



サウンドスペクトログラムにおける連鎖パターンマイニング—個人識別への応用—

メタデータ	言語: jpn 出版者: 計測自動制御学会 公開日: 2016-02-17 キーワード (Ja): 話者認識, 連鎖パターンマイニング, 頻出パターン, マイニング, スペクトログラム キーワード (En): 作成者: 三澤, 慶太, 後藤, 廉太, 岡田, 吉史 メールアドレス: 所属:
URL	http://hdl.handle.net/10258/3859

サウンドスペクトログラムにおける連鎖パターンマイニング—個人識別への応用—

著者	三澤 慶太, 後藤 廉太, 岡田 吉史
雑誌名	計測自動制御学会システム・情報部門学術講演会講演論文集
巻	2015
発行年	2015-11-18
URL	http://hdl.handle.net/10258/3859

サウンドスペクトログラムにおける連鎖パターンマイニング -個人識別への応用-

○三澤慶太 後藤廉太 岡田吉史 (室蘭工業大学)

概要 我々は以前、複数の系列データに跨って繰り返される頻出パターンである連鎖パターンを抽出する方法を提案した。現在我々は、この手法を母音のサウンド スペクトログラムに適用し個人識別を試みている。本発表では、個人間の判別に寄与する周波数点を特定し、それを用いた場合の個人識別結果について 報告する。

キーワード: 話者認識, 連鎖パターンマイニング, 頻出パターン, マイニング, スペクトログラム

1. はじめに

話者認識技術は、ニュースや対談における内容の書き起こし、対話を行うロボット、個人認証などへの応用が期待され、盛んに研究されている¹⁾。

本研究では、母音のサウンドスペクトログラム(以下、スペクトログラム)に含まれる時系列情報を用いた話者認識手法を提案する。既存のスペクトログラムを対象とした話者認識手法として、スペクトログラムの画像解析を用いた方法がある²⁾。これは、発話者の特徴量としてスペクトログラムの画素値を用いている。一方、本研究で提案する手法は、連鎖パターンマイニング^{3,4)}と呼ばれる方法を用いて、個々人のスペクトログラムの時系列パターンを特徴量として話者認識を行う。連鎖パターンマイニングにより、異なる周波数点間に跨って繰り返されるスペクトルパワーの時間的推移パターン(連鎖パターン)を発見できる。

本稿では、母音「あ」のスペクトログラムから連鎖パターンを抽出し、これを特徴量とした話者認識の結果を示す。

2. 方法

Fig1 にスペクトログラムからの連鎖パターン抽出とユーザプロフィールデータベース (UPD) の構築方法を、Fig2 に話者認識の概要を示す。

2.1 ユーザプロフィールデータベース (UPD) の構築

2.1.1 スペクトログラムの作成

スペクトログラムは、縦軸と横軸にそれぞれ周波数と時間を取り、色の濃さでスペクトルパワーを表現した3次元グラフである。スペクトログラムは、音声デー

タに対して窓処理を行い、切り出された分析区間に周波数分析を繰り返すことにより生成される。

2.2.2 周波数選択

スペクトログラムから個人の識別に寄与すると思われる周波数点のみを選択する。事前に各被験者から音声サンプルを収集し、それらのスペクトログラムを生成しておく。続いて、周波数点ごとにすべての被験者ペアでクラスカル・ウォリス検定を行い(有意水準 $\alpha = 0.05$)、有意差があったペアの個数をカウントする。なお、ここでのp値としてボンフェローニ法による補正後の確率が用いられた。最終的に、カウント数が全ペア数の60%以下の周波数点をスペクトログラムから削除し、それ以外を被験者間に差別化に有効な周波数点として使用することとする。

2.2.3 スペクトログラムからの連鎖パターン抽出

スペクトログラムからの連鎖パターンマイニングは次の手順で行われる。まず、周波数点ごとに、繰り返し現れるスペクトルパワーの時系列パターン(頻出パターン)を抽出する。次に、異なる周波数点に跨って同時刻帯に繰り返し現れる頻出パターンの集合を連鎖パターンとして抽出する。Fig1では連鎖パターン $\{(f1, A), (f2, B), (f4, C)\}$ などが抽出されている。ここで、例えば $(f1, A)$ は、周波数点 $f1$ で頻出パターン A を持つことを意味している。

2.2.4 連鎖パターンへのスコア付け

本研究では、種類数と出現頻度の高い頻出パターンを多く含む連鎖パターンほど、被験者を特徴付ける情報であるとみなす。そこで、スペクトログラムから抽出された連鎖パターン L のスコアを、(1)式に基づいて算出する。

