

SPデータを用いた交通需要予測のための マーケット・セグメンテーションに関する研究

MARKET SEGMENTATION FOR TRAVEL DEMAND ANALYSIS USING STATED PREFERENCE DATA

森川高行*・白水靖郎**

By Takayuki MORIKAWA and Yasuo SHIROMIZU

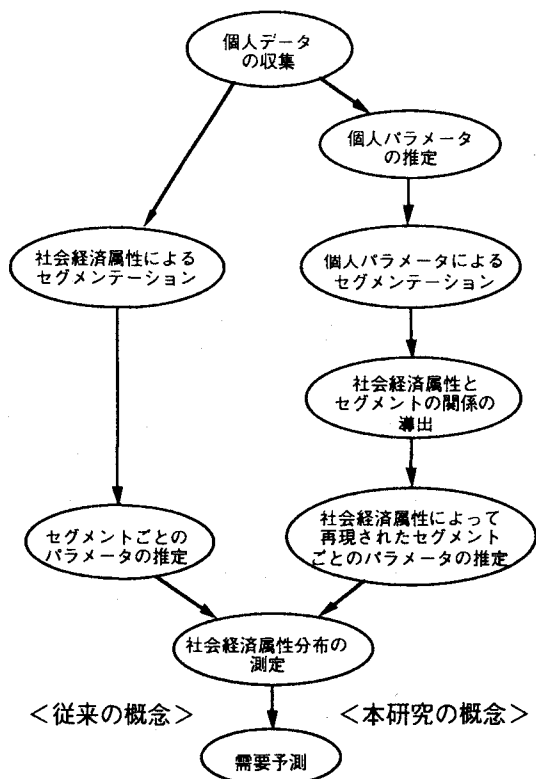
This paper presents a methodology for market segmentation for disaggregate travel demand models using stated preference (SP) data. Multiple responses from each individual in SP data enable one to estimate individual-specific parameters of choice models. Clustering techniques such as visual one and cluster analysis are employed against the space of individual-specific parameter vectors in order to divide the whole sample into segments which have relatively homogeneous taste. The segments obtained are reproduced using socio-economic variables. Then, choice models are estimated for each reproduced segment using revealed preference (RP) data to evaluate effectiveness of the segmentation. Two empirical analyses are also presented.

1. はじめに

交通需要予測に用いられる非集計モデルでは、効用関数のパラメータは通常対象とする母集団の中で同一であると仮定している。このようなモデルでは、同質的な集団は、全員が選択対象の属性に基づく効用評価において同一の構造を有していると仮定しているため、どのようにして母集団から同質的な集団（セグメント）を抽出するかが問題となる¹⁾。

従来は、個人の社会経済属性によって事前にセグメントを規定したりあるいは母集団のままでも分類（セグメンテーション）を行わずに需要予測を行なうことが多かったが、この方法では得られたセグメントは本当に同質的であるのか、つまり「嗜好」の似たものが集まっているのかが不明であった。

そこで本研究では、個人の「嗜好」は効用関数のパラメータで表されることに着眼し、個人ごとに推定されたパラメータ値の類似性によって選択構造の似たセグメントに分類する手法を提案する。ここで、マーケット・セグメンテーションに対する従来の概念と本研究の概念の違いは図-1のように表される。



* 正会員 Ph.D. 名古屋大学助教授 工学部土木工学科
(〒464-01名古屋市中千種区不老町)

** 正会員 中央復建コンサルタンツ株式会社
(〒530 大阪市北区梅田1-2大阪駅前第2ビル404-7)

図-1 マーケット・セグメンテーションの考え方

ところで、従来非集計モデルでは、市場における実際の行動結果であるRP (Revealed Preference) データを用いることが多かったのだが、政策の変化や新しいサービスに対する需要予測を行なうために、仮想の状況下における選好の意思表示であるSP (Stated Preference) データを用いる考え方が、1970年代からマーケティング・リサーチの分野で発達してきた。近年、交通行動モデルにおいても、新しい交通サービスに対する需要予測にSPデータを用いることが多くなってきている²⁾。SPデータは、一種の実験データであることから操作性が高く、また個人ごとに選好に関する情報を多く得ることができるため個人ごとに未知パラメータを推定することが可能になる。本研究においても、個人単位のパラメータを推定する際にSPデータを用いる。ただしSPデータを用いる際の注意点として、その信頼性に関する問題を考慮しなくてはならない。すなわち仮想の状況下と実際の行動の意思決定構造の乖離の問題である。そこで本研究では、SPデータを用いたセグメンテーションが実際の行動における「嗜好」を正しく表しているかどうかをRPデータを用いて検証することによって、その信頼性に関する問題を検討する。

