

私の仕事

筆者が大学に入ったのは、大学紛争が終わった直後で、入学式も卒業式もなかったくらいだからまだ、きちんとした体系だった教育システムはどここの学部でも整っていなかった。おかげで、研究に必要な常識とか基礎的技術もほとんど知らないままに、研究者を職業に選んでしまった。研究者としての最初の仕事は水産実験所の職員だった。水産実験所の職員の仕事の内容は、大学本部の教員や研究所に所属する研究者の仕事とは少し違っている。自分の研究テーマはあるし、研究テーマも自分で考えて良いという自由ももちろんあるが、実験所に様々な人が来る。研究者・学生・院生を含めて、外部からくる研究者の手伝いも業務の内だ。同じ学部の人もあるが、他学部からも来るし、他の大学や研究機関からも来る。利用者が多い実験所は良い実験所なのだ。施設や設備の良さに加えて、使いやすさが実験所の評価ということになる。水産実験所には技術職員（当時は技官という職名）がいて、彼らが主としてその外部利用者に対するサービスをするが、彼らだけでは間に合わないこともある。そうすると、研究所の助手（今の助教）や助教授（今の准教授）の知識・技量が問われる。小型船舶を操船して、調査やサンプリングの手伝いをしたり、実験所の施設（特に飼育設備・ほとんどが水槽）を使った実験の手伝いをしたり、頼まれた試料を集めて化学分析をしたり、自分の専門以外のこともする。外部の研究者に代わって採集許可を取って生物採集をしたりもする。技術職員（当時は技官）も含めて、どのくらいの情報、知識、技術があって、どのくらいのことを頼めるのか、ということが重要なのだ。利用者の信頼を得ると、飼育実験など実験そのものをこちらに丸投げしてくることもある。自分の研究テーマではないのに、長期にわたる飼育実験を頼まれることもあって、割の合わない仕事のようなのだが、うまくいって、投稿論文になれば、高名な研究者が連名者に入れてくれることもあるし、自分の専門外のことも含めて、いろいろなことを学べるから、若手の研究者にとっては悪い職場ではない。特に、基礎的知識・技術がない筆者は、実験所で、様々なことを学んだ。ラテン方格というのも、他人の仕事の手伝いで覚えた言葉だ。

シラスウナギとラテン方格

頼まれた実験の内容はシラスウナギの捕獲するためのトラップ（籠網）の性能を評価してくれというものだった。水産実験所には野外水槽がいくつもあるから、何とかなるだろうと思って、例によって、安請け合いしてしまった。そのころシラスウナギの値段はすでに高騰していたが、まだ、1kg100万円はしなかった。それでも、普通に買えば、80万円/kgくらいはした。養鰻組合に事情を話して、安く売ってくれるように手配してもらって、1kg 数十万円で手に入れることができた。シラスウナギの購入費も含めて、まとまった金額の費用も概算でもらったから、シラスウナギを安く買えば、差額が、自由に使える研究費として手元に残る。自分の研究費を持たない助手としてうまい話だ。

どこで作ったのかは知らないが、シラスウナギのトラップ（籠網）が送られてきた。なんで

も、シラスウナギは目立つものによって来るというので、色、大きさ（高さ）、形の異なるトラップを作ったということだった、トラップは4つあった。材料はそろったのだが、ここで、ハタと困ってしまった。4つのトラップを水槽に放り込んで、一晩ぐらいでとり上げて、トラップの中のシラスウナギの数を数える。これをいく晩か繰り返せばよいだろうと、漠然と甘く考えていたのだが、そうはいかないということが、池の中のシラスウナギの分布を見たらわかった。シラスウナギは池の中に一様の分布していない。シラスウナギは池の中央部にはほとんどいない。池のへりに固まって分布していて、隅の角の部分には特に多くにシラスウナギがいる。また、夕方の光の加減なのか、4つの隅ごとに、シラスウナギの数が違う。トラップの色や形よりも、角との位置関係の方が、明らかに、シラスウナギの捕獲量に重要な要素になるだろうと予想がついた。トラップは隅に入れるしかないが、隅の角との位置関係も問題になるから、一つの隅には一つのトラップしか入れられない。そうすると、隅が、4つだから、第一日目には、一つの隅に一つずつトラップを入れて、翌日、トラップを引き上げて、トラップのなかのシラスウナギの数を数え、シラスウナギを別のところへ移したうえで、一つずつ、トラップをローテーションして、隣の隅に入れて、翌日、取り上げて、シラスウナギの数を数える、これを、4日間繰り返せば、すべての、トラップを、異なる4つの隅でシラスウナギの捕獲効率の試験をしたことになる。4回の平均値は、隅の違いと実験日の違いによる効果を打ち消したものになるだろうと、直感的に考えた。これを、もう少し数学的に、何がどのように分析可能なかを考える。

表1にこの実験配置をまとめた。

表1. シラスウナギトラップの性能評価の実験配置

		隅			
		a	b	c	d
実験日	1	A	B	C	D
	2	D	A	B	C
	3	C	D	A	B
	4	B	C	D	A

4×4の個々のマス目に入るデータ（捕獲されたシラスウナギの数）を $D_{(trap, day, corner)}$ と表すことにする。例えば、実験日第一日のaの隅に入れられた、Aのトラップで漁獲されたシラスウナギの数は、 $D_{(A, 1, a)}$ と表す。

それぞれのデータは、トラップの効果と、実験日の効果と、隅の効果と、その他ここで取り上げていない他の効果、これら4つの和によってできているから、これらを

I_{corner} , I_{day} , I_{trap} , e とすると、データは

$$D_{(A, 1, a)} = I_{trap=A} + I_{day=1} + I_{corner=a} + e_{A, 1, a}$$

のように表現することができる。

A,B,C,Dのトラップごとにデータを整理すると、

トラップ A について

$$D_{(A,1,a)} = I_{trap=A} + I_{day=1} + I_{corner=a} + e_{A,1,a}$$

$$D_{(A,2,b)} = I_{trap=A} + I_{day=2} + I_{corner=b} + e_{A,2,b}$$

$$D_{(A,3,c)} = I_{trap=A} + I_{day=3} + I_{corner=c} + e_{A,3,c}$$

$$D_{(A,4,d)} = I_{trap=A} + I_{day=4} + I_{corner=d} + e_{A,4,d}$$

左辺、右辺をそれぞれ足し合わせると

$$\sum_{i=1,j=a}^{4,d} D_{(A,i,d)} = 4I_{trap=A} + \sum_{i=1}^4 I_{day=i} + \sum_{j=a}^d I_{corner=j} + \sum_{i=1,j=a}^{4,d} e_{A,i,j} \quad \text{式 1}$$

両辺を 4 で割ると

$$\frac{1}{4} \sum_{i=1,j=a}^{4,d} D_{(A,i,d)} = I_{trap=A} + \frac{1}{4} \sum_{i=1}^4 I_{day=i} + \frac{1}{4} \sum_{j=a}^d I_{corner=j} + \frac{1}{4} \sum_{i=1,j=a}^{4,d} e_{A,i,j} \quad \text{式 2}$$

となる。 $\frac{1}{4} \sum_{i=1}^4 I_{day=i}$ は I_{day} の平均値、 $\frac{1}{4} \sum_{j=a}^d I_{corner=j}$ は I_{corner} の平均値だから、これを

$$\frac{1}{4} \sum_{i=1}^4 I_{day=i} = M_{day}, \quad \frac{1}{4} \sum_{j=a}^d I_{corner=j} = M_{corner}$$

とあらわすと

$$\frac{1}{4} \sum_{i=1,j=a}^{4,d} D_{(A,i,d)} = I_{trap=A} + M_{day} + M_{corner} + \frac{1}{4} \sum_{i=1,j=a}^{4,d} e_{A,i,j} \quad \text{式 3}$$

ところで、全体の平均は、trap、実験日、隅の効果の平均値の和のはずだから、

$$M = M_{trap} + M_{day} + M_{corner}$$

となるので、

$$M_{day} + M_{corner} = M - M_{trap}$$

これを式 3 に代入すると

$$\frac{1}{4} \sum_{i=1,j=a}^{4,d} D_{(A,i,d)} = I_{trap=A} + M - M_{trap} + \frac{1}{4} \sum_{i=1,j=a}^{4,d} e_{A,i,j}$$

$$\frac{1}{4} \sum_{i=1,j=a}^{4,d} D_{(A,i,d)} - M = I_{trap=A} - M_{trap} + \frac{1}{4} \sum_{i=1,j=a}^{4,d} e_{A,i,j}$$

何の条件もなければ、

$$\frac{1}{4} \sum_{i=1,j=a}^{4,d} e_{A,i,j} = 0$$

だから、

$$\frac{1}{4} \sum_{i=1, j=a}^{4, d} D_{(A, i, d)} - M = I_{trap \text{ A}} - M_{trap} \quad \text{式 4}$$

となる。

式の意味を日本語で表現すると、

「同じトラップをつかったデータ平均値とデータ全体の平均値の差は、そのトラップの効果とトラップの効果の平均値の差（偏差）に等しい。」

となる。

もとにもどって、何をしたかったのかを考えると、トラップの性能の違いを、捕獲されたシラスウナギの数をデータとして評価したいのだけれど、捕獲されるシラスウナギの数（データ）は、実験日の天候状態、水槽の隅の条件（方向や回りの構造物）影響を受けていて、これらの効果とトラップの効果を足し合わせたものであり、トラップによる効果を単独で取り出すことはできない。

そこで、「トラップ A を使ったすべてのデータを集めて、全体の平均との「偏差」を計算すれば、それは、トラップ A とトラップについての平均との偏差になる。これをデータとして使えばよい。」ということである

偏差の 2 乗を全トラップについて計算し、その総和を求めれば平方和（sum of square:SS）となる。平方和を自由度で割ったものが分散である。トラップの分散とは別に、残渣の分散が計算できれば、残渣分散に対して、トラップの差による分散がどのくらい大きいかを評価すれば、トラップの偏差が統計的に有意であるか否かを論ずることもできる。

