

Feature Extraction Using Mel-Frequency Cepstral Coefficients (MfCC) Technique For A Tajweed Guess Based on Android Application Development

Khodijah Hulliyah^{1*}, Lilik Umami Kultsum², Wahyu Hendarto Wibowo³ Anif Hanifa Setianingrum⁴, Arini⁵, Yusuf Durachman⁶

^{1,3,4,5,6} Department of Informatics Faculty of Science and Technology, Syarif Hidayatullah State Islamic University Jakarta

²Department of Quran and Tafsir Faculty of Ushuluddin, Syarif Hidayatullah State Islamic University Jakarta

^{1,2,3,4,5,6}Jl. Ir H. Juanda No.95, Ciputat, Kec. Ciputat Tim., Kota Tangerang Selatan, Banten Indonesia

ABSTRACT

Article:

Accepted: February 10, 2025

Revised: October 02, 2024

Issued: April 30, 2025

© Hulliyah et al, (2025).



This is an open-access article under the [CC BY-SA](https://creativecommons.org/licenses/by-sa/4.0/) license

*Correspondence Address:

khodijah.hulliyah@uinjkt.ac.id

The development of information and communication technology today has had a significant impact on various aspects of life, including education. One notable example is the increasing number of applications designed for learning to recite the Quran with proper tartil. The growing trend of tahfidz (Quran memorization) is undoubtedly a positive development from a religious perspective. However, many individuals focus solely on memorization without acquiring the ability to recite the Quran properly and accurately. One discipline that supports proper Quran recitation is the knowledge of tajweed. Numerous applications have been developed in this field, especially on Android platforms. However, applications that utilize artificial intelligence (AI) to recognize tajweed rules and involve users in guessing tajweed readings are still in need of further development. The aim of this research is to develop a tajweed learning application using the concept of Automatic Speech Recognition (ASR). This study employs data collection methods such as literature review, quantitative methods, and testing. The design is represented using Unified Modeling Language (UML), while the application is tested using the Black Box Testing method. For data analysis and testing of the speech recognition model, the Hidden Markov Model (HMM) algorithm is employed, with Mel-Frequency Cepstral Coefficients (MFCC) used for feature extraction. The output of this research is an Android-based tajweed learning application that integrates speech recognition and allows users to guess tajweed rules interactively.

Keywords : *android; learning apps; artificial intelligence; hidden markov model; speech recognition.*

1. INTRODUCTION

The rapid development of technology and information has enabled humans to interact with computers or gadgets using voice as input. One of the technologies facilitating this interaction is Automatic Speech Recognition (ASR), also known as speech recognition. ASR enables programs to convert human speech into written text [1]. In the home automation sector, ASR is widely used to control household devices through voice commands [2].

In the field of education, particularly for Muslims, learning to read the Qur'an holds great significance for those who wish to recite it with *tartil* (proper articulation and precision), especially when learning Tajweed rules. There are several ways to learn Tajweed, including studying with teachers, reading books, or utilizing electronic applications. However, self-learning without guidance from a teacher has its limitations, as learners may struggle to identify and correct mistakes in their Qur'anic recitation, particularly in adhering to Tajweed rules. Many existing applications still require further development to address this issue.

The term Tajweed, linguistically, means "to make better." In Islamic terminology, it is defined as "pronouncing each letter from its articulation point while giving it its due and required rights." Reciting the Qur'an without observing the rules of Tajweed is considered sinful because Allah Almighty revealed the Qur'an with proper Tajweed. Learning Tajweed is a commendable act for Muslims. Those who study and apply it in their Qur'anic recitation are elevated in rank by Allah, both in this world and in the Hereafter [3].

Tajweed encompasses a variety of recitation laws, some of which govern the elongation and shortening of certain verses or words, while others dictate the pronunciation rules when one letter meets another [4].

To make learning Tajweed more accessible and engaging, this study aims to develop a Tajweed guessing application that integrates speech recognition technology. The application provides Tajweed learning materials and interactive exercises. Users will listen to Qur'anic verses and identify the type of Tajweed present in the recitation.

One of the algorithms used in speech recognition technology is the Hidden Markov Model (HMM). This algorithm processes voice

input and compares it with predefined Tajweed terms in the application. If the user's guess is correct, the application will display a message confirming the correct answer. Conversely, if the answer is incorrect, a message indicating the error will be shown.

