

クジ比較における反事実への懸念と認知ポテンシャル

Worrying About Counterfactuals in the Lottery Comparison and the Cognitive Potential

犬童 健良[†]

Kenryo Indo

[†] 関東学園大学

Kanto Gakuen University

kindo@kanto-gakuen.ac.jp

概要

本研究では反事実がクジ選におけるリスク態度に与える変化として認知的ポテンシャルをモデル化し、仮想的なクジのペアから選択する質問において選ばなかったクジの結果がどの程度気になるかを聞いた実証的データを用いて分析した。

キーワード: 反事実, 気になる程度, アレの背理, 12-6ポテンシャル

1. はじめに

選択問題においてある一つの代替案が選ばれるとき、反事実(counterfactuals)としての選ばれなかった案はどのように影響しているだろうか。仮想的なクジ選択において選ばなかったクジの結果がどの程度気になるかについて、反事実条件文を用いて直接回答者に聞く方法は、[2]で最初に報告された。[2]ではアレの背理(Allais' paradox) [1]を対象としたリスク態度への影響が調べられ、後悔・安堵に基づく解釈が試みられた。実験データは[3]の付録として公開されている。また同実験データを再解釈した[3][4]では、認知的な最適化のメカニズムが推測された。[5][6][7]はデータ分析によってアレの背理と関連する阻止関係を見出し、ネットワークの混雑現象との類比が考察された。先行研究のモデルでは主な変数として気になる度合いの差が用いられたが、本研究では[3]のデータから反事実の気になる度合いを直接用いてアレの背理の発生を阻止するマーカーを抽出した。阻止マーカーによってよりシンプルな予測が可能となると同時に、化学ポテンシャルとの類比からリスク態度に影響する相反する認知的要因が示唆される。

2. アレの背理と反事実的結果の質問

先行研究[2]では、アレの背理の例題を用いたウェブアンケート方式の予備的実験が行われた。この実験は、仮想的なクジペアの選択問題を問うと同時に、選んだクジの可能な結果と選ばなかったクジの結果の組み合わせ

わせがそれぞれの程度気になるかを質問した。

より具体的には表 1a に示した 4 つのペア比較問題の A~H の計 8 オプションにおけるペア比較の選択、および表 1b に示した仮想的に選んだオプションの各結果 X_i と選ばなかったオプションの各結果 Y_j の間の組み合わせ計 26 ペアについての気になる程度を、表 1 における順序で質問した(回答者数 96 名)。表 1b 中の Ck. X_i : Y_j は、仮想的に選択したオプションの結果 X_i と選ばなかったオプションの反事実的な結果 Y_j との組み合わせである(カッコ内は賞金金額を万円単位で示す)。

表 1a 実験で用いられた選択問題と反事実の質問

問	オプション 1	オプション 2	反事実
Q1	A. 400 万円(80%)	B. 300 万円(100%)*	4 ペア
Q2	C. 400 万円(20%)*	D. 300 万円(25%)	8 ペア
Q3	E. 500 万円(10%) 100 万円(89%)	F. 100 万円(100%)**	6 ペア
Q4	G. 500 万円(10%)**	H. 100 万円(11%)	8 ペア

注 *印は共通比効果, **印は共通結果効果の選択パターン

表 1b 実験で用いられた反事実の質問

選択問題	仮想選択 1	仮想選択 2
Q1	C1. A(400): B(300) C2. A(0): B(300)	C3. B(300): A(400) C4. B(300): A(0)
Q2	C1. A(400): B(300) C2. A(0): B(300) C5. A(400): B(0) C6. A(0): B(0)	C3. B(300): A(400) C4. B(300): A(0) C7. B(0): A(400) C8. B(0): A(0)
Q3	C1. E(500): F(100) C2. E(100): F(100) C3. E(0): F(100)	C4. F(100): E(500) C5. F(100): E(100) C6. F(100): E(0)
Q4	G. 500 万円(10%)**	H. 100 万円(11%)

注 Ck. $X: Y$ は仮想的な選択の結果 X と反事実的な結果 Y とのペア。カッコ内は結果の金額 (単位は万円)。

以降、選択問題 Q_h における k 番目の反事実の質問を Q_h_Ck と書く。また各ペア比較 Q_h における反事実 Ch_k についての「気になる程度」の回答変数 ρ_{hk} とする。[2]では各選択問題の後に ρ_{hk} の実証値を取得するための質問を行っている。例えば、反事実 $Q2_C1$ の質問は、「Q2 で 20% の 400 万円のクジの方を選んで、400 万円が当たったとします。このとき選ばなかった 25% の 300 万円のクジを別の人が選んで 300 万円が当たったとすると、どのくらい気になりますか？」