以上のように、本研究ではSPデータを用いて個人単位に効用関数の係数を推定し、その推定値よりサンプル全体を同質の選好構造を持つ幾つかのグループに、マーケット・セグメンテーションを行なう手法を提案する。そしてこの手法の有効性を、東京で行なわれた深夜交通機関選択に関するアンケート調査データと、オランダで行なわれた都市間旅行交通機関選択に関するアンケート調査データを用いて検証する。

本稿は次のように構成される。まず2.において本研究におけるSPデータを用いたマーケット・セグメンテーションの考え方を述べる。続いて3.、4.では、2.で提案した方法の有効性を検証するため、実際のデータを用いて実証的研究を行なう。最後に5.では、3.、4.の結果から得られた知見と今後の展望について述べる。

2. SPデータを用いたマーケット・セグメンテーション

(1) マーケット・セグメンテーションとその必要性

本研究で提案する手法は、マーケティング・リサーチの分野におけるマーケット・セグメンテーションの

概念を根底としている。マーケット・セグメンテーションとは、人は異質的であるという認識の下で何らかの意味で同質的なセグメントに市場を分割することによって、より有効できめ細かなマーケティング活動を行なおうというものである。セグメンテーションを行なうためには、同質的なセグメントを識別するための基準を選ばなければならない。その識別方法は、事前に基準を定めそれによってセグメンテーションを行なうアプリアリ・セグメンテーションと、多次元的な変数を基準としてセグメンテーションを行なうクラスタリング・セグメンテーションの2つに大別される³⁾。従来交通需要予測で行なわれていた個人の社会経済属性を用いたセグメンテーションは前者のアプリアリ・セグメンテーションに相当し、本研究で用いる手法は後者のクラスタリング・セグメンテーションに相当する。

ところで、マーケット・セグメンテーションは、効用関数のパラメータが個人ごとに異なるが母集団内のあるセグメント内では十分にそれが均一であるとみなせるときに有効である。もしそれぞれのパラメータが個人の社会経済属性で表されていることがアプリアリにわかっている場合には、効用関数の係数を個人属性の線形関数として次のように表すことができる。

$$\begin{aligned} u_{in} &= ([\alpha]y_n)'x_{in} + \varepsilon_{in} \\ &= v_{in} + \varepsilon_{in} \end{aligned} \quad (1)$$

u_{in} : 個人 n に提示された代替案 i に対する効用

v_{in} : 個人 n に提示された代替案 i に対する効用の確定項

ε_{in} : 個人 n に提示された代替案 i に対する効用のランダム項

$[\alpha]$: 個人 n の未知パラメータ・マトリクス

x_{in} : 個人 n に提示された代替案 i に対する説明変数ベクトル

y_n : 個人 n の社会経済属性ベクトル

(1)式のより厳密な定式化は、効用関数の係数の定式化にもランダム項を導入することによって次のように表される。

$$\begin{aligned} u_{in} &= ([\alpha]y_n + \zeta_n)'x_{in} + \varepsilon_{in} \\ &= v_{in} + \xi_{in} \end{aligned} \quad (2)$$

$$v_{in} = ([\alpha]y_n)'x_{in} \quad (3)$$

$$\xi_{in} = \zeta_n'x_{in} + \varepsilon_{in} \quad (4)$$

ξ_{in} : 個人 n に提示された代替案 i に対する効用のランダム項

ζ_n : 未知パラメータのランダム項ベクトル

ε_{in} : 個人 n に提示された代替案 i に対する効用のランダム項のうち ζ_n で表せないもの

(1)式, (2)式で表されるモデルは, 効用関数の係数で表される「選好」のばらつきを表現していることから Taste Variation Model と呼ばれ, さらに(2)式は係数のランダム性を表していることから Random Coefficient Model と呼ばれる⁴⁾.

