

母音の合一と混同の理論

—津軽、出雲方言を例として—

前川 喜久雄（上智大学博士課程）

ディスクリプタ： 母音の混同 音韻変化 方言
フォルマント周波数 確率モデル
一様分布 中心極限定理 仮説検定

0. はじめに

日本語の諸方言のなかには標準語的な5母音体系とは異なる母音体系を有する方言も多い。そのなかで特に標準語的体系において認められる母音間の対立の一部が失われるか又は非確定的な方言は、従来、母音の合一ないし混同を示す方言と呼ばれてきた。合一と混同という術語には共に管見の及ぶ限りでは明確な定義が与えられていないようである。しかし、両者は概略次のような含意を伴って使い分けられているといえよう。母音の合一とは複数の（ほとんどの場合、ふたつの）母音が通時的変化の結果、対立を解消して唯ひとつの母音に統合された状態をさす。従って合一の変化以前に問題とする2母音が構成していたminimal pairsは変化後ではすべて同音異義語となる。一方、母音の混同とは共時的にみてひとつの言語共同体内に特定の母音対の間に対立を有する成員とそうでない成員とが混在する状態をさす。更に同一の話者が音韻体系そのものを全面的に切り換える（switchする）ことなしに社会言語的な場面に応じて2母音を対立させたり、させなかったりという状況が存在するならば、これも混同に含めることにしたい。¹ このように考えれば合一と混同とは音韻変化現象の相異なるふたつの段階として理解される。混同は共

MAEKAWA Kikuo (The Graduate School of Languages and Linguistics, Sophia University) — A Theory for Vowel Confusion: Based on the data from two Japanese dialects, Tsugaru and Izumo.

注1 本稿は日本方言研究会第35回研究発表会で行った発表（文献1）を改訂したものである。発表の折に鳥取大学今石元久、東京外語大学井上史雄の両先生からいただいた出雲方言の母音体系の解釈についてのコメントは本稿の作製に大変貴重であった。記して感謝をあらわしたい。

時的にみて言語共同体がひとつの安定した母音体系から他のもうひとつの体系へ移行する過渡的狀態であり、それが収束した姿が合一である。無論、合一の狀態から新たな母音の分離への過程としても混同は存在するだろう。

このような見解は常識的にみて妥当なものといえようが、実際に諸方言の音声資料にあたってどの方言のどの場面が合一の狀態にあり、どれが混同の過程にあるかを客觀的に判別するのは容易なことではなかった。混同から合一へ（あるいはその逆）の移行は時間領域で連続と考えられる。又、その移行の姿を把握するために我々が觀察することのできる音声は物理的連続体である。就中母音は特に調音上の連続性が著しい。従って同一個人内ないし個人間での母音の音質の差の散らばりが問題となる場合、本質的に離散的な音声記号による転写にはおのずから限界が生じたからである。しかし、今日、音響音声学的な分析により母音の音質を連続性のあるパラメータ（フォルマント周波数）で数量的に表現することは比較的容易となっている。そうすると残る問題は、このパラメータを利用して合一と混同とを識別するための数理モデルを構成することである。本稿では2母音の合一狀態を表現する確率モデルを提示する。このモデルに基づく統計的仮説検定を施すことにより、共時的な観点において合一をその他の狀態（混同の狀態、及び2母音間の対立が社会的に確定され安定している狀態）から識別することが可能となるのである。

1. 基礎データと問題の呈示

図1, 2, 3に今回の研究の出発点となった基礎データを示す。図1は東京地方（発話者数 $N=18$ ）、図2は青森県津軽地方（ $N=14$ ）、図3は島根県出雲地方（ $N=26$ ）の成人男性話者（東京地方以外はいずれも各地はえぬきの60歳以上の人物。東京は都内出身の大学学部生）が発音した方言母音をサウンド・スペクトログラフ（RION社SG-07）により分析し、第1、第2フォルマント周波数を測定したものである。² 尚ここで方言母音と呼ぶのは、クイズ方式で録音された調査語「息」、「駅」、「秋」、「沖」、「浮き」の名語頭母音であり、以下ではこれらを夫々イ、エ、ア、オ、ウと略記する。この記法は単なるmnemonic symbolizationであり、各母音の音質や音韻論上の属性については何も含意していない。

