

検索語の説明文による音声内容検索を利用した 音声検索語検出

南條 浩輝^{1,a)} 川口 達也²

概要： 音声の中で検索語がそのまま現れる部分を特定する音声検索語検出 (Spoken Term Detection: STD) の研究を行う。一般的に STD では、検索語と検索対象の音声ドキュメントをそれぞれをシンボル列に変換し、音声認識誤りを考慮しつつそれらの一致度に基づいて検出を行う。このため、同音異義語や近い文字列などを誤検出する問題がある。本論文では、検索語候補が含まれる音声ドキュメントの話題を調べ、検索語が出現しにくい話題である時は、その候補は当該検索語である可能性が低いと仮定して誤検出に対応する。具体的には、検索語が与えられたとき、その説明文と音声ドキュメントとの意味的な類似性を音声内容検索 (Spoken Content Retrieval: SCR) に基づいて求めて、検出候補をリスコアリングする誤検出抑制手法を提案する。種々の STD 検索タスクで評価したところ、全てのタスクで検索精度の向上が得られ、提案手法の有効性および汎用性を示した。また、本手法は、意味的な類似度情報を利用する誤検出抑制手法であり、我々がこれまでに提案している文字列の一致度に基づく検出候補のリスコアリング手法と併用する効果もあることがわかった。

キーワード： 音声検索語検索, 音声内容検索, 意味的マッチング, リスコアリング

Spoken Term Detection using Spoken Content Retrieval by Term Explanation Text

HIROAKI NANJO^{1,a)} TATSUYA KAWAGUCHI²

Abstract: This paper addresses Spoken Term Detection (STD), which finds speeches including a specified query term. Typically, terms are detected based on string matching, which causes false detections for phonetically similar terms. In this paper, we propose a novel STD method which combines string matching and semantic matching. Specifically, we perform Spoken Content Retrieval (SCR) with a term descriptive text and combine string matching-based STD score and SCR score. As for STD from lecture corpus, we showed the effectiveness of the proposed method. We achieved a STD performance improvement for several STD tasks, which showed a validity and robustness of the proposed method.

Keywords: Spoken Term Detection, Spoken Content Retrieval, Semantic Similarity, Rescoring

1. はじめに

デジタル化された大量の音声や動画 (音声ドキュメント)

から、特定の区間を取り出す音声ドキュメント検索技術が求められている。本研究では、検索語そのものが出現する位置を見つける音声検索語検出 (Spoken Term Detection: STD) タスク [1][2][3][4] に焦点を当てる。

音声認識が完全でない状況で、音声中の検索語の位置を漏れなくみつけるには、音声認識誤りを考慮するために単語よりも小さい単位 (サブワード) で、文字列の近さに基づいて照合を行うことが一般的である。その際、検索語で

¹ 京都大学学術情報メディアセンター
Academic Center for Computing and Media Studies, Kyoto University

² 龍谷大学理工学部
Faculty of Science and Technology, Ryukoku University

^{a)} nanjo@media.kyoto-u.ac.jp

はないものを検出する誤検出問題が発生する。

この問題に対して、本論文では、検索語と検索対象の意味的な近さを考慮して誤検出を抑制する手法について述べる。具体的には、これらの誤検出への対応を目的とした、2パスの検索語検出アルゴリズムを提案する。これは、第1パスで得られた検索語候補に対して、第2パスで検索語の説明文と検索語候補を含む音声ドキュメントの意味的類似性を参照してスコアを補正することで、検索語の誤検出を抑制する方法である。第2パスでのスコアの補正は、検索語の説明文と意味的に似ていないドキュメント中の検索語候補は不確かという仮定に基づいて行う。このような2パスまたはそれ以上のパスを用いてSTDを行う研究には[5][6][7][8]などが挙げられる。これらは検索語と異なる情報は用いず、種々の検出モデルや検出システムで検索語の検出を行って、結果を統合する。これに対し、提案手法は、検索語そのものだけでなく、別の情報を使って1パ目の検出結果の高精度化を目指すものである。本手法はこれらの手法と組み合わせることも可能といえ、先行研究の高精度化も期待できる。

