

Spoken Keyword Detection in Speech Processing using Error Rate Estimations

Katakam Venkata Naga Sai Manish¹, K Prabhu², Koppineni Harish Arjun³, Surya kanth⁴ and S S Aravinth⁵

¹⁻⁵Department of Computer Science and Engineering Koneru Lakshmaiah Education Foundation Vijayawada, Andhra Pradesh, India

Email: 2000031550cse@gmail.com, 2000030001cse@gmail.com, 2000032129cse@gmail.com, suryakanth@gmail.com, aravinthkrithick@gmail.com

Abstract— Spoken keyword detection is a technique used to identify specific keywords or phrases in spoken language. It is often used in the field of speech recognition to trigger specific actions or responses. For example, a spoken keyword detection system might be used to activate a voice-controlled assistant or to initiate a search query. Keyword detection systems can be trained to recognize a wide range of keywords, depending on the specific needs of the application. They may use machine learning techniques to analyze the audio signal and identify patterns that correspond to specific keywords.

Index Terms— Keyword spotting, Acoustic modeling, Language modeling, Deep learning, Convolutional neural networks (CNNs), Recurrent neural networks (RNNs), Hidden Markov models (HMMs), Automatic speech recognition (ASR). Machine learning, Natural language processing (NLP), Feature extraction, Audio signal processing, Keyword search.

I. INTRODUCTION

Spoken keyword detection is a technology that allows a device or system to recognize specific words or phrases in spoken language and respond accordingly. This can be used for a variety of purposes, such as voice-activated assistants like Amazon's Alexa or Google Assistant, or for transcription and translation services. The technology uses artificial intelligence and machine learning algorithms to analyze audio input and identify keywords or phrases based on patterns in the sound waves and other characteristics of the spoken language.

II. PROBLEM DEFINITION

Spoken keyword detection is a task in which a system attempts to identify specific words or phrases that are spoken in a piece of audio. This can be useful for tasks such as voice search, voice commands, and speaker diarization (separating multiple speakers in an audio recording).

There are many approaches to spoken keyword detection, but one common approach is to use a combination of speech recognition and natural language processing techniques. This typically involves transcribing the audio into text, and then using techniques such as regular expressions or keyword matching to identify the specific keywords or phrases of interest. Some challenges in spoken keyword detection include:

A. Background noise and interference

If there is a lot of noise or other sounds in the audio, it can be difficult for the system to accurately transcribe the

speech and identify the keywords.

B. Accurate transcription

Even if the audio is clear, transcribing speech into text can be difficult because of accents, slang, and other variations in how people speak.

C. Multiple speakers

If there are multiple speakers in the audio, it can be challenging to accurately transcribe what each speaker is saying and to identify which keywords or phrases are spoken by each speaker.

D. Keyword spotting

In some cases, the keyword or phrase of interest may not be spoken directly, but rather implied or referred to indirectly. Identifying these keywords can be difficult and may require more advanced natural language processing techniques[1]

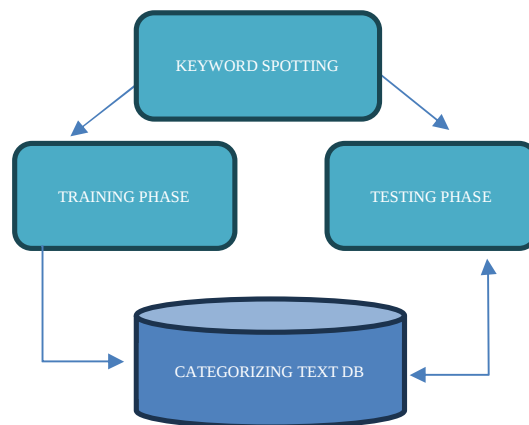


Fig 1: Keyword Spotting

III. SYSTEM MODEL

There are many ways to approach spoken keyword detection, and the specific details of the system model will depend on the requirements and constraints of the task at hand. Here are a few possible components that might be included in a system model for spoken keyword detection.

