

恐怖条件づけにおけるABA復元効果 のベイジアン統計モデリング

二瓶正登^{1,3}・北條大樹^{2,3}・澤幸祐⁴

¹ 専修大学大学院文学研究科 ² 東京大学大学院教育学研究科

³ 日本学術振興会 ⁴ 専修大学人間科学部

序文

- 古典的条件づけにおけるABA復元効果とは
 - ① CSとUSの対呈示を文脈Aで行う = CRが獲得される
 - ② CSのみの呈示を文脈Bで行う = CRが減弱する
 - ③ 文脈Aに戻ってCSを呈示する = 減弱したCRが再び生じる

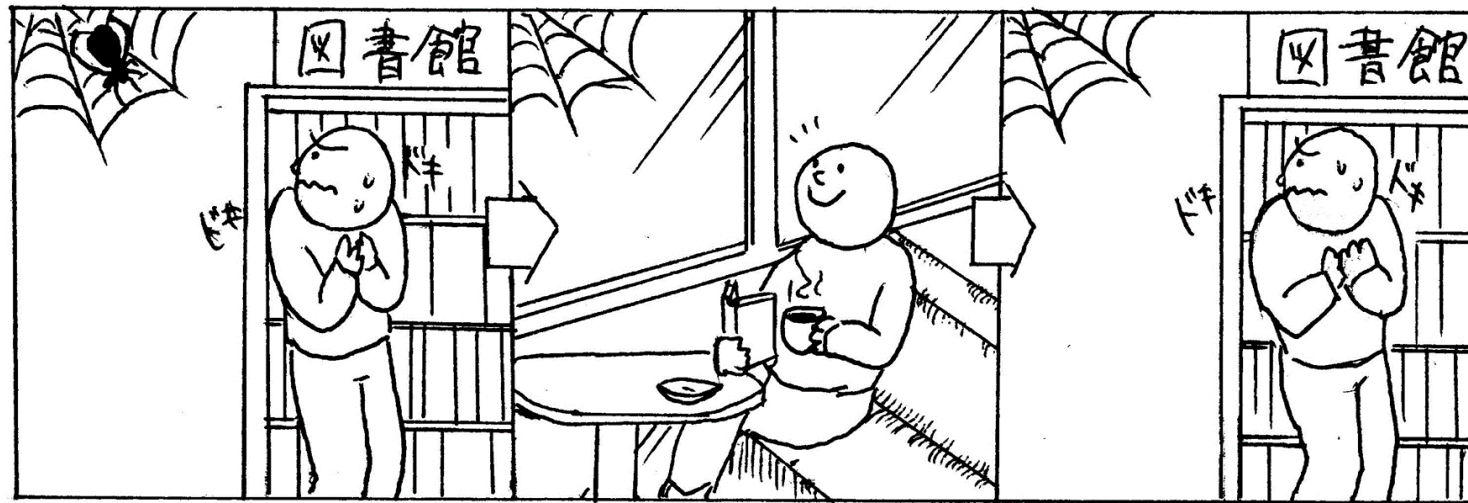


図1. 場所を文脈、クモをUS、クモの巣をCSとした場合のABA復元効果

- この現象の説明として多くの連合学習理論が用いられ、その妥当性が検証されてきた

主なABA復元効果の説明として

- **Rescorla-Wagnerモデル (Rescorla & Wagner, 1972)**
- **Boutonのモデル (Bouton, 1993)**

