



Analytical Review for Sentiment Recognition – Machine Learning

¹ **Prof. Rahul Raut**

² **Ms. Siddhali Raut**

³ **Ms. Sneha Ambhore**

⁴ **Ms. Shubhada Thalkar**

¹Assistant Professor, Department of Information Technology, Sandip Institute of Technology and Research Center, Nashik, India,

²Department of Information Technology, Sandip Institute of Technology and Research Center, Nashik, India,

³Department of Information Technology, Sandip Institute of Technology and Research Center, Nashik, India,

⁴Department of Information Technology, Sandip Institute of Technology and Research Center, Nashik, India

ABSTRACT:

Sentiment recognition from audio signal requires module to extract feature and classifier for training. Here we define the feature vector consists of sound packets signal which is differentiate speaker specified features such as pitch, energy, tone, which is important to train the sentiment recognition model to detect and recognize a particular sentiment perfectly. We consider here open source dataset for data training, which is available in English language. Sound vocal tract information, represented by Mel-frequency cepstral coefficients (MFCC), was extracted from the audio samples in training dataset. The pitch or that particular sound extracted, anger, happiness, sad, neutral, etc sentiments. The test data extract procedure followed for which of the classifier would make a prediction regarding the underlying sentiment in the test audio.. The paper details the two methods applied on feature vectors and the effect of increasing the number of feature vectors fed to the classifier. Our module provides an accuracy of analytical classification for an Indian English speech. The accuracy for Indian English was 73 percent.

Keywords: [Machine Learning, SVM, Classifier, K-means, Matplotlib, sk-learn, pandas, data analysis]

I. INTRODUCTION

“To Simplify the machine understanding about the sentiments behind the speech or any kind of conversation using machine learning. So, the machine can understand people conversation and working for them.”

Speech sentiment recognition, The best example of it can be seen at call centers. If you noticed, BPO centers employees never talk in the expected sound of communication, their way of communication to the client changes with clients [1]. This is also happen with common people too, how can they tackle the calls in BPO centers? Here is the answer, the employees recognize clients sentiments from speech, so they can improve their service and convert more people. This is proper way to communicate; they are using speech sentiment recognition [2]. Speech Sentiment Recognition, abbreviated as SER, is the act of attempting to recognize human sentiment and affective states from speech. This is bold part of the fact that voice/sound often reflects underlying sentiment through tone and pitch. This is also the scenario of that animals like dogs and horses employ to be able to understand human sentiment [3].

Have you ever called a phone service? How was your experience? It is usually very irritating, when you have a bot asking you a lot of questions. Now suppose you are upset, and you decide to call the company, and still get a robot on the other end of your call. That is an example in which you could try to recognize speech sentiment with machine learning and improve customer services. Adding sentiments to machines has been

recognized as a critical factor in making machines appear and act in a human-like manner [4]. Speech Sentiment Recognition (SER) is the task of recognizing the sentiment from speech irrespective of the semantic contents. However, sentiments are subjective and even for humans it is hard to notate them in natural speech communication regardless of the meaning. The ability to automate the task and conduct constantly, it's very difficult and still an ongoing part of research. We measure the height of this wave at equal time of intervals. The exactly what is sampling rate means. You are basically doing more that thousands of samples every second by second. Each point represents the amplitude of that audio/sound file at the instant and will be turned into an list of numbers [5].

II LITERATURE REVIEW

Support Vector Machines to Classify Data Attentiveness for the Development of Personalized Learning Method

There have been many methods in which researchers have attempted to classify student attentiveness. Many of these approaches depended on a sound quality analysis and lacked any quantitative analysis. We focused on bridging the gap between qualitative approaches and quantitative approaches to classify student attentiveness [2].