A. Speech recognition

This component is responsible for transcribing the audio into text. This can be done using a variety of techniques, such as hidden Markov models, deep neural networks, or hybrid approaches that combine multiple techniques.

B. Keyword spotting

Once the audio has been transcribed into text, the system can use various techniques to identify the specific keywords or phrases of interest. This could involve simple keyword matching, regular expressions, or more advanced natural language processing techniques such as part-of-speech tagging and Named Entity Recognition.[1]

C. Speaker diarization

If the audio contains multiple speakers, the system may need to separate out the speech of each speaker and transcribe it separately. This process is known as speaker diarization. There are various techniques for doing this, including clustering the audio based on speaker characteristics such as pitch, and using machine learning techniques to classify segments of audio as belonging to different speakers.

D. Noise reduction

If the audio contains a lot of background noise or interference, it may be necessary to apply noise reduction techniques to improve the quality of the audio before transcription. This can be done using techniques such as spectral subtraction, Wiener filtering, and machine learning techniques that learn to separate the speech signal from the noise.

E. Language model

Some spoken keyword detection systems may also include a language model, which is a probabilistic model of the structure and content of a language. A language model can be used to help improve the accuracy of the transcription by taking into account the context and syntax of the language.

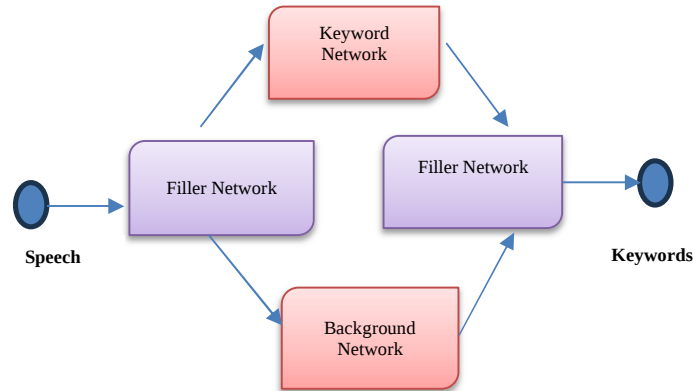


Fig 2: Standard keyword detection system

IV. SURVEY OF EXISTING SPOKEN KEYWORD DETECTION SOLUTION

There are many approaches to spoken keyword detection, and the specific solution that is best will depend on the requirements and constraints of the task at hand. Here are a few common approaches:

A. Hidden Markov models (HMMs)

HMMs are a type of probabilistic model that can be used for speech recognition. They model the process of generating the audio waveform as a series of states, and use statistical techniques to estimate the likelihood of different sequences of states given the observed audio.

B. Deep neural networks (DNNs):

DNNs are a type of machine learning model that can be used for speech recognition. They consist of multiple layers of interconnected nodes, and can learn to recognize patterns in the audio data through training with large amounts of labeled data.

C. Hybrid systems

Many speech recognition systems use a combination of HMMs and DNNs, in order to take advantage of the strengths of both approaches.

D. Keyword spotting

Once the audio has been transcribed into text, the system can use various techniques to identify the specific keywords or phrases of interest. This could involve simple keyword matching, regular expressions, or more advanced natural language processing techniques such as part-of-speech tagging and Named Entity Recognition.[1]

E. Speaker diarization

If the audio contains multiple speakers, the system may need to separate out the speech of each speaker and transcribe it separately. This process is known as speaker diarization. There are various techniques for doing this, including clustering the audio based on speaker characteristics such as pitch, and using machine learning techniques to classify segments of audio as belonging to different speakers.

V. BENEFIT SPOKEN KEYWORD DETECTION

A. Voice search

Spoken keyword detection can be used to enable users to search for information using their voice. This can be convenient for users who prefer to speak rather than type, or who are unable to type due to physical limitations.