Random Coefficient Model は, 個人 n の代替案 i に対するランダム項 (ξ_{in}) が(4)式のようになり, 独立で同一な分布をしていないため, より一般的なプロビット・モデルの推定を必要とし計算上実用的ではない。そこで本研究では, セグメンテーションを行わないモデルとして, 係数のランダム項を考慮していない Taste Variation Model ((1)式) を考え, マーケット・セグメンテーションの必要性についても吟味した。

(2) 個人パラメータ推定モデル

クラスターリング・セグメンテーションの範疇で, 本研究のように選好構造の似たものをセグメンテーションする手法は, ベネフィット・セグメンテーションと呼ばれる。いわゆる便益(ベネフィット)の共通な人々によるセグメントを構成する手法である。この種のセグメンテーションの際に問題となるのは個人ごとの行動モデルの構築に関する問題である。

近年この個人モデルの問題について, コンジョイント分析の考えを導入する研究が盛んである⁵⁾⁶⁾⁷⁾⁸⁾。しかし個人レベルの小サンプルで Logit モデルなどの離散型モデルを推定する際には, 「解の不定性」の問題が生起する。これは, 「矛盾のない序列」の問題とも言われるもので, コンジョイント分析において, 与えられた順序関係が加法的表現と矛盾しない場合には, その解空間は凸多面錐を構成し, 解がメトリックな意味で不定になるというものである⁷⁾。従来の研究では, この問題に対してパラメータを基準化し誤差分散の相対的な大きさを外生的に与える方法⁷⁾や各選択肢の相対的選択確率をアンケートで各回答者に答えさせそれを

用いて各個人ごとに誤差分散を推定する方法⁹⁾が提案されているが, 「解の不定性」の問題は現在のところ完全な解決法はないと言える。

そこで本研究では「個人モデルによって, 個人ごとの選好構造の概形がわかればよい」という視点に立ち, 一対比較質問を段階評定で回答させ, その評定を連続変数と見なすことによって最小2乗法により個人パラメータを求める。このように順序尺度である段階評定を比例尺度の近似として未知パラメータを推定することは理論的には問題があるが, 質問方法によってはかなり良い近似を与えられると思われる。これは次式のように定式化される。

$$y_n = \beta_n' \Delta x_n + \varepsilon_n \quad (5)$$

y_n : 個人 n に提示された一対比較質問の評定

ε_n : 個人 n の線形個人モデルのランダム項

β_n : 個人 n の未知パラメータベクトル

Δx_n : 個人 n に提示された代替案1と代替案2に対する説明変数の差のベクトル

(3) 個人パラメータによるセグメンテーション

個人パラメータを用いたマーケット・セグメンテーションの手法としては, クラスタ分析と視覚的分割法を用いる。クラスタ分析は最もポピュラーな自動的分類手法であるが, 手法が強力である故に珍妙なセグメントを抽出したり, アルゴリズムによって異なるセグメント形成するなど, いくつかの問題点がある。視覚的分割法はデータをヒストグラムや散布図など視覚的に接近しやすいように表現し, 視覚的にセグメントを探す手法であり, 分析者の主観の影響が大きい欠点があるがクラスタ分析では判断しにくい湾曲したり細長いセグメントに対応することができる。しかしどちらの手法に関してもセグメント数を決定する厳密な方法はないので本研究ではいくつかの方法を試みる。

また個人パラメータの扱い方にも, 個人パラメータをそのまま用いる方法と個人パラメータを加工してから用いる方法の2通りの方法が考えられる。前者は, 上記の視覚的分割法などにおいて軸の意味合いがはっきりしているので, セグメントの解釈がしやすいが同時に多くのパラメータを評価することはできない。後者は, 主成分分析や因子分析などを用いてパラメータ空間の次元を減らしてからセグメンテーションを行なう。この手法では同時に多くのパラメータを評価する

ことができるが、主成分や因子の解釈ができないと意味のわかりにくいセグメントが形成される。

このような様々な手法によるセグメンテーションを評価するために、各セグメントごとに未知パラメータを推定した後、求められたセグメント間に有意な差が見られるかを検証し、セグメンテーションを行なわないモデルと適合度を比較する。

また2.で述べたように、本研究では求められたセグメントが実際の行動におけるセグメントと一致しているかを検証する。そこで、実際の行動結果であるR Pデータを用いて求められたセグメントごとに未知パラメータを推定し、上記のS Pデータを用いて推定された未知パラメータと対比・吟味を行なうことによってS Pデータの信頼性に関する事後検証を行なう。

