

『計量国語学』アーカイブ

ID	KK300603
種別	研究資料
タイトル	リジット解析 —計数データを用いた言語研究への適用—
Title	Ridit Analysis: Its Application to Linguistic Count Data
著者	石井 正彦
Author	ISHII Masahiko
掲載号	30巻6号
発行日	2016年9月20日
開始ページ	357
終了ページ	377
著作権者	計量国語学会

研究資料

リジット解析

—計数データを用いた言語研究への適用—

石井 正彦 (大阪大学)

要旨

医学・医療統計の分野で用いられることの多い「リジット解析」という統計手法が、言語研究における計数データの分析において、各群あるいは各変数カテゴリーの相対的な大きさやその差を直観的に評価する手法として有用であることを述べる。リジット解析は、本来、群×質的変数の分割表があるとき、変数の順序カテゴリーをリジットという尺度値に変換し、それをもとに各群の平均値（平均リジット）を求める（数量化する）もので、その考え方と計算法はきわめて平易である。小稿では、リジット解析を先行調査研究のデータに適用することによって、本来の質的データの群間比較のほかにも、量的データの群間比較、さらには、2変数の分割表におけるそれぞれの順序カテゴリーの数量化にも適用できることを示し、その（探索的な手法としての）有用性を確認する。

キーワード：リジット解析，平均リジット，数量化，リッカート法，交互平均法，探索的データ解析

1. はじめに

分割表（クロス集計表，度数分布表）にまとめられるような計数データの分析では、各群あるいは各変数（カテゴリー）の度数分布を要約・数量化して、それらの相対的な大きさやその差を直観的に評価しようとすることがある（渡部ほか 1985:145）。言語研究でもこうした作業の重要性は指摘されているが（荻野 1980:13-14）、小稿では、この種の作業において「リジット解析」（「リジット分析」とも）という簡便な手法が有用であることを、いくつかの事例をもって紹介する。

リジット解析 (ridit analysis) は、Bross (1958) によるものとされ（フライス 1973: 111）、主として医学・医療統計の分野で利用されているが、言語研究における利用例は管見にして知らない。なお、「リジット (ridit)」とは、Bross (1958:19) によれば、“Relative to an Identified Distribution” の頭文字に、「プロビット」や「ロジット」と同様に変換値であることを示す接尾辞的要素 “-it” を後接させた造語であるという。

2. リジット解析の考え方と方法

リジット解析は、本来、群×質的変数（順序カテゴリー）の分割表に適用し、群間の比

較を行う統計手法で、上述したように、専ら医学・医療統計の分野で用いられている。その考え方と計算法はきわめて平易で、フライス（1973）、富永（1982）、遠藤・山本（1992）、石村ほか（2003）など、この分野の教科書・参考書でも広く紹介されている。本節では、まず、これらの記述にもとづき、遠藤・山本（1992:98）による「薬効試験の例」を使って、リジット解析の基本的な方法・手順を説明する。この例は、対照薬 A と実験薬 B について薬効試験を行い、それぞれの効果について、「a）著効」から「e）悪化」まで5段階の順序カテゴリーから成る分割表データ（表 1）が得られたとき、A と B の薬効の間にどの程度の差があるかを評価しようというものである。

表 1：薬効試験の例

	対照薬 A	実験薬 B	A + B
a) 著効	5 人	10 人	15 人
b) 有効	10 人	15 人	25 人
c) やや有効	15 人	12 人	27 人
d) 不変	15 人	10 人	25 人
e) 悪化	5 人	3 人	8 人
計	50 人	50 人	100 人

この種の分割表の場合、統計的な有意差の有無を確認するだけであれば標準的な χ^2 検定を行えばよいが、それでは「カテゴリーに含まれる自然の順序づけという決定的な情報が失われてしま」（フライス 1973:110）、群間の相対的な大きさを比較することにならない。

このような時よく行われるのが、カテゴリーの最も低位のものから最も高位のものまで順に等間隔の点数を与え、それを比率尺度とみなして平均値などを求めることにより、群間の比較を行う「リッカート法」¹という方法である。ただし、こうした見かけ上の数量化には、「真実に存在する姿よりもはるかに正確であるかのような印象を与える、どのような点数づけの方法をとるかによって結果が変わる、といった多くの欠点が認められる」という（フライス 1973:110）。

さて、リジット解析では、最初に、比較の基準となる群と順序カテゴリーの並べ方とを決める必要がある。薬効試験の例では、当然、対照薬を基準群とすることになるが、多群の比較などでどの群を基準群とするかが明らかでない場合は、群全体の合計を基準群とみなすことになる。以下では、対照薬 A を基準群、実験薬 B を比較群とする場合の手順を紹介し、合計を基準群とする手順については次節で紹介する。また、順序カテゴリーは、順序を崩さぬ限り、どちらの方向に並べてもよいが、この例では、後述する平均リジットの大きさと薬効の大きさ（程度）とが対応するように、表 1 とは逆に、「e）悪化」から「a）著効」の方へ並べることにする²。

リジット解析の手順は、大きく、リジット変換、平均リジットの算出、平均リジットの

1 「リッカート法」については、西里（2010:10-25）による詳しい紹介と批判がある。

2 遠藤・山本（1992:98）では、「a）著効」から「e）悪化」に並べたリジット解析を紹介している。

有意差検定，という3つのステップから成る．このうち，リジット変換から平均リジットの算出までの計算については，以下のような，表を使って行う簡便な方法が提案されており（表2），Excelなどの表計算ソフトを使っても簡単に行うことができる．なお，上述したように，順序カテゴリーは「e）悪化」から「a）著効」へと並べている．

- (1) 表2の列①に，基準群の度数を置く．
- (2) 列②で，①を半分にする．
- (3) 列③で，①を1行だけ下にずらした上で，その度数を累積する．
- (4) 列④で，②と③を足し合わせる．
- (5) 列⑤で，④を①の合計（ N_A ）で割る．この数値を「リジット」という．
- (6) 列⑥で，⑤に①を掛けて合計する（ T_A ）．
- (7) ⑨で， T_A を N_A で割って，基準群の「平均リジット」が0.5になることを確認する．
- (8) 列⑦に，比較群の度数を置く．
- (9) 列⑧で，⑤に⑦を掛けて合計する（ T_B ）．
- (10) ⑩で， T_B を⑦の合計（ N_B ）で割って，比較群の平均リジットを求める．

表2：リジット変換から平均リジット算出までの計算法

	① 基準群 A	② ①÷2	③ 累積度数	④ ②+③	⑤ ④÷ N_A	⑥ ⑤×①	⑦ 比較群 B	⑧ ⑤×⑦
e) 悪化	5	2.5	0	2.5	0.05	0.25	3	0.15
d) 不変	15	7.5	5	12.5	0.25	3.75	10	2.50
c) やや有効	15	7.5	20	27.5	0.55	8.25	12	6.60
b) 有効	10	5.0	35	40.0	0.80	8.00	15	12.00
a) 著効	5	2.5	45	47.5	0.95	4.75	10	9.50
計	$N_A=50$					$T_A=25.00$	$N_B=50$	$T_B=30.75$
平均リジット						⑨ $T_A \div N_A$		⑩ $T_B \div N_B$
						0.50		0.615

(1)～(5)では，基準群の各カテゴリーの度数（列①）を「リジット」（列⑤）に変換している．リジット（「リジット尺度」とも）とは，基準群について，そのカテゴリーよりも下位のカテゴリーの全度数（列③）と，そのカテゴリーの度数の半分すなわち折半数（列②）を加えたもの（列④）を総度数（ N_A ）に対する割合として求めた「累積相対度数」であり，その分布は基準群の期待的（標準的）な分布とみなし得る³．