つまり、トラップの効果を単独で取り出さなくても、その効果の統計的な有意性を論じることができるという意味である。

実際に計算してみる方が、理解しやすいかもしれないので、データを入れて計算してみる。データは、思いつく数字を適当に入れたものである。

表 2. シラスウナギトラップの性能評価の実験の結果（シラスウナギの捕獲尾数）

	a		b		c		d	
1	A	219	B	183	C	199	D	165
2	D	186	A	197	B	160	C	219
3	C	210	D	141	A	150	B	152
4	B	152	C	168	D	96	A	148

最初に全体の平均 M と分散を計算する。あまり、こういう計算をしたことがない読者のために、計算手順を説明しておく。分散を計算するときには平方和 (SS) をまず計算しなければならない。平方和の定義は $\sum_{i=1}^n (x_i - \bar{x})^2$ だが、定義式の通りに計算するのは大変だ。次のように定義式を変形すれば、総和と二乗和を計算すればよいので、計算が楽になる。

$$\begin{aligned}\sum_{i=1}^n (x_i - \bar{x})^2 &= \sum_{i=1}^n x_i^2 - 2 \sum_{i=1}^n \bar{x}x + \sum_{i=1}^n \bar{x}^2 \\ &= \sum_{i=1}^n x_i^2 - 2\bar{x} \sum_{i=1}^n x + n\bar{x}^2\end{aligned}$$

$\bar{x} = \frac{1}{n} \sum_{i=1}^n x$ をつかって

$$\begin{aligned}&= \sum_{i=1}^n x_i^2 - 2 \frac{1}{n} \left(\sum_{i=1}^n x \right)^2 + n \left(\frac{1}{n} \sum_{i=1}^n x \right)^2 \\ &= \sum_{i=1}^n x_i^2 - 2 \frac{1}{n} \left(\sum_{i=1}^n x \right)^2 + \frac{1}{n} \left(\sum_{i=1}^n x \right)^2 \\ &= \sum_{i=1}^n x_i^2 - \frac{1}{n} \left(\sum_{i=1}^n x \right)^2\end{aligned}$$

EXCEL による計算シートを例示しておく。

図 1. EXCELL による計算手順 1. (全平均・全分散の計算)

総和の計算

	a	b	c	d	Sum	square
1	219	183	199	165	766	
2	186	197	160	219	762	
3	210	141	150	152	653	
4	152	168	96	148	564	
Sum					2745	7535025
mean					171.5625	470939.1

↑

平均

二乗和の計算

	a	b	c	d	Sum
1	47961	33489	39601	27225	148276
2	34596	38809	25600	47961	146966
3	44100	19881	22500	23104	109585
4	23104	28224	9216	21904	82448
Sum					487275
mean					30454.69

↑
平均

左の表がデータの総和の計算のシートで、右がデータの 2 乗の総和を計算するための計算シートである。左の表で、実験日ごとに、すべての隅のデータの合計を計算したのが Sum の列であり、その合計が $766+762+653+564=2745$ となる。これをデータの総数で割った $2745 \div 16 = 171.5625$ が全データの平均値 M である。総和の 2 乗が $(\sum_{i=1}^n x)^2 = 7535025$ であり、これをデータ総数 $N=16$ で割った値が $\frac{1}{n} (\sum_{i=1}^n x)^2 = 470939.06$ である。

右の表で、Sum の列の総和、 $148276+146966+109585+82448=487275$ が二乗の総和である。二乗総和から、総和の二乗の N 分の 1 を差し引いたものが、全 SS (total Sum of square) であり、これを全自由度 (tdf) $=16-1=15$ で割ったものが全分散 (total

variance) である。

$$\sum_{i=1}^n x_i^2 - \frac{1}{n} \left(\sum_{i=1}^n x \right)^2 = 487275 - 470939.06 = 16335.94$$

$$tv = \frac{16335.94}{15} = 1089.063 \cong 1089$$

これ計算しておいて、データを、trap について整理しなおす。

図 2. EXCELL による計算手順 2. (trap の効果の計算)

wo

	1	2	3	4	Sum	mean	mean-M	square
A 1a	219	A2b 197	A3c 150	A4d 148	714	178.5	6.9375	48.1289063
B 1b	183	B2c 160	B3d 152	B4a 152	647	161.75	-9.8125	96.2851563
C1c	199	C2d 219	C3a 210	C4b 168	796	199	27.4375	752.816406
D 1d	165	D2a 186	D3b 141	D4c 96	588	147	-24.5625	603.316406
					2745			1500.54688
					171.5625			

図 2 は表 2 を A,B,C,D のデータごとにわけて整理しなおしたものである。Sum の列が trap ごとの総和、mean の列が平均値である。これらを、そのまま trap ごとのデータとして記述しても、間違いではないが、それぞれのトラップの効果という意味では、平均からの偏差として記述したほう (mean-M) が、それぞれの要因の効果と比較するには便利かもしれない。個の偏差の 2 乗の総和が SS_{trap} だが、表計算上の 2 乗和は、1500.54688 だが、これは 4 つのデータの平均だから、SS はその 4 倍で、

$$SS_{trap} = 4 \times 1500.54688 \cong 6002$$

トラップについてのデータ数が 4 つだから、 $n_{trap} = 4$ で、自由度は $df_{trap} = n_{trap} - 1 = 4 - 1 = 3$ だから、trap の分散は、

$$v_{trap} = \frac{SS_{trap}}{df_{trap}} = \frac{6002}{3} = 2001$$

となる。

話の核心はここからである。まず、やってはいけないことを書く。この結果をもとに、

$$SS_{total}=16336$$

$$SS_{trap} = 6002$$

で、

$$df_{total}=15$$

$$df_{trap}=3$$

だから、残渣の平方和、自由度、分散を

$$SS_{residual}=16336-6002=10334$$

$$df_{residual} = 15 - 3 = 12$$

$$v_{residual} = \frac{10334}{12} = 861$$

とし、残渣の分散と trap の分散の比を

$$F_{\frac{trap}{residual}} = \frac{2001}{861} = 2.32$$

$$P_{trap}=0.126 \quad (\text{計算は EXCEL の関数 F.DIST.RT})$$

したがって、trap の違いが捕獲量に影響するとは言えない。

(たまたまそれ以上の値になる可能性が 12.6%ある。)

というような、F 検定してはいけない。分析しようとする要因以外は残渣と考えて、その分散で割って、分散比 F を計算して、何が悪いと開き直られると、一応、理屈が通っていないことはないのだが、せっかく、ラテン方格にした意味がない。

$$SS_{residual}=16336-6002=10334$$

は残渣(説明できないランダムの変動)の SS とは言えない。この中には、実験日 (day) の効果と隅(corner)の効果がふくまれている。これらをコントロールすることはできないが、説明可能な変動である。だから、これによる分散は取り除ける。取り除けるものは取り除かなければならない。

Day の効果と SS_{day} を計算する。

図 3. EXCELL による計算手順 3. (実験日 day の効果の計算)

	a			b			c			d			Sum	mean	mean-M	square
1	A1a	219		B1b	183		C1c	199		D1d	165		766	191.5	19.9375	397.503906
2	D2a	186		A2b	197		B2c	160		C2d	219		762	190.5	18.9375	358.628906
3	C3a	210		D3b	141		A3c	150		B3d	152		653	163.25	-8.3125	69.0976563
4	B4a	152		C4b	168		D4c	96		A4d	148		564	141	-30.5625	934.066406
													2745			1759.29688
													171.5625			

Trap の効果と同様に計算して、

$$SS_{day} = 4 \times 1759.29688 = 7037$$

$$df_{day}=3$$

$$v_{day} = \frac{SS_{day}}{df_{day}} = \frac{7037}{3} = 2346$$

次に隅 corner の効果と SS_{day} を計算する。

Trap, day と同様に計算して

$$SS_{corner} = 4 \times 820.921875 = 3284$$

$$df_{corner}=3$$

$$v_{corner} = \frac{SS_{corner}}{df_{corner}} = \frac{3284}{3} = 1095$$

となる。これらを全 SS から差し引くと

図 4. EXCELL による計算手順 3. (隅 corner の効果の計算)

	1		2		3		4		Sum	mean		
a	A1a	219	D2a	186	C3a	210	B4a	152	767	191.75	20.1875	407.535156
b	B1b	183	A2b	197	D3b	141	C4b	168	689	172.25	0.6875	0.47265625
c	C1c	199	B2c	160	A3c	150	D4c	96	605	151.25	-20.3125	412.597656
d	D1d	165	C2d	219	B3d	152	A4d	148	684	171	-0.5625	0.31640625
									2745			820.921875
									171.5625			273.640625

$$SS_{residual}=16336-6002-7037-3284=13$$

$$df_{residual} = 15 - 3 - 3 - 3 = 6$$

$$v_{residual} = \frac{SS_{residual}}{df_{residual}} = \frac{13}{6} = 2.17$$

これに対する、trap, day, corner の分散比 F を求める。

$$F_{\frac{trap}{residual}} = \frac{2001}{2.17} = 922$$

$$P_{trap} = 2.23E - 08$$

それ以上の値になることは、0.00000002 の確率しかない

$$F_{\frac{day}{residual}} = \frac{2346}{2.17} = 1081$$

$$P_{trap} = 1.37E - 08$$

それ以上の値になることは、0.00000001 の確率しかない

$$F_{\frac{corner}{residual}} = \frac{1095}{2.17} = 505$$

$$P_{trap} = 1.30E - 05$$

それ以上の値になることは、0.00001 の確率しかない

以上の結果を、分散分析表にまとめる。

表 3. シラスウナギのトラップの性能試験の結果 (分散分析表)

Factor	SS	df	V	F	P
Trap	6002	3	2001	922	0.00000002
Day	7037	3	2346	1081	0.00000001
Corner	3284	3	1095	505	0.00001
Residual	13	6	2.17		
Total	14336	15			

