

Speech Recognition System Using Machine Learning Algorithm

Jonnakuti Pooja¹, J. Aswini², N. Deepa³

¹UG Scholar, ^{2,3}Assistant Professor, Saveetha School of Engineering, Saveetha Institute of Medical and Technical Sciences, Chennai, India

¹Poojareddy7821@gmail.com, ²aswinij.sse@saveetha.com, ³deepa23narayanan@gmail.com

Article Info

Volume 82

Page Number: 10793 – 10796

Publication Issue:

January-February 2020

Abstract

By improvement of data innovation as well as prevalence of gadgets, voice detection method, which is introduced in numerous gadgets like mobiles as well as programmed control hardware, is assuming a undeniably significant job in individuals' lives. Moreover, the Security problems ought to likewise get adequate consideration. People impression of discourse signals is influenced by earlier information as well as setting. That is normal for humans to unknowingly made the dark piece of signals that comes from speech he heard. So it won't use along these lines for speech detection application. At particular situation, Inaccurate outcomes which are long way from ordinary individuals detection will be back from application. This event displays the speech detection methods has weakness that made conceivable the structured risky directions might be compiled not with administrator's cognizant. Here in this project introduces a system to pre-process speech sign as well as make a great deal of prepared signs that are comparative in people sound-related however unique in real. At that point, it introduces a procedure to choose prepared speech signs focusing on an objective voice detection method.

Keywords: Voice detection, Security, Signals, risky directions, speech.

Article History

Article Received: 18 May 2019

Revised: 14 July 2019

Accepted: 22 December 2019

Publication: 19 February 2020

1. Introduction

Voice identification is one of the top most rated things in the current generation, the main purpose of the voice recognition is to work smartly and give an helping hand to the people who are disabled like for those people who are blind and to the people who cannot type. With introduction of voice recognition it helped a lot of people by getting the results just by the voice, the main of this paper is to recognise the voice and converting it into text, the difference between other papers and this paper is that other papers listen to the voice of the user just to an certain extent , but in this paper we propose an system which is capable of listening to the long conversations and making it as an text. In this paper we use an MFC (mell-frequency cepstrum) for the longer duration identification of speech, the MFC algorithm is an algorithm which is used for the identification of the voice given by the user, the MFC recognises the voice which is in certain frequency and make it as an machine understandable one. So with the help MFC algorithm we

are going to propose model for converting the speech into text of longer durations.

2. Literature Review

This project mainly focuses on building an set of instructions to upgrade voice detection of Stammering voice. Stammering is a Disorganisation that influences familiarity of voice by automatic continuous processing, extensions of letters or words, or automatic quiet interims. Present voice detection methods neglect to perceive stammering voice. Techniques to identify shatter has described for in writing however proficient procedures for stammer verification has not described for. Current project describes the problem as well as introduces process to identify as well as verification shatter inside time span. For removing extensions to an extent from example, abundancy marked value through neural systems is created. Redundancies that are expelled along with reiteration of string removing set of instructions utilizing a current TTS method. In this manner sign removing all shattering, creates voice detection.

Voice expression detection is a well known content all over the world. Some problems that need to be improved the presentation of that voice expression detection is the way of choosing of voice expression characters. To provide a great voice expression detection method, that is necessary to grab the characters that suits for voice expression features. Scientists did a tremendous work with great effort introduces a expression characters as well as done robust result. Even every attributes which has been implemented which has shown that need to be efficient in the work, various techniques depends on mono kinds. This project it introduced that techniques of various attributes depends on machine learning, joining some characters. This demonstration describes the systems that enhance the presentation of voice expression detection method.

3. Proposed System:

Mel Frequency Cepstral Coefficient got large fame in Speech Detection method. It is the most recognised as well as more repeatedly utilised for voice as well as speech detection. At the point when recurrence groups are set substantially in MFCC, it assesses human method reaction cautiously than some other methods. Technique for preparing MFCC depends on less period research, as well as in this way from casing a MFCC vector is processed. So as to get collaboration, voice trial has been extracted as sample file as well as hamming window is selected to diminish interruption of sign. At that point Discrete Fourier Transform (DFT) that need to be utilized.