図1, 2, 3からただちに見てとれることは、東京地方（図1）では5母音イ、エ、ア、

注2 測定は300Hzフィルタのパタンによる位置決めと、45Hzフィルタのセクションの共鳴曲線あてはめによる。調査地点等の詳細は文献2, 3を参照されたい。

図1. 東京地方のデータ

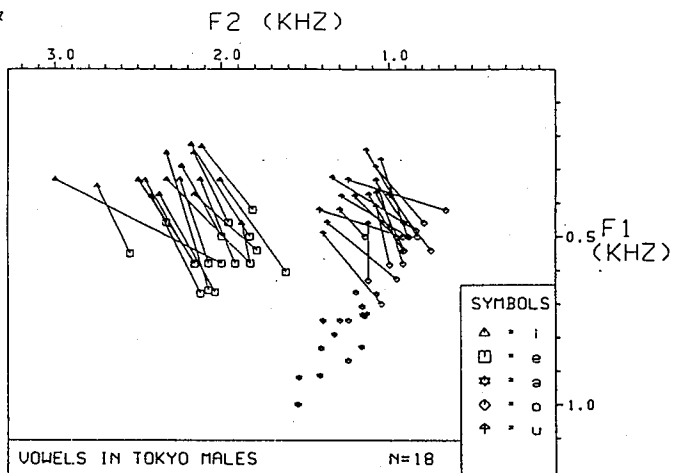


図2. 津軽地方のデータ

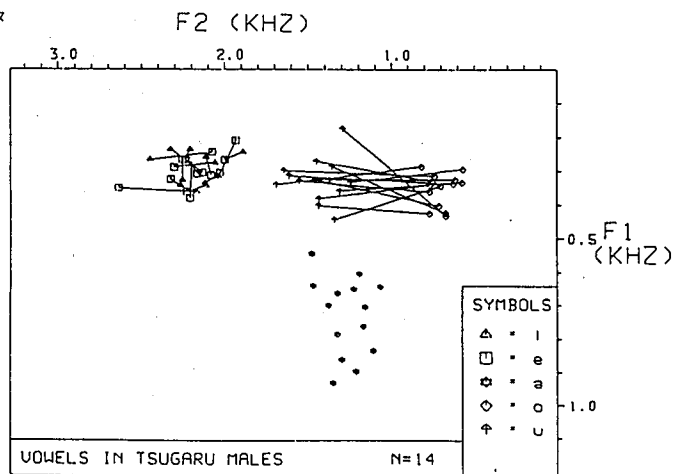


図3. 出雲地方のデータ

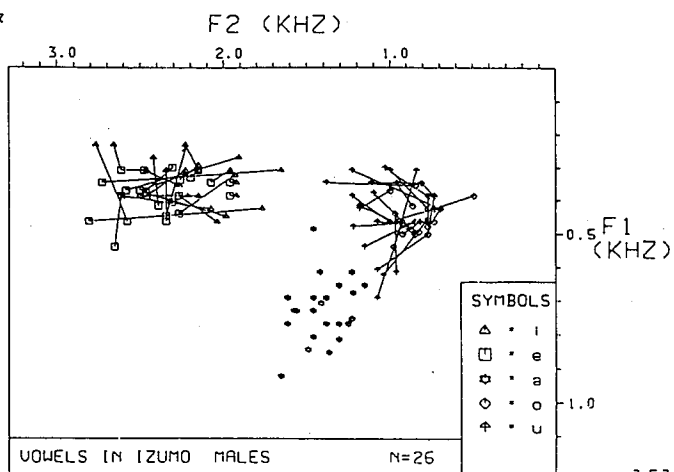


表1. 津軽地方基礎統計量 (Hz)

		イ	エ	ア	オ	ウ
F1	Mean	290	300	730	330	340
	S. D.	40	50	110	100	50
	Range	130	170	390	170	170
F2	Mean	2,180	2,190	1,270	650	1,460
	S. D.	130	170	120	190	130
	Range	440	700	410	250	440

表2. 出雲地方基礎統計量 (Hz)

		イ	エ	ア	オ	ウ
F1	Mean	330	380	710	430	430
	S. D.	63	64	100	58	100
	Range	230	260	500	230	390
F2	Mean	2,170	2,400	1,390	830	1,010
	S. D.	275	220	130	100	194
	Range	1,150	850	370	360	730