となる。つまり、トラップの効果も、実験日の効果も、角の効果もすべて、統計的に有意なのである。

注：余計な解説を入れると、Residual（残渣）の変動の大きさ（分散）と何かの要因による変動の大きさ（分散）の大きさを比べるという感覚が、分散分析の感覚で、この感覚が分散の計算式なんかよりはるかに重要で、残渣変動の中に説明できる変動が混ざっていると、分母が大きくなって、比が小さくなるから、有意差が出ないという話です。

結論として書けば、

「トラップも実験日も隅のいずれも、明確にシラスウナギの捕獲にかかわる大きな要因であった。したがって、ラップによる漁獲量は、周辺の環境やその日の天気等によって大きく変化すると考えられる。

その中であって、平均的な漁獲量 171 尾に対して。トラップ C は 27 尾、トラップ A は 7 尾、漁獲尾数を増加させ、トラップ D は 25 尾、トラップ B は 10 尾漁獲を低下させた。この違いは統計的に有意であった。」

となる。

こんな計算をして報告書を作って、依頼者に報告したら、「きちんと丁寧に、ラテン方格でやってくれたんですね。ありがとう。」と言われた。その時、ラテン方格という言葉を知って、いったい何だろうと思って、それからラテン方格というのを勉強した。

ラテン方格

ラテン方格とは何かというと、 n 行 $\times n$ 列のマス目に n 個の異なる記号が、各行、各列に一つずつ現れる表（正方行列）のことである。この記号の部分に、 n 個の実験条件の水準を割り当てるとというのが、ラテン方格を使った実験計画である。

ラテン方格の作り方は、いくらでも考えられるが、最も簡単の方法の一つは、図2の左に示したように、1 行目の第一列のものを、2 行目では最後の行にして、1 列づつ左にずらす。これを、2 行目以降も繰り返すという方法である。この方法だと、1 列目は 1 行目と同じ並び順（この例では A,B,C,D,E の順）になる。

図2．ラテン方格の作り方（例）とラテン方格の性質

A	B	C	D	E
B	C	D	E	A
C	D	E	A	B
D	E	A	B	C
E	A	B	C	D

C	D	E	A	B
D	E	A	B	C
A	B	C	D	E
E	A	B	C	D
B	C	D	E	A

E	B	D	C	A
A	C	E	D	B
C	E	B	A	D
B	D	A	E	C
D	A	C	B	E

ラテン方格の定義から考えると、当然のことなのだが、ラテン方格になっている行列の、行ごとにランダムに入れ替えたり（図の左から中への変換）、列ごとにランダムに入れ替えてたり（図の中から右への変換）しても、 n 個の異なる記号が、各行、各列に一つずつ現れるというラテン方格の持つ性質は変わらない。このようなやり方で、ラテン方格を作るとすると、その組み合わせの数はとても多いだろう、左の図の作業は、違うやり方もあるから、それらを含めると、もっと多いだろう。また、結果的におなじ並び方になるものもあるから、実際、どのくらいあるのか、私にはわからない。実験科学の場合、実験の順番や水槽の隅を意図的に動かすことはできないし、その影響を実験者がコントロールすることはできない。実験日の天候とか、水槽の位置とか、コントロールできない要因をブロック因子と呼ぶ。ブロック因子という名称は、実験計画法の乱塊法という、実験条件の構成法に由来する。実験計画法という学問は、主として、農学分野の圃場実験のやり方に関する研究によって発達してきた。実験圃場の環境は決して一様ではない。水はけ、土質、給水、日当たり、こういうものを、できるだけ一様に管理しようとしても、完全にすべての場所の環境を一様にはできない。さらに、たいていの場合、その環境には傾斜があって、どちらかが高く、どちらかが低いというような、傾きを持っている。その傾きの影響が実験結果に影響しないように、実験配置を考えなければならない。まず考えられるのは、その圃場内にランダムに、レベルの違う実験群を配置することだ。比較すべきものが A,B,C, の 4 つあるならば、圃場を 4 の倍

数になるように区分して、比較すべきものを図3のように、ランダムに区分することを考えるだろう。

図3。ランダムな実験配置

C	B	B	A
A	C	D	C
A	D	C	D
B	B	D	B

しかし、もし、図4に示すように、圃場に何か環境傾斜（水はけ、日当たり、作業員の注意の届きやすさなど）が、矢印の方向にあった場合には、傾斜に沿って、4ブロック

図4。乱塊法による配置

LowHigh

ブロック1	ブロック2	ブロック3	ブロック4
B	B	D	C
A	C	A	B
D	A	C	D
C	D	B	A

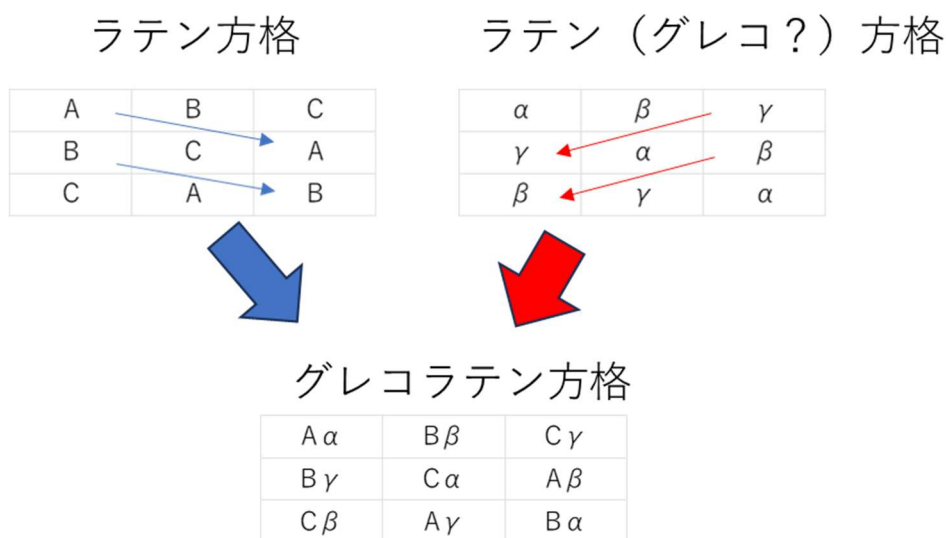
に分けて、ブロック内でランダムに配置する方が、確実に環境傾斜の影響を取り除くことができるだろう。現実的に考えると、環境傾斜がどのようになっているのか、わからないことの方が多いだろうから、どんな場合でも、ブロックを作って、ブロック内で、ランダムに配置しておくというのが、実践的な考え方だろう。ラテン方格では、ブロック因子が2つあって、それらが直交的な因子となっている。乱塊法的に考えれば、ラテン方格を造ったら、行ごと列ごとにランダムにシャッフルするべきだ。その方が、投稿論文などにした時には安全だろう。シラスウナギのトラップの性能試験を行ったときは、筆者も若くて未熟だったから、ラテン方格に並べるところまでは思いついたし、分散分析的にどのようにすればよいのかまでは考えたが、シャッフルするところまでは、考えが至らなかった。教えてもらっていないというのは言い訳で、不勉強というべきだろう。

オイラーとグレコ・ラテン方格

ところで、ラテン方格という名称は、あの大数学者のオイラー（Leonhard Euler 1707-1783）に由来するという説明が、ネットなどで出てくる。これを、ラテン方格を考えたのが、オイラーだと思っではいけない。オイラーは、その生涯で800以上の論文を書き、現代数学のあらゆる部分のその業績を残した大天才である。オイラーの公式 ($e^{i\theta} = \cos\theta + i\sin\theta$) なんか、普通の人には思いつかないだろう。それに比べると、ラテン方格はちょっと考えれば誰でも思いつく。彼はこんなものを論文にしていない。オイラーが論じたのは、グレコ・ラテン方格である。こっちは、かなり難しいパズルだ。

グレコ・ラテン方格とはどんなものかという説明の前に、例として3次のラテン方格を作ってみる。

図 5. 3×3のグレコ・ラテン方格の作り方（例）



第5図の左上のラテン方格は、前の行の先頭の記号を、次の行の最後にもっていくという作りで作った。右上のラテン方格は、前の行の最後をつぎの行の先頭に持ってくるという作り方をした。グレコ・ラテン方格とは、方形の行列の中に描かれているシンボルの特性が直交しているということである。この例で具体的に確認すると、A,B,C について、各行列の組み合わせを整理すると

A: A α A β A γ
B: B α B β B γ
C: C α C β C γ

となっていて、 α, β, γ について整理すると。

$$\alpha : A\alpha \quad B\alpha \quad C\alpha$$

$$\beta : A\beta \quad B\beta \quad C\beta$$

$$\gamma : A\gamma \quad B\gamma \quad C\gamma$$

となっていて、確かに、A,B,C と α, β, γ は直交している。

これを、表を作ったときの、ブロック因子まで含めて考えると

$$1aA\alpha \quad 1bB\beta \quad 1cC\gamma$$

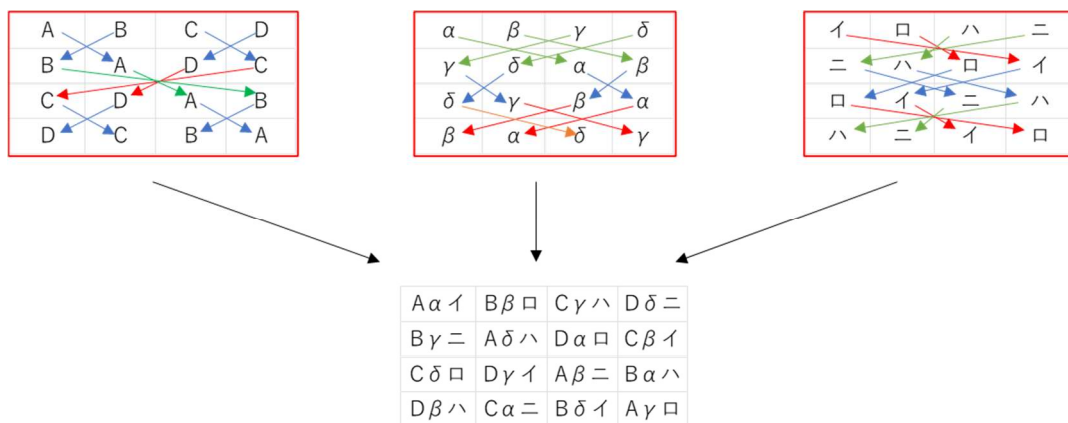
$$2aB\gamma \quad 2bC\alpha \quad 2cA\beta$$