3 竹内ほか（1989:458）による「リジット変換」の解説は，以下のとおり．

リジット解析は順序カテゴリーごとに与えられる頻度に対してある変換を行い，この変換値が一様分布に従うことを利用し，検定や推定を行う．（略）たとえば， I 個の順序カテゴリー応答と J 個の群（処理）からなる $I \times J$ の2元分割表の形の観測度数（ N_{ij} ）が与えられたとする．ここに， $i=1, \dots, I$; $j=1, \dots, J$ である．

基準群（基準処理）に対して各順序カテゴリーをリジット尺度（0,1）に変換する．ただし，基準群が与えられないときには，群（処理）の合計を基準群とする．このとき，第 i 順序カテゴリーのリジット R_i は

$$R_i = (N_{i/2} + \sum_{k=1}^{i-1} N_k) / N, \quad N_i = \sum_{j=1}^J N_{ij}$$

で与えられる．ここに， N_i は第 i 順序カテゴリーの観測度数，および N は総度数である．

(6)~(10)では、基準群のリジットを順序カテゴリーの尺度得点として、各群の「平均リジット」(⑨・⑩)を求めている。ある群の平均リジットとは、各カテゴリーのリジットと(その群の)度数との積(列⑥・列⑧)を求め、その総和をその群の総度数($N_A \cdot N_B$)で割った値であり、リジットを各カテゴリーの尺度値(重み)として求めた、それぞれの群の平均得点と言えるものである。

リジットは、平均 1/2、分散 1/12 の一様分布の確率変数であり、0~1 の間の値をとることが知られている(石村ほか 2003:130-131)。平均リジットはこの確率分布の平均値であり、基準群の平均リジットは常に 0.5 となることから、これと比較群の平均リジットとを比べればよい。この例では、比較群の平均リジットは 0.615 (表 2 の⑩)となり、基準群の平均リジット 0.5 (同⑨)より大きい。表 2 では「e」悪化」から「a」著効」に並べたので、比較群(実験薬 B)は基準群(対照薬 A)よりも薬効が大きいことになる。また、平均リジットは確率として解釈できるので、比較群が基準群に比べて薬効が大きくなる見込み(オッズ)は 8 対 5 ($=0.615/(1-0.615)$)であることもわかる(フライス 1973:112-113)。

リジット解析では、3 番目のステップとして、こうした平均リジットの差の有意性を検定する。この場合は、比較群の平均リジットが基準群の平均リジット(0.5)と有意に異なっているかどうかを検定する(基準群ではない 2 群の間の平均リジットの有意差検定については、次節で紹介する)。

いま、基準群の平均リジットを $\bar{R}_A (=0.5)$ とし、比較群の平均リジットを $\bar{R}_B$ 、その総度数を N_B とすると、 $\bar{R}_B - \bar{R}_A$ の分布は、“両群の母集団分布は等しい”という帰無仮説が成立するとき、近似的に平均値 0、標準誤差 $1/\sqrt{12N_B}$ の正規分布に従う(富永 1982:127)。したがって、 $\bar{R}_B - \bar{R}_A$ の有意性は、

$$z_0 = \frac{|\bar{R}_B - \bar{R}_A|}{1/\sqrt{12N_B}} = \sqrt{12N_B} \times |\bar{R}_B - \bar{R}_A|$$

の値を標準正規分布の棄却限界値 z と比べて検定できる(フライス 1973:113-114)。この例の場合、

$$z_0 = \sqrt{12 \times 50} \times |0.615 - 0.5| = 2.817 > z \left(\frac{0.01}{2} \right) = 2.576$$

であるから、有意水準 1 % で帰無仮説を棄却でき、基準群と比較群の平均リジットには有意差が認められる。

以上がリジット解析の手順である。この例では、リジット解析によって、実験薬と対照薬の薬効が平均リジットという形で数量化され、その相対的な大きさと差が直観的に評価し得ることが理解されよう。

3. 質的データの群間比較

このように、リジット解析は、順序カテゴリーをリジットという尺度値(重み)に変換し、それをもとに各群の平均値(平均リジット)を求めて比較するという、「数量化」の一手法であることがわかる。これが、冒頭に述べた「計数データの分析で、各群あるいは各変数(カテゴリー)の相対的な大きさやその差を直観的に評価しようとする」言語研究

にリジット解析が利用できる理由である。そのことを、まずは、質的データの群間比較を行った言語研究にリジット解析を再適用することによって確かめてみよう。

岡崎和夫(1980)は、いわゆる「ラ抜き言葉」の使用がどのような動詞に多くみられる(みられない)かを、411人の中学生を対象にしたアンケート調査(1979年)によって、調べている。アンケートでは、30の動詞(「載せる」「かける」のみ文例が2つ)について、「一レルの言い方」の文例と「一ラレルの言い方」の文例とを示し、「気のおけない知人・友人たちなどとのふだんの会話のなかで」どちらを使うかを、以下の選択肢から選ばせている(以降の記述では、岡崎(1980)の「一レルの言い方」を「ラ抜き」、「一ラレルの言い方」を「標準形」と記す)。

- A. ラ抜きのみを用いる
- B. ともに用いるがラ抜きの方をよく用いる
- C. ラ抜き・標準形ともに同程度に用いる
- D. ともに用いるが標準形の方をよく用いる
- E. 標準形のみを用いる

このA～Eの選択肢は、ラ抜き使用の程度を5段階で一次的に示す順序カテゴリーであり、調査の結果は、30の動詞を群として、それぞれの各カテゴリーへの回答者数を度数とする(30群×5カテゴリーの)2元分割表(表3の太線の枠内。ただし、群を行に、順序カテゴリーを列に配置)にまとめられる。岡崎は、まず、これら30動詞のラ抜き使用の程度を、以下のような3つの基準によって、それぞれ順位づける。

基準1：以下の順に並べる。

- (1) A の場合のみで全回答者数の過半数 206 を越えるもの(をその順に)……1語
- (2) A+B の場合で同じく 206 を越えるもの(をその順に)……5語
- (3) A+B+C の場合で同じく 206 を越えるもの(をその順に)……3語
- (4) A+B+C+D の場合で同じく 206 を越えるもの(をその順に)……4語
- (5) A+B+C+D の場合でも 206 を越えないもの……上以外のすべて

基準2：Aの降順に並べる。

基準3：Eの昇順に並べる。

表3には、太枠の右側に、各動詞のこれら3つの基準による順位が示してある。岡崎は、これらの「順位のあいだには若干のズレが認められる」としながらも、このズレは、表3のように全体を4つのグループに分ければ、それぞれのグループ内での変動にとどまるとして、30の動詞をこれら4グループに分類し、それぞれを次のように特徴づける。

第1グループ：いずれも語幹⁴が1音節のものであり、その中では、上一段活用のも

4 岡崎(1980:69)では、学校文法に従い「語幹と活用語尾との別のないもの」と記しているが、ここでは、一段動詞を母音活用動詞ととらえて語幹の音節数を数え、記述を明確化した。

