

ネットワーク構造がもたらす協調の頑健性と脆弱性

石田 芳文 山本 仁志 岡田 勇 太田 敏澄

相互作用関係がスケールフリーネットワークである繰返し囚人のジレンマにおいて、協調は安定的に維持されるのであろうか？また、協調を安定的に維持するためにはどのような方策が望ましいだろうか。我々は、エージェントの相互作用関係がスケールフリーであるネットワークに対し、協調の進化と崩壊の動的なメカニズムを明らかにするために、エージェントベースシミュレーションをおこなった。我々の成果は、以下の通りである（１）スケールフリーネットワーク上の囚人のジレンマにおいて、協調は達成されるが脆弱であり崩壊しやすい（２）ネットワークの接続次数の高いエージェントが寛容になることが協調の崩壊を招き、また厳格になることで協調は最も高いレベルで維持される（３）ネットワークの接続次数の高いエージェントが、一定の確率で裏切り行為をとることで、集団全体の寛容性を抑制し、裏切り戦略の侵入を防ぐ効果がある。

Is cooperation evolutionally stable in the iterated prisoner's dilemma (IPD) game in which individuals interact on a scale-free network? Or what is an appropriate strategy for keeping cooperation stably in the game? We simulate a scale-free network of interactions among individuals in the IPD game and use an agent-based approach to investigate dynamics of the emergence of cooperation and the cooperation collapses. Our results are the following: 1) On the scale-free network, we can achieve cooperation to be sustainable, but it is vulnerable to collapse. 2) The reason cooperation collapses is that agents who have high node degrees are tolerant of others. A high level of cooperation is sustainable when those agents are strict to others. 3) What agents with high node degrees act non-cooperatively with a slight probability has effects to inhibit tolerance among individuals and to prevent invasion of defections.

1 はじめに

個々の相互作用においては非協調が優位であるにも拘らず、集散的な協調現象が出現するメカニズムは、社会学的にも生態学的にも重要な課題であり、協調の

進化に関する多くのシミュレーション研究がなされている。協調の進化に関する Axelrod(1984) の古典的研究[2] は、互惠主義の優位性を示している。しかし、相互作用関係に総当たりの完全グラフを想定しているのみで、ネットワーク構造が協調の進化に与える影響に関する議論は不完全である。ネットワーク構造の特徴は、例えば接続次数の非均質性で表現できる。現実の生物学的ネットワークや社会的関係は非均質な構造を有しており、スモールワールド性を持つことや[14]、スケールフリー構造を持つこと[3] が示されている。Joeng et.al. は、タンパク質ネットワークがスケールフリーネットワークであることを示している[7]。社会科学の領域においても、Ebel et al. は、Email のトラフィック量がスケールフリーネットワークであることを示し[5]、Barabasi and Albert は、俳優の共演関係においてスケールフリーネットワークが見出せることを示している[3]。我々は非均質な構造

これは草稿バージョンです。最終稿は以下の論文誌を参照ください。コンピュータソフトウェア, Vol.24, No.1(2007), pp.70-80.

Robustness and Vulnerability of Cooperation: Emergence on Network Structures.

Yoshifumi Ishida, 電気通信大学大学院情報システム学研究科, Graduate School of Information Systems, University of Electro-Communications.

Hitoshi Yamamoto, 立正大学経営学部, Faculty of Business Administration, Rissho University.

Isamu Okada, 創価大学経営学部, Faculty of Business Administration, Soka University.

Toshizumi Ohta, 電気通信大学大学院情報システム学研究科, Graduate School of Information Systems, University of Electro-Communications.

を持つネットワークの一例としてスケールフリーネットワークに着目する．

スケールフリーネットワークは協調の進化にどのような影響を与えるのだろうか？Abramson and Kuperman はネットワークポロジとネットワーク密度とが，協調割合に与える影響について調査した[1]．また，Ifti et al. は，隣人のサイズが増大すると非協調が増加することを示している[6]．Vukov and Szabo は，ネットワークに階層構造を仮定すると，階層の Lower level で非協調が増加することを見出した[13]．このように接続次数やネットワーク構造の違いが協調の進化に影響を与えることはよく議論されている．しかし，なぜネットワーク構造がトラスから非均質に変わることによって非協調が増加するかの原因とメカニズムはわかっていない．更に，どのような戦略や方策が脆弱といわれるネットワーク構造においても，協調を高いレベルで維持することができるかがわかっていない．我々は，この二つの疑問に，スケールフリーネットワークを用いた繰り返し囚人のジレンマゲーム実験を行い答える．

Cohen et.al. は，相互作用関係を 2 次元格子トラス上のノイマン近傍に固定することにより協調関係が安定することを示している[4]．我々は，2 次元格子トラスとスケールフリーネットワークにおいて協調のレベルに差があるのかをシミュレーション実験する．続いて，協調の崩壊が生じるメカニズムを明らかにするために，ネットワークの構造を規定する接続次数ごとにエージェントを分類し，それぞれの振る舞いを観察する．最後に，それぞれの接続次数のエージェントの戦略を操作することで，エージェントの接続次数と戦略が，集合的な協調の脆弱性と頑健性にどのような影響を及ぼすかを明らかにする．

2 モデル

本節では，ネットワーク構造と協調の進化の相互関係を明らかにするため，スケールフリーネットワークを用いた繰り返し囚人のジレンマゲームの，シミュレーションモデルを構築する．

2.1 囚人のジレンマを用いた協調の進化シミュ

表 1 PD の利得行列

		プレイヤー 2 の行為	
プレイヤー 1 の行為		協調	非協調
	協調	(R, R)	(S, T)
	非協調	(T, S)	(P, P)

