

海外特別研究員最終報告書

独立行政法人 日本学術振興会 理事長 殿

採用年度 平成29年度

受付番号 845

氏 名 八賀 洋介

(氏名は必ず自署すること)

海外特別研究員としての派遣期間を終了しましたので、下記のとおり報告いたします。

なお、下記及び別紙記載の内容については相違ありません。

記

1. 用務地（派遣先国名）用務地： ワシントンD. C. （国名： アメリカ合衆国）2. 研究課題名（和文）※研究課題名は申請時のものと変わらないように記載すること。“共感”のラットモデルの妥当性の実験的検討3. 派遣期間：平成29年 4月 26日 ～ 平成31年 4月 11日

4. 受入機関名及び部局名

American University Department of Psychology5. 所期の目的の遂行状況及び成果…書式任意 書式任意（A4判相当3ページ以上、英語で記入も可）

(研究・調査実施状況及びその成果の発表・関係学会への参加状況等)

(注)「6. 研究発表」以降については様式10-別紙1～4に記入の上、併せて提出すること。

まず、本研究の目的の概要を説明し、次に研究実施・遂行状況及び成果について述べていく。

本研究テーマの背景および目的

人間は他人が困っている場合や苦しんでいる場合には手を差し伸べることがある。時には、その行動は自らに益がなくとも、また時には、自らの命を賭してまで行うことがある。このような行動の原因ないし背景には共感の働きが関係すると素朴に考えられている。逆に共感に基づく行動を全く示さない場合には、自閉症、反社会性パーソナリティ障害、社交不安障害などといった病的な診断が下ることもある (Belzung et al., 2005; Lukas et al., 2011; Preston et al., 2007)。このような臨床上的判断指標として共感の有無が関係を持つことから、人間の共感に対する動物モデルの策定が試みられ、そのモデルが医療の進展へ寄与すると見込まれてきた。関連した動物研究はアリ (Nowbahari et al., 2009) からチンパンジー (e.g., Silk et al., 2005) まで、幅広い種で進められてきた。その中でも最も利用されてきた種は、ラットと類人猿、サルである。このうち、ラット等げっ歯類は安価で多くの研究者にとって研究実施（および現象の確認）が容易であり、心理学のみならず遺伝学、生理学など様々な次元でよく調べられてきた実験動物である。このラットが共感の動物モデルとして確立するならば、その価値は大きく、素朴に仮定されている共感の働きについての検討は促進されるものと思われる。そのような見通しは、ラットを利用した共感研究を魅力的なテーマとしている。また、ポジティブな結果への一般メディアの反応は大きい。例えば、本申請研究の基幹論文となる Ben-Ami Bartal et al. (2011) の Science 誌への掲載を受けて、多少調べただけでも、以下のような一般メディアに関連記事が見つかった。

<https://www.wired.com/2011/12/rat-empathy/> (Wired.com)

<https://m.huffpost.com/us/entry/1138675/amp> (Huffinton post)

<https://www.nationalgeographic.com/science/phenomena/2011/12/09/empathic-rats-spring-each-other-from-jail/> (National geographic)

<https://www.bbc.com/earth/story/20150514-rats-save-mates-from-drowning> (BBC)

このように、人間の共感能力の解明に向けた動物モデルは、社会的に大いに注目されている。

ただし、当然のことながら、実証の過程では慎重な検討が重ねられていくべきである。行動レベルで

観察される現象は、脳や関連生理の働きと対応させ、それらの働きの意味を明確にし、さらに、人間の臨床や医療応用への基礎的知見となることが見込まれるならばなおさらである。その動物モデルから得られた知見を人間へ適用する前に、そのモデルが十分な確からしさを持つことを示していくことが基礎研究者に求められている。

ラットの共感に関する研究自体は早くも 1960 年代から検討されてきたが、近年になりいっそう注目を集めるテーマとなってきた。その発端となった先行研究として、Ben-Ami Bartal et al. (2011) が発表したラットの救助行動の実験が挙げられる。彼らの最初の実験条件で、円筒状の箱に拘束されたラット（拘束ラット）に対し、もう一匹のラット（フリーラット）が、その箱の正面扉を開放する割合が、各種の統制条件に比べ高いことを示した。その実験において、拘束ラットが拘束箱から解放された後は、広い実験箱の中で 2 個体のラットは共に時を過ごすことができた。