1.1. Related Works

Research on the application of Automatic Speech Recognition (ASR) in Qur'anic learning has been gaining attention in recent years, particularly with the integration of mobile applications and speech recognition technology. A study conducted by Kamarudin, Al-Haddad, Hassan, and Abushariah[5] highlights the widespread acceptance of mobile applications worldwide, with a focus on developing apps tailored for different user backgrounds. One of the prominent areas of development is learning the Qur'an, particularly Tajweed and the pronunciation of *makharijul* letters, on Android platforms due to their broad user base. The study found that 81% of respondents agreed that using speech recognition on mobile devices made learning easier and more accessible.

In addition, research by Nada, Ridhuandi, Santoso, and Apriyanto (2019) [6][7] explored the application of ASR with the Hidden Markov Model (HMM) for pronouncing Hijaiyah letters in Qur'anic learning. This study emphasized that improper pronunciation can alter the meaning of the Qur'anic text. By leveraging speech recognition technology, the study demonstrated its ability to detect errors in the pronunciation of Hijaiyah letters. The results showed that the accuracy for testing individual Hijaiyah letters was 100%, while for testing different letters, the accuracy dropped to 54.6%.

Another significant study by Rahmantara, Wardhani, and Saf (2018) [8] and by Sujjada (2024) [5] developed a Qur'anic surah name recognition application for the 30th Juz using speech recognition technology. The algorithm used in this application was based on the Markov Model (Markov Chain), which processed the user's voice input to determine the probability of correctly identifying the surah name. The study found that the highest accuracy rate of the application occurred in noiseless environments, achieving 100% accuracy when the user was positioned approximately 50 cm from the device. Furthermore, the questionnaire

results indicated that 87.74% of respondents believed that the application was effective in helping users better recognize surah names in the 30th Juz.

Meanwhile, Sudiarjo, Mariana, and Nurhidayat [9] developed an application focused on Tajweed, *waqaf* (pauses in recitation), and *makharijul* letter learning, utilizing the Luther developer method with Java programming. This research addressed the challenges faced by many Muslims in memorizing and applying Tajweed rules during Qur'anic recitation. The application was designed to assist users in overcoming these difficulties, particularly those who had limited knowledge of Tajweed or rarely applied it when reading the Qur'an [5].

In another study, Anggraini., et., al[10] and [11] worked on the development of an Android-based Qur'an memorization correction application integrated with Google Speech recognition. This application was designed to support the memorization of specific surahs, namely Al-Ikhlâs, Al-Kautsar, and An-Naas. Although the Google Speech API achieved 100% accuracy for short verses, the study found that the API struggled with longer verses or unclear pronunciations. Despite these limitations, the application showed promising results in converting voice input into text and providing feedback for Qur'anic memorization.

These previous studies highlight the potential of ASR technology in enhancing Qur'anic learning, especially in improving pronunciation and recitation accuracy. However, many existing applications focus mainly on the recognition of individual letters or surahs, and there is still a need for tools that specifically address the rules of Tajweed. The integration of ASR with the Hidden Markov Model, as demonstrated in these studies, offers a promising approach for improving the accuracy and engagement of Tajweed learning applications.

Although previous studies have successfully integrated ASR into Qur'anic learning, they primarily focused on letter pronunciation, surah identification, or general Tajweed applications without providing an interactive guessing-based approach. This research introduces a Tajweed Guessing Game Application that incorporates speech recognition with HMM and MFCC-based feature extraction. The novelty of this study lies

in the interactive Tajweed recognition mechanism, where users listen to verses and determine the correct Tajweed rule, enhancing engagement and learning outcomes.

1.2. Theoretical Foundation

a. Android

Android is an operating system based on Linux, designed to be installed on smartphones and tablet devices. This operating system is adaptable, ranging from low-end to high-end specifications. Android is a mobile operating system based on Linux. In addition, plains that Android is an OS for Linux-based mobile phones. This OS provides an open platform for developers to create applications that can be used on various mobile devices[12].

b. Speech Recognition

Speech recognition, commonly known as automatic speech recognition (ASR), is a development of techniques and systems that allow computers to receive input in the form of spoken words. This enables devices to recognize and understand spoken words by digitizing the words and matching the digital signals with certain patterns stored within the device. There are two modes of speech recognition: the dictation mode, where the user speaks a word that is then recognized by the computer and converted into text, and the command-and-control mode, where the user speaks predefined words in the database to execute certain commands [13]. The sounds we hear actually have specific signal characteristics, while speech consists of words spoken aloud. In speech recognition, voice signals are extracted to analyze the characteristics of incoming voice, obtaining information that can be analyzed for each variation of the voice signals. These recognized characteristics are then converted into text.

c. Hidden Markov Model

In the development of the application, the Markov Model algorithm is used. This is a mathematical technique employed to model various systems and business processes. It can be used to forecast future changes based on dynamic variables in past events and observations.