オ、ウが相互によく分離した分布を示すのに比べ、津軽地方（図2）ではイとエの分布が大幅に重なりあっていること。又、出雲地方（図3）ではイとエに加え更にオとウの分布にも同じ現象が生じているということである。このような重複現象の存在は、当該する方言において問題の母音対の間に言語学上の対立が存在しないことを予想させるものである。しかし、同時に図2、図3をよく比較すると津軽地方におけるイ、エの関係と出雲地方におけるイ、エおよびオ、ウの関係とでは、かなり相違する点もあることが看取される。津軽地方ではイの分布する範囲とエの分布する範囲がきわめて似通っている。それに対し出雲地方ではいずれの母音対においても、ふたつの母音が多少とも相互に分離する傾向をみせている。そしてこのことは表1、2に示した両地方のデータの記述統計量からも或る程度読みとることができる。ここでただちに問題となるのは津軽、出雲両地方での母音の重複現象には果たして言語学上の質的な差異が認められるか否かという点である。この問題

は通常の統計的データ解析の立場から眺めると、2次元母集団分布の同一性の検定の問題に帰着されやすいだろうが、ここではそのようなアプローチは採用しない。我々が興味をもつのは各母音の母集団分布の性質ではなく、個人レベルでの母音の音質の差が集団としてどの程度一貫して保持されているかなのである。以下ではこのような言語データとしての特殊性を考慮して筆者が考案した方法により解析を進めてゆく。

2. 2母音間の音質の差の数量的表現

従来、母音の音質を音響パラメータにより表現するには、図1, 2, 3のように第1, 第2フォルマント周波数(以下、夫々 f_1 , f_2 と略記)を直交座標に配した音響空間にすることが多い。ここでも任意の母音 V の音質を $f_1 - f_2$ 平面上の2次元ベクトルとして表現しよう。つまり話者 i の発音した2母音 j, k を夫々、

$$V_{ij} = [f_{1ij} \quad f_{2ij}]^T$$

$$V_{ik} = [f_{1ik} \quad f_{2ik}]^T$$

と表わす。ここで f_1, f_2 はフォルマント周波数を表わすスカラー。イタリックの添字 i は話者の別を示し、又 j, k は共に母音の別を示す名義尺度上の添字である($j, k = \text{イ, エ, ア, オ, ウ}$)。ベクトル右肩の T はベクトルの転置を示している。今、ひとりの話者 i が特定音声環境下でふたつの母音 j, k を発音したとき、その2母音間の(tokenとしての)音質の差を下記の差ベクトル D_{ijk} で定義する。

$$(1) \quad D_{ijk} = V_{ij} - V_{ik} \\ = [d_{1ijk} \quad d_{2ijk}]^T$$

いうまでもなくスカラーとして、 $d_{1ijk} = f_{1ij} - f_{1ik}$, $d_{2ijk} = f_{2ij} - f_{2ik}$ である。定義式(1)より明らかなように D_{ijk} は幾何学的には $f_1 - f_2$ 平面上のふたつの(矢線)ベクトル V_{ij} と V_{ik} の先端を結ぶ線分を、原点に V_{ij} の側が重なるように平行移動したものである。ちなみに図1, 2, 3では同一話者の発音したイ, エとオ, ウの対を実線で結んである。これらの線分を夫々イとオが原点に重なるように平行移動したものが、 D_{iei} , D_{uou} ($i=1, 2, \dots, N$)と見做して大略さしつかえない。³ この D_{ijk} が表現しているのは、

注3 図1, 2, 3では作図上の要請から F_1 軸と F_2 軸のスケーリングファクタが異っている。後述する図4, 5, 6では両軸ファクタが等しいので、前3図を単に平行移動しただけでは後者に厳密に一致しない。ただし、ここで必要な視覚的理解にさしきわることはないと思う。

(a) $f_1 - f_2$ 平面上で V_{ijk} と V_{ilk} が相対的にどのような位置関係にあるのか、(b) V_{ijk} と V_{ilk} は平面上でどのくらい離れているか、の二種の情報であると考えられる。つまり(a)方向と(b)距離の情報であるが、以下の解析ではこのうち(a)の情報にのみ着目してゆくことにする。その理由を簡単に述べよう。図1の東京地方のデータから明らかなように、2母音間の対立(音質の差)が社会的に確定している場合には、上述の差ベクトルに等価な線分は $f_1 - f_2$ 平面上でN本がおおよそ平行的な関係にある。問題を2母音間の関係に限定する限りでは、この平行関係こそが母音間の対立の確定性の表現となっているのであって、線分の長さ自体は声道長の個人差に代表される調音器官の生理的個人差に大きく影響され易く、言語学的にはさほど重要でない変動を示し易いと考えられる。⁴ 又、方向と量との二元的データが、方向のみの一元データに縮約されてしまう点もデータ解析上みのがせない利点である。