回答は「1: とても気になる」～「5: まったく気にならない」から選択され、その回答を ρ_{hk} 値とする。

ちなみに $Q1$ における B 選択、および $Q3$ における F 選択の両回答は、いわゆる確実性効果[1]である。アレの背理は確実性効果を示した回答者が、 $Q2$ や $Q4$ ではより大きな金額のオプションを取る現象であり、表 1 では BC と FG の両選択パターンに相当する。選択パターンの組み合わせ BC は、 C と D の当選確率が A と B の $1/4$ であることから共通比効果(common ratio effect; CRE)とも呼ばれる。また FG の組合せを選択するパターンは、 G と H では、 E と F からそれぞれ 89% 分の 100 万円の部分クジが消去されていることから、共通結果効果(common consequence effect; CCE)とも呼ばれる。

付録 A として[2]の追試結果を再集計した。なおこの実験では選択結果として「無差別」を許しており、集計では“tie”と記されている。本論文では B と $Q2$ における tie の組み合わせや F と $Q4$ における tie の組み合わせもアレの背理を示すケースとして数えることとする。

3. 共通比効果を阻止する反事実

以降の分析では、選択問題の回答を数値化し、数値化された選択問題間の回答変数の差 ω をリスクシフトと呼ぶ。正のリスクシフトは賞金金額の絶対値のより大きなオプションを選ぶようになる変化を表す。例えば $Q1$ で「400 万円当たるクジを選ぶ」回答を 1、「現金 300 万円を選ぶ」回答をマイナス 1 とする。同様に $Q2$ において「400 万円当たるクジを選ぶ」回答を 1、「300 万円のクジを選ぶ」回答をマイナス 1、無差別の回答 (tie) は 0 とした。

図 1 に $Q1$ と比べた $Q2$ におけるリスクシフト ω_{12} の平均値のグラフを示した。横軸は反事実の質問 $Q2_C1$ における気になる度合いの回答値 ρ_{21} である。

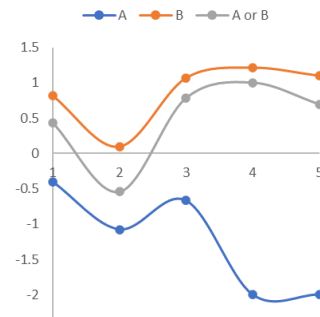


図1 $Q2_C1$ の気になる度合い ρ_{21} と $Q1$ - $Q2$ 間のリスクシフト ω_{12} 。

図 1 中の 3 系列 A , B , A or B は $Q1$ での実際の選択に対応する。図 1 の横軸は反事実の質問 $Q2_C1$ (表 1b 参照) の気になる度合い ρ_{21} である。図 1 の縦軸は $Q1$ と比べた $Q2$ でのリスクシフト ω_{12} である。

$\omega_{12} > 0$, つまり $Q1$ と $Q2$ の間のリスクシフトが正のときかつそのときだけ共通比効果が観測される。ちなみに A 選択者については $Q2$ でリスク取得が増大することはないため、つねに $\omega_{12} \leq 0$ であり、縦軸は正の値にならない。各質問のリスクシフトを付録 C にグラフ化しておく。

以上をまとめると、図 1 から分かることとして、(1) $Q2_C1$ における回答値 2 ではアレの背理 (共通比効果) が起きにくい。(2) とくに $Q1$ における B 選択者については強い抑制が見られ、(3) 逆に、反事実の質問 $Q2_C1$ が「気にならない」とき、あるいは「気になりすぎるとき」にアレの背理が観測される。

本研究の以降の部分では、小さい ω 値を誘導する反事実の質問を阻止マーカーとして活用する。 $Q2_C1=2$ 回答は共通比効果の阻止マーカーである。付録 A 表 2a の行 B のグループ ($B \& C$, $B \& D$, $B \& \text{tie}$) より、 B を回答した 71 件中 40 件が $Q2$ では C または tie であり、このうち $Q2_C1=2$ 回答 11 件中では「 $B \& \text{tie}$ 」($Q1$ で B かつ $Q2$ で無差別)の 1 件のみである。逆に、阻止マーカーの不在から共通比効果を予測すると、同上表より $Q1$ で B 回答 71 件の 49 件 (69%) が該当する。