この結果について、3 つの解釈が可能である。第 1 に、ラットは社会性動物であることは周知のことであるので、拘束ラットとの社会的接触（社交）がフリーラットの行動への報酬となっていた（社交仮説）。第 2 に拘束ラットから発せられた苦痛による悲鳴が嫌悪的であったためそれを止めるために扉を開放した（悲鳴音嫌悪仮説）。第 3 に、フリーラットが拘束ラットに共感し、拘束される苦痛を緩和することに動機づけられていた（共感仮説）。

Ben-Ami Bartal et al. (2011) は悲鳴音嫌悪仮説を検討するためにラットが発する高周波音を記録したが、それはあまり頻繁に発せられることはなかったため、フリーラットの行動の原因ではなかっただろうと結論づけている。また、社交仮説と共感仮説を分離するために、実験 2 では、フリーラットが前面扉を押すと拘束箱の後面扉が開き、拘束ラットは実験箱の別区画へと解放される手続きへと変更したところ、フリーラットは引き続き拘束箱を開ける行動を行うことを示した。このような次第で、この先行研究は“拘束され助けを求める”ラットに対しフリーラットが人間と類似した“共感”を示した、行動レベルにおける経験的根拠として頻繁に引用されてきた。

しかし、本申請者の受け入れ研究室において、Ben-Ami Bartal et al. (2011) について組織的に追試を行ったところ（Silberberg et al., 2014）、もし、拘束ラットを開放すると社交が可能となるという状況を過去に一度も経験していない場合には、フリーラットは Ben-Ami Bartal et al. の実験 2 のような状況で、拘束箱の扉を開ける行動をほとんど行わないことを確認している。また、解放された後の 2 匹のラットの観察を続けると、頻繁に拘束箱へ出入りする姿が見られ、そもそも拘束箱に入れられること自体が嫌悪的であるという想定について、その一貫性へ疑義を挟む余地もあると指摘された。

ラットが他個体への共感に動機づけられて救助行動を行うという経験的事実を報告する他の研究（Rice & Gainer, 1962; Sato et al., 2015）においても、それぞれに得られた事実と導かれる結論との間になんらかの飛躍を指摘することができるため、本申請者は未だラットが他個体へ共感を示したという経験的事実は、実験室実験において確立していないと考えている。本申請においては、ラットは“共感”に基づいて行動を起こすのか、他の仮説、とりわけ社交仮説を排除することは可能であるのかを、厳密かつ批判的に実験的検討を加えることで、共感仮説の真偽を見極めていくことが主な目的であった。

実験 1：E 字型迷路における選択場面を利用した、社交仮説、共感仮説の実験的検討の研究

いくらか手続きの詳細において変更点があるものの、概ね当初の目的どおりに 1 つ目の実験を実施し、初年度の秋頃にそれを完了した。その後、投稿論文の準備を行い、2018 年度に Journal of the Experimental Analysis of Behavior 誌へその成果が掲載された。

この実験の概要と得られた結果をまとめる。実験の目的は E 字型迷路中の選択場面において、社交仮説と共感仮説がそれぞれの予測に一致する結果が得られるかを確認することだった。フリーラットは E 字型迷路の中央のスタートボックスへ配置され、走路を走った先の T 字路で左右いずれかの小道を選択し先へ進むと、その走路の終点でゴールボックスへ達した。条件 1 では一方のゴールに拘束箱に入れられた拘束ラット（Restraint 選択肢）、もう一方は空のゴールボックス（Empty 選択肢）を配置した。条件 2 では一方のゴールには拘束されたラット（Restraint 選択肢）、もう一方は、開放された拘束箱とゴールボックスで自由に動き回ることのできるラット（Open 選択肢）が配置された。条件 3 では一方のゴールには開放された拘束箱とラット（Open 選択肢）、もう一方のゴールは空のゴールボックス（Empty 選択肢）とした。18 個体の被験体が、飼育室では 3 個体 1 ケージで飼育され、ケージごとにランダムにフリーラットおよび、Restrained ラット、Open ラットが選ばれ、実験終了まで役割は固定された。もしフリーラットが Restraint 選択肢を選択したならば、そのラットがゴールボックスへ進入する際に拘束箱の正面扉が開放された。フリーラットはゴールボックスに到達後そこで 1 分間過ごした後、別ケージに隔離し試行間隔 20 秒の後に再びスタートボックスへ入れられた。1 セッションの構成は強制試行 6 試行（左右各 3 試行をランダムな順に実施）に続く選択試行 6 試行となる計 12 試行であった。各条件を 8 セッション実施した。