4. Results and Discussion

The initial step of implementation process starts with providing or uploading the input as audio. So it can take the audio and helps in processing of further steps. Next step is the preprocessing that can be utilized to remove noisy data which removes all the unclear data that surrounds all the data which covers under a audio. It helps in focusing on the context or speech available in the input. While managing the sound signs, these procedure converts into removing the sound carvings that don't contain audio. It is said to be as VAD ie. Voice activity Detection.

Received audio can be considered in a wave format. This can be analysed and can be segmented using windowing frame concept. These frames can be occurred all over the length of the audio that help in the configuring the content in the audio. These various frame helps in identification and classification with predefined values in trained set of algorithm. Finally the feature extraction takes place helps in producing the exact output regarding the given input in the form of audio. It gives out the result that can be carried out after previous step of classification and identification. During the final step as feature extraction it helps in producing output by the above recommended steps.

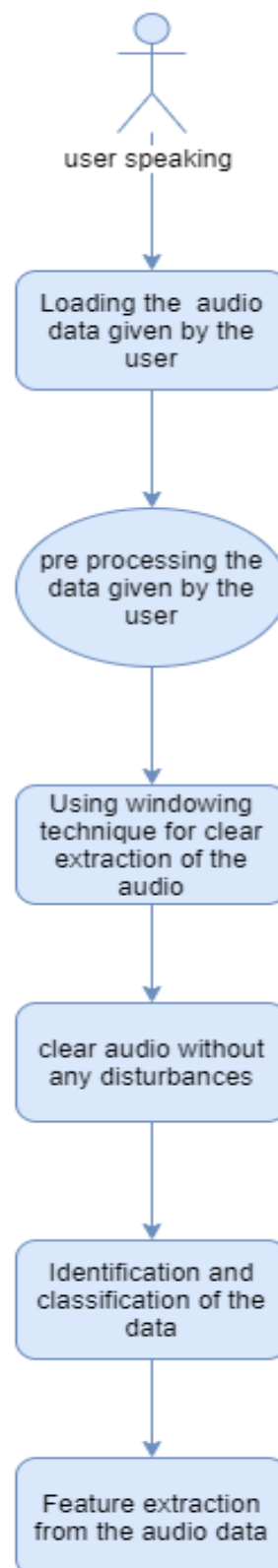


Figure (i): Implementation Process

```
C:\Users\SREEKANTH\Desktop\speaker_recognition>python Speech-to-text.py
listening ...
You said: hello hello hello hello hello
listening ...
You said: hi I am Pooja
listening ...
```

Figure (ii):Output Screenshot

To extract the output, initially we need to compile the python program. Here i have compiled my python file and got the output file. First off all, it will listen the audio that we have spoke for a while and it will listen the audio until the speaker speaks. And next it will print what we have spoke at the previous time. After printing the statements that we have spoke before again the process repeats continuously untill we close the current window.

5. Conclusion

Here regarding this project it has been handled issue of speech detection in mobile vehicular structure, by perceiving personality of one who speaks in Closed as well as Open-Set, then Sound obtaining is presented at various separations as well as that various kinds of natural commotion are taken.

It has introduced a vigorous speaker identification steps, that inserts a pre-handling technique for VAD as well as it exactly explored ideal appearance as well as edge determination foundation so as to get a strong as well as viable voiced/quietness sound recognition. It was additionally exhibited definite Exploration of exhibitions of regular Speaker voice methods, contrasted which are acquired by introducing VAD announcer detection, in, testing conditions that are skilled.