$Q2_C1$ に加えて $Q4$ における反事実の質問 $Q4_C2$ を加えると予測力が向上する。また $Q2_C1$ は共通結果効果にかんしても予測力を持つ。この予測モデル改良については、この後続く二つの節で述べる。

しかし一体なぜ、ごく少数の反事実的質問の値だけでアレの背理を 7 割程度以上予測でき、しかも直接関係しない選択問題にまでその予測力が及ぶのだろうか。次節ではポテンシャルモデルを導入しこれを解釈する。

4. ポテンシャル

図1は、反事実の「気になる程度」 ρ を認知的な「距離」と考えたとき、化学ポテンシャルのグラフ形状とよく似たものになる。これによってアレの背理の発生にかかわる相反する認知的効果をモデル化する。

Lennard-Jones[11][12]によって提案された関数形は

$$f(\rho) = \frac{1}{\rho^n} - \frac{1}{\rho^m} \cdots (1)$$

の形であり、ここで ρ は気体分子間の距離を表す。その化学結合の強さ $f(\rho)$ は、引力項 ρ^{-n} と斥力項 ρ^{-m} の和として表現される。とくに $n = 12$, $m = 6$ とした場合は 12-6 ポテンシャルとも呼ばれ、このケースを Lennard-Jones モデル (あるいは L-J ポテンシャル) とする文献もある(SklogWiki “Lennard-Jones_model”参照)。

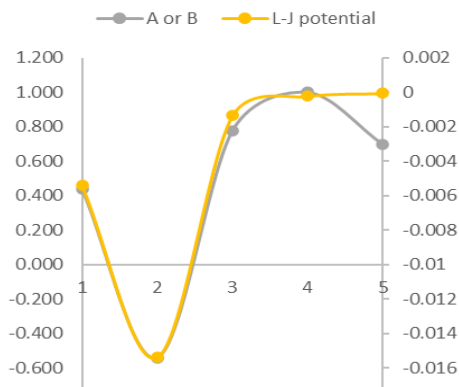


図2 ρ_{21} に対するリスクシフト ω_{12} を近似する 12-6 ポテンシャル。

図2では横軸を $r = \rho_{21} + 0.0009$ 、左縦軸を96件全体(系列名 A or B)のリスクシフト ω_{12} の平均値、右縦軸を(1)式の 12-6 ポテンシャル $f(r)$ とした。

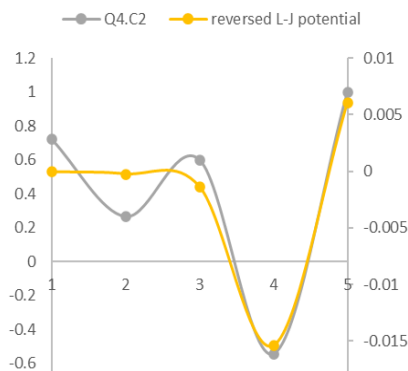


図3 リスクシフト ω_{12} を近似する 12-6 ポテンシャル $f(r)$ (横軸 $r = \omega_{12} = \rho_{21} + 0.0009$)。

図2から分かるように、スケールを無視すれば、 $f(\rho)$ は $\omega_{12}(\rho)$ のグラフ形状をよく近似する。なお、 $r = \rho_{21}$ として(1)の右辺の分数の分子を1.0085, 1.0179に変えることによっても図2と相似のグラフが得られる。マイナス0.3未満のピークを有する反事実には、Q2_C1の他、Q2_Q8とQ4_C2がある(付録C参照)。

具体的にはQ4_C2では次のように質問される。「Q4で10%の500万円の方を選んで、0円の結果だったとします。このとき、選ばなかった方の11%の100万円のクジを別の人が選んで100万円当たったとすると、どのくらい気になりますか？」

Q4_C2の気になる度合い ρ_{42} についてもポテンシャル近似が可能である(図3参照)。図3の右縦軸は、横軸を反転し、 $r = (5 - \rho) + 0.001$ とした $f(r)$ である。図3から分かるように回答値4がその負のピークである。反事実の質問Q4_C2=4と組み合わせて阻止マーカーとして予測すると、付録A表2bから、B選択者のうち55件、約77%について共通比効果の有無を言い当てる(5節の決定木による分析も参照)。ちなみにロジスティック回帰を用いると84%まで改善する(付録B)。