(4) 社会経済属性とセグメントの関係

S Pデータを用いたマーケット・セグメンテーションを交通需要予測に活用するには、求められたセグメントと個人属性の関係を明らかにする必要がある。というのは、個人属性データから個人の所属セグメントが判別できれば、セグメントごとの需要予測が行なえるからである。本研究ではセグメントを再現する関数として、数量化2類と多項ロジットモデルを採用し両者を比べてみる。数量化2類はこのような判別関数としてマハラノビスの汎距離による判別分析と並んでよく用いられる。多項ロジットモデルを判別関数として用いることは過去例にないが、個人属性によるランダム効用モデルを構築し、その効用値が最大のセグメントに個人が属すると考えることによって判別することが可能である。

(5) 個人属性によって再現されたセグメントの信頼性

さて次に、前節で述べた手法を用いて再現されたセグメントが需要予測に有効であるか、もとのセグメントの性質を再現しているかその信頼性を検証する必要がある。そこで再現されたセグメント（以下フィッティッド・セグメントと呼ぶ）ごとに未知パラメータを推定する。フィッティッド・セグメントは、次の性質を持っているほど信頼性が高いと思われる。

- a) セグメントの再現率が高い。
- b) フィッティッド・セグメントのオブザベーション数と、もとのセグメントのオブザベーション数の変化が少ない

- c) フィッティッド・セグメントの未知パラメータと、もとのセグメントの未知パラメータの傾向が似ている。

このように本研究では、S Pデータを用いたマーケット・セグメンテーションの信頼性と、需要予測における有効性の検証を核としているところが従来の研究とは異なる。

3. ケーススタディー1

ケーススタディー1では次のことを中心に検討した。

- a) マーケット・セグメンテーションの必要性
- b) 個人パラメータを主成分分析した結果によるセグメンテーション
- c) マーケット・セグメンテーションを行なうことによってどれだけモデルの適合度が向上するか
- d) フィッティッド・セグメントを用いることによってどれだけモデルの適合度が向上するか

(1) データの概要

ケーススタディー1では、東京急行電鉄株式会社交通事業部と株式会社三菱総合研究所によって昭和63年に東京で行なわれた、「都市の24時間化に伴う将来交通体形の在り方に関する調査」データを用いる。

この調査のうち、ケーススタディー1で用いるデータは、次の2種類のデータである。

- a) 社会経済属性に関するアンケート
性別・年齢・職業・住所・深夜交通機関利用場所等
- b) 深夜交通機関選択に関するアンケート
深夜交通機関の代替案として、新交通機関となる「終夜バス（現行の深夜バスよりも遅い時間帯に運行するバス）」とその競合相手として「タクシー」を設定している。そして、終夜バス運行条件を変化させることによって16通りのシナリオを設定し、それぞれのシナリオについて一対比較質問を3段階評定法で行ない16個のS Pデータを得ている。ここで運行条件とは、「全所要時間」、「アクセス時間」、「イグレス時間」、「運行間隔」の4条件である。また3段階評定法は、「タクシーを利用する」、「甲乙付けがたい」、「深

夜バスを利用する」という形式で質問された。

(2) マーケット・セグメンテーションの結果

まず、このデータに対してマーケット・セグメンテーションが必要であるか検討するため、2. (1)で述べた、係数のランダム項を考慮していないTaste Variation Model (1)式)を用いて未知パラメータを求めた。この結果未知パラメータのうち社会経済属性の関数になりそうなものは定数項を除いてほとんどなく、マーケット・セグメンテーションは必要だと思われる。

そこで従来のように個人の社会経済属性を用いてア prioriにセグメンテーションを行ない、その結果得られたセグメントごとに未知パラメータを推定した結果が表-1である。この結果、各セグメント間で推定されたパラメータに大きな差異は見られず、また $\bar{p}^2$ よりモデルの適合度もあまり向上していないことがわかる。このように、従来の方法では個人の嗜好の違いを反映したセグメントが形成されているとはいえない。

そこで、本研究で提案する個人パラメータを用いたマーケット・セグメンテーションを行なう。まず(5)式によって個人パラメータを推定する。ここで説明変数は、「全所要時間」、「アクセス時間」、「イグレス時間」、「運行間隔」の4条件である。