レーション

繰り返し囚人のジレンマ (Iterated PD) は，協調の進化を扱う基盤としてのモデルの 1 つである．PD は，生物学，社会学の多くの領域で，ミクロの相互作用から集合的秩序が創発するメカニズムを理解するために分析されている．PD は 2 人のプレイヤーが対戦する．ゲームにおいて，2 人のプレイヤーは 2 つの行動が可能である (協調，非協調)．利得行列を表 1 に示す．

PD の必要条件は，以下の不等式である．

$$T > R > P > S$$

$$2R > T + S$$

このゲームでは，プレイヤーは短期的には相手を裏切る (非協調) ほうが利得は高くなる．しかし協調関係が長期的には維持されると指摘されている[2]．PD は，なぜ生物が他者に対して協調的に振舞うことがあるのかという問いに答えるために有効なツールである．ゲームはエージェント，ネットワーク構造，PD，適応プロセスの 4 つのパートを持つので順に記述する．

第一に，我々は PD におけるプレイヤーを，内部に自律的な戦略を持ち他者と相互作用する主体であるエージェントとして定義する．第二にネットワーク構造は，エージェントとエージェント間のリンクで規定され，リンクの存在するエージェント間で対戦がおこなわれる．第三に，エージェントは対戦相手と対戦することによって利得を得る．ゲームの結果，各々のエージェントが得た利得は，環境における「適応度」とみなされる．最後に，各々のエージェントは適応プロセスを行う．エージェントは最も高い利得を得た隣人の戦略を模倣することある．

我々は，エージェントの戦略における適応プロセスのダイナミクスを観察することで，協調の進化のメカニズムを解明する．我々は，次のセクションでモデルの詳細を記述する．

2.2 モデル詳細

エージェントは内部に「戦略」を持ち「戦略」に従い他エージェントと対戦する．対戦の結果，高い利得を得た戦略は次世代まで生き残り，低い利得の戦略は淘汰される．その結果，利得獲得に有利な戦略が支配的となりその社会における行動秩序が形成される．

エージェントは，行動戦略 $Act^x = (a_i^x, a_p^x, a_q^x) \in [0, 1]^3$ を持つ．エージェントは同じ対戦相手と 1Period あたり 4 回の対戦を行う． a_i^x は第一回目の対戦で協調する確率であり， a_p^x は前の対戦で相手が協調行動をとったときに協調行動をとる確率， a_q^x は前の対戦で相手が非協調行動をとったときに協調行動をとる確率である．本論文では単純化のため $i = p$ としている．この記述法は，PD ゲームにおける基本的な戦略を含んでいる． (a_p^x, a_q^x) が $(1, 1)$ であることは常に協調 (All-C)， $(1, 0)$ はしつべ返し (TFT)， $(0, 0)$ は常に非協調 (All-D) を表すことになる．

シミュレーションは「PD ゲーム」と「進化」のフェーズで記述される．PD ゲームフェーズでは，エージェントは対戦可能な相手と既定回数対戦（本研究では 4 回）する．エージェントの対戦可能な相手はネットワーク構造に依存する．ネットワーク構造は，2次元格子トラス (2DK: 2-Dimensions Keeping)，スケールフリーネットワーク (SF: Scale-Free Network) の 2 通りである．各エージェントの対戦相手の数をそのエージェントの接続回数と呼ぶ．2DK において，エージェントは格子上のノイマン近傍に配置された 4 人の固定されたエージェントと対戦する．SF においては，エージェントはスケールフリー構造を持つネットワークで他のエージェントとリンクしており，リンクが存在するエージェントと対戦する．

進化フェーズでは，エージェントは対戦したエージェントで最も高い利得を獲得したエージェントの戦略を模倣する．模倣の際には進化過程における突然変異が生じる．突然変異は確率 ϵ_1 の比較エラーと確率 ϵ_2 のコピーエラーがある．比較エラーは，利得の高い相手の選択を誤ることである．対戦した相手のうちランダムな相手の戦略を模倣する．コピーエラーは戦略のコピー時に生じるエラーであり，模倣した戦略に平均 0，標準偏差 0.4 のガウスノイズが加わる．シ

表 2 シミュレーションプロセス

PD ゲーム	対戦	エージェントは対戦可能な相手と，それぞれ 4 回ずつ対戦を行う．利得は各エージェントの総利得を対戦回数で除算し，1 対戦あたりの利得を算出する．算出した値を各エージェントの利得とする．
	模倣	対戦可能なエージェントで最も高い利得を獲得したエージェントの戦略をコピーする．
進化	突然変異	比較エラー（利得の高い相手の選択を誤る）とコピーエラー（ A_p, A_q に，平均 0，標準偏差 0.4 のガウスノイズが加わる）

ミュレーションの手順を表 2 にまとめる．

3 ネットワーク構造が協調の進化に与える影響

本節では，2DK と SF において協調の達成に差があるのかを議論するために，それぞれのネットワークにおいてシミュレーション実験をおこなう．

3.1 協調の達成と崩壊

本シミュレーションでは，ネットワークのエージェント数 $N = 256$ とする．また，2DK の接続回数 $k = 4$ とする．SF は，Barabasi and Albert の提案した”成長と優先的選択アルゴリズム (BA)” [3] を用いて構築する．本モデルでは，BA によってノード数を 256 まで成長させたネットワークを用いる．突然変異については ϵ_1 及び ϵ_2 を共に 0.1 とする．利得パラメータを $(T, R, P, S) = (5, 3, 1, 0)$ としている．エージェントの戦略 (A_p, A_q) の初期値は $[0, 1]^2$ の乱数である．観察指標は，Payoff, Retention と FracLow である．それぞれの意味を表 3 に示す．基本モデルとして，ネットワーク構造を持たず，対戦相手がランダムに選択される IPD ゲームを考える．このモデルにおける集団の平均利得は 1.7 となり，最大利得は 2.3 を超えることがない [11]．よって，協調が達成される利得の閾値

