

人の集まり方の確率的見方

高見 章・田中孝文・中村 理

大都市に人が集まったり、農村の人口が減少したり、都市内部においてある一部の地域だけ人がかたまっているなど、人は偏って分布するように見える。それではなぜ人は集まる傾向にあるのか、あるいは、どう集まる傾向にあるのか、どんなときに集まる傾向にあるのか。

この問題は、人に限らず工場や商店等がどうして集まるか、いろいろな分野で研究が重ねられている。その中では“利益を得るために”という理由で説明するものももっとも多いと思われる。いわゆる集積の利益はそのうちの1つであって、たしかにそれは現実をよく説明しているし説得力もあるようである。しかしながら、すべての人の動きをそのような経済的理由によって説明してしまうことはできない。

ちょっと身の回りを見ても「なんとかして、もうけよう」とか「どこにどう動いたらいちばんもうかるだろうか」などと考えながら動くことはあまり多くなくて、だいたい経済的要因以外のなんらかの理由で動くのではないだろうか。すると、人が都市に集まったり、都市のある地域にかたまったりする現象も“利益を得るために”という理由以外のなんらかの理由でそうなるのだ、と考えてもよいのではないかと思われる。

もちろん、集積の利益とか規模の経済といわれるように、人が大勢集まることは大きな利益を得ることになるのだが、それでは都市に住むひとりひとりが本当にそれを意識して集まっているのか

ということを考えると、かならずしもそうでないのではないか。教育や仕事の関係、住宅の都合など、ありとあらゆる原因で人は動き、経済的要因はそのうちからあらわれた1つの側面であろう。

本稿は以上のような考えから、経済的要因にとられずに、人の集まり方や動き方について、いくつかのモデルを紹介する。

1. 都市人口規模の分布

——ランクサイズルール——

都市に人が集まるのはなぜかという問題は一応置いて、人口の多い少ないが出てくる様子を統計数理的に眺めてみよう、というのが基本的な考え方である。つまり、1,000万都市東京が1つだけあり、2万人くらいの町は全国にたくさんあるというような都市人口の分布の様子を、統計数理的にどういう視角から見たらよいのか、という大ざっぱな議論をしてみようというわけである。

さて、ランクサイズルール(順位法則)とは、人口の多い順に都市を並べたときに、ある都市の順位を R 、その人口を P としたときに

$$PR^{\alpha} = \text{一定} \quad (\alpha \text{ は定数}) \quad (1)$$

が成り立つというものである。Zipfの法則ともいわれており、都市人口の分布だけでなく、他のいろいろな分野で成り立つらしいことが経験的にいわれている。

たとえば、ある小説に出てくる単語の頻度、ある科学雑誌に論文をのせた人の論文数、地震の規

模とその発生回数(つまり大きい地震ほど数が少ない), 日本人の姓の数など上げればきりのないところである。

ところで, (1)式の表現は, はじめに述べた問題をあつかうにはあまり適切ではない。そこで, 確率論の土俵に持ち込むために, ある属性の度数が i であるものの個数が $f(i)$ であるという表現をとることにする。たとえば, ある小説に i 回登場する単語の数(種類)は $f(i)$, マグニチュード i の地震の起こる回数は $f(i)$ 回, というようになる。この表現を用いると, 人口 i 万人の都市が $f(i)$ 個あるとして都市人口の分布が記述されるのである。そうすると, i 人の人口を持つ都市の順位(多いほうからの) R_i は

$$R_i = \sum_{j=i}^{\infty} f(j) \quad (2)$$

となる。これを用いると Zipf の法則は

$$\sum_{j=i}^{\infty} f(j) = \text{const.} \cdot i^{-(1/\alpha)} \quad (1)'$$

と変形され, これを満たす $f(i)$ を考えればよいことになる。その有力な候補の 1 つは

$$f(i) \sim c \cdot i^{-(\rho+1)} \quad c: \text{定数}, \rho: \text{パラメータ} \quad (3)$$

のパレート分布であるといわれる。こうすると, 都市人口の分布を, 小説の中に登場してくる単語の度数分布や所得の分布, 日本人の姓の分布などと関連づけることができ, 結局, パレート分布をみちびく確率モデルの構成をどうするのかという問題に置き換えることができるわけである。