条件 1 では Restraint 選択肢への選択割合がチャンスレベルから有意に高くなる結果を得た。この結果は社交仮説と共感仮説ともに予測していたもので、いずれの仮説とも一致する。条件 2 では Restraint と Open の間で選択割合は同程度となった。この結果は、共感仮説には否定的な結果であ

るが、両選択肢間で差が見られなかったことは第 2 種の過誤である可能性もある。一方、社交仮説は積極的な予測を持たない。なぜなら、どちらの選択肢を選んでも他個体と接触可能だからである。条件 3 では Open 側への選好割合がチャンスレベルから有意に高くなった。条件間比較を行ったところ、条件 2 の Restraint 選択割合より条件 1 の、また条件 2 の Open 選択割合より条件 3 の選択割合が有意に高かったが、条件 1 の Restraint 選択割合と条件 3 の Open 選択割合に差はみられなかった。したがって、本実験結果はいずれの条件においても社交仮説の予測と整合性を保つものとなった。その一方で、条件 1 は共感仮説と整合的であったものの、条件 2 の結果は否定的な結果となった。

本実験結果を受けてラットは共感に基づいて行動をしなかったと主張するわけにはいかない。そもそも“共感”が「無い」を実験的に例示することは論理的に困難である。しかし、共感仮説は条件 1 と 2 が同程度の選択割合を示すことを予測するが、結果は両条件間の選択割合に差があることを示した。一方、本実験のすべての結果は、社交仮説のみによって儉約的に説明が可能である。つまり、少なくとも本実験の結果は、フリーラットが拘束されているラットに共感を示したから解放へ向かったと解釈しなくとも、単に仲間のラットに会える選択肢をフリーラットは選んでいた、つまり社交の強化力という説明で十分であることを示している。このことは、人間の共感の動物モデルとしてのラットに疑問を投げかける結果である。

実験 1-2：E 字型迷路下における電気ショックを受けるラットへの救助行動の生起

実験 1 の完了後に、E 字型迷路を利用し、関連データを更に取得することを検討した。ここでは、さらに強力な嫌悪的場面として、ゴールボックス内に置かれた個体は拘束箱に拘束される代わりに、毎試行開始時に電気ショックを与えられた。

実験 1 で使用されたラットについて、各飼育ケージにおけるフリーラットの役割をそのケージ内の別のラットへ変更し、以下に説明する予備実験を行った。この実験では一方のゴールボックスにラットを入れ、各試行開始時に金属製の床板から電気ショックを与えた (Shocked 選択肢)。他方のゴールボックスは空であった (Empty 選択肢)。フリーラットが Shocked 選択肢を選んだ場合はゴールボックスの扉が開き、Shocked ラットは走路へと逃避することが可能となった。この実験では、Shocked ラットは明らかに嫌悪的環境に置かれており、その状況から逃げるための救助を必要としている。もし共感仮説が正しければ、フリーラットがそのゴールボックス側を選択する割合は極めて高くなるであろう (ただし、社交仮説においても、この条件における予測は同じ)。この予備実験を実験 1 と類似の手続きにより 10 セッション実施した。その結果は、はじめの数セッションでは Shocked ラット側の選択肢への選好が高かったが、セッションの進行とともにその選好が減少していき、最終的にはおよそチャンスレベルへと達するというものであった。

この結果はいずれの仮説とも整合的とは言えないため解釈が難しい。1 つの可能性は、フリーラットは社交ないし共感によって確かに動機づけられていたが、この実験では、それに加えて、電気ショックを受けたラットが発する悲鳴がフリーラットへの嫌悪刺激となっており、嫌悪刺激からの逃避だけでなく回避の動機づけも存在したかもしれない。当初、仲間のラットに会いたい (社交)、あるいは救いたい (共感) という動機づけが強く、または悲鳴音を止めるために、選択を行っていたが、セッションが進行するにしたがって、社交ないし共感の動機づけに馴化が起きて弱まり、嫌悪性回避の動機づけと拮抗するようになったのかもしれない。これは後付けの解釈以上ではなく、かつ嫌悪刺激の回避動機づけや馴化の可能性など、多くの仮定を据えなくては成り立たない複雑なものである。拘束箱へ拘束されることよりも明確にラットが救助を必要とする事態として電気ショックの使用を検討したわけであるが、多くの別要因を遡上に載せてまで利用する価値はないと判断し、この実験は 1 条件終了後に中止した。