が多く用いられてきた

Rescorla-Wagnerモデルによる説明

① CRはCSと文脈刺激の連合強度の合計によって決定

② 連合強度の変化量は

- ・ 個体の持つ予期と現実間の差
- ・ 刺激の物理的特性

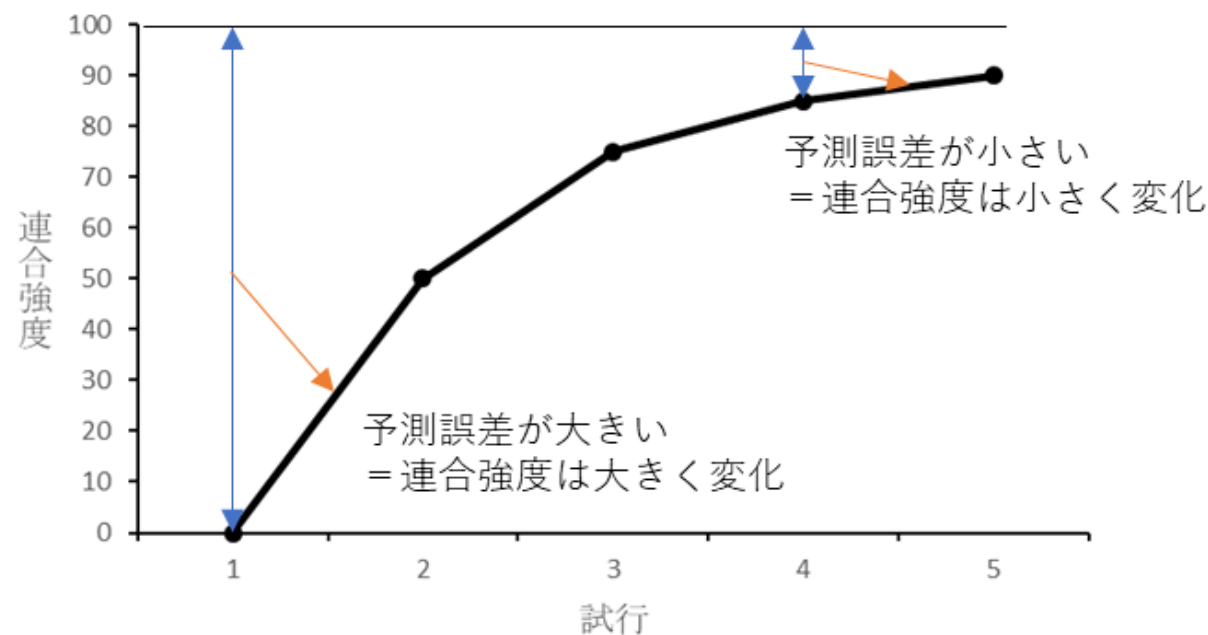
の2点により決定



① $CR \propto \sum V$

② $\Delta V_{CS} = \alpha_{CS} \beta (\lambda - \sum V)$

$\Delta V_{Context} = \alpha_{Context} \beta (\lambda - \sum V)$



Rescorla-Wagnerモデルによる定性的説明

獲得期：CSと文脈Aが強化

= CSと文脈Aが連合強度を獲得

※試行間間隔中における文脈刺激の非強化の効果は非常に小さいと仮定

消去期：CSと文脈Bが非強化

= CSの連合強度は減弱し文脈Bは条件性制止子となる

つまり

- ・ 消去期で文脈Aは非強化を受けない = 連合強度は減弱しない
- ・ 消去期でCSの連合強度が0にならない

(protection from extinction現象 (Lovibond et al., 2000))