つまり, ランクサイズルールというのは, 確率的には, ある確率変数にある確率分布を仮定したときに, その順序統計量の分布を対象としたもので, 順序統計量の期待値が順位の関数になることを示したにすぎないわけだから, もっともあてはまりの程度がよいと思われるパレート分布の成立機構を都市人口のタームで考えていくのが, 問題設定として適当だと思われるわけである。現に, 今日までの研究では, Yule 以来, Champarnown, Simon, Mandelkrot, Vining, Hill らがこの問

題に取り組んでいる。ここではその中で Simon のモデルを簡単に紹介して, 数理的な問題点に触れておくことにしよう。Simon のモデルの仮定はつぎのとおりである。

1) 総人口 k 人が n_k 個の都市にいま分布しているとして, つぎの $(k+1)$ 人目の人が人口 i 人の都市にあらわれる確率は, $i \cdot f(i, k)$ に比例する。ここに, $f(i, k)$ は総人口 k 人のときの人口 i 人の都市の数をあらわすから, $i \cdot f(i, k)$ は人口 i 人の都市の人口の合計である。つまり, つぎの 1 人が出現する確率は人口に比例するという仮定である。

2) $(k+1)$ 人目がいままでも都市でなかったところにあられ, そこが都市になる確率は一定値 α である。この仮定は, 新しい都市が生まれる(都市以外のところで人口が増加して“都市”になる)確率が一定であるというものであるが, わが国のような場合, その確率が総人口 k によらないというのは疑問であるので, 後に α が k の減少関数であるとして修正を加えることにしよう。

仮定 1) 2) とともに, ある単位の人口をひとまとめにして取りあつかうことができるので, $(k+1)$ 人というのを $(k+1)$ 万人というように置き換えて考えれば現実的であろう。

さて, Simon のモデルはつぎのようになる。

$$f(i, k+1) - f(i, k) = K(k) \{ (i-1) \cdot f(i-1, k) - i \cdot f(i, k) \} \quad (4)$$

$$f(1, k+1) - f(1, k) = \alpha - K(k) f(1, k) \quad (5)$$

$$\text{制約条件} \quad K = \sum_{i=1}^k i \cdot f(i, k)$$

である。(4)の意味するところは, 総人口 k 万人から $(k+1)$ 万人に人口が増加したときに, 人口 i 万人の都市の増減がどうなるかということ, 増加する前に $(i-1)$ 万人であった都市に 1 万人が入るときには 1 個増え, i 万人であった都市に 1 万人が入る場合には 1 個減少するというものである。(5)は, (4)と同じように考えて, 最小人口の都市(この例では 1 万人)の増減がどうなるか, ということを記述したものである。

これを解けばいいのだが、そのためにはつぎのような意味で「安定状態」にあると仮定する必要がある。(安定状態については参考文献[7]参照)

$$3) f(i, k+1)/f(i, k) = k+1/k \quad (6)$$

以上のモデルの解はつぎのようになる。

$$\alpha = \frac{n_k}{k} \left(= \sum_i f(i, k) / \sum_i i \cdot f(i, k) \right) \quad (7)$$

$$f(1, k) = n_k / \{2 - \alpha\} (= k\alpha / (2 - \alpha)) \quad (8)$$

$$\frac{f(i, k)}{f(i-1, k)} = \frac{(1-\alpha)(i-1)}{1+(1-\alpha)i} = \frac{1-i}{\rho+i}$$

$$\left(\rho = \frac{1}{1-\alpha} \right) \quad (9)$$

こうして、人口 i 万人の都市の数 $f(i, k)$ は

$$f(i, k) = (1+\rho)f(1, k)B(i, \rho+1)$$

(B はベータ関数) (10)

となる。 i がある程度以上大きいとしてスターリング近似を行なうと

$$f(i, k) \sim c \cdot i^{-(\rho+1)} \quad c = (1+\rho)f(1, k)$$

$$\Gamma(\rho+1) \quad (11)$$

とあらわされる。人口 P 人の都市の順位 R_P は、人口 R_P 人を越える人口を持つ都市の数と同じと考え

$$R_P = c \cdot \sum_{i=p}^{\infty} f(i, k) = C_0 / P^\rho \quad (c_0 \text{ 定数}) \quad (12)$$

