>	<i>Recall</i>
Anak-anak (0-14 Tahun)	0,77	0,80
Dewasa (15-40 Tahun)	0,70	0,69
Usia Lanjut (40+ Tahun)	0,78	0,74

4. Kesimpulan

Penelitian ini menggunakan suara dari 100 orang narasumber dalam membaca teks cerita yang sama yang masing-masing suara berdurasi kurang lebih dua menit. Data tersebut dibagi berdasarkan kategori usia yang ditetapkan yakni anak-anak (0-14 tahun), remaja (15-40 tahun) dan usia lanjut (40+ tahun). Kemudian data dilakukan proses *preprocessing* sehingga diperoleh 1500 segmen suara berdurasi tiga detik yang dihasilkan dari 100 data awal. Setelah dilakukan ekstraksi ciri MFCC dengan tiga metode augmentasi data suara yakni metode *noise*, metode *shift* dan metode *dynamic change*, data segmen suara diduplikasi dan diubah menjadi beragam dari 1500 segmen suara menjadi 6000 segmen suara. Lalu, dilakukan pencarian parameter terbaik dari metode SVM yang menghasilkan parameter terbaik *kernel* = *poly*, nilai konstanta $C = 0.1$, $\gamma = 1$ dan nilai *derajat* = 2. Dari masukan parameter terbaik untuk metode SVM ini didapatkan hasil klasifikasi dengan akurasi sebesar 75,8%, dengan tingkat keberhasilan sistem dalam menemukan kembali yang tertinggi terdapat pada kategori usia anak-anak dengan hasil sebesar 80%.

Daftar Pustaka

- [1] Mirza Ardiana, Titon Dutono, and Tri Budi Santoso, "Identifikasi Jenis Kelamin Secara Real Time Berdasarkan Suara Pada Raspberry.Pi," *J. Politek. Caltex Riau*, vol. 8, no. 1, (2022), pp. 158–167.
- [2] M.Tri Satria Jaya, Diyah Puspitaningrum, and Boko Susilo, "Penerapan Speech Recognition Pada Permainan Teka-Teki Silang Menggunakan Metode Hidden Markov Model (HMM) Berbasis Desktop," *J. Rekursif*, vol. 4, no. 1, (2016), pp. 119–129.
- [3] A. Abdillah Alwi, P. Pandu Adikara, and Indriati, "Pengenalan Jenis Kelamin dan Rentang Umur berdasarkan Suara menggunakan Metode Backpropagation Neural Network," *J. Pengemb. Teknol. Inf. Dan Ilmu Komput.*, vol. 4, no. 7, (2020), pp. 2083–2093.
- [4] N. Kumar Goel, M. Sarma, T. Singh Kushwah, D. Kumar Agrawal, Z. Iqbal, and S. Chauhan, "Extracting speaker's gender, accent, age and emotional state from speech," *Interspeech*, (2018), pp. 2384–2385, doi: 10.1088/1742-6596/1410/1/012073.
- [5] Heriyanto, S. Hartati, and A. Eko Putra, "Ekstraksi Ciri Mel Frequency Cepstral Coefficient (MFCC) dan Rerata Coefficient Untuk Pengecekan Bacaan Al- Qur 'an," *Telematika*, vol. 15, (2018), pp. 99–108.
- [6] A. Abdulsatar, V. Davydov, V. Yushkova, A. Glinushkin, and V. Rud, "Age and gender recognition from speech signals," *J. Phys. Conf. Ser.*, vol. 1, (2019), doi: 10.1088/1742-6596/1410/1/012073.
- [7] D. Katerenchuk, "Age Group Classification with Speech and Metadata Multimodality Fusion," *Aclanthology.org*, (2017), pp. 188–193.
- [8] Teguh Puji Laksono, *Speech To Text Untuk Bahasa Indonesia*, (2018).
- [9] D. Putra and A. Resmawan, "Verifikasi Biometrika Suara Menggunakan Metode MFCC Dan DTW," *LONTAR Komput.*, vol. 2, no. 1, (2011), pp. 8–21.
- [10] Totok Chamidy, "Metode Mel Frequency Cepstral Coeffisients (MFCC) Pada klasifikasi Hidden Markov Model (HMM) Untuk Kata Arabic pada Penutur Indonesia," *J. MATICS*, vol. 8, no. 1, (2016), pp. 36–39, , doi: 10.18860/mat.v8i1.3482.
- [11] Angga Setiawan, Achmad Hidayatno, and R. Rizal Isnanto, "Aplikasi Pengenalan Ucapan dengan Ekstraksi Mel-Frequency Cepstrum Coefficients (MFCC) Melalui Jaringan Syaraf Tiruan (JST) Learning Vector Quantization (LVQ) untuk Mengoperasikan Kursor Komputer," *TRANSMISI*, vol. 13, no. 3, (2011), pp. 80–86,.
- [12] Torkis Nasution, "Metoda Mel Frequency Cepstrum Coefficients (MFCC) untuk Mengenali Ucapan pada Bahasa Indonesia," *J. Sains Dan Teknol. Inf.*, vol. 1, no. 1, (2012), pp. 22–31, doi: <https://doi.org/10.33372/stn.v1i1.309>.
- [13] P. Singh, R. Srivastava, K. P. S. Rana, and V. Kumar, "A multimodal hierarchical approach to speech emotion recognition from audio and text," *Knowl.-Based Syst.*, vol. 229, (2021), [Online]. Available: <https://www.sciencedirect.com/science/article/abs/pii/S0950705121005785>
- [14] A. Rahman Fauzi, *Simulasi Control Smart Home Berbasis Mel Frequency Cepstral Coefficient Menggunakan Metode Support Vector Machine (SVM)*. UIN Sunan Ampel, (2020). [Online]. Available: http://digilib.uinsa.ac.id/43041/2/Arif%20Rahman%20Fauzi_H06216003.pdf
- [15] Ryan Rifkin and Aldebaro Klautau, "In Defense of One-Vs-All Classification," *J. Mach. Learn. Res.*, vol. 5, (2004), pp. 101–141.
- [16] Renaldy Irfan, *Analisis Perbandingan Algoritma K-nearest Neighbor Dengan Algoritma Support Vector Machine Pada Pengklasifikasian*. (2012). [Online]. Available: <https://repository.uinjkt.ac.id/dspace/bitstream/123456789/55999/1/RENALDY%20IRFAN-FST.pdf>
- [17] I. Syarif, A. Prugel-Bennett, and G. Wills, "SVM Parameter Optimization Using Grid Search and Genetic Algorithm to Improve Classification Performance," *TELKOMNIKA*, vol. 14, no. 4, (2016), pp. 1502-1509. doi : 10.12928/TELKOMNIKA.v14i4.3956.