= A文脈に戻ると反応が生じる

Rescorla-Wagnerモデルによる定量的説明

獲得期

CSと文脈Aが連合強度を獲得する

例：CS = 0.8, 文脈A = 0.2 $CR = 0.8 + 0.2 = \mathbf{1}$

消去期

CSの連合強度が減少し、文脈Bは制止性連合を獲得する

例：CS = 0.2, 文脈B = -0.2 $CR = 0.2 + -0.2 = \mathbf{0}$

テスト期

CSと文脈Aの連合強度が加算される

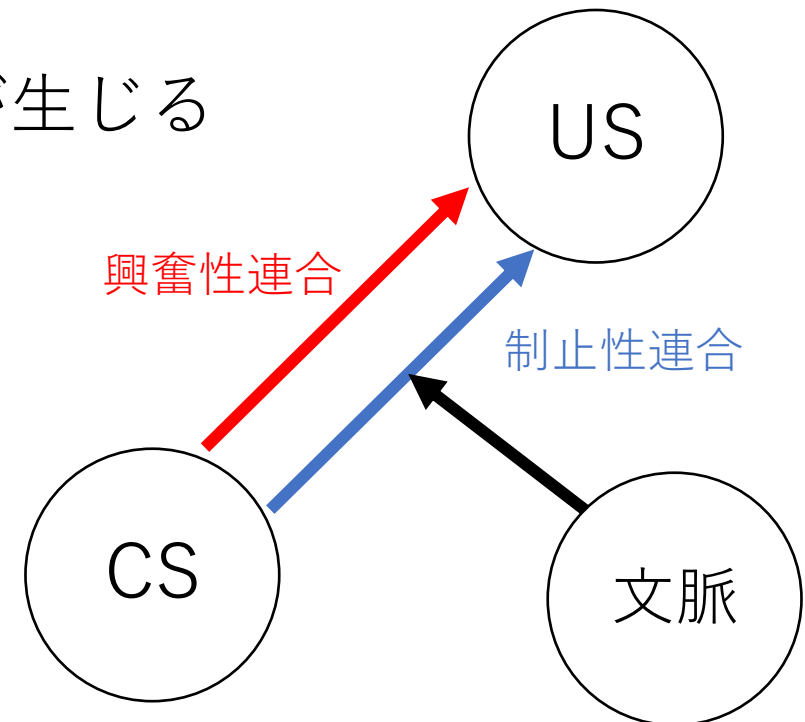
例：CS = 0.2, 文脈A = 0.2 $CR = 0.2 + 0.2 = \mathbf{0.4}$



消去期には生じなかった
反応が文脈を移動するだけ
で生じるようになる

Boutonのモデルの定性的説明

- CSは強化時に興奮性連合、非強化時に抑制性連合を獲得する
- CRは両連合がどの程度検索されたかによって決定される
- 抑制性連合の検索は文脈変化に脆弱である
→ 消去を行った文脈から離れると
抑制性連合の検索が弱まるため反応が生じる



Boutonのモデルの定量的説明

- 条件づけ時

CSが興奮性の連合強度を獲得する

例：興奮性連合 = 1 CR = **1**

- 消去時

CSが制止性の連合強度を獲得する

例：興奮性連合 = 1, 制止性連合 = -1 CR = 1 + (-1) = **0**

- テスト時

制止性連合の検索が文脈移動によって減衰される

例：興奮性連合 = 1, 制止性連合 = -1 * 0.6 CR = 1 + (-0.6) = **0.4**



文脈変化によって減衰

- これまで多くの研究で両モデルの妥当性が検証されてきた
→ 研究の多くはBoutonのモデルを支持
(e.g., Bouton & Swartzentruber, 1986)
- ただしほとんどの研究はモデルの定性的な予測
(= 有意な差が生じるかどうか) に基づいた議論であった
→ 定量的な予測 (= 具体的な数値の予測) はデータと一致するか？

目的

- 両モデルから得られる**定量的な予測がどの程度データと適合しているか**をベイズ統計モデリングを用いて検証
- また探索的な目的で、推定されたパラメータおよび不安の個人差との相関も検討

方法

- 対象者

大学生および大学院生64名（事前の例数設計により決定）に実験を行った対象者は無作為に実験群（ABA群）と統制群（AAA群）に割り当てられたデータ測定時のエラーがあった1名（ABA群）を除く63名を分析対象とした（男性23名，女性40名，平均年齢 21.63 ± 3.67 ）