となり、これは Zipf の法則 ($PR^\alpha = \text{一定}$) と同じものである。ここで、とくに $\alpha \equiv 0$ (つまり新しい都市のできる確率が小さい) とすると $\rho \equiv 1$ となり

$$PR = \text{一定} \quad (13)$$

という Auerbach の法則をみちびける。

さて、先に述べたように、新しい都市の出現する確率 α が総人口 k によらず一定であるという仮定を修正してみよう。その前に、仮定3)の安定状態について考察を加えなければならない。

安定状態とは、人口 k を大きくしても、人口 i 人の都市数あるいは人口 i 人の都市に属する人口の合計の、全体に対する割合が一定であるような成長状態をいう。人口 k 人のときの人口の異なる都市数 (つまり都市の種類数) を n_k とすると

$$\frac{i \cdot f(i, k+1)}{k+1} = \frac{i \cdot f(i, k)}{k} \quad (\text{人口について}) \quad (14)$$

$$\frac{f(i, k+1)}{n_{k+1}} = \frac{f(i, k)}{n_k} \quad (\text{都市について}) \quad (15)$$

のどちらかが、任意の k について成り立つことである。両式はよく似ているが、それがまったく同じ条件であるのは n_k が k と比例関係にあるとき、つまり $n_k/k = \text{const.}$ のときであり、これは $\alpha = \text{const.}$ の Simon の仮定2) である。では、もし α が k の関数であるときはどうなるか。モデルは

$$f(i, k+1) - f(i, k) = K(k) \{ (i-1)f(i-1, k) - if(i, k) \} \quad (16)$$

$$f(1, k+1) - f(1, k) = \alpha(k) - K(k)f(1, k) \quad (17)$$

$$k = \sum_{i=1}^{\infty} i \cdot f(i, k)$$

$$\alpha(k) = \frac{dn_k}{dk}$$

となり、安定状態の条件として(15)をとると解は

$$\frac{f(i, k)}{f(i-1, k)} = i-1 / \left(\frac{k}{1-\alpha(k)} \cdot \frac{n_{k+1}-n_k}{n_k} + i \right) \quad (18)$$

である。ここで(18)をよく見ると、右辺に k が入っており、左辺は k の変化によって変わってしまう。それを防ぐには

$$\frac{k}{1-\alpha(k)} \cdot \frac{n_{k+1}-n_k}{n_k} = \rho(\text{const.}) \quad (19)$$

でなければならない。これが成り立てば

$$f(i, k) = (\rho+1)f(1, k)B(i, \rho+1) \sim ci^{-(1+\rho)}$$

となるのである。(19)を満足するような α は $\alpha = \text{const.}$ はもちろんであるが、他に Mandelbrot のように、

$$n_k = ck^{\rho-1} \quad (20)$$

でも $\rho \equiv \text{const.}$ になる。したがって

$$\alpha(k) = c_0 \cdot k^{\rho-1} \quad (\rho < 1) \quad (21)$$

もその1つである。

最後に Simon のモデルの検証を行なってみよう。表1は日本の DID 地区人口の分布、表2は筆者の学科の同窓会名簿による姓の頻度の分布で

ある。

表 1 DID 地区人口による Simon モデルのあてはめ

[$K=5691$, $n_k=2225$, $\alpha=0.39$, $\rho=1.64$, $f(1,k)=1382$]

人口数(千人台)	5～9	10～14	15～19	20～24	25～29	30～34	35～39	40～44	45～49
Observed	348	115	84	63	40	24	22	17	14
Expected	455	151	75	45	30	21	16	12	10

50～54	55～59	60～64	65～69	70～74	75～79	80～84	85～89	90～94	95～99	100～
17	11	11	4	1	5	6	12	6	6	106
8	6	5	5	4	3	3	3	2	2	53

2. 分布交通と住宅分布

本節では、ランダムという概念を用いて人の集まり方や動き方を説明するモデルを紹介する。基本的な考え方は、ものをランダムに分布させるというものであり、そこに制限をいろいろ加えることによって、偏りのある分布をみちびこうとするものである。