のが上位に、次いでカ変、下一段のものが下位に位置している。

第2グループ: 「出る」を除いてすべて上一段の動詞で、語幹が2音節以上のもの。

第3グループ: 10語のうち8語までが下一段の動詞で、いずれも語幹が2音節のもの。

第4グループ: すべて下一段の動詞で、とくに下位には語幹の音節数の多いものが集中している。

表3: 30動詞の素データ(度数分布)と順位

	動 詞	A	B	C	D	E	順 位		
							基準 1	基準 2	基準 3
第1グループ	煮る	245	63	19	28	56	1	1	4
	見る	191	107	42	39	32	2	2	1
	着る	148	116	69	43	35	3	5	2
	射る	185	75	27	43	81	4	3	6
	来る	155	104	69	36	47	5	4	3
	寝る	115	101	65	72	58	6	6	5
第2グループ	出る	79	91	76	47	118	7	9	9
	起きる	83	92	65	72	99	8	8	7
	降りる	71	72	72	88	108	9	11	8
	借りる	72	87	41	88	123	10	10	10
	生き延びる	96	72	30	71	142	11	7	11
第3グループ	投げる	64	66	38	95	148	12	13	12
	食べる	51	58	55	89	158	13	14	13
	逃げる	32	56	11	87	225	14	17	14
	染める	43	27	40	57	244	15	15	15
	生きる	28	36	24	73	250	16	20	16
	居る	27	34	31	65	254	17	21	17
	攻める	65	28	21	40	257	18	12	18
	載せる①	36	28	20	68	259	19	16	19
	受ける	31	27	12	64	277	20	18	20
	載せる②	31	24	11	66	279	21	18	21
第4グループ	越える	23	27	22	38	301	22	22	22
	かける①	22	12	12	51	314	23	23	23
	かける②	12	15	10	52	322	24	26	24
	考える	12	19	10	21	349	25	26	25
	覚える	8	7	6	31	359	26	28	26
	伝える	19	11	4	16	361	27	25	27
	入れる	21	12	3	12	363	28	24	28
	比べる	8	9	8	17	369	29	28	29
	忘れる	2	4	0	11	394	30	30	30
	計	1975	1480	913	1580	6382			

岡崎は、この結果から、語幹の音節数が少ない(語長が短い)動詞ほど抜きで用いられやすく、同じ音節数の動詞では語幹末尾の音節がエ列(下一段活用)よりイ列(上一段活用)の方が用いられやすいという傾向を見出している。岡崎のこの方法は、3つの基準

による順位のズレという誤差的な現象を、4つのグループにまとめることの根拠として使うという点で巧みであり、これによってラ抜き使用の程度が、これらのグループごとに共通する動詞のカテゴリカルな特徴（語幹の長さや末尾音節）によってもたらされることが強く示唆される。

ただし、この集計方法では、そうした動詞グループごとの特徴・異同はとらえられても、グループごとのラ抜き使用の程度やその隔たりがどれほどのものか、また、グループ内の個々の動詞のラ抜き使用の程度が（順位だけでなく）どれほどのものであるか、ということを知ることはできない。そもそも、個々の動詞のラ抜き使用の程度の相対的な大きさがとらえられなければ、これら 30 動詞が岡崎の言うように 4 グループに分かれることを確認することもできない。

さて、上述したように、先の選択肢 A～E は一次元的な順序カテゴリーであるから、リジット解析が適用できる。その際、まずは基準群を決めてリジット変換をする必要があるが、この例のように基準となる群（動詞）が明確でない場合は、群全体の合計を基準とする方法が推奨されている。表 4 は、表 3 最下行の「合計」欄の数値列を基準群とみなし、表 2 と同じ手順で、そのリジット（列⑤）を求め、さらに、例として動詞「見る」および「煮る」の平均リジット（⑫⑬）を求めたものである。なお、順序カテゴリー（列①）は、平均リジットの大きさとラ抜き使用の程度とが対応するように、選択肢 E → A の順に並べている。

表 4：合計を基準としたリジット変換と平均リジット算出までの計算法

選択肢	①基準群 (合計)	② ①÷2	③ 累積度数	④ ②+③	⑤ ④÷N	⑥ ⑤×①	⑦ 「見る」	⑧ ⑤×⑦	⑨ 「煮る」	⑩ ⑤×⑨
E	6382	3191	0	3191	0.259	1651.7	32	8.3	56	14.5
D	1580	790	6382	7172	0.582	919.0	39	22.7	28	16.3
C	913	456.5	7962	8418.5	0.683	623.4	42	28.7	19	13.0
B	1480	740	8875	9615	0.780	1154.1	107	83.4	63	49.1
A	1975	987.5	10355	11342.5	0.920	1816.8	191	175.7	245	225.4
計	N=12330					T=6165	N _T =411	T _T =318.8	N _T =411	T _T =318.3
平均リジット						⑪ T÷N		⑫ T _T ÷N _T		⑬ T _T ÷N _T
						0.50		0.776		0.774

表 5 は、このようにして求めた 30 動詞の平均リジットを降順に並べたものである。上述したように、表 4 では選択肢 E → A の順に並べたので、平均リジットの値が大きい動詞ほどラ抜き使用の程度が大きいことを表している。これにより、各動詞のラ抜き使用の程度が数量化され、順位だけでなく⁵、その相対的な大きさもわかるようになる。前述したように、平均リジットは確率と解釈できるので、たとえば、最上位の「見る」と最下位の「忘れる」とのオッズ比は $9.087 \left((0.776/(1-0.776))/(0.276/(1-0.276)) \right)$ となり、「見る」

5 順位については、表 3 の基準 A による順位とはほぼ同様の結果となった。表 5 では、表 3 の 3 か所で上下に隣り合う動詞が入れ替わり（1・2 位、4・5 位、7・8 位）、1 か所で連続する 3 語の順位が逆転した（26・27・28 位）ほかは、「攻める」が 18 位から 15 位に変わっただけで、4 つのグループを越境する変動はない。

は「忘れる」の約 9 倍ラ抜きで使われやすい、などということもわかる。

表 5：30 動詞の平均リジット

順位	動 詞	平均リジット	有意差なしの範囲
1	見る	0.776	}
2	煮る	0.774	
3	着る	0.749	
4	来る	0.739	}
5	射る	0.713	
6	寝る	0.695	}
7	起きる	0.633	
8	出る	0.617	
9	降りる	0.608	}
10	借りる	0.596	
11	生き延びる	0.591	
12	投げる	0.559	}
13	食べる	0.541	
14	逃げる	0.461	}
15	攻める	0.452	
16	染める	0.448	
17	生きる	0.432	}
18	居る	0.428	
19	載せる①	0.426	
20	受ける	0.406	}
21	載せる②	0.402	
22	越える	0.383	
23	かける①	0.362	}
24	かける②	0.348	
25	考える	0.329	}
26	入れる	0.320	
27	伝える	0.320	
28	覚える	0.311	}
29	比べる	0.305	
30	忘れる	0.276	}

ただし、各動詞の平均リジットを比べる場合には、その差について検定を行う必要がある。この場合のように、合計の数値列を基準群とみなすとき、すなわち、同一対照（合計）に対して異なる 2 群の平均リジットがあるとき、その有意差検定は以下の方法により行うことができる。

いま、ある群の平均リジットを $\bar{R}_1$ 、別の群の平均リジットを $\bar{R}_2$ 、それぞれの総度数を N_1 、 N_2 とすると、 $\bar{R}_2 - \bar{R}_1$ の分布は、“両群の母集団分布は等しい”という帰無仮説が成立するとき、近似的に平均値 0、標準誤差 $\sqrt{\frac{1}{12} \left(\frac{1}{N_1} + \frac{1}{N_2} \right)}$ の正規分布に従う（富永 1982: 129）。

したがって、 $\bar{R}_2 - \bar{R}_1$ の有意性は、

$$z_0 = |\bar{R}_2 - \bar{R}_1| / \sqrt{\frac{1}{12} \left(\frac{1}{N_1} + \frac{1}{N_2} \right)}$$

