

論文 / 著書情報
Article / Book Information

題目(和文)	意思決定における限定合理性の数理モデル－リスク回避と推論不完全性の記述と解析－
Title(English)	
著者(和文)	今野直樹
Author(English)	naoki konno
出典(和文)	学位:博士(理学), 学位授与機関:東京工業大学, 報告番号:甲第9268号, 授与年月日:2013年9月25日, 学位の種別:課程博士, 審査員:木嶋 恭一,出口 弘,金子 宏直,猪原 健弘,中丸 麻由子
Citation(English)	Degree:Doctor (Science), Conferring organization: Tokyo Institute of Technology, Report number:甲第9268号, Conferred date:2013/9/25, Degree Type:Course doctor, Examiner:,,,,
学位種別(和文)	博士論文
Type(English)	Doctoral Thesis

学位請求論文

意思決定における限定合理性の数理モデル
ーリスク回避と推論不完全性の記述と解析ー

今野直樹

東京工業大学大学院社会理工学研究科

価値システム専攻

指導教員 ： 木嶋恭一教授

2013 年 6 月

目次

第 1 章 序論	2
1.1 本論文の背景と目的	2
1.2 本論文の構成	4
第 2 章 限定合理性と不確実性	6
2.1 限定合理性	6
2.1.1 限定合理性の分類	7
2.2 問題固有の不確実性	8
2.2.1 期待効用理論	8
2.2.2 期待効用理論の拡張	10
2.2.3 生態学における不確実性と進化ゲーム	10
2.3 意思決定者の能力により生起する不確実性	11
2.3.1 均衡精緻化の文脈におけるエラーの考慮	11
2.3.2 戦略形ゲームにおける有限量の行動エラーのモデル化	11
2.3.3 展開形ゲームにおけるの推論エラーのモデル化	12
2.3.4 実際の実験結果との比較	12
2.4 本研究の展開	13
第 3 章 新たな不確実性評価指標の提案	15
3.1 リスク回避的評価	15
3.2 複製モデル	16
3.3 タイプの支配	17
3.4 期待到達時間	19
3.4.1 支配と ERT の関係	21
3.4.2 有限量に対する ERT	22
3.5 3 章のまとめ	25

第 4 章 不完全な推論に関する新たな数理モデルの提案	26
4.1 問題状況の複雑性と限定合理性	26
4.2 限定合理的な意思決定のモデル	26
4.2.1 エラーに関する数理的モデルの先行研究	26
4.3 推論不完全性を考慮した展開形ゲームにおける意思決定のモデル化	28
4.4 二つの具体的な粗い推論モデル	32
4.4.1 logit 関数に基づく粗い推論モデル	33
4.4.2 指数エラーに基づく粗い推論モデル	34
4.5 ムカデゲームへの適用	35
4.6 現実の意思決定における推論不完全性	39
4.7 4 章のまとめ	42
第 5 章 結論と今後の課題	43
5.1 結論	43
5.2 今後の課題	44
参考文献	46
謝辞	50

第1章 序論

1.1 本論文の背景と目的

本論文の目的は、以下の三点である。第一には、従来は問題固有の要素とみなされていた不確実性に関して、意思決定者の能力という異なる分類軸を提案することである。第二には、意思決定問題における不確実性に関して新たなリスク回避的な評価指標を提案することである。第三には、意思決定の推論の際に限定合理性を考慮したモデル化を行い、合理的な解と実験結果が乖離しているとされるムカデゲーム状況 [41] に関して、推論能力の限定性という視点から従来よりも強力な説明を与えることである。

人間の意思決定は通常、現実直面した多くの要素をはらむ問題状況に対して（１）問題状況のモデル化を行い（２）問題状況の構造を踏まえて計算したうえで代替案を比較検討し（３）選んだ代替案を実行する。という三つのプロセスからなる [46]。これらの意思決定の各段階において、様々なタイプの困難が存在する。

まず、（３）の行動過程においては、通常の数理的な意思決定モデルにおいては、最適な行動を確率 1 でとると仮定されることが多い。しかしながら、一例として車の運転で、曲がることのできるつもりでハンドルを切ったにもかかわらず事故につながることもある。このように一般には、意思決定の計算結果を確実に遂行できるとは限らない。このような行動の際のエラーに関する視点は、経済学においては、震える手完全均衡 [45] やプロパー均衡 [34] など、エラーに対する頑健な均衡の解析という文脈で研究されている。また Quantal response model [31] では、さらに踏み込んで、有限確率でのエラーが不可避の状況における均衡の存在を議論している。

以上の点を踏まえて不確実性という言葉をもう一度考察してみる。不確実性は数理的な意思決定においては、Knight [28] の分類に代表されるように、問題固有の性質として自然の状態として表現されることが多い。しかしながら、問題固有の性質という面だけではなく、それを問題として取り扱おうとする意思決定者の行動精度にも依存する面も存在すると考えている。さらに、意思決定者の能力が極めて高い時は、エラーの起こる確率が無視できるほど小さいため、結果として従来のモデルの結果とほぼ一致するとみなせる。このような認識論的な立場のほうが、問題に関するより深い理解が可能になると考えている。

これに加えて、意思決定の認識や推論のフェイズにおいても意思決定者は誤りを起こすこともある。この認識や推論フェイズにおけるエラーを具体的に扱った数理的なモデルは現時点では非常に少ない。しか

しながらこれらのエラーにより、意思決定者は目の前の問題に対して、客観的に記述されている状態から得られる解と、異なった行動を合理的と認識してしまう可能性もある。これまではあたかも客観的な神の視点から見た場合のように、このような結果は単なるエラーとして記述されてきた。しかしながら、人間が神のような完全な認識能力及び推論能力を持つことは原理的にあり得ないと思われるので、これらフェイズにおけるエラーも、問題の不確実性の部分に入れて記述する立場のほうが、より一般性の高い視点となっていると考えている。

このようなエラーもモデルに盛り込むならば、エラーに関する数理的な性質について議論を行うことができる。ここで重要なのは、エラーが起こりうるが、それがどのような方向に起こるかわからないという立場で止まってしまうならば、単に変数を増やしただけで分析に新たな深さを加えることはできない。たしかに一般的なすべての状況に適用できるエラーに関する定理を確立することは極めて困難である。ただ、問題状況がある程度絞り込んだ場合には、エラーの方向性は完全にランダムな分散となるとは限らず、エラーの起きやすい方向に共通性がみられる場合もあると考えている。一例として、認識フェイズにおいては、矢印の向きを変えることに同じ棒の長さが異なって見えるような錯視が起こりやすい図形があることが知られている。また推論フェイズにおいては、一般に演算の桁数が大きくなる方が間違いが起きやすい。タイピングを行う際は、タイピングの速度を上げるほどミスの頻度が高くなる。

このように目の前の問題が原理的には不確実性が全くない状況であったとしても、結果に不確実性をもたらす場合があり、安全工学では、「人間はミスするもの」「人間の注意力はあてにならない」という前提から、フールプルーフが信頼性工学における安全設計の基本的で重要な概念となっている [17][37]。このように人間のエラーに関しては、どのような傾向が起こりやすいのかという分析とともに、またそれに対するシステムの安定性も重要な点として議論されている。

このように限定合理性が引き起こす影響も不確実性の一つとみなすならば、それらの不確実性に対するマネジメントも必要となる。従来の経済学における不完備情報の理論では、認識の不完全性に対して、世界に対する真の状態がわからなくても、可能性のある状態とその確率が正確に分かっているという条件の下では数学的に豊かな結果を得ている [21]。その一方で、不確実性に対する評価に関しては、von Neumann の提案した期待効用理論 [50] を代表に、そこから派生した様々な公理系が提案されている。

しかし経済活動の際には、保険などに代表されるように、リスク中立よりもリスク回避的な選好を仮定することが一般的である。また特によくわからないものや大きな脅威に対するほどリスク回避度が高まるという実験報告もある。そこで、この不確実性をマネジメントするにあたり、どの程度リスク回避をすべきかという疑問が出てくる。本論文では、この点に関してまずは最も基本的な確率過程に分析対象の変動が従う状況に注目する。この際に余分な複雑性を回避するために、まずは上で述べたような意思決定者によ

る影響が存在しない状況に限定する。実際に天候に関する不確実性は、意思決定者が仮に完全に合理的であったとしても排除できないような、意思決定問題が固有に抱える重要な要素である。この状況下で、一つの直感的なリスク回避的指標を提案し、その上で、その評価基準が従来のリスク回避的な評価基準の拡張となっていることを示す。

次に、オセロやチェスといった完全情報ゲームを考えてみると、これらのゲームは、原理的には上で説明したような、問題が固有に抱える不確実性は存在せず、原理的には両者が最善を尽くした場合には、どちらかのプレイヤーに必勝手順が存在するか、あるいは引き分けになることが知られている。しかし、それが具体的にどのような戦略の組になるのか計算することや、その手順を実際に誤りなく実行することは現実の人間の能力では不可能である。従ってこのようなタイプの計算に関する困難さは、意思決定者の能力に依存するものである。このような問題状況においては、現実の人間のプレイヤーは、例えば将棋や囲碁などのプロも、最善を常に追求した戦略を実行しているわけではなく、自らが読み間違いする可能性を織り込んだ上で、安全性を重視する戦略を採用する場合もある。このように推論リスクは、これまで数理的な研究は非常に少ないにもかかわらず、現実の意思決定に於いては、各分野のプロと言えるほどの推論能力をもったプレイヤーにとっても、無視できないほど重要な要因である。

そこで、本論文では意思決定者が自分では合理的に行動しようと推論しているつもりでも、推論エラーが起こりうるような状況のモデル化を行い、その妥当性を現実の被験者を用いた実験結果と比較する。

以上を改めてまとめると、

1. 対象が増殖をしていくような複製状況において、直感的な一つの意味決定基準を提案し、それが従来の議論の自然な拡張となっていることを示す。
2. 完全情報のゲームにおいて、意思決定者の推論能力が不完全であった場合にはどのような事が起こるのかということを記述するモデルを提案し、従来のゲーム理論の解概念と実際の人間に対する実験結果の乖離が大きいムカデゲーム [41] に適用し、既存のモデルよりもよく適合していることを示す。

1.2 本論文の構成

本論文の構成は以下のとおりである。第2章では、これまでの先行研究を通じて限定合理性に関する描写の議論を紹介した上で、不確実性に関してどのような視点からの研究が行われているのかということを紹介する。

第3章では、対象が複製していく状況において、期待到達時間という概念を提案する。その上で期待到達時間最小化は、期待値最大化と一般には異なることを示すと共に、目標を無限大として場合には期待到達

時間を最小化する戦略と、他のタイプを支配する戦略が一致するということを示す。さらに投資問題に適用することによりポートフォリオの有用性をこの評価基準が確かに数値的にも評価できることを示す。

第4章では、自分の推論能力が完全であると信じているような状況に関して、完全情報の展開形ゲームにおいてモデル化を提案する。その際には利得の差が小さければ小さいほど、また遠い未来の状態になればなるほど推論エラーが起こりやすいという仮定を採用する。さらにその上でこのモデルをムカデゲームに適用する。ムカデゲームの実験においては、非合理的な行動が高頻度で観察されることが知られている。これまでは、利他性などで説明されることが主流だったが推論エラーによっても非合理的な行動の説明が可能であることを示す。

第5章では結論と今後の課題を述べる。

第2章 限定合理性と不確実性

2.1 限定合理性

限定合理性とは、合理的であろうと意図するけれども、認識能力の限界によって、限られた合理性しか経済主体が持ち得ないことを表す。この概念は、1947年に H.A.Simon が『Administrative Behavior』[46]で提出した人間の認識能力についての概念である。

現実の意思決定問題では、状況認識が最初の課題となる。この過程において状況認識を正確に誤りなく行っていて、その後の推論や行動においてもエラーが無いと仮定するならば、効用最大化原理に従って行動することは極めて合理的である。

しかし、現実の意思決定ではそれらの仮定が完全に満たされることはないため、不完全な目的定義上での最大化探索になり、結果的に望ましい選択とは必ずしもならない。また最適化が意味を持つためには、全ての代替案が同一の効用関数で計測されて、それぞれの効用の値に、首尾一貫性と推移性が求められる。しかしながら D.Kahnemann と A.Tversky の実験 [51] によると、それらの仮定の現実的妥当性も疑問が投げかけられている。

このような知見に対して Simon[46] は満足化基準を提唱し、意思決定者は、ある一定の希求水準以上の効用を得られる代替案があれば、それを選択すると仮定した。全ての代替案を正確に比較した上でそれらの中で最大なものを選択するという最大化基準とは異なる哲学に基づいている満足化基準による意思決定が現実的な行動原理であると Simon は主張している。

現実の意思決定において、これらの違いを考えてみる。代替案を選択することにより発生する結果の効用関数を考えることや、それらの比較を行うことは意思決定者の脳に負荷がかかりまた実際の時間も消費する。従ってそれらが大きくなるような複雑な問題では無視できないコストであると考えられる。それらのコストも考慮したより大きな問題として定式化した場合には、全検索が必要な最大化は、原理的に不可能な場合もあるであろうし、また仮に可能であったとしても、コストに見合うほど期待値の向上が見られず、適当なところで探索を打ち切った意思決定よりも望ましい結果になるとは必ずしも限らない。

ある程度満足のいく代替案が発見できたならばそこで思考を打ち切り決定するという行動は、追加的な探索から得られた最適解と元の最適解を比較した際の期待値の向上分と探索コストを比較した上での最適

化行動として解釈しなおすことも原理的には可能である。しかしながらこのようなモデル化は、主観的な経験により今後の探索の影響を評価する必要があるため一般的な分析は困難である。そこで限定合理的な状況すべてに適用できる理論の構築よりも、限定合理性の分類を行う動きが出てきた。

2.1.1 限定合理性の分類

Simon は合理性を、実質的合理性と手続き的合理性に分類している。(図 2.1 参照) 実質的合理性とは、与えられた制約条件の下で、目的に照らして行動が適切である時に、その行動は実質的に合理的であると呼んでいる。手続き的合理性とは、行動が十分に考えられて計画されているときに、その行動は手続き的に合理的であると呼んでいる。

Simonによる合理性の分類

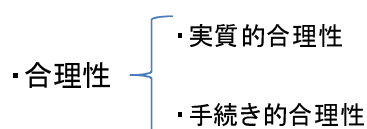


図 2.1: Simon の限定合理性の分類

一方、O.Williamson[54]によると、合理性には3段階の強さがある。最も強いものは最大化計算ができるという意味の合理性であり、通常のミクロ経済学が仮定しているものである。逆に最も弱いものは有機的合理性と呼ばれるものであり、貨幣、市場、所有権など、誰も計画的にそれを作り出そうとしなかったにも関わらず発生した制度に関する合理性である。限定合理性はその中間である。さらに、これらの合理性の外には非合理性がある。

