



HUMAN SPEECH RECOGNITION FOR RECOMMENDATION IN CONVERSATIONS USING KEYWORD EXTRACTION AND CLUSERING

Aarti korade¹, Gauri Jaydeokar², Nikita Kangane³, Juili Kedari⁴

Prof. Manjusha Tatiya⁵

¹ aarukorade@gmail.com, ² gaurijaydeokar97@gmail.com, ³ nikikangane1998@gmail.com, ³ juilik157@gmail.com,
⁵ manjusha.tatiya@indiraicem.ac.in

¹Computer Engineering, Indira College of Engineering and Management, Pune University

²Computer Engineering, Indira College of Engineering and Management, Pune University

³Computer Engineering, Indira College of Engineering and Management, Pune University

⁴Computer Engineering, Indira College of Engineering and Management, Pune University

⁵Computer Engineering, Indira College of Engineering and Management, Pune University

Abstract — speech Recognition is the process of automatically recognizing a word spoken by a specific speaker based on that information which are included in speech waves. Now days there are rapid growth of wireless communications, hence the need for voice recognition techniques has been increased greatly.

In this Paper we are going to describe the problem of keyword extraction from conversations, with this paper user can access the keywords which are document relevant for each short conversation fragment, which can be recommended to users. Sometimes, even a short fragment includes different types of words, which are potentially related to different topics. However, errors can be introducing into system or conversation by using automatic speech recognition system. To extract keyword from the output of ASR technique we use one algorithm. It is based on the sub modular functions which gives range of different words and reduces the noise. Then use a method to obtain a multiple topic from this keyword set only for to access the one relevant recommended document

Keywords: information retrieval, keyword extraction, Speech to Text, speech recognition.

I. INTRODUCTION

Every-day large amount of information is generated such as the documents, databases, or multimedia resources. Access to this information is conditioned by the availability of suitable search engines, but the todays humans are very much busy in their activity since they should not get this relevant information, they are not aware that relevant information is available. Through this paper we proposed a method using to this user can easily access the relevant recommended documents. For instance, when users participate in a meeting, their information needs can be modeled as implicit queries that are constructed in the background from the pronounced words, obtained through real-time automatic speech recognition (ASR). To retrieve and recommend documents from the Web or a local repository implicit queries are used, which users can choose to inspect in more detail if they find them interesting. The main purpose of this paper is on formulating implicit queries to a just-in-time-retrieval system for use in meeting rooms. Diversity and Relevance can be started at three stages: from extracting the keywords; from building one or several implicit queries; or when we are going to re-ranking their results. The first two things are the main point of this paper. The re-ranking results of a single implicit query is not able to improve satisfaction with the recommended documents. Methods which are proposed previously for formulating implicit queries from text rely on word frequency. The paper sequence is as follows. we consider existing just-in-time-retrieval systems and the policies they use for query formulation. We used previous methods for keyword extraction. we are going to define our proposed technique for implicit query formulation, which is based on keyword extraction algorithm and a topic-aware clustering method. For instance, while a method based on word

frequency would retrieve the following Wikipedia pages: „Light“, „Lighting“, and „Light My Fire“ for the above-mentioned fragment, users would prefer a set such as “Lighter“, „Wool“ and „Chocolate“. Diversity and Relevance can be started at three stages: from extracting the keywords; from building one or several implicit queries; or from re-ranking their results. Previous methods for formulating implicit queries from text rely on word frequency or TFIDF (Term Frequency -Inverse Document Frequency) weights to rank keywords and then select the highest-ranking ones. Other methods perform keyword extraction by using topical similarity but do not set a topic diversity constraint. In this paper, we introduce a novel keyword extraction technique from ASR output, which maximizes the coverage of potential information needs of users and reduces the number of irrelevant words. When all the set of keywords is got extracted, it is clustered in order to build several queries related to Topics, which are run independently in next phase, offering better precision than a larger, topically-mixed query. The final Results are merged into a ranked set before showing them as recommendations to users. In this Paper we are going to address the topics such as just-in-time-retrieval systems and the policies which are used for query formulation, In previous methods for keyword extraction as well as the proposed technique for implicit query formulation, which depends on a novel topic-aware keyword extraction algorithm and a topic-aware clustering method.