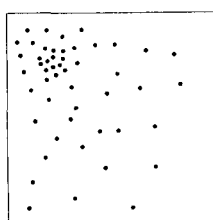
2-1 “集まっている”とはなにか

モデルの紹介の前に、どういう状態だと集まったと判断できるのかを考えてみよう。なぜなら、人がちょっとかたまっている、もしかすると、それは偶然にその偏った状態が起こったのかもしれないからである。それを直観的に確かめてみる。図1の3つのうち、左図はたぶん集まっている状態であり、逆に右図はだれが見てもちらばっている状態であるといえよう。ところがまん中の図となると問題で、集まっていないといえそう見えるし、中央部が少なくても右上に集まっているようにも見える。実は、左はわざと集めて点を打ったもの、右は5×5にわけて各マス目に2つずつ入るように点を打ったもの、まん中は100×100にわけて縦横それぞれ乱数を発生させて点を打ったものである。つまり、まん中はなんの制限もなしにデタラメに点を打ったときの1つの状態を示したものである。だから、このような分布は、集まっているとも集まっていないとも判定することはできない。ちょっと見ると、点がかたまっているし、なにもないところも広いがこの程度の偏りでは集まっているといっけいではないのである。

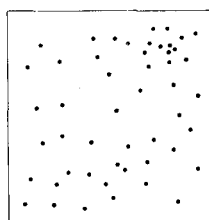
さて、これを都市に人口をランダムに配るといふ例にあてはめてみよう。このときには各都市を1つのマス目と見て、そこにデタラメに人を配る

表 2 学科名簿による Simon モデルのあてはめ

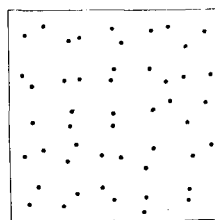
人数(人)	1	2	3	4	5	6	7	8	9
Observed	321	44	18	7	1	3	3	1	1
Expected	312	56	17	7	3	2	1	1	0



集まった分布



ランダムな分布



ちらばった分布

図 1

という考え方をすればよい。かりに、 N 個の都市に M 人の人口を配るとすれば、 $M/N=\lambda$ 、つまり1都市あたりの平均人口が λ 人であるとする、人口 x 人の都市の数が

$$N \times p(x) = N \cdot e^{-\lambda} \lambda^x / x! \quad (22)$$

にしたがうことになる*。もしある地域における人口分布が、(22)式にしたがっていないと判定できなければ、その人口は集まる傾向にあるとかちらばる傾向にあるとかいう判断もできなくなる。(11)の分布がどのくらい偏っているか100個の都

* 人がデタラメに各都市に属するという事は、どの都市にも等しい確率で属すること(等確率の仮定)と、他の人がどこに属するかに関係なくある人が属すること(独立性の仮定)の2つの仮定が必要である。したがって、 M 人のうちある都市に x 人属する確率は、二項分布

$$p(x) = {}_M C_x (1/N)^x (1-1/N)^{M-x}$$

になる。ここで、 M , N が大きいたして $M/N=\lambda$ と置くと

$$p(x) = e^{-\lambda} \lambda^x / x!$$

というポアソン分布が得られる。

表 3 1,000万人を100個の都市にランダムに分布させたとき

人口(万人台)	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19~
都 市 数	1	1	2	3	4	6	9	11	12	12	11	9	7	5	3	2	1	1	0

市に1,000万人を分布させるといって例で考えてみよう。都市というものを形成しなければならないから、ある程度の人口がまとまっていることが必要で、かりにその最小単位を1万人として、1万人のかたまりをそれぞれの都市に配るというように考える。したがって100個のマス目に1,000個の点を落とすことになる。(11)式において $N=100$, $\lambda=10$ を代入すると、期待される分布は表3のようになる。表3によれば、人口5万人以下の都市が10個を上回り、人口15万以上の都市も7個くらいできるのが普通で、この程度の偏りでは、集まる傾向にあると判断してはならないのである。

2-2 起こりやすさと場合の数

2-1においてランダムに点をバラまくことについて考察した。そこでは、なんの制限もなしに点を落とすときどのような状態になるかについて述べ、点を人間とみなして人の集まり方について考えたのである。本節では、その考え方の応用として、いろいろな制約のもとでランダムに点を落とすことを考え、分布交通量と郊外住宅地の分布というまったく異なったモデルを同じような手法でみちびいてみよう。

出発点はどちらも、マス目で区切った平面にランダム点を落としたときに、各マス目の中に落ちる点の数はどうなるのがもっとも起こりやすいかを総費用一定という制限のもとで考えてみようというところにある。すなわち、分布交通モデルにおいては、図2のようなOD表において、総トリップ数 N をランダムに落とし、住宅分布モデルにおいては図3のように、ある鉄道沿線の駅に属する住宅地を考え、そこに N 人の通勤者をランダムに分配する。そして、前者においては総トリップ数 N の総交通費が一定、後者では都心を往復する N 人の総通勤費が一定という制限がそれぞれ課さ

