

状態空間の部分的次元化手法を用いたマルチエージェント強化学習法

藤田 和幸[†] 松尾 啓志[†]

[†] 名古屋工業大学電気情報工学科, 名古屋市

E-mail: †matsuo@nitech.ac.jp

あらまし マルチエージェント強化学習では、エージェントは自分以外のエージェントも環境の一部として観測する。そのため、エージェント数の増加に伴い状態空間が指数的に増加し（次元の呪い）、学習速度を著しく低下させるという問題が生じる。この問題を解決する手法として提案された Modular Q-learning には、知覚の不完全性により学習性能が低下する問題がある。本研究では状態空間の部分的次元化手法を用いて Modular Q-learning の学習性能を改善する手法を提案する。

キーワード マルチエージェント, 強化学習, Modular Q-learning, 状態空間

Multi-Agent Reinforcement Learning with the Partly High-dimensional State Space

Kazuyuki FUJITA[†] and Hiroshi MATSUO[†]

[†] Department of Computer Engineering, Nagoya Institute of Technology, Gokiso-cho, Showa-ku, Nagoya,
466-8555, Japan

E-mail: †matsuo@nitech.ac.jp

Abstract In Multi-Agent Reinforcement Learning, each agent observe a state of other agents as a part of environment. Therefore, the state space is exponential in the number of agents and learning speed significantly decrease. Modular Q-learning [6] needs very small state space. However, the incomplete observation involves a decline in the performance. In this paper, we improve Modular Q-learning's performance with the partly high-dimensional state space.

Key words Multi Agent, Reinforcement Learning, Modular Q-learning, State Space

1. はじめに

近年、分散協調システムの一つとしてマルチエージェントシステムが注目されている。マルチエージェントシステムは、一般に中央制御を必要としないため複雑なタスクを行う大規模なシステムの制御への応用が期待されている。マルチエージェントシステムを構成するエージェントの設計方法の一つに強化学習 [1] がある。強化学習は教師なし学習の一種であり、システムの設計者が環境に応じた設計を施す必要がないため、大規模で複雑なシステムを設計する場合などに非常に有効な手法であるといえる。しかし、マルチエージェント強化学習には、荒井 [2] が示しているように不完全知覚問題（一般にマルチエージェントシステムは大規模であり、系全体を観測することは困難である。そのため、エージェントの誤った状態認識を招く）、同時学習問題（マルチエージェント環境において、エージェントは他のエージェントも環境の一部として観測する。そのため、他

のエージェントの政策が変化することで環境の状態遷移確率が変化し、MDP 環境としてモデル化することができない）、報酬配分問題（個々のエージェントの政策を評価し、それに見合った報酬を与えることは困難である。そのため不適切な政策が強化学習されてしまう可能性がある）、次元の呪い（エージェント数が増加すると、状態空間が指数的に増大する。そのため学習速度が著しく低下する）などの問題点がある。

本研究では、次元の呪いについて議論する。シングルエージェント環境においても次元の呪いは存在する。実環境に強化学習を適用する場合、とりわけロボット制御などでは行動出力が連続値であり、状態空間が連続値で与えられる。そのため行動空間・状態空間が爆発する（次元の呪い）。次元の呪いが生じる環境では、エージェントが学習空間を探索するために多くの時間を必要とし、学習速度を著しく低下させてしまう。シングルエージェント環境では、連続的な行動空間に対して Actor-Critic 法が提案されており、連続的な状態空間に対しては、SVM(Support

Vector Machine) を用いた状態一般化法 [3], 状態空間の階層化 [4] などがある.

シングルエージェント環境においても, 強化学習を実環境に用いる場合には学習速度の低下が問題とされている. したがって, シングルエージェント環境下での学習よりも多くの学習時間を必要とするマルチエージェント強化学習では, 学習速度の低下は一層重要な問題といえる. マルチエージェント強化学習における次元の呪とは, エージェント数の増加に伴い状態空間が指数的に増大する問題である. 高橋ら [4] の手法は, 下位層で状態空間を幾つかの小さな領域に分割し, 各領域において独立して学習を行うことで膨大な状態空間に対応している. また, 高位層では下位層学習器間の移動を定義しており, 下位層において目標状態への到達が不可能な場合には, 高位層の学習器を用いて目標状態に近い下位層の学習器への移動を行うというものである. しかし, 状態遷移に他エージェントの行動という不確定な要素が含まれるマルチエージェント環境において, 抽象化された高次元空間における状態遷移を定義することは非常に困難である. また, Fuzzy システムを用いた状態数の削減手法 [5] があるが, これはエージェント数の増加に伴う状態空間の増大の根本的な解決にはならない.