の値を標準正規分布の棄却限界値 z と比べて検定すればよい. 表 4 に例としてあげた「見る」と「煮る」を例にとれば,

$$z_0 = |0.776 - 0.774| / \sqrt{\frac{1}{12} \left(\frac{1}{411} + \frac{1}{411} \right)} = 0.064 < z \left(\frac{0.05}{2} \right) = 1.960$$

であるから, 帰無仮説は棄却できず, この 2 つの動詞の平均リジットの差は統計的に有意であるとはいえない. 表 5 の最右列には, このようにして 30 動詞のすべての組み合わせについて検定を行い, 危険率 5 % 水準で有意差が認められなかった動詞の範囲を (リーグ戦表のように示すとスペースをとるので) 同じ “{” の中に示した. これによって, 同じ “{” のの中に入っていない動詞であれば, その平均リジットには有意差があることがわかる. たとえば, 第 1 位「見る」は, 「煮る」「着る」「来る」との “{” の中にしか入っていないので, これら以外の動詞との間には有意差があり, 第 3 位「着る」は 2 つの “{” のの中に, 第 10 位「借りる」は 3 つの “{” のの中に入っているので, それぞれ, いずれの “{” のの中にも入っていない動詞との間には有意差がある, ということになる.

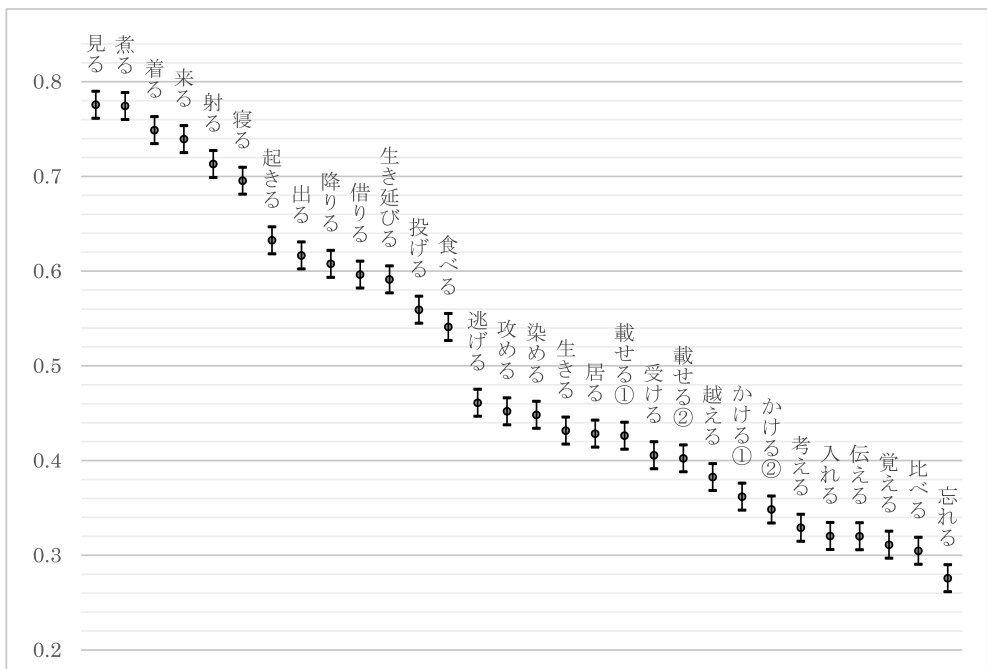


図 1: 30 動詞の平均リジット ± 標準誤差

とはいえ, 表 5 のような平均リジットの数値だけでは, 各動詞のラ抜き使用の程度の相対的な大きさやその隔たり具合がわかりにくいので, これをグラフ化してみよう. 図 1 は, 30 動詞の平均リジットを “●” で示し, その上下に (基準値 0.5 に対する) 標準誤差

($1/\sqrt{12 \times 411}$) の幅をもつ、いわゆるエラーバー⁶を付けて表示したものである（遠藤・山本 1992:100）。これによって、30 もの動詞のラ抜き使用の程度の相対的な大きさとその差がより把握しやすくなる。

図 1 を見ると、岡崎の言うように、30 動詞を 4 つのグループに分けることはほぼ妥当であることがわかる。すなわち、最も左上にある「見る」～「寝る」の 6 動詞は岡崎の第 1 グループと、その右下に並ぶ「起きる」～「生き延びる」の 5 動詞は第 2 グループと完全に一致し、しかも、両グループの間にはグラフ上でも（平均リジットの）明確な隔たりがある（表 5 最右列から有意差も確認できる）。さらにその右下には残りの 19 動詞が、岡崎の第 3 グループの 10 動詞を先（左上）に、第 4 グループの 9 動詞を後（右下）にして並んでいる。

ただし、図 1 では、岡崎の第 3 グループの上位 2 動詞「投げる」「食べる」は、むしろ第 2 グループに近づいていること、また、第 3 グループと第 4 グループとの間には大きな隔たりがないことなど、岡崎の結果とは異なる、あるいは、岡崎の結果からはわからなかったことも示されている。

いずれにしても重要なことは、リジット解析を行うことで、30 もの動詞の「ラ抜き使用の程度」（の相対的な大きさ）が数量化でき、それによって、岡崎の 4 グループ説を検証するとともに、グループ間の、そして、グループ内の個々の動詞の、ラ抜き使用の程度やその差を評価することが可能になるということである。また、その結果として、たとえば、「投げる」「食べる」が同類の下一段動詞より異質な上一段動詞のグループに近づいているのはなぜかといった、新たな研究課題や仮説の設定が可能になるということも重要である。

4. 量的データの群間比較

リジット解析は、群×量的変数の分割表、すなわち、多群の一変量度数分布表を作り、群間の比較を行うような場合にも適用できる⁷。ただ、連続的な量的データの群間比較は、各群の平均値や標準偏差をもって行うことが一般的であり、多くの場合、リジット解析を行うまでもない。しかし、第三者による調査の報告などでは、度数分布表のみが示されていて正確な平均値・標準偏差を求めることが難しいことも多く、また、以下で例とするような語の使用頻度（使用率）分布の場合などは、いわゆる L 字型分布（ベキ分布）をなしてしまって、平均値を（代表値として）比較することにあまり意味がないことも多い。このような場合には、度数分布をリジットという尺度値に変換し、平均リジットによって群間の比較を行うリジット解析が有用である。

国立国語研究所が 1956 年に行った「雑誌九十種の語彙調査」では、（人名・地名、助詞・助動詞等を除く β 単位の）語種別の異なり語数表・延べ語数表（国立国語研究所 1964:58）が作成され、語種の下位類（和語・漢語・外来語・混種語）を群とする使用頻度（使用率）の比較が行われている。表 6 は、そのうちの異なり語数表にもとづく度数分

6 一般に、 N が十分大きければ、平均値 $\pm$ 標準誤差の範囲に母集団平均があてはまる確率は 68%，平均値 $\pm 2 \times$ 標準誤差なら 95% となる。

7 渡部ほか（1985:143-145）は、リジット解析を量的データの群間比較に用いた例として、2 つの群（現役と浪人）の間のテスト得点の比較をあげている。

布表である。なお、頻度の階級幅は 2^n ($n=0,1,2,3,\dots$) を上限とするものに設定されている。

表 6：雑誌九十種調査の語種別度数分布表

頻度	和語	漢語	外来語	混種語	計
1	4475	5433	1563	1033	12504
2	1746	2317	481	307	4851
3～4	1565	2204	395	231	4395
5～8	1169	1668	251	111	3199
9～16	784	1162	169	74	2189
17～32	550	767	66	37	1420
33～64	375	459	25	18	877
65～	470	397	14	15	896
計	11134	14407	2964	1826	30331