Restricted Boltzmann Machine for Linear and Nonlinear System Modelling

The use of a deep learning method restricted Boltzmann machine, for nonlinear system identification. The neural-network model has deep architecture and is generated by an auto random search method. The initial weights of this deep neural-network model are obtained from the restricted Boltzmann machines. To identify nonlinear systems model, we propose special unsupervised learning methods with user input data. The normal supervised learning is used to train the weights with the output dataset.[6]

Classification Techniques used in Machine Learning as Applied on Vocational Guidance Data model

The developments in information systems as well as computerization of business processes by so many organizations have led to a faster, easier and more powerfully accurate data analysis. Machine learning techniques and algorithms make it possible to deduct meaningful further information from those data processed by data mining [7].

Vision Transformer Backbones with Machine Learning Technology

We explore the plain, non-hierarchical Vision Transformer as a backbone network

for sound detection. This design enables the original architecture and system to be fine-tuned for object detection without needing to redesign a hierarchical backbone for pre-training. Surprisingly, we observe: (i) it is proper to build a simple feature of pyramid from a single-scale feature map (without the common design) and (ii) it is proper to use window attention without shifting need, aided with very few cross window propagation blocks. With plain backbones pre-trained as Masked Auto encoders, our detector.

Neural networks in JAX via callable PyTrees and filtered transformations-ML

PyTorch are popular Python auto differentiation frameworks. JAX is based around pure range of a functions and functionality programming. PyTorch has popularized the use of an object-oriented programming class-based syntax for defining parameterized functions, such as neural networks. That this seems like a fundamental difference means current libraries use to define module for building parameterized functions in JAX have either rejected the object oriented approach entirely (Stacks) or have introduced object oriented -to-functional transformations, multiple new abstractions, and been limited in the extent to which they integrate with JAX. Parameterized functionalities are themselves represented as PyTrees, which means that the parameterizations of functions are transparent to the JAX framework module.