マルチエージェント強化学習における次元の呪を解決する手法として Modular Q-learning がある. この手法は, 2 体エージェントから構成される部分状態を用いるため, 状態空間の大きさがエージェント数に影響されない. しかし, 部分状態のみを観測するため, 不完全知覚による学習性能の低下を招く. そこで本研究では, 学習性能低下の要因となる状態を高次元化することで不完全知覚を取り除き, Modular Q-learning の学習性能を改善する手法を提案する.

以下, 2 章では次元の呪と Modular Q-learning の問題点について述べる. 3 章では, 提案手法における高次元化手法について説明し, 4 章で追跡問題を用いて提案手法の性能評価を行う.

2. 次元の呪

エージェント k が観測する状態を S_k , エージェント i の状態を s_i , エージェント数を N としたとき, 状態 S_k は全エージェントの状態の集合 $\langle s_1, \dots, s_N \rangle$ から成る. したがって, 状態空間の大きさ $|S_k|$ は $|s|^N$ となり, エージェント数が増加すると $|S_k|$ は指数的に増大する. 強化学習において状態空間の増大は学習速度を急激に低下させ, また膨大なメモリ量を必要とさせる.

2.1 Modular Q-learning

Modular Q-learning は状態空間の爆発を防ぐ手法として Ono ら [6] によって提案された手法である. Modular Q-learning では, 自分と他の 1 体のエージェントから構成される部分状態空間を用いる. そのため状態空間の大きさは, エージェントの数に関わらず常に $|s|^2$ となり, 状態空間の爆発を防ぐことが可能となる.

エージェントが N 体存在する場合, 部分状態空間は $N-1$ 個存在し, それぞれに対して一つずつ学習器が割り当てられて

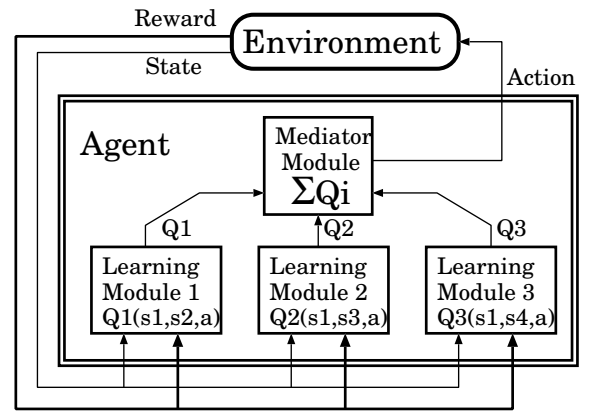


図 1 Modular Q-learning におけるエージェントの構成
Fig. 1 A structure of Agent in Modular Q-learning

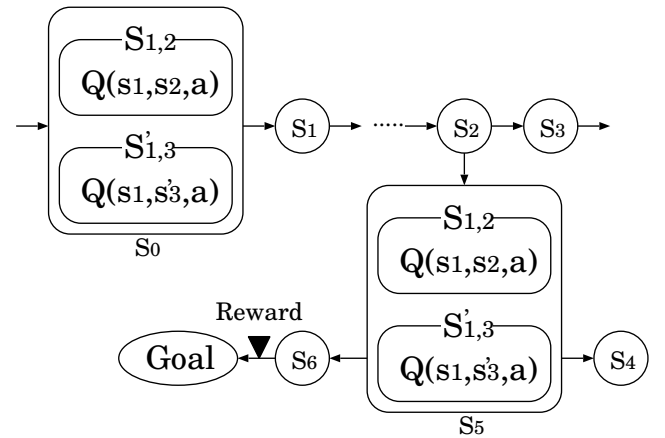


図 2 Modular Q-learning における学習性能低下の例: 状態 S_0 と状態 S_5 では同一の部分状態 $S_{1,2}$ を観測する
Fig. 2 An example that Modular Q-learning's performance down.