この度数分布について、調査の報告書は、最上位の階級（頻度区分）では和語が漢語を上回り、それより下位の階級ではすべて漢語の方が多くなっていること、また、外来語・混種語については、語数が少ないものの、下位の階級ほど多くなっていることから、「高い使用率をもつ語の性格としては和語が漢語などより多いこと、（中略）低い使用率をもつ語の性格としては外来語、混種語の占める率が多いこと、（中略）その中間の段階では漢語の占める率が大い」(国立国語研究所 1964:61-62)と結論づけている。

このような、使用頻度の階級を高・中・低に分けて群間を比較し特徴づける解釈は、度数分布表の分析として妥当なものだが、一方では、各群の使用頻度の全体を要約し、その大きさを直観的に評価することも考えられる。たとえば、(1956 年当時の雑誌では)和語と漢語とでは総体としてはどちらが多く使われているのか、といった比較である。そのためには、(量的データの場合)上述したように、平均値による方法と(リジット解析の)平均リジットによる方法とがある。

まず、平均値による比較を行う。上述の異なり語数表・延べ語数表から、各群（語種）の総語数だけを取りだすと、表 7 のようになる。ここで、各群の異なり語数を延べ語数で割った値、いわゆる NK 値は、1 語あたりの平均使用頻度であり、要するに各群の使用頻度の平均値である。これをみると、和語の平均使用頻度が漢語のそれよりもかなり大きいことがわかる。

表 7：雑誌九十種調査の語種別語彙量（異なり語数・延べ語数）

	和語	漢語	外来語	混種語	計
異なり語数 (K)	11134	14407	2964	1826	30331
延べ語数 (N)	221875	170033	12034	8030	411972
NK 値 (N/K)	19.9	11.8	4.1	4.4	13.6

次に、平均リジットによる比較を行う。表 8 は、先の度数分布表（表 6）にリジット解析を施したもので、前節の例と同様、群全体の合計（表 6 の横計欄）を基準群とみなしてリジット変換を行い、得られたリジット（列⑤）をもとに各群の平均リジットを算出して

いる。ここでは、使用頻度の小さい階級から大きい階級へと並べているので、平均リジットの大きい群ほど、使用頻度が大きいということになる。

表 8：語種別の平均リジット算出までの計算法

頻度	①基準群 (合計)	② ①÷2	③ 累積度数	④ ②+③	⑤ ④÷N	⑥ ⑤×①	⑦ 和語	⑧ 漢語	⑨ 外来語	⑩ 混種語
1	12504	6252	0	6252	0.206	2577.4	922.4	1119.9	322.2	212.9
2	4851	2425.5	12504	14929.5	0.492	2387.8	859.4	1140.5	236.8	151.1
3~4	4395	2197.5	17355	19552.5	0.645	2833.2	1008.9	1420.8	254.6	148.9
5~8	3199	1599.5	21750	23349.5	0.770	2462.7	899.9	1284.1	193.2	85.5
9~16	2189	1094.5	24949	26043.5	0.859	1879.6	673.2	997.7	145.1	63.5
17~32	1420	710	27138	27848	0.918	1303.8	505.0	704.2	60.6	34.0
33~64	877	438.5	28558	28996.5	0.956	838.4	358.5	438.8	23.9	17.2
65~	896	448	29435	29883	0.985	882.8	463.1	391.1	13.8	14.8
計	N=30331					15165.5	5690.3	7497.1	1250.2	727.9
平均リジット						0.500	0.511	0.520	0.422	0.399

平均リジットは、和語が 0.511、漢語が 0.520 となり、わずかな差ではあるが、漢語の方が和語より大きいという結果になった。有意差検定を行うと、

$$z_0 = |0.511 - 0.520| / \sqrt{\frac{1}{12} \left(\frac{1}{11134} + \frac{1}{14407} \right)} = 2.554 > z \left(\frac{0.05}{2} \right) = 1.960$$

となり、有意水準 5 % で有意差が認められる。

以上のように、平均値による比較では和語の方が漢語より使用頻度が大きいという結果となったが、平均リジットによる比較ではわずかに漢語の方が和語より大きいという結果となった。これは、平均値による比較の方が「抵抗性」(渡部ほか 1985:2) が低い、すなわち、全体の傾向から大きく逸脱した一部データの影響を受けやすいために、和語の使用頻度が過大に評価されたものと考えられる。

上述した計算方法からわかるように、使用頻度の平均値は総延べ語数を、また、平均リジットは各階級のリジットに度数(その階級の異なり語数)を掛けたものの総和を、いずれも総異なり語数で割って求める。いま、この総異なり語数で割る 2 つの値(延べ語数、リジット×度数)が各階級にどのように分布しているかを、和語と漢語のみについてまとめると表 9 のようになる。

これを見ると、延べ語数の方は、和語・漢語ともに最上位の階級に大きく偏っており、表 6 と合わせて見れば、和語では異なり語数の 4.2 % (最上位の階級にある 470 語) が延べ語数全体の 73 % ほどを占め、漢語でも異なり語数の 2.8 % (同じく 397 語) が延べ語数全体の 53 % 余りを占めていることがわかる。つまり、平均値は、総延べ語数の多くを占めるごく少数の高頻度語の影響を大きく受け、より大きな値の方向に引き寄せられているわけである。和語の平均値が漢語のそれより大きくなるのも、この最上位の階級の差がもっぱら反映したものである。これに対して、リジット×度数の値は、累積相対度数であ

るリジットを尺度値とするため、特定の階級の値が突出するようなことはない。表9の百分率を見れば、「3～4」の階級から「17～32」の階級で漢語が和語を上回っており、このあたりの分布に厚いことが漢語の平均リジットが和語のそれをわずかではあるが上回ることに作用したものと考えられる。このように、量的データの群間比較であっても、語の使用頻度（使用率）分布などでは、平均値よりもリジット解析の方が抵抗性が高く、より妥当な方法であるといえる。

表9：平均値と平均リジットの比較（カッコ内は百分率）

頻度	延べ語数 (総和は平均値の分子)		リジット×度数 (総和は平均リジットの分子)	
	和語	漢語	和語	漢語
1	4475 (2.02)	5433 (3.20)	922.4 (16.21)	1119.9 (14.94)
2	3492 (1.57)	4634 (2.73)	859.4 (15.10)	1140.5 (15.21)
3～4	5293 (2.39)	7497 (4.41)	1008.9 (17.73)	1420.8 (18.95)
5～8	7174 (3.23)	10258 (6.03)	899.9 (15.81)	1284.1 (17.13)
9～16	9272 (4.18)	13783 (8.11)	673.2 (11.83)	997.7 (13.31)
17～32	12648 (5.70)	17708 (10.41)	505.0 (8.87)	704.2 (9.39)
33～64	17394 (7.84)	20359 (11.97)	358.5 (6.30)	438.8 (5.85)
65～	162127 (73.07)	90361 (53.14)	463.1 (8.14)	391.1 (5.22)
計	221875 (100.00)	170033 (100.00)	5690.3 (100.00)	7497.1 (100.00)

ところで、上述したように、雑誌九十種調査の報告書は、語種別の使用頻度についてこのような総体的な比較を行っていない。あくまで、使用頻度の階級を高・中・低に分けて各語種の比較を行い、それぞれの特徴を解釈しているだけなのだが、その過程で、図2のようなグラフを示している（国立国語研究所 1964:61）。これは、表7に示した各語種の総語数を使って、異なり・延べそれぞれの構成比を二重円グラフにして比較したものである。これについて、報告書は、「異なり語数では総語数についてみると漢語の方が多いのに、延べ語数では和語の方が多い」としつつも、「これは（略）度数65以上のところで和語が多いということと関連する」として、和語と漢語とを総体として比較することはしていない。