- 刺激

条件刺激（CS+とCS-）：中性的な顔画像

無条件刺激（US）：怒り表情と嫌悪的な言葉からなる動画刺激

中性刺激（NS）：中性的な表情と中性的な言葉からなる動画刺激

※動画刺激の音量は全て80db以下に設定された

※用いるUSとNSは対呈示されるCSと対応する同じ人物（男 or 女）を用いた
文脈刺激：スクリーン画面の背景色（赤・青）

測度

- ・ US予期

全試行のCS呈示後に、USが次にどの程度到来すると感じるかを0（絶対に来ない）から9（絶対に来る）の間で評定を求めた

- ・ 感情価

各期の終わりにCSに関する感情価について1（不快）から9（快）までの間で評定を求めた

※本分析では感情価データは使用しなかった

- ・ Short version of Fear of Negative Evaluation (SFNE: 笹川ら, 2004)

社交不安の個人差を測定するため、実験開始前に回答を求めた

※本尺度は順向項目のみが社交不安の測度として妥当であることが報告されているため（二瓶ら, 2018）、本研究では順向項目のみを分析に使用した

手続き

手続き

- 全ての参加者は獲得期，消去期，テスト期を経験した

獲得期（9試行）：CS+とUSの対呈示

消去期（18試行）：CS+とNSの対呈示

テスト期（3試行）：CS+とNSの対呈示

※CS-は全ての期においてCS+と同じ試行だけNSと対呈示された

- AAA群の参加者はこの手続きを全て同じ文脈下で行う一方，ABA群の参加者は消去期のみを異なる文脈下で行った

手続きの模式図

	獲得期 (9試行)	消去期 (18試行)	テスト期 (3試行)
AAA群	CS-US	CS-NS	CS-NS
ABA群	CS-US	CS-NS	CS-NS

※用いる刺激および文脈の種類は全てカウンターバランスされた。

分析

- Rescorla-WagnerモデルおよびBoutonモデルを統計モデルとして記述し、ベイズ統計モデリングを使用してパラメータの推定を行った
 - ※iterationは20000回, burn-inは5000に設定した
 - $\hat{R}$ が1.10以下だった場合にモデルが収束したと判断した
- 両モデルによって推定された値の傾向がデータの傾向と一致しているかを確認するため, 推定された連合強度の推定値とデータとの適合の比較を行った
 - ※両モデルともCS-では反応に変化が生じないことを予測するため, 分析にはCS+のみを用いた
- モデル比較のために事後予測分布, WAICおよびWBICを用いた
- 探索的な目的で社交不安の個人差と推定された両モデルのパラメータの個人差に関係があるかを相関分析によって検討した

結果

- 両モデルは収束したと判断した。
- 得られたデータを図1に、両モデルによって推定された連合強度の平均値を図2と3に示した
- AAA群のテスト期以外においてはデータと推定値平均の間に大きな乖離は見られなかった。