実験 2：拘束箱への進入は嫌悪的ではなく、むしろ強化的であることの実験的例示

前項の実験 1-2 を発案した背景には Ben-Ami Bartal et al. (2011) が利用した拘束箱へ拘束されたラットに対する救助行動は、本当に拘束ラットが嫌悪的な状況に置かれていたから生じたのかについて十分に検討されてこなかったという認識がある。Ben-Ami Bartal et al. では拘束されることは嫌悪的であり、他個体からの救助が必要な状況であると仮定しているが、彼ら自身が拘束ラットから頻繁に悲鳴が上がることはなかったと報告している点は、拘束されることが嫌悪的ではなかったという解釈と矛盾しない。また、時折ラットが開放された拘束箱の中へ探索しに戻って来る姿は、先行研究 (Silberberg et al., 2014) においても指摘されていたが、本報告の実験 1 の実施時に申請者自身も観察した。これらの傍証は拘束箱の嫌悪性へ疑義を挟む余地を与えている。そこでこの点を調べるために本申請実験 2 を実施した。実験は 2018 年度に実施され、冬頃に終了し、論文にまとめ、学術誌に現在投稿中である。

実験 2 ① では、Sprague-Dawley ラットを用い、6 つのホームケージに 3 匹ずつ飼育されていたラットを、各ホームケージでランダムに 3 群へ割り当てた。0-min 群は、実験箱内で仕切られた小區画に 10 分間置かれ、その後仕切りが外され、解放された拘束箱を含む実験箱全体で 15 分の間、動き回ることができた。その際に、拘束箱に訪れた回数と時間を測定した。10-min 群と 20-min 群は、セッション冒頭の 10 分ないし 20 分の間、拘束箱に拘束され、その後、拘束箱の片ドアが開放され、15 分

間の観察が行われた。実験は 7 セッション実施した。その結果、10-min と 20-min 群は 0-min 群よりも頻繁に拘束箱へ入ることが確認された。また最初の 2 セッションでは 20-min 群ではより長く拘束箱内に滞在した。この滞在時間の差はセッション経過とともに失われていった。しかし、すべての群で全セッションを通して、拘束箱内で滞在した時間の割合は、ランダムに移動した際得られる割合より大きかった。この結果は拘束箱へ入ることに強化力があることを示唆している。

実験 2 ② のフェイズ 1 では、セッション開始時にラットは解放された拘束箱が置かれた実験箱に放たれ、一度、拘束箱に入り、そこから外に出てくると自動的に拘束箱のドアが閉じた。実験箱に設置されたレバーを押すと、拘束箱のドアが開き、自由に中へ入ることができた。フェイズ 2 では、セッション開始時にラットは拘束箱の中に 10 分間拘束され、その後、ドアが開き、そこから外に出てくると自動的に拘束箱のドアが閉じた。フェイズ 1 と同様に、レバーを押すとドアが開き、拘束箱の中へ入ることができた。結果、レバーを押してから拘束箱の中に入るまでの潜時はセッションの進行とともに短くなっていき、その反応と結果の随伴性を学習したことを示した。つまり、レバー押し行動の後続事象として、拘束箱へ入ることが強化的であったことを示唆している。

この一連の結果はこれまでの先行研究が、拘束箱に拘束されることがストレスであり、救助を必要としているとする前提を否定し、むしろ、拘束箱内に入ることが少なくともいくらかの強化力を持つという、全く異なる主張を生み出した。拘束箱内に入ることが強化的であることは解釈として妥当性があるのだろうか。生物学的に妥当であることを示唆する先行研究がある。本来、野生環境では、ラットは巣穴を掘って、その中で暮らす。Boice (1977) は実験室のラットを“人工的な”野生環境へ放し、その生態を観察した。Sprague-Dawley ラットは野生のラットと同様に巣穴を掘り、その中で過ごすことが確認された。掘られた穴の直径は 6 cm または 7 cm 程であった。この直径は、拘束箱の直径より狭い。もし巣穴を掘ったラット達はその直径では狭すぎてストレスとなり、嫌悪的であるとするれば、もう少し広めに掘ればいいのであるが、そのようにはしなかった。拘束箱の中は野生環境における巣穴と似ており、またラットにとって十分な広さであったことが示唆される。もしそうであるならば、本実験 2 の結論は何も不思議なものではない。