5. 決定木帰納による予測モデル抽出

反事実の質問のペアQ2_C1とQ4_C2は、決定木学習アルゴリズムを用いると自動抽出される(図4)。図4にRのrpartパッケージを用いデフォルトの条件で適用した結果を示す(作図はpartykit)。

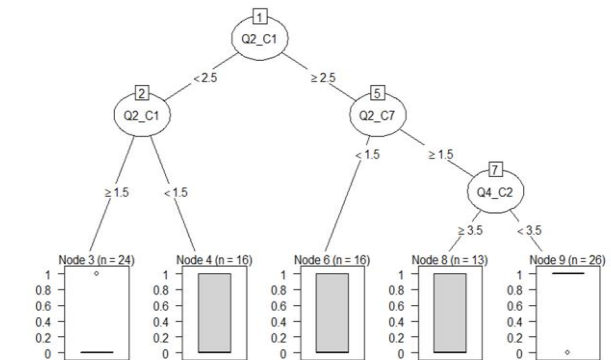


図4 共通比効果を予測する決定木

反事実のマイナスピークを検出するマーカーを説明変数とすることによっても、前述の予測ルールが得られる(図5)。図5でmarker_2は $\rho_{21}=2$ のケースを検出し、marker_5は $\rho_{42}=4$ のケースを検出する。

なお図4はRのrpartパッケージを用いた($cp=0.06$)。

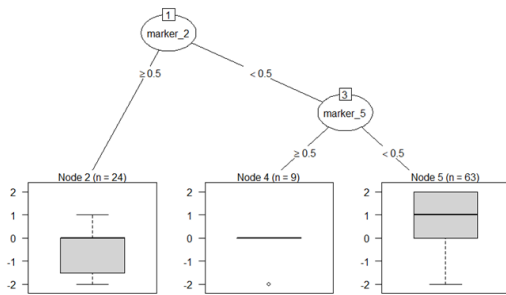


図5 マーカーを用いた予測：共通比効果

図5と同様の負のポテンシャルマーカーによる決定木抽出を共通結果効果の予測に対して適用した結果を図6に示す。

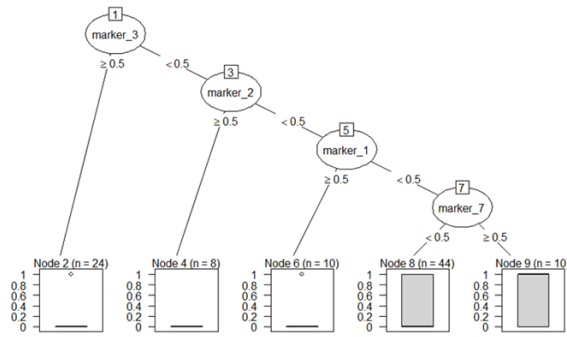


図6 マーカーを用いた予測：共通結果効果

図6の決定木における先頭3マーカーは、 $\rho_{21} \leq 2$ (marker_3), $\rho_{14}=1$ (marker_2), $\rho_{13}=4$ (marker_1) である。これら3マーカーの選言を共通結果効果の阻止マーカーとすると、Q3のF選択者が共通結果効果を起こすかどうか45件中38件(82%)で正しく判定する(付録A表2c)。ちなみにロジスティック回帰では79%である(付録B)。

6. 反事実と統計的因果推論

政策などの効果検証の文脈において介入の効果を推定する際、RCT(ランダム化比較試験)が実施できない場合、選択バイアス(交絡)に対処する必要がある。介入確率がサンプルの特微量に依存すると、選択バイアスが生じるからである。例えば結果の変数が傾向スコアと順相関すると選択バイアスが生じ、介入効果が過大に評価される。そこで傾向スコア、すなわち介入確率をサンプルの特微量から推定した値を用い、傾向スコアマッチング(PSM)や逆確率ウェイト付け