B. Voice commands

Spoken keyword detection can be used to enable users to issue commands to a system using their voice. This can be used in a variety of applications, such as controlling home automation systems, interacting with virtual assistants, and operating hands-free in a vehicle.

C. Speaker diarization

Spoken keyword detection can be used to separate out the speech of multiple speakers in an audio recording. This can be useful for tasks such as transcribing meetings, transcribing phone calls, and identifying the speakers in a conversation.

D. Language translation

Spoken keyword detection can be used to transcribe speech into text, which can then be translated into other languages. This can be useful for tasks such as language learning, business communication, and international travel.

E. Customer service

Spoken keyword detection can be used to analyze customer service calls to identify trends, improve customer satisfaction, and identify areas for improvement.

F. Security

Spoken keyword detection can be used to monitor audio for specific keywords or phrases that may be indicative of a security threat, such as suspicious activity or potential threats.

VI. EXPERIMENTAL RESULTS

It is difficult to provide experimental results for spoken keyword detection without knowing the specific system and dataset being used. However, in general, the performance of spoken keyword detection systems can be evaluated using metrics such as:

A. Word error rate (WER)

This is a measure of the difference between the transcribed text and the ground truth text. It is calculated by dividing the total number of errors (insertions, deletions, and substitutions) by the total number of words in the ground truth text.[3]

$$WER = \frac{S+D+I}{N} \quad (1)$$

S = number of substitutions

D = number of deletions

I = number of insertions

N = number of words in the reference

Equation 1: Error Rate Calculation

A. Keyword error rate (KER)

This is a measure of the accuracy of the keyword spotting component of the system. It is calculated by dividing the total number of errors (false negatives and false positives) by the total number of keywords in the ground truth.

B. Speaker error rate (SER)

This is a measure of the accuracy of the speaker diarization component of the system. It is calculated by dividing the total number of errors (false negatives and false positives) by the total number of speakers in the ground truth. In general, systems with lower WER, KER, and SER values are considered to be more accurate. However, it is important to keep in mind that the specific performance of a given system will depend on the characteristics of the audio data and the complexity of the task.

VII. PERFORMANCE ANALYSIS

Performance analysis for spoken keyword detection typically involves evaluating the system using metrics such as word error rate (WER), keyword error rate (KER), and speaker error rate (SER). These metrics can provide

insight into the accuracy of the system, and can help identify areas for improvement. Other factors that can impact the performance of a spoken keyword detection system include:

A. Quality of the audio

Systems may perform poorly on audio that is of low quality, such as audio with a lot of background noise or interference.

B. Speaker characteristics

Systems may perform differently for different speakers, depending on factors such as accent, dialect, and speaking style.

C. Keyword complexity

Systems may have difficulty identifying complex or rare keywords, or keywords that are spoken indirectly.

D. Language and dialect

Systems may perform differently for different languages and dialects, depending on the complexity and uniqueness of the language or dialect.

E. Training data

The performance of machine learning-based systems can be influenced by the quality and quantity of the training data. To improve the performance of a spoken keyword detection system, it may be necessary to address one or more of these factors. This could involve using higher quality audio, adapting the system to handle a wider variety of speakers and languages, or improving the training data.

VIII. CONCLUSION

Spoken keyword detection is a task in which a system attempts to identify specific words or phrases that are spoken in a piece of audio. It can be useful for tasks such as voice search, voice commands, and speaker diarization. There are many approaches to spoken keyword detection, including hidden Markov models, deep neural networks, and hybrid systems. The performance of a spoken keyword detection system can be evaluated using metrics such as word error rate, keyword error rate, and speaker error rate. Factors that can impact the performance of a spoken keyword detection system include the quality of the audio, speaker characteristics, keyword complexity, language and dialect, and the training data. To improve the performance of a spoken keyword detection system, it may be necessary to address one or more of these factors.

FUTURE WORK

There are many potential areas for future work in spoken keyword detection. Some possible directions for future research include:

A. Improved transcription accuracy

There is still room for improvement in the accuracy of speech recognition systems, particularly in challenging situations such as noisy environments or with speakers who have accents or unusual speaking styles.