これより各個人につき5つの個人パラメータが推定された(定数項を含む)。2. (3)より5つのパラメータを同時に評価することは困難であるので、個人パ

表-1 性、年齢によるセグメンテーション

	pooled	Male -Old	Male -Young	Female -Old	Female -Young
Constant	2.27	2.19	2.52	2.37	2.26
Cost	-0.0654	-0.0625	-0.0609	-0.206	-0.0628
Total time	-0.0907	-0.0944	-0.0899	-0.0964	-0.0878
Access time	-0.187	-0.193	-0.186	-0.212	-0.180
Egress time	-0.186	-0.190	-0.179	-0.218	-0.190
Interval	-0.0755	-0.0807	-0.0739	-0.0850	-0.0723
N	19216	5312	7920	1440	4544
$\bar{p}^2$	0.334	0.348	0.333	0.381	0.316
$\bar{p}^2$ of unrestricted models	0.337				

表-2 主成分分析

	Principal 1	Principal 2
Constant	0.02181	-0.5505
Total time	0.03110	0.2870
Access time	1.511	-0.04984
Egress time	0.5725	-0.3924
Interval	0.7858	0.3857
Eigenvalues	3.231	0.6907
Percent variance(%)	64.61	13.81
Cumulative percent variance(%)	64.61	78.42

ラメータを主成分分析することによってパラメータ空間の次元を減らしてから、マーケット・セグメンテーションを行なうことにした。ここで主成分分析を行なった結果が表-2である。この表より第1主成分と第2主成分を合わせると寄与率は80%弱であり、この主成分分析が有効であることを示している。

さて2. (3)で述べた通り、主成分分析の結果をマーケット・セグメンテーションに利用するためには、各主成分の解釈を明確にする必要がある。第1主成分は、アクセス時間、運行間隔、イグレス時間の順に正の相関が高いので、「交通機関に乗車するまでの時間を重要視した端末時間(Terminal Time)」を表していると解釈できる。一方第2主成分は、バス定数項、イグレス時間と負の相関が高く、運行間隔、全所要時間と正の相関が高いので、「バスタクシー」軸を表していると解釈できる。

この結果を用いて、2. (3)で述べたクラスター分析や視覚的分割法で様々なマーケット・セグメンテーションを行なった。ここではその中でも視覚的分割法の結果を掲載する(他は紙面の都合で割愛)。

図-2は、第1・第2主成分の散布図を3本の直線で3セグメントに分割したものである。ここでセグメント1は第1主成分、第2主成分ともに負のセグメン

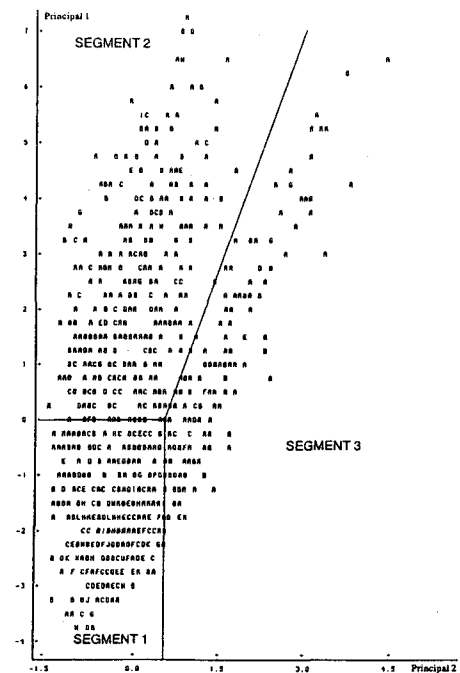


図-2 視覚的分割法

トと考える。ただし、第2主成分を0.5で区切ったのは、0.5を誤差と考えることによって視覚的な要因を尊重するためである。またセグメント2とセグメント3の境界線は、第1主成分と第2主成分の比率による境界と考えた。

求められたセグメントごとに未知パラメータを推定した結果が表-3である。この結果より各セグメント間でパラメータの推定値が有意に異なっていることが尤度比検定から確かめられる。また従来の手法に比べて、推定されたパラメータが、各セグメント間で異なっていることが確認される。さらにセグメンテーションすることによって $\bar{p}^2$ 値が著しく向上しておりモデルの適合度が向上したといえる。

以上より個人パラメータによって選好構造の似たものに分類されたセグメントは、従来の方法で求めたセグメントに比べて、個人の嗜好の違いを反映しているといえる。

(3) 社会経済属性と求められたセグメントの関係

このケーススタディー1のデータには、個人属性として「性別」、「年齢」、「職業」、「帰宅距離」の4種類の情報が含まれている。これらのデータは「帰宅距離」を除いてカテゴリカルデータであるので、2.