A Markov chain is a sequence of random variables such as $X_1, X_2, X_3, \dots, X_n$, which has the Markov property, meaning the future state

depends only on the current state and is independent of past states. In other words:
 $P(X_{n+1} = j | X_1 = x_1, X_2 = x_2, X_3 = x_3, \dots, X_n = i)$
 $= P(X_{n+1} = j | X_n = i) = P_{ij}$

Where i is the variable we are looking for and j is the known variable. Thus, the probability of the occurrence of variable i can be determined if the value of variable j is known. In Hidden Markov Models (HMM), there are two types of states: observable states and hidden states [6]. Research conducted by Li X et.al [14] defines HMM as:

$$\lambda = (A, B, \pi)$$

Elements of the Hidden Markov Model include:

1. N is the number of hidden states, $S = \{S_1, S_2, \dots, S_N\}$ as a sequence of long states t as Q_t .
2. M is the number of observable symbols at the deepest state, $V = \{v_1, v_2, \dots, v_M\}$ as a sequence of long states t as O_t .
3. A is the probability distribution of state transitions, $A = \{a_{ij}\}$, where $1 \leq i, j \leq N$.
4. B is the probability distribution of observable symbols, $B = \{b_j(V_k)\}$, where $1 \leq h \leq N, 1 \leq k \leq M$.
5. Π is the distribution of initial states, $\pi = P(Q_1 = s_i)$, where $1 \leq i \leq N$.

1.3. MFCC (Mel Frequency Cepstrum Coefficients)

The MFCC method is used for feature extraction, which involves obtaining parameters and information about the characteristics of a person's voice. The MFCC method was first introduced by Davis and Mermelstein around 1980. It is considered one of the most effective methods for speech recognition in the field of automatic speech recognition. MFCC is widely used for feature extraction in speaker recognition and speech recognition applications[15].

MFCC is a feature extraction method used to obtain cepstral coefficients and frames, which can then be processed for speech recognition to achieve better accuracy[16].

a. Pre-emphasis

Pre-emphasis, high-frequency components while aligning low and high frequencies. Pre-emphasis reduces noise to improve the Signal-to-Noise Ratio (SNR) and mitigate unwanted sounds. Pre-emphasis is a

simple signal processing technique that acts as a linear filter and is still in the time-domain.

Pre-emphasis is the initial stage in the MFCC process. This step is necessary because the signal often experiences noise interference, so noise must be reduced. Pre-emphasis addresses this issue using a very simple filtering method. The goal of pre-emphasis is to maintain a good signal quality at higher frequencies while ensuring that the baseband level remains consistent.

The pre-emphasis process, according to Proakis and Manolakis (1996), is described using an α value between 0 and 1, or $0.9 \leq \alpha \leq 1.0$, and follows this equation:

$$y(n) = s(n) - \alpha \cdot s(n-1) \quad (1)$$

Where $y(n)$ is the pre-emphasized signal, $s(n)$ is the original signal, and α is the pre-emphasis filter constant between 0.9 and 1.0.

b. Frame Blocking

After pre-emphasis, the signal undergoes a frame blocking process, where the signal is divided into frames containing NN samples, and each frame is shifted by MM samples, with $N = 2MN = 2M$ and $M < NN < N$ (Abriyono & Harjoko, 2012). The frame width is represented by NN , and the frame shift width is denoted by MM . The overlap width between frames is calculated as the difference $N - MN - M$. Frame blocking, according to Holmes (2003), is the process of analyzing speech signals by dividing them into frames. Each frame is represented by a single feature vector, which is the average spectrum for the time interval of that frame. The frame duration typically ranges between 20-40 milliseconds. The frames are taken to achieve good frequency resolution while maintaining the shortest possible time to optimize the time resolution. The number of frames is calculated using the following equation:

$$fl(n) = \text{and}(Ml+n) \quad (2)$$

Where $fl(n)$ is the result of frame blocking, nn is from 0.1 to $N-1N-1$, NN is the number of samples, MM is the frame length, and ll is from 0.1 to $L-1L-1$, where LL is the entire signal.

c. Windowing

Windowing, according to Proakis and Manolakis (1996), smooths the spectrum after the frame blocking process. The purpose of windowing is to reduce the discontinuity effects at the edges of each frame, which result from the frame blocking process. Common windowing techniques include Rectangular Window, Hamming Window, and Hanning Window (Chamidy, 2016). Researchers have used the Hanning window in their work because it is smoother compared to the others (Putra, 2008a). The windowing function is represented by the following equation:

$$X(n)=fl(n) \cdot w(n) \quad X(n)=fl(n) \cdot w(n) \quad (3)$$

Where $X(n)$ is the windowed signal, $fl(n)$ is the result of frame blocking, n ranges from 0.1 to $N-1N-1$, N is the number of samples in each frame, and $w(n)$ is the window function.