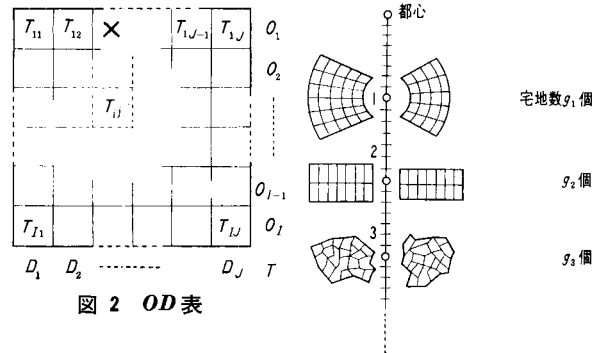


図 2 OD表

図 3 郊外住宅地

れ、ポアソン分布とは異なる分布をみちびくのである。

その前に、もっとも起こりやすいということはどういうことをいうのかについて、若干考察することにしよう。

直観的には、もっとも起こりやすいというのは、そうなる確率をもっとも大きいというふうに考えられる。たとえば、100枚のコインをランダムに投げると、そのうちの50枚が表になる確率をもっとも高いから、50枚が表になるという状態をもっとも起こりやすいといえる。投げる枚数を1,000, 10,000と増やしていくと、このもっとも起こりやすい状態に近い状態が、図4のように非常に大きくなり、いつももっとも起こりやすいことが起こるのである。

もし、コインが4枚だったら半分が表になるのがもっとも起こりやすいが、いつもその状態が起

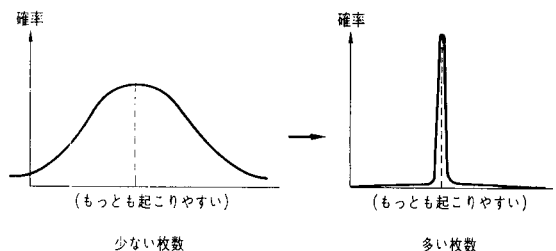


図 4 起こりやすい状態の確率密度の様子

目の合計

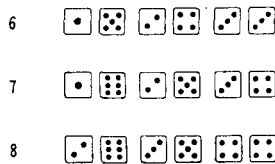


図 5 2つの区別ないサイコロの目

こととはいえない。しかし、コインが1万枚だとちょうど5,000枚表が出ることは少ないにしても、5,000枚くらいが表になる確率は1に近いであろう。したがって、細かいことをいわずに「1万枚のコインを投げると半分が表になる。」という状態がいつも起こるのである。

それでは、もっとも確率が高いというのはどうやって求めるか。まず、確率密度関数を求めてその最大化を行なうことが考えられる。つぎに、場合の数という概念を持ってきて、〈場合の数最大＝確率最大〉であることから、場合の数の最大化を行なうことが考えられる。

たとえば、 N 枚のコインのうち a 枚表になるという確率は

$$p(x) = {}_N C_a \cdot \left(\frac{1}{2}\right)^a \cdot \left(\frac{1}{2}\right)^{N-a} \quad (23)$$

として $p(x)$ の最大化を行なっても、 a 枚表になる場合の数が

$$W = {}_N C_a = N! / (N-a)! a! \quad (24)$$

として、 W の最大化を行なっても、 $a = N/2$ (N が偶数)、 $a = (N \pm 1)/2$ (N が奇数) という答を得ることができる。

ただし、場合の数を用いるときには、その1つ1つ場合が等確率に起こるのだ、ということをしっかり吟味しないと、とんでもない間違いをするおそれがあるので注意しなければならない。

たとえば、2個のまったく区別のつかないサイコロの目の合計を考えてみよう。2～12の11とおりの目の出方のうち、6、7、8、の出る場合の数はそれぞれ3とおらずであるからそれらももっとも起こりやすい、と考えてしまうのは明らかに誤りである。実際には、7は6とおりに、6と8

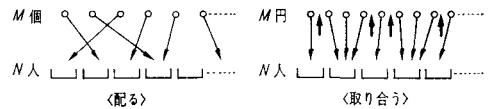
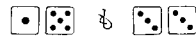