ところが、このグラフは、後に多くの文献で引用されるのだが、総語数の語種構成比であったため、（異なりと延べとを比べることで）語種間の使用頻度を総体的に比較したものとして

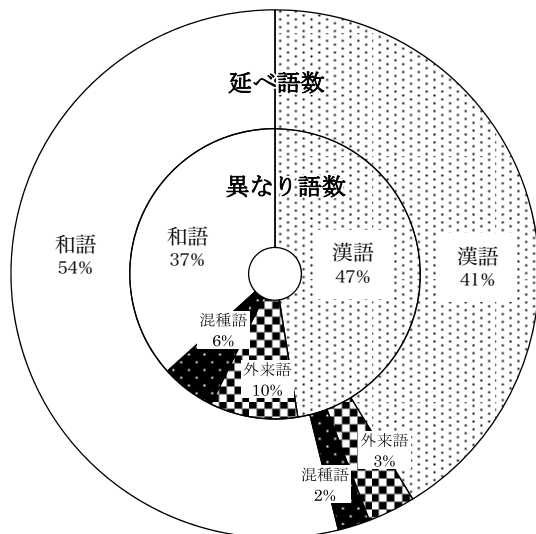


図2：異なり語数・延べ語数の語種構成比
（雑誌九十種調査のデータをもとに作成）

誤解された可能性がある。たとえば、田中章夫(1978:173)は、このグラフについて、「この事実、和語は、漢語に比べて、語の種類は少ないが、繰り返し使用される基本的な語を含んでいるのに対して、漢語はバラエティーには富んでいるが、総体的な使用頻度では和語に劣っていることを物語っている」(傍点は引用者)と記している。上述したように、延べ語数の語種構成比は、ごく少数の高頻度語、とくに和語のそれを過大に評価してしまっているから、「漢語は総体的な使用頻度では和語に劣っている」と結論づけることは適切ではない。リジット解析の結果が示すように、和語と漢語の使用頻度は総体としてはわずかに漢語の方が大きいとすべきである。

以上、連続的な量的データの群間比較にリジット解析が有用であることを述べたが、その際、度数分布表の階級設定が任意となることには注意が必要である。質的データの分析ではあらかじめ設定された順序カテゴリーの区分に従えばよいが、量的データの場合にはどのように階級区分を設定するかで、平均リジットの値が異なる可能性があるからである。ただし、階級設定の仕方には絶対的な基準はなく、「ステージスの公式」をはじめとする経験的な方法が開発されてはいるものの、それぞれの専門分野あるいは個別の事象ごとに経験的に妥当とされる方法が採られているのが実情であろう。本節の例でも、階級幅の上限を 2^n ($n=0,1,2,3,\dots$) とする区分が採用されているが、これも語の使用頻度分布の階級設定としてはよく行われるものである。階級設定の問題は、リジット解析とは切り離して考えるべき、独立した問題である。

5. 2変数の順序カテゴリーの数量化

リジット解析は、本来、表 10 のような、群×変数の分割表があるとき、変数の順序カテゴリー(量的変数の場合は順序性をもつ階級区分)を尺度値化することによって各群を数量化し、群間の相対的な大きさを直観的に評価しようとする手法であるが、表 11 のような、2つの変数からなる分割表でも、列行両変数のカテゴリーがともに順序性を持つ場合には、適用することができるものと考えられる。たとえば、2節で紹介した薬効試験の例で、実験群を病状の程度によってさらに複数の群に分けたような場合、あるいは、年齢によってさらに複数の群に分けたよ

表 10: 群×変数の分割表 (n は度数)

		群			
		A	B	...	Z
変数 X の カテ ゴリ (階級)	X_1	n_{1A}	n_{1B}	...	n_{1Z}
	X_2	n_{2A}	n_{2B}	...	n_{2Z}
	$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$
	X_i	n_{iA}	n_{iB}	...	n_{iZ}
	計	n_{+A}	n_{+B}	...	n_{+Z}

表 11: 変数×変数の分割表 (n は度数)

		変数 Y のカテゴリー				計
		Y_1	Y_2	...	Y_j	
変数 X の カテ ゴリ	X_1	n_{11}	n_{12}	...	n_{1j}	n_{1+}
	X_2	n_{21}	n_{22}	...	n_{2j}	n_{2+}
	$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$	$\vdots$
	X_i	n_{i1}	n_{i2}	...	n_{ij}	n_{i+}
	計	n_{+1}	n_{+2}	...	n_{+j}	n_{++}

うな場合、それらは群であると同時に、病状の程度という質的変数の順序カテゴリー、あるいは、年齢という量的変数の順序性をもつ階級区分でもある。このような場合は、これらを順序カテゴリーとして尺度値化し、薬効のカテゴリー(著効～悪化)を群とみなして数量化することが可能になる。つまり、2つの変数についてリジット解析を交互に行うと

いうことである。表 11 でいえば、順序はどちらでもよいのだが、まず、列変数 Y の各カテゴリ ($Y_{i \sim j}$) を群とみなし、行変数 X の各カテゴリ ($X_{i \sim i}$) を順序カテゴリとしてリジット解析を行い、次いで、行変数 X の各カテゴリ ($X_{i \sim i}$) を群とみなし、列変数 Y の各カテゴリ ($Y_{i \sim j}$) を順序カテゴリとしてリジット解析を行うわけである。こうすることで、2つの変数の順序カテゴリがともに数量化でき、それぞれの相対的な大きさを評価できることになる。

ところで、質的変数の 2 元分割表におけるカテゴリの数量化については、その「計算方法は多数存在するが、そのうちのひとつである交互平均法は歴史的に古く、かつ計算過程が最も理解しやすい」(西里 2010:27)といわれる。日本語研究では、荻野綱男(1980)が交互平均法と同じ計算法(荻野の数量化の方法)を独自に考案し、質的変数の 2 元分割表に適用して、それぞれのカテゴリを数量化(間隔尺度に変換)している。以下では、この、質的変数の 2 元分割表における順序カテゴリの数量化について、交互平均法とは別にリジット解析も適用できることを、荻野(1980)による交互平均法の解析結果と対照させることにより、述べていく。

表 12:「知っている」の分割表

場 面 語 形	V	I	II	III	IV	VI	計
シッテル類	438	371	164	90	10	4	1077
シッテマス類	39	110	280	279	243	195	1146
シッテオリマス類	7	17	39	88	142	143	436
ゾンジテオリマス類	1	3	9	31	98	148	290
計	485	501	492	488	493	490	2949

表 12 に、荻野(1980:14)が交互平均法を適用した分割表を示す。これは、「いろいろな話し相手 6 人に対して「知っている」という場合、どのような語形を使うか」を 503 人のインフォーマントに質問した結果である。ここで、「場面」とは「話し相手(聞き手)」のことで、右に示す 6 類が設定され、「語形」は回答から得られた 107 種が(「その他」を除く) 4 類にまとめられたものである。

場面Ⅰ＝同じ年頃の親しい友人
 場面Ⅱ＝あまり親しくない人で、少し年下の人
 場面Ⅲ＝親しい人で、少し年上の人
 場面Ⅳ＝あまり親しくない人で、少し目上の人
 場面Ⅴ＝ふだんことばづかいをいちばん気にしない
 いで話ができる相手
 場面Ⅵ＝ふだんいちばん丁寧なことばで話をする
 相手