表 3 観察指標

Payoff	全エージェントの平均利得． 協調の達成を表す．
Retention	シミュレーション期間中に 利得が閾値 $TH(2.3)$ を超えた 確率．協調の達成を表す．
FracLow	シミュレーション期間中に 利得が閾値 $TL(1.7)$ を下回っ た回数．協調の崩壊数を表す． (100 回試行における累積値)

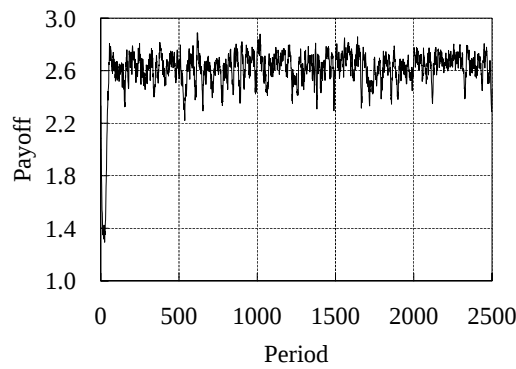


図 1 平均利得：2DK

表 4 実験結果

ネット ワーク	Payoff (S.D)	Retention (S.D)	FracLow
2DK	2.61(0.008)	0.98(0.005)	0
SF	2.52(0.016)	0.93(0.018)	20
t 値	48.9***	28.6***	-
p 値	0.000	0.000	-

*** $p < .001$

$TH = 2.3$ とし、協調が崩壊する閾値を $TL = 1.7$ とした。

各々のネットワークにおいて 100 回のシミュレーションを実行した結果を表 4 に示す。1 回のシミュレーションにおける期間は 2500 とする。

2DK と SF の Payoff および Retention に対して、平均値の差の検定をおこなった (表 4)。表 4 から、2DK は、SF よりも高い Payoff を持つ。Retention に関しては、2DK が SF より高い安定性を保持している。また、FracLow を見てみると、2DK では一度も協調の崩壊が生じていないが、SF では 20 回の崩壊が観察される。これらの結果は、2DK は協調が安定で、SF は協調が脆弱であることを示している。

図 1 及び図 2 は、2DK においても SF においても、ごく初期 (50 期程度) に非協調が支配的になるが、その後、協調が達成され、高い水準で協調が達成されることを示している。これは、以下のようなメカニズムによって説明できる。初期のランダムな状態では、All-C を搾取できる All-D が一番有利であるため、All-D が支配的となる。All-D が支配的な環境下で、突然変異

により TFT エージェントが発生するときを考える。あるエージェントの対戦相手のうち、TFT エージェントの割合を α 、All-D エージェントの割合を $1 - \alpha$ とする。このとき、TFT エージェントの利得 $E(T)$ は $E(T) = \alpha 4R + (1 - \alpha)(S + 3P)$ となり、All-D エージェントの利得 $E(D)$ は $E(D) = \alpha(T + 3P) + (1 - \alpha)4P$ となる。本研究では、 $(T, R, P, S) = (5, 3, 1, 0)$ であるので、 $E(T) - E(D) > 0$ となる条件は、 $\alpha > 0.2$ である。すなわち、All-D が支配的な環境下で、突然変異で生じる TFT エージェント同士が上記の条件を満たすようにお互いに隣り合うことができれば、TFT 戦略は All-D より高い利得を得ることができる。この結果、TFT のコロニーが発展することで TFT が支配的になり協調の安定が発現する。

2DK においては、協調が継続されるため、Payoff が TH を上回る確率が高い。このため Retention が高く FracLow が生じていない。SF の Payoff は、しばしば閾値 TH を下回り、時々 TL を下回る。すなわち、協調は達成されているが、不安定であり脆弱であるといえる。

X 軸、Y 軸が A_p, A_q である 2DK を図 3 に示し、同様に SF を図 4 に示す。矢印は、 A_p 及び A_q 値の母集団平均値の変化の方向を示す。円の大きさは、母集団がその領域の中に残る期間数を示す。各々の図で最も大きな円は、母集団が最も多くの時間を過したセルと一致する。他の円の領域は、そのセルと比較したスケールとなる。シミュレーションの初期において、両方のグラフ共に、 $A_p - A_q$ 空間の中央の領域から

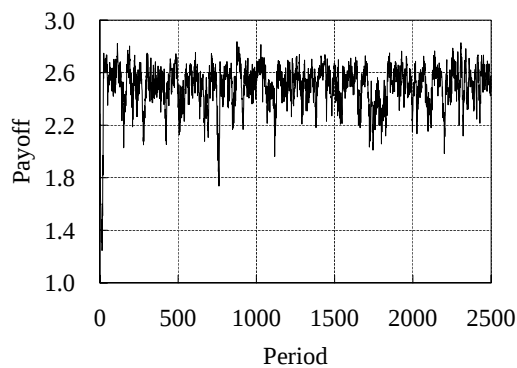


図 2 平均利得 : SF

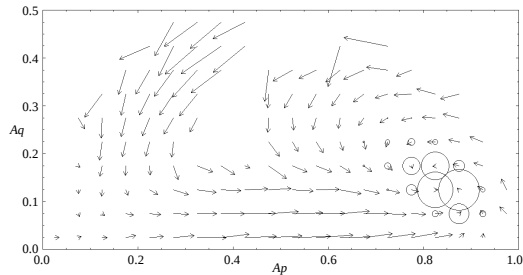


図 3 戦略の推移 : 2DK

All-D 領域 ($A_p = 0, A_q = 0$) に向かう。しかしその矢印はすぐに TFT 領域 ($A_p = 1, A_q = 0$) へ進む。これらの推移は、協調の出現を示す。2DK においては、協調の達成後 A_p は 0.5 以上の領域で推移している。しかし、SF においては、運動の軌跡は、 A_p が 0.5 を下回る領域を含んでいる。この推移は、協調の一時的な崩壊を示す。2DK と SF は、共に協調を達成することが出来る。しかしながら、両者を比較すると、SF は、脆弱な協調であり、崩壊する傾向がある。