誤検出を抑制する2パス検索方法には、関連語や拡張語の利用[3][4]がある。これらは、検索語と関連する語や共起しやすい語が同時に見つかった場合は、当該検出候補は確からしいという仮定に基づくものである。これらは文脈などの意味的な類似性を間接的に捉えていると見ることできる。これに対し、本研究では意味的な類似性を充分に利用する方法を提案する。具体的には、検索語が与えられたとき、その説明文から音声ドキュメントとの意味的な類似性を求め、意味的に類似している/していないドキュメントに含まれる検索語候補は確からしい/確からしくないとするものである。このような研究はこれまでに見られず、提案手法はこの点において新規性を有する。本論文では、種々の検索タスクで提案手法を評価し、その有効性を示す。

本論文の構成は次のとおりである。2章では、一般的な連続DPマッチングに基づく音声検索語検出について述べる。3章では、提案手法である音声検索語検出における意味的類似度の利用法を示す。4章では、提案手法の評価を行い、提案手法の有効性を示す。5章では、種々のSTDタスクに対して提案手法を適用し、提案法が広く適用可能であることを示す。6章で結論を述べる。

2. 音声検索語検出

2.1 概要

音声検索語検出(Spoken Term Detection: STD)とは、音声の中で検索語がそのまま現れる部分を特定する処理である。

STDの方法として様々な研究がなされているが、最も一般的かつ広く利用されているのは、音声を音声認識システムにより音声認識結果、すなわち文字列に変換してお

き、これと検索語の文字列との間で文字列マッチングを行う方法である。検索語は一般的に単語列であるため、音声ドキュメントに認識誤りがあると単語単位のマッチングは行えない。なお、音声認識誤りは本質的に避けられないため、この問題には何らかの対処が必要不可欠である。

2.2 連続DPマッチングによる音声検索語検出

STDでは音声認識誤りに対応するために、単語より小さな単位であるサブワード(音素など)の単位でサブワード系列どうしを誤りを許容して照合(連続DPマッチング距離に近いものを検索結果として選ぶ)することが一般的である。

本研究で用いる連続DPマッチングに基づくSTDのアルゴリズムは次のとおりである。

連続DPマッチングに基づくSTD

- (1) ユーザーから検索語 Q (長さ L_Q)を受け取る。
- (2) 各音声ファイルを検索(STD)する。音声ファイル S の u 番目の発話には、検索語との文字列距離(編集距離) $\text{dist}(Q, S_u)$ を付与しておく。
- (3) 検索語との距離の近さを表すスコア($1 - \frac{\text{dist}(Q, S_u)}{L_Q}$)を付与し、その順に出力する。

このような文字列の近さのみで検出を行うと、軽微な音声認識誤りが起きても検索語を検出できる反面、同音異義語や別の語の一部分、もともとサブワード列としてみたときに近い語などを誤って検出する問題がおきる。

本論文では、このような誤検出を抑制することを目的とする手法を提案する。

3. 検索語説明文を用いた音声検索語検索

3.1 提案法の概要

本論文では、STDにおける誤検出を抑制する手法を提案する。具体的には検索語と検索対象の意味的な近さを考慮した抑制方法を提案する。誤検出を抑制する方法には、関連語や拡張語の利用[3][4]がある。これらは、検索語と関連する語や共起しやすい語が同時に見つかった場合は、当該検出候補は確からしいという仮定に基づくものである。これらは文脈などの意味的な類似性を間接的に捉えていると見ることできるが、意味的な類似性を充分に利用しているとは言えない。これに対し、本論文では、意味的な類似性を充分に利用するSTD手法を提案する。

3.2 提案法のアルゴリズム

検索語説明文を利用したSTDのアルゴリズムは次のとおりである。概要を図1に示す。

検索語説明文を利用した STD のアルゴリズム

- (1) ユーザーから検索語 Q (長さ L_Q) を受け取る.
- (2) 連続 DP マッチングに基づく STD を行う. 音声ファイル S の u 番目の発話には, 検索語との文字列距離 (編集距離) $\text{dist}(Q, S_u)$ を付与しておく.
- (3) 検索語の説明文 text_Q を生成する.
- (4) text_Q との意味的な類似度に基づいて音声ドキュメントを検索する (SCR; Spoken Content Retrieval を行う). ここでの検索は音声ファイル S ごとに行い, 説明文との意味的類似度を表すスコア $\text{sim}(\text{text}_Q, S)$ (0 から 1) を付与しておく.
- (5) 音声ファイル S の各発話 u の編集距離を調整し,

$$\begin{aligned} \text{moddist}(Q, S_u) = \\ \lambda * \text{dist}(Q, S_u) + (1 - \text{sim}(\text{text}_Q, S)) \quad (1) \end{aligned}$$

を得る. λ は二つのスコアの調整重みである.