(IPW)による調整が行われる[13]。効果検証の介入群／非介入群の各ケースはそれぞれ非介入／介入の条件が直接観察されない反事実である。本研究の対象とする問題は効果検証の問題ではなく、直接的な適用はできないが、本研究で考察したデータは反事実を直接質問しており、CRE予測値の近いケースの平均を調べてみることは興味深いと思われる。反事実の質問と事実的条件の質問は予測値の級ごとにばらついており、素朴にPSMのような反事実の補正を適用してアレの背理を予測することはむずかしいと考えられる。

7. まとめ

本研究では反事実の「気になる程度」 ρ を認知的な「距離」として、リスクシフト ω の反応傾向をポテンシャル関数で近似した。これによってアレの背理についてのよりシンプルな予測と直観的解釈が同時に得られた。すなわちポテンシャルの引力項が「気にならないこと」と斥力項が「気になりすぎること」に対応する。相反する2つの作用間のバランスが崩れて不安定化することによって、正のリスクシフトを阻止できなくなる。直接関係しない設問であっても、両作用が拮抗する反事実の質問における ρ 値はリスク態度の正変化を阻止するマーカーないし特微量として解釈しうる。先行研究[5][6][7]では反事実間の気になる度合い ρ の差が分析され、上記の事実が見落とされた。その理由は、マーカーとしての反事実の質問Q2_C1やQ4_C2が後悔や安堵のように直観的な解釈をしにくいからであろう。また一点で不連続に反応するため、ロジスティック回帰において有意になりにくい(興味深いことに、両者の混合は予測を改善する)。なお[6][7]でもポテンシャル概念を論じているが、本研究とは異なり、混雑ゲームと類比された。両モデル間には直観的な関係があると考えられるが別の機会に論じたい。ちなみにQ2_C1ポテンシャルではむしろ共通結果効果の方が予測がよい。一方[10]の累積プロスペクト理論では共通結果効果の予測は共通比効果に比べて難しい[8]。また[9]で論じられたランダムウォークによる確率ウェイト関数の構成との関連も論じてみたい。

付録A. アレの背理と認知的変数

本付録は、認知的な変数で選択問題回答を集計した表2abc, およびQ3-Q4間のリスクシフト図6とそれを12-6ポテンシャルで近似した図7を掲載する。

表 2a Q1～Q4 の選択と反事実の質問 Q2_C1

選 択 パ タ ン	Q2_C1 の回答値 ρ_{21}					計
	1	2	3	4	5	
Q1&Q2/Q3&Q4						
A	5	13	3	1	3	25
A & C	4	5	2			11
A & D	1	6	1	1	3	12
A & tie		2				2
B	11	11	15	14	20	71
B & C	3		6	7	9	25
B & D	5	10	5	4	7	31
B & tie	3	1	4	3	4	15
E	8	11	3	9	14	45
E & G	6	5	3	8	13	35
E & H	2	6		1		9
E & tie					1	1
F	7	10	14	5	9	45
F & G	2	1	10	4	5	22
F & H	5	9	3	1	3	21
F & tie			1		1	2
Tie	1	3	1	1		6
tie & G				1		1
tie & tie	1	3	1			5
計	16	24	18	15	23	96

注 行を選択問題ペアの回答の組み合わせでグループ化した。

表 2b Q1 と Q2 の選択と反事実の質問 Q4_C2

選 択 パ タ ン	阻止マーカー-x, y				計
	x y	x	y	-	
Q1, Q2					
A	9	3	12	1	25
C	5	1	5		11
D	4	2	5	1	12
tie			2		2
B	54	6	10	1	71
C	25				25
D	15	6	9	1	31
tie	14		1		15
計	18	34	25	11	96

注 x は $\rho_{21}=2$, y は $\rho_{42}=4$.

表 2c Q3 と Q4 の選択と認知的阻止マーカー

選 択	認知的阻止マーカー 1, 2, 3							計
	-	1	2	3	13	23	123	
Q3, Q4								
E	16	6	4	14	2	3		45
G	15	5	4	10		1		35
H	1			4	2	2		9
tie		1						1
F	24	3	1	11	4	1	1	45
G	18	1		3				22
H	4	2	1	8	4	1	1	21
tie	2							2
tie	2			4				6
G	1							1
tie	1			4				5
計	42	9	5	29	6	4	1	96

注 マーカー1, 2, 3 は $\rho_{13}=4$, $\rho_{14}=1$, $\rho_{21}\leq 2$ である。

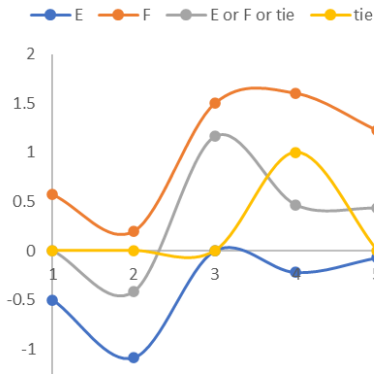


図 7 Q2_C1 の気になる度合い ρ_{21} に対する Q3-Q4 間のリスクシフト ω_{34} .