図1, 2, 3のデータから D_{ijk} を計算し、ついで各差ベクトルの長さを1に正規化した結果を夫々図4, 5, 6に示す。この正規化は下式により施される。

$$D_{ijk}^* = D_{ijk} / |D_{ijk}|$$

ここで D_{ijk}^* は長さ1に正規化された差ベクトル。また記号 $||$ はベクトルのノルム(長さ)を表わすもので、次のように計算される。

$$|D_{ijk}| = \{(d_{1ijk})^2 + (d_{2ijk})^2\}^{\frac{1}{2}}$$

さて、ここで各図に示されたN本のベクトルが単位ベクトル(unit vector)

$$e_1 = [1 \quad 0]^T$$

と平面上でなす角度を、ラジアンを単位として、 θ_{ijk} と定義しよう(尚以下では誤解のおそれがない限り、 ijk 等の添字を省略する)。この θ を求めるための単純な代数式は存在しないので、以下に述べる手順を踏むことが必要である。まず、 e_1 と任意のひとつの D (D^* を用いても等価である) の内積を $\langle e_1, D \rangle$ で表わせば、

$$\langle e_1, D \rangle = \cos \theta \cdot |e_1| \cdot |D|.$$

第二余弦定理を用いて

$$\langle e_1, D \rangle = d_1 \cdot 1 + d_2 \cdot 0 = d_1$$

を得るから、これより

注4 本稿のデータは成人男性のみであるので、図1等ではこの点がさほど強調されない。しかし、女性や子供のデータを含めた場合は大きな問題となる。

図4. 正規化された差ベクトルの分布
東京地方（左イ, エの対, 右オ, ウの対）

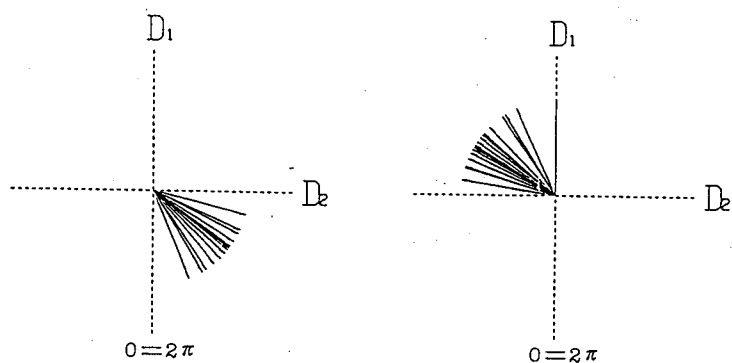


図5. 正規化された差ベクトルの分布
津軽地方（左イ, エの対, 右オ, ウの対）

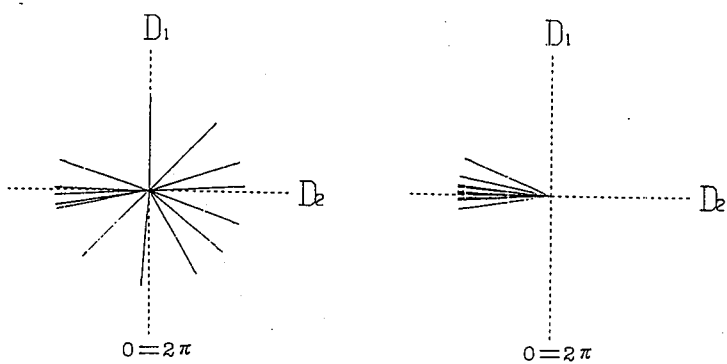
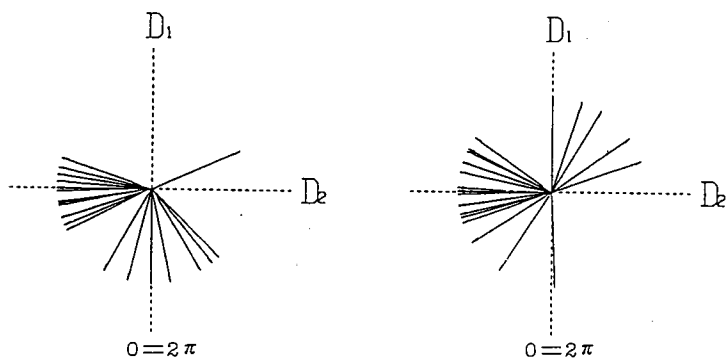