$$3aC\beta \quad 3bA\gamma \quad 3cB\alpha$$

となって、1,2,3、a,b,c、A,B,C、 α, β, γ の4つの因子が、それぞれ、互いに直交関係になっていることが確認できるだろう。

何故、これをグレコ・ラテン方格というのかというと、このような性質を持つ n 次元の正方行列について興味を持って研究していたオイラーが、直交性を持つ2つの正方行列の説明で、一方の行列の因子を表すのにアルファベット（ローマ字）を使い、他方の行列の因子を表すのに、ギリシャ文字（Greek alphabet）で表したので、これを、グレコ・ラテン方格（Greco Latin Square）と呼んだ。何故、グレコ・ローマン方格と言わないのか知らないが、多分それだとレスリングと区別がつかなくなるからだろうと思う。いずれにしても、それで、合成する前の正方行列をラテン方格と呼ぶことになったということらしい。ラテン方格の方は誰でも思いつくが、グレコ・ラテン方格の方は、オイラーが苦労しただけあって難しく、 n 次のグレコ・ラテン方格の、一般的な作り方を考えるのは大変だ。それでも、3次のグレコ・ラテン方格は、いとも簡単に作ることができた。だが、このやり方では、4次のグレコラテン方格は作れない。4次のグレコ・ラテン方格を作るもっとも簡単なやり方は、左右2セットずつ、動かして、必要に応じて、左右を入れ替えるというやり方で、同じ動きにならないように、3つのラテン方格を作って、この3つを、重ね合わ

図6. 4×4のグレコ・ラテン方格の作り方（例）



せる。このやり方を、具体的に図6に示しておいた。

出来上がった、グレコ・ラテン方格の、因子が、相互に直交していることは、読者の方で、確かめてもらいたい。5×5のグレコラテン方格の例を図7に示しておく。

図7. 5×5のグレコ・ラテン方格の作り方 (例)

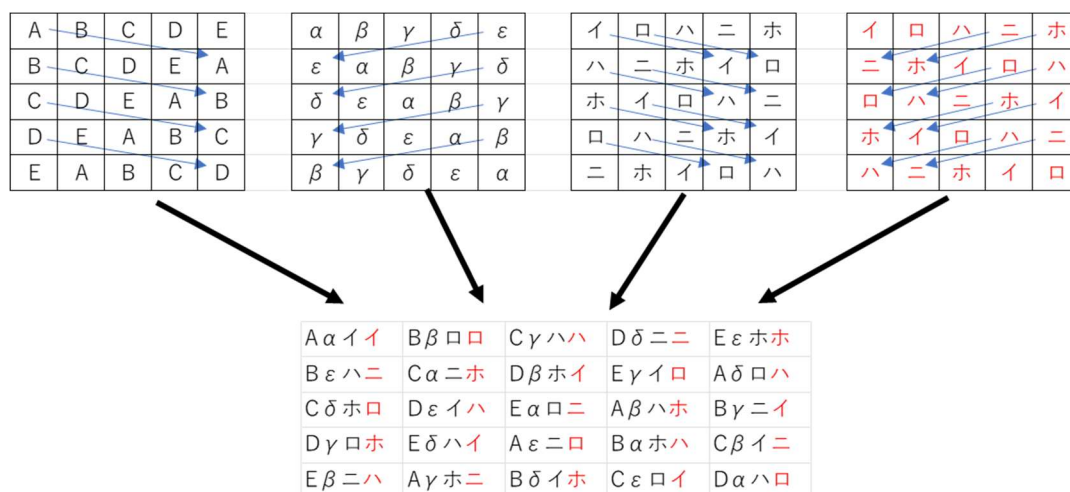


図7の作り方は、一例で、これ以外にもグレコラテン方格はいくらでもある。図7に示した方法では、3次の時と同様に、上の行の先頭列の記号を、下の行の最後部の列に持ってくるというやり方でラテン方格を作り、次に上の行の最後部の列の記号を、下の行の先頭の列にもっと来るという方法で次のラテン方格を作る。3つ目からは、2列をセットにして動かす。3つ目のラテン方格は、上の行の先頭から2番目までの列の記号を下の行の最後部から2番目までの列に動かし、4つ目のラテン方格は、上の列の最後尾から2番目までの列の記号を、下の列の先頭から2番目までに移動させる。

奇数次の方形行列の場合は、この方法で、グレコ・ラテン方格が作れるから、奇数次のグレコ・ラテン方格は、いくら次数が高くても作れるということは、筆者でもわかる。偶数次の方格だと、この方法では、グレコ・ラテン方格はつくれない。この方法で、4×4、6×6のラテン方格を作ってみると、すぐにわかる。偶数次の時は、図6に示した通りで、順序を入れ替えたりするので面倒なのだが、6×6ではこの方法でも、グレコ・ラテン方格は作れない。これもやってみればすぐわかる。奇数次に使える方法は、偶数次のグレコ・ラテン方格には使えない。偶数次用の方法も6×6の場合には使えない。ではどうしたらよいのかというと、筆者にはわからない。物の本によると、6×6のグレコ・ラテン方格は存在しないことが、すでに分かっているらしい。

それについては、面白い話がある。オイラーもこのことに気づいて、これをもっと一般化して、次のkのグレコ・ラテン方格($k = 4n + 2, n = 0, 1, 2, 3, \dots$)は存在しない(つまり、2×2, 6×6, 10×10, 14×14などのグレコラテン方格は作れない)と予想した。確かに2×2

ではグレコ・ラテン方格が作れないことは、やってみればわかる。 6×6 でグレコ・ラテン方格が作れないことは、古くからパズルのような遊びで知られていて、オイラーは、36人の士官問題 (36 officers problem) というゲーム (6 につの連隊があって、それぞれの連隊に6つの階級がある。行・列内で、連隊・階級が同じものがないように、方形の行列に並べるゲーム) の形でこれを紹介している。その後、オイラーの予測の妥当性の検証はされてこなかった (検証した人はいたかもしれないけれど、何せ、作業量が多くて大変だ。)。200 年近くたってから、1901 年に Gastn Tarry によって 6×6 では、グレコ・ラテン方格が作れないことが、確認された。Gastn Tarry はすべての組み合わせを枚挙して、直交性持つものが存在しないことを証明したらしい。大変な作業量だ。その後、コンピュータが発達してくると、その計算速度の速さを使って、しらみつぶしに調べていくことが可能になった。Gastn Tarry の仕事の 181 年後の 1959 年になって、 10×10 、 14×14 について、グレコ・ラテン方格が存在することが示され、オイラー予の予想は否定された。現在、 2×2 、 6×6 以外は、グレコ・ラテン方格があることが証明されている (どうやって証明したか知らない。多分、群論とかそういう難しい数学を使うのだと思う。))。

筆者も、グレコ・ラテン方格の作り方を一般化することはできなかったが、かのオイラーでさえ、グレコ・ラテン方格についての一般化し法則について誤った予測をした。

「オイラーができないことを、オイラができなくてもあたりまえジャン。」

実験計画法とグレコ・ラテン方格

グレコ・ラテン方格の話は面白いので、実験計画法という本来の目的を離れて、余計な話をしてしまった。しかし、知っておいて悪い話ではないだろう。実験計画法というのは、簡単に見えて手ごわいのだ。そう認識したうえで、本来の実験計画法に話を戻す。シラスウナギの例では、実験者が管理できないブロック因子が二つあった。その問題をどのように解決するかが、そもそもの課題だった。しかし、ブロック因子がなくて、すべてが、実験者が管理できる条件であったとしても、グリコ・ラテン方格には、実験計画法的に重要な意味がある。私は、水生生物の卵や精子などの生殖細胞や、受精卵、幼生などの生きたまま凍結保存する方法を開発していた。難しそうに思うかもしれないが、原理的には、生細胞から速やかに熱エネルギーを奪って、氷結晶を作らせないで、分子運動を制限（ガラス化する。）すればよいというだけのことだ。大雑把に言えば、小さな塊（体積に対して表面積が大きい）だと、凍結保存は容易だが、大きな塊（体積に対して表面積が小さい）は凍結が難しい。無茶苦茶簡単だ。だから、研究競争となるのは、生殖細胞（精子・卵子）、受精卵、幼生、成体など、どれだけ大きな塊を、生きたまま凍結保存できるかになる。このあたりの境界的な大きさの生物だと、細胞表面の特性や、水分含量、構造、体成分の組成などがかかわって来るから、それらに応じて、凍結速度、凍結保護物質の種類、凍結保護物質の濃度、保存液の浸透圧など、様々な凍結条件を検討していくことになる。凍結保存の成功度にかかわる因子が、5つあって、その水準を5段階に分けた場合、総当たり的に、あり得るすべての組み合わせについて、実験しようとする $5^5=3125$ 通りの組み合わせ数の実験をしなければならないことになる。実験材料、時間、費用がそれを許すならば、すべての組み合わせで、実験するかもしれない。その方が、交互作用も明らかになってよいだろう。しかし、実際の研究はそうはいかない。多くの場合、水生生物の成熟期は限られているから、その時期でなければ、精子も卵子も手に入らないし、手に入ったとしても、実験に十分な量ではないかもしれない。季節や量が限られている対象生物では、それだけの実験を終えるのに、何シーズンもかかってしまう。グレコ・ラテン方格を使えば、一要因に5つのレベルがある場合、 $5 \times 5 = 25$ 組みに、組合せの数を減らすことができる。これが、グレコ・ラテン方格の実験計画学的な意味である。