(4)で述べたように数量化2類と多項ロジットモデルを用いて社会経済属性と求められたセグメントの関係を求めた。

これらの結果より、どちらの手法を用いても個人属性データによるセグメントの再現率は低く(25%~60%)予測に用いる数字としては信頼性に欠ける。

しかし再現率は低い、個人属性によって再現されたセグメントごとに選択モデルを推定し、そのパラメータがフィッティッド・セグメントごとに有意に異なれば、このようなセグメンテーションの手法には、意味があるといえる。そこでフィッティッド・セグメントごとに、未知パラメータを推定した。

その結果、フィッティッド・セグメントごとにパラメータの差異は認められるが、元のセグメントに比べてその差異の大きさはずいぶん小さく、個人の社会経済属性に基づいた従来のセグメンテーションと同じ程度であった。また $\bar{p}^2$ 値に関しても、従来の手法と同程度もしくはそれ以下しか向上しておらずモデルの適合度も上がっていないといえよう。さらにパラメータの性格がもとのセグメントのものと違っているため予測に用いるには、問題があると思われる。

表-3 視覚的分割法によるセグメンテーション

	pooled	Segment1	Segment2	Segment3
Constant	2.27	3.63	2.99	1.54
Cost	-0.0654	-0.0398	-0.0467	-0.0714
Total time	-0.0907	-0.142	-0.0910	-0.0686
Access time	-0.187	-0.306	-0.156	-0.210
Egress time	-0.186	-0.293	-0.160	-0.244
Interval	-0.0755	-0.119	-0.0771	-0.0567
N	19216	10208	6192	2816
$\bar{p}^2$	0.334	0.507	0.441	0.364
$\bar{p}^2$ of unrestricted models			0.465	

5. ケーススタディー2

ケーススタディー2では、次のことを中心に検討した。

- 個人パラメータをそのまま用いるマーケット・セグメンテーション
- マーケット・セグメンテーションを行なうことによってどれだけモデルの適合度が向上するか
- フィッティッド・セグメントを用いることによってどれだけモデルの適合度が向上するか
- SPデータより得られたセグメントをRPデータに適用した時どの程度のバイアスが生じるか

(1) データの概略

ケーススタディー2では、Netherlands Railwaysの依頼でHague Consulting Groupによって1987年にオランダで行なわれた「都市間旅行における列車と車の機関選択に影響する要因に関する調査」データを用いた。

この調査のうちケーススタディー2で用いるデータは、次の4種類のデータである。

- 都市間旅行に関するデータ(RPデータ)

このRPデータは、「選択モード(列車or自動車)」、「トリップ目的」、「幹線交通時間」、「端末交通機関(アクセス・イグレス時間)」、「乗り換え回数」、「定性的サービスの主観適評価」、「社会経済属性」の7項目で構成される。

- 2種類の異なるサービスの列車に対する選好意識調査(SP1データ)

このSP1データは、「列車vs列車」について、フル・プロファイル対比較質問を5段階評定法で行ない、平均して一人当たり12~13個データを得た。このときの質問で用意された仮想のシナリオは、「旅行費用」、「幹線交通時間」、

「乗り換え回数」，「列車の快適性」の4条件を変化させることによって作成された。

c) 自動車と列車に対する選好意識調査 (SP2データ)