3.2 SF における協調の脆弱性

シミュレーション結果から、SF では協調の達成は可能だが協調は脆弱であり崩壊が生じやすいことがわかった。なぜこのような脆弱性が生じるのであろうか？どのエージェントが協調の崩壊をもたらすのだろうか？エージェントの特徴を示すために、本論文では、接続次数が最も高い 5 つのエージェントをハブエージェント (Hub) とする。接続次数が最も低い 5 つの

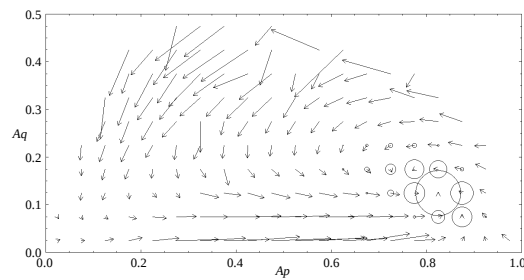


図 4 戦略の推移 : SF

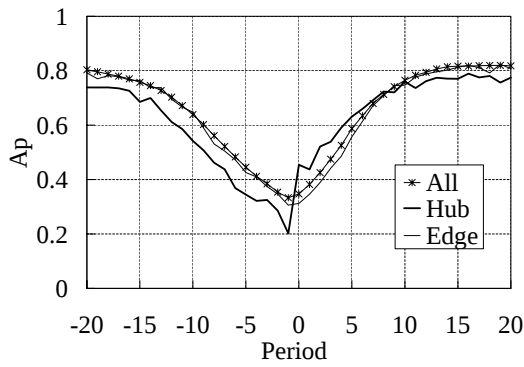
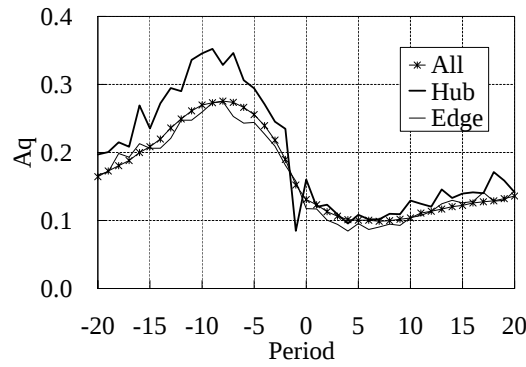
エージェントを末端エージェント (Edge) とよぶ。

まず、協調が崩壊するときに、どのような現象が生じているのかを明らかにする。100 回の協調の崩壊を抽出するために、SF のシミュレーション実験を 617 回おこなった。協調の崩壊とは利得が閾値 $TL = 1.7$ を下回ることをいう。協調が崩壊した期のうち平均利得が最小の期を $Period = 0$ とし、その前後の期において戦略の挙動を観察した。

図 5, 6 で示すように、 $Period = 0$ 前後のエージェントの戦略を観察する。図の横軸はシミュレーション期間で、縦軸は A_p, A_q の 100 回試行の平均である。図 5 において、 $Period = 0$ の周辺で A_p がもっとも低くなっている。これは非協調行動が増加し、協調が崩壊していることを示している。一方図 6 において、崩壊直前のハブエージェントの戦略は特徴的であり、全体の平均や末端エージェントに比べて A_q が高くなっている。

つまり、協調の崩壊はハブエージェントが寛容になることで発生すると考えられる。我々は、 A_q が高いことを寛容な戦略と呼ぶ。 A_q は、対戦相手が前回の対戦で非協調をとったとき、エージェントが今回の対戦で協調する確率である。つまり、非協調に対する寛容性を表す。寛容とは非協調に対して懲罰をおこなわないことを意味する。一方、エージェントの戦略が TFT ($(A_p, A_q) = (1, 0)$ に近いこと) を厳格なエージェントと呼ぶ。厳格とは、協調に対しては協調で応じるが、非協調に対しては非協調をとることを意味する。

協調の崩壊時にハブエージェントの A_q はなぜ上昇するのか。協調が達成されているときは、集団の平

図5 エージェントの特徴と協調の崩壊時の A_p 図6 エージェントの特徴と協調の崩壊時の A_q

均 A_p は 1 に近く, A_q は 0 に近い. しかし, 突然変異の影響により, A_p, A_q には常に 0.5 に近づく力が働いている. 1 に近づいていた A_p が 0.5 に近づこうとすると, 周囲は TFT 戦略になっているため, 懲罰を受け, A_p は 1 に留まる. 0 に近づいていた A_q が 0.5 に近づこうとすると, そのエージェントは裏切りの侵入を許すので淘汰される. ただし, 侵入した裏切りは, 周囲の TFT に懲罰を受け, 結果として集団の協調は維持される.

しかし, ハブエージェントの A_q が上昇すると, 裏切りへの誘因を受ける周辺エージェントが多数存在するため, 同時に大量の裏切りが発生し, TFT の懲罰が追いつかない. よって, ハブエージェントの A_q が上昇することを発端に, 協調が崩壊していることが考えられる.

協調が崩壊した状態は, 侵入した裏切りが A_q の高いハブエージェントを搾取し裏切りが優位となっている状態である. このように裏切りが優位な環境下でも, 突然変異で生じる TFT が隣り合うことや, 裏切りが浸透していない TFT エージェントの小集団が存在することで, 3.1 節で論じたように, TFT が All-D に対して優位となるので, 協調は再度達成される.