図 6 “配る”型と“取り合う”型

は5とおりが場合の数の正しい計算であって、7がもっとも起こりやすいとしなければならない。このまちがいがなぜ起こったかという点、等確率の事象がなんであるかを誤ったからである。

たしかに



も1つの場合であるが(サイコロは区別できない)



になる確率は



になる確率の2倍であることを見落としてはならないのである(図5)。

サイコロの例ぐらいだとわかりやすいが、所得の分布を考えるとむずかしくなってくる。これは、 M 円を N 人に与えるという試行によって考えられる。考え方としては2つある。1つは M 円を1人1人に配るとするもの、1つは N 人が M 円を取り合うとするものである。この違いは言葉ではわかりにくい、図6によって、配る考え方は N 個の箱に M 円を1円ずつ落とし、取り合う考え方は M 円に切れ目を入れていくと考えると違いが明らかとなる。配るときには人は受動的であり、切るときには人は能動的である。

A	①②③	①②	①②	①③	……	ノ	(配る)
B	×	③	×	②	……	×	
C	×	×	③	×	……	①②③	

〈それぞれ1/27で起こる〉

A	3	2	2	1	1	1	0	0	0	0	(取り合う)
B	0	1	0	2	0	1	1	2	3	0	
C	0	1	1	0	2	1	2	1	0	3	

〈それぞれ1/10で起こる〉

図 7 3円を3人に与える方法

			配る	取り合う
0	0	3	0.11 (3)	0.3(3)
0	1	2	0.66(18)	0.6(6)
1	1	2	0.33 (6)	0.1(1)
			(27)	(10)

図 8 分布の起こる確率

さて、両方の場合の数はどうなるか、簡単のために $N=3$, $M=3$ のときを考える。配るときには M 個のお金は 1 つ 1 つ区別され、取り合うときには区別されないで場合の数が数えられている。前者では 27 とおり、後者では 10 とおりが可能で、それぞれ 1 つ 1 つは $1/27$, $1/10$ の等確率で起こる。

さてどういう分布がどのような確率で起こるかという、図 8 のように配るときと取り合うときとはかなり異なる。ただ、この例ではたまたま $(0, 1, 2)$ という分布がもっとも起こりやすくなっているが、 M と N を大きくするといちじるしく異なることに注意しなければならない。結論だけ述べると、配るときにはポアソン分布型、取り合うときには指数分布型がもっとも起こりやすい状態になるのである。

2-3 分布交通量モデル

1) Wilson の考え方

図 2 の OD 表を考えて、総数 T のトリップをおのおのマス目に分布させる。たとえば、×印のマスに 3 個のトリップが落ちたとすると、ゾーン 1 からゾーン 3 へ 3 個のトリップがあったことになる。各ゾーンの総発生量、総集中量はそれぞれ $O_i D_j$ とさまっており、この点でつぎに述べる Choukroun のモデルと異なる。

さて、場合の数の数え方であるが、 T 個の点をランダムに OD 表のマス目に落とすと考えるのだから、先に述べた 2 つのうち“配る”考え方をとればよい。したがって、ある OD パターン $\{T_{ij}\}$ が起こる場合の数は

$$W = T! / \prod_{ij} T_{ij}! \quad (T_{ij} \text{ は } i \rightarrow j \text{ のトリップ数}) \quad (25)$$

$p_1 q_1$	$p_1 q_2$		$p_1 q_{j-1}$	$p_1 q_j$	p_1
					p_2
$p_i q_j$					
					p_i
					p_I
$p_j q_1$				$p_j q_j$	
	q_1	q_2			q_j

図 9 Choukroun の OD 表

になる。この W を最大にするような $\{T_{ij}\}$ を求めればよいのだが、ある $i \rightarrow j$ のトリップに要するコストを c_{ij} とすると

$$\sum_{ij} c_{ij} T_{ij} = C \text{ (一定)} \quad (26)$$

という制限が課せられる。その他に

$$\sum_j T_{ij} = O_i \quad \sum_i T_{ij} = D_j \quad (27)$$

という 2 つの制限がある。すなわち、求める解は 3 つの制約条件のもとで W を最大にするような $\{T_{ij}\}$ である。これはラグランジュ乗数法で解くと簡単に得られ、結果は

$$T_{ij} = A_i B_j O_i D_j \exp(-\beta c_{ij}) \quad (28)$$

となる。(詳細は参考文献[8])