いる. 図 1 にエージェントが 4 体の場合のエージェント 1 の構成を示す. 各学習器は Q-learning を行い, Q 値を行動選択器へと渡す. 式 (1) に示すように行動選択器は受け取った Q 値の合計値がより大きな値となる行動を優先して選択する [7].

$$\arg \max_{a \in A} \sum_{i=1}^{N-1} Q^{k,i}(S_{k,i}, a) \quad (1)$$

しかし, 部分状態のみを観測することで不完全知覚を招き, 学習性能を低下させてしまうという欠点がある. 図 2 に性能低下の例を示す.

図 2 では, 状態 S_0 と状態 S_5 において同一の部分状態 $S_{1,2}$ を観測する. 状態 S_5 において最適な行動は状態 S_6 へ遷移する行動であり, その行動の Q 値 $Q(s_1, s_2, a)$ が状態 S_5 において最大となっていると仮定する. 状態 S_0 における $Q(s_1, s_2, a)$ は状態 S_5 における $Q(s_1, s_2, a)$ と同一のものであり, また状態 S_0 はゴール状態から遠いため, $Q(s_1, s'_3, a)$ は $Q(s_1, s_2, a)$ に比べて非常に小さい. したがって状態 S_0 においても, 状態 S_5 における最適行動と同じ行動を選択する可能性が高く, この例では同じ行動を選択した場合を考える.

状態 S_0 から遷移可能な状態は状態 S_0 と同様にゴールから遠く, Q 値も小さい. したがって, Q 学習において状態遷移先の価

値として用いられる $\max_a Q(s, a)$ の値は $Q(s_1, s_2, a)$ より小さい。このため、学習を行うことで $Q(s_1, s_2, a)$ の値は減少する。これは、状態 S_5 において最適な行動を選択する可能性が低下することを意味し、状態 S_6 ではなくゴール状態から遠い状態 S_4 へと遷移する可能性が高くなる。また、状態 S_2 から状態 S_5 へと遷移した場合、学習に用いられる $\max_a Q(s, a) = Q(s_1, s_2, a)$ の値が本来の値よりも小さいため、状態 S_5 へと遷移する行動の Q 値も本来の値よりも小さな値となり、状態 S_3 へと遷移する可能性が高くなる。状態 S_3 、状態 S_4 への遷移は最適な政策ではないため、学習性能が低下したといえる。

3. 提案手法

Modular Q-learning では部分状態のみを観測しているため、不完全知覚が生じ学習性能が低下した。そこで、図 2 の状態 S_5 のようにゴール状態に近く、 Q 値が低下すると学習性能に大きな影響を与える状態では、部分状態 $S_{1,2}, S_{1,3}$ を用いるのではなく、それらの状態を高次元化して得られる状態 $S_{1,2,3} = \{s_1, s_2, s_3\}$ を用いることで学習性能の低下を防ぐことができると考えられる。なぜなら、状態 $S_{1,2,3}$ を用いることは、状態 S_5 では $Q(s_1, s_2, a), Q(s_1, s_3, a)$ ではなく $Q(s_1, s_2, s_3, a)$ にしたがって行動を選択することを意味し、状態 S_0 からの遷移によって $Q(s_1, s_2, a)$ の値が低下しても、状態 S_5 における行動選択は影響を受けないためである。また、状態 S_2 から状態 S_5 へと遷移した場合も、 $\max_a Q(s, a) = Q(s_1, s_2, s_3, a)$ を用いることで、状態 S_5 へと遷移する行動の Q 値の低下を防ぐことができる。状態 S_0 を高次元化しても同様のことが言えるが、そもそもゴールから遠い状態は価値が低く、遷移する確率も小さい。したがって、そのような状態を高次元化しても、状態 S_5 を高次元化した場合ほど性能改善の効果は期待できない。また、多くの状態を高次元化することによって扱う状態数が増えすぎると、Modular Q-learning の利点である速い学習速度が失われてしまう。

そこで本手法では、図 2 における状態 S_5 のような状態のみを高次元化し、不完全知覚を取り除くことで Modular Q-learning の学習性能を改善する。

3.1 高次元化手法

高次元化を行うためには、まず現在の状態が図 2 の状態 S_5 のような状態にあるか否かを決定する必要がある。また上述したように、本手法では、可能な限り少ない高次元状態を用いることで、学習速度を損なわない効率的な性能改善を目的としている。したがって、状態 S_0 と状態 S_5 を識別する必要がある。しかし、部分状態を観測して得られる情報だけでは決定することはできない。そこで、部分状態の価値を表す状態価値関数 V_i を用いて、状態の識別を行う手法を提案する。自分の状態 s_{self} と他エージェント i の状態 s_i から構成される部分状態の価値 $V_i(s_{self}, s_i)$ の学習式を式 (2) で与える。

$$V_i^t(s_{self}^{t-1}, s_i^{t-1}) = V_i^{t-1}(s_{self}^{t-1}, s_i^{t-1}) + \beta(\max(r_t, \gamma V_i^t(s_{self}^t, s_i^t)) - V_i^{t-1}(s_{self}^{t-1}, s_i^{t-1}))(2)$$