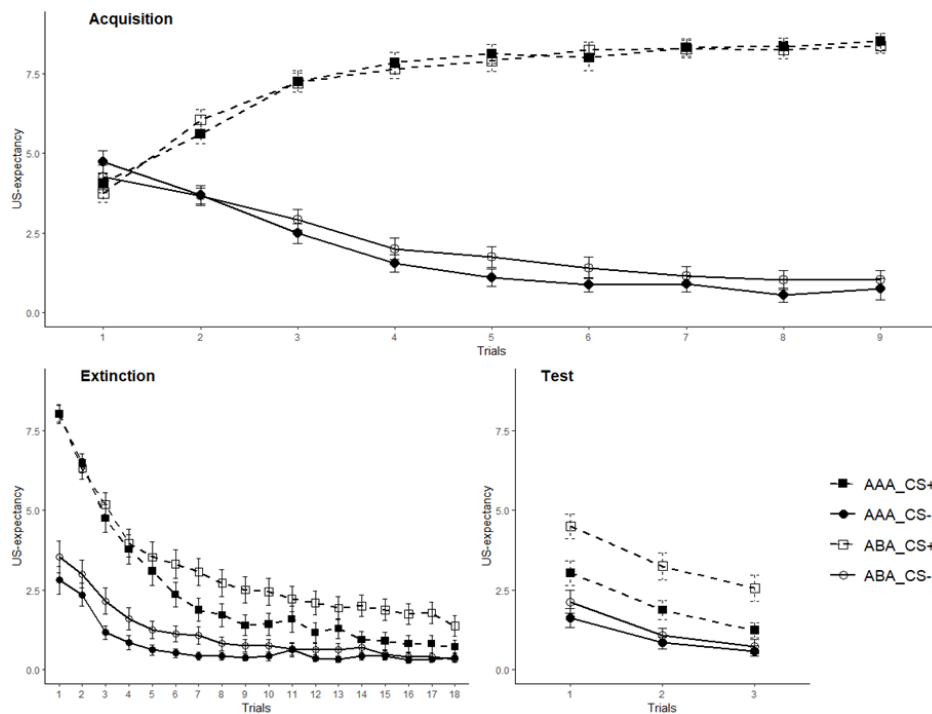


図1. 各期における両群のUS予期の平均値（エラーバーは標準誤差）

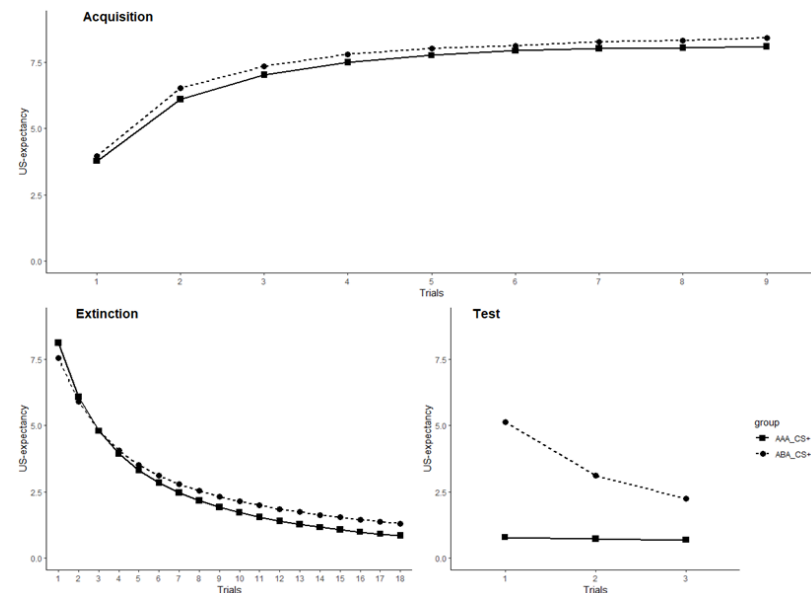


図2. Rescorla-Wagnerモデルによって推定された各期における両群の連合強度の平均値

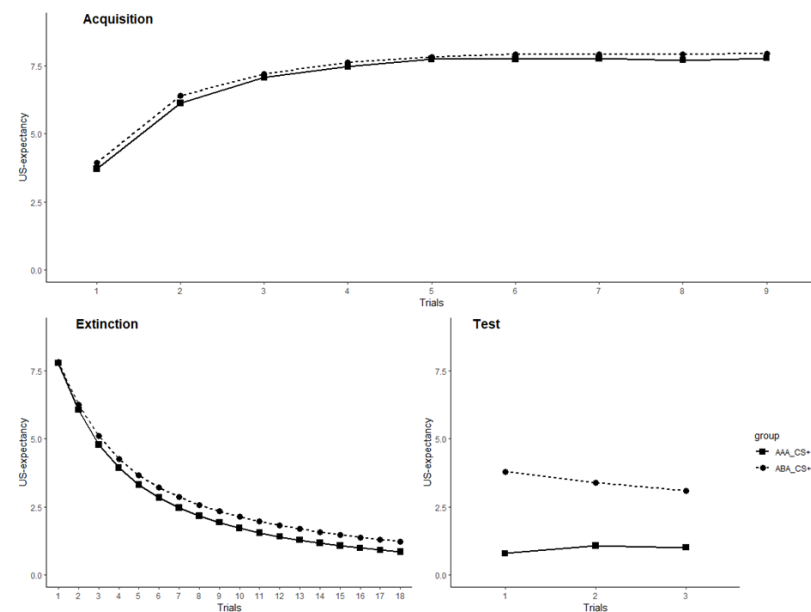


図3. Boutonのモデルによって推定された各期における両群の連合強度の平均値

結果：事後予測分布（図4）

- 両群ともにデータとほぼ一致した事後予測分布が得られた

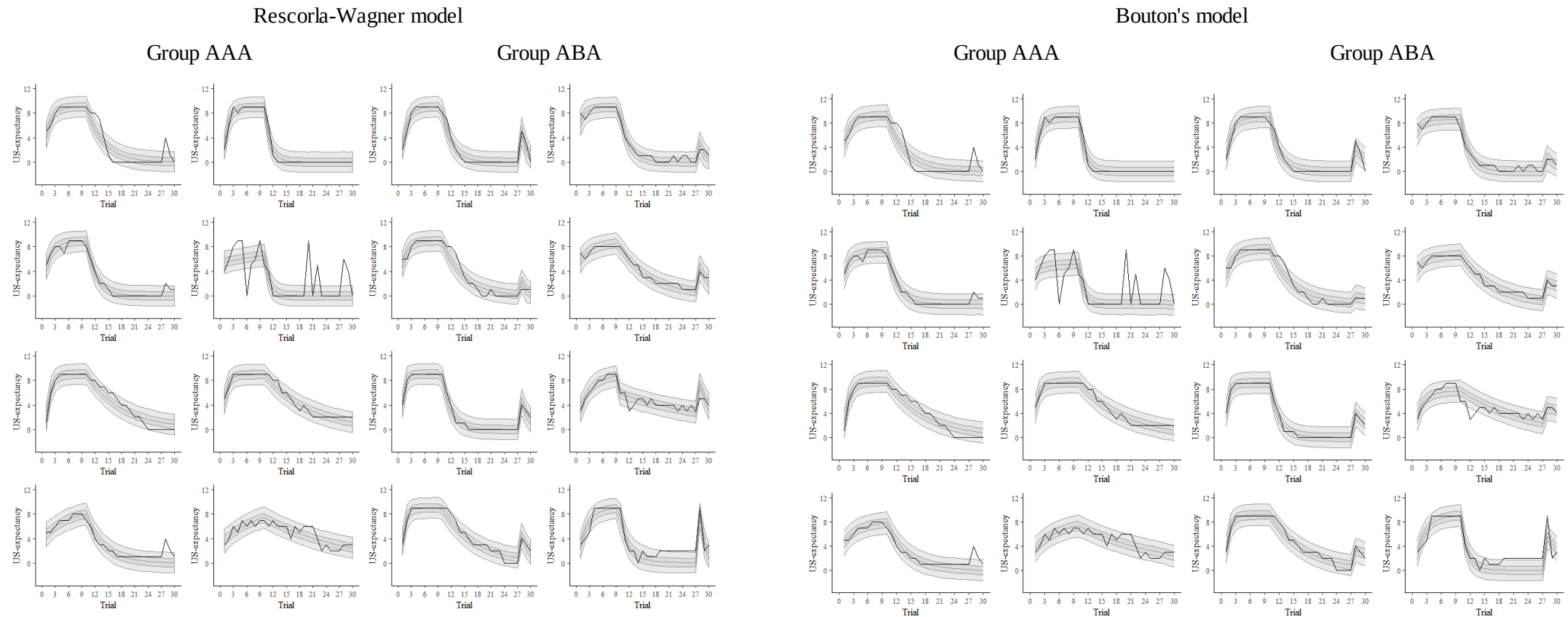


図4. 両モデルによって得られた16名の事後予測分布。直線はデータ，薄いグレーは95%信用区間，濃いグレーは50%信用区間を表す。

結果：適合度指標

- WAIC

Rescorla-Wagnerモデル:-2.22

Boutonのモデル:-2.18

- WBIC

Rescorla-Wagnerモデル:-117.19

Boutonのモデル:-122.89

= 予測の観点からはRescorla-Wagnerモデル，データとの適合の観点からはBoutonのモデルの方が良いことが示された

結果：相関分析（表1・表2）

- 両モデルにおけるパラメータ間には中程度以上の相関関係がある組み合わせがある一方で、社交不安の個人差とパラメータ間の相関関係はかなり小さいことが示された

表1. Rescorla-Wagnerモデルにおいて推定されたパラメータおよびSFNEの平均値、標準偏差およびそれらの間の相関係数

	<i>M</i>	<i>SD</i>	1	2	3	4
1 α^{acq}	0.46	0.24	-			
2 α^{ext}	0.25	0.20	0.15	-		