III. PROPOSED SYSTEM

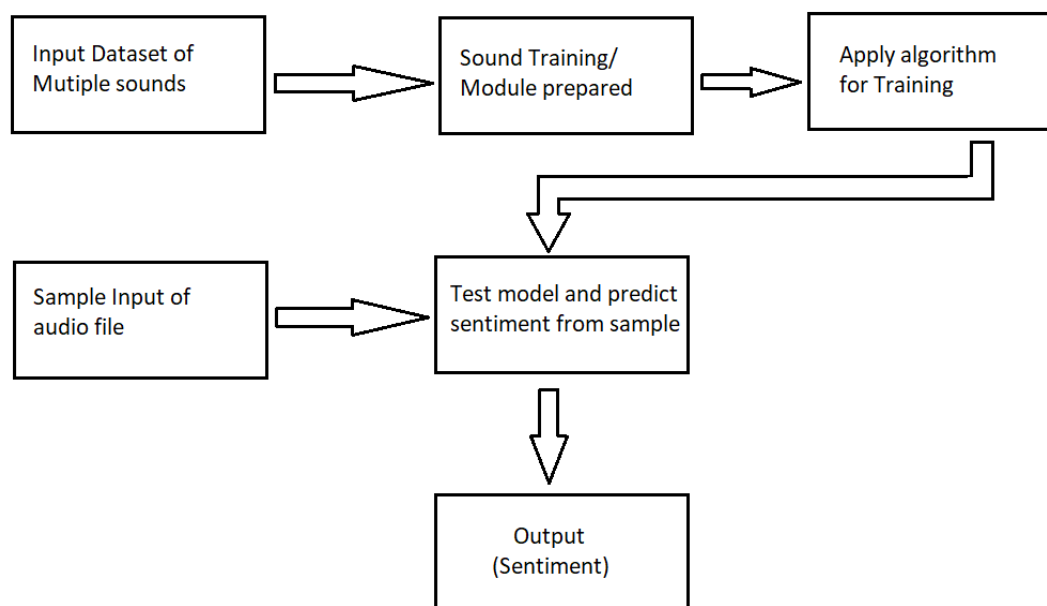


Fig. 1 Proposed System flow

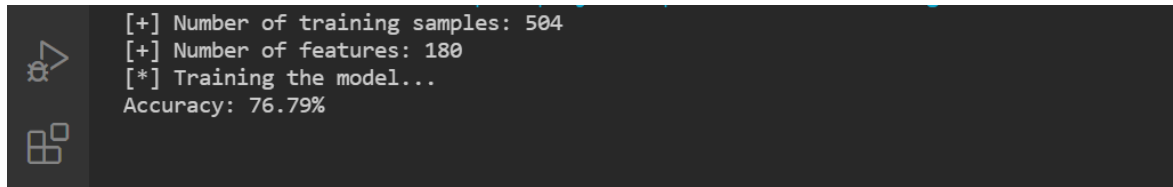
i. About Dataset

Here, we use an open source dataset for model training. In this dataset, we covered some common sentiments such as *happy*, *sad*, *angry*, *neutral*, etc. For a data training module, we created some addition words of sound and merge into the previous dataset for more accuracy. A Panda is an important module to manage and visualize data. We use VLC player to play our sound files, the VLC is also an open source tool and anyone can easily access this.

ii. Module training

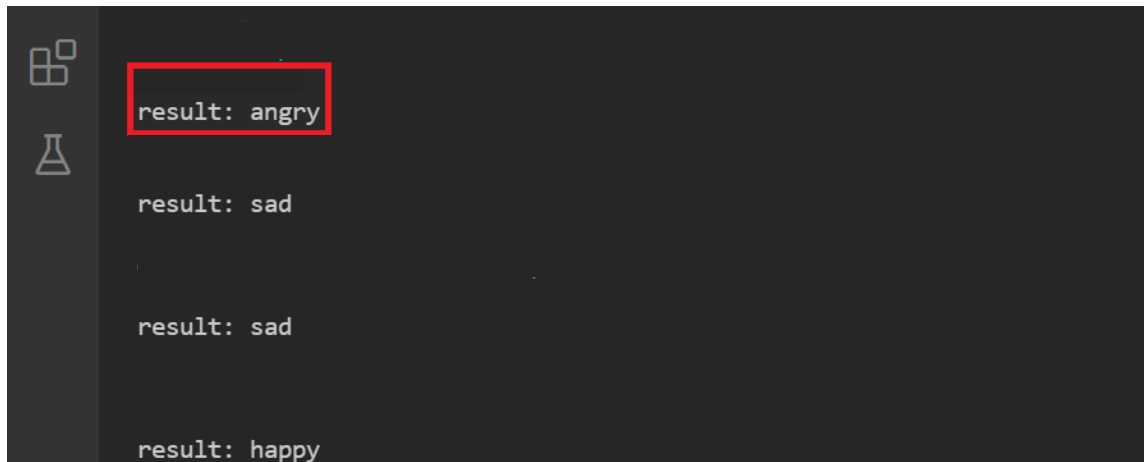
In this section, we use a multiple python modules to process and train the data. The **SKLEARN** is our first and main module; in this module we use *MLPClassifier* from neural network class. To load the model from our dataset file and start working on using *MLPClassifier*. Finally, the **pickle** module created a model extension file as a result.

iii. Results



```
[+] Number of training samples: 504
[+] Number of features: 180
[*] Training the model...
Accuracy: 76.79%
```

Image 1. Accuracy of our model



```
result: angry
result: sad
result: sad
result: happy
```

Image 2. Testing and detected sentiments

IV. SIMULATION MODULE

- q₀ = Start
- q₁ = Data Cleaning
- q₂ = Process Data (Apply Classifiers to prepare for training)
- q₃ = Data Training (To Prepare model)
- q₄ = Verify model
- q₅ = Model Testing
- q₆ = End node