上記二つの分類は合理性という言葉の意味について注目した分類となっている。

これに対し、金子 [7] は、分類した後のセグメントすべての事象について当てはまる性質を求めようとする場合には、上記の分類は適さないため、限定合理性を分類したうえで有力な知見を導くためには、プレイヤーの計算、推論能力の限界に限定して用いるべきであるとしている。(図 2.2 参照)

本論文ではこの立場から、特に完全情報の展開型ゲームにおいてプレイヤーの計算、推論能力が限定的である場合に、不確実性を用いたモデル化を行う。このほかに、不確実性の中には、従来から扱われてきた問題自身が固有に保持している不確実性も存在する。

金子による合理性の分類

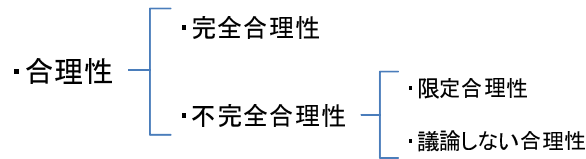


図 2.2: 金子の限定合理性の分類

2.2 問題固有の不確実性

不確実性に関して F.Knight[28] は、起こりうる状態の候補と確率はわかっているが、具体的に何が起こるのかわからない場合（リスク）と、起こりうる候補すらわからない状態（真の不確実性）を分類した。厳密な意味では現実的には後者のタイプの不確実性も多いが、確率を仮定しない場合には深い意思決定分析が困難である。前者のリスク状況に関しては、幅広い分析がなされているが、最も代表的なものは、以下に紹介する期待効用理論である。

2.2.1 期待効用理論

一例として宝くじを買うべきか否かという問題を考えてみると、購入した場合にどのような結果が生起するのかということに関しては一点で予測することはできず、不確実性が存在する。このような不確実性を伴う問題状況に関して最も広く受け入れられている概念は期待値や期待効用である。このような期待効用に関する議論は von Neumann と O. Morgenstern の『Games and Economic Behavior』 [50] の中でゲーム理論のための基礎概念として提案された。

期待値は非常にわかりやすく強力な概念であるものの、金銭に対してそのまま適用した場合に以下のような問題も生じる。これはサンクトペテルスブルグのパラドックスとして良く知られる例であるが、いまコイン投げを何度も繰り返すゲームの参加費としていくら支払うことが妥当なのか考えてみる。まず一回目に投げられたコインのが表ならば 2 円をもらってゲームは終了する。裏ならばもう一度コインを投げて 2 回目に表ならば 4 円をもらってゲームを終了する。以下同様にして n 回目に表がでた場合は 2^n 円支払うものとする。ここでゲームの回数の上限がない場合にはこのゲームの期待値は無限大となるので、有り金全てはたいても、期待値的にはこのゲームに参加するのは合理的という結果となる。しかしながら現実的にはこのような意思決定を行う人は多くは無いと考えられる。その理由としてベルヌーイが与えているのは人

間の効用は金銭に対して逓減するからであるという説明である。一例としては、確実に1万円をもらうほうが、確率1/2で2万円、確率1/2で0円という籤よりも望ましいと考える人間が多いという状況である。

このような効用の概念に関する議論の立場としては基数的効用と序数的な効用という2つの立場がある。前者は効用は実数で計測できて2倍の効用というような演算が可能であるという立場である。後者は効用は代替案同士の比較によってどちらが大きいのかは決定できるが基数的な場合のような演算はできないという立場である。序数的な効用表現のほうが一般的であるもののその一般性のために多くの命題を導くことは、基数的な効用の場合と比較すると困難である。

なお先のサントペテルスブルグのパラドクスの例においては、 x 円に対する効用が $\log x$ で与えられる基数効用を仮定すると、このゲームをプレイする時の期待効用は有限の $\log 4$ となる。なおこの対数の関数形の効用は典型的なリスク回避的効用関数の一例である。

なお X を確率変数として、 U を効用とするときに、確率変数の期待値を取ったものの効用である $U(E(X))$ と確率変数それぞれの効用の期待値である $EU(X)$ を比較して、

1. $U(E(X)) > EU(X)$ ならばリスク回避的
2. $U(E(X)) = EU(X)$ ならばリスク中立的
3. $U(E(X)) < EU(X)$ ならばリスク愛好的

と呼ばれる。また

$$U(E(X) - \alpha) = EU(X)$$

となる α に対し、 α をリスクプレミアムと呼ぶ。リスク回避的選好には正の α が、リスク中立的選好にはゼロの α が、リスク愛好的選好には負の α が対応する。効用曲線の曲率が大きいほどリスクプレミアムの絶対値は大きくなるので、効用曲線の曲率をリスク回避やリスク愛好の強さの測度として採用することができる。通常 $-U''/U'$ を絶対的危険回避測度、 $-E(X)U''/U'$ を相対的危険回避測度と呼ぶ。

なお期待効用理論において採用されている、独立性の公理に関してはM.Allais[8]やD.Elsberg[19]の実験による反例などから、現実的には必ずしも成り立たないという意見もある。以下はAllaisの反例としてよく知られる例である。

1. くじA: 確率1で1万円
2. くじB: 確率0.89で1万円、確率0.1で5万円、確率0.01で0円
3. くじC: 確率0.89で0円、確率0.11で1万円

4. くじ D : 確率 0.9 で 0 円, 確率 0.1 で 5 万円

このような 4 つのくじに対して, A と B のどちらのくじを好ましいと思うのか, また C と D のくじのどちらを好ましいと思うのかという二組の比較を被験者にアンケートをとってみると, A と D を好ましいと回答する被験者が多かった.

独立性公理に従うならば, くじ A のほうがくじ B よりも望ましいと考えるならば, C と D の比較においては C を選好しなければならない. また同様にくじ D のほうがくじ C よりも望ましいと感じるならば, A と B の比較においては B を選好しなければならない. この問題に関して回答するために, 期待効用理論のいくつかの公理を緩めることにより, 新たな理論が様々な形で提案されている.

2.2.2 期待効用理論の拡張

この問題に対して Tversky[51] は人間の効用は基準点の前後で傾きが急になり, それから離れると傾きは逓減していく. さらにその効果は負の効用に関する方が大きいというプロスペクト理論を提案している. 上の例においては A と B の比較においては高確率で 1 万円がもらえるのでそこが基準になり, そこから 0.01 という微小な確率でも 1 円ももらえないという損失をかなり重要な要素として評価するため, A を B よりも選好したが. しかし, C と D の比較においては, 0 が基準となるため, そこからの増分に関する比較においては期待値的な判断に近く D のほうを選好すると説明される.

これ以外にも期待効用理論の拡張は様々な形で試みられていて, Einhorn[18] は利得の線形性は保持するものの, 確率に関する線形性を緩和した主観的確率論を提案している. また Bell[11] はリグレット理論で, 実現しなかった利得の影響を考慮することにより特徴づけを行っている. Machina[29] は選好関数という概念を導入することにより結果と確率の相互依存性を考慮する一般化された期待効用理論を提案した. Gilboa[23] は期待効用基準と辞書式順序を折衷することにより maxmin 期待効用理論を提案した. これらの理論は説明力を優先するために期待効用理論よりも勝る一方, 規範的な予測能力を持たない上, 扱いが複雑という問題も他方で抱えているため現実問題に関する適用は今の時点ではまだそれほど多くはない.

2.2.3 生態学における不確実性と進化ゲーム

以上の議論は経済学の分野から発生したものが多く, 他方で生態学の分野においても Maynard Smith[47] はゲームにおける利得を子孫の数の期待値と解釈し, 戦略が子孫に遺伝すると仮定した際にどのようなタイプが進化の過程において最も適応していくのかという問題に関する定式化を行った. どのタイプと次期にゲームを行うのかということに関しては, 対象全体に占める割合に比例すると仮定されるために不確実

性が存在するものの、このとき獲得した利得に応じて次期に採用する個体がそれぞれ決定されるため高い期待利得を上げる戦略ほど次期の採用率は大きくなると仮定される。これを微分方程式で定式化したレプリケーターダイナミクスにより種の挙動を記述した結果として他の戦略に駆逐されないような進化的安定戦略 (Evolutionary Stable Strategy 以下 ESS) の概念を提案した。さらに Weibull[53] はそれらの概念を数理的に厳密な形で定式化を行った。進化的なゲーム理論においては各個体は詳しく相手のタイプを先読みして行動するというよりは、自分のいままでの経験や周囲の個体の選択により次期の行動を決めるという、限定合理的な意思決定スタイルを仮定されることが多い。

2.3 意思決定者の能力により生起する不確実性

2.3.1 均衡精緻化の文脈におけるエラーの考慮

以上は不確実性という意思決定問題が固有に抱える困難さに着目した議論であった。これに対して意思決定モデルとしては厳密に定義できて合理的な行動が原理的には定義できるような状況においても、意思決定者の能力の限界により、実際には必ずしもそれを推論することができるとは限らない。

まずは合理的な行動に関する概念としては経済学などでは均衡概念に代表される。特にゲーム理論では先読みの立場からナッシュ均衡 [36] が提案され、さらにそれらが複数ある場合にはどれがより妥当な均衡なのかという問題意識から均衡解の精緻化というパラダイムの研究成果が続いた。L.Selten は、震える手完全均衡 [45] で、相手が他の戦略を誤って選んでしまう可能性を考慮した場合の安定性という観点から均衡を比較して安定性を議論した。R.Myerson は、プロパー均衡 [34] で、均衡からの逸脱は、利得が低くなるような逸脱ほど起こりづらいという点をから安定性の更なる絞込みを行った。これらは、相手がエラーする確率が微小でも存在するならばそれに対する最適性を考慮することにより、より合理性の高い均衡を絞込む立場であった。しかしながらこのようなエラーというのは極限を取った議論となるため、自分の戦略の安定性を議論する際には相手の微小な確率でのエラーを考えるために支配される戦略が採用される確率は理論的には 0 となるものである。

2.3.2 戦略形ゲームにおける有限量の行動エラーのモデル化

これに対して、実験経済学の分野から Mackelvey and Palfrey[32] は、「利得が高い戦略ほど高確率で選択したがるが、最適な戦略を確率 1 で選ぶことはできない」という仮定のもとで、各プレイヤーが行動する Quantal Response 均衡 (QRE) を提案した。このモデルにおいては震える手完全均衡やプロパー均衡の場

合とは異なり、支配されている戦略でもある程度の確率で採用することを仮定している。さらに Quantal Response が logit 関数で表される場合には、合理性のパラメーターが 0 の際には完全にランダムなプレイヤーと一致して全ての戦略を等確率で取り、また合理性パラメーターが無限大の場合には完全に合理的なプレイヤーと一致してでナッシュ均衡が取られるモデルとなっている。また展開形ゲーム状況においては Mackelvey and Palfrey[33] はゲームをエージェント正規形により標準形ゲームに書き換えることにより QRE を適用し、ムカデゲームにける実験結果 [31] に適用しモデルの妥当性を検討した。ここでムカデゲームとは、部分ゲーム完全均衡の妥当性へ疑問を投げかけるために Rosenthal によって提案された二人完全情報の有限展開型ゲームで、部分ゲーム完全均衡がパレートの最適ではない結果となる代表的な例である [41]。

このもともとの QRE のモデルでは、パラメーターにより原理的にはほとんどの結果を表現できるため説明能力は非常に高い。しかしながら、その状況が与えられた際に具体的に何が起こるのかということは決定できないため、正規 Quantal Model など特定化のモデルも提案されている。

2.3.3 展開形ゲームにおけるの推論エラーのモデル化

以上で紹介したようなエラーに関するモデル化は基本的には全て（3）の決定の際に起こるエラーと解釈される。一方 Heifetz[25] らは、各意思決定ノードでは選択肢が 2 つしかないような完全情報の展開形ゲームにおいて推論エラーのモデル化を行った。彼のモデルにおいては推論エラーの確率は一定であるとしている。この条件の下でプレイヤーは通常と同様に後ろ向き帰納法により合理的な戦略の探索を試みるものの、その過程においてエラーが発生したまま推論を続けた上で決定を選択するモデル化である。推論に関するエラーは決定に関するエラーとは異なり、個々が微小な確率であっても、推論過程の中で蓄積されて結果的には高確率で非合理的な行動を取ることもあるということを示している。

2.3.4 実際の実験結果との比較

このような理論の成果の一方で囚人のジレンマのような支配戦略が存在するようなゲーム状況においても、実際に被験者を用いて実験を行ってみると [13]、均衡以外の戦略が少なからぬ割合で選択されることが知られている。一回限りのゲームにおいては両者が裏切ることが唯一の均衡である。しかしながら、実際には協力戦略が、無視できないほどの確率で観察される。さらに非合理的な協力戦略の選択確率は、たとえ拘束力がなかったとしても選択前に事前に会話を許すことで上昇することが知られている [13]。この実験結果に関しては明らかにエラーの影響だけとは思えないほどの頻度で、非合理的な戦略が採用さ

れている。このような事前の拘束力のない交渉がゲームの結果に与える影響は、ゲーム理論の現在の枠組みでは表記できない要素なので、このような経験的な事実をどのように扱うべきなのかということは重要な問題と考えられている。

また囚人のジレンマの利得の本質的な構造が保存されていたとしても、両者が協力した場合の利得が非対称な場合には、利得が少なくなる側のプレイヤーのほうが、合理的な戦略である裏切りを選択する頻度が高いことなども知られている。しかしながらこのような実験において実際の人間行動を描写するアプローチにおいては、通常獲得した利得に応じて金銭報酬が支払われるものの、それでもなお実験状況において実際に体感する効用と完全には一致するとは限らないため、理論事態の問題なのか実験環境の問題なのか議論がある場合も多い。一例として報酬が100円単位なのか、100万円単位なのかでは、理論的には効用のアフィン変換からの独立性より、両者に差がないが、現実には少なからぬ違いが発生することが想像される。この原因にはゲームをプレイする際に利他性などの感情によって、100円単位以上の外部的な効用が発生することもあるという問題 [31] や、同じ作業の繰り返しには被験者が飽きを感じて違う行動をとろうとするなど設計的な問題 [10] もあるといわれている。

ムカデゲームにおいては McKelvey and Palfrey [31] が実験を行った結果、均衡解とは著しく異なり、非合理的な協力行動が少なくない頻度で観察された。この話題に関しては第4章で詳しく紹介する。

2.4 本研究の展開

図 2.3 は、限定合理性と不確実性に関する分類をまとめたものである。

新たな軸 既存の軸	限定合理性の影響なし	限定合理性により発生
不確実性無	計算量理論	4章
リスク	3章	これからの分野
真の不確実性	支配, Maxmin, など	これからの分野