4 協調の脆弱性と頑健性

3 節の結果から, ハブエージェントが寛容になることが協調の崩壊を招いていることがわかった. では協調が頑健に維持されるためには, どのエージェントが, どのような戦略を採用するのがよいのであろう

か. また, どのような戦略が協調を最も脆弱にするのであろうか. 我々は, ハブエージェントに着目し, 寛容性 (A_q) の影響と協調性 (A_p) の影響を明らかにする.

4.1 寛容性の影響

我々は, 寛容性の影響を調査するために, エージェントの戦略のうち寛容性を表現する A_q に対して以下 4 点の操作を行う. ハブエージェントを寛容に固定する ($A_q^{Hub} = 1$), 末端エージェントを寛容に固定する ($A_q^{Edge} = 1$), ハブエージェントを厳格に固定する ($A_q^{Hub} = 0$), そして末端エージェントを厳格に固定する ($A_q^{Edge} = 0$) である.

それぞれの条件における利得について検討した (表 5, 図 7). 利得は, 各条件でシミュレーション実験を 100 回おこなった平均値である. 戦略の操作を独立変数に二要因の分散分析をおこなったところ主効果が有意であった (表 6). このことから, $A_q^{Hub} = 1$ になることでもっとも集団全体の利得が低くなることがわかる. これはハブエージェントの寛容性が協調の崩壊を招くことを示している. 更に, $A_q^{Hub} = 0$ で最も利得が高くなり, 協調が安定的に維持されていることがわかる.

続いて, 協調の進化と崩壊のダイナミクスを観察する. ここでは特徴的な操作である, $A_q^{Hub} = 1$ と $A_q^{Hub} = 0$ に関して議論する.

図 8, 図 9 は, 戦略を $A_q^{Hub} = 1$ と $A_q^{Hub} = 0$ に操作したときの, 集団全体の利得の変化を示してい

表 5 各操作における平均利得及び標準偏差

	Edge		Hub		Total	
	平均	S.D.	平均	S.D.	平均	S.D.
$Aq = 0$	2.54	0.02	2.61	0.01	2.57	0.01
	N	100	N	100	N	200
	平均	2.43	平均	2.00	平均	2.22
$Aq = 1$	S.D.	0.02	S.D.	0.18	S.D.	0.10
	N	100	N	100	N	200
	平均	2.49	平均	2.31	平均	2.40
Total	S.D.	0.02	S.D.	0.09	S.D.	0.06
	N	200	N	200	N	400

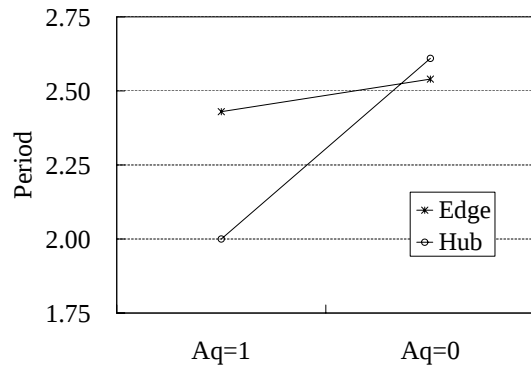
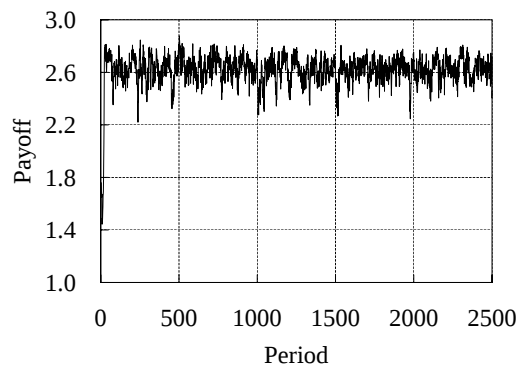
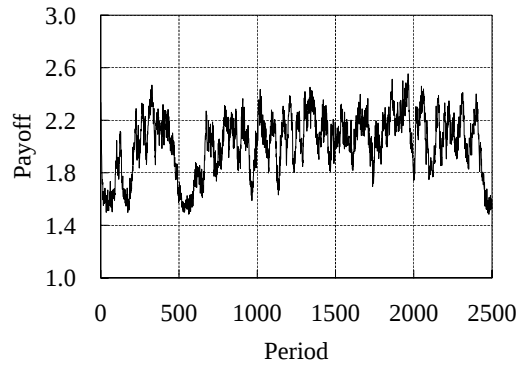


図 7 各操作における平均利得

表 6 利得の分散分析

要因	残差平方和	自由度	平均平方	F 値	P 値
Hub Edge	3.21	1	3.21	401**	0.00
$Aq = 0$ $Aq = 1$	12.77	1	12.77	1595**	0.00
交互作用	6.33	1	6.33	791**	0.00
誤差	3.21	396	0.01		
全体	3.21	399			

** $p < .01$ 図 8 平均利得: $Aq^{Hub} = 0$ 図 9 平均利得: $Aq^{Hub} = 1$

る。図 10, 図 11 は X 軸, Y 軸が A_p, A_q で, それぞれの戦略の変化を示している。矢印は, A_p 及び A_q 値の母集団平均値の変化の方向を示す。円の大きさは, 母集団がその領域の中に残る期間数を示す。各々の図で最も大きな円は, 母集団が最も多くの時間を過したセルと一致する。他の円の領域は, そのセルと比較したスケールとなる。図 8 ($Aq^{Hub} = 0$) にて, 縦軸が利得, 横軸がシミュレーション期間を示す。図 9 では, 同じく $Aq^{Hub} = 1$ を示す。 $Aq^{Hub} = 0$ の場合,