1.4. Tajweed Rules

Tajweed is the science of Qur'anic recitation that ensures proper pronunciation and articulation of Arabic letters. Tajweed rules can be categorized into several key principles:

1. Makharijul Huruf: The correct articulation points of letters.
2. Sifatul Huruf: Characteristics of letters, such as heaviness (tafkhim) and lightness (tarqiq).
3. Ahkam Nun Sakinah wa Tanwin: Rules for pronouncing Nun Sakinah and Tanwin, including Idgham, Iqlab, and Ikhfa'.
4. Ahkam Mim Sakinah: Rules for Mim Sakinah, including Idgham Mutamathilain, Ikhfa Syafawi, and Izhar Syafawi.
5. Madd (Elongation): Various types of elongation in recitation, such as Madd Wajib Muttasil and Madd Jaiz Munfasil.

2. METHODS

This stage is carried out to understand the research methods used and analyze system needs. The type of research carried out is experimental research by conducting experiments on control variables in the form of *inputs* to obtain results in the form of *outputs*.

2.1. Requirements Definition and Analysis

a. References

Study several references, books, journals or other literature materials related to research.

b. Data Collection

Conduct data collection by summarizing some data from various sources for use in material features. Tajweed material is taken from the summary of the book Tajweed Science Studies by Marwan Hadidi, M.Pd.I.

c. Observation

Observe similar applications to find out about the advantages, and disadvantages that can be improved.

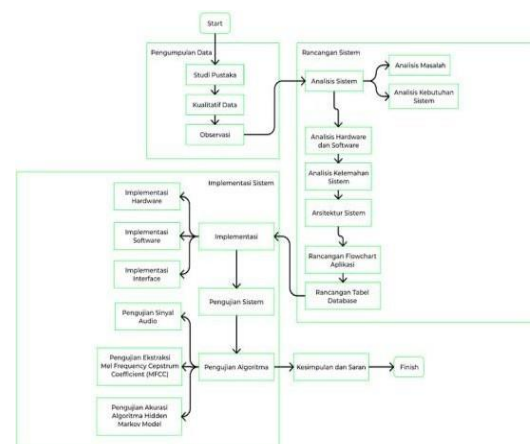


Figure 1. The flow of thinking framework

2.2. System Architecture Design

This application is designed with an architectural scheme as shown below with data used locally and cloud-based so that users must also be connected to the internet.



Figure 2. System architecture design

2.3. Application Flowchart Planning

In the process of designing this application design, it is necessary to have a design sketch that describes the series of application processes. The sketch is formed in a flowchart.

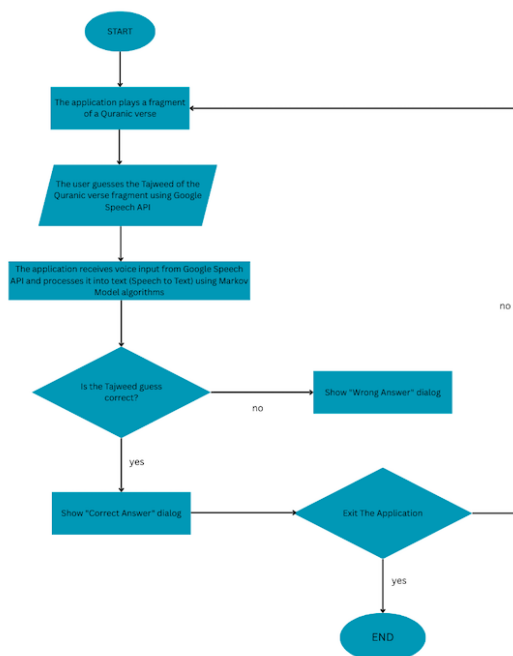


Figure 3. Application flowchart design

2.4. System and Application Design

This stage is formed by the structure of the system based on the requirements that have been processed in the previous phase.