図6. 正規化された差ベクトルの分布
出雲地方（左イ, エの対, 右オ, ウの対）



$$\begin{aligned}\cos \theta &= d_1 / |e_1| \cdot |D| \\ &= d_1 \left\{ (d_1)^2 + (d_2)^2 \right\}^{-\frac{1}{2}}.\end{aligned}$$

次に逆余弦関数を利用して θ の値そのものを求めたいのだが、一価関数としての逆余弦関数の定義域の関係から、本論のように θ が半開区間 $(0, 2\pi]$ に分布する場合、一般には

$$\cos^{-1}(\cos \theta) = \theta$$

なる関係が成立しない。そこで問題とするベクトル D^* が図4, 5, 6において第1から第4までのどの象限に属するかを知らねばならない。この目的のため、与えられた D の値から、それが図4, 5, 6等で属する象限を判定する計算機プログラムのためのアルゴリズムを付録として示すことにした。一旦、この判定が行われたならば、 θ は次のように決定される。

$$\theta = \cos^{-1}(\cos \theta) : \text{第1, 第2象限にある場合}$$

$$\theta = 2\pi - \cos^{-1}(\cos \theta) : \text{第3, 第4象限にある場合.}$$

このようにして決定された θ は任意の D から先述の(a)の情報のみを抽出したデータと見做せるものである。

3. 2母音の合一の確率モデル

ここで再度、差ベクトル D の定義を考えてみよう。(1)式はふたつの母音(j と k)が $f_1 - f_2$ 平面上で同一の点となる場合にのみ、 $D = [0 \ 0]^T$ つまり2母音間に音質の差が存在しないことを含意している。無論これは言語学的にみて強過ぎる仮定であって、よく知られているように、人間の音声はたとえ同一話者の発話内でも通常個々の発音毎に何らかの誤差を伴うものである。この誤差に加え、実際には更にフォルマント周波数測定上の誤差が加わる。結局、 D の個々の値を問題とするべきではなく、(1)式が意味をもつのは N 個の D の分布についての統計的性質が問題とされる場合なのである。或る地域で2母音 j, k 間に一定の音質上の差が社会的言語習慣として保持されているならば、 N 個の差ベクトルの分布には上述の誤差成分と共に θ が一定の範囲に集中するという形でその習慣が反映されているはずである。反対に2母音が合一の状態にあるならば、つまり音韻論上ひとつの母音を2回繰り返して発音した状態に等しいならば、 θ の分布には何ら組織的な傾向は見出されないはずである。このことは、2母音の合一の状態につき、次のような θ の確率モデルを想定することに等しい。

$$(2) \quad \theta \sim U(0, 2\pi]$$

(2)式の主張するところは、 N 個の θ の分布が区間 $(0, 2\pi]$ で一様分布(uniform distribution)に従うということであり、換言すれば、誤差成分の分布として一様分布を想定するわけである。いうまでもなく、この場合の密度関数は、

$$f(x) = \Pr \{ \theta = x \} = 2\pi^{-1}$$

で与えられ、その期待値 $E(x)$ 、分散 $V(x)$ は下記のとおりである。

$$E(x) = \int_0^{2\pi} x f(x) dx = \pi$$

$$V(x) = \int_0^{2\pi} \{ x - E(x) \}^2 f(x) dx = \frac{1}{3} \pi^2 \\ \simeq 3.29$$

4. モデルに基づいたデータの解析

θ の分布に関する我々の確率モデル(2)は音声学的言語学的な考察から理論的に想定されたものである。このモデルを実際のデータの解析に利用する時、(2)は通常の統計的仮説検定における帰無仮説(H_0)の役割を果たすことになる。これを改めて下記のように述べておこう。

$$H_0 : \theta \sim U(0, 2\pi]$$

尚、ここで H_0 に対する対立仮説としては、 H_0 以外のすべての場合を考えることにする。つまり言語学上の対立が全く認められないか(H_0 採択)、そうとはいえないか(棄却)を問うわけである。この場合、 H_0 の棄却がただちに言語学上明確な対立の存在を積極的に肯定するわけではないことはいうまでもない。