- (6) 調整後の編集距離 $\text{moddist}(Q, S_u)$ を用いて, 元の検索語 (長さ L_Q) と各音声ファイルとの距離の近さを表すスコア $(1 - \frac{\text{moddist}(Q, S_u)}{Q_L})$ を計算し, その順に結果を出力する

ここで, SCR システムとしてベクトル空間モデルに基づくドキュメント検索システム [9] を採用する. 類似度は, 説明文とドキュメントのベクトル間距離には SMART[10] を用いる.

手順 5 での調整は, 文字列同士が似ており (編集距離が小さく), かつ, 意味的類似度が高い (1 に近い) ほど確からしいという仮定に基づいて行われている. なお, $\text{sim}(\text{text}_Q, S)$ は 0 以上の任意の値を設定できる. 本研究では, 全ての音声ファイル S の類似度を text_Q と最も意味的に類似している音声ファイル $\hat{S}$ との類似度で割って, すべての $\text{sim}(\text{text}_Q, S)$ が 0 から 1 の範囲に収まるように正規化を行っている. 編集距離 $\text{dist}(Q, S_u)$ は正の整数値をとるので, 正規化された $\text{sim}(\text{text}_Q, S)$ との単純なスコア結合では編集距離の差が 1 より大きい場合は調整できない. このため $0 \leq \lambda \leq 1$ の調整重みを導入している.

なお, SCR システムとしては, 文献 [9] に書かれてある方法で構築した講演単位検索システムを用いる. SCR と STD には同じ音声認識結果を用いる.

3.3 検索語説明文の自動生成

次に, 提案法で必要となる検索語の説明文の生成方法について述べる. 生成方法には様々なものが考えられるが, 本研究では, 自動生成手法を採用する. 具体的には, はじめに与えられた検索語について, それが Wikipedia の見出し語になっているかを調べる. 見出し語になっていた場合

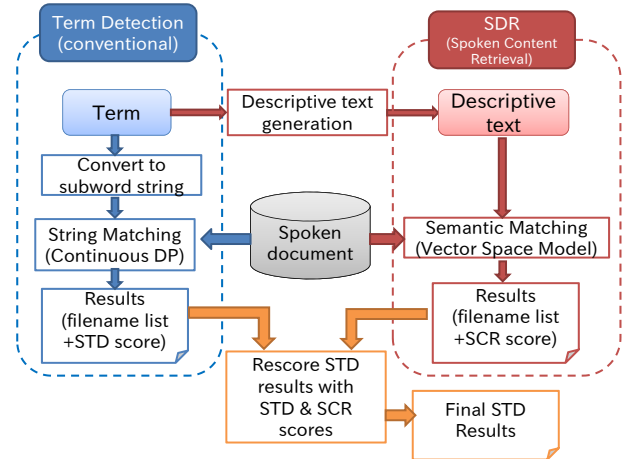


図 1 説明文との意味的類似度を利用した STD の概要
Fig. 1 Overview of STD with its descriptive text

は, Wikipedia の冒頭部分 (見出し部分) と「概要」項目の文章を抽出し, 参考文献を表す数字を削除した上でもとの検索語を加えて説明文とする. 見出し語になっていない場合は, 検索語をそのまま説明文とする.

なお, ユーザに調べたい検索語と見出し説明文を入力してもらうような STD システム (例えば, 検索クエリとして「鯛: 魚の鯛」「愛知県: 中部地方の県で名古屋があるところ」など) を採用した場合は, 検索語説明文の自動生成は不要であり, 本手法をそのまま適用できる.

4. 評価実験

4.1 評価尺度

本研究では, STD の性能の評価尺度として MAP (Mean Average Precision) を採用する. これは, どれほど正解を見つけたかを表す再現率 (recall) と, 検索結果中の正解の割合を示す精度 (precision) の両者を考慮した尺度である. 具体的には各クエリ (検索語) に対して検索結果出力数を変化させて様々な再現率レベルのときの精度を求め, それらの平均をとった値 (平均精度) のクエリ間平均として求める.