状態価値関数 V_i の値は、報酬が与えられるか、より価値

の高い状態へと遷移した場合に増加する。したがって、ゴールに近い状態ほど価値 V の値が大きく、ゴールから遠い状態ほど価値 V の値は小さくなる。状態 S_5 はゴール状態に近い状態 $V(s_1, s_2), V(s_1, s_3)$ の値は大きくなる。一方、状態 S_0 はゴール状態から遠く $V(s_1, s_3')$ の値は小さい。状態 S_0 から状態 S_1 へと遷移することで $V(s_1, s_2)$ の値は減少するが、少なくとも $V(s_1, s_3')$ よりは大きくなる。そこで、 $V(s_1, s_2) > \eta$, $V(s_1, s_3) > \eta$, $V(s_1, s_3') < \eta$ となる閾値 η を設け、状態を構成する部分状態の価値が共に閾値を超える状態を状態 S_5 として認識し、状態 S_0 と区別することが可能となる。このようにして状態 S_5 にあると判定された状態を高次元化する。価値 V の値は学習し続けるため、学習途中で減少し $V(s_1, s_2) < \eta$ となる可能性もある。本手法では、一度高次元化された状態は性能改善に重要な状態である考え、以後 V が低下しても高次元状態として扱う。

閾値 η は式 (3) で与える。

$$\eta = \gamma^\lambda \cdot \text{Reward} \quad (3)$$

λ はパラメータ、 γ は状態価値関数 V の学習式に用いた割引率、 Reward は環境から与えられる報酬を表す。 λ の値を無限大に近づけると閾値 η の値は 0 となる。そのため、全ての状態が高次元化され、学習アルゴリズムは Q-learning と等しくなる。したがって、 λ の値を大きくすることで学習性能は大きく改善されるが、学習速度が低下する。また、 λ の値を非常に小さな値に近づけると閾値 η の値は無限大となる。そのため、高次元化される状態は存在せず、学習アルゴリズムは Modular Q-learning と等しくなる。したがって、 λ の値を小さくすることで学習速度は速くなるが、学習性能が低下する。

本手法は、この λ を適切な値に設定することで、十分な学習速度を保ったまま、Modular Q-learning の学習性能を改善することが可能である。

3.2 エージェントの構成

図 3 にエージェントの構成を示す。

• Check Module:

次元数に応じた学習器を選択するモジュール。各状態ベクトルの次元数を記憶しており、観測した状態ベクトル (s_1, s_2, s_3) が 2 次元状態として扱われるならば、2 次元学習器に部分状態 $\{s_1, s_2\}$ と $\{s_1, s_3\}$ を渡し、3 次元状態として扱われるならば、状態 $\{s_1, s_2, s_3\}$ を 3 次元学習器に渡す。また、2 次元部分状態を State-Value Module に渡す。

• State-Value Module:

高次元化の判定と高次元状態の生成を行うモジュール。Check Module から 2 次元部分状態を受け取り、閾値 η を用いて高次元化の判定を行う。また、高次元化の条件を満たしているならば高次元状態を生成する。

• N 次元学習器 (Learning Module)

N 次元状態の学習を行うモジュール。学習には Q-learning を用いる。学習を行うのは Check Module によって選択された学習器のみで、選択された学習器は行動選択に必要となる Q 値を行動選択器に渡す。また、報酬も選択された学習器のみに与え