2) Choukroun の考え方

Wilson のモデルでは、どのマス目に入る確率も等しいと考えられている。しかし、これは現実的でなく、あるゾーン間のトリップが起こりやすいという偏りがあるはずだ、として組み立てられたのが Choukroun の考え方である。

すなわち図 9 のように、OD 表の 1 つ 1 つのマス目に先験的にある確率が与えられており、その確率は発生側の確率 p_i と集中側の確率 q_j の積と表現できるとしている。今度は Wilson のときのような場合の数最大化によって解を得ることはできない。なぜならば、1 つ 1 つの場合は等確率ではないからである。そこで、Choukroun のモデルでは、ある $\{T_{ij}\}$ が起こる確率を求めて、それを最大化するという手法をとる。終トリップ T のうちある OD $i \rightarrow j$ に x 個のトリップが落ちる確率は

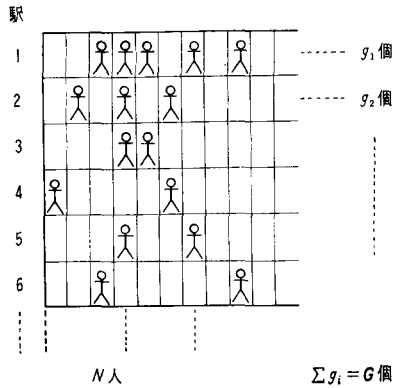


図 10

$$p(x) = r C_x (p_i q_j)^x (1 - p_i q_j)^{T-x} \quad (29)$$

である。したがって、ある OD パターン $\{T_{ij}\}$ の起こる確率は

$$p(T_{ij}) = \frac{T!}{\prod_{ij} T_{ij}!} \prod_{ij} (p_i q_j)^{T_{ij}} \quad (30)$$

であらわされる。制約条件は

$$\sum_i p_i = 1, \sum_j q_j = 1 \quad (31)$$

と、費用制約

$$\sum_{ij} c_{ij} T_{ij} = C$$

である。求める解は

$$T_{ij} = \alpha p_i q_j \exp(-\beta c_{ij}) \quad (32)$$

である。ここでは $O_i (= \sum_j T_{ij})$, $D_j (= \sum_i T_{ij})$ は Wilson のモデルのように制約条件ではなく結果である。

2-4 住宅分布モデル

都市の概念は図 3 であらわされるが、図 10 のような描き方をすれば、OD のときと同じ考え方があつかえる。図 10 の合計 G 個のマスキに N 人をランダムに分布させると考えることは、いままでとまったく異ならないが、住宅分布の場合では図 10 にもあるとおり、1 つの住宅地には 1 つしか入ることを許されない、という制限を考えるのが妥当だろう。むりやり 1 区画に 2 戸建てるとこのような状況はモデルの対象とはしないことにする。さて、このようにして N 人の通勤者が、1 駅に n_1 人、2 駅に n_2 人、 $\dots$, i 駅に n_i 人分布しているときの場合の数を数えるのである。(どのマスキに

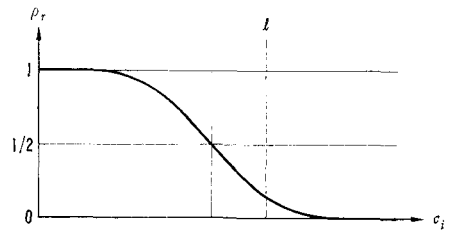


図 11 $p_r = 1 / \{1 + \exp(\alpha + \beta c_i)\}$

入るかはまったく等確率であるという Wilson の考え方を採用) まず、 N 人を $n_1, n_2, \dots, n_i, \dots$ にわける場合の数は