How to identify and guess tajweed from voice to text and its relationships.

a. System planning with UML (Unified Model Language)

The Speech Recognition method in this study through the design process and application of guess tajweed conceptually using usecase diagrams, activity diagrams and sequence diagrams.

b. User interface design (User Interface)

Communication facilities between users and systems are provided by user interfaces. User Interface provides information facilities and various information with the aim of tracing the problem until the solution is provided. In addition, the user interface is designed as an illustration at the time of system creation.

c. Google Cloud Speech API implementation

A means of coding using Python implemented in the Google Speech API. At this stage also carried out the conversion of audio data, storage on Google Cloud Storage. In addition, it is also the process of designing speech recognition and integrating data with Google Cloud Storage.

2.5. System Testing Techniques

This phase is the testing phase of the system that has been created. The focus of this test is to test the accuracy of the output results by measuring the recognition rate and testing the system using blackbox testing. The test is carried out to minimize errors and ensure the output produced is as expected.

The Black Box Testing method performs direct testing on android applications installed on android and tested to get the expected results.

a. Layout Switching Testing

This test is done to observe the movement from one layout to another.

b. Input Menu Testing

This test is performed to test that the input menu created in the application is able to function properly.

c. Process Menu Testing

This test is done to test the function processes created in the application.

d. Output Menu Testing

This test is carried out to test the output obtained in accordance with the input.

e. Design Use Case Diagram

This process is used to show the relationship between the user and the system designed to be able to perform activities within the application.

f. Database Table Design

The design of this database table is used to manage data that will be interrelated to be integrated with the application system. The database design used is as follows:

Table 1. Database table design

Attribute Name	Data Type	Information
No_tajwid	Int (3)	Primary Key
Name_tajwid	Varchar (100)	Name of Tajweed
Img_tajwid	Int(10)	Path Icon Tajwid
Material_tajwid	Int(10)	Path Materi Tajwid
Audio_quiz	Int(10)	Path Audio Ayat
Img_quiz	Int(10)	Path Image Ayat
Answer_quiz	Varchar (100)	Practice Answers

3. RESULTS AND DISCUSSION

This stage is carried out to discuss the implementation of the system design that has been analyzed in the design into the form of programming to produce applications that are made based on needs, as well as the results of application testing using the black-box testing method. Interface Implementation

Here are some interfaces of the application that have been implemented:

Test Cases and Results			
Action/Input Data	Expected Response	Observation	Conclusion
Pressing The <i>Tajweed</i> Material button	Move to the <i>Material List</i> page	The choice of action is as expected	Succeed
Press the <i>Exercise</i> button	Move to the <i>List of Training Levels</i> page	The choice of action is as expected	Succeed
Press the <i>Search</i> button	Move to a <i>Page Search</i>	The choice of action according to what to expect	Succeed
Press knob <i>About We</i>	Move to yard <i>About We</i>	Choice Action accordingly with that Expected	Succeed

Table 2. Main menu layout switching test

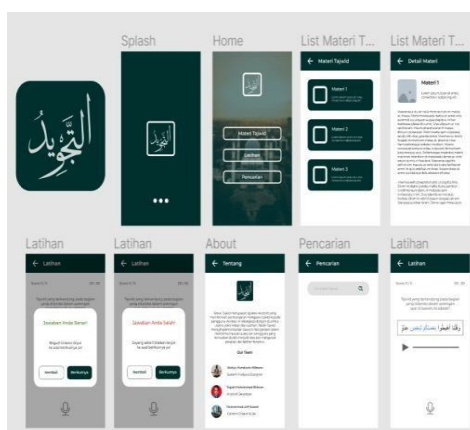
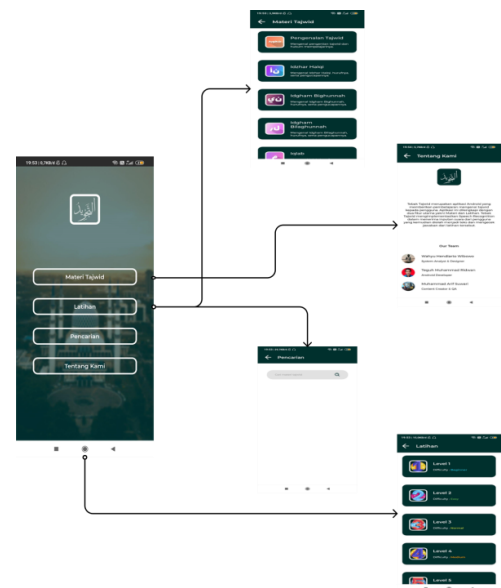


Figure 4. All tebaktajweed application pages

3.1. System Testing

Testing is mandatory in application development in obtaining information about the quality of the application that has been built to find out whether the function of the application has run as planned. The method of system testing in this writing is testing with Black Box Testing.