図 2.3: 限定合理性と不確実性に関する分類

本論文の第3章では、プレイヤーの合理性が完全で、問題自身が固有の不確実性を持つ状況において、一つの直感的な意思決定的基準を提案し、その基準と先行研究の基準を明らかにする。さらに第4章では、プレイヤーが限定合理的であり、問題自身が固有の不確実性を持たない状況において、推論エラーに関するより一般的なモデルを提案し、その上で先行研究で扱われたムカデゲームの実験結果との対応を比較する。

これ以外の箇所に関しては、まずプレイヤーが完全合理的であって、問題状況が不確実性を持たない場合には、解くためにはどの程度の時間や計算量が必要なのかということが焦点となる。この点については情報科学における計算量理論で研究がおこなわれている。また限定合理的かつ問題状況に不確実性がある場合は、限定合理性についての数理的なモデルが不確実性のない状況の分析がまだ発展途上のため、それより難しいモデル化となるために、これからの分野である。また、それ以外の合理性の場合についても、少なくとも数理的に統一的に扱うモデルはまだ提案されていない。この分野を数理的に扱うには、また別の条件を加えて問題の特定化を行う必要があると考えている。

第3章 新たな不確実性評価指標の提案

3.1 リスク回避的評価

2.2.3 で紹介したように、不確実性下の意思決定における代表的な評価基準である期待効用とは別に、進化ゲーム理論でも生態学的な視点から不確実性下の状況で、どのようなタイプが優位となるのかという議論が行われている。また特に生命の複製状況に注目すると環境要因により全個体が一律に影響を受けるタイプのリスクと、各個体に内在し個体ごとに独立して発生するリスクがあると考えられる。Robson[39] は個体数の期待値の最大化と種の支配が必ずしも一致しないことを示し、個体リスクに関しては期待値最大化戦略が支配戦略となる一方、環境リスクに関しては、対数を取ったものの期待値最大化戦略というリスク回避的なタイプが支配的になることを示した。

また経済学的な問題設定状況では、複製状況において環境リスクに対して、Blume and Easley[12] は進化的な視点から市場のリスクに対して進化的に最も適応度が高いのはリスク中立なタイプとは限らず、対数をとったものの期待値最大化戦略がやはり最も適応することを示した。

投資問題を考える際には、通常はリスクのない定期預金のような商品と、リスクのある株式に対してポートフォリオを張ることが一般的な認識となっている。このとき株式の期待値が定期預金の期待値を上回っているならば全額株式に対して投資という選択を行ったほうが良いのであろうか？この考え方は一般には否定されている。期待値が正でも高確率で破産してしまうような数学的な例が知られているため、期待値的な意味での望みしさだけではなく最頻値としての望みしさも重要視されている。一般には期待値の最大化と最頻値の最大化は一致するとは限らないが、この両者の違いに最初に言及したのは Kelly[27] である。

Kelly は、好みの金額を賭け、勝てば賭け金の R 倍がもらえ、負ければ賭け金が没収というゲームにおいて、1 回ごとに勝つ確率を P としたとき、初期資金 A 円をどのような戦略で賭ていったらいいのか？という問題を定式化し分析を行なった。期待値最大化では、常に所持金全額を賭け続けるのが最適となる。しかしながら、この戦略は賭けの期待値が正であったとしても、長期的には確率 1 で破産してしまう。ここで

$$f = \frac{(R + 1)P - 1}{R}$$

とするとき、常に所持金の f 倍をかけ続ける戦略が、最頻値最大化戦略となることを示し、この f は

$optimal f$ と呼ばれている．この $optimal f$ に従い賭け続ける戦略も対数の期待値最大化戦略となる．

これらの先行研究は，対象となるもの（個体，金銭）が増減を行いつつ増えていく状況で，どのような評価や戦略が望ましいのかということを議論をしている．本論文では，これらと共通の土台である複製状況において，上記とは一見異なる形の評価指標を提案する．その上で，本論文で提案した評価指標と上の評価指標に共通して出てくる支配タイプや最頻値最大化戦略（対数の期待値最大化戦略）との関係を議論していく．

3.2 複製モデル

本章で議論の土台となる複製状況とは，生物や金銭などのある対象に関して対象の変化が確率過程によって規定されているモデルである．特に本論文では増減ルールが，期によらず一定の場合について限定した議論を行う．本論文で提案する評価指標は，増減ルールが期に依存して変化する場合も適用可能ではあるが，増減ルールが変化する場合は，実際の計算がかなり煩雑になるために，一般的な議論が困難となるための制約である．例えば生態学においては種の爆発や絶滅などの問題は適用範囲として考えられる．また物理学においては，核分裂などで臨界状態に達するまでかかる時間なども表現することが可能である．

ここで対象がそれぞれ個体を区別可能で整数の離散値しかとりえないような場合の状況と，全体のボリュームのみに絞ることにより連続値を取ることができる場合ではかなり数理的に異なる性質を持つことが知られている．例えば生態学のモデルにおいては 1.5 個体というのは意味を持たない．一方実数値全体を対象として取ることを仮定すると個体ごとに発生するタイプの不確実性の記述は出来ない．本論文では複数のタイプが混在する状況でどのような複製ルールに従うタイプが最も優位に立つ可能性が高いのかということに主たる興味がある．さらに個体ごとのリスクよりもむしろ環境リスクに関するマネジメントに興味があるので，連続環境における複製モデルについて分析を行っていく．

ここで $z_{t(i)}$ を type i の時刻 t における量とする．また z_t を時刻 t における全タイプの合計量とする．このとき 複製状況は以下で定義される

定義 1

$$D = (I, \Xi, \{\pi^s\}_{s \in S}, P^i)$$

ここで I はタイプの集合， $\Xi = \{\xi^s\}_{s \in S}$ は環境の集合， S は有限の添え字集合とする．なお ξ^s が生起する確率を π^s として，ここではこの確率はこれまでの履歴からは独立であるものとする． $P^i : \Xi \rightarrow \mathbf{R}$ は i についての倍率を定める関数で任意の t に対して次期の個体数は $z_{t+1(i)} = P^i(\xi^s)z_{t(i)}$ で与えられるものと仮定する．

また $(\forall s)(\pi^s > 0), (\sum_{s \in S} \pi^s = 1)$ で $z_{0(i)} = 1$. とする. 最初の条件は確率の条件であり, 2 番目のものはあくまで計算の便宜上のものであるが実際の適用対象の最初の個体数が仮に 1 でなくとも, 扱う対象が実数値を取りうるので, 実数倍の変換を行うことにより初期値を 1 として計算した上で後で引き戻しても同一の結果となる. また $z_{t(i)}$ は必ずしも整数とは限らず実数値を取ることもある.

3.3 タイプの支配

このような複製状況の評価するための基準はいくつか考えられるが, 最も有名なものは期待値である. しかしながら, 期待値での評価は目的によっては奇妙な結果を生むこともありうることを次の例で紹介する.

例 1 環境は 1 と 2 の二つの状態があると仮定する. これらは各期ごとに独立してそれぞれ確率 $1/2$ とに発生するものとする. また対象としては二つのタイプが存在してタイプ 1 は環境に敏感で, 環境 1 が発生したならば 1.8 倍に, 環境 2 が発生したならば 0.5 倍になると仮定する. 一方タイプ 2 はどちらの環境の下でも 1 のまま一定と仮定する. このときに t 期後のタイプ 1 の量の期待値は $(1.15)^t$ となるため, 期待値での評価では無限期経過後には無限大に期待値は発散する. 一方 $n^1(t)$ と $n^2(t)$ を t 期までに環境 1 と 2 がそれぞれ観察された数とする. このとき大数の法則から

$$\lim_{t \rightarrow \infty} n^1(t)/t = 1/2$$

かつ

$$\lim_{t \rightarrow \infty} n^2(t)/t = 1/2.$$

となることから

$$P\left(\lim_{t \rightarrow \infty} (1/t) \ln(z_{t(i)}) = \frac{1}{2}(\ln 1.8 + \ln 0.5)\right) = 1.$$

ここで $(\ln 1.8 + \ln 0.5) = \ln 0.9$, なのでタイプ 1 が無限期の後に量が 0 に収束する確率は 1 となる. したがってタイプ 2 のほうが期待値では劣るものの, 全体に対する 2 の割合が 1 に収束する確率も 1 となっていることがわかる.

このように期待値最大化は必ずしも全体に対して支配的であるとはいえないことがわかる. ここで支配という概念は数学的には以下のように定義される

定義 2 タイプ i が支配的であるとは以下が成り立つことを言う.

$$P\left(\lim_{t \rightarrow \infty} (z_{t(i)}/z_t) = 1\right) = 1.$$

それではどのようなタイプが長期的な複製の後には支配的となるのだろうか．これを議論するためにいくつかの概念を導入する．

定義 3 α_i は以下を満たすときタイプ i の長期の平均成長率と呼ぶ

$$P\left(\lim_{t \rightarrow \infty} (1/t) \ln(z_{t(i)}) = \ln(\alpha_i)\right) = 1$$

補題 1 $\rho^i = \sum_{s=1}^S \pi^s \ln(P^i(\xi^s))$. とする．このとき ρ^i が長期平均成長率と一致する．

(Proof) ここで ξ_t を t 期目に生じた環境とする．また $n^s(t)$ を t 期目までに環境 ξ^s が起こった回数とする．このとき

$$\frac{1}{t} \sum_{t=0}^{t-1} \ln(P^i(\xi_t)) = \sum_{s=1}^S (n^s(t)/t) \ln(P^i(\xi^s)),$$

なので

$$P\left(\lim_{t \rightarrow \infty} (n^s(t)/t) = \pi^s\right) = 1$$

従って大数の法則より

$$P\left(\lim_{t \rightarrow \infty} \frac{1}{t} \sum_{t=0}^{t-1} \ln(P^i(\xi_t)) = \rho^i\right) = 1,$$

$$W_t = z_{t(i)} / \prod_{t=0}^{t-1} P^i(\xi_t)$$

よって $\lim_{t \rightarrow \infty} W_t$ となる確率が 1, となるので

$$P\left(\lim_{t \rightarrow \infty} \frac{1}{t} \ln(z_{t(i)}) = \rho^i\right) = 1.$$

従って定義より $\ln(\alpha_i) = \rho^i$.

Q.E.D.

この補題より他の任意のタイプ j に対してタイプ i が以下を満たすならば支配的である $\rho^i > \rho^j$.

この結果はタイプの支配というのは期待値ではなく対数を取ったものの期待値の最大化で決定されていることを示している．期待値はリスクニュートラルな意思決定基準であることはよく知られているが、このことからタイプの支配というのはリスク回避的な性質を持つことがわかる．

以上の議論は無限期の複製に注目した議論であったがこれを有限期にした場合はどうなるのであろうか．次に有限期の複製に対して適用できるリスク回避的な評価基準である期待到達時間という概念を提案し、これが目標無限の際にはタイプの支配と一致する概念であるということを次に紹介する．

3.4 期待到達時間

本論文で新たに提案する、「期待到達時間 (ERT)」はある目標値まで $z_{t(i)}$ が到達するまでどのぐらいの時間がかかるのかということに関する評価基準である。ERT は今回扱う連続環境のモデルだけではなく個体ごとの複製の状況に対しても適用できる広い概念である。しかしながら個体ごとの複製を考える状況では数理的な扱いが非常に困難になるため本論文では対象を連続環境に絞った上で分析を行う。

定義 4 あたえられた目標値 $K > 0$, に対して $Et(i, K)$ を $z_{t(i)}$ が K に到達するまでかかる時間の期待値としてこれを期待到達時間 (*Expected Reaching Time*, 以下 ERT と省略) と呼ぶ。

ここで $p^t(i, K)$ をタイプ i の量が時刻 t においてはじめて K 以上となる確率とする。このとき仮に

$$\lim_{T \rightarrow \infty} \sum_{t=1}^T p^t(i, K) = 1$$

と

$$\sum_{t=1}^{\infty} t p^t(i, K) < \infty$$

が成り立つならば $Et(i, K)$ は

$$Et(i, K) = \sum_{t=1}^{\infty} t p^t(i, K).$$

として有限の値として well-defined となる。それ以外の場合には $Et(i, K)$ は定義できないとする。

ここで ERT の数理的性質を考える。 $z_{t(i)}$ それ自身よりもその対数を取ったものに注目すると

$$\ln(z_{t(i)}) = \sum_{t=0}^{t-1} \ln(P^i(\xi_t)).$$

となっていることがわかる。ここで各期の変動は i.i.d に従うことを仮定しているため、この過程はランダムウォークとみなすことができる。

いま $K > 1$, とした場合にタイプ i の ERT は $\rho^i > 0$ のときに限り well-defined となる。また $K < 1$, ならば $\rho^i < 0$ のときに限り well-defined となる。前者の状況においてはなるべく速く到達したほうが望ましいと考えられるので、小さい ERT のほうが望ましいということになる。対称的に後者の場合においてはなるべく長い間減少をとどめたいような状況と考えられるので大きい ERT のほうが望ましいと考えられる。

なおこの条件の直感的な意味を図で表したのが図 3.1 である。ここで過程の対数を取っている。ここで目標が $K > 1$ ならば $\ln K$ は 0 より大きくなる。このとき ρ^i はこのランダムウォークの期待値と一致する。

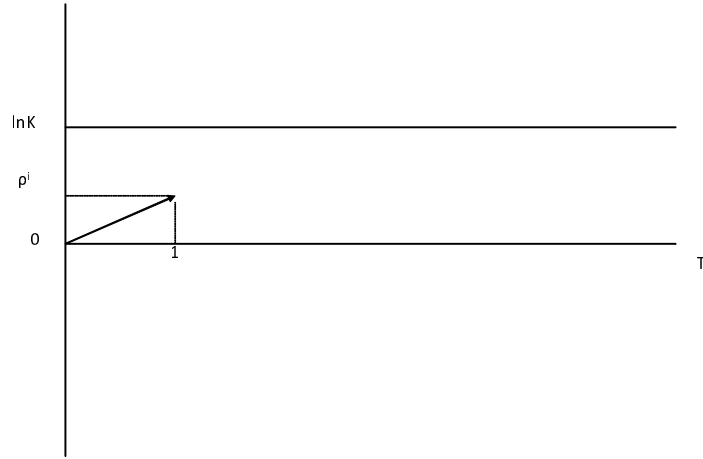


図 3.1: ERT 定義条件の直感

この期待値が正のときに限り, ERT は well-defined となる. また逆の場合も同様に目標が $K < 1$ ならば, ランダムウォークの期待値 $\rho^i < 0$ のときに限り well-defined となる.