「“共感”のラットモデルの妥当性の実験的検討」のまとめ

本申請では、Ben-Ami Bartal et al. (2011) に始まるラットの共感に基づく行動の研究パラダイムをテーマとした。そのテーマは、人間の行動のメカニズムの理解に関わり、医療的応用へとつながる可能性があり、本当に共感のラットモデルとなり得るなら、その意義は非常に大きいだろう。また、直感的にわかりやすいテーマであるため学術分野を離れた実社会でも注目を集めやすい。しかし、それゆえ本当にラットが共感に基づく救助行動を行うという行動レベルの証拠を例示する必要がある、それは解釈に疑義のない明白なものである必要がある。これを検討することは、紛れもなく実験心理学者の仕事である。

本申請の実験 1 では、拘束箱に入れられ、身動きが自由にできないという“ストレスを感じる嫌悪的な”状況に仲間が置かれていることに対し、フリーラットが“救助に”向かうかについて、E 字迷路を利用して検討した。確かに、何もない空のゴールボックスが対置された場合には、拘束箱に拘束された仲間のいるゴールボックスへの選択割合は高かった。しかし、ゴールボックス内に解放された拘束箱と自由に動き回る仲間が対置された場合には、拘束箱に拘束された仲間のゴールボックスを選択する割合はチャンスレベルへと下がった。また、拘束箱を自由に動き回る仲間のいるゴールボックスと空のゴールボックスとの選択では、仲間のいるゴールボックスを選択する割合は拘束箱にいる仲間のいるゴールボックスを選択する割合と同程度であった。このことは、必ずしもフリーラットが共感に基づき“救助に”向かったのではなく（共感仮説）、単に仲間と社交の場で接触を持つことに動機づけられていた（社交仮説）という解釈と整合性がある。

本申請の実験 2 では、そもそも拘束箱に入ることが本当に“ストレスを感じる嫌悪的な”状況であるのかを実験的に検討した。実験 2 ① では、拘束時間経験の異なる 3 群（0 分、10 分、20 分）での、拘束箱への再訪とその際の滞在時間を測定したところ、拘束時間経験の長い 2 群の再訪回数がより多く、またいずれの群でも拘束箱への滞在時間はチャンスレベル以上であった。このことは、拘束箱への滞在が強化力を持つことを示唆した。実験 2 ② では、拘束箱へ入るためには、レバーを押して拘束箱のドアを開けなければならなかったが、ラットはレバー押しを獲得し、獲得後は速やかに拘束箱へ入っていった。この行動は、ラットが 10 分間の拘束を経験した後でも変わらず生じた。これは、拘束箱への滞在が強化子として機能し得ることを示した。

以上の実験的検討より、Ben-Ami Bartal et al. (2011) の実験パラダイムは、これまでのその論文への評価と反し、ラットが共感に基づき救助行動をすることを例示したものではないことを示唆した。換言すると、ラットが人間のような共感を示す行動を示したという実験報告は未だ得られていないといえる。しかし、このことから、ラットが人間の共感の動物モデルになりえないと結論づけるわけではない。ただ、その主張のためには、基礎的事実が不十分であり、おそらくは拘束箱を利用した実験パラダイムとは別の方法で行動レベルの証拠を示すことが必要となると思われる。本申請で示した研究結果は人間の共感の動物モデルとしてのラットに対しネガティブな証拠を重ねた。しかし、ややもすると熱を帯びがちなテーマへ、いくらか減速を迫ることは悪いことではない。あやふやな事実の上

に応用研究や学際研究を展開していくことは多くの無駄を生み出すことだろう。

その他の実験に関する遂行の報告

以上の報告のように、申請書に記載した実験計画に基づく研究へ専心し、海外派遣された2年間の間に、申請テーマへの結論を得た。その一方で、動物実験は否応なく各々の実験準備に時間を要するため、その間に別の研究を並行して実施してきた。本研究計画は比較心理学の領域の研究となるが、申請者の第一の専門は学習と行動の基礎研究であり、また受入教授もその長い研究歴において、比較心理学だけでなく学習の研究、行動経済学的研究などの業績を上げてきた。受入教授のその多岐にわたる研究領域に業績を持つことに敬意を抱き、この機会に申請書記載のテーマだけでなく、幅広く教えを乞いたいと考えていた（申請書における派遣先の選択の項へ記載した）。申請書記載研究以外にも多岐にわたる研究に携わり、非常に生産的な2年間を過ごしたことを、ここで簡単に各研究の概要を紹介することで報告したい。