表3, 4, 5に3地方の可能なすべての2母音の組み合わせ(${}_3C_2 = 10$ 通り)について計算した θ の標本値分布の基礎統計量を示す。これらの数値のなかのどれが H_0 に示された母集団分布からの標本と見做されうるものかを客観的に判定したい。そのために、ここで中心極限定理(central limit theorem)と呼ばれる数学上の著名な定理を利用しよう。この定理が成立するための十分条件は、 N 個すべての標本が独立に同一の分布から採られた標本(i, i, d の標本)であることだが、我々の仮説 H_0 に従えば、これは保証される。中心極限定理によれば、 θ の標本平均を $\bar{\theta}$ 、標本分散を S^2 とすると、

$$(3) \quad Z = (\bar{\theta} - \pi) \cdot \sqrt{N} / S$$

なる Z は、 $N \rightarrow \infty$ の時、標準正規分布 $N(0, 1)$ に従って分布する。この Z を利用して

表3. 東京地方 θ の分布

母音対	$\bar{\theta}$	S. D.	N. V.	Z
イ, エ	5.42	0.22	0.01	44.9
エ, ア	5.03	0.11	0.01	72.8
ア, オ	4.08	0.14	0.01	27.8
オ, ウ	2.20	0.35	0.04	-11.4
ウ, イ	1.61	0.06	0.00	-108
イ, ア	5.15	0.09	0.00	94.7
イ, オ	4.86	0.05	0.00	145
エ, オ	4.68	0.06	0.00	108
エ, ウ	4.45	0.15	0.01	37.0
ウ, ア	6.48	0.36	0.04	39.3

(Radian)

*表中N. V. (Normalized Variance) とは、 θ の標本分散値を、一様分布モデルより得た理論値 $\pi^2/3$ で除した値、表4, 5も同じ。

表4. 津軽地方 θ の分布

母音対	$\bar{\theta}$	S. D.	N. V.	Z
イ, エ	2.92	1.83	1.02	-0.45
エ, ア	5.15	0.12	0.00	62.6
ア, オ	4.12	0.19	0.01	19.3
オ, ウ	1.61	0.13	0.00	-44.1
ウ, イ	1.62	0.11	0.00	-51.8
イ, ア	5.16	0.10	0.00	75.5
イ, オ	4.75	0.03	0.00	201
エ, オ	4.75	0.04	0.00	151
エ, ウ	4.76	0.12	0.00	50.5
ウ, ア	5.95	0.47	0.07	22.4

(Radian)

表5. 出雲地方 θ の分布

母音対	$\bar{\theta}$	S. D.	N. V.	Z
イ, エ	2.10	1.71	0.88	-3.11
エ, ア	5.05	0.12	0.00	81.1
ア, オ	4.26	0.12	0.00	47.5
オ, ウ	2.20	1.23	0.46	-3.9
ウ, イ	1.65	0.11	0.00	-69.1
イ, ア	5.20	0.20	0.01	52.5
イ, オ	4.79	0.06	0.00	140
エ, オ	4.76	0.05	0.00	165
エ, ウ	4.75	0.08	0.00	103
ウ, ア	0.86	0.43	0.06	-27.1

(Radian)

H_0 の妥当性を3地方、各10組の母音対について検定してみよう。⁵

(3)式によって計算されたZの値は、表3, 4, 5のなかに示しておいた。全体として眺めると、予想通りに津軽地方のイ, エ 出雲地方のイ, エとオ, ウの3対のZ値がその他の対に比べて非常に小さくなっている。このなかで、津軽地方イ, エの対に関しては検定の水準を両側5%にとっても H_0 は棄却されない(片側5%の検定でも H_0 は採択されるが、このような棄却域の設定は今回の検定の主旨からみて明らかに不自然

である)。一方、出雲地方のイ, エとオ, ウの対に関しては、検定水準を行動科学では最も厳しい水準のひとつといえる両側1%にとっても H_0 は棄却されてしまう。他の27組

注5 $\bar{\theta}$ やZの分布の正規型への接近は標本の採られる分布の型とNの大きさに影響される。文献4 P125 f に一様分布からN=10の標本を採った場合の実験が報告されているが、近似はかなり良好のようである。