a. Main Menu Layout Switching Testing



5. Main menu layout switching test results

Figure

b. Input Menu Testing



Figure 6. Input menu test results

c. Audio Testing Questions

Table 3. Audio test questions

Test Cases and Results			
Action/Input Data	Expected Response	Observeran	Slaughter
Press the <i>play</i> button on the problem	There was a voice that indicted the fish about	The choice of action is as expected	Succeed
Press the <i>pause</i> button on the problem	The sound of the question stopped ringing	The choice of action is as expected	Succeed

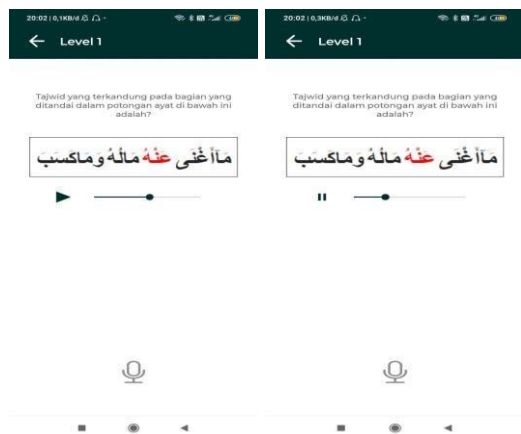


Figure 6. Audio question test results

d. Answering Testing with Speech Recognition

Table 4. Test answering with speech recognition

Test Cases and Results			
Action/Input Data	Expected Response	Observation	Conclusion
Press the microphone	Bring up a popup that listens to the user's speech	The choice of action is as expected	Succeed
Answer questions correctly	Bring up the correct answer popup	The choice of action is as expected	Succeed
Answer question incorrectly	Bring up a false answer popup	The choice of action is as expected	Succeed

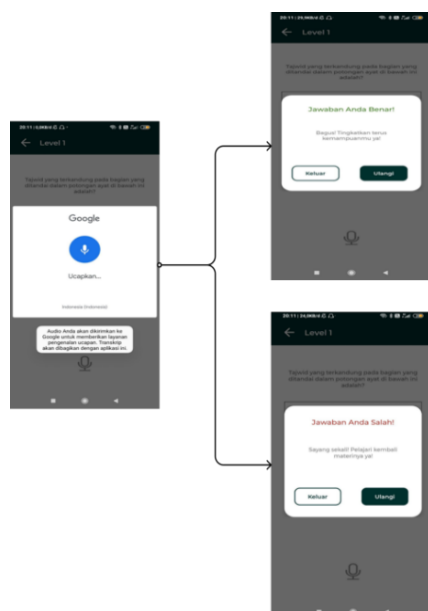


Figure 7. Test results answer with speech recognition

3.2. Testing the Hidden Markov Model Algorithm

a. Audio Signal Testing

The results of the voice recordings were used to find out how accurate the words entered into the hidden markov training model were. Audio created in *.wav format and recorded 9 times.

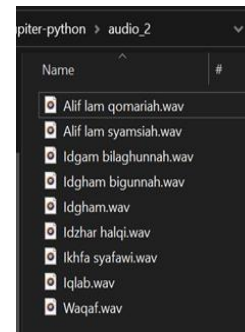


Figure 8. Data file recording with *.wav format

After the audio recording is made, it will then be displayed in the form of word snippets with

There is a chunk in the audio file, in this case using the identity of the audio data path to display the audio file in the form of a string.

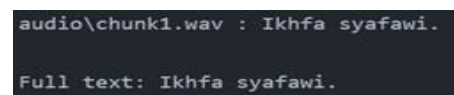


Figure 9. Results of audio data split

Furthermore, to find out the shape of the signal from the audio input, use a plot for sampling rate.

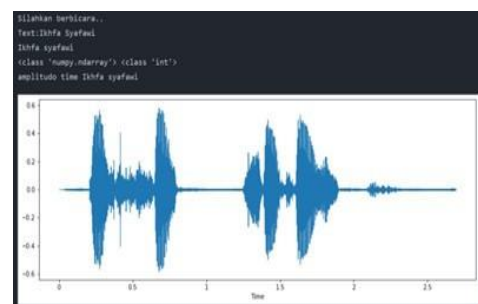


Figure 10. The form of the sampling rate signal in ikhfa syafawi.wav

b. Mel Frequency Cepstrum Coefficient (MFCC) Feature Extraction Testing

Audio file data is assigned an id paired with a label using the word2id dictionary.