このSP2データは，「列車 vs 自動車」について，フル・プロファイル対比較質問を5段階評定法で行ない，平均して一人当たり7個の回答を得た。このときの質問で用意された仮想のシナリオは，「旅行費用」，「幹線交通時間」，「乗り換え回数」，「列車の快適性」の4条件を変化させることによって作成された。

d) 社会経済属性に関するデータ

性，年齢，職業，旅行人数，旅行目的等

(2) マーケット・セグメンテーションの結果

まず，このデータに対してマーケット・セグメンテーションが必要であるか3. (2)と同様に検討した結果，必要であることが分かった。

さて，ケーススタディー2では，上記のSP1データとSP2データからそれぞれ別々に未知パラメータを推定した。というのは，両データから得られる情報が互いに補完的であると考えられるからである。すなわちSP1データでは，同じ交通機関である列車どうしのサービスを変化させることによってデータを得ているため，純粋な意味での，費用，時間，快適性などのパラメータが得られると考えた。一方SP2データでは，列車と自動車を比較しているため，信頼性の高い列車定数項 (Rail Constant) のパラメータを得ることができると考えた。

さて，前述したようにSP1データからは，信頼性の高い「費用」，「幹線旅行時間」，「乗り換え回数」，「快適性」のパラメータが得られた。そこで，幹線旅行時間のパラメータを費用のパラメータで除したもので定義される時間価値パラメータに注目し，「時間価値の高いもの」，「時間価値が特に高くはないが正のもの」，「時間価値が負のもの」の3セグメントにセグメンテーションすることを考えた。

一方SP2データは，列車定数項でセグメンテーションするのが妥当だと考えた。というのは，SP2データの調査では，どのようなシナリオにおいても一つの交通機関しか選択しない人，すなわち「列車キャプティブ層」と「自動車キャプティブ層」が存在しており，残りの人も列車そのものに対する選好，つまり列車定数項でセグメンテーションを行なうことが自

表-4 性，年齢によるセグメンテーション

	pooled	Male -Old	Male -Young	Female -Old	Female -Young
Rail constant	-7.22		-2.11		-10.6
Cost per person	-0.828	-0.981	-0.934	-0.712	-0.812
Line-haul time	-0.967	-0.544	-1.08		-1.54
N	4452	732	1281	741	1698
$\bar{p}^2$	0.341	0.376	0.364	0.439	0.298
$\bar{p}^2$ of unrestricted models	0.350				

表-5 時間価値によるセグメンテーション

	pooled	High	Low	Negative
Rail constant	-7.22		-5.59	
Cost per person	-0.828	-0.368	-1.32	-0.685
Line-haul time	-0.967	-2.83	-1.25	0.402
Value of time	1.18	7.69	0.947	-0.587
N	4452	801	2703	948
$\bar{p}^2$	0.341	0.334	0.362	0.418
$\bar{p}^2$ of unrestricted models	0.365			

表-6 列車定数項によるセグメンテーション

	pooled	Positive	Negative
Rail constant	-7.22	-2.57	-15.6
Cost per person	-0.828	-0.829	-0.863
Line-haul time	-0.967	-1.07	-1.06
N	4452	1706	774
$\bar{p}^2$	0.341	0.211	0.312
$\bar{p}^2$ of unrestricted models	0.243		

然だと考えられるからである。よって前述した2つのキャプティブ層以外に，列車定数項の正負によってセグメンテーションを行なった。

まず，従来のように個人の社会経済属性を用いて求められたセグメントごとに未知パラメータを推定した結果が表-4である。次に時間価値を基準にセグメンテーションを行なった結果が表-5であるが，未知パラメータの推定値が各セグメントで大きく異なっており，そのうえ推定値から求められる時間価値パラメータが，セグメントの性質を反映したものとなっている。さらに $\bar{p}^2$ 値も従来の手法に比べより向上しておりモデルの適合度が向上したといえる。また，列車定数項を基準にセグメンテーションを行なった結果が表-6であるが，未知パラメータの推定値は各セグメントそれほど異なっていないものの，列車定数項のパラメータは大きく異なっており，セグメントの性質を反映したものとなっている。また $\bar{p}^2$ 値は低下しており一見モデルの適合度が落ちたかのように見えるが，この分析はノン・キャプティブ層のみを対象としたモデルであり (キャプティブ層は，効用関数の独立変数が一定であるため推定できない)，全サンプルを用いたpooled