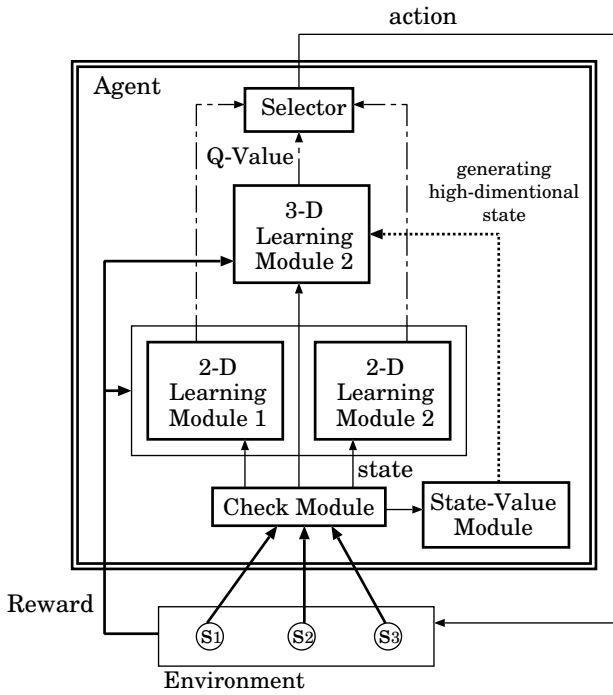


図 3 エージェントの構成
Fig. 3 A structure of Agent

られる。

- 行動選択器 (Selector)

各学習器から受け取った Q 値を用いて行動選択を行うモジュール。受け取った Q 値を行動ごとに合計し、その合計値がより大きな値を持つ行動を優先した戦略を取る。

3.3 学習アルゴリズム

本手法の学習アルゴリズムを以下に示す。

- (1) 各学習モジュールが持つ Q 値の値を初期値に設定する。また、Check Module が記憶している全状態ベクトルの次元数を、本手法で扱う最小次元 (2 次元) に設定する。
- (2) 環境を観測する
- (3) 観測して得られた知覚をもとに、チェックモジュールを用いて高次元状態が存在するか調べ、現在の状況に対応した次元の学習モジュールを選択し状態を与える。
- (4) 選択された学習モジュールは Q 値を行動選択器に渡す。
- (5) Q 値を受け取った行動選択器は Q 値の合計値をもとに行動を選択する。
- (6) 環境を観測し、与えられた報酬に基づいて Q 値・状態価値を学習する。
- (7) 高次元化条件に基づいて状態の高次元化を行う。
- (8) ゴール状態なら終了。そうでなければ (3) へ。

4. 評価実験

4.1 追跡問題

提案手法の性能評価を行うために、追跡問題を用いる。追跡問題とは、ハンターが獲物を追跡し捕獲する問題である。本研究では複数のハンターと獲物が存在する環境において、ハンターをエージェントと見なし、エージェントに獲物捕獲行動を

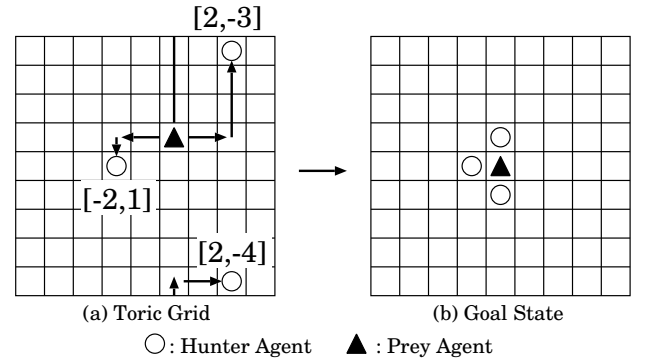


図 4 追跡問題: (a) トーラス状のグリッド, 矢印で示すように相対位置を計算する. (b) ゴール状態
Fig. 4 pursuit domain: (a) toric grid, The arrow explain relative position. (b) goal state

学習させる。以下に、本研究で用いた追跡問題の環境設定を示す (長行ら [8] が評価実験に用いた問題と同様の環境設定である)。

- 2次元トーラス状グリッド (5x5 もしくは 9x9) に、3体のハンターと1体の獲物が存在する。
- ハンターと獲物は、各時間ステップに同時に行動を決定する。また、ハンターが選択可能な行動は5通りあり、上、下、左、右の隣接するグリッドへの移動、現在の位置に停止である。獲物が選択可能な行動は3通りあり、上へ移動 (40%)、右へ移動 (40%)、現在の位置に停止 (20%) である。また、括弧内の数字は行動選択確率を表す。
- 獲物は学習を行わず、行動選択確率は一定とする。
- 獲物とハンターが同じグリッドに移動することを許可する。
- 初期状態から獲物を捕獲する (ゴール状態) までも1エピソードとする。獲物とハンターの初期配置はランダムに決定する。また、ゴール状態は4(b)に示すように、ハンターが異なる3方向から獲物に隣接した状態とする。
- エージェントの状態は、獲物との相対位置で表現する。例えば、図4に示すように、獲物との相対位置は $[-1, 1]$, $[1, -3]$, $[1, -4]$ のように表現する。
- 獲物捕獲した場合、エージェントに報酬 1.0 を与え、それ以外の場合には報酬 -0.05 を与える。