ここで T を目標 K を始めて超える時刻とすると T は停止時間となる. このとき $X(T) = E \left(\sum_{t=1}^T \ln(P^i(\xi_t)) \right)$ とすると, 以下の関係が成り立つことが知られている. (Wald, 1944)

$$E(T) = \frac{X(T)}{\rho^i}$$

(証明)

$$Y_n = \sum_{t=1}^n \ln(P^i(\xi_t)) - n\rho^i$$

とする. このとき Y_n はマルチンゲールとなり, また optional stopping theorem を満たすことから,

$$E \left(\sum_{t=1}^T \ln(P^i(\xi_t)) - T\rho^i \right) = E(Y_T) = E(Y_0) = 0$$

従って

$$E \left(\sum_{t=1}^T \ln(P^i(\xi_t)) \right) = \rho^i E(T)$$

よって

$$E(T) = \frac{X(T)}{\rho^i}$$

Q.E.D.

なお一般には特殊な場合を除いては $E(T)$ の値を直接計算することができないことが知られているものの (Feller, 1966). 不等式による評価を行うことは可能である.

また T の期待値である $E(T)$ は, 定義より本モデルの $Et(i, K)$ と等しくなるので

$$Et(i, K) = \frac{X(T)}{\rho^i}$$

という関係が期待到達時間と ρ^i の間に成り立つことがわかる.

このときタイプ i の変動の係数 $P^i(\xi^s)$ の中で最大となるものを $P_{max}^i(\xi^s)$ とする. このとき次が成り立つ

補題 2

$$\frac{\ln K}{\rho^i} \leq Et(i, K) < \frac{\ln K + \ln P_{max}^i(\xi^s)}{\rho^i}$$

これは K が十分大きい場合には ERT は概ね目標値の対数である $\ln K$ に比例して, さらに ρ^i に反比例するということを意味している.

3.4.1 支配と ERT の関係

支配は無限期の後の状態に注目した概念である. 一方 ERT は有限期に注目して定義された概念である. ここで目標値 K が無限大に近づく際に, 両者の関係を明らかにする.

定理 1 $\rho^i > 0$, で タイプ i が支配的であるとする. このときに限り以下が成り立つ.

$$\lim_{K \rightarrow \infty} \frac{Et(i, K)}{Et(j, K)} < 1 \text{ なお } j \neq i.$$

(証明) ($\Rightarrow$)

補題 1 よりタイプ i が支配的であるならば, 任意の他のタイプ j に対して $\rho^i > \rho^j$ となる. 補題 2 より

$$\frac{Et(i, K)}{Et(j, K)} < \frac{(\ln K + \ln P_{max}^i(\xi^s))/\rho^i}{\ln K/\rho^j}.$$

となるので $K \rightarrow \infty$, ならば右辺は $\rho^j/\rho^i < 1$ に収束する.

($\Leftarrow$)

補題 2 は $K \rightarrow \infty$ の際に

$$\frac{\ln K/\rho^i}{(\ln K + \ln P_{max}^j(\xi^s))/\rho^j} < Et(i, K)/Et(j, K).$$

となる．右辺は 1 より小さいので

$$\lim_{K \rightarrow \infty} \frac{\ln K / \rho(i)}{(\ln K + \ln P_{max}^j(\xi^s)) / \rho(j)} = \rho(j) / \rho(i) < 1.$$

となることから $\rho^i > \rho^j$, これはタイプ i が支配的であることを意味する．

Q.E.D.

定理 1 により支配と ERT が目標無限大という条件の下では厳密に一致するということが示された．実際に二つの基準は十分時間がたったあとの対象の量を評価しているのに対し, ERT はある目標に対してどれだけの期間で到達することが出来るのかという意味で対称的な評価ということが出来るがこれらが一致するかどうかは直感だけでは断定が困難であるがこの証明により一致することが示された．

3.4.2 有限量に対する ERT

ERT のもともとの問題意識は有限の目標量 K に対する時間というものであった．そこで今度は有限の値に対しての振る舞いを検討してみる．また投資におけるポートフォリオを ERT を用いて評価することにより ERT が期待値と分散を統一的に扱う一つの指標であることを示す．まずは以下の例を考える．

例 2 いま A, B という二つの商品があると仮定する．それぞれの変動は独立であって 2% 上昇する確率が 60%, 2% 減少する確率が 40% という同じ確率分布にしたがっているものとする．またこの確率は期によらず一定であると仮定する．このとき 1 と 2 という 2 つの投資戦略を考える．戦略 1 は単純に毎期ごとに収益を全て A につぎ込むという単純な投資戦略である．一方戦略 2 は毎期ごとに得られた収益の半分ずつをそれぞれ A と B に配分して投資するポートフォリオ戦略である．この戦略は 2% 上昇する確率が 36%, 変動しない確率が 48%, 2% 減少する確率が 16% であるという商品に投資していることと等しくなる．(図 3.2 参照．)

まず最初に K が変動に対して十分大きい場合を考える．このとき補題 2 が適用できるので 2 仮に $K = 3$ とするならば,

$$\rho^1 = 0.6 \ln 1.02 + 0.4 \ln 0.98 \approx 0.003800, \quad P_{max}^1(\xi^s) = \ln 1.02 \approx 0.01980,$$

$\ln 3 \approx 1.0986$, ここで補題 2 より

$$1.0986 / 0.003800 \leq Et(1, 3) < (1.0986 + 0.01980) / 0.003800,$$

となるので

$$289.1 \leq Et(1, 3) < 294.3.$$

ポートフォリオ問題への適用

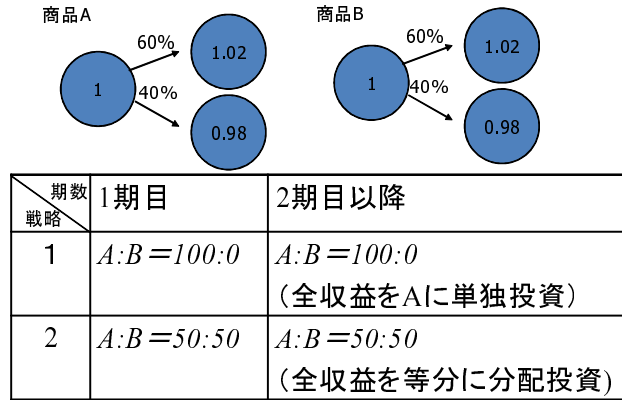


図 3.2: 2つの戦略

$$282.0 \leq Et(2, 3) < 287.1.$$

となることがわかる。このことからわかるようにポートフォリオ戦略による分散の減少効果は期待値による評価には反映されていないものの ERT ではそれも総合的に組み込んだ評価となっていることがわかる。

いまは変動に対してかなり大きい目標を設定したため、補題 2 の不等式による評価で戦略の優劣を議論することができたが目標が小さい場合にはこれは難しい。そこでもっと小さい目標の場合に関してはモンテカルロシミュレーションを行う。

ここで ERT のモンテカルロシミュレーションによる近似値を ERT' と表記する。以下は 100 万回の平均値である。図 3.3 と図 3.4 はそれぞれ例 3 の 1 と 2 に関する結果である。これを見るとわかるように、 K はそこまで大きい値ではなくても K の対数に ERT はほぼ比例しているとみなしてよいと思われる。

また図 3.5 は二つのタイプの比を求めたものである。ここで $\ln K$ が極めて 0 に近い範囲では比はほぼ 1 ということになるため ρ^i による評価の結果とは異なってくることがわかる。しかしながらある程度大きいところでは、解析的な極限值である

$$\lim_{K \rightarrow \infty} Et(2, K)/Et(1, K) = \rho(1)/\rho(2) \approx 0.9754.$$

に比較的早い段階で収束していることがわかる。

ERT の近似値が ρ^i にほぼ反比例するということからわかるように、ERT は厳密な値を解析的に求めることは簡単ではないものの、モンテカルロシミュレーションは比較的簡単に実行できるため、シミュレー

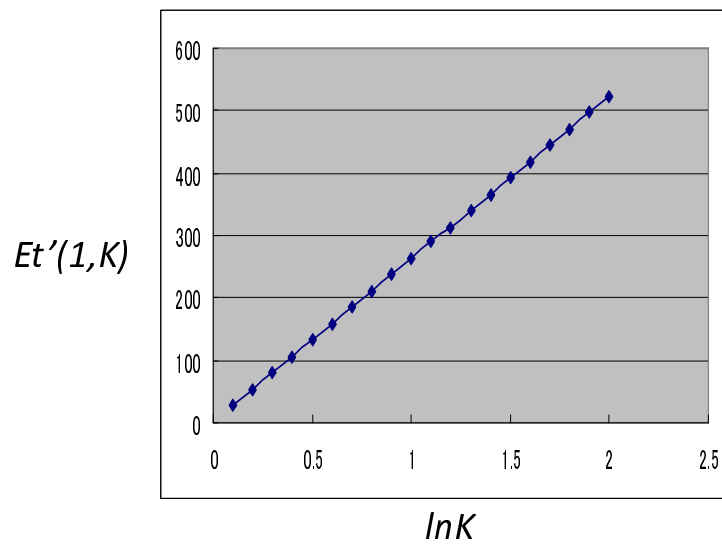


図 3.3: タイプ 1 の ERT の近似値

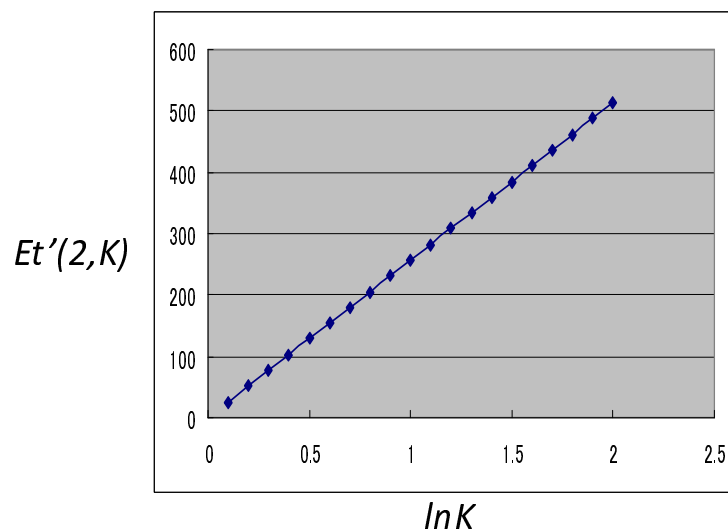


図 3.4: タイプ 2 の ERT の近似値

ションにより値を比較的容易に求められる上近似値ならば期待値と大差ない計算量により求められることから、単に直感的なだけでなく、期待値と分散を総合的に評価した十分に扱いやすい評価基準となっていることがわかる。

なお、金融などにおいては割引率が問題になる場合もある。割引率が時間によらず一定の場合は、変動が割引を換算した後の数字と解釈すると、期待到達時間同一は置き換えた分布の結果と同一になる。

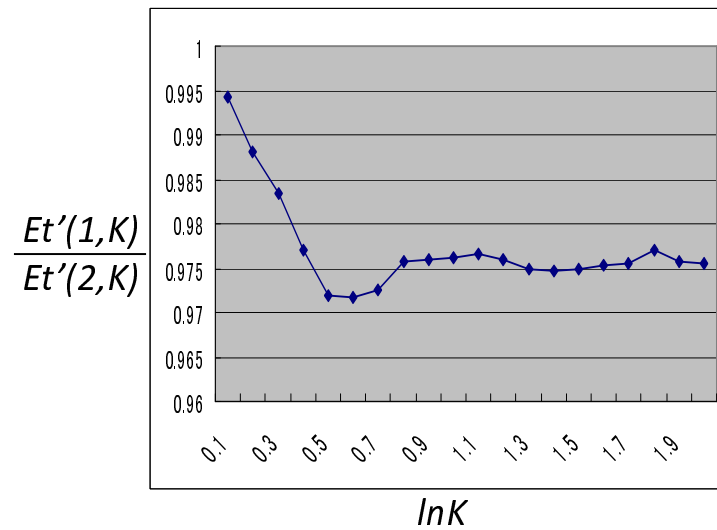


図 3.5: 両タイプの ERT の近似値の比

例えば, 1.1 倍の確率が 70%, 0.9 倍の確率が 30%, 割引率が 2% の場合は, この商品は割引率を換算すると, $(1.1 \div 1.02)$ 倍になる確率が 70%, $(0.9 \div 1.02)$ 倍になる確率が 30% の商品に投資している状況と同一になる.

3.5 3 章のまとめ

本章では不確実性下の意思決定状況の中で特に複製状況に注目し, その上で期待到達時間というリスク回避的な性質を持つ評価指標を提案した. この期待到達時間は, 目標が有限の場合にも定義できる指標であり, 目標が無限大の極限の場合には, 進化的に最も安定したリスクに対する態度であるタイプの支配の基準と一致することを数学的に示した. また期待到達時間が, 実際にリスク低下効果を数値的に望ましい指標として評価することを確認するために, 投資状況に適用を行った. その上で単独投資戦略と分散投資戦略に関して期待到達時間を解析的及びシミュレーションを行うことにより計算した. その結果として近似値ならば容易な計算量で求められることを確認した.