あるクエリ q に対する平均精度 AP_q は, 式 (2) で与えられる.

$$AP_q = \frac{1}{\#cor(q)} \sum_{t=1}^{N_q} \text{IsTrue}(q, t) \cdot P(q, t) \quad (2)$$

ここで, $\#cor(q)$ はクエリ q に対する正解文書数, N_q は検索システムがクエリ q の答え (検索結果) として出力した文書数, $\text{IsTrue}(q, t)$ はクエリ q での検索結果の t 番目が正解であれば 1, そうでなければ 0 を返す関数であり, $P(q, t)$ は q の検索結果の t 番目までを評価したときの精度 (precision) である.

この AP_q を全クエリで平均したもの (式 (3)) が, MAP (Mean Average Precision) である. MAP は 0 から 1 をとり, 1 に近いほど平均的に精度が高いことを表す.

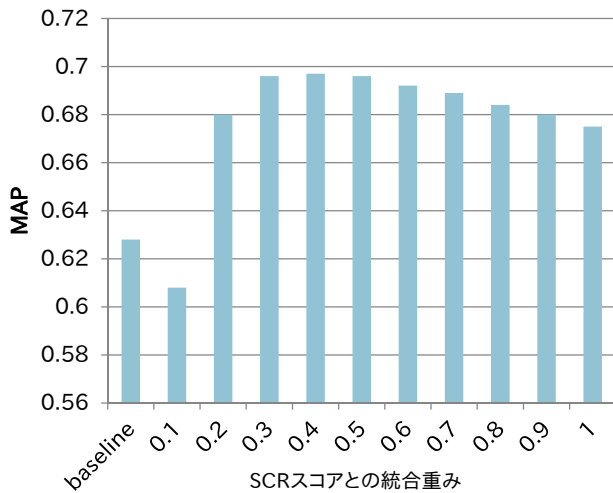


図 2 NTCIR-9 dry run タスクでの評価

Fig. 2 Evaluation on NTCIR-9 dry run task

$$\text{MAP} = \frac{1}{N} \sum_q^N AP_q \quad (3)$$

4.2 実験結果

実験データには、NTCIR-9 SpokenDoc[11] のテストコレクションを用いた。これは日本語話し言葉コーパス (CSJ: Corpus of Spontaneous Japanese) [12] の講演音声を対象とした音声ドキュメント検索のためのテストコレクションである。NTCIR-9 SpokenDoc では、STD タスクとして、検索対象を全講演 (2702 講演) とする ALL タスクと一部 (177 講演) を対象とする CORE タスクが設定されており、本研究では ALL タスクを用いた。クエリには dry run 用クエリ (100 件) と formal run 用クエリ (50 件) があり、ここでは dry run 用クエリ (100 件) を用いた。検索対象の講演音声の認識結果には、タスクオーガナイザから配布されているマッチドモデルによる単語音声認識結果 (Word Corr.=74.1%, Word Acc.=69.2%, Syll. Corr.=83.0%, Syll. Acc.=78.1%) [11] を用いた。

はじめに連続 DP マッチングのみに基づく検出 (baseline) の結果について述べる。検索精度 (MAP) は 0.628 であった (図 2 (baseline))。次に、提案法の結果について述べる。説明文は、2016 年 9 月 14 日に Wikipedia から取得した。この結果も図 2 に示されている。図中の統合重みは、検索語説明文を利用した STD アルゴリズムにおける式 (1) の統合重み λ であり、本実験では、0.1 から 1.0 まで 0.1 刻みで変えて実験を行った。 λ の値が 0.1 の時を除き、baseline よりも高い精度が得られた。0.3~0.6 程度で高い精度が得られており、0.4 の時に最も高い MAP 値 0.697 が得られた。

$\lambda = 0.4$ の時の結果を baseline の結果と検索語ごとに比較した。結果を図 3 に示す。100 件の検索語のうち、80 件の検索語 (図中、青四角) で精度向上が 16 件の検索語 (図