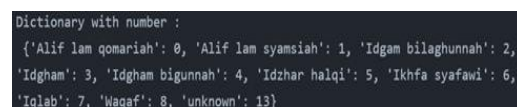


Figure 11. Results of word2id audio file

```
Train Sample Data from Dict:
      folder      file \
1 Idgam bilaghunnah audioIdgam bilaghunnah/Idgam bilaghunnah.wav
2 Idgham bigunnah audioIdgham bigunnah/Idgham bigunnah.wav
3 Ikhfa syafawi audioIkhfa syafawi/Ikhfa syafawi.wav
4 Alif lam syamsiah audioAlif lam syamsiah/Alif lam syamsiah.wav
      Waqaf      audioWaqaf/Waqaf.wav

      label
1 Idgam bilaghunnah
2 Idgham bigunnah
3 Ikhfa syafawi
4 Alif lam syamsiah
      Waqaf
```

Figure 12. Shorting results with word2id

Then the data that has been shorted with word2id will be used to train samples with one hot encode, which will produce 5 labels for the train.

```
One Hot from train data:
[[0. 1. 0. 0. 0. 0. 0. 0. 0.]
 [0. 0. 1. 0. 0. 0. 0. 0. 0.]
 [0. 0. 0. 0. 1. 0. 0. 0. 0.]
 [0. 0. 0. 0. 0. 0. 1. 0. 0.]
 [0. 0. 0. 0. 0. 0. 0. 0. 1.]]
```

Figure 13. Results of one hot encode with binary values of 1 and 0

Furthermore, to see the file data that has been trained, validated and tested, a list filename must be set using vector y on the ground truth label.

```
Break Dataset:
FileNames_val:
('Ikhfa syafawi.wav',)
FileNames_test:
('Ikhfa syafawi_test.wav',)
FileNames_train:
('Idgham.wav', 'Alif lam syamsiah.wav', 'Idzhar halqi.wav', 'Iqlab.
wav', 'Waqaf.wav', 'Idgham bigunnah.wav', 'Idgam bilaghunnah.wav',
'Alif lam qomariah.wav')
Break Value Y:
Value Y filenames_val:
(6.0,)
Value Y filenames_test:
(0.0,)
Value Y filenames_train:
(3.0, 1.0, 5.0, 7.0, 8.0, 4.0, 2.0, 0.0)
```

Figure 14. Validation, test, and train filename results

Audio sample data that has been trained, will be classified with Mel Frequency Cepstrum Coefficient. With MFCC, you can also see what percentage of the sample data is problematic and get a result of 0.016%.

```
Value MFCC from dataset:
[[-5.61334929e+01 -5.66461711e+01 -4.94706475e+01 -2.92383669e+01
 -1.90881867e+01 -1.64016039e+01 -1.92806627e+01 -2.31104152e+01
 -2.22756641e+01 -2.14788503e+01 -2.30657721e+01 -2.51154162e+01
 -2.53547912e+01 -2.75536870e+01 -3.55592134e+01 -4.43165142e+01
 -2.91387854e+01 -1.88764803e+01 -1.54697369e+01 -1.60046447e+01
 -2.20042223e+01 -3.06959133e+01 -3.51444286e+01 -3.70683806e+01
 -3.99344529e+01 -4.16819318e+01 -3.92175164e+01 -3.89195369e+01
 -4.12572727e+01 -3.67726469e+01 -2.74708414e+01 -2.38454326e+01
 ...
 2.18201975e-01 9.04190508e-02 3.19248909e-01 8.54499909e-02]]
Percent of Probability Problem Data from MFCC: 0.016 %
```

Figure 15. The result of one of the audio signals mel frequency cepstrum coefficient

3.3. Testing the Accuracy of the Hidden Markov Model Algorithm

One of the accuracy calculation methods used in this case is the provision of data values in training of 8 data samples including:

```
filenames_train:
('Idgham.wav', 'Alif lam syamsiah.wav', 'Idzhar halqi.wav', 'Iqlab.
wav', 'Waqaf.wav', 'Idgham bigunnah.wav', 'Idgam bilaghunnah.wav',
'Alif lam qomariah.wav')
```

Figure 16. Data results that have been trained

With the final result of calculating the accuracy value in this Hidden Markov Model is:

$$\begin{aligned} \text{Accuracy} &= (\text{data train} / \text{all data}) * 100\% \\ &= (8/9) * 100\% \\ &= 88.9\% \end{aligned}$$

From result accuracy to Identity

Words obtained by 88.9% of the 9 words that provided with 8 words trained by Hidden Markov Model.