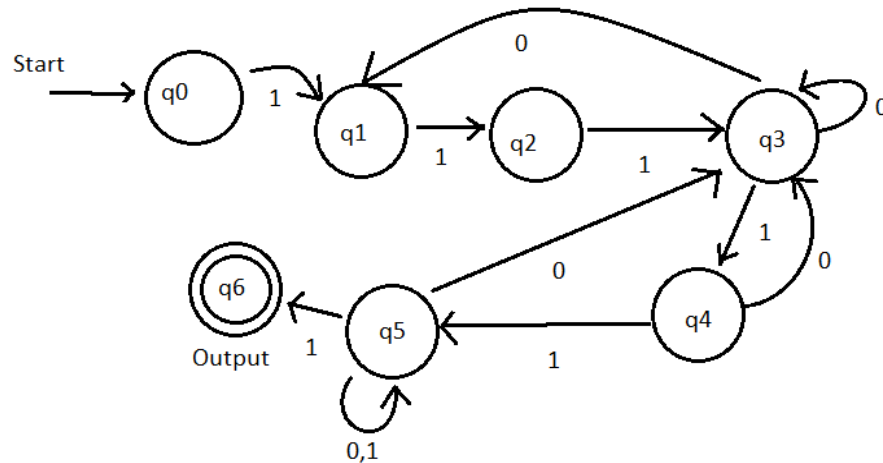


Fig. 2 Simulation model

CONCLUSION

In this paper we proposed a method that used the deeply checked sentiments behind of audio, one of the deep neural networks, to extract the sentimental characteristic parameter from sentimental speech signal automatically. We combined deep ML network and support vector machine (SVM) and proposed a classifier model which is based on machine learning and support vector machine (SVM). In the practical training process, the model has small complexity and 7% higher final recognition rate than traditional artificial extract, and this method can extract sentiment characteristic parameters accurately, improving the recognition rate of sentimental speech recognition obviously. But the time cost for training DBNs feature extraction model was 136 hours, and it was longer than other feature extraction methods.

In future work, we will continue to further study speech sentiment recognition based on DBNs and further expand the training data set. Our ultimate aim is to study how to improve the recognition rate of speech sentiment recognition.

REFERENCES

- [1]. Burkhardt, F.; Ajmera, J.; Englert, R.; Stegmann, J.; Burleson, W. Detecting anger in automated voice portal dialogs. In Proceedings of the Ninth International Conference on Spoken Language Processing, Pittsburgh, PA, USA, 17–21 September 2006.
- [2]. Hossain, M.S.; Muhammad, G.; Song, B.; Hassan, M.M.; Alelaiwi, A.; Alamri, A. Audio–Visual Sentiment-Aware Cloud Gaming Framework. *IEEE Trans. Circuits Syst. Video Technol.* **2015**, *25*, 2105–2118. [[CrossRef](#)]
- [3]. Oh, K.; Lee, D.; Ko, B.; Choi, H. A Chatbot for Psychiatric Counseling in Mental Healthcare Service Based on Sentimental Dialogue Analysis and Sentence Generation. In Proceedings of the 2017 18th IEEE International Conference on

Mobile Data Management (MDM), Daejeon, Korea, 29 May–1 June 2017; pp. 371–375.

[4]. Yenigalla, P.; Kumar, A.; Tripathi, S.; Singh, C.; Kar, S.; Vepa, J. Speech Sentiment Recognition Using Spectrogram & Phoneme Embedding. In Proceedings of the INTERSPEECH, Hyderabad, India, 2–6 September 2018.

[5]. Deriche, M.; Abo absa, A.H. A Two-Stage Hierarchical Bilingual Sentiment Recognition System Using a Hidden Markov Model and Neural Networks. Arab. J. Sci. Eng. **2017**, 42, 5231–5249. [[CrossRef](#)]

[6]. Pravena, D.; Govind, D. Significance of incorporating excitation source parameters for

improved sentiment recognition from speech and electroglottographic signals. Int. J. Speech Technol. **2017**, 20, 787–797. [[CrossRef](#)]

[7]. Bandela, S.R.; Kumar, T.K. Stressed speech sentiment recognition using feature fusion of teager energy operator and MFCC. In Proceedings of the 2017 8th International Conference on Computing, Communication and Networking Technologies (ICCCNT), Delhi, India, 3–5 July 2017; pp. 1–5.

[8]. Koolagudi, S.G.; Murthy, Y.V.S.; Bhaskar, S.P. Choice of a classifier, based on properties of a dataset: Case study-speech sentiment recognition. Int. J. Speech Technol. **2018**, 21, 167–183. [[CrossRef](#)]