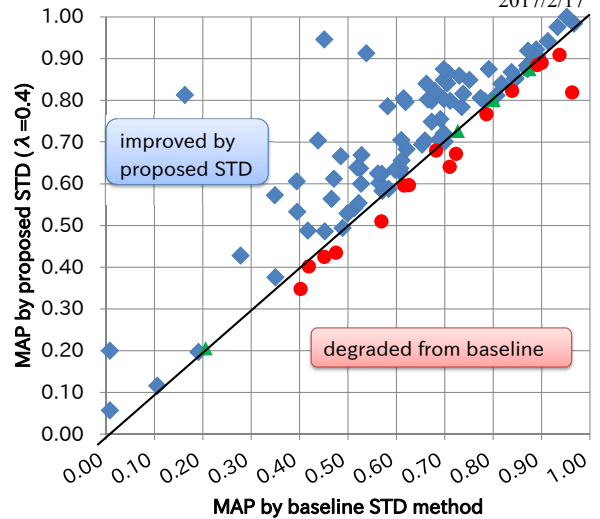


図 3 NTCIR-9 dry run タスクでの評価。ベースラインと提案法でのクエリごとの比較結果

Fig. 3 Evaluation on NTCIR-9 dry run task. Comparing baseline and SCR integrated STD ($\lambda = 0.4$) for each query.

中、赤丸) で精度低下がみられた。4 件の検索語 (図中、緑三角) では変化がなかった。また、大きく精度向上している検索語が見られる一方、大きく精度低下した検索語はないことがわかる。これらの結果は、提案手法の有効性を示している。

5. 他タスクでの評価実験

4 章の結果より、STD において意味的類似度を用いることに効果があることがわかった。本章では、最も精度が高かった統合重みパラメータ ($\lambda = 0.4$) を用いて、種々のタスクで提案法の効果を調べる。

5.1 NTCIR-9 SpokenDoc formal run

NTCIR-9 SpokenDoc formal run で評価を行った。これは、NTCIR-9 SpokenDoc dry run と検索対象が同じ (CSJ の全ての講演のマッチドモデルによる単語音声認識結果: Word Corr.=74.1%, Word Acc.=69.2%, Syll. Corr.=83.0%, Syll. Acc.=78.1%) であり、検索語セットが異なるタスクである。

結果を図 4 に示す。説明文は、2016 年 9 月 14 日に Wikipedia から取得した。結果の傾向は dry run の結果と同じであり、 λ が 0.1 以外の場合で精度改善が得られること、0.3~0.6 程度で高い精度が得られることがわかった。0.4 の時の MAP 値は 0.561 であった。検索語セットを変えても提案手法に効果があることが確認できた。

5.2 NTCIR-10 SpokenDoc2 formal run large-size task

NTCIR-10 SpokenDoc2[13] formal run で評価を行った。CSJ の 2702 講演を検索する Large-size タスクで評価を行った。Large-size タスクは 4 章の実験と検索対象が同

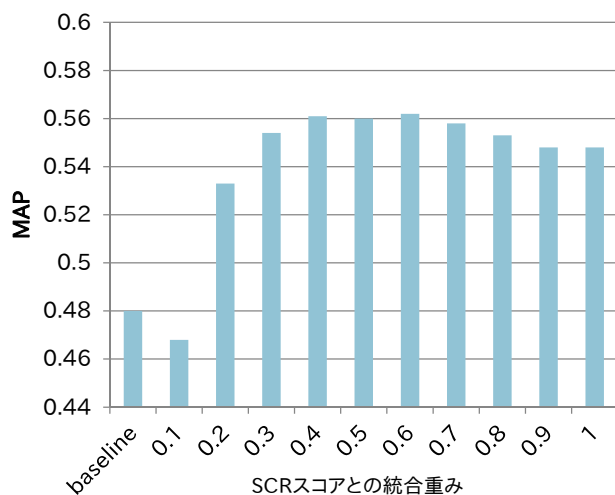


図 4 NTCIR-9 formal run タスクでの評価

Fig. 4 Evaluation on NTCIR-9 formal run task

表 1 NTCIR-10 SpokenDoc2 formal run Large-size task での評価

Table 1 Evaluation on NTCIR-10 SpokenDoc2 formal run Large-size task

全検索語		既知語の検索語		未知語の検索語	
従来法	提案法	従来法	提案法	従来法	提案法
0.471	0.558	0.562	0.625	0.394	0.500

じ (CSJ の全ての講演のマッチドモデルによる単語音声認識結果: Word Corr.=74.1%, Word Acc.=69.2%, Syll. Corr.=83.0%, Syll. Acc.=78.1%) であり, 検索語セットが異なるタスクである。