References

- [1] "A Vulnerability Test Method for Speech Recognition Systems Based on Frequency Signal Processing", Honghao Yang, Dong Liang, Xiaohui Kuang, Changqiao Xu, IEEE 2018.
- [2] "Dropout Approaches for LSTM Based Speech Recognition Systems", Jayadev Billa, IEEE 2018.
- [3] "Speech Recognition and Correction of a Stuttered Speech", Ankit Dash, Nikhil Subramani, Tejas Manjunath, Vishruti Yaragarala, Shikha Tripathi, IEEE 2018.
- [4] "Robust Audio-Visual Speech Recognition System based on Gabor Features and Dynamic Stream Weight Adaption", Ali Saudi, Mahmoud I. Khalil, Hazem Abbas, IEEE 2018.
- [5] "Fuzzy Evaluation System: Intelligent Speech Assessment System", Alireza Zourmand, Ting Hua Nong, IEEE 2018.
- [6] "Applying Machine Learning Techniques for Speech Emotion Recognition", K. Tarunika, R.B Pradeeba, P. Aruna, IEEE 2018.
- [7] "Adaptive Speech Recognition System Using Data From Keyboard", Sai Ajay Modukuri, Gaurav Kumar, IEEE 2018.
- [8] "Feature Fusion of Speech Emotion Recognition Based on Deep Learning", Gang Liu, Wei He, Bicheng Jin, IEEE 2018.
- [9] "A Speech Enhancement System for Automotive Speech Recognition with a Hybrid Voice Activity Detection Method", Haikun Wang, Zhongfu Ye, Jingdong Chen, IEEE 2018.
- [10] "Two-Stage Enhancement of Noisy and Reverberant Microphone Array Speech for Automatic Speech Recognition Systems Trained with Only Clean Speech", Qundong Wang, Sicheng Wang, Fengpei Ge, Chang Woo Han, Jaewon Lee, Lianghao Guo, Chih-Hui Lee, IEEE 2018.
- [11] "A Segmental CRF Speech Recognition System", Geoffrey Zweig, Patrick, Nguyen, MSR-TR-2009-54 | May 2009.
- [12] "Large vocabulary Speech Recognition System", YASUHISA FUJII, KAZUMASA YAMAMOTO, SEIICHI NAKAGAWA, Toyohashi University of Technology Department of Computer Science and Engineering, JAPAN.
- [13] "A Review on Speech Recognition Technique", Santosh K. Gaikwad, Bharti W. Gawali, Pravin Yannawar, Dr. Babasaheb, Associate Professor Department of CS& IT, Ambedkar Marathwada University, Aurangabad.
- [14] "An introduction to the application of the theory of probabilistic functions of a Markov process to automatic speech recognition", S. E. Levinson, L. R. Rabiner, M. M. Sondhi.
- [15] "Heterogeneous Acoustic Measurements and Multiple Classifiers for Speech Recognition", Andrew K. Halberstadt B.M, Eastman School of Music (1991) B.S., University of Rochester (1992) M.S., University of Rochester, Department of Electrical Engineering and Computer Science, MASSACHUSETTS INSTITUTE OF TECHNOLOGY, November 1998.
- [16] "A Maximum Likelihood Approach to Continuous Speech Recognition", Lalit R. Bahl, Frederick Jelinek, Robert L. Mercer, T. J. Watson MEMBER, IEEE.

- [17] “Maximum mutual information estimation of hidden Markov model parameters for speech recognition”, Bahl, P. Brown, P. de Souza, R. Mercer IBM Thomas J. Watson.
- [18] “Review of TDNN architectures for speech recognition”, M. Sugiyama, H. Sawai, ATR Interpreting Telephony, Kyoto, Japan.
- [19] “SPEECH RECOGNITION WITH WEIGHTED FINITE-STATE TRANSDUCERS “Mehryar Mohri, 31 Courant Institute 251 Mercer Street New York, NY 10012, Fernando Pereira, University of Pennsylvania 200 South 33rd Street Philadelphia, PA 19104.
- [20] “Speech recognition with deep recurrent neural networks”, by Alex Graves, Abdel-rahman Mohamed, Geoffrey Hinton Department of Computer Science, University of Toronto, Canada.
- [21] “Speech Recognition Technology”, by Ed Marcato, Kurzweil A. I., 411 Waverley Oaks Road, Waltham, MA 02154.
- [22] “Design of speech recognition system”, by Chao Wang, Ruifei Zhu, Hongguang Jia, Qun Wei, Huhai Jiang Tianyi, Zhang, Linyao Yu, Changchun Institute of Optics, Fine Mechanics and Physics, Chinese Academy of Sciences, CO 130033 China.
- [23] “Development of automatic speech recognition system for voice activated Ground Control system”, by Kavitha S, Veena S, R Kumaraswamy, SIT Tumakuru, CSIR - National Aerospace Laboratories Bengaluru, India.
- [24] “Research on Speech Recognition Technology and Its Application”, by Youhao Yu, Dept. of Electron. & Inf. Eng., Putian Univ., Putia, china.
- [25] “A speech recognition system”, by J. Guizol, H. Meloni, Groupe Intelligence Artificielle, Faculté des Sciences de Luminy, Marseille.