

(19) 日本国特許庁(JP)

(12) 特 許 公 報(B2)

(11) 特許番号

特許第4853997号

(P4853997)

(45) 発行日 平成24年1月11日(2012.1.11)

(24) 登録日 平成23年11月4日(2011.11.4)

(51) Int.Cl.

F I

G 0 6 N 3/00 (2006.01)

G 0 6 N 3/00 5 5 0 E

請求項の数 6 (全 23 頁)

(21) 出願番号	特願2005-236604 (P2005-236604)	(73) 特許権者	393031586
(22) 出願日	平成17年8月17日(2005.8.17)		株式会社国際電気通信基礎技術研究所
(65) 公開番号	特開2007-52589 (P2007-52589A)		京都府相楽郡精華町光台二丁目2番地2
(43) 公開日	平成19年3月1日(2007.3.1)	(74) 代理人	100067828
審査請求日	平成20年3月27日(2008.3.27)		弁理士 小谷 悦司
(出願人による申告)平成17年度独立行政法人情報通信研究機構、研究テーマ「人間情報コミュニケーションの研究開発」に関する委託研究、産業活力再生特別措置法第30条の適用を受ける特許出願		(74) 代理人	100096150
			弁理士 伊藤 孝夫
		(74) 代理人	100109438
			弁理士 大月 伸介
		(72) 発明者	杉本 徳和
			京都府相楽郡精華町光台二丁目2番地2
			株式会社国際電気通信基礎技術研究所内
		(72) 発明者	鮫島 和行
			京都府相楽郡精華町光台二丁目2番地2
			株式会社国際電気通信基礎技術研究所内
			最終頁に続く

(54) 【発明の名称】 エージェント学習装置、エージェント学習方法及びエージェント学習プログラム

(57) 【特許請求の範囲】

【請求項1】

環境に基づいて学習し且つ当該学習の結果に基づいて当該環境に対して働きかけるエージェント学習装置であって、

観測関数を通した観測変数として環境の状態を観測し、当該観測変数に基づいて環境を連続状態から離散状態へと抽象化する環境抽象化手段と、

前記環境抽象化手段により抽象化された後の離散状態の中から、現在時刻において最も環境をよく抽象化している一の状態を指定するインデックスを決定する状態決定手段と、

前記状態決定手段により決定された一の離散状態において学習を行うことで取るべき一の行動を指定するインデックスを決定する行動決定手段と、

下位報酬として連続関数である複数の下位報酬関数の中から、前記状態決定手段により決定されたインデックスと、前記行動決定手段により決定されたインデックスとを有する一の下位報酬関数を選択する下位報酬選択手段と、

前記下位報酬選択手段により選択された下位報酬関数を最大にするように環境への制御出力を決定し、当該制御出力を用いて環境に対して働きかける制御出力決定手段とを備え

る。

前記環境抽象化手段は、現在時刻における環境の状態と環境への制御出力とに基づいてモデルごとに異なる環境の状態変化の予測値を算出する複数の状態予測モデルにより環境を抽象化し、前記状態予測モデルにより算出された環境の状態変化の予測値と観測変数から推定された環境の状態変化の推定値との差に基づく状態予測誤差が小さいものほど大き

10

20

な値を取る責任信号を算出し、当該責任信号を抽象化された環境ごとに出し、

前記状態決定手段は、前記責任信号を最大にするインデックスを、現在時刻において最も環境をよく抽象化している一の状態を指定するインデックスとすることを特徴とするエージェント学習装置。

【請求項 2】

前記行動決定手段は、行動価値関数の状態を指定するインデックスとして前記責任信号を最大にするインデックスを有し且つ前記責任信号により重み付けされた当該行動価値関数の値を最大とする行動を指定するインデックスを、取るべき一の行動を指定するインデックスとすることを特徴とする請求項 1 記載のエージェント学習装置。

【請求項 3】

前記制御出力決定手段は、下位報酬関数を最大にする際に、前記状態予測モデルと下位報酬関数との組み合わせにより構成される下位価値関数の値を最大にするように制御出力を決定することを特徴とする請求項 1 または 2 に記載のエージェント学習装置。

【請求項 4】

前記下位報酬関数は、環境の状態と制御出力とによって張られる関数空間において、前記状態決定手段により決定されたインデックスと、前記行動決定手段により決定されたインデックスとにより指定される離散的な領域内に頂点を持つ凸関数であることを特徴とする請求項 1 乃至 3 のいずれかに記載のエージェント学習装置。

【請求項 5】

環境抽象化手段、状態決定手段、行動決定手段、下位報酬選択手段および制御出力決定手段を備えるエージェント学習装置を用いて、環境に基づいて学習し且つ当該学習の結果に基づいて当該環境に対して働きかけるエージェント学習方法であって、

環境の状態を観測関数を通した観測変数として観測し、当該観測変数に基づいて環境を連続状態から離散状態へと抽象化する環境抽象化ステップと、

前記環境抽象化ステップにより抽象化された後の離散状態の中