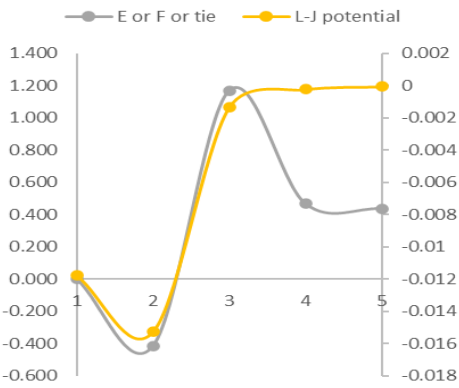


図 8 リスクシフト ω_{34} と 12-6 ポテンシャル (横軸 $r = \omega_{34} = \rho_{21} + 0.002$ とした)。

付録 B. ロジスティック回帰による予測

本付録は予測値の級ごとに反事実の回答平均が実際の選択によってどの程度変わるかを調べる。

なお全反事実の質問を用いた Q1 のロジスティック回帰 (R の glm 関数) の回帰係数及び統計量は表 3 のようである。なお Q2_C1 は有意ではない。表 4 に示すように選択予測は 84%の的中率である。CCE についての同様の結果を表 5 と表 6 に示す (的中率は 79%)。

表 3 CRE のロジスティック回帰係数

Coefficients CRE model:

Estimate	Std.	Error	z value	Pr(> z)	Signif.
Intercept	-2.73224	1.2057	-2.266	0.02344	0.05
Q1_C1	-0.06411	0.36578	-0.175	0.86088	
Q1_C2	-1.29491	0.67965	-1.905	0.05675	0.1
Q1_C3	0.56949	0.51683	1.102	0.27051	
Q1_C4	0.22464	0.37076	0.606	0.54458	
Q2_C1	-0.03644	0.67222	-0.054	0.95677	
Q2_C2	-1.35583	0.74558	-1.818	0.06899	0.1
Q2_C3	-1.96708	0.75467	-2.607	0.00915	0.01
Q2_C4	-0.74125	0.59964	-1.236	0.2164	
Q2_C5	0.97405	0.57547	1.693	0.09052	0.1
Q2_C6	0.99112	0.63588	1.559	0.11908	
Q2_C7	2.48555	0.89858	2.766	0.00567	0.01
Q2_C8	0.69868	0.39558	1.766	0.07736	0.1
Q3_C1	2.32513	0.83456	2.786	0.00534	0.01
Q3_C2	-0.34839	0.5952	-0.585	0.55833	
Q3_C3	2.66317	0.82043	3.246	0.00117	0.01
Q3_C4	0.17067	0.47083	0.362	0.71699	
Q3_C5	-0.93138	0.63885	-1.458	0.14486	
Q3_C6	-0.74827	0.50768	-1.474	0.14051	
Q4_C1	0.70853	0.75278	0.941	0.34659	
Q4_C2	-1.90049	0.77316	-2.458	0.01397	0.05
Q4_C3	0.20313	0.4312	0.471	0.63759	
Q4_C4	-1.2229	0.82989	-1.474	0.1406	
Q4_C5	-0.40557	0.62776	-0.646	0.51825	
Q4_C6	-0.83189	0.49894	-1.667	0.09546	0.1
Q4_C7	0.21602	0.57081	0.378	0.7051	
Q4_C8	0.20969	0.60416	0.347	0.72853	
*AIC: 123.38 Fisher Scoring: 7 iterations					
deviance: 130.405 (95 d.f.) for null					
69.383 (69 d.f.) for residual					

表 4 CRE 選択の予測(glm)

class id	glm	Q1=A	Q1=B	total	hit
1	0.10	13	13	26	25
2	0.20	5	9	14	11
3	0.30	5	2	7	6
4	0.40	1	4	5	3
5	0.50	1	3	4	3
6	0.60	0	7	7	5
7	0.70	0	7	7	3
8	0.80	0	6	6	6
9	0.90	0	5	5	4
10	1.00	0	15	15	15
total		25	71	96	81