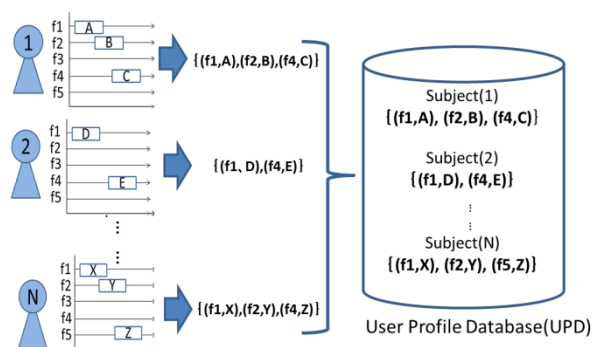


Fig 1: Linkage Pattern Mining and Creation of UPD

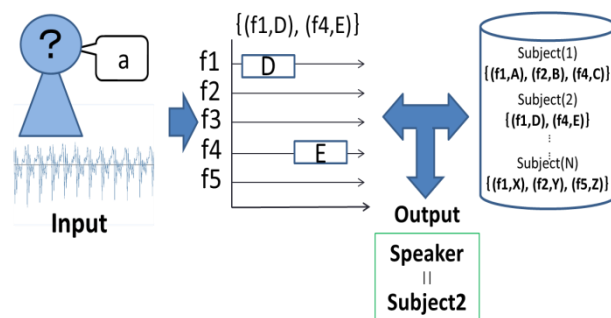


Fig 2: Speaker recognition

$$\text{Score}(L) = (\text{no} * \text{np} * r) / C \quad (1)$$

ここで, noは連鎖パターンLの出現頻度, npは連鎖パターンLに含まれる頻出パタンLの数, rは連鎖パターンLと同じサイズの連鎖パターンをもつサンプルの割合を表す. また, Cは定数であり, 本研究ではC=5とした.

2.2.5 UPDへの連鎖パタンの格納

2.2.3節から2.2.4節の操作をすべての被験者の音声サンプルに対して行う. なお, 本研究では, 2.2.4節のスコアが上位10%の連鎖パターンのみをユーザプロファイルデータベース (UPD: User Profile Database) に格納した.

2.2 話者認識

まず, 認識対象となる発話者の入力音声に, 2.2.3節から2.2.4節の操作を施す. 次に, 抽出された連鎖パターンとUPDに格納された連鎖パターンを照合し, 最も類似した連鎖パターンを保持している被験者を発話者として出力する.

3. 実験

3.1 音声データの収集とスペクトログラム作成

男子大学生5名の母音「あ」の音声データを収集した. 母音の録音に使用した機材はOLYMPUS製リニアOCMレコーダー(型番:LS-11)である. 各被験者には, 母音「あ」を1秒以上発音してもらい, 録音した音声から1000ミリ秒を切り出し, これを4分割した250ミリ秒の音声を1サンプルとした. 上記の作業は, 各被験者に対して3回ずつ行われた. その結果, 被験者ごとに12サンプル, 計60サンプルの音声データを収集した. 続いて, 各サンプルからスペクトログラムを作成した. なお, スペクトログラム算出のパラメータは, フレーム幅5ms, シフト幅0.5msと設定した.

3.2 話者認識

本手法の話者認識精度を評価する. 実験手順は以下のとおりである. まず, 被験者5名の12サンプル(計60サンプル)から連鎖パターンを抽出し, スコア付けを行い, 上位10%の連鎖パターンをUPDに格納する. 次に, UPDから1サンプルを取り出して入力音声とし, 2.5節の話者認識により出力結果の正誤を記録し集計する. 上記の実験をUPD内の全てのサンプルに関して実行する.

4 結果・考察

Table1は本実験による話者認識の結果を示している. 行が入力に用いた被験者番号, 列が出力結果の被験者番号である. 行列内の各数値は結果として出力された回数を示す. Table1の右端の列は入力被験者ごとの正解率を示している. 全体での識別率は約36%となっており, 被験者の中では被験者1が一番高く約67%, 被験者3, 5が一番低く約17%となった. 認識結果の傾向とし

Table1:Recognition Result

		Output					Recognition rate
		subject(1)	subject(2)	subject(3)	subject(4)	subject(5)	
Input	subject(1)	8	2	1	1	0	0.67
	subject(2)	5	3	1	3	0	0.25
	subject(3)	3	2	2	4	1	0.17
	subject(4)	2	1	1	7	1	0.58
	subject(5)	2	2	0	6	2	0.17

て, 正解率の低い話者は正解率の高い話者に誤認されている.

このような, 被験者間の識別率の差は, 抽出された連鎖パターン数や種類数の違いによるものと考えられる. 高い識別率が得られなかった理由として, 被験者間で共通の連鎖パターンが多数含まれることにより, 互いに誤認してしまう場合, 同じサンプル内で似通った連鎖パターンが残っており, その連鎖パターンがまとめて誤認されていることなどが考えられる.

5 まとめ

本稿では, 母音「あ」のスペクトログラムから得られる連鎖パターンを用いた話者認識手法を示した. 実験として, 被験者5名の音声を用いた話者認識を行った. 結果, 全体的な識別率は約36%と低く, 被験者間での識別率にもばらつきがあった.

今後は, 被験者間を差別化する連鎖パターン抽出の方法を考案する. また, 被験者間で抽出される連鎖パターン数にバラツキが出ないよう被験者ごとにパラメータ設定を変更して実験を行う. また, 類似した連鎖パターンを1つの連鎖パターンとしてまとめる処理を追加することにより, 連鎖パタンの質を上げて話者認識を試みる.

参考文献

- 1) 御船 正樹, 鈴木 基之, 人 福継ぎ, 北 研二: “クラスタリングに基づく GMM 学習法による話者モデルの構築”, 電子情報通信学会技術研究報告. SP, 音声, 2011
- 2) 和田 はるか, 西 仁司: “音声スペクトログラムの画像解析による話者識別の研究”, 2012
- 3) 三浦貴大, 岡田吉史: “複数の系列データにまたがる頻出パターン抽出方法の提案”, 第14回日本感性工学会, 日本感性工学会, 2012
- 4) Saerom Lee, Takahiro Miura and Yoshifumi Okada: “A New Method for Improving the Performance of Linkage Pattern Mining”, Proc. of IMECS 2014, International Association of Engineers, 36-40, 2014.