例えば、図7で作ったグレコ・ラテン方格をつかうと5つのレベルを持つ、少なくとも6個の因子を、互いに直交的に配置することができる（図8）。こ

図 8. 実験条件の組み合わせの数を減らすためにグレコラテン方格を使う方法（例）

① 5 次のグレコラテン方格

	a	b	c	d	e
1	A α イイ	B β □□	C γ ハハ	D δ ニニ	E ε ホホ
2	B ε ハニ	C α ニホ	D β ホイ	E γ イ□	A δ □ハ
3	C δ ホ□	D ε イハ	E α □ニ	A β ハホ	B γ ニイ
4	D γ □ホ	E δ ハイ	A ε ニ□	B α ホハ	C β イニ
5	E β ニハ	A γ ホニ	B δ イホ	C ε □イ	D α ハ□

② これをブロック要因の表に入れる。

	a	b	c	d	e
1	A α イイ	B β □□	C γ ハハ	D δ ニニ	E ε ホホ
2	B ε ハニ	C α ニホ	D β ホイ	E γ イ□	A δ □ハ
3	C δ ホ□	D ε イハ	E α □ニ	A β ハホ	B γ ニイ
4	D γ □ホ	E δ ハイ	A ε ニ□	B α ホハ	C β イニ
5	E β ニハ	A γ ホニ	B δ イホ	C ε □イ	D α ハ□

③ 列ごと、行ごとにランダムにシャッフルする。

	a	b	c	d	e
1	D β ホイ	A ε ニ□	E α □ニ	B δ イホ	C γ ハハ
2	C α ニホ	E δ ハイ	D ε イハ	A γ ホニ	B β □□
3	B ε ハニ	D γ □ホ	C δ ホ□	E β ニハ	A α イイ
4	A δ □ハ	C β イニ	B γ ニイ	D α ハ□	E ε ホホ
5	E γ イ□	B α ホハ	A β ハホ	C ε □イ	D δ ニニ

④ ブロック因子も他の因子と同様に扱って、条件の組み合わせをつくる。

1aD β ホイ	1bA ε ニ□	1cE α □ニ	1dB δ イホ	1eC γ ハハ
2aC α ニホ	2bE δ ハイ	2cD ε イハ	2dA γ ホニ	2eB β □□
3aB ε ハニ	3bD γ □ホ	3cC δ ホ□	3dE β ニハ	3eA α イイ
4aA δ □ハ	4bC β イニ	4cB γ ニイ	4dD α ハ□	4eE ε ホホ
5aE γ イ□	5bB α ホハ	5cA β ハホ	5dC ε □イ	5eD δ ニニ

⑤ 必要ならば、全体の並び方をシャッフルする。

それぞれの要因が、どんな水準で構成されているのかを整理したいのだが、これまでの議論で、使えそうな記号を使いつくしてしまったので、それらの因子を、ギリシャ文字の大文字で表すことにする。各因子のレベルは次のようになる。

A; 1, 2, 3, 4, 5

B; α, b, c, d, e

Γ ; A, B, C, D, E

Δ ; $\alpha, \beta, \gamma, \delta, \varepsilon$

E; $\iota, \rho, \pi, \sigma, \tau$

Z; $\iota, \rho, \pi, \sigma, \tau$

これらをグレコ・ラテン方格で配置して、分散分析表の形で、データを集約すると

要因	平方和	自由度	平均平方 (分散)	分散比 (F)
A	SSA	4	SSA/4	
B	SSB	4	SSB/4	
Γ	SS Γ	4	SS Γ /4	
Δ	SS Δ	4	SS Δ /4	
E	SSE	4	SSE/4	
Z	SSZ	4	SSZ/4	
r(残渣)	0	0		
total	SST	24	SSZ/24	

となる。データの数 (n) が 25 だから、全自由度は (n-1=24) で、一つの要因に含まれているレベルが (= 5) だから、それぞれの自由度は (k - 1 = 4) である。要因は 6 つだから、それらを合計した自由度は 24 であり、残渣の自由度は n-1-6(k-1)=0。残渣の平方和も 0 になる。つまり残渣がないので、残渣分散が計算できない。要因のレベルを k (k は 2, 6 以外の正の整数) とした場合、 $k \times k$ のグレコ・ラテン方格を作って、要因のレベルを配置すれば、 $k + 1$ 個の要因の効果を比較することができるのだが、 $k + 1$ 個にすると、残渣分散が残らないから、ランダムな変動と、要因による変動を比較して、要因による変動の統計的な有意性を論ずることができない (分散分析ができない)。かなり特殊な場合を除けば、科学実験で、有意差検定は必要だろう。だから、2, 6 以外の k 個のレベルのある要因の場合、k 個の要因までは、グレコ・ラテン方格を使って、組み合わせ数を減らすことができると覚えておいた方がよい。

グレコ・ラテン方格を使った、多要因の分析の分散分析の一般的な形を書いておく。

表 4. グレコ・ラテン方格の分散分析表

要因	平方和	自由度	平均平方 (分散)	分散比 (F)
F_1	SSF_1	$k - 1$	$VF_1 = SSF_1/k - 1$	$VF_1/VF_{residual}$
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$
F_n	SSF_n	$k - 1$	$VF_n = SSF_n/k - 1$	$VF_n/VF_{residual}$
$F_{residual}$	$SS_{residual}$	$df_{residual}$	$VF_{residual} = SS_{residual}/df_{residual}$	k
Total	SS_{total}	$k^2 - 1$		

$$SS_{residual} = SS_{total} - \sum_{i=1}^n SSF_i$$

$$df_{residual} = k^2 - 1 - n(k - 1) = (k - n + 1)(k - 1)$$

実験計画法の応用としての、コンジョイント分析

コンジョイント分析とは何かと問われると、よく知らない。これも、きちんと勉強していないからだ。それでも、商品の選択実験で、商品の内容をいくつかの特性に分けて、その徳江氏の価値を評価する、場合によっては、それを金銭的な評価としてあらわす。そういうことは、良く行われていて、私も見よう見まねでやったことがある（というか、データを渡されて分析だけを依頼される。）。選択実験をどのようにやるか、というのは、いろいろあって、ネットで、アンケート調査に参加してもらう代わりに、何かのポイントを与えるという形で、参加してもらうとか、謝金を払って、会場に来てもらい、説明を聞きながらアンケートに答えてもらうとか、いくらでも考えられるが、アンケートの回答項目が多いと、時間がかかって、いやになってしまい、まじめに答えてもらえなくなる。だから、質問項目はできるだけ少なくしたい。そんな時に、グレコ・ラテン方格は使える。

例えば、素焼きの団子とみたらし団子と磯辺団子があって、団子の大きさをどのくらいにして、1串何個にして、いくらの値段で売ればいいのかを考える。団子を作ったことも売ったこともないから、よくわからないから、たとえ話だと思って、読んでほしい。仮に、基準となる素焼き団子1個の値段を15円と考えて25円、35円という、値段設定にする。すると、以下のような、団子の種類と大きさ、非特使の個数、値段という要因があることになる

要因	0	1	2
A 種類	素焼き	みたらし	磯辺
B 大きさ	小	中	大
C 個数	2	3	4
D 値段	35 円/個	25 円/個	15 円/個

(他が昇順なのに、値段だけ降順にしたところに注意)

グレ・ラテン方格を知っていて、表計算をよくやっている人がこの表を見ると、これを、表4のようにまとめなくなる。

表4：団子の書か買う県的にかかわる要因と水準

	水準		
要因	0	1	2
A	A0	A1	A2
B	B0	B1	B2
C	C0	C1	C2
D	D0	D1	D2

要因

A 種類	A0:素焼き	A1:みたらし	A2:磯辺
B 大きさ	B0:小	B1:中	B2:大
C 個数	C0: 2	C1: 3	C2:4
D 値段	D0:35 円/個	D1: 25 円/個	D2:15 円/個

水準の番号を、1 から始めずに 0 から始めるのは、その方が後々、いろいろと便利だからである。表計算の時は、そういう習慣にしておいた方が良い。

これらを、総当たりの、すべての要因の組み合わせを作ると、 $3^4 = 81$ 個の組み合わせができてしまう。これらをすべてについて、買うか買わないかを問うことは、現実的でないだろう。回答者はすぐに嫌になって、でたために答え始める。選択肢を減らさなければならぬ。こういう時にグレコ・ラテン方格が登場する。

3 次のグレコ・ラテン方格は図 5 で作った。

図 5 で作ったグレコラテン方格

A α	B β	C γ
B γ	C α	A β
C β	A γ	B α

これを、ブロック因子の表に入れて、

	a	b	C
1	A α	B β	C γ
2	B γ	C α	A β
3	C β	A γ	B α

次のように書き換える

1aA α	1bB β	1cC γ
2aB γ	2bC α	2cA β
3aC β	3bA γ	3cB α

さらに、これを、1 列に並べると

1aA α
1bB β
1cC γ
2aB γ
2bC α
2cA β
3aC β
3bA γ
3cB α

こういうのを、 3×3 の直交表 ($L_9^{3^4}$) という (それぞれが、他の要因の水準を一つずつ含んでいることを確認してもらいたい。)

これを利用して、1文字目を、団子の種類 A として、1:A0 素焼き、2:A1 みたらし、3:A2 磯辺、2文字目を、団子の大きさ B として、a:B0 小、b:B1 中、c:B2 大、3文字目を1串の団子の数 C として、A:C0 2つ、B:C1 3つ、C:C2 4つとし、4文字目を価格 D として、 α :D0 15 円/個、 β :D1 25 円/個、 γ :D2 35 円/個とすると

表 5. 団子の商品スペックリスト

A0:素焼き	B0:小	C0:2 個	D0: 35 円/個
A0:素焼き	B1:中	C1:3 個	D1: 25 円/個
A0:素焼き	B2:大	C2:4 個	D2: 15 円/個
A1:みたらし	B0:小	C1:3 個	D2: 15 円/個
A1:みたらし	B1:中	C2:4 個	D0: 35 円/個
A1:みたらし	B2:大	C0:2 個	D1: 25 円/個
A2:磯辺	B0:小	C2:4 個	D1 25 円/個
A2:磯辺	B1:中	C0:2 個	D2 15 円/個
A2:磯辺	B1:大	C1:3 個	D0 35 円/個