戦略の推移は短期的に TFT 領域を周回する (図 10)。対照的に, $Aq^{Hub} = 1$ は, TFT と ALL-D 戦略の間を往復する (図 11)。このダイナミクスは, 協調がしばしば崩壊することを意味する。図 8~11 をまとめると, 協調の崩壊は寛容なハブエージェントに起因し, 厳格なハブエージェントは協調を継続させる。

4.2 協調性の影響

我々は, 協調性の影響を調査するために, ハブエージェントの戦略のうち A_p を操作する。 A_p は, 対戦相手が前回の対戦で協調をとったときエージェントが今回の対戦で協調する確率, および, 第 1 回目の対戦で協調する確率である。つまり, A_p はエージェントの協調性を表す。ここでは, ハブエージェントの A_p に対して, 常に非協調 ($A_p^{Hub} = 0$) から常に協調 ($A_p^{Hub} = 1$) まで刻み幅 0.05 で連続に変化させ,

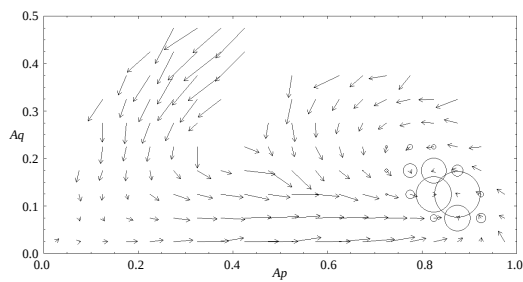


図 10 戦略の推移: $A_q^{Hub} = 0$

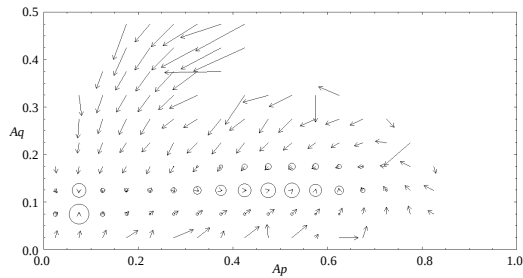


図 11 戦略の推移: $A_q^{Hub} = 1$

集団全体の平均利得を観察する．

図 12 は、横軸がハブエージェントの A_p であり、縦軸が集団全体の利得である．ハブエージェントが協調的になるに従い、集団全体の利得が上昇する ($A_p < 0.85$)．しかし、ハブエージェントの $A_p = 0.95$ をピークとして A_p の増加に対して利得が下降している．

自分からは裏切らない $A_p^{Hub} = 1$ のときより、少ない確率だが自ら非協調を取ることもあるほうが集団全体の利得が高くなっている．このメカニズムは、集団全体の寛容度 (A_q) を観察することで説明できる．図 13 は、ハブエージェントの戦略 A_p を操作したときの、集団全体の寛容度 (A_q) の変化を示している．ハブエージェントの戦略 $A_p^{Hub} > 0.8$ の時、ハブエージェントの戦略 A_p の増加に対して、集団全体の寛容度 (A_q) が急に上昇している．なぜなら、ハブエージェントが自ら裏切らない常に協調戦略に近づくと、集団全体から非協調行動が排除されていく．このとき、懲罰をおこなう必要がなくなるので、 A_q が高いという戦略は、 $A_q = 0$ である TFT 戦略と判

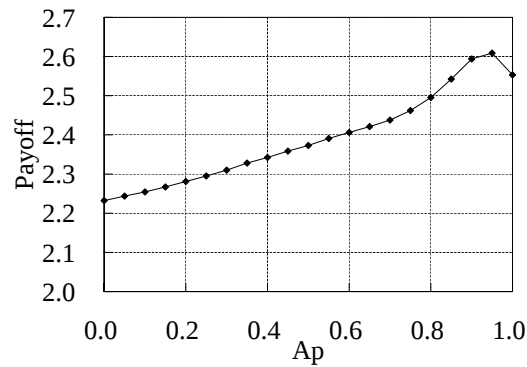


図 12 平均利得: $A_q^{Hub} = 0 \sim 1$

別がつかない．よって、集団全体の A_q は上昇する．

次に $A_p^{Hub} = 0.95$ と $A_p^{Hub} = 1$ のときを比較する． $A_p^{Hub} = 0.95$ では、少ない確率ではあるが、ハブエージェントが非協調を選択するので、多数のエージェントが懲罰をおこなう必要が生じる．よって集団全体で A_q の上昇が抑制され、非協調戦略の侵入を防ぐことができる．一方、 $A_p^{Hub} = 1$ では、侵入した非協調戦略を持つ末端エージェントとリンクする少数のエージェントしか、懲罰をおこなわない．なぜなら、突然変異による非協調戦略は、確率的に集団の大多数を占める末端エージェントに発生するからである．結果として集団全体の A_q が 0 に近づく誘因は小さい．このことが頻繁な非協調戦略の侵入を許すこととなり、結果的に集団全体の利得は減少する．

このように、自分からは裏切らない $A_p^{Hub} = 1$ のときより、少ない確率だが自ら非協調を取ることもあるほうが集団全体の利得が高くなる効果を我々は「引締め効果」と表現する．また、2DK においても引締め効果は確認できたが、その効果は微少であった．

5 結論

我々は、相互作用関係がスケールフリーネットワークである囚人のジレンマにおいて協調は安定的に存続可能なのだろうか、という課題に答えるためにシミュレーション実験を行い、協調が進化的に安定となる条件を探索した．シミュレーション結果から、スケールフリーネットワークにおいて、協調は達成されるが脆弱であり崩壊しやすいことがわかった．協調の脆弱性