の対に関しては、いうまでもなく棄却である。

さて上に述べた確率論上の事実は、次のような解釈を許すものである。すなわち同じく $f_1 - f_2$ 平面上での母音の重複現象を示すとはいっても、津軽地方と出雲地方とでは現象に質的な相違が認められる。前者（イ、エの対）では我々の θ で表現される2母音の位置関係が数学上の意味でランダムに決定されており（つまり2母音が相互にいかなる位置関係にある確率も等しい）、これは言語学的には次のような主張を可能とする。津軽方言共

図7. 3地方における θ の分布
母音対イ、エの場合

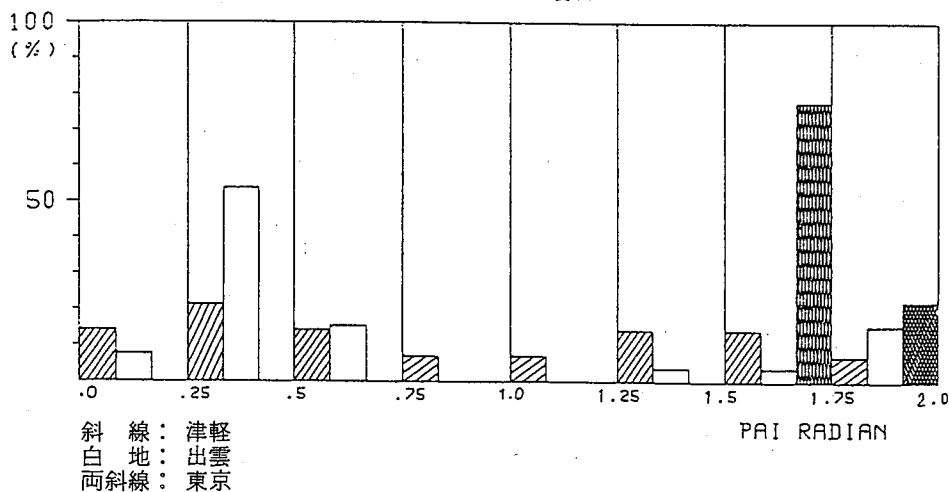
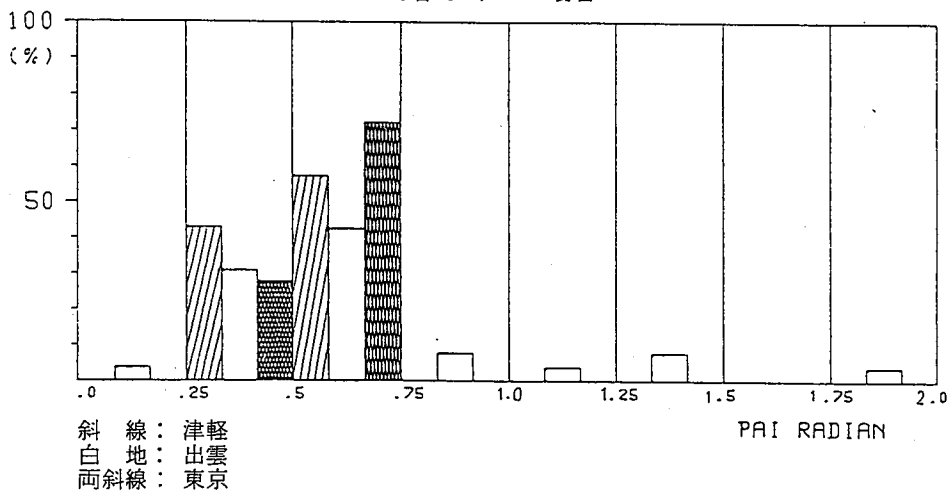