第4章 不完全な推論に関する新たな数理モデルの提案

4.1 問題状況の複雑性と限定合理性

この章では複雑な意思決定において実際の人間はどのような意思決定を行っているのかということに関して、数理的なモデルを提案する。ここで問題が複雑であるといった場合には、問題を数学的なモデルに落とし込む際の認識レベルの複雑性や、問題が定式化された後にどの代替案を選ぶべきなのかという推論過程に関する複雑性や、楽器の演奏などに代表されるように何をすればいいのかはわかったとしても実際に行動する際の行動レベルの複雑性などいくつかの側面が存在する。本章では特に推論過程に関するモデル化について議論を進めていく。

2つの数字のうち大きいほうを選ぶといったような、問題状況が単純な場合は推論能力に関する限定合理性が意識されることは少ない。しかしながら状況が複雑になると、正しいと思っていたことが間違っていたという読み間違いの発生頻度が高くなる。この読み間違いがどのような要因で起こりやすいのかということに関して注目することにより、意思決定のモデルをメタ的な視点からの分析を行う。さらに、原理的に問題が解けるということと、実際の意思決定者がその解を実際に求められるということは大きな隔たりのあり、限定合理性を考慮して、実際には自分のミスの可能性を考慮したプレイが実際には行われていることも議論する。

4.2 限定合理的な意思決定のモデル

4.2.1 エラーに関する数理的モデルの先行研究

一例として確率 0.7 で利得 10、確率 0.3 で利得 -10 と予想したときに、実際に真のゲーム状況がどうなっているのかということに関して、利得分布と真の決定状況に関する関係付けを行わないと分析ができない。そこで本節では数理的なモデルとして分析を行うためにその極端な状況で確率 1 で自分の推論を信じている状況のモデル化を行う。

そのためにまずはエラーの特性についての数理的な仮定に関する議論のために、有限量のエラーを仮定している Quantal Response Model を紹介する。このモデルにおける均衡概念である QRE は各プレイヤーが有限量のエラーは避けられない条件の下で、自分以外のプレイヤーの戦略に関する信念に対して整合的であることを要請している。QRE では表現できる範囲が広すぎるため実際に何が起こるのか予想することが困難であるため以下ではそれに条件を加えた正規 Quantal 均衡の定義を紹介する。

そのためにまずは以下の記号を準備する。

$G = (N, \{S_1, \dots, S_n\}, \pi_1, \dots, \pi_n)$ を戦略型ゲームとする。ここで $N = \{1, \dots, N\}$ はプレイヤーの集合、 $S_i = \{s_{i1}, \dots, s_{iJ(i)}\}$ はプレイヤー i の戦略集合、 $S = S_1 \times \dots \times S_N$ を戦略プロファイルの集合、 $\pi_i : S_i \rightarrow \mathbf{R}$ をプレイヤー i の利得関数とする。さらに $\Sigma_i = \Delta^{J(i)}$ を S_i 上の確率分布とする。ここで $\sigma_i \in \Sigma_i$ を S_i から Σ_i への写像である混合戦略とする。また $\sigma_i(s_i)$ をプレイヤー i が純戦略 s_i を選択する確率とする。 $\Sigma = \Sigma_1 \times \dots \times \Sigma_N$ を混合戦略のプロファイルとする。与えられた混合戦略のプロファイル $\sigma \in \Sigma$ に対して、プレイヤー i の期待利得は σ によって導かれる純戦略上の確率分布を $p(s) = \prod_{i \in N} \sigma_i(s_i)$ とするとき、 $\pi_i(\sigma) = \sum_{s \in S} p(s) \pi_i(s)$ によって与えられる。

ここで P_{ij} をプレイヤー i が戦略 j を選ぶ確率とする。このとき

定義 5 $P_i : R^{J(i)} \rightarrow \Delta^{J(i)}$ が以下の 4 条件を満たすとき正規 *quantal response* 関数という。

1. 任意の $j = 1, \dots, J(i)$ と $\pi_i \in R^{J(i)}$ に対して $P_{ij}(\pi_i) > 0$. (*Interiority*)
2. 任意の $\pi_i \in R^{J(i)}$ に対して $P_{ij}(\pi_i)$ が連続で微分可能. (*Continuity*)
3. 任意の $j = 1, \dots, J(i)$ と $\pi_i \in R^{J(i)}$ に対して $\partial P_{ij}(\pi_i) / \partial \pi_{ij} > 0$. (*Responsiveness*)
4. 任意の $j, k = 1, \dots, J(i)$ に対して、 $\pi_{ij} > \pi_{ik}$ ならば $P_{ij}(\pi_i) > P_{ik}(\pi_i)$. (*Monotonicity*)

Interiority はモデルが完全な定義域を持つことを要請している。Continuity は数学的な理由であるが、 P_i が非空かつ一点に定まるための要請である。また微小な期待利得の変化が劇的な選択確率のジャンプを生まないようにというように、現実的にも妥当な仮定であると考えられる。Responsiveness は、他の条件が一緒である限り期待利得の上昇した場合には選択確率も上昇する必要があることを要請している。Monotonicity は合理性の弱い形の表現で、一対比較において高い期待利得となる行動は、低い期待利得を与える行動よりも高確率で選択されなければならないことを要請している。

$P(\pi) = (P_1(\pi_1), \dots, P_n(\pi_n))$ が正規であるとは、それぞれの P_i が上の 4 つの条件を満たすことを言うものとする。ここで $P(\pi) \in \Sigma$ かつ $\pi = \pi(\sigma)$ が任意の $\sigma \in \Sigma$ に対して定義されるので、 $P \circ \sigma$ が Σ から自身への写像として定義される。

定義 6 P が正規であると仮定する．戦略形ゲーム G において混合戦略プロファイル σ^* が, $\sigma^* = P(\sigma)$ を満たすときに, 正規 *Quantal response* 均衡と呼ぶ．

P の正規性は連続性を含むことから, $P \circ \sigma$ は連続写像となる．従って Brouwer の不動点定理より直接正規 *Quantal response* 均衡の存在が示される．

定理 2 戦略形ゲーム G の任意の正規の P に対して, *Quantal response* 均衡が存在する．

このように高い利得を与える戦略ほど高確率で採用され, さらに一点で定まるという定式化の仮定の下でも均衡が存在するという結果がこのことからわかる．これらの仮定は次節の推論のエラーのモデル化においても採用している．

4.3 推論不完全性を考慮した展開形ゲームにおける意思決定のモデル化

QRE モデルはとりわけ実験ゲーム理論においてさまざま方面からの研究が進められている．しかしながらエラーの種類に関してはあまり注意を払われてきていなかった．人間の意思決定過程は,(1) 認識 (2) 推論 (3) 行動という 3 つのフェーズから構成されていると考えられるが, これまでのゲーム理論ではこれらのフェーズの役割の違いとそこで起こるエラーの性質の違いに注意を払われてくることが少なかった．多くの研究はそれらをすべて一くくりにするか, あるいは震える手完全均衡に代表されるように (3) の選択フェーズにおけるエラーのみに注目したものが殆どだった．これはとりわけ戦略形ゲームの意思決定状況では, 推論過程において何をしているのかモデル化することが困難であるということが大きな理由の一つであると考えられる．ここでは推論フェーズの役割がよりクリアに分離できる展開形ゲーム状況においてエラーの役割の分析を紹介するために提案したモデルを紹介する．ここでは意思決定主体は, 通常の完全情報の展開形ゲーム状況において, 不完全な精度で推論しながら意思決定をしている状況のモデルとなっている．ここで考える意思決定状況は, 分析者にとっては以下のような完全情報の有限展開形ゲーム状況である．

定義 4 真の客観的完全情報展開ゲームは以下で与えられる

$$G = (I, N, A, \alpha, (N_i)_I, (r_i))$$

ここで I はプレイヤーの集合, N はノードの分割, N_T と N_D はそれぞれ N の分割で N_T は終点ノード全体の集合で N_D は決定ノード全体の集合である． A は行動全体の集合で, $\alpha: N - \{n_1\} \rightarrow N_D$ は始点ノード n_1 以外のノードから前のノードを定める関数である． $P: N_D \rightarrow I$ はそのノードで誰が決定するのか定める関数で $(r_i): N_T \rightarrow \mathbf{R}^I$ はそれぞれのプレイヤーの利得関数である．

ここで G は完全情報の展開形ゲームなので、後ろ向き帰納法により部分ゲーム完全均衡が得られる。しかしながら、現実的にはすべてのノードにおいて正確な比較をエラーなしに行うことは多くの場合において現実的ではない上、実際に将棋やチェス等の完全情報ゲームでも、最終盤の特殊な状況以外ではそもそもそういう計算を人間のトッププレイヤーでさえ行っていないことが知られている。

そこで本論文では、エラーを含む推論のヒューリスティクスを以下のように提案する。そのために必要な

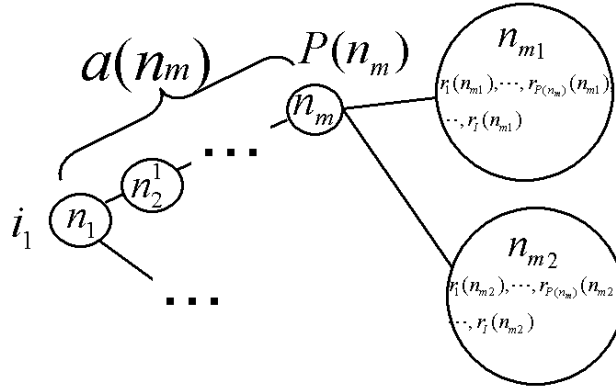


図 4.1: 展開形ゲームの図

概念をいくつか紹介する。

N_2 : n_1 から到達可能なノードの集合。具体的には

$$N_2 = \{n | n \in N, \alpha(n) = n_1\}.$$

N_m^1 : G における最後の意思決定ノードの集合。具体的には

$$N_m^1 = \{n | n \in N_D, \exists n_t \in N_T, \text{ s.t. } \alpha(n_t) = n, \text{ かつ } \neg(\exists n_d \in N_D, \text{ s.t. } \alpha(n_d) = n)\}.$$

n^* : プレイヤー $P(n)$ にとってノード $n \in N_D$ において最適なノードである。

$$\text{具体的には } n^* \in \operatorname{argmax}_{n'} \{r_{P(n)}(n') | \alpha(n') = n\}$$

ここで $r(n_d) = (r_1(n_d), \dots, r_j(n_d), \dots, r_I(n_d))$ を決定ノード $n_d \in N_d - \{n_1\}$ においてそのあと最善であろう意思決定の列が行われる場合に各プレイヤーが得られる利得ベクトルを並べたものに関する推論結果とする。

このヒューリスティクスを「粗い推論のモデル」と呼ぶ。

1. まず i_1 を始点ノード n_1 で意思決定を行うプレイヤーと仮定する. i_1 は $n_2 \in N_2$ の各ノードに関して後ろ向き帰納法で仮想利得ベクトルを求めようとする.
2. 実際に i_1 がこのベクトルを計算するために, 終点ノードからその一つ前の意思決定ノード $n_m \in N_m^1$ への推論による利得の引き戻しから開始する. ここで $a(n_m)$ を始点ノードから n_m までの深さとする. また $b(n'_m)$ をこの意思決定における最適な行動 $r_{P(n_m)}(n_m^*)$ から ノード (n'_m) を選択した場合のプレイヤー $P(n_m)$ にとっての損失とする. すなわち

$$b(n'_m) = r_{P(n_m)}(n_m^*) - r_{P(n_m)}(n'_m)$$

ここで $n'_m \in \{n | \alpha(n) = n_m\}$.

3. i_1 は正しく推論できた場合には $r(n_m^*)$ を $r(n_m)$ における仮想利得ベクトルとして引き戻すが, ある確率でエラーが起こるものとする. ここではエラーの確率は $a(n_m)$ の増加関数かつ $b(n_m)$ の減少関数と仮定する. これは遠い未来の意思決定に関する予想ほど精度が落ち, また結果の利得の差が大きい判断ほど間違えづらいという直感を表現したものになっている. またもし最適応答が複数ある場合はそれらを等確率で引き戻すと仮定する.
4. 以上の引き戻しが全ての $n_m \in N_m^1$ に関して終了したならば, i_1 は $n_m \in N_m^1$ を得られた仮想利得を持つ終点ノードと新たにみなすことにより縮約されたゲームを新たに作る. ここで縮約されたゲームにおける最後の意思決定ノードの集合 N_m^2 が同様に定義することができる. これらのプロセスは n_2 にいたるまで遡っての引き戻しが行われる. 以上の操作により i_1 はノード n_2 に対するエラーを含んだ上での推論結果による評価が得られる.
5. 最後に i_1 は各 $n_2 \in N_2$ の評価ベクトルを比較することにより最適な行動を決定する. (注意としてこのプロセスは推論過程におけるエラーを表現している一方最後の決定フェーズにおけるエラーは無い物と仮定している.)
6. ここで i_2 を i_1 の後に決定を行うプレイヤーとする. ここで i_2 も i_1 と完全に独立に上記の推論プロセスを行うことにより意思決定を行う.
7. 以上のプロセスによりにして得られた行動を各プレイヤーが取ることによりある決定ノードがゲームの結果として定まる.

このプロセスは各プレイヤーが推論能力の範囲で可能な限り合理的に振舞おうとしている状況のモデルである. またこのプロセスによる結果は一般的には非決定的で, 終点ノード N_T 上の確率分布が得られる.

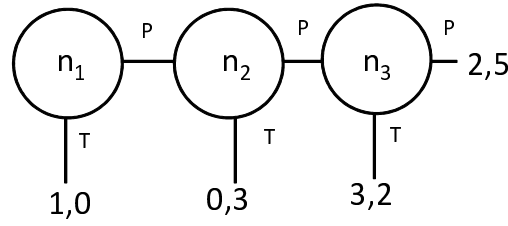


図 4.2: 基本ゲーム

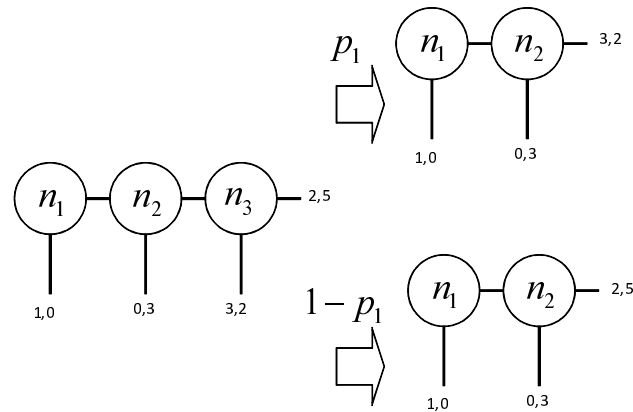


図 4.3: 基本ゲームの1段階目の推論結果

図 4.2 – 4.5 はこのプロセスにおける推論結果を図で表したものである．いま客観的な基本ゲームが図 4.2 のものであったとする．このとき 1 段階目において正しい推論をする確率が p_1 だった場合には，推論結果は図 4.3 のようになる．ここで 1 段階目において正しい推論を行った場合の，2 段階目において正しい推論を行う確率を p_2 とするときに，この推論結果は図 4.4 の上の場合に相当する．同様に 1 段階目において誤った推論を行った場合の，2 段階目において正しい推論を行う確率を p_3 とするときに，この推論結果は図 4.4 の下の場合に相当する．結果的にこのゲームの推論による結果は図 4.5 で与えられる．

ここでプレイヤー間の推論能力が実際には異なっている場合も考えられるが，本モデルでは自分の能力で相手の意思決定に関する推論を行うと仮定している．なお一人のプレイヤーがゲーム中に二回以上行動する場合も一般には考えられるが，このモデルでは前のノードにおける推論結果と後のノードでの推論結果は独立であると仮定している．これはあるシナリオが最善だと意思決定の初期には判断したものの，時間がたち状況の細部が見えてくるにつれて初期のシナリオが不正確であったことが判明した場合には意思決定を修正しながら行うといった状況も描写できるようなモデルとなっている．