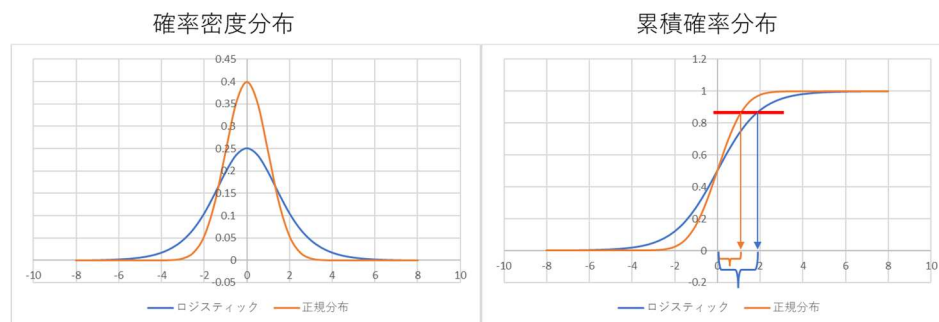
これが、コンジョイント分析の、商品の組み合わせを作るときの作業なのだが、一番上の組み合わせは、値段が高く、商品としての魅力も最も小さいと考えられるから、これを、基準と考えて、選択してもらう商品の中に含めず、8つの組み合わせにすることが一般的だろうと思う。これが、基本的なやり方なのだが、実際には、要因の水準の数が違って、正方形にならず、グレコ・ラテン方格が使えないことがある。これは実践的な問題だが、ダミーの水準を入れるなどして、無理無理、正方形にするという方法も考えられるが、実際に、どのようにしているのか知らない。これについては、あとから考える。

とは言ったものの、正直に告白すると、この先どのようにするのか、あまり詳しくない。というか知らない。詳しい人に聞くか、何か解説書のようなものがあるだろうから、それを参照のこと。以下は、私だったら、こうするだろうという意見のようなものだ。教科書を読むときに、何を考えているのか参考になるかもしれない。

最も、実際の買い物に近いのは、この8つの商品の一つずつ、それを買うか買わないか問うというやり方だろう。これを集団でやると、その商品が選択される率がわかる。これは、その集団にその商品を売ろうとした場合、売れる確率と考えて良いだろう。確率分布というのは、もちろん線形ではない。その確率となる元の値を知るにはどうしたらよいのかというのが、次の課題だ。私の言語能力では、具体的に何するのか説明するのが難しいので、図で説明する。図9の左が赤い曲線が正規分布の確率密度分布である。右の図の赤い曲線

が正規分布の累積確率曲線で、図には累積密度 $p = 0.85$ ぐらいのところに赤い線を引いた。赤い線と、正規分布の累積密度関数の曲線（青い線）との交点が、横軸上で原点 0 からどのくらい離れているのか、その距離が知りたいということである。つまり、ガウス関数（正規関数）の逆演算で、Excel の関数だと、中心が 0、標準偏差が 1、正規分布確率 p の逆演算 $\text{norminv}(p,0,1)$ である。これをプロビット変換という。

図 9.正規分布とロジスティック分布



プロビット変換を使っても良いのだが、よく使われるのは、ロジット変換である。ロジットというのはオッズの対数である。オッズは、あることが起きることの確率と、あることが起こらない確率の比で、賭け事で勝つ可能性を比較したり、病気になるリスクの比較に使われる。ロジットとは、その対数のことである。

$$\text{オッズ: } \frac{p}{1-p}$$

$$\text{ロジット: } \log_e \frac{p}{1-p}$$

ロジット t として、 p を t で表す。

$$t = \log_e \frac{p}{1-p}$$

$$e^t = \frac{p}{1-p}$$

$$e^t(1-p) = p$$

$$e^t = p + pe^t = p(1 + e^t)$$

$$p = \frac{e^t}{1 + e^t} = \frac{1}{e^{-t} + 1} = \frac{1}{1 + e^{-t}}$$

これが、標準標準ロジスティック曲線で、図 9 の右側の青い線だ。この曲線は

$$n = \frac{k}{1 + e^{-bt}}$$

の形で、生態学などでは、個体群の増殖曲線としてよく使われる。私としてはなじみのある数式である。これを微分すれば

$$\frac{dp}{dt} = \frac{e^{-t}}{(1+e^{-t})^2}$$

となって、左側の図の青い線になる。正規分布とロジスティック分布は、よく似た分布で、

$$\frac{e^{-\frac{t}{s}}}{s\left(1+e^{-\frac{t}{s}}\right)^2}$$

として、Sを調節して、正規分布と重ねてみると、ほとんど変わらない。

では、なぜ、プロビット変換ではなくて、ロジット変換が使われるのかというのは、よくわからない。多分、ロジット変換の方が式の形が単純で、計算しやすいからだと思う。確率変数を線形にして回帰分析に持っていきたいときに、このロジット変換をよく使う。これをロジスティック回帰という。被験者全員による選択率を買ってもらえる確率と考えて、これをロジット変換したものを、目的変数として、表5のスペックリストを、説明変数として、分析するという、作戦である。

次のような式が出来上がる。

$$L_i = \alpha + \beta_0 A0_i + \beta_1 A1_i + \beta_2 A2_i + \beta_3 B0_i + \beta_4 B1_i + \beta_5 B2_i + \beta_6 C0_i + \beta_7 C1_i + \beta_8 C2_i + \beta_9 D0_i + \beta_{10} D1_i + \beta_{11} D2_i$$

L_i はそれぞれのスペックのロジット値、 α は切片、 β_i はそれぞれの説明変数の係数

スペックリストは変数は金額を除いて数値変数ではないから、0、1のダミー変数にする
と、スペックS0からスペックS8の説明変数のリストは表6のようになる。

表6.スペックごとの説明変数

	A0	A1	A2	B0	B1	B2	C0	C1	C2	D0	D1	D2
S0	1	0	0	1	0	0	1	0	0	1	0	0
S1	1	0	0	0	1	0	0	1	0	0	1	0
S2	1	0	0	0	0	1	0	0	1	0	0	1
S3	0	1	0	1	0	0	0	1	0	0	0	1
S4	0	1	0	0	1	0	0	0	1	1	0	0
S5	0	1	0	0	0	1	1	0	0	0	1	0
S6	0	0	1	1	0	0	0	0	1	0	1	0
S7	0	0	1	0	1	0	1	0	0	0	0	1
S8	0	0	1	0	0	1	0	1	0	1	0	0

ここで、困ったことに気づく。S0を除いたロジット値の数は8個なのに、説明変数が12個もある、これでは回帰はできない。そこで、テクニックを使う。そもそも、A0、B0、C0は基準で、それに該当したとしても、0しか入らないから、式から取り除いても構わないだろう。そうすると、説明変数は9個になる。これでも、一つ多い、そこで、最後の、金額は数値変数だから、一つにまとめて、金額を入れることにする。そうすると、表7ができて、説明変数の数が7つとなって、これで重回帰ができることになる。このまま重回帰しても問題ないのだが、一応、一個35円を基準にしたから、価格のところは基準との差にしておいた方が、論理的かもしれない(重回帰の結果は正負が入れ替わるだけにすぎないけれど)。その方が、品が良い感じがする。

表7.基準 (S0, A0, B0,C0) を取り除いた変数リスト

	A1	A2	B1	B2	C1	C2	D
S1	0	0	1	0	1	0	25
S2	0	0	0	1	0	1	15
S3	1	0	0	0	1	0	15
S4	1	0	1	0	0	1	35
S5	1	0	0	1	0	0	25
S6	0	1	0	0	0	1	25
S7	0	1	1	0	0	0	15
S8	0	1	0	1	1	0	35

表8. 価格を基準との差に値で表現する。

	A1	A2	B1	B2	C1	C2	D
S1	0	0	1	0	1	0	-10
S2	0	0	0	1	0	1	-20
S3	1	0	0	0	1	0	20
S4	1	0	1	0	0	1	0
S5	1	0	0	1	0	0	-10
S6	0	1	0	0	0	1	-10
S7	0	1	1	0	0	0	-20
S8	0	1	0	1	1	0	0

これで式を作ると、

$$L_i = \alpha + \beta_1 A1_i + \beta_2 A2_i + \beta_3 B1_i + \beta_4 B2_i + \beta_5 C1_i + \beta_6 C2_i + \beta_7 D$$

という形になって、具体的には

$$\begin{aligned}
 L_1 &= \alpha && + \beta_3 && + \beta_5 && - 10\beta_7 \\
 L_2 &= \alpha && && + \beta_4 && + \beta_6 && - 20\beta_7 \\
 L_3 &= \alpha &+ \beta_1 && && \beta_5 && - 20\beta_7 \\
 L_4 &= \alpha &+ \beta_1 && + \beta_3 && && + \beta_6 \\
 L_5 &= \alpha &+ \beta_1 && && + \beta_4 && - 10\beta_7 \\
 L_6 &= \alpha && + \beta_2 && && + \beta_6 && - 10\beta_7 \\
 L_7 &= \alpha && + \beta_2 &+ \beta_3 && && - 20\beta_7 \\
 L_8 &= \alpha && + \beta_2 && + \beta_4 &+ \beta_5
 \end{aligned}$$

という式の回帰分析をすればよいことになる。

被験者の回答の個票から集計表を作って、表9のロジット (L) を計算する。これができたら、エクセルで次のように被説明変数と説明変数を指定して、回帰分析すればよい（実際には、R か Python を使うと思うから必要ないけど、エクセルだと、上のところにデータをクリックすると、右の方にデータ分析があって、そこを開けると、一番下に回帰分析というのが出てくる。）