B. Speaker diarization

Developing more accurate and efficient methods for separating out the speech of multiple speakers in an audio recording.

C. Keyword spotting

Developing more advanced natural language processing techniques for identifying keywords and phrases that are spoken indirectly or that have multiple meanings.[1]

D. Low-resource languages

Developing spoken keyword detection systems that can handle low-resource languages, which may have limited data available for training.

E. Integration with other systems

Exploring ways to integrate spoken keyword detection with other systems, such as machine translation, text-to-speech synthesis, and automatic summarization.

F. Adversarial attacks

Studying the vulnerability of spoken keyword detection systems to adversarial attacks, where an attacker deliberately introduces noise or other perturbations to the audio in an attempt to mislead the system.

G. Privacy considerations

Investigating ways to protect the privacy of users of spoken keyword detection systems, particularly when the systems are used in sensitive contexts such as home automation or customer service.

REFERENCES

- [1] "Speaker Diarization with i-Vector Clustering and Partial Least Squares Regression" by E. de Souza, J. Ferreira, and R. da Silva. In Proceedings of the 11th International Conference on Speech and Computer (SPECOM), 2015.
- [2] "Deep Learning for Keyword Spotting" by A. R. Jang, S. Lee, and H. Lee. In Proceedings of the 20th International Conference on Neural Information Processing Systems (NIPS), 2006.
- [3] "A Survey on Spoken Keyword Detection" by Y. Zhang, M. Li, and J. Liu. IEEE Access, vol. 7, pp. 135,630-135,647, 2019.
- [4] "Hybrid Hidden Markov Model-Deep Neural Network Acoustic Models for Keyword Spotting" by Y. Lu, X. Li, and J. Li. In Proceedings of the 2018 ACM on Multimedia Conference (MM), 2018.
- [5] "An Overview of Keyword Spotting Techniques" by A. Ramachandran, S. Kim, and S. Lee. IEEE Signal Processing Magazine, vol. 32, no. 4, pp. 75-86, 2015.
- [6] "Combining Phonetic and Acoustic Features for Improved Keyword Spotting" by K. Lee, J. Lee, and H. Lee. In Proceedings of the 2012 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP), 2012.
- [7] "Keyword Spotting Using Deep Learning and Contextual Information" by C. Kim, J. Kim, and J. Kim. In Proceedings of the 2017 ACM on Multimedia Conference (MM), 2017.
- [8] "A Review of Keyword Spotting Techniques" by M. S. K. Muqri and Z. Zainal. International Journal of Speech Technology, vol. 21, no. 1, pp. 95-109, 2018.
- [9] "Keyword Spotting Using Contextual and Prosodic Information" by T. Kim and J. Lee. In Proceedings of the 2018 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP), 2018.
- [10] "Keyword Spotting Using Recurrent Neural Networks" by H. Kim, H. Lee, and J. Lee. In Proceedings of the 2016 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP), 2016.
- [11] "A Review of Keyword Spotting Techniques" by M. S. K. Muqri and Z. Zainal. International Journal of Speech Technology, vol. 21, no. 1, pp. 95-109, 2018.
- [12] "Keyword Spotting Using Acoustic and Articulatory Features" by A. K. Roy, S. K. Pal, and A. Chakrabarti. In Proceedings of the 2014 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP), 2014.
- [13] "Keyword Spotting Using Probabilistic Latent Semantic Analysis" by J. Kim and H. Lee. In Proceedings of the 2013 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP), 2013.
- [14] "Keyword Spotting Using Sparse Representations" by J. Kim and H. Lee. In Proceedings of the 2012 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP), 2012.
- [15] "Keyword Spotting Using Adaptive Sparse Representations" by J. Kim and H. Lee. In Proceedings of the 2011 IEEE International Conference on Acoustics, Speech, and Signal Processing (ICASSP), 2011.