3 λ	8.97	0.09	0.10	0.04	-	
4 <i>FNE</i>	25.59	6.85	0.12	-0.08	-0.01	-

	<i>M</i>	<i>SD</i>	1	2	3	4	5	6	7
1 α_{con}^{acq}	0.11	0.09	-						
2 α_{CS}^{acq}	0.40	0.24	-0.06	-					
3 α_{con}^{ext}	0.08	0.10	-0.01	0.24	-				
4 $\alpha_{CS}^{ext \cdot test}$	0.13	0.12	0.06	0.41	0.10	-			
5 α_{con}^{test}	0.31	0.19	-0.01	0.29	0.51	0.36	-		
6 λ	8.94	0.01	-0.28	-0.09	-0.34	-0.43	-0.44	-	
7 <i>FNE</i>	25.65	7.67	-0.05	0.08	0.01	-0.04	-0.25	0.10	-

表2. Boutonのモデルにおいて推定されたパラメータおよびSFNEの平均値、標準偏差およびそれらの間の相関係数

	<i>M</i>	<i>SD</i>	1	2	3	4
1 α^{Ve}	0.49	0.22	-			
2 α^{Vi}	0.25	0.20	0.20	-		
3 λ	8.84	0.38	0.26	0.07	-	
4 <i>FNE</i>	25.59	6.85	0.08	-0.08	0.16	-

	<i>M</i>	<i>SD</i>	1	2	3	4	5
1 α^{Ve}	0.51	0.23	-				
2 α^{Vi}	0.23	0.17	0.28	-			
3 λ	8.72	0.43	0.22	0.55	-		
4 <i>Renewal</i>	0.60	0.21	0.30	0.27	0.26	-	
5 <i>FNE</i>	25.65	7.67	0.003	-0.06	0.04	0.004	-