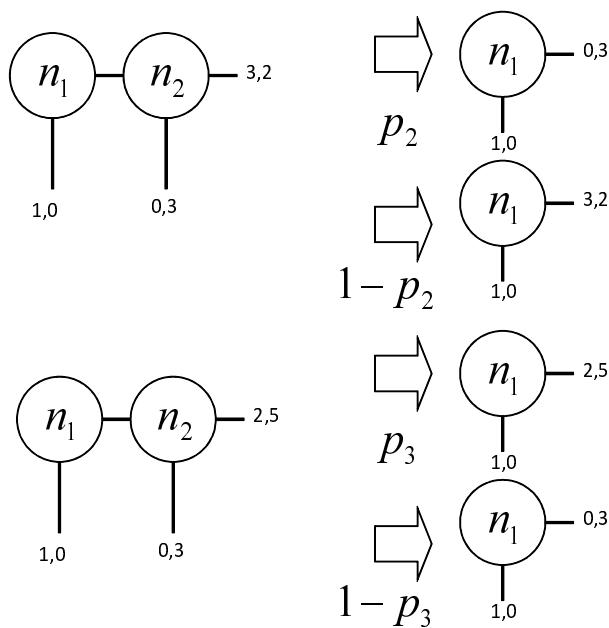


図 4.4: 基本ゲームの 2 段階目の推論結果

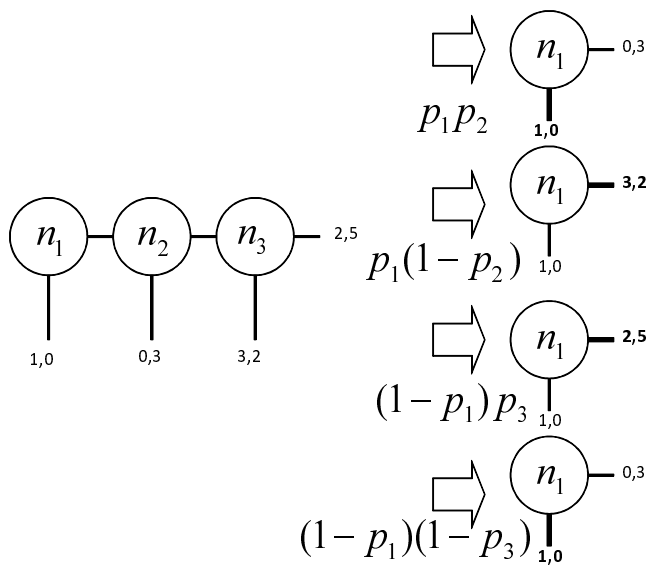


図 4.5: 基本ゲームの最終的な推論結果

4.4 二つの具体的な粗い推論モデル

前節で提案した一般的な粗い推論モデルは関数形を自由度の高い形の仮定にとどめているため、かなり多くの状況を説明可能である。その一方具体的にどのような結果が起こるのか予想しようとする場合には

一般性故に候補が広すぎるという問題がある．そこで推論のエラー率が利得差の増加関数、深さの減少関数という仮定を満たすものの中で代表的な二つの関数形のモデルに注目してみる．なおエラー率はまず利得差に対しては線形よりも急激な影響を持つと考えられるので、ここでは利得差に関しては震える手完全均衡や QRE の場合と同様に指数関数的に効くものを採用した．また深さに関しても、5 手先の状況に比べると 10 手先の状況に関する比較の精度は半分よりももっと急激に効いてくると考えられるため、差に対しても指数関数的な分布を仮定した．

4.4.1 logit 関数に基づく粗い推論モデル

まずプレイヤー i がノード n_k から決定ノード n_l に関する推論を定式化する．そのために次のような記号の準備をする：

j : ノード n_l において決定を行うプレイヤー．

n_{sx} : n_l から到達可能なノード．

N_s : n_l から到達可能なノードの集合． i.e. $N_s = \{n | n \in N, \alpha(n) = n_l\}$ ．

σ : 推論能力に関するパラメーター．

ここで注意として σ は単に推論能力に関するパラメーターを表すだけではなく利得の単位に関する調整パラメーターの役割も同時に果たしている．例えば表示単位がドルからセントに変わったならば、同じ推論能力を表す σ も 1/100 になる．従って、もし他の条件が一緒であるならば推論能力が上昇するならば、それに相当する σ も上昇する．

定義 7 logit 関数に基づく粗い推論モデルでは、 $r(n_l)$ に $r(n_{s1})$ を割り当てる確率を以下で定義する．

$$\frac{\exp\left(\frac{r_j(n_{s1})}{a}\sigma\right)}{\sum_{n_s \in N_s} \exp\left(\frac{r_j(n_s)}{a}\sigma\right)}$$

これはほぼ利得差の指数関数比に選択確率が比例して、深さはその比率を緩やかにする方向の影響を与える設定となっている．また σ は大きいほうが結果の差が際立つことから合理性の高いことを意味している．

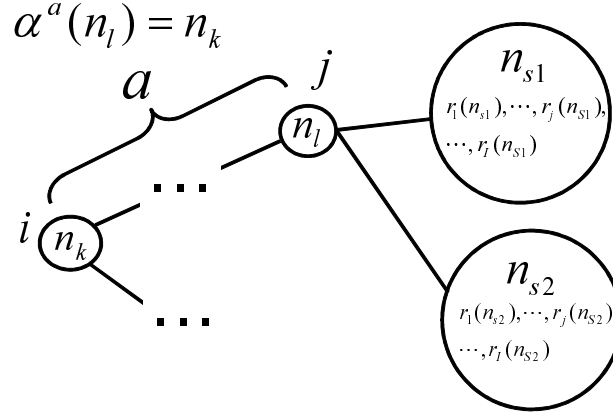


図 4.6: 2つの具体的なモデルのための記号

4.4.2 指数エラーに基づく粗い推論モデル

同様にプレイヤー i がノード n_k から決定ノード n_l に関する推論を定式化する．そのために以下の記号を定義しておく．

N_{s*} : n_{s*} の集合． i.e. $N_{s*} = \operatorname{argmax}_{n'} \{r_{P(n_s)}(n') | \alpha(n') = n_s\}$

$b(n_s)$: $b(n_s) = r_j(n_{s*}) - r_j(n_s)$ $n_s \in N_s$ and $n_s \notin N_{s*}$ これは最適な選択と n_s を選択した場合の利得の差をあらわしている．

k : $k = |N_s|$ ．

c : $c = |N_{s*}|$ ．

ϵ : a reasoning ability parameter.

ここで ϵ も logit モデルにおける σ と同様に, 単に推論能力に関するパラメーターを表すだけではなく利得の単位に関する調整パラメーターの役割も同時に果たしている．なお注意することとしてこのモデルにおいては ϵ が大きいほどエラー率が高くなることからわかるように ϵ が小さいほうが合理性の程度が高いことを意味している．

定義 8 指数エラーに基づく粗い推論モデルでは, $r(n_l)$ に $r(n_{s1})$ を割り当てる確率を以下で定義する．

$$\begin{cases} \min \left\{ \frac{1}{k}, \exp \left(\frac{-b(n_{s1})}{a} \epsilon \right) \right\} & \text{if } n_{s1} \notin N_{s*} \\ \frac{1 - \sum_{n_{s1} \notin N_{s*}} \min \left\{ \frac{1}{k}, \exp \left(\frac{-b(n_{s1})}{a} \epsilon \right) \right\}}{c} & \text{if } n_{s1} \in N_{s*} \end{cases}$$

この定義ではより直接的にエラーの生起確率が利得差の指数関数に比例することを仮定している。また深さは *logit* 関数のときと同様に差が緩やかになる方向に指数関数的に効いていることを仮定している。

どちらのモデルに関してもエラー率は a の増加関数で b の減少関数でありまたその効き方は概ね指数関数的である。したがって両者に共通の特徴が観察されるならば、このような性質を持つ関数形で表現される意思決定のより広いクラスに対して、同様の性質が成り立っていることが期待できる。次にこのモデルをナッシュ均衡と現実の被験者を用いた実験結果が大きく食い違うことで有名な Rosenthal のムカデゲームに適用することにより、ナッシュ均衡からのどのようなシステム的な逸脱が推論エラーによって生ずるのか検討する。

4.5 ムカデゲームへの適用

ムカデゲームはいろいろな派生形のゲームがあるもののもともとは有限の 2 人完全情報の展開形ゲームで、プレイヤー 1 と 2 が交互に戦略パス (P) またはテイク (T) を選択するゲームである。パスを選択すると選択したプレイヤーの利得は下がるものの、その減少分より大きい利得の増加が相手プレイヤーに与えられる。テイクを選択するとゲームは終了し、その時点での利得を両者が得ることによってゲームが終了する。ここでゲームが n 個の決定ノードを持つときにここではそれを n 階のムカデゲームと呼ぶ。

このゲームは有限ゲームであるため最後の決定節が存在する。その決定節において、その期の決断はそれ以降にまったく影響を与えないのでテイクを選択することが合理的な行動になる。したがってその一つ前の決定節において、仮に最終の決定節に到達したならば相手プレイヤーはテイクを取るであろうと合理的なプレイヤーは推論を行う。したがって今自分の決断する期が事実上の最終の決定節と同一視できることから合理的なプレイヤーは、最後から一つ前の決定節においてもテイクを選択することが合理的な行動となる。同様の議論が次々に成り立つため結果として全ての決定ノードにおいてテイクが選択されゲームは一回目で終了するという結果が唯一の部分ゲーム完全均衡による結果となる。

しかしながら下記の例のようにパスした際の利得の変分を自分が -1 、相手が 3 であるようなゲームにおいて、仮にゲームが 3 階であった場合の例を考える。すると均衡結果の利得 $(1, 0)$ は両者が P を選択し続けた場合の $(5, 4)$ という結果にパレートの意味で大きく劣ることとなる。実際に被験者を用いて指数関数的に利得が大きくなる状況で実験を行った McKelvey and Palfrey による実験結果を図 4.13 で紹介している。6 階のゲームの場合には、一回目には 9 割以上の被験者がパスを選択するという均衡とは著しく異なるものだった。しかしながらパスが選択される確率はゲームが進むごとに減少していき、最終期においては 9 割以上の被験者が均衡戦略であるテイクを選択するというものだった。

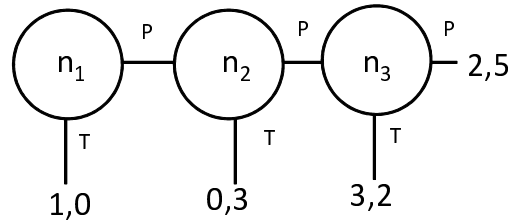


図 4.7: 3 階のムカデゲーム

ここでは n 階のゲームに対して一回目での協力行動の確率 (FCPF) に注目する. これはゲームの一回目でパスが選択された場合には, 残りの部分ゲームは $(n-1)$ 階のゲームと同一視できるので n 階ゲームにおける FCPF が計算できればあとは掛け合わせることで結果の分布を計算することが可能である. 実際に解析的な計算を行うのは $n=5$ 前後ですらかなり煩雑な式になるので今回はモンテカルロシミュレーションを用いて各 100 万回の平均値をグラフにプロットしたものが以下のものである. なお以下のシミュレーションでは全て両プレイヤーの推論能力は同一で, またノードによっても変化はしないものと仮定している.

ここで図 4.8 においては σ が大きいほど合理性の高いことを意味し, 図 4.9 においては ϵ が低いほうが合理性が高いことを意味している. まず目に付くこととし, てゲームの階数と FCPF に関して言えば, ゲームが長くなるにつれて協力行動の頻度は増大している. これは McKelvey and Palfrey の実験結果を非常によく説明していると考えられる.

次に FCPF と推論能力の関係に関して注目してみると, ここで上げた全ての n に対して両方のモデルにおいて極値が唯一つだけ存在することがわかる. 極値に至るまでは合理性の程度が増大するにつれて FCPF も上昇傾向にある. しかしながら合理性の程度が極値を越えた後では協力行動の頻度が減少することがわかる.

このことからわかることは最も合理性の低い主体による決定結果が最も非合理的な結果を最高の頻度で取るということに必ずしもつながるわけではなく, 適度な合理性を持つプレイヤーによる推論結果が却って完全なランダムな推論よりも非合理的な行動を高い頻度で取ることにつながることもあるということもわかった.

次に両者の平均利得と合理性の程度に注目する. なお $APP1$ はプレイヤー 1 の平均利得を, $APP2$ はプレイヤー 2 の平均利得をあらわしている. 図 4.10 と図 4.11 では 10 階のムカデゲームについて図 4.10 は logit 関数のモデルによる結果で 図 4.11 は指数エラーのモデルによる結果である. 図 4.10 においては $\sigma=5$ の辺りで, 図 4.11 においては $\epsilon=10^{-3}$ の付近で平均利得が最大となることがわかる.

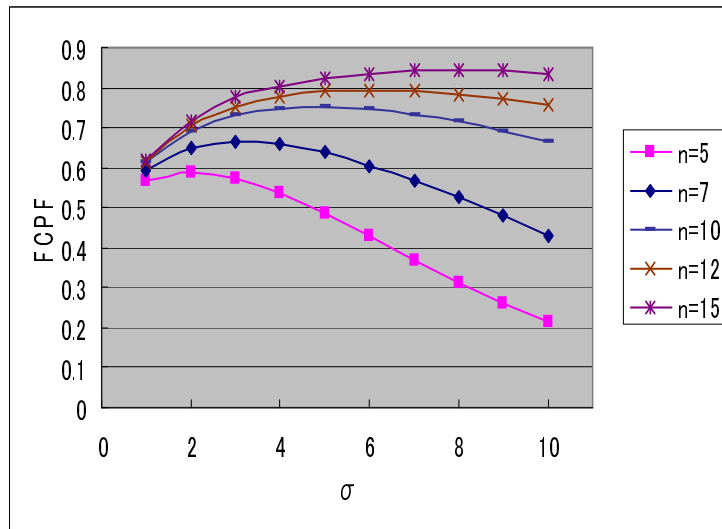


図 4.8: logit 関数モデルにおける FCPF

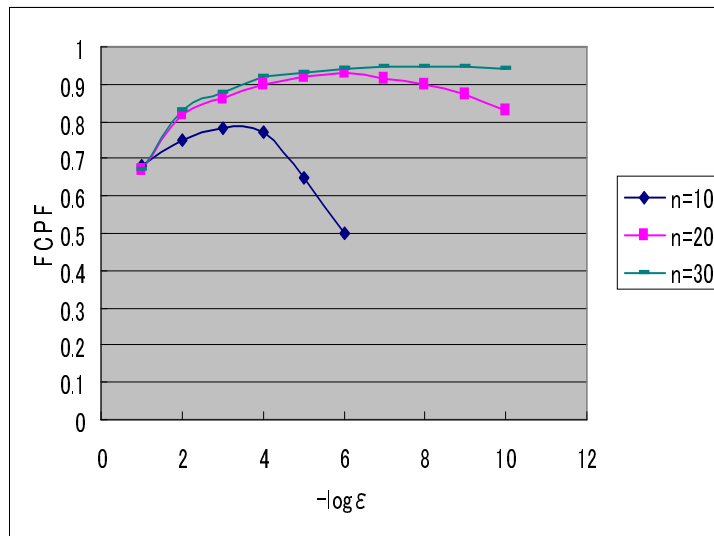


図 4.9: 指数エラーモデルにおける FCPF

次にこれらの結果を Heifetz の結果と比較してみる．彼のモデルにおいては枝が 2 本の場合で，エラー率は深さや利得の差には依存せず一律であると仮定している．彼の結果は図 4.12 に紹介している．四回目までは確かに回数が伸びるほど非合理的なパスの頻度が高くなるという結果になるものの，それ以上回数が増えてもパスの率はエラー率の値によっても異なるものの 0.7-0.9 程度の間をループするというやや奇妙な結果になっている．したがって少なくともムカデゲームの場合においては本モデルの割り当て処理のほうが描写能力が高い結果を得ている．