表 9 コンジョイント分析の集計表

被験者	S1	S2	S3	S4	S5	S6	S7	S8
1	0	0	1	0	1	1	1	1
2	0	0	1	1	0	1	1	0
3	0	1	1	0	0	1	0	0
⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮
n	0		0	1	0	1	1	0
Sum	2	5	30	25	10	50	70	30
n	100	100	100	100	100	100	100	100
選択率	0.02	0.05	0.3	0.25	0.1	0.5	0.7	0.3
オッズ	0.020408	0.052632	0.428571	0.333333	0.111111	1	2.333333	0.428571
ロジット	-3.89182	-2.94444	-0.8473	-1.09861	-2.19722	0	0.847298	-0.8473

図 10 回帰分析のためのエクセルシートと解

-3.89182		0	0	1	0	1	0	-10
-2.94444		0	0	0	1	0	1	-20
-0.8473		1	0	0	0	1	0	-20
-1.09861		1	0	1	0	0	1	0
-2.19722		1	0	0	1	0	0	-10
0		0	1	0	0	0	1	-10
0.847298		0	1	1	0	0	0	-20
-0.8473		0	1	0	1	1	0	0

概要								
回帰統計								
重相関 R	1							
重決定 R2	1							
補正 R2	65535							
標準誤差	0							
観測数	8							
分散分析表								
	自由度	変動	分散	割された分散	有意 F			
回帰	7	16.93609	2.419441	#NUM!	#NUM!			
残差	0	0	65535					
合計	7	16.93609						
	係数	標準誤差	t	P-値	下限 95%	上限 95%	下限 95.0%	上限 95.0%
切片	-7.28774	0	65535	#NUM!	#VALUE!	#VALUE!	-7.28774	-7.28774
X 値 1	3.326955	0	65535	#NUM!	3.326955	3.326955	3.326955	3.326955
X 値 2	4.708	0	65535	#NUM!	4.708	4.708	4.708	4.708
X 値 3	1.330635	0	65535	#NUM!	1.330635	1.330635	1.330635	1.330635
X 値 4	0.715359	0	65535	#NUM!	0.715359	0.715359	0.715359	0.715359
X 値 5	1.017084	0	65535	#NUM!	1.017084	1.017084	1.017084	1.017084
X 値 6	1.531539	0	65535	#NUM!	1.531539	1.531539	1.531539	1.531539
X 値 7	-0.10482	0	65535	#NUM!	-0.10482	-0.10482	-0.10482	-0.10482

それぞれの係数を、支払金額の係数で割れば、それぞれの団子の特性の一つ一つの、基準との差が金額で出てくる。

表 10.基準との金額差の計算

	係数			金額の差
素焼き	0	-0.10482	0	0
みたらし	3.326955	-0.10482	-31.7396	31.73964
磯辺	4.708	-0.10482	-44.915	44.91501
小	0	-0.10482	0	0
中	1.330635	-0.10482	-12.6944	12.69445
大	0.715359	-0.10482	-6.82463	6.82463
2 個串	0	-0.10482	0	0
3 個串	1.017084	-0.10482	-9.70313	9.703128
4 個串	1.531539	-0.10482	-14.6111	14.6111

つまり、みたらしは、素焼きより 32 円、磯辺は 45 円高い。中は小よりも、13 円高いが、それ以上大きくなっても根は上がらず、大と小の差は 7 円。3 個串は 2 個串よりも 10 円高く、4 個串は、2 個串よりも 15 円高いという結論になる。

これで、一つ一つの要素について、価値を金額で表すことはできたが、この計算結果の妥当性には、少し不安がある。自由度が 7 で、求める係数が 6 だから、回帰というよりは、連立方程式を解いているような形になって、図 10 では、標準誤差が 0 になって、分散分析ができず、p が計算できない。そう考えると、もともと、3 水準の時に、4 要因で検討するのは、無理があるかもしれない。3 要因にして、3 水準で回答してもらった方が、分散分析が可能で、統計的に確かな値が得られると判断すべきかもしれない。団子の大小などは、適当だから、そんなもの最初からとってしまった方が良いだろう。

最後に、途中で思考を放棄して、先送りした問題について考える。要因の数が多かったり、水準の数が違って、グレコ・ラテン方格が作れない時にどうすべきなのかという問題だ。一人で、独自に考えると作業量が多すぎるから、R のお世話になる。Conjoint という、そのものずばりのパッケージを使う。要因と水準を与えると、直交的な組み合わせを提言してくる。

パッケージのインストールと読み込み

```
install.packages("conjoint")
```

```
library(conjoint)
```

初めに機能を確認するために、我々が使った、4 要因、3 水準について、組み合わせを提言させる。

属性と水準の設定

```
experiment <- expand.grid(
```

```
  A = c("A0", "A1", "A2"),
```

```
  B = c("B0", "B1", "B2"),
```

```
  C = c("C0", "C1", "C2"),
```

```
  D = c("D0", "D1", "D2")
```

```
)
```

```
# 直交表の作成
```

```
design <- caFactorialDesign(data = experiment, type = "orthogonal")
```

```
print(design)
```

```
      A  B  C  D
5  A1 B1 C0 D0
10 A0 B0 C1 D0
27 A2 B2 C2 D0
34 A0 B2 C0 D1
42 A2 B1 C1 D1
47 A1 B0 C2 D1
57 A2 B0 C0 D2
71 A1 B2 C1 D2
76 A0 B1 C2 D2
```

確かに、グレコラテン方格で回答してくる。3 要因 2 水準にすると

```
experiment <- expand.grid(
```

```
+   A = c("A0", "A1", "A2"),
```

```
+   B = c("B0", "B1", "B2"),
```

```
+   C = c("C0", "C1", "C2")
```

```
+ )
```

```
>
```

```
> # 直交表の作成
```

```
> design <- caFactorialDesign(data = experiment, type = "orthogonal")
```

```
> print(design)
```

```
      A  B  C
2  A1 B0 C0
6  A2 B1 C0
7  A0 B2 C0
12 A2 B0 C1
13 A0 B1 C1
17 A1 B2 C1
19 A0 B0 C2
27 A2 B2 C2
```

この例でもグレコ・ラテン方格で直交表を返してくるが、A1、B1、C2 を基準として、取り除いている。5 条件 2 水準だと、

```

experiment <- expand.grid(
+   A = c("A0", "A1"),
+   B= c("B0", "B1"),
+   C= c("C0", "C1"),
+   D= c("D0", "D1"),
+   E= c("E0", "E1")
+ )
>
> # 直交表の作成
> design <- caFactorialDesign(data = experiment, type = "orthogonal")
> print(design)
      A  B  C  D  E
3  A0 B1 C0 D0 E0
6  A1 B0 C1 D0 E0
9  A0 B0 C0 D1 E0
16 A1 B1 C1 D1 E0
18 A1 B0 C0 D0 E1
23 A0 B1 C1 D0 E1
28 A1 B1 C0 D1 E1
29 A0 B0 C1 D1 E1

```

グレコ・ラテン方格ではないが、8通りの直交的な組み合わせをかえしてくる。確かに、総当たりにした時の32通りよりも少ない。

```

> experiment <- expand.grid(
+   A = c("A0", "A1", "A2"),
+   B= c("B0", "B1"),
+   C= c("C0", "C1"),
+   D= c("D0", "D1")
+ )
>
> # 直交表の作成
> design <- caFactorialDesign(data = experiment, type = "orthogonal")
> print(design)
      A  B  C  D
2  A1 B0 C0 D0
6  A2 B1 C0 D0
9  A2 B0 C1 D0

```

```
10 A0 B1 C1 D0
13 A0 B0 C0 D1
18 A2 B1 C0 D1
21 A2 B0 C1 D1
23 A1 B1 C1 D1
```

3 水準 * 2 水準 * 2 水準 * 2 水準でも

```
experiment <- expand.grid(
+   A = c("A0", "A1", "A2"),
+   B = c("B0", "B1"),
+   C = c("C0", "C1"),
+   D = c("D0", "D1")
+ )
>
> # 直交表の作成
> design <- caFactorialDesign(data = experiment, type = "orthogonal")
> print(design)
```

```
   A  B  C  D
2  A1 B0 C0 D0
6  A2 B1 C0 D0
9  A2 B0 C1 D0
10 A0 B1 C1 D0
13 A0 B0 C0 D1
18 A2 B1 C0 D1
21 A2 B0 C1 D1
23 A1 B1 C1 D1
```

8通りで収まり、

3 水準 * 3 水準 * 2 水準 * 2 水準では

```
experiment <- expand.grid(
+   A = c("A0", "A1", "A2"),
+   B = c("B0", "B1", "B2"),
+   C = c("C0", "C1"),
+   D = c("D0", "D1")
+ )
> # 直交表の作成
> design <- caFactorialDesign(data = experiment, type = "orthogonal")
> print(design)
```

	A	B	C	D
2	A1	B0	C0	D0
6	A2	B1	C0	D0
16	A0	B2	C1	D0
19	A0	B0	C0	D1
22	A0	B1	C0	D1
26	A1	B2	C0	D1
27	A2	B2	C0	D1
30	A2	B0	C1	D1
32	A1	B1	C1	D1

>

9通りの組み合わせで、カバーできるのだが、D1、C0が6回登場するのに、D0、C1が3回しか登場しないというように、対称性という意味では、少し不満があるかもしれない。式の数を変数の数より多ければ、回帰分析は可能だから、これで計算できないわけではないから、これでよいという考え方もあるだろう。

グレコ・ラテン方格にならない場合は、できるだけ偏りがなく、できるだけ、少ない組み合わせ数で、スペック表を作らなければならないので、大変だが、Rのconjointなど、パッケージに組み込まれた機能を使えば、おそらく問題はないだろう。ということで、コンジョイント分析についての話は、これでおしまいにする。

BWS に関する考

BWS (Best worst scaling) 分析でも、コンジョイントと同様に、質問の組み合わせが問題になる。こちらは、さらに厄介なパズルを解かなくてはならない。普通はそんなパズルを解いている時間はないから、R か Python のパッケージのお世話になることになる。質問項目の設計法などどうでもよいと言えば、どうでも良いのかもしれない。しかし、何をしているのかを知っている方が、設計を誤ったり、分析結果の解釈を間違えることが少ないだろう。ということで、BWS の質問項目の設計について、考察する。と言っても、BWS というのを知ったのはごく最近のことだから、あまり詳しくない。例によって、自分勝手に考察する。