II. EXISTING SYSTEM

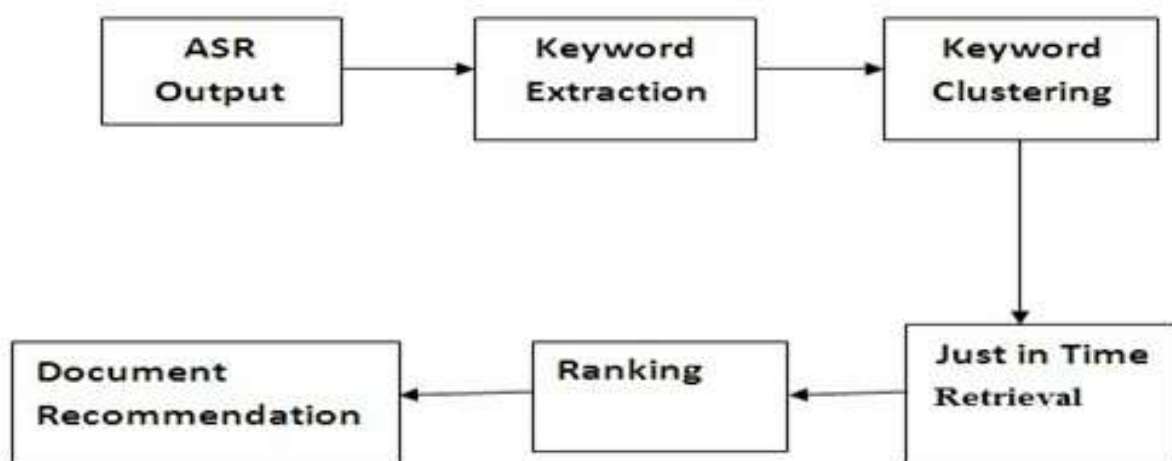
A short piece contains a mixture of words, which are possibly identified with a few subjects, in addition, utilizing a programmed discourse acknowledgment (ASR) framework presents blunders among them. In this way, it is hard to induce accurately the data needs of the discussion member.

Disadvantages in Existing System:

- 1)Each short conversation fragment contains a small number of potentially relevant documents, which can be recommended to participants.
- 2)Automatic speech recognition (ASR) system introduces errors.

III. PROPOSED SYSTEM

Propose a calculation to separate decisive words from the yield of an ASR framework (or a manual transcript for testing), which makes utilization of theme demonstrating methods and of a sub modular prize capacity which supports differing qualities in the catchphrase set, to coordinate the potential assorted qualities of points and lessen ASR commotion. At that point, we propose a technique to infer different topically isolated questions from this magic word set, with a specific end goal to expand the shots of making no less than one important suggestion when utilizing these inquiries to pursuit over the English Wikipedia.



Advantages in Proposed System:

- 1) Represents a promising solution for a document recommender system to be used in conversations.
- 2) Formulating implicit queries to a just-in-time-retrieval system for use in meeting rooms.

RELEVANT MATHEMATICS ASSOCIATED WITH THE PROJECT

Let S is the Whole System Consist of:

$S = U, D, ASR, DKE, KC, QF, O.$

$U = \text{User.}$

$U = u_1, u_2, \dots, u_n$

$D = \text{Dataset.}$

$D = d_1, d_2, \dots, d_n.$

(ASR)= Automatic Speech Recognition

(DKE)= Diverse keyword extraction.

KC = Keyword Clustering

QF = Query Formulation

O = Output.

Procedure:

Keyword Extraction:

ASR: automatic speech recognition converts the speech and provides output to algorithm that extract keywords from the output of an ASR system

Selection of Configurations:

By Using the RBO- rank biased overlap as a similarity metric, which is based on the fraction of keywords, overlapping at different ranks.

$$RBO(S, T) = \frac{1}{\sum_{d=1}^D \left(\frac{1}{2}\right)^{d-1}} \sum_{d=1}^D \left(\frac{1}{2}\right)^{d-1} \frac{|S_{1:d} \cap T_{1:d}|}{|S_{1:d} \cup T_{1:d}|}$$

Where,

RBO is nothing but rank biased overlap

S and T be two ranked lists,

and S_i is the keyword at rank i in S The set of the keywords up to rank d in S.