CONCLUSION

The overall features of the built application have functioned well through Black Box Testing. The Hidden Markov Model algorithm has been well implemented in this application. The application uses Google's Cloud API in implementing the model so it requires an internet connection for the user. And the Hidden Markov Model algorithm implemented has an accuracy of 88.9% where 8 data have been successfully trained from 9 data samples.

REFERENCES

- [1] A. Abdulkareem, T. E. Somefun, O. K. Chinedum, and F. Agbetuyi, "Design and implementation of speech recognition system integrated with internet of things," *International Journal of Electrical and Computer Engineering*, vol. 11, no. 2, pp. 1796–1803, Apr. 2021, doi: 10.11591/ijece.v11i2.pp1796-1803.
- [2] U. I. Ibrahim, H. Ohize, U. A. Umar, and Y. Aliyu, "Design and Construction of Voice Controlled Home Automation using Arduino," *ABUAD Journal of Engineering Research and Development (AJERD)*, vol. 7, no. 1, pp. 144–152, Apr.

- 2024, doi: 10.53982/ajerd.2024.0701.14-j.
- [3] M. J. Aqel and N. M. Zaitoun, "Tajweed: An Expert System for Holy Qur'an Recitation Proficiency," in *Procedia Computer Science*, Elsevier, 2015, pp. 807–812. doi: 10.1016/j.procs.2015.09.029.
- [4] "H" SAYUTI JWID." [Online]. Available: www.tedisobandi.blogspot.com
- [5] A. Sujjada, G. P. Insany, and M. F. Nugraha, "Implementasi Automatic Speech Recognition Pada Penilaian Hafalan Al-Quran Dengan Metode Murojaah Berbasis Android," *CICES*, vol. 10, no. 2, pp. 216–228, Aug. 2024, doi: 10.33050/cices.v10i2.3247.
- [6] Q. Nada, C. Ridhuandi, P. Santoso, and D. Apriyanto, "Speech Recognition dengan Hidden Markov Model untuk Pengenalan dan Pelafalan Huruf Hijaiyah," 2019.
- [7] A. Al Harere and K. Al Jallad, "Quran Recitation Recognition using End-to-End Deep Learning."
- [8] D. S. Rahmantara, K. D. K. Wardhani, and M. R. A. Saf, "Aplikasi Pengenalan Nama Surah pada Juz ke 30 Kitab Suci Al-Qur'an Menggunakan Speech Recognition," *Jurnal RESTI (Rekayasa Sistem dan Teknologi Informasi)*, vol. 2, no. 1, pp. 345–353, 2018, doi: 10.29207/resti.v2i1.285.
- [9] A. Sudiarjo, A. R. Mariana, and W. Nurhidayat, "Aplikasi Pembelajaran Ilmu Tajwid , Waqaf dan Makharijul Huruf Berbasis Android," *Jurnal Sisfotek Global*, vol. 5, no. 2, pp. 54–60, 2015, [Online]. Available: <http://journal.stmikglobal.ac.id/index.php/sisfotek/article/view/80>
- [10] N. Anggraini, A. Kurniawan, L. K. Wardhani, and N. Hakiem, "Speech recognition application for the speech impaired using the android-based google cloud speech API," *Telkomnika (Telecommunication Computing Electronics and Control)*, vol. 16, no. 6, pp. 2733–2739, Dec. 2018, doi: 10.12928/TELKOMNIKA.v16i6.9638.
- [11] A. Akbar, A. Y. Husodo, and A. Zubaidi, "The Implementation of the Google Speech on Qur'an Recitation Correction Application Based on Android)." [Online]. Available: <http://jtika.if.unram.ac.id/index.php/JTIKA/>
- [12] N. K. Ekanayake and N. K. Ekanayake, "Android Operating System," 2018. [Online]. Available: <https://www.researchgate.net/publication/325257105>
- [13] A. Abdulkareem, T. E. Somefun, O. K. Chinedum, and F. Agbetuyi, "Design and implementation of speech recognition system integrated with internet of things," *International Journal of Electrical and Computer Engineering*, vol. 11, no. 2, pp. 1796–1803, Apr. 2021, doi: 10.11591/ijece.v11i2.pp1796-1803.
- [14] X. Li, T. Henning, and C. Camerer, "Estimating Hidden Markov Models (HMMs) of the cognitive process in strategic thinking using eye-tracking," *Frontiers in Behavioral Economics*, vol. 2, Oct. 2023, doi: 10.3389/frbhe.2023.1225856.
- [15] "Automatoc Speech Recognition: Systematic Literature Review," *IEEE Access*, 2021.
- [16] "Mel Frequency Cepstral Coefficient and Its Application: A Review," *IEEE Access*, 2022.