図8. 3地方における θ の分布
母音対オ、ウの場合



同体の成員は2母音イ、エの間に一定の音質上の差を生じさせる社会的言語習慣を有しておらず、イとエは合一の状態にある。この地方では「息」と「駅」は完全な同音異義語である。一方、出雲地方（イ、エの対、オ、ウの対共に）では2母音の位置関係は数学上の意味でランダムということはできない。つまり、ここでは θ が特定の範囲（あらかじめひとつとは限らない）にある確率がそれ以外よりも高いはずである。このことは、既出の図4、5、6における D^* の散布からも確認できるが、この点をより明瞭にするため、東京を含めた3地方での θ の散布状態をヒストグラムにしたのが、図7、8である。両図からは、まず津軽地方イ、エの θ が8つの階級すべてにわたって、ほぼ一様に散らばっていることが見てとれる。しかし、それに加え、更に興味深いのは出雲地方の θ がイ、エとオ、ウのいずれの場合にも単峰型（uni-modal）といってよい分布を示していることである。ヒストグラムから読みとれるように出雲地方いずれの母音対でも θ の70%前後が最大 0.5π ラジアン以内の連続した区間に集中している。これは言語共同体成員のかんりの部分が、各母音対の2母音間でよく似通った音質の差を共有していることを示しており、その点では出雲地方の両母音対は、津軽地方のイ、エよりは、むしろ東京地方のイ、エとオ、ウや津軽地方のオ、ウなどの θ の分布型によく類似している。しかし、他方で東京地方の両母音対や津軽地方のオ、ウでは100%の θ が最大 0.5π ラジアン以内の連続区間に存在しており、この点では出雲地方と好対照をなしている。結局、出雲地方イ、エとオ、ウの両母音対は合一の状態と確定的な母音対立の中間態ということになり、本稿の最初に規定した混同の状態にあるものと見做されるのである。

5. まとめと展望

本稿では、まず日本語の諸方言中には東京方言などには認められない母音のフォルマント平面上での重複現象を示すものがあることを津軽、出雲両方言の分析資料により示した。ついでこの現象の本質を抽象的に数量化して表現する方法を述べ、その上で2母音が完全に対立を解消した合一状態に関する確率モデルを呈示した。このモデルに基づいた解析の結果、同じ母音の重複現象中にも質的に区別すべきふたつの状態、合一と混同が存在することが立証された。津軽方言のイとエは合一、出雲のイとエ、オとウは混同の状態にある。

本稿で提案されたデータ解析の方法は、任意の2母音間の対立をフォルマント平面上での位置関係についての社会的言語習慣として捉えるものであり、一種のdirectional

data analysisに相当する。⁶ 現時点における、この方法の弱点は各発話者が各調査語を一回ずつ発音したデータ（もしくは平均値のデータ）しか取り扱えないことである。この点はいずれ改良しなければならない。

又、言語学的見地からは次のような点が問題となるであろう。第一に調査語の数が5つと限定されている点であるが、これに関しては本稿の主旨が必ずしも津軽方言や出雲方言の母音体系全体に関係するものではないことを断っておきたい。第二に本論では使用したデータが各方言共同体を代表する標本であるように看做して議論を進めたが、実はこのデータはランダム・サンプリングによるものではない。ただし、方言研究において一体何を母集団として規定すべきかは今日まで明快な解答を与えられていない根本問題である。最後に最も興味深い問題として、出雲方言の現在の混同状態が将来合一へと向かう過程にあるのか、それとも合一から新たな母音分離への過程にあるのかという通時論的な疑問がある。この問題の解明には文献資料の援用と同時に、より広い年齢層、社会言語的場面にわたる新たな臨地調査を計画せねばならない。更に若年層を中心とする全国共通語化への動向をもあわせた重層的な言語生活の把握が必要不可欠となろう。今後ただちには十全な議論を期し難い問題であるが、将来の課題として第一級の重要性を有するものである。

文献

1. 前川喜久雄, 2母音間の混同を定量的に把握する方法, 『日本方言研究会第35回研究発表会発表原稿集』(1982) pp. 1-9.
2. 菅原勉, 小島慶一 他, 津軽, 下北両地方の母音, *Sophia Linguistica* VIII / IX (1981), pp. 196-208.
3. 前川喜久雄, 井上美穂 他, 出雲地方の母音, *Sophia Linguistica* XI (1982), pp. 75-82.
4. ホーエル・P. G. 『入門数理統計学』浅井, 村上訳(培風館: 1978).
5. Mardia, K. V., et al. *Multivariate Analysis* (Academic Press: 1979).

(受付: 1983年9月22日)

注6 球面(円周)上の点の分布を論じるdirectional data analysis(D. D. A)に関しては、例えば文献5の15章“Directional Data” PP. 424-451を参照されたい。本稿で論じた諸問題はD. D. Aの基礎的分布であるvon Mises-Fischer分布によれば、より一般的に処理できそうである。

付録： D_{ijk} の属する象限を決定するアルゴリズム