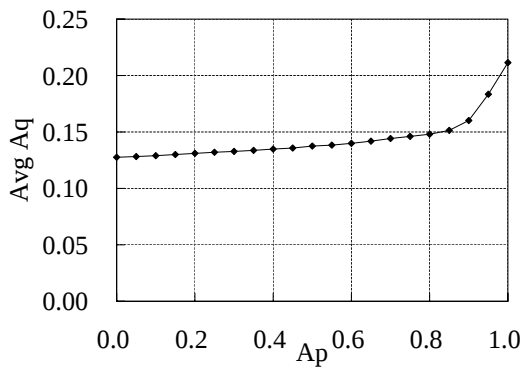


図 13 集団の平均 Aq : $Aq^{Hub} = 0 \sim 1$

は、ネットワーク内で接続次数が高いハブエージェントが非協調に対して寛容になることで生じることがわかった。ハブエージェントは、協調の安定性に3つの重要な影響を及ぼす。第一に、ハブエージェントが寛容になることで協調は崩壊する。第二に、ハブエージェントが厳格であることで協調を最も安定的に維持可能となる。第三に、ハブエージェントがごく小さい確率で採用する非協調戦略は、集団全体の寛容性を抑制する引締め効果をもたらす、全体としてより高い平均利得を獲得可能にする。

5.1 本研究の位置付け

ネットワーク構造の違いによる協調の進化を扱った研究は多くなされている。

Santos and Pacheco は、レギュラーネットワークやランダムネットワークと比較して、スケールフリーネットワークの方が、利得行列の値やネットワークサイズ、ネットワーク密度に関わりなく、協調が安定することを示した[12]。Wu et.al は、Santos and Pacheco のモデルが様々な学習パターンでも成立すると述べている[15]。Abramson and Kuperman は、ネットワーク密度とトポロジーが、非協調の存在割合に与える影響について調査した[1]。Masuda and Aihara は、スモールワールドネットワークが最も協調の達成速度が速いことを示した[8]。これらの研究は、ネットワーク構造の特徴に着目するため、エージェントの戦略を、純粋な協調戦略 (All-C) と純粋な裏切り戦略 (All-D) に限定し、それぞれの戦略の集団

全体における存在割合を観察している。

しかし、このような観察指標では、ネットワーク構造の違いによる集団全体の協調の達成度合いの違いは議論できるが、エージェント間のどのような局所的なダイナミクスが、協調の進化をもたらすのかを議論することができない。このダイナミクスを詳細に議論するためには、個々のエージェントの戦略を多様に表現する必要がある。そのためにエージェントの戦略は All-C と All-D に限定せず、しつぺ返しとよばれる TFT 戦略を含んだ連続的な戦略空間を表現しなければならない[4]。

そこで我々は、個々のエージェント内部に連続的な戦略空間を設定した。その結果、ネットワーク構造の違いがもたらす協調の進化と崩壊を、個々のエージェントの相互作用のメカニズムから説明することができた。さらに、どのエージェントのどのような戦略が協調の進化と崩壊に影響を与えるのかを明らかにすることができた。

5.2 協調の進化と崩壊

2次元格子トラスにおいてもスケールフリーネットワークにおいても、初期のランダムな状態から一度 All-D が支配的になるが、その後、TFT が支配的となり協調が集合的な秩序として成立する。これは次のメカニズムで説明可能である。初期状態のランダムな状態では、All-C を採取できる All-D が一番有利であるため、All-D が支配的となる。しかし、All-D が支配的な環境下で、突然変異で生じる TFT エージェント同士がお互いに隣り合うことで、TFT 戦略は All-D より高い利得を得ることができ、TFT のコロニーが発展する。この結果、TFT が支配的になり協調の安定が発現する。TFT が支配的な環境下では、All-C と TFT が対戦しても、常に協調行動をとることになり All-C は排除されない。2次元格子トラスでは、この状況で発生する All-D は、All-C を採取することができる。しかし、All-D は周囲のその後多数を占める TFT に懲罰されて結果的に TFT で安定する。その結果2次元格子トラスでは、TFT が維持される。しかしスケールフリーネットワークでは、ハブが All-C になると、周囲に多数の All-D を発生さ

せることになる．その結果，All-D が多数派になってしまい，TFT の懲罰が及ばない．これが協調の崩壊につながっている．

5.3 スケールフリーネットワークにおける協調の脆弱性

本論文で設定した囚人のジレンマにおいて，スケールフリーネットワークにおいてはハブエージェントが寛容だと協調が崩壊し，厳格だと協調が安定することがわかった．集団において TFT 戦略が支配的になった状態では，ハブエージェントが非協調的になること，すなわち A_p が低くなることは，周囲のエージェントによる懲罰が働くため，協調の崩壊は招かない．しかし，ハブエージェントの寛容性 (A_q) が高くなることは，TFT 戦略に対して中立的なミュータントであるため，排除機能が働かない．ハブエージェントが寛容になることで，多数のエージェントがハブエージェントを搾取することが可能となる．このため，集団全体で非協調戦略が増加することとなる．これがスケールフリーネットワークで協調が脆弱であるメカニズムである．ハブエージェントの寛容性を社会的文脈で解釈すると，ハブエージェントは社会における中心人物であると考えられる．現実社会は中心人物が寛容になることに対して脆弱性を持つことを示唆している．本論文は相対取引における囚人のジレンマを想定している．これは，例えば消費者間オンライン取引における売り手と買い手の行動などが該当する[16]．市場において積極的な参加者が評価の局面で寛容になることは市場の協調的な規範を脆弱にする危険性があることを示している．

5.4 引締め効果

以上の考察から，ハブエージェントが厳格である必要性はシミュレーション結果によって示されていると思われる．ハブエージェントの戦略で興味深い効果は，引締め効果である．直観的に，自らは裏切る確率が低いほうが集団全体の協調性は高くなると考えられる．しかし，図 12 が示すように，ハブエージェントの戦略は，自ら裏切ることがない $A_p^{Hub} = 1$ であることより，時折裏切りを選択する $A_p^{Hub} = 0.95$ 付