以上の環境で、獲物捕獲に要するステップ数・状態空間の大きさについて、Modular Q-learning, Q-learning との比較を行う。また、各手法における行動選択にはボルツマン選択を用いた。

4.2 実験結果

a) 5x5 グリッド環境

図5に、提案手法, Modular Q-learning, Q-learning を用いた実験結果を示す。横軸はエピソード数, 縦軸は獲物捕獲に要したステップ数を表している。提案手法と Modular Q-learning の結果は 100 試行, Q-learning の結果は 20 試行の平均値である。また、図5の値は 10 エピソード毎の移動平均である。各学習法におけるパラメータの値は、学習率 $= 0.3 \times 0.998849^{e_{pi}}$,

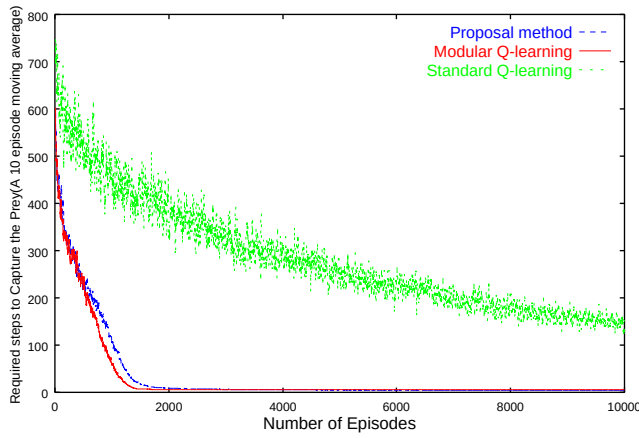


図 5 5x5 グリッド環境における実験結果
Fig. 5 An experimental result in 5x5 grids

割引率 = 0.9, 温度パラメータ = $0.1 \times 0.998849^{epi}$ とした。ここで, epi はエピソード数を表し, $0.998849^{10000} \approx 0.1$ となるように温度パラメータの減衰率を決定した。また, 閾値パラメータ λ の値は 4.0 とした。

Q-learning における状態空間の大きさは 15625 状態であり非常に大きい。そのため学習速度も遅く, 10000 エピソード学習を行っても獲物捕獲に 150 ステップ程必要としている。一方, 提案手法, Modular Q-learning における状態空間の大きさはそれぞれ 625 状態, 1100 状態であり, Q-learning に比べて非常に小さく, 図 5 が示すように学習速度も Q-learning に比べて非常に速い。

次に, 図 6 に提案手法と Modular Q-learning が獲物捕獲に要したステップ数を 0 ~ 20 ステップの範囲で示す。提案手法では, 高次元化手法によって扱う状態数が増えるため, Modular Q-learning よりも学習速度が若干低下している。しかし, Modular Q-learning のステップ数がほぼ収束したと思われる 3500 エピソード付近で, Modular Q-learning と提案手法のステップ数が同じ値となっている。さらにその後のエピソードでは, 提案手法のステップ数が Modular Q-learning のステップ数よりも小さな値となっており, 試行の終了時では Modular Q-learning のステップ数を約 17%改善できている。このことから, 高次元状態を用いて不完全知覚を取り除くことにより, Modular Q-learning の学習性能を改善可能であることが確認できる。

b) 9x9 グリッド環境

グリッドの大きさを 9x9 として実験を行った。9x9 グリッドでは Q-learning における状態空間の大きさは 531441 状態と非常に大きい。また, 5x5 グリッド環境の実験において明らかに Q-learning の学習速度は遅く, 9x9 グリッド環境においても同様の結果になると考えられる。また, 学習に要するメモリ量も膨大となるため, 9x9 グリッド環境においては Modular Q-learning とのみ比較を行った。閾値パラメータ λ の値は 7.0 とした。図 7 に提案手法と Modular Q-learning を用いた実験結果を示す。