結果を表 3 に示す。説明文は, 2016 年 9 月 15 日に Wikipedia から取得した。検索語には, 単語音声認識時に単語辞書に含まれる既知語の検索語と, 含まれない未知語の検索語が存在するため, それぞれでの評価も行っている。既知語, 未知語に関わらず提案法に効果があることがわかる。特に未知語に対して, 精度の大きな向上 (0.394 から 0.500) がみられた。未知語は音声認識において別の語に認識されるため, もとの検索語と文字列としての一致度は低くなる傾向がある。これらを検出しようとする別の近い文字列の語の誤検出が多くなる傾向がある。一方, 検索語が含まれる音声ドキュメント中では検索語以外の語, 特に既知語はある程度正しく音声認識されていると考えられ, 検索語候補が含まれる音声ドキュメントと検出語の説明文との意味的類似度が誤検出抑制に効果を発揮したと考えられる。

次に未知語について詳細に調査した。具体的には, Wikipedia の見出しに見つかった場合 (42 件) と見つからなかった場合 (12 件) について分析した。本実験では, 検索語の説明文を, 検索語そのものと Wikipedia 文書で作成した。Wikipedia の見出しになっていない検索語については, 検索語そのものを説明文として意味的マッチング

表 2 Wikipedia の見出し語の検索語とそうでない検索語の比較 (NTCIR-10 formal run Large-size タスク 未知語検索語)

Table 2 Comparison of Wikipedia-entry words and not-entry words (NTCIR-10 formal run Large-size task OOV terms)

	baseline	提案法
見出し語 (42 件)	0.369	0.495
見出し語でない (12 件)	0.480	0.520

(SCR) を行っている。未知語は音声認識されていないので, 検索語そのものの SCR は効果が得られにくい, 悪影響もないと予想される。結果を表 2 に示す。

見出し語になっていない場合は, 効果が小さいことがわかる。このことから, 説明文による SCR の効果がわかる。見出し語になっていない場合の結果 (12 件) を詳細に分析したところ, 精度向上が見られた検索語が 5 件, 変化がなかったものが 5 件, 精度低下がみられたものが 2 件あった。

精度の変化があった未知語検索語のうち, 精度向上したものは「ウェアラブルコンピュータ (MAP:0.380→0.448)」「周波数ワーピング (MAP:0.793→0.794)」「チャイルドトランスミッション (MAP:0.399→0.472)」「花束贈呈 (MAP:0.364→0.487)」「ポートフォリオ評価 (MAP:0.223→0.492)」, 精度低下したものは「楕型フィルタ (MAP:0.379→0.337)」「有毛細胞 (MAP:0.710→0.695)」であった。これらは複合語からなる検索語であった。一部分が音声認識辞書に存在して認識が行われており, SCR が行われた結果と考えられる。また精度低下の悪影響は小さいこともわかった。

このことから, 提案手法は有効といえる。

5.3 検索語文字列拡張法との併用

これまでに我々は STD において, 検索語の前後に格助詞を付加した拡張検索語 (拡張クエリ) による検索結果を統合して誤検出を抑制する STD をこれまでに提案している [3]。この手法は効果はあるものの, 本質的には文字列マッチング情報のみを用いており, 意味的マッチング情報を用いる提案法との併用効果は期待できる。本論文では, この手法と提案法の併用を行い, 効果があることを示す。

5.3.1 STD における検索語文字列拡張法の概要

[3] で提案した拡張検索語による STD における誤検出抑制法について述べる。これは, 検索語の前後につきやすい語を付加して検索語を長くして拡張語とし, もとの検索語と拡張語が同時に出現しなかった場合は誤検出らしきが高いと判定する方法である。検索語は基本的に名詞, 特に固有名詞となることが多いことに着目し, 検索語の前または後に格助詞 (10 種類) を付加して関連語としている。この手法には, どのような検索語が与えられても拡張語を得る

表 3 検索語文字列拡張法と提案法の併用の効果

Table 3 Combination of query string expansion method and proposed method

	baseline	拡張	提案法	拡張+提案
NTCIR-9 dry run	0.628	0.702	0.697	0.716
NTCIR-9 formal run	0.480	0.587	0.561	0.611
NTCIR-10 formal run Large-size	0.471	0.569	0.558	0.606