近のほうが集団全体の利得は高くなっている．

なぜなら，ハブエージェントがわずかな確率で自ら裏切る戦略であるとき ($A_p^{Hub} = 0.95$)，もし，周囲の多数のエージェントの戦略が All-C であると搾取されることがあるが，TFT であれば裏切りに対して懲罰を行うので，All-C より優位となる．よって，多数のエージェントの A_q は低く抑えられ，集団全体の戦略は TFT が支配的となり，突然変異による裏切りの侵入に対しても頑健となる．しかし，ハブエージェントが自らは裏切らない戦略であると ($A_p^{Hub} = 1$)，周囲の多数のエージェントにとって，懲罰を行う必要が生じるのは，突然変異による裏切りに対してのみとなる．突然変異は，多数存在する末端エージェントに発生しやすいので，懲罰を行わなければならないエージェントは末端エージェントとリンクする少数のエージェントに留まる．この結果，大多数のエージェントは， A_q を 0 にとどめておく誘因を持たないので，集団の A_q は上昇することになる．そのため，集団全体に TFT である誘因が働く $A_p^{Hub} = 0.95$ の時と比べ，頻繁に非協調戦略が侵入してしまい，集団全体の利得を下げることになる．

このように自らは決して裏切らない戦略より，少ない確率だが自ら非協調を取るものがあるほうが集団全体の利得が高くなる効果を，我々は「引締め効果」と表現する．

5.5 今後の課題

本研究では，スケールフリーネットワークにおける協調の頑健性と脆弱性を議論した．しかし，我々の実験結果をもたらすネットワーク構造は，スケールフリーネットワーク以外にも存在しうると考えられる．今後，本研究の知見がスケールフリー性に起因するかどうか，またネットワークのノード総数が及ぼす影響を明らかにする必要があるだろう．林は様々なネットワークの構造を分析し，口コミや連鎖倒産といった社会システムと，インターネットや分子モータなどの技術・生物システムでは，同じスケールフリーネットワークであっても結合のパターンが異なることを指摘している[10]．社会システムは，ハブとハブが直接リンクを貼りやすい Assortative 結合という特性を

持ち, 技術・生物システムは, Disassortative 結合という特性を持つ [9]. 今後の課題は, 本研究の知見をオンライン消費者間取引や組織内のコミュニケーションの分析に応用することである. その際には, 社会システムにおけるスケールフリーネットワークの特性である Assortative 結合を考慮することで, より精緻なモデルを構築ができると考えられる. また, 引締め効果は突然変異により非協調戦略が侵入する程度と, ハブエージェントの接続次数といったネットワーク構造に依存すると考えられる. これらの特徴に基づいて適切な引締め効果をもたらすハブエージェントの戦略を定式化することで, より詳細な戦略の分析が可能になると考えられる.

参 考 文 献

- [1] Abramson, G. and Kuperman, M.: Social games in a social network, *Phys. Rev. E* 63(2001), 030901.
- [2] Axelrod, R.: *The Evolution of Cooperation*, New York, Basic Books, 1984.
- [3] Barabasi, A. L. and Albert, R.: Emergence of scaling in random networks, *Science*, 286(1999), pp.509-512.
- [4] Cohen, M. D., Riolo, R. L. and Axelrod, R.: The Role of Social Structure in the Maintenance of Cooperative Regimes, *Rationality and Society*, Vol. 13, No. 1(2001), pp.5-32.
- [5] Ebel, H., Mielsch, L.I. and Bornholdt, S.: Scale-free topology of e-mail networks, *Phys. Rev. E* 66(2002), 035103-1-4.
- [6] Ifti, M., Killingback, T. and Doebeli, M.: Effects of neighbourhoodsize and connectivity on the spatial Continuous Prisoner's Dilemma, *Journal of Theoretical Biology* 231(2004), pp.97-106.
- [7] Jeong, H., Mason, S., Barabasi, A. L., and Oltvai, Z. N.: Lethality and centrality in protein networks, *Nature* 411(2001), pp.41-42.
- [8] Masuda, N. and Aihara, K.: Spatial prisoner's dilemma optimally played in small-world networks, *Physics Letters A* 313(2003), pp.55-61.
- [9] Newman, M.E.: Assortative mixing in networks, *Phys Rev Lett* 89 (2001):208701.
- [10] 林 幸雄, 招待講演: Scale-free ネットワークの研究動向, 情報処理学会 知能と複雑系研究会 (SIG-ICS) 人工知能学会 知識ベースシステム研究会 (SIG-KBS) 共催, 「ネットワークが創発する知能」, ICS-136(2004), pp.1-8,
- [11] Riolo, R.: The Effects of Tag-Mediated Selection of Partners in Evolving Populations Playing the Iterated Prisoner's Dilemma, Santa Fe Working Paper 97-02-016(1997). Santa Fe, NM.
- [12] Santos, F. C., and Pacheco, J. M.: A new route to the evolution of cooperation, *J. of Evolutionary Biology*, 19(3), pp.726-732.
- [13] Vukov, J. and Szabo, G.: Evolutionary prisoner's dilemma game on hierarchical lattices, *Phys. Rev. E* 71(2005), 036133.
- [14] Watts, D. J. and Strogatz, S. H.: Collective dynamics of 'small-world' networks. *Nature* 393(1998), pp.440-442.
- [15] Wu, Zhi-Xi, Xu, Xin-Jian, and Wang, Ying-Hai: Does the scale-free topology favor the emergence of cooperation?, [oai:arXiv.org:physics/0508220](https://arxiv.org/abs/physics/0508220)
- [16] Yamamoto, H., Ishida, K. and Ohta, T.: Modeling Reputation Management System on Online C2C Market, *Computational & Mathematical Organization Theory*, Vol. 10, No. 2(2004), pp.165-178.