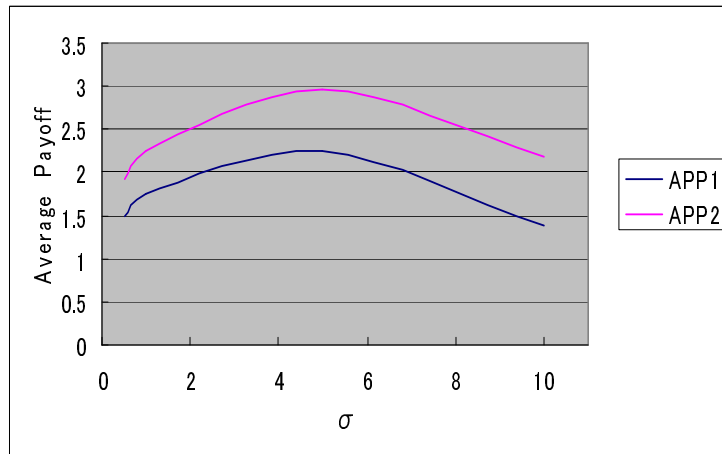


図 4.10: Logit 関数モデルにおける平均利得

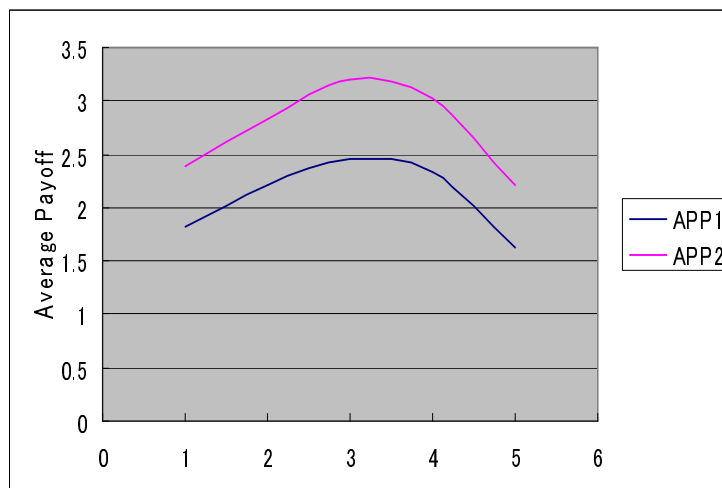


図 4.11: 指数関数モデルにおける平均利得

一般に両者の利得の和を最大化するための行動はしばしば自分の利得を減らすことも少なくない。ムカデゲームは拘束力のある合意が作ることができない状況で協力行動が選択肢にあるような社会の一つのモデル化とも考えられる。このような状況においては、完全に合理的なプレイヤーからなる社会においては、仮に罰則システムのようなものの監視が不十分であるならばパレート効率的な状態からの逸脱を十分には阻止できずに結果として社会厚生が低下してしまう可能性もある。しかしながら合理性の程度が不十分であれば、必ずしも厳密な監視や罰則システムなしにもパレート効率的な結果が維持されうることから推論能力に関する合理性の増大は必ずしも社会厚生を上げることと一致するとは限らず、適度な合理性の社会が最大の社会厚生を達することになることもありうるという示唆を得た。

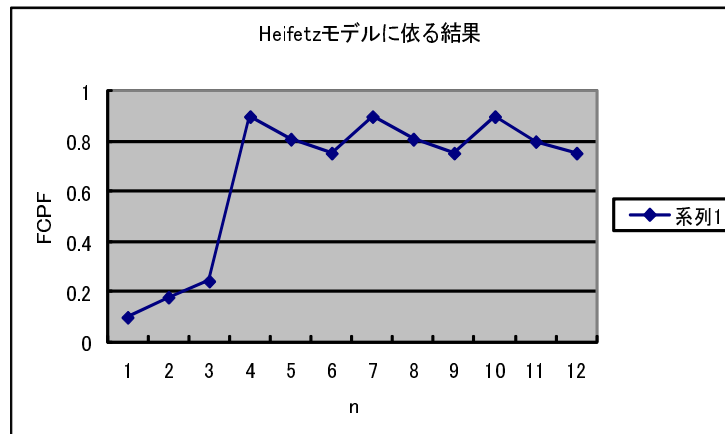


図 4.12: Heifetz の結果

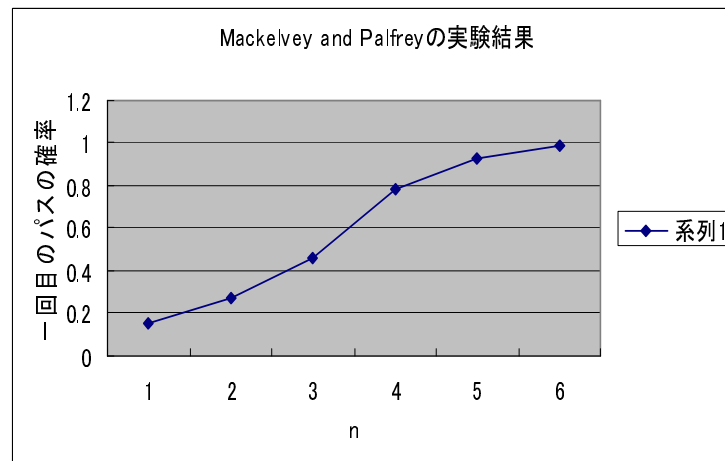


図 4.13: ムカデゲームの実験結果

4.6 現実の意思決定における推論不完全性

問題が数学的には厳密に定式化できるにもかかわらず、その上で実際に何をしたら良いのか推論を正確に行うことが困難な状況の例としては、チェスや将棋などの2人ゲームが典型的に挙げられる。人工知能などの分野の研究者などがこの問題に取り組んでいてゲーム用のプログラムの研究も進められている。実際にオセロなどでは完全な解析はまだ終わっていないものの、人間のチャンピオンを凌駕するレベルのアルゴリズムが提案されている。この理由として言われているのは人間の意思決定は計算能力の精度ではコンピューターに劣りミス完全に排除することは不可避だから、特に単純な計算の積み重ねの正確性を要求されるような状況以外の部分でリードしない限りは勝てないということである。特に推論の積み重ねていく仮定では、例え一箇所に関する推論ミスの可能性が低かったとしても推論ノードを深く読もうとすると

エラー率は急激に増大していくと考えられる。

この事実に関して実際のプレイヤーはどのように振舞うのかというと、とりわけ有利な側はチェスで言うならばお互いの駒が出来る限り少なくなるようにして可能な変化の数を減らすような単純化戦略が有効であると言われている。これは局面が単純になるならば自分だけではなく相手のミスも減ると考えられるが初期状態で有利ならばそのままリードを保って逃げ切れる可能性が高いからということで説明できる。一方不利な側は、できるだけ駒を多く残し変化の数を増やすことによって、自分のミスの可能性も増大するものの相手のミスの可能性も上がるため挽回出来るチャンスは増えると説明することができる。このように状況が有利ならば複雑性を減らしリスクを回避する方向に、状況が不利ならば複雑性を増大させリスクを求める方向への行動ということがこういったゲームにおけるセオリーとして一般的に認識されている。

一例として将棋では、自分の王が安泰の場合は、長手数かかる詰みがあったとしても、その手順を採らずに安全な手を選ばれることが少なくない。また不利な側は、自分の王が詰むことがわかっていても、それが複雑な場合は、相手のミスに期待してそれを無視して相手への攻めを採用する場合もある。

実際には上に挙げた例のような二人完全情報零和ゲームは、ゲーム内に不確実性の要素は全く存在せず、従って理論的にはリスクという要素は出てこない。しかし構造の複雑性を実際の人間が取り扱う際には、自身の推論能力が完全でないことを知っているプレイヤーは限定合理性を認識したうえで、問題の複雑性に対する自身推論の能力の限界を不確実性に変換し取り扱うような行動の例と考えられる。

それではなぜそれほど複雑な状況において、計算量に劣る人間のほうがコンピューターよりも相対的に優れた判断を下すことができるのであろうか。それに対する一つの要因としては、自分の推論能力を過信せずに、それに対してリスク回避的な判断を加味した上で意思決定を行っているということがあげられる。自分の推論能力が不完全であることを十分に認識しているプレイヤーの場合は、結果を一点まで絞りきらずにいくつかの候補をあげそれらの上に確率分布を振るというような処理をする可能性が高いと考えられる。たとえ厳密に最適な手の列がわからなくても最善に近い手を尽くした場合の結果の範囲が比較的狭い範囲に収まる場合には、仮想利得を一点ではなく範囲で置き換えた上で他の手との比較を行うと考えられる。そこである手を選んだ場合の評価の最小値が他の手の最大値を上回っているならばどの手を選択すべきか合理的に決定することが可能である。

このようなプロセスは複雑性の不確実的な変換と解釈できる。囲碁や将棋といった完全情報のゲームにおいては不確実性が本来は全く存在しない。しかしながら時間投入量と計算精度の関係は収束すると考えられるため、有限の時間で解くことができないこともまた明らかである。従ってどこかの時点で割り切って意思決定しなければならない状況に迫られる場合も少なくない。その際に自身の推論能力を過信せずにこのように問題の変換を行ったうえで意思決定するほうがリスクを回避できる可能性があることを次の例

で示す.

選択肢 A を選んだ場合には確率 0.7 で利得 10 得られるものの、確率 0.3 で利得 -10 という結果になる判断したとする. 一方選択肢 B を選ぶと確率 0.5 で 8.1, 確率 0.5 で 7.9 になると判断したとすると、不確実性下の意思決定として主観的な期待利得を計算すると選択肢 B のほうが高くなるため、選択肢 B を選ぶということも一つの判断方法であると考える. しかしながらここで注意する必要があるのは、ここで出した確率というのは真の意味での確率ではなく、自身の推論能力に対する信念とも関係のあるパラメーターであるということであり、特に一般に最適応答の箇所が決定しづらい一方最終的な結果は大きくばらつくような状況においてはこのような評価を行ったとしても不安定になりやすい.

逆に最適応答を決定することが困難でも、多少のずれが最終的に大きな違いにはつながらないような状況ならば、このような変換手法の精度は高いと考えられる. 現実的な意思決定状況に戻って考えてみると、現実的な変化の絶対数が多かったとしても、局所的に評価がある程度まとまっているような状況においては、このような変換で意思決定を評価したとしても、真の値からのずれはあまり大きく無いと考えられる. 従って上の例においては、選択肢 B のような状況のほうが推論精度が低かったとしても最終的な判断が入れ替わってしまうようなリスクは低いと考えられる.

実際にチェスなどでも言われているのは、人間は対コンピューターに対しては、一寸間違えたときに利得が大きく振れるような展開に持ち込むよりは、このように局所的にはあまり変化が大きい変化に誘導するほうが勝ちやすいといわれていて、例え手数が長い変化であったとしても、このような結果の位相が近いような集合をまとめてうまく判断することができるということが人間推論のコンピューターに対する優位性となっていると考えている.

またもう一つの例としては我々が数学的な論文の証明を意思決定モデルで表現することを考える. このとき仮説が証明できるかどうかということは、仮に複雑な仮説で一見では判断が困難であったとしても、論理的には命題が真であるか偽であるかあるいは問題として閉じてないかのどれかが完全に決まっていて、不確実性は存在しない. しかしながら特に難しい問題であれば証明ができるかどうかということはたとえ一線級の研究者の推論能力でも簡単にはわからないものも存在する. どの問題に取り組むべきなのかということに関しては研究者本人の興味などの影響は受けるが一般には問題が解けた場合の社会的な評価と解けるのかどうかという主観確率によってあたかも不確実性下のような意思決定を行う場合も多いと考えられる.

これとは別に複雑でなにもわからないからランダムで決定するという哲学も考えられると思われるが、仮に不正確であっても近似的に複雑性のある種の不確実性に落とし込むことによってランダムよりはいい結果を期待値的に得られる場合も少なくないと考えられる. 複雑性をそのまま合理的な方法で取り扱うことは難しい場合も多いので、不確実性への変換ということも複雑性を実践的に有効に取り扱う一つの手法で

あると考える。

4.7 4章のまとめ

第4章では、まず始めに状況が複雑である場合にどのような意思決定を行うのかということに関して、自分の推論能力が完全ではないことを認識した上で、結果の予測を一点ではなくある程度幅を持たせて推論することに関する数理的な性質を議論した。これは問題状況の複雑性をそのままでは扱うことができないので、あたかも不確実性のように変換することで、推論エラーによるリスクを回避しようとするような態度である。このような態度が効果的なのは、結果の位相がある程度まとまっていて、全体的な効用が部分ゲームの総和にある程度近いような形で表現できる場合に特に有効になりがちであるということを議論した。

次に数理的なモデルを提案するために、完全情報の展開形ゲーム上で、自分の推論能力を完全に信じているけれども、実際には推論エラーを犯す可能性のあるようなプレイヤーによる意思決定のモデル化を行った。その際に特に推論エラーに関して、高い利得を生み出す戦略ほど高確率で正しいものと推論し、遠い将来の段階の決定に関する推論ほど、判別精度が落ちるという二つの仮定により特徴づけを行った。

さらにより詳細な性質を探るために、これらの二つの仮定に関して、指数関数的に影響を受けるような二つの具体的なモデルである、logit 関数エラーモデルと指数エラーモデルを提案し、それらのモデルをゲーム理論における均衡解の概念に依る結果と実際に被験者を用いて実験を行った結果が著しく乖離している例であるムカデゲームに適用した。ムカデゲームにおける協力行動は、実験においては回数が長くなるほど単調に増加するという結果である。本論文の二つのモデルによる結果はそれと非常に良く適合している。また回数によってピークの位置は変わってくるものの、いずれの回数の場合においても推論能力が中程度の場合が行動レベルでは非合理的な行動の採用確率が最大となるという結果を得た。このことは、推論能力が0の状態からたとえ推論能力が向上したとしても、かえって非合理的な行動の採用確率が上昇してしまうこともあるという皮肉な結果を表している。また期待利得を計算すると両プレイヤーにとって、中程度の推論能力のときに共に期待利得が最大となった。このことから中程度の推論能力をもつプレイヤーからなる社会のほうが、かえって推論能力が完全なプレイヤーからなる社会よりも社会厚生が大きくなることもありえるという結果を得た。