Diverse Keyword Extraction:

The advantage of diverse keyword extraction is that the main topics maximization coverage of the conversation fragment. There are Three Steps for the proposed method - diverse keyword extraction.

1. Used to represent the distribution of the abstract topic for each word.
2. These topic models are used to determine weights for the abstract topics in each conversation fragment represented by z

3. The keyword list $W = w_1, w_2, \dots, w_k$. Which covers a maximum number of the most important topics are selected by rewarding diversity, using an original algorithm introduced in this section.

IV. CONCLUSION

In this Paper we have considered a form of just-in-time retrieval systems for conversational environments, in which they are going to recommend to user's documents that are related to their information needs. Basically, we focused on the user's information which needs by deriving implicit queries from short conversation. These queries are related to or based on sets of keywords which are extracted from the conversation. We have our proposed technique called as a novel diverse keyword extraction technique which is basically used for covering the maximal count of important topics in a fragment. Then, for reducing the noise from the conversation, we going to propose a clustering technique which is basically used for dividing the set of keywords into smaller subsets (topically-independent) constituting implicit queries. We compared the diverse keyword extraction technique with existing methods depended on word frequency or topical similarity in terms of priority of the keywords and the relevance of recommended documents. In Our proposed system our current aim is to process external queries, and giving rank to document results with the intension of maximizing the coverage of all the information, while minimizing redundancy in a short list of documents. Integrating these techniques in a working prototype should.

V. REFERENCES

- [1]. M. Habibi and A. Popescu-Belis, Enforcing topic diversity in a document recommender for conversations, in Proc. 25th Int. Conf. Comput. Linguist. (Coling), 2014, pp. 588599.
- [2]. S. Ye, T.-S. Chua, M.-Y. Kan, and L. Qiu, Document concept lattice for text understanding and summarization, Inf. Process. Manage, vol. 43, no. 6, pp. 16431662, 2007.
- [3]. Harwath and T. J. Hazen, Topic identification based extrinsic evaluation of summarization techniques applied to conversational speech, in Proc. Int. Conf. Acoust., Speech, Signal Process. (ICASSP), 2012, pp. 50735076.
- [4] A. Popescu-Belis, E. Boertjes, J. Kilgour, P. Poller, S. Castronovo, T. Wilson, A. Jaimes, and J. Carletta, The AMIDA automatic content linking device: Just-in-time document retrieval in meetings, in Proc. 5th Workshop Mach. Learn. Multimodal Interact. (MLMI), 2008, pp. 272283.
- [5] B. J. Rhodes and P. Maes, Just-in-time information retrieval agents, IBM Syst. J., vol. 39, no. 3.4, pp. 685704, 2000.
- [6] D. Traum, P. Aggarwal, R. Artstein, S. Foutz, J. Gerten, A. Katsamanis, A. Leuski, D. Noren, and W. Swartout, Ada and Grace: Direct interaction with museum visitors, in Proc. 12th Int. Conf. Intell. Virtual Agents, 2012, pp. 245251.
- [7] A. S. M. Arif, J. T. Du, and I. Lee, Examining collaborative query reformulation: A case of travel information searching, in Proc. 37th Int. ACM SIGIR Conf. Res. Develop. Inf. Retrieval, 2014, pp. 875878.
- [8] A. S. M. Arif, J. T. Du, and I. Lee, Towards a model of collaborative information retrieval in tourism, in Proc. 4th Inf. Interact. Context Symp., 2012, pp. 258261.
- [9] J. Zaino, MindMeld makes context count in search, [Online]. Available: <http://semanticweb.com/mindmeld-makes-context-countsearch> b42725 2014
- [10] M. Habibi and A. Popescu-Belis, Using crowdsourcing to compare document recommendation strategies for conversations, Workshop Recommendat. Utility Eval.: Beyond RMSE (RUE11), pp. 1520, 2012.