モデルと比較することはできない。キャプティブ層に対する仮想モデルの適合度を1と考えるとやはりセグメンテーションを行なった方が良い適合を示しているといえよう。

次に、時間価値を基準としたセグメントごとにRPデータを用いて未知パラメータを推定した結果が、表-7である。この結果、SPデータによるセグメンテーションの性格が反映されたパラメータが各セグメントにおいて推定されていることが分かる。また列車定数項を基準としたセグメンテーションにおけるキャプティブ層は、RPデータにおいてもほぼキャプティブ層であった。これらのことより、SPデータを用いたセグメンテーションが実際の行動における「嗜好」をほぼ正しく表していると言え、その信頼性が示されたと思われる。

(3)社会経済属性と求められたセグメントの関係

さて多項ロジットモデルを用いて3.(3)と同様に社会経済属性と求められたセグメントの関係を求めた結果、ケーススタディー2においても社会経済属性によるセグメントの再現率は低かった。また、フィッティッド・セグメントの有効性も3.(3)と同様に認められなかった。

6. 得られた知見と今後の展望

個人の社会経済属性から求められる従来のセグメントは、予想していたとおり個人の嗜好の違いをあまり反映したものではなく、モデルの適合度もそれほど向上しなかった。

一方、個人パラメータを用いて得られたセグメントは、前述したとおりセグメント間の差異を非常に良く表しており、モデルの適合度も向上した。また求められたセグメントごとにRPデータを用いてパラメータを推定したところ、もとのセグメントと同じ傾向のパラメータが推定された。このことより、SPデータを用いて個人ごとに選好パラメータを推定し、その類似性に基づいて選好構造の似たものをクラスタリングするマーケット・セグメンテーションの信頼性・有効性を示すことができた。

ところが、個人の社会経済属性を用いて、個人パラメータによってクラスタリングされたセグメントを再現することは、ケーススタディー1、ケーススタディー2のどちらの場合でも困難であった。この理由としては、セグメントを再現するのに適した手法がないこ

表-7 時間価値によるセグメンテーション (RP)

	pooled	High	Low	Negative
Rail constant	0.463	0.446	0.704	0.648
Cost per person	-0.274	-0.109	-0.398	-0.391
Line-haul time	-0.313	-0.820	-0.241	-0.0611
Value of time	1.14	7.52	0.606	0.156
N	228	40	138	50

と、個人の選好構造を社会経済属性を用いて分類すること自体に若干無理があること等が考えられる。

しかし、個人の社会経済属性を用いて再現されたセグメントは需要予測には有効でないが、個人パラメータから求められた本来のセグメントは十分信頼できることから、これを利用した需要予測は考えることができる。たとえば、短期需要予測において予測対象の構成にほぼ変化がないとみなせる局面では、アンケートの被験者を選ぶ際、需要予測の対象区域からランダムサンプリングを行なうことによって、本研究の手法によってえられたセグメントの構成比を全対象区域に割り当て、有効なセグメンテーションを活かした需要予測が行なえる。

今後の課題としては、社会経済属性とセグメントの関係を求める分析手法の開発や個人モデルの改善を行なう一方、今回提案した有効なセグメンテーションの実務的な適用が挙げられる。

参考文献

- 1) 片平秀貴：離散型選択モデルと選好の異質性，経済学論文集，Vol.3，No.50，pp.31～45，1987.10.
- 2) 森川高行：ステイティッド・プリファレンス・データの交通需要予測モデルへの適用に関する整理と展望，土木学会論文集，No.413，1990.1.
- 3) 片平秀貴：マーケティング・サイエンス，東京大学出版会，1987.4.
- 4) Daganzo, C. : *Multinomial Probit, The Theory and Its Application to Demand Forecasting*, Academic Press, 1979.
- 5) Chapman, R. G. and Staelin, R. : Exploiting Rank Ordered Choice Set Data within the Stochastic Utility Model, *Journal of Marketing Research*, Vol.19, pp.288～301,1982.
- 6) 小川孔輔：「コンジョイント尺度」を与える最尤推定量について，経営志林，Vol.18，pp.37～52，1981.
- 7) 片平秀貴：多属性消費者選択モデル，経済学論文集，Vol.2，No.50，pp.2～18，1984.7.
- 8) Ogawa, K. : An Approach to Simultaneous Estimation and Segmentation in Conjoint Analysis, *Marketing Science*, Vol.6, No.1, pp.66～81, 1987.
- 9) 湯沢・須田・高田：コンジョイント分析の交通機関選択モデルへの適用に関する諸問題，土木学会論文集，No.419，pp.51～60，1990.12.