これらの類はいずれも順序カテゴリであり、分割表の度数分布からみて、場面は「Ⅴ→Ⅰ→Ⅱ→Ⅲ→Ⅳ→Ⅵ」の順に、語形は「シッテル類→シッテマス類→シッテオリマス類→ゾンジテオリマス類」の順に、より丁寧になることがわかる。荻野は、こうしたカテゴリの順序だけでなく、それらが「どれくらい離れて並んでいるのか(どれくらい丁寧なのか)」を知りたい(荻野 1980:14)として、この分割表に交互平均法による数量化を施す。

交互平均法とは、行変数のカテゴリと列変数のカテゴリに、分割表の度数分布を最もよく説明する最適の重みを割り出して与え(＝数量化し)、それぞれの変数のカテゴリ

一間の相対的な大きさを評価する統計手法である。その計算法は、以下の通り（荻野 1980:15）。

- (1)各語形に次のような初期値（重み）を与える。

シッテル類	シッテマス類	シッテオリマス類	ゾンジテオリマス類
1	2	3	4

- (2)この語形の重みを使って、各場面の平均得点を求める（場面ごとに、語形の度数と重みをかけて合計したものを場面の総度数で割る）。

場面Ⅴ	場面Ⅰ	場面Ⅱ	場面Ⅲ	場面Ⅳ	場面Ⅵ
1.115	1.305	1.783	2.123	2.665	2.888

- (3)この場面の平均得点を、最小が 1，最大が 6 になるように比例変換する。

場面Ⅴ	場面Ⅰ	場面Ⅱ	場面Ⅲ	場面Ⅳ	場面Ⅵ
1.000	1.536	2.882	3.842	5.372	6.000

- (4)この点数を場面の重みとして、各語形の平均得点を求める（語形ごとに、場面の度数と重みをかけて合計したものを語形の総度数で割る）。

シッテル類	シッテマス類	シッテオリマス類	ゾンジテオリマス類
1.768	3.981	4.827	5.397

- (5)この語形の平均得点を、最小が 1，最大が 4 になるように比例変換する。

シッテル類	シッテマス類	シッテオリマス類	ゾンジテオリマス類
1.000	2.830	3.529	4.000

- (6)以上の(2)から(5)を 1 サイクルとして、重みの値が収束するまで（通常は 5～7 サイクル）繰り返す。最終結果は、以下の通り。

場面Ⅴ	場面Ⅰ	場面Ⅱ	場面Ⅲ	場面Ⅳ	場面Ⅵ
1.000	1.757	3.655	4.550	5.717	6.000

シッテル類	シッテマス類	シッテオリマス類	ゾンジテオリマス類
1.000	3.015	3.619	4.000

このように、交互平均法は、行変数・列変数の各カテゴリーの重みを求めて交互平均の過程を繰り返すものだが、これにより、場面・語形の順序カテゴリーは（(6)のように）間隔尺度に変換＝数量化され、

- [1] 場面ⅠとⅡの差が大きく、場面ⅣとⅥの差が小さい。また、シッテル類とシッテマス類の差がかなり大きく、シッテオリマス類とゾンジテオリマス類の差が小さい（荻野 1980:15）

という結果（評価）が得られる。西里（2010:31）によれば、交互平均法により最終的に得られた各カテゴリーの重みは、データへの回帰が線形になるような変換値であり（完全線形回帰）、分割表のデータを説明するための最適な数量となる、という。

続いて、表 12 の分割表にリジット解析を施す。ここでは、場面・語形とも「合計」の数値列を基準群として解析を行う。表 13 が場面の、表 14 が語形の解析結果である。

表 13: 場面のリジット解析

語形	①合計 の度数	②リジ ット	②×①	②× 場面 V	②× 場面 I	②× 場面 II	②× 場面 III	②× 場面 IV	②× 場面 VI
シッテル類	1077	0.183	196.7	80.0	67.7	29.9	16.4	1.8	0.7
シッテマス類	1146	0.560	641.2	21.8	61.5	156.7	156.1	136.0	109.1
シッテオリマス類	436	0.828	360.9	5.8	14.1	32.3	72.8	117.5	118.4
ゾンジテオリマス類	290	0.951	275.7	1.0	2.9	8.6	29.5	93.2	140.7
計	2949		1474.5	108.5	146.2	227.4	274.9	348.5	368.9
平均リジット			0.500	0.224	0.292	0.462	0.563	0.707	0.753
リジット解析				1.000	1.643	3.254	4.208	5.565	6.000
交互平均法				1.000	1.757	3.655	4.550	5.717	6.000

表 14: 語形のリジット解析

場面	①合計の 度数	②リジット	②×①	②×シッ テル類	②×シッ テマス類	②×シッテ オリマス類	②×ゾンジテ オリマス類
V	485	0.082	39.9	36.0	3.2	0.6	0.1
I	501	0.249	125.0	92.5	27.4	4.2	0.7
II	492	0.418	205.5	68.5	117.0	16.3	3.8
III	488	0.584	285.0	52.6	162.9	51.4	18.1
IV	493	0.750	369.9	7.5	182.3	106.5	73.5
VI	490	0.917	449.3	3.7	178.8	131.1	135.7
計	2949		1474.5	260.8	671.6	310.1	231.9
平均リジット			0.500	0.242	0.586	0.711	0.800
リジット解析				1.000	2.850	3.525	4.000
交互平均法				1.000	3.015	3.619	4.000

表 13・14 の下部には、場面・語形それぞれの平均リジットを荻野（1980）と同様に 1～6、1～4 の範囲に比例変換し、交互平均法の結果とともに示した。これをみると、リジット解析の結果はおおよそ交互平均法の結果に近く、上の [1] はリジット解析によっても導くことができる。したがって、リジット解析を表 12 のような分割表に適用して順序カテゴリーの数量化に用いることは、おおむね妥当であるものと考えられる。

さらに、リジット解析では、前述したように、平均リジットの差の有意性を検定することができる。表 13・14 で求めた平均リジットについて、3 節で紹介した、合計の数値列を基準群とみなすときの、異なる 2 群の平均リジットの有意差検定を行うと、場面および語形のいずれにおいても、すべての組み合わせで統計的な有意差が認められる。ただ、インフォーマントや語形をさらに細分して分析するような場合は、データ数が少なくなって標準誤差の値が大きくなり、結果として有意差が認められないこともあるようである。表

15 は、荻野 (1980:21) の表 5 をもとに、出現頻度の高い 14 種の語形について、交互平均法による「丁寧さ」とリジット解析による平均リジットを総度数とともに記したものであるが、丁寧さと平均リジットの順位は完全に並行している。荻野は、このうち、表 16 に示す、

表 15: 代表語形 14 種の数量化

語形	総度数	丁寧さ	平均 リジット
シッテル	183	1.00	0.230
シッテルヨ	465	1.09	0.237
シッテルワヨ	120	1.12	0.246
シッテルワ	121	1.30	0.250
シッテイルヨ	40	1.86	0.273
シッテイル	71	2.42	0.289
シッテマスヨ	250	7.27	0.486
シッテイマスヨ	49	7.90	0.511
シッテオリマスヨ	37	9.09	0.566
シッテマス	512	10.27	0.618
シッテイマス	297	10.71	0.641
シッテオリマス	376	12.68	0.735
ゾンジテマス	74	13.84	0.803
ゾンジテオリマス	142	14.00	0.810