表 5 CCE のロジスティック回帰係数

Coefficients for CCE model:

Estimate	Std.	Error	z value	Pr(> z)	Signif.
Intercept	-2.95156	1.2124	-2.434	0.0149	0.05
Q1_C1	-0.50522	0.35532	-1.422	0.1551	
Q1_C2	-0.31741	0.44927	-0.706	0.4799	
Q1_C3	-0.01218	0.34559	-0.035	0.9719	
Q1_C4	0.38099	0.2955	1.289	0.1973	
Q2_C1	0.14992	0.50918	0.294	0.7684	
Q2_C2	-0.01766	0.54557	-0.032	0.9742	
Q2_C3	0.32701	0.46709	0.7	0.4839	
Q2_C4	-0.82522	0.44854	-1.84	0.0658	0.1
Q2_C5	0.52657	0.40949	1.286	0.1985	
Q2_C6	0.60136	0.4395	1.368	0.1712	
Q2_C7	0.04793	0.52377	0.092	0.9271	
Q2_C8	0.12751	0.32787	0.389	0.6974	
Q3_C1	-0.2139	0.51927	-0.412	0.6804	
Q3_C2	0.03256	0.51283	0.063	0.9494	
Q3_C3	0.36469	0.49312	0.74	0.4596	
Q3_C4	0.31746	0.41155	0.771	0.4405	
Q3_C5	-0.48646	0.59718	-0.815	0.4153	
Q3_C6	0.0523	0.46878	0.112	0.9112	
Q4_C1	0.64513	0.62261	1.036	0.3001	
Q4_C2	-0.30283	0.5134	-0.59	0.5553	
Q4_C3	-0.72105	0.42222	-1.708	0.0877	0.1
Q4_C4	-0.4156	0.58641	-0.709	0.4785	
Q4_C5	0.2432	0.50987	0.477	0.6334	
Q4_C6	0.26136	0.41349	0.632	0.5273	
Q4_C7	0.12602	0.53104	0.237	0.8124	
Q4_C8	0.04217	0.51093	0.083	0.9342	
*AIC: 137.48 Fisher Scoring: 6 iterations					
deviance: 110.111 (95 d.f.) for null					
83.484 (69 d.f.) for residual					

表 6 CCE 選択の予測(glm)

class id	glm	Q3≠F	Q3=F	total	hit
1	0.10	18	16	34	20
2	0.20	10	2	12	10
3	0.30	7	5	12	7
4	0.40	7	9	16	4
5	0.50	4	1	5	2
6	0.60	3	5	8	4
7	0.70	2	3	5	6
8	0.80	0	2	2	6
9	0.90	0	0	0	4
10	1.00	0	2	2	13
total		51	45	96	76

付録 C リスクシフト

本付録は各反事実的質問の回答値 ρ ごとのリスクシフト ω_{12} と ω_{34} のグラフを掲載する。

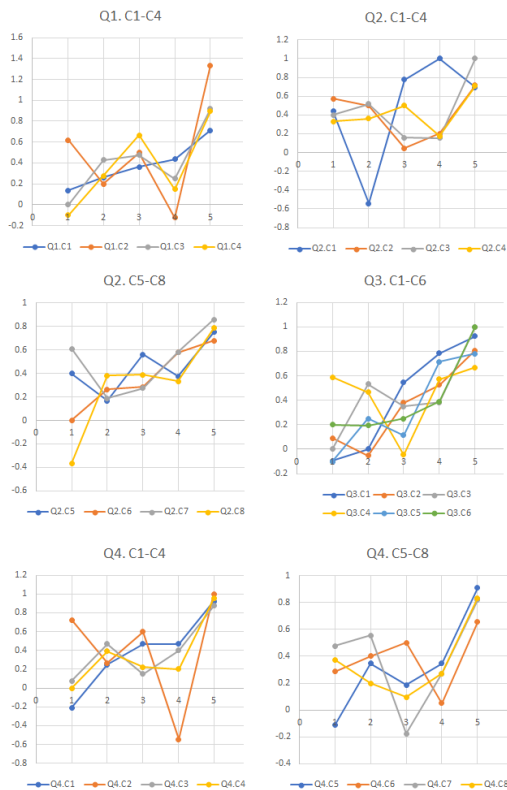


図9 Q1 と Q2 の間のリスクシフト ω_{12} .