考察

- 両モデルともAAA群のテスト期の結果以外は本実験のデータとほぼ一致する定量的な予測を行うことが可能

※AAA群のテスト期の結果は多くの先行研究（e.g., Vansteenwegen et al., 2005）と一致しないため、実験手続き上の問題（感情価評価が文脈操作として機能していた、など）である可能性

- モデル比較の結果は予測と適合の観点のそれぞれで、良いとされるモデルが異なる

→本研究の結果だけでABA復元に関する「より妥当な」モデルを決定することは困難

- 社交不安の個人差と両モデルで推定されたパラメータ間に相関がほとんどない

※多くの研究で不安と恐怖条件づけの消去の速度には関連があることが示されている

(e.g., Duits et al., 2015)

→実験事態を先行研究と揃えた上で同様の検討が必要？

これらの統計モデルの妥当性をさらに検証するためには

- 異なる実験事態で同様の検証を行う
- モデルが仮定するパラメータの操作と対応した実験手続きを行った場合において量的な予測とデータが一致するか

といったことを検証する必要がある

→従来あまり行われてこなかった「**量的な予測**」の**妥当性**を多様な事態で検証していく必要

注：発表時未掲載スライド

- 本発表の内容が論文化されました

Nihei, M., Hojo, D., & Sawa, K. (2021). The renewal effect in fear conditioning with aversive facial expression and negative sentences as unconditioned stimuli. *Learning and Motivation*, 74(March), 101725.
<https://doi.org/10.1016/j.lmot.2021.101725>