のような対では、一貫して「イ」のついた語形の方がつかない語形より丁寧であるものの、ぞんざいな語形に「イ」がつくと丁寧さが大幅に上がるが、丁寧な語形についてもそれほど大幅には上がらないことから、「イ」が丁寧さを増すとしても、その働きはどのような語形に続くかによって異なるものとする。一方、これらの対の平均リジットの差について有意差検定を行うと、いずれの z_0 値も有意水準 5 % の限界値 (1.960) に達しないため、そもそも「イ」のついた語形の方が丁寧であると結論づけることはできないという結果になる (表 16)。もちろん、有意差がないとはいえ、上のような (丁寧さ・平均リジットのいずれで見ても) 一貫した傾向があることを重視すれば、「イ」のついた語形の方がつかない語形より丁寧であるというのは、有力な

表 16: 語形 4 対の有意差検定

語形	総度数	丁寧さ	平均 リジット	z_0 値
シッテル シッテイル	183 71	1.00 2.42	0.230 0.289	1.471
シッテルヨ シッテイルヨ	465 40	1.09 1.86	0.237 0.273	0.738
シッテマスヨ シッテイマスヨ	250 49	7.27 7.90	0.486 0.511	0.555
シッテマス シッテイマス	512 297	10.27 10.71	0.618 0.641	1.106

仮説であるといってよい。ただし、その仮説は、「イ」の働きはどのような語形に続くかによって異なるという見方とあわせて、なお (データ数を増やすなどして) 検証していく必要があるということを、リジット解析の結果は示していると言える。

以上のように、リジット解析は、本来、群×変数の分割表における群間比較のための統計手法ではあるが、カテゴリー変数の 2 元分割表における順序カテゴリーの数量化にも用いることができるものと考えられる。その計算法は交互平均法に比べてきわめて平易であり、数量化の結果も交互平均法に近似する。また、カテゴリー間の平均リジットの差の有意性を検定することもできる。交互平均法に加え、こうした特長をもつリジット解析を用いることで、この種の分割表における順序カテゴリーの数量化の検討は、より多面的に行うことができるものと考えられる。

6. おわりに

以上、リジット解析という統計手法が、言語研究における計数データの分析でも、各群あるいは各変数カテゴリーの相対的な大きさやその差を直観的に評価する手法として有用であること、具体的には、本来の質的データの群間比較のほかにも、量的データの群間比較、さらには、2変数の分割表におけるそれぞれの順序カテゴリーの数量化にも適用できることを述べた。

リジット解析の特長は、その考え方と計算法がきわめて平易なことである。また、求める平均リジットは、確率であり平均値でもあって、それ自体が意味をもつ数値であるからわかりやすいし、有意差検定もできる。さらに、4節の量的データの群間比較でも述べたように、リジット解析は（平均値による比較よりも）抵抗性が高く、探索的データ解析の手法⁸としての側面ももっている。

ただし、小稿では、ここで例としてあげた先行調査研究の方法よりリジット解析の方が優れているとか、リジット解析の方を使うべきだとか、主張したいわけではない⁹。リジット解析は、基本的には、（確認的ではなく）探索的な解析法の一つとして、これまでに行われている他の統計手法などと併用していくことが望ましいのではないかと考えている。

なお、リジット解析では、2節で紹介したような基準群と比較群との解析で、基準群の分布に（たとえば、連続するいくつかのカテゴリーにそろって度数がないような）極端な偏りがあったり（橋本・長谷 1992:4）、3節で紹介したような合計（の数値列）を基準群として行う解析でも、比較しようとする異なる2群の分布が鏡像関係になったり（フライスほか 2003:172-173）すると、正しい解析が行えない場合のあることが指摘されている。ただし、これらの問題は、前者では合計の数値列を基準群とみなすことで、後者の場合は、比較する2群のどちらか一方を事後的な基準群とすることで回避できることが多いとされ、リジット解析の有用性を損なうものではないと考えられる。

文献

Bross, I.D.J. (1958) How to use ridit analysis. *Biometrics*, 14(1):18-38.

石村貞夫・謝承泰・久保田基夫（2003）『SPSSによる医学・歯学・薬学のための統計解析』東京図書。

遠藤和男・山本正治（1992）『医統計テキスト』西村書店。

8 渡部ほか（1985:143-145）によれば、探索的データ解析で行われる「リジット解析」では、リジットに代えて「ss比」が用いられる。ss比とは、データをある境界値で区切ったときに、その境界値未満のデータ数を n 、その境界値に等しいデータ数を m 、始数（データ数が0にならないために適当に加える小さな数）を s としたときの、その境界値までの累積度数（ss数）について、全データ数を N としたときのss数の全体に対する比率 p であり、以下の式で表すことができる。

$$p = \frac{n+m/2+s}{N+2s}$$

要するに、対数変換などの再表現に備えて始数を（リジットの分子に s 、分母に $2s$ ）加える点だけが異なるのだが、データ数が大きい場合は始数の影響は無視し得るので、両者はほとんど同じと考えてよい。

9 とはいえ、フライス（1973:110）が指摘するように、リッカート法よりは優れた数量化の手法であると考えられる。

- 岡崎和夫（1980）「『見レル』『食ベレル』型の可能表現について—現代東京の中学生・高校生について行った一つの調査から—」『言語生活』340:64-70.
- 荻野綱男（1980）「敬語における丁寧さの数量化—札幌における敬語調査から（2）—」『国語学』120: 左 13-24.
- 国立国語研究所（1964）『現代雑誌九十種の用語用字 第3分冊 分析』秀英出版.
- 竹内啓ほか編（1989）『統計学辞典』東洋経済新報社.
- 田中章夫（1978）『国語語彙論』明治書院.
- 富永祐民（1982）『治療効果判定のための実用統計学—生命表法の解説—』蟹書房.
- 西里静彦（2010）『行動科学のためのデータ解析—情報把握に適した方法の利用—』培風館.
- 橋本修二・長谷文雄（1992）「リジット解析の紹介—なぜ適切でないか—」『医薬安全性研究会会報』35:1-5.
- J.L. フライス（1973）（佐久間昭訳 1975）『計数データの統計学—医学・疫学を中心に—』東京大学出版会.
- J.L. フライスほか（2003）（Fleiss 愛好会訳 2009）『計数データの統計学 Third Edition』(株)アーム.
- 渡部洋・鈴木規夫・山田文康・大塚雄作（1985）『探索的データ解析入門—データの構造を探る—』朝倉書店.

付記

小稿は、平成 15～17 年度科学研究費補助金（基盤研究(C)）「探索的データ解析による日本語研究法の開発」（研究代表者：石井正彦）の成果の一部、および、計量国語学会第 59 回大会（2015 年 9 月 26 日、神戸大学）における口頭発表「交互平均法とリジット解析」の内容に加筆修正してまとめたものである。口頭発表の際には、有益な教示・助言を多数いただいた。記して感謝申し上げる。

（2016 年 6 月 3 日受付・2016 年 6 月 17 日改稿受付）

Resource

Ridit Analysis: Its Application to Linguistic Count Data

ISHII Masahiko (Osaka University)

Abstract:

In this essay, I introduce "ridit analysis" commonly used in the field of medical statistics, and show that it is also useful for the analysis of linguistic count data. The ridit analysis is originally a statistical method for comparing among groups in a frequency table by means of the average ridit value. The average ridit is extremely easy to calculate and available for the statistical significance test. By applying the ridit analysis to the count data of some previous researches in Japanese linguistics, it is revealed that the analysis is applicable not only to the comparison among groups of qualitative or quantitative data but also to the quantification of each ordinal variable in a two-dimensional contingency table. The ridit analysis that is also regarded as one of the methods of exploratory data analysis (EDA) has a certain effectiveness in quantitative linguistic researches.

Keywords: ridit analysis, average ridit, quantification, Likert scaling method, reciprocal averaging method, exploratory data analysis