$$W_1 = N! / \prod_i n_i \quad (33)$$

つぎに、ある r 駅に属する g_r 個の住宅地群に n_r 人が属するときの場合の数は順列

$$g_r P_{nr} = g_r! / (g_r - n_r)! \quad (34)$$

であらわせるから 1 駅ずつ考えると

$$W_2 = \prod_i \{g_i! / (g_i - n_i)!\} \quad (35)$$

したがって、全体の場合の数は

$$W = W_1 \times W_2 = N! \prod_i g_i! / n_i! (g_i - n_i)! \quad (36)$$

である。 i 駅までの費用を c_i とあらわすとすると総通勤費一定の制約条件は

$$\sum_i c_i n_i = C \quad (37)$$

である。この条件のもとに W を最大にするような $n_1, n_2, \dots, n_i, \dots$ を求めると

$$n_i = g_i / \{1 + \exp(\alpha + \beta c_i)\} \quad (38)$$

である。密度をあらわす指標 $\rho_i = n_i / g_i$ を用いて

$$\rho_i = 1 / \{1 + \exp(\alpha + \beta c_i)\} \quad (39)$$

という分布がみちびかれる。

人口密度が稀薄で $\rho_r \ll 1$ のときは近似を行なうと

$$n_i = \alpha_i \exp(-\beta c_i) \quad \beta = N/C \quad (40)$$

という式をみちびくことができる。これは分布交通量モデルでみちびいた結果と似ている。その理由は、 g_r 大きいので n_r 人の 1 人がある駅の住宅地を占めたときにも、つぎの 1 人が同じ駅の住宅地に住む確率はほとんど変わらない、というところにある。つまり、ある駅 r に属する確率

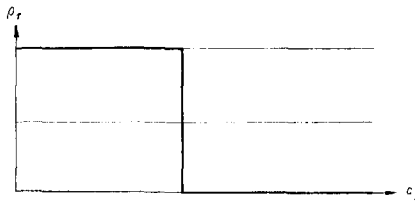


図 12 $\beta \rightarrow \infty$ のとき

$$p(r) = g_r / G \quad (41)$$

が n_r によってほとんど変化しないためである。したがってこのときには、Choukroun のモデルと同じ方法で、式をみちびくことができる。このときには

$$N^2 / C g_r \ll 1 \quad (42)$$

になるが、これは総通勤費 C が非常に大きく、遠くから通う人が多いことを示し、人口密度が稀薄であることと一致する。また、 $\alpha > 0$ になるので図11において l よりも右側を考えたことになり(なぜなら $c_i > 0$ のときのみ意味がある)、その部分の(39)式を(40)式で近似したともいえる。

ところでパラメータ β^{-1} は、通勤者 1 人あたりの平均費用の大小を示すが、 $\beta \rightarrow \infty$ のときの式のグラフは図12のようになる。これは平均費用をできるだけ小さくしながら分布させると、ある c_i より安い地域に全員が集まって飽和した状態をつくればよいことを考えると、妥当な結果である。

参 考 文 献

- 1) Auerbach F. (1954) "Das Gesetz der Bevölkerungskonzentration" *Determinant's Geographische Mitteilungen* vol. 59-I pp. 74~77
- 2) Champernowne D. G. (1953): "A Model of Income Distribution". *Econometrica*, vol. 63, pp. 318-351.
- 3) Choukroun, J. M. (1975) "A General Frame-

work for the Development of Trip-distribution Models" *Regional Science & Urban Economics* vol. 5 pp. 177-202

- 4) Gurney R. W. (1949): *Introduction to Statistical Mechanics*. The McGraw Hill Books Co. Inc.
- 5) Hill B. M. (1970): "Zipf's Law and Prior Distributions for the Composition of a Population." *Journal of the American Statistical Association*, vol. 65, pp. 1220-1232.
- 6) Mandelblat B. (1965) "A Class of Long-Tailed Probability Distributions and the Empirical Distribution of City Size" *Mathematical Explorations in Behavioral Science*
- 7) Simon H. A. (1955): "On a class of skew Distribution Functions". *Biometrika*, vol. 42, pp. 425-440.
- 8) Wilson A. G. : *Entropy in Urban and Regional Models*.
- 9) Yule G. U. (1924) "A Mathematical Theory of Evolution, based on the Conclusions of Dr. J. C. Wills" *Philosophical Transactions of the Royal Society Series B* vol. 213 pp. 21-87
- 10) Zipf G. K. (1949): *Human Behavior and the Principle of Least Effort*. Cambridge Mass. Addison-Wisley.
- 11) 高見 章 "Rank Size Rule 研究"
- 12) 田中孝文 "「群がり」分析の序論的考察"
- 13) 中村 理 "都市現象における場合の数の適用"
((11), (12), (13)は昭和50年度修士論文)

たかみ・あきら	1952年生	日本不動産銀行
たなか・たかぶみ	1952年生	経済企画庁調査局
なかむら・おさむ	1952年生	東京大学都市工学科大学院博士課程