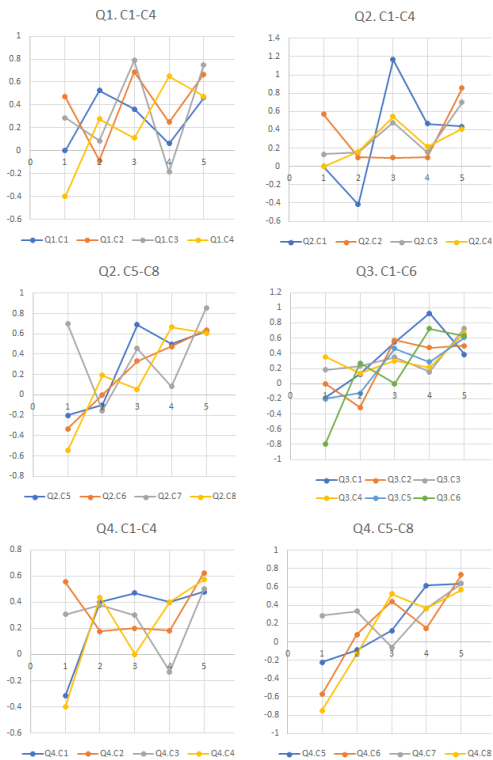


図10 Q3 と Q4 の間のリスクシフト ω_{34} .

文献

[1] Allais, M. (1953) “Le comportement de l'homme rationnel devant le risque: critique des postulats et axiomes de l'école américaine”, *Econometrica*, 21, 503–546.

[2] 犬童健良 (2013) “アレの背理における注目と注目の流れ”, *行動経済学*, Vol. 6, pp. 70-73. doi:10.11167/jbef.6.70

[3] 犬童健良 (2014a) “アレの背理における反事実的注目とリスク選好の認知的安定性”. *関東学園大学経済学紀要*, Vol. 39, pp. 53–80. doi:10.20589/kantogakueneconomics.39.0_53

[4] 犬童健良 (2014b). “リスク下の選択における認知的資源配分: 注目の枠組みの最適性”, *日本認知科学会第 31 回大会論文集*, 864–872.

[5] 犬童健良 (2016b) “認知的ネットワークによるアレの背理におけるリスク態度変化の説明”, *行動経済学*, Vol. 9, pp. 81-84. doi:10.11167/jbef.9.81

[6] 犬童健良 (2016a). “ポテンシャルゲームを用いてクジ比較における認知をモデル化する”, *日本認知科学会第 33 回大会論文集*, pp. 864–872.

[7] 犬童健良 (2017) “ポテンシャルゲームを用いたリスク選択の認知モデル”, *関東学園大学経済学紀要*, Vol. 42, pp. 1-20. doi:10.20589/kantogakueneconomics.42.0_1

[8] 犬童健良 (2018) “共通比効果と共通結果効果を共に予測するプロスペクト理論のシミュレーション研究”, *関東学園大学経済学紀要*, Vol. 44, pp. 19-43. doi:10.20589/kantogakueneconomics.44.0_19

[9] 犬童健良 (2021) “ベータ分布を用いた累積プロスペクト理論における確率ウェイト関数についての一考察”, *関東学園大学経済学紀要*, Vol. 47, pp. 1-29. doi:10.20589/kantogakueneconomics.47.0_1

[10] Lennard-Jones, J.E. (1924) “On the Determination of Molecular Fields. I. From the Variation of the Viscosity of a Gas with Temperature”, *Proceedings of the Royal Society of London. Series A, Containing Papers of a Mathematical and Physical Character* 106, pp. 441–462.

[11] Lennard-Jones, J.E. (1924) “On the Determination of Molecular Fields. II. From the Equation of State of a Gas”, *Proceedings of the Royal Society of London. Series A, Containing Papers of a Mathematical and Physical Character* 106, pp. 463–477.

[12] Tversky, A., & Kahneman, D. (1992) “Advances in prospect theory: Cumulative representation of uncertainty. *Journal of Risk and Uncertainty*”, Vol. 5, No. 4, pp. 297-323.

[13] 安井翔太(2020). “効果検証入門”,技術評論社.