ことができる利点がある。

5.3.2 評価実験

NTCIR-9 dry run および formal run, NTCIR-10 formal run Large-size の 3 タスクで検索語文字列拡張手法 [3] と提案法の併用を行った。具体的には, [3] の検索語文字列拡張手法^{*1}で得られた修正編集距離を式 (1) の $\text{dist}(Q, S_u)$ として用いた。

結果を表 3 に示す。全てのタスクで併用に効果が見られ, 検索語文字列拡張手法と提案法それぞれを単独で用いる場合よりも高い検索精度を得ることができた。

これらのことは, 提案手法の有効性を示している。

6. まとめ

STD において, 検索語の説明文を自動生成したうえで, その説明文と音声ドキュメントの類似度を STD に用いる方法を提案し, 効果を示した。

謝辞 本研究は科研費 (15K00254) の助成を受けた。

参考文献

- [1] 伊藤慶明, 西崎博光, 中川聖一, 秋葉友良, 河原達也, 胡新輝, 南條浩輝, 松井知子, 山下洋一, 相川清明: 音声の中の検索語検出のためのテストコレクションの構築と分析, 情報処理学会論文誌, Vol. 54, No. 2, pp. 471–483 (2013).
- [2] Noritake, K., Nanjo, H. and Yoshimi, T.: Image Processing Filters for Line Detection-based Spoken Term Detection, *Proc. INTERSPEECH*, pp. 2125–2128 (2011).
- [3] 南條浩輝, 前田 翔, 吉見毅彦: 音声検索語検出のためのクエリ拡張の検討, 情報処理学会研究報告, 2014-SLP-101-16 (2014).
- [4] 小田原一成, 山下洋一: 音声の中の検索語検出における単語共起情報の利用, 情報処理学会研究報告, 2016-SLP-110, pp. 1–6 (2016).
- [5] 三浦成一, 桂田浩一, 入部百合絵, 新田恒雄: Suffix Array を用いた高速 STD のための検索閾値の調整手法, 第 8 回音声ドキュメント処理ワークショップ, No. 6 (2014).
- [6] Takahashi, J., Hashimoto, T., Kon'no, R., Sugawara, S., Ouchi, K., Oshima, S., Akyu, T. and Itoh, Y.: An IWAPU STD System for OOV Query Terms and Spoken Queries, *NTCIR-11 Workshop Meeting*, pp. 384–389 (2014).
- [7] 坂本伊織, 工藤祐介, 山下洋一: ベクトル量子化に基づいた音声の中の検索語検出における検索結果の統合, 第 8 回音声ドキュメント処理ワークショップ, No. 8 (2014).
- [8] 牧野光晃, 山本直樹, 甲斐充彦: 分布間距離ベクトル表現による音響的類似度を利用したテキスト及び音声クエ

リからの音声検索語検出の改善, 第 8 回音声ドキュメント処理ワークショップ, No. 10 (2014).

- [9] 西尾友宏, 南條浩輝, 吉見毅彦: 講演音声ドキュメント検索のための擬似適合性フィードバック, 情報処理学会論文誌, Vol. 55, No. 5, pp. 1573–1584 (2014).
- [10] 小作浩美, 内山将夫, 井佐原均, 河野恭之, 木戸出正継: WWW 検索における複数検索結果の統合処理とその評価, 情報処理学会論文誌, Vol. 44, No. SIG 8(TOD 18), pp. 78–91 (2003).
- [11] Akiba, T., Nishizaki, H., Aikawa, K., Kawahara, T. and Matsui, T.: Overview of the IR for Spoken Documents Task, *NTCIR-9 Workshop Meeting*, pp. 223–235 (2011).
- [12] Maekawa, K.: Corpus of Spontaneous Japanese: Its design and evaluation, *Proc. ISCA & IEEE-SSPR*, pp. 7–12 (2003).
- [13] Akiba, T., Nishizaki, H., Aikawa, K., Hu, X., Itoh, Y., Kawahara, T., Nakagawa, S., Nanjo, H. and Yamashita, Y.: Overview of the NTCIR-10 SpokenDoc-2 Task, *NTCIR-10 Workshop Meeting*, pp. 573–587 (2013).

^{*1} 前に 10 種類の格助詞, 後ろに 10 種類の格助詞を付与した 20 拡張語使用, $\text{penalty}=2.5$: [3] 内での最も精度の高かった手法