・ハトの選択行動はチンパンジーと同様にナッシュ均衡し、かつマッチング関係を成立させる。

Martin et al. (2014) はマッチングペニーゲーム (Matching-pennies Games) で競合するチンパンジーの対戦結果がナッシュ均衡へ近似する一方で、同ゲームにおいてヒト同士では結果がナッシュ均衡から離れる傾向を報告した。彼らは、次の2つの説で結果を説明した：(1) チンパンジーの生態系がこのゲーム事態で必要な認知能力にアドバンテージを与える進化を引き起こした、(2) ヒトは言語能力の発達と引き換えにこのゲーム事態で最適に選択行動を行うために必要な認知能力を衰えさせる進化的なディスアドバンテージを被ってきた。

本研究では、ハトを利用して同様のゲーム事態で実験を行い、チンパンジーと同程度にナッシュ均衡へ近似することを示した。全く異なる生態系にあったと思われるハトが同程度の反応遂行を示したことは、チンパンジーの進化的アドバンテージで説明することが困難であることを示唆する。そのみならず、ハトの結果は並立強化スケジュールで一般的にみられるマッチング法則(2つの選択肢への反応比は各々から得られる強化比と共変する)の予測とも一致していた。また、Martin et al.のデータの再分析からチンパンジーにおいても同様にマッチング法則に従う傾向が得られた。さらに並立スケジュールにおけるマッチングを生じることが知られている累積効果モデル (the cumulative-effects model, Davis, et al., 1993) を利用して、マッチングペニーゲーム事態でシミュレーションを行ったところナッシュ均衡に達することが示された。このことは並立スケジュール事態で得られるマッチングと、ゲーム事態で得られるナッシュ均衡はいずれも共通した過程(例えば報酬の最大化)の働きから生じた派生的結果であることを示唆している。また、マッチングが短時間で急速に生じることを受けて、それは生得的な過程であると主張する先行研究があるが(Gallistel et al., 2007)、もしそれが正しいのであれば、ナッシュ均衡に達することはマッチングの付帯現象である可能性もある。また、様々な動物種が選択事態で一般的にマッチングを示すことは、ゲーム事態におけるチンパンジーが他の種に比べて認知的アドバンテージを示していることにはならないと考えられる。また、ヒトは他の種と異なり一般的にマッチングが見られにくいことが知られているが、多くの研究者はこれを、認知的な複雑さのためとみなしてきたことは、ヒトの認知的ディスアドバンテージ説と相反している。

本研究は2018年に Journal of Comparative Psychology へ掲載された。

・単純な Win-Shift, Lose-Stay 戦略が生み出す選択データが Wald-Wolfowitz の連検定を通過する。

Martin et al. (2014) では、ヒトを実験参加者とするインスペクションゲームを実施し、得られたデータについて、Wald-Wolfowitz の連検定に基づき、繰り返される各選択間は独立でランダムであることが保証され、その上で、個々人のデータが概して(チンパンジーのデータよりも)ナッシュ均衡から離れていることを示した。本研究では、インスペクションゲームを Martin et al. と同様に実施し得られたデータを分析した。データは Martin et al. と同様に連検定によって分析をしたところ、2選択肢間の選択交替数と、勝ち負けの事象生起数はいずれも、ランダムネスから生じる予測の範疇に大まかに収まった。しかし、次試行での選択肢交替確率は現試行までの勝ちの連の関数として高くなり、負けの連の関数として低くなっていった。この選択傾向は賭博者の過誤 (Gambler's fallacy) と整合的であった。そこで、Win-Shift, Lose-Stay 選択ルールを持つ両プレイヤーをプログラムし、インスペクションゲームをシミュレーションした。得られたデータで連検定を実施したところ、その選択データは確実に依存性を持ち、ランダムではないはずだが、その依存性を検出することができなかった。これは、ゲーム理論のパラダイムを用いた行動実験で、繰り返し選択間の独立性を示すためには連検定では不十分であることを示している。