まず、BWS とは何かだが、質問の仕方によって、Case 1, Case2, Case3 の 3 つがあるのだが、ここで取り上げるのは、Case2 で、次のような質問をする。

あなたが、旅行の目的地を選定するときに、どんなことを重視しますか。次の 4 つの中から、最も重視すること、最も重視しないことを選んで、チェックしてください

質問 1

最も重視する		最も重視しない
()	景観の美しさ	()
()	食べ物のおいしさ	()
()	歴史的建造物	()
()	宿泊施設	()

質問 2

最も重視する		最も重視しない
()	食べ物のおいしさ	()
()	アクセスの良さ	()
()	地元の人の人柄	()
()	費用	()

こんな質問をいくつか回答してもらおう。BWS の良さは、それぞれの項目評価採点してもらおうと、多くのひとは、曖昧で、中間的なこたえをしがちであるが、BWS では、比較になっているので、はっきりとした答えになり、何回か繰り返し比較されるので、回答者にあまり負担をかけずに、明確な数値的評価が得られることである。BWS をやるときには、いくつの項目を比較するのか、評価してもらおう項目の数と、質問の数の関係がどうすべきかを考えなくてはならない。これが、かなり難しいパズルになっている。このパズルの解法を一般化して、答えてみたかったのだが、できなかったのが、例によって、具体的な例について考える。まず、設計のルールを確認しておく、すべての選択肢が、全質問

で、同じ数だけ、比較対象になっていること、必ず、一つ一つの選択肢が、のすべての選択肢と1回ずつ比較されていることである。

実際に、このようなことが可能な、選択肢の数と一回に比較する選択肢数と質問数の関係は、どうなっているのか、どのようにして、質問票を設計するのか、よくわからない。ネットで調べても、作り方について、詳述したものは見つからなかった。Chat GPT で調べたら、BIBD (Balanced Incomplete Block Designs) でやると、書いてあったが、その内容については、よくわからない返事が返ってくるだけで、まともな答えは得られなかった。そこで、BIBD で検索したら、

[https://math.libretexts.org/Bookshelves/Combinatorics_and_Discrete_Mathematics/Combinatorics_\(Morris\)/04%3A_Design_Theory/17%3A_Designs/17.01%3A_Balanced_Incomplete_Block_Designs_\(BIBD\)](https://math.libretexts.org/Bookshelves/Combinatorics_and_Discrete_Mathematics/Combinatorics_(Morris)/04%3A_Design_Theory/17%3A_Designs/17.01%3A_Balanced_Incomplete_Block_Designs_(BIBD))

というのが出てきたが、実に間抜けなことが書いてある。病気に対する処置法が7つあって、その処置法について総当たりに比較実験すると、どういう組み立て方をするのか知らないが(あれとこれをやって、あれをやらないとか、そういうことだと思う)、 $2^7 - 1 = 127$ 匹のマウスが必要だが、頭を使って、7つの処理から、2つずつ取り出して、比較実験しても、7つから2つを取り出す組み合わせの総数 ${}_7C_2 = 21$ (アメリカ式に書くと $\binom{7}{2} = 21$) 匹のマウスが必要になる。マウスが7匹しかいない時にどうするのかという、問いかけがあって、その答えとして、7匹のマウスの一匹ずつに、1から7までの処理法の3つを次のように選んで、処理すればよい。という答えが示される。

$$\{1,2,3\} \{1,4,5\} \{1,6,7\} \{2,4,6\} \{2,5,7\} \{3,4,7\} \{3,5,6\}$$

こういうのが、BIBD だというのだが、その後、一般化して、どういうときにこれが可能なのかという、ことを数式を使って解説がついている。数学的には確かに、面白い問題かもしれないが、知りたいのは、これをどうやって作るのかだ。それが書いてない。しょうがないから、自分でやる。

最初の考えたのは、ラテン方格のところでやった、前の列の最後尾の者を前にもってくるというやりかただ。これならば。間違いなく網羅的に全部の者が含まれて、その回数が均等にするということでこれを、3 選択肢から選ぶ場合は次のような組み合わせになる。

図 1 1. ラテン方格的なやり方、3 選択肢から選ぶ場合

1,2,3

1	2	3
2	3	1
3	1	2

1,2,3,4

1	2	3
2	3	4
3	4	1
4	1	2

1,2,3,4,5

1	2	3
2	3	4
3	4	5
4	5	1
5	1	2

1,2,3,4,5,

1	2	3
2	3	4
3	4	5
4	5	6
5	6	1
6	1	2

このやり方が、全く駄目だということはすぐにわかる。1 番左は、全選択肢数が 3 の場合だが、確かにラテン方格になっているが、比較だから、並び順は関係ないから、どの行の内容も同じで、質問を 3 回繰り返す意味がない。左から 2 番目は列が足りないので、ラテン方格ではないが、網羅的にはなっている。しかし、それぞれが、2 回も同じものと比較されて、大変無駄な構造になっている。右から 2 番目は、全選択肢が 5 つの場合だ。この例では、1 は 3,4 とは、1 回だけ比較されているが、2,5 とは 2 回も比較されている。一番右の例では、1 は 2,6 とは 2 回も比較されているのに、3,5,とは 1 回だけ、4,とは 1 回も比較されていない。このやり方では、条件を満たすものは作れない。もう少し丁寧に、一つずつ条件を満たすように作ってみる。

図 12. BIBD 的やりかた

Diagram illustrating three sets of 3x3 grids with numbers 1-4, 1-5, and 1-6. Each set consists of a 3x3 grid on the left and a 1x3 grid on the right. The 1x3 grids are highlighted in yellow.

Set 1 (Numbers 1-4):

1	2	3
2	4	
3	4	

1	4	
---	---	--

Set 2 (Numbers 1-5):

1	2	3
2	4	
3	4	

1	4	5
---	---	---

Set 3 (Numbers 1-6):

1	2	3
2	4	5
3	4	6

1	4	5
2	5	6
3	5	

1	6	
---	---	--

Diagram illustrating a single set of 3x3 grids with numbers 1-7. The 3x3 grid is on the left and the 1x3 grid is on the right. The 1x3 grid is highlighted in yellow.

Set 4 (Numbers 1-7):

1	2	3
2	4	6
3	4	7

1	4	5
2	5	7
3	5	6

1	6	7
---	---	---

図 12 に示したのは、1 行目は、先頭の 1 と比べることが可能な組み合わせを順番に、横に並べて書く。1 と比べるものがなくなったら、2 行目に移って、2 と比べることが可能なものを、順番に書く、2 と比べるものがなくなったら、3 行目に移って、3 と比べることが可能なものを書くというやりかただ。

図 12 の上の一番左は、全選択肢が 1,2,3,4 の 4 つの場合だが、一度に比較できるのは 3 つだから、1,2,3 を比較したのちに、となりのブロックに移って、1 と 4 を比較しなければならないが、それ以上比べるものがないから、このブロックの最後のボックスが埋まらない。2 と 4、3 と 4 を比べるブロックの最後のボックスも埋まらない。中央は全選択肢が 1,2,3,4,5 の場合だが、1 行目は、横に並んだ、2 つのブロックが埋まるが、2 行目、3 行目の、2 と 4、3 と 4 のブロックは、残りのボックスが埋まらない（4 と 5 はすでに、1 行目の 2 番目のブロックで比較されている。）。一番右は、全選択肢が 1,2,3,4,5,6 の場合だが、1 行目は 3 つのブロックが埋まり、2 行目も 2 つのブロックが埋まるのだが、3 行目のブロックの 3 つ目のボックスが埋まらない。つまり、3 選択肢の場合、全選択肢数 6 までは、条件を満たすものはない。図 12 の下の段は、全選択肢が 7 の場合である。1 行目、2 行目、3 行目のすべてのブロックが埋まっている。この組み合わせで、すべての選択肢が、全質問で、同じ数だけ、比較対象になっていること、必ず、一つ一つの選択肢が、のすべての選択肢と 1 回ずつ比較されていることという、条件を満たしていることを、確認してもらいたい。面倒ならば、ネットの記事で紹介された組み合わせと、見比べてくても良い。

図 12 に示したやり方で、選択肢が 4 の場合について、可能な組み合わせを作ったので、図 13 に示した。

図 13. 一度に比較するものが 4 つの場合の、BIBD

	1	2	3	4	5	6	7	8	9	10	11	12	13						
k	1	2	3	4		1	5	6	7		1	8	9	10		1	11	12	13
	2	5	8	13		2	6	9	11		2	7	10	12					
	3	5	9	12		3	6	10	13		3	7	8	11					
	4	5	10	11		4	6	8	12		4	7	9	13					

</

一度に比べる選択肢の数を、4 にすると、総選択肢数 13 の時に、条件を満たす全組み合わせができる。図 12 で選択肢が 4 になると、行数が 4 になる。列の数をブロック単位で数えると、1 行目が 4、それ以下が 3 である。だから、組み合わせた質問の数は $4 \times 3 + 1 = 13$ となる。これを一般化して、一度に比べる選択肢数を k とすると、総選択肢数は

$$N = k(k-1) + 1$$

であると書きたくなるのだが、やめておく。何しろ、あのオイラーだってこういうことを一般化して間違えた。もっと高次になると、何が起こるのかよく知らない。もっとも、我々が、使うのは、せいぜい 4 つの選択肢ぐらいだから、この式を一般化して覚えておいても、あまり問題はないだろう。ということで、知る人ぞ知る知識ということでよいだろう。それよりは、作り方の方を身に着けた方が良い。特に、ブルーでマークしたところ

テクニックは、グレコラテン方格のところで使っている。

以上が、BWS について、私が考えたことだ。BWS の結果の得点化などについては、その辺にいくらでも、解説が転がっているから、そちらを参考にしてもらいたい。