第5章 結論と今後の課題

5.1 結論

本論文の主たる知見は次のようなものである。

1. 不確実性の問題に対して複製状況において期待到達時間最大化という概念を定義し、それがタイプの支配という概念と目標無限大の基準の下で同一であるということを示した。さらに実際にポートフォリオ問題に適用することにより、ポートフォリオの分散減少効果は期待到達時間の低下という形で正確に現れることを示したうえで期待到達時間が有限の場合でも十分に扱いやすい指標であることも明らかにした。
2. 複雑性に関しては、まず自分自身の推論能力が完全ではないような状況に関して、自分の推論能力の不完全性を認識しているプレイヤーの、複雑性の不確実性への変換について議論を行い、それが効果的となるための条件を探った。
3. さらに完全情報の展開形ゲームの例において、推論精度が利得の差の減少関数で現在からの遠さの減少関数と仮定することにより、従来のモデルよりも一般的な形で、不完全な推論に関する数理的なモデルの定式化を行った。そのうえで従来いろいろな形で説明がなされてきたムカデゲームにおける協力行動の発生に関してこれまでのモデルよりもより適合性の高い結果を得た。また中途半端な合理性による推論能力が、ランダム推論よりも高い確率で非合理的な行動を導く可能性や、完全推論よりも高い社会厚生を生み出す可能性などの結果も得た。

現在の意思決定理論では数理的に厳密に扱うことができる状況に絞ったモデル化が中心である。しかしながらそこで仮定する合理性は、人間の意思決定コストや能力の問題などに関しては超越的な仮定をしているため、狭い意味での合理性の表現となっている。これに関しては現実的な仮定を置こうとすると、とりわけ数理的な分析が困難になりがちという理由も大きいからと考えられる。

現実には問題状況を必ずしも十分に理解できないまま意思決定を決断せざるを得ないことは現実にも多いが、十分理解できないということ定式化するのは非常に困難である。しかしよくわからないからといって

適当に決定しているわけではなく、能力の限りにおいて最善を尽くそうとしている場合が多く、またそのような態度のほうがランダムな意思決定よりは良い結果を生むことが一般には多いと考えられる。

Ashby は問題の複雑性が自身の能力を上まわっているときには必ずしもうまく決定ができるとは限らないと最小多様度の法則で主張しているが、それでも能力なりに問題に取り組むことによりどのような結果が得られるのかということは重要な問題であると考えている。またそのような限定的な能力により問題に対処することによるリスクをどのようにマネジメントするのかという点は今後現実的にも一層重要な問題となると考えている。

第4章でも議論したように、人間の経験的な処理はしばしばコンピューターに依る計算量を超える成果を出すこともある。このような進化的に紡がれてきたリスクに対する対処法を複製状況における一つの評価基準として表現したものが期待到達時間であるということができると考えている。

5.2 今後の課題

今後の課題として不確実性に関しては、期待到達時間は今回扱った範囲のモデルよりもより広い範囲で定義は可能な一方数理的な扱いの困難さゆえ今回は限定的な複製状況に絞っての議論を行った。今後は数理的にどのような範囲まで同様の議論が定義できるのか考察していきたい。

また複雑性に関しても今回は完全情報の展開形ゲームに関して扱った。これが一般の完備情報や同時意思決定の標準形のゲーム状況の場合にはどのようなモデル化が妥当なのかということを考察していきたい。今回の完全情報のゲームに関しては、完全に合理的な推論モデルから崩す形で不完全な推論モデルの定式化を行った。しかしながら同時意思決定のモデルにおいては、そもそも何を考えて意思決定しているのかということに関しては未だ決定的な見解が統一されておらず、ナッシュ均衡が実際に存在しても、それが確実に起こるとは限らないという議論もある。従って完全な推論モデルが定義できないとなると、そこを基本として不完全な推論モデルの構築ということもできないため、これに関してはまずはナッシュ均衡が一個だけしか存在しない場合など状況を絞ることにより、議論をしていきたいと考えている。

さらに複雑性の不確実的な取り扱いということに関しても、近年 PC 市場や携帯電話市場などにおいても、機能の多種性よりもむしろわかりやすさをセールスポイントにした商品のヒットが続くなどの例が見受けられる。これは意思決定に関するコストが下がることは適応できる範囲の増大よりも望ましいと感じる消費者も少なく無いということだと考えられる。この問題に関して複雑性は単純なコストだと仮定して分析を行うような研究も当然有力であるとする。しかしながらそのコストがなぜ発生したのかということとを突き詰めて考えると、複雑な操作が要求されるということとはどのような結果が起こるのかということ

に関して正確な予測が不可能になり、結果としてリスクを回避したがるという解釈も可能であると考える。

このような情報の少ない商品に関する意思決定に関してはマーケティングの分野においては消費者は複数のタイプが存在して、目あたらしいものであれば商品の性能がわからないうちからでも飛びつきたがるような少数の先行的な消費者と、商品の性能が明らかになってから行動を決定する多数の消費者とセグメンテーションを行うことによる分析が一般的となっている。

複雑な状況で、状況の評価がはっきりしない場合に最終的にどのような行動が観察されるのかということに関する問題はこの例などからもわかるように個人差もかなり大きく誰にでも当てはまるような定式化というのは確かに困難かもしれない。しかしながら現代の我々の生活において、一人一人が触れる情報量は確実に増大しているために、今後は一つの問題に関して割くことのできる情報収集時間や分析時間は減少しがちになることが予想される。従ってこのような社会状況の分析は一層重要になっていくものと考えている。

参考文献

- [1] 伊藤邦武, 人間的な合理性の哲学, 勁草書房, 1997
- [2] 猪原健弘, 限定された視野を持つプレーヤーが行う展開型ゲームの分析, 経営情報学会 2000 年秋季全国研究発表大会プログラム, 2000
- [3] 巖佐庸, 数理生物学入門—生物社会のダイナミクスを探る, 共立出版, 1998
- [4] 宇沢弘文, ケインズ『一般理論』を読む, 2008
- [5] 宮川公男, 意思決定論, 丸善, 1975
- [6] 依田高典, 不確実性と複雑性の経済学, 日本評論社, 1997
- [7] 金子守, ゲーム理論における合理性と限定合理性, 筑波大学ディスカッションペーパー No.1100, 2004
- [8] M. Allais, "Le comportement de l 'homme rationnel devant le risque: critique des postulats et axiomes de l 'ecole Americaine", *Econometrica* 21, 503-546, 1953
- [9] R. Aumann, "Correlated Equilibrium as an Expression of Bayesian Rationality", *Econometrica* 55, 1-18, 1992
- [10] R. Aumann, "On the Centipede game", *Games and Economic Behavior* 23, 97-105, 1998
- [11] D. Bell, "Disappointment in Decison Making under Uncertainty", *Operations Research* 33, 1-27, 1985
- [12] L.Blume, and D.Easley, "Evolution and Market Behavior" *Journal of Economic Theory* 58, 9-40, 1992
- [13] R. Axelrod, "The Complexity of Cooperation", Princeton, 2001
- [14] L. Blume and D. Easley, "Optimality and Natural Selection in Markets", *Journal of Economic Theory* 107, 95-135, 2002

- [15] J. Doob, "Stochastic Processes", John Wiley, 1953.
- [16] R. Durrett, "Essentials of Stochastic Processes", Springer 1999
- [17] E. Ebeling, "An Introduction to Reliability and Maintainability Engineering", McGraw-Hill Companies Inc., Boston, 1997
- [18] H. Einhorn, and R. Hogarth, "Decision Making under Ambiguity Rational Choice", The University of Chicago Press 41-66.
- [19] D. Ellsberg, "Risk, Ambiguity, and the Savage Axioms" Quarterly Journal of Economics 75, 643-69, 1961.
- [20] W. Feller, "An Introduction to Probability Theory and Its Applications", John Wiley, 1966.
- [21] D. Fudenberg, and D. Tirole, "Game Theory", MIT Press, 1993
- [22] H. Watson and F. Galton, "On the probability of the extinction of families," J. Anthropol. Inst. Great Britain Ireland 4, 138-144, 1874
- [23] I. Gilboa, and D. Schmeidler, "Maximin Expected Utility with Non-Unique Prior", Journal of Mathematical Economics 18, 141-153, 1989
- [24] J. Harsanyi and R. Selten, "A General Theory of Equilibrium Selection in Games", MIT Press, 1988
- [25] A. Heifetz, A. Pauzner "Backward Induction with Players who Doubt Others' Faultlessness" Mathematical Social Sciences 50, 252-267, 2005
- [26] D. Kahneman, and A. Tversky, "On the Psychology of Prediction", Psychological Review 80, 237-251, 1973
- [27] L. Kelly, "A New Interpretation of Information Rate," Bell Systems Technical Journal 35, 917-926, 1956
- [28] F. Knight. "Risk, Uncertainty, and Profit", Orlando, Signalman Publishing, 1921
- [29] M. Machina, "Dynamic Consistency and Non-Expected Utility Models of Choices under Uncertainty" Journal of Economic Literature 27, 1622-68, 1989
- [30] H. Markowitz, "Portfolio selection", The Journal of Finance 7 77-91, 1952

- [31] R. McKelvey and T. Palfrey, "An Experimental Study of the Centipede Game", *Econometrica* 60, 803-836, 1992
- [32] R. McKelvey and T. Palfrey "Quantal Response Equilibria for Normal Form Games", *Games and Economic Behavior* 10, 6-38, 1995
- [33] R. McKelvey and T. Palfrey "Quantal Response Equilibria in Extensive Form Games", *Experimental Economics* 1, 9-41, 1998
- [34] R. Myerson "Refinements of the Nash Equilibrium Concept." *Int. Journal of Game Theory* 7, 73-80, 1978
- [35] I. Nonaka and H. Takeuchi, "The Knowledge-Creating Company: How Japanese Companies Create the Dynamics of Innovation", Oxford University Press, 1995
- [36] J. Nash "Non-Cooperative Games" *The Annals of Mathematics* 54, 286-295, 1951
- [37] K. Neubeck, "Practical Reliability Analysis", Prentice Hall, 2004.
- [38] M. Osborne and A. Rubinstein "A Course in Game Theory", MIT Press, 1994
- [39] A. Robson, "The Evolution Attitudes to Risk: Lottery Ticks and Relative Wealth" *Games and Economic Behavior* 14, 190-207, 1996
- [40] A. Robson, "A biological basis for expected and non-expected utility", *Journal of Economic Theory*, 68, 397-424, 1996
- [41] R. Rosenthal, "Games of Perfect Information, Predatory Pricing and the Chain-Store Paradox" *Journal of Economic Theory* 25, 92-100, 1981
- [42] A. Rubinstein, "Modeling bounded rationality", MIT Press, 1998
- [43] P. Samuelson, "Foundations of Economic Analysis", Harvard University Press, 1947
- [44] R. Selten "The Chain-Store Paradox" *Theory and Decision* 9, 127-159, 1978
- [45] R. Selten "A reexamination of the perfectness concept for equilibrium points in extensive games." *International Journal of Game Theory* 4, 25-55, 1975
- [46] H. Simon, "A Behavioral Model of Rational Choice", New York: Wiley, 1957

- [47] J. Smith, "Evolution and the Theory of Games," Cambridge Univ. Press, 1982
- [48] F. Spitzer, "Principle of Random Walk", D.Van Nostrand, 1964
- [49] W. Viscusi, "Sources of inconsistency in societal response to health risks", American Economic Review. 80, 257-261, 1980
- [50] J. von Neumann and O. Morgenstern. "Theory of games and economic behavior", Princeton University Press, 1944.
- [51] F. Craig and A. Tversky, "Ambiguity Aversion and Comparative Ignorance", Quarterly Journal of Economics 110, 585-603, 1995
- [52] A. Wald, "On cumulative sums of random variables", Annals Mathematical Statistics 15, 283-296, 1944
- [53] J. Weibull, "Evolutionary game theory," MIT Press, 1995
- [54] O. Williamson, "The Economics of Organization: The Transaction Cost Approach", American Journal of Sociology 87, 548-577, 1981
- [55] P. Young "An Evolutionary Model of Bargaining" Journal of Economic Theory 59, 212-220, 1993

謝辞

本研究の過程では非常に多くの方のお世話になった。とりわけ指導教官である木嶋恭一先生には、修士課程からゼミでの熱心な指導にとどまらず、投稿論文の作成から国際会議の発表まで多くのポイントにおいて、有益なコメントを頂いた。また博士論文の作成に関しては筆者のモチベーションが低下しがちな時も丁寧に暖かく見守っていただき、それを支えになんとか纏め上げることができたもので、木嶋先生なしにはこの論文の完成もなかったと断言できるほど本当にお世話になったと思っている。

また猪原健弘先生にも、修士課程時代のゼミから博士課程にわたるゼミで大きくお世話になっただけでなく、先生の多岐に渡る研究範囲のなかでも将来の行動を期待值的に評価する研究は本論文の複雑性に関する研究の基となるアイデアとなるなど研究者としてインスパイアされる部分が大きかった。

出口弘先生には、博士課程在学中から共同ゼミで非常に広範な知識に基づくコメントを頂くことに始まり、21世紀 COE プログラムの代表者として、様々なプログラムの運営から世界の広い分野の研究者と交流する機会を作っていただきモチベーションを高められることが多かった。

松村良平先生には、修士課程から博士課程に至るゼミのなかで兄貴分として暖かいコメントを頂くとともに、プリンシパルエージェント問題に関する議論などを通じて研究者としての守備範囲の拡大に大きなお世話になった。

オーストラリアのニューサウスウェールズ大学の Norman Foo 先生には博士課程中の在学中に2ヶ月ほどの滞在させていただき論理的な視点からのコメントをいただき、視野の拡大に大変役立った。

小林憲正先生には修士課程時代から研究に関して哲学的なところから、数学的な定理のところまで非常に幅広い視野に渡り議論していただきまた英語論文の際などのコメントでもお世話になった。小山祐介先生には共同ゼミや国際会議の現場などで、経済学やシミュレーションの専門家としての視点からコメントを通じて新たな視点の発掘に刺激を受けた。荒井祐介先生には COE 関係におけるプログラム運営などの面でマネジメントの手法を学ばせていただいた。増田浩通先生には研究室のゼミからシミュレーションに関するコメントまでお世話になった。その他の木嶋研究室の修士博士課程時代における先輩や後輩には暖かい刺激から新たな視点での研究などでモチベーションを度々回復させていただいた。松丸美智子さんには修士課程のころから研究室内での事務的なことには大変お世話になりました。また勢川聡美さんと矢内恵美子さんに

は,COE プログラムの事務に関して大変お世話になりました.

最後に私事になりますが父昌雄と母悦子,妹有希子には研究生活に集中させて頂くことの謝意を表すことをお許しいただきたい.