5x5 グリッド環境における実験結果と同様に, 学習途中から

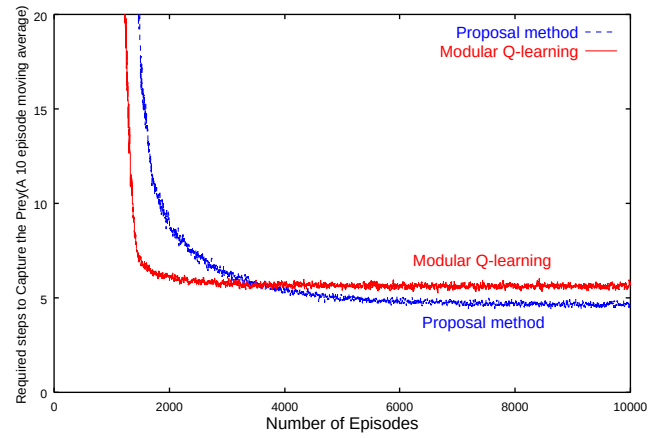


図 6 5x5 グリッド環境における実験結果：提案手法と Modular Q-learning の結果のみを表示

Fig. 6 An experimental result in 5x5 grids: A proposal method and Modular Q-learning.

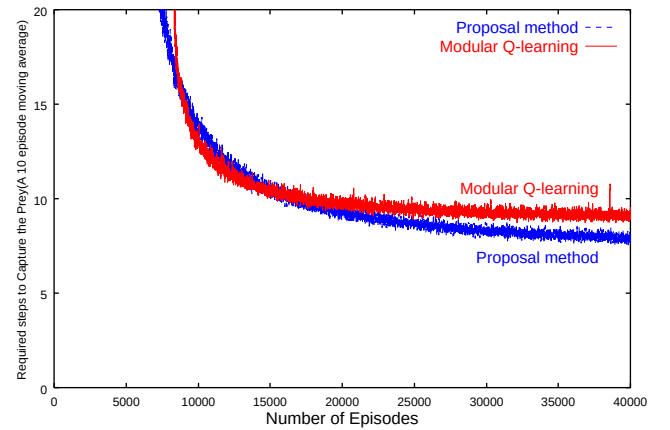


図 7 9x9 グリッド環境における実験結果

Fig. 7 An experimental result in 9x9 grids

提案手法のステップ数が Modular Q-learning のステップ数よりも小さな値となっている。また, Modular Q-learning, 提案手法における状態空間の大きさはそれぞれ 6561 状態, 7300 状態であり, 提案手法の状態空間の方が大きい。しかし, 図 7 が示すように, 提案手法の学習速度と Modular Q-learning の学習速度には 5x5 グリッド環境の時ほど差はない。この理由として以下のことが考えられる。

Q-learning において, 報酬はゴール近傍の状態からより遠くの状態へと伝播する。また, 割引率によってより遠くの状態ほど伝播する報酬は少なくなる。したがって, 割引率があまりに小さな値である場合や, ゴールに到達するまでに非常に多くのステップ数を必要とする場合には, ゴールから遠い状態における Q 値はほとんど学習されない。ただし, 割引率を適切な値に設定することでこのような問題は生じない。しかし Modular Q-learning では, これ以外の要因でゴールから遠い状態に伝播する報酬が減少し, 学習速度が低下してしまう。例えば, 図 8 のような状況を想定する。獲物が以後停止し続けると仮定した場合, 丸で囲まれた 2 体のエージェントと獲物からなる部分状態 (以下, 部分状態 S) において, 各エージェントは現在位置に

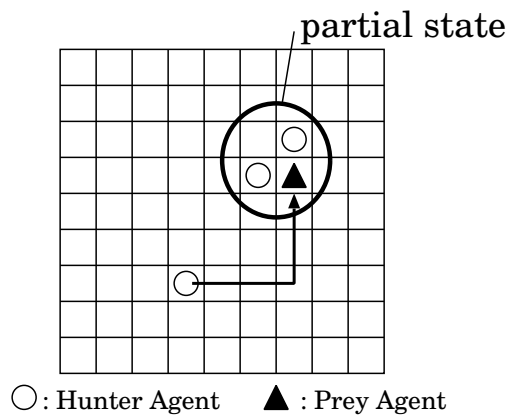


図 8 最適行動の Q 値が低下する状況:獲物は停止していると仮定した場合、丸で囲まれた 2 体のエージェントの最適行動「停止」の Q 値は、左下のエージェントが獲物に隣接するまで減少し続ける。
Fig. 8 The Q-Value of an optimal action decrease.

表 1 各手法における獲物捕獲に要したステップ数の比較

	Modular Q-learning	提案手法	Q-learning
5x5	5.60	4.65	4.85
9x9	9.10	7.90	—

停止する行動が最適な行動となる。また、残り 1 体のエージェントは単純に獲物に近づく行動が最適な行動となる。本来最適な行動を選択した場合、その Q 値は増加しなければいけない。しかし、この状況において左下のエージェントが最適な行動（5 ステップで獲物に隣接）を選択しても、その間獲物付近の 2 体のエージェントの停止行動に対する Q 値は減少し続ける。よって、その後のエピソードにおいて部分状態 S に遷移した場合、 $\max_a Q(S, a)$ として用いられる Q 値は本来期待される値よりも小さい。Modular Q-learning ではこのような状況が多く存在し、報酬伝播が遅延することにより学習速度が低下する。

一方、提案手法では高次元化手法により不完全知覚を取り除くことで、図 8 のような状況においても Q 値の値は減少せず、報酬伝播の遅延による学習速度の低下を防ぐことができたと考えられる。

表 1 に、各手法における試行終了時の獲物捕獲に要したステップ数を示す。ただし、Q-learning の学習エピソード数は 500 万エピソードとした。提案手法のステップ数が最も小さな値となっていることから、Modular Q-learning の学習性能改善に成功したと言える。また、本来なら Q-learning のステップ数が最も小さな値となるべきであるが、これは 500 万エピソードでは学習期間が足りなかったためと考えられる。

5. ま と め

マルチエージェント環境では、エージェント数が増加するに伴い状態空間が指数的に増大する。このため、マルチエージェント強化学習では学習速度が急激に低下する、膨大なメモリ量が必要となる、といった問題が生じる。Ono らは、この状態空間の爆発を防ぐ手法として Modular Q-learning を提案した。しかし、Modular Q-learning では部分状態のみを観測するた

め、不完全知覚が生じ学習性能が低下してしまうという欠点がある。

そこで、本研究では Modular Q-learning を基盤としたマルチエージェント強化学習法を提案し、学習性能低下の要因となる状態を高次元化することで不完全知覚を取り除くことで Modular Q-learning の学習性能を改善した。また、追跡問題を用いた実験により、十分な学習速度、少ないメモリ量で性能改善が可能であることを確認した。

本手法は Modular Q-learning を基盤としており、エージェントの学習に Q-learning を用いている。しかし、マルチエージェント環境は、他エージェントの政策変化によって環境の状態遷移確率が変わるため MDP 環境としてモデル化することができない。よって、学習に状態遷移先の Q 値が必要となる Q-learning をそのまま用いることの合理性に疑問が残る。今後の課題は、このような問題にも対応できるように学習法を改良することである。

文 献

- [1] Richard S. Sutton and Andrew G. Barto, "Reinforcement Learning: An Introduction", A Bradford Book, The MIT Press, 1998
- [2] 荒井 幸代, "マルチエージェント強化学習 -実用化に向けての課題・理論・諸技術との融合-", 人工知能学会誌, Vol.16, No.4, pp.476-481, 2001
- [3] Ryo Goto, Toshihiro Matsui and Hiroshi Matsuo, "State Generalization with Support Vector Machines in Reinforcement Learning", 4th Asia-Pacific Conference on Simulated Evolution and Learning
- [4] 高橋 泰岳, 浅田 稔, "階層型学習機構における状態行動空間の構成", 日本ロボット学会誌, Vol.21, No.2, pp.164-171, 2003
- [5] Irfan Gultekin and Ahmet Arslan, "Modular-Fuzzy Cooperation Algorithm for Multi-agent Systems", Advances in Information Systems: 2nd International Conference, pp.255-263, 2002
- [6] N. Ono and K. Fukumoto, "Multi-agent Reinforcement Learning: A Modular Approach", in Proceedings of the 2nd International Conference on Multi-agent Systems (ICMAS-96), AAAI Press, 1996
- [7] Whitehead, S., Karlsson, J. and Tenenber, J., "Learning Multiple Goal Behavior via Task Decomposition and Dynamic Policy Merging", Robot Learning, Kluwer Academic Publishers, 1993
- [8] 長行 康男, 伊藤 実, "2 体エージェント確率ゲームにおける他エージェントの政策推定を利用した強化学習法", 電子情報通信学会論文誌, Vol. J86-D-I No.11, 2003