本研究は Tripoli, Hachiga, & Silberberg として論文にまとめ学術雑誌に投稿中である。

Win-Increase ルールによるシミュレーションでマッチング関係が急速に獲得される

並立強化スケジュールの下でマッチング関係（各選択肢への反応比が各々から得られる強化比と共変する）が得られることは報酬の最大化を学習した結果であると推測される。しかし、2000年代以降非常に短期間のうちにマッチングが成立するために、マッチングは学習の結果ではなく生得的な過程であるという主張がなされるようになった（Gallistel et al., 2007）。本研究では報酬を最大化する反応パターンを学習する以前の訓練初期の選択反応は Win-Increase ルールによって生じると提案を行った。これは強化子を獲得するたびに、それを生み出した選択肢への選好が確率的に増加する、いわゆる効果の法則（the law of effect）のみを仮定した単純なルールである。シミュレーションを実施したところ、このルールに従った選択反応はわずか4強化子程度を取得する間にマッチング関係を成立させることが示された。シミュレーションの結果から、急速にマッチングが成立するためには、Win-Increase ルールに従いさえすれば可能であることが示唆された。

本研究はすでに論文にまとめ学術雑誌に投稿中である。

学習訓練初期は Win-Increase ルールによって、訓練後期は報酬の瞬時最大化によってマッチングが生じる

Silberberg et al. (1978) は並立スケジュール下でハトの訓練を60セッション行い、最後の10セッションのデータから系列統計量を分析したところ、瞬時最大化の予測と類似した傾向が生じたと報告した。Nevin (1979) は彼の1969年発表した論文から8-14セッションのデータを再分析したところ、瞬時最大化の予測と矛盾した結果となったことを報告した。本研究ではラットを利用し75セッションにわたる訓練を行い、そのうちの8-14セッションのデータと最後7セッションのデータの系列統計量を比較した。訓練初期(8-14セッション)の結果はNevinの傾向と類似の傾向を示し、訓練末期(最後7セッション)までにSilberberg et al. の報告するような傾向へと変化していった。また、Win-Increase ルールを利用したシミュレーションによって生成したデータの系列統計量はラットの訓練初期の結果と類似のものであった。この結果はラットが瞬時最大化反応パターンを学習する前の訓練初期には、Win-Increase ルール（強化された選択肢への選好が高まる）にしたがって選択を行い、訓練末期では、ラットは瞬時最大化する反応パターンを学習したことを示唆している。したがって、マッチング関係は訓練の初期から末期まで得られ続けたが、それは反応系列を生み出す2種類のルールが関係していたことによると主張を行った。

本研究は投稿準備中である。

強化子省略後の選好パルスは選択行動の連構造の副次効果か、強化省略による局所的効果か。

選好パルスは強化子が提示された直後にその強化子が随伴した選択肢への相対反応率が一時的に高まり、時間経過とともに減少していく現象である。しかし、McLean, et al. (2014) は、その現象は、強化子の局所的効果ではなく、選択反応の連構造の副次的効果として起こる可能性を指摘した。もしそれが正しければ、強化子の局所効果が明白に存在しない場面である強化子省略の後にも選好パルスが生じるはずである。ただし、強化省略の後の局所的選択へ、もし強化子省略自体の局所効果が生じる可能性があれば、これは妥当な解釈とはいえない。本研究では、時折の強化子省略がその後の選択行動へ持つ影響を評価するために、並立 VI 30 s VI 30 s スケジュールの後に並立 VI 30 s 消去 スケジュールが2回続く多元スケジュールを繰り返し、その下でラットの選択行動を測定した。並立 VI VI スケジュール下の50%を強化省略し、強化子提示時との選択行動の違いを分析した。その結果、強化されたレバーが次成分では消去レバーである場合、強化子提示時と省略時のいずれにおいても選好パルスが観察された。しかし、強化されたレバーが次成分で再びVIレバーとなる場合には、強化子省略後には逆パルス（強化省略の結果が随伴していない側のレバーへ一時的に相対反応率が高まる）が観察された。加えて、強化省略後10秒までの選択は、その後の選択よりも変動的であり、また強化後10秒までの選択と比較してもより変動的であった。また、強化省略後10秒までの絶対反応率はその後の反応率よりも高かった。これらの特徴は強化子省略の操作自体が選択行動に対して局所的効果を持ち、それゆえに消去選択肢への選好パルスとVI選択肢への逆パルスが生じたことを示唆する。

本研究は慶應義塾大学・三田哲学会出版の「哲学」142集の坂上貴之教授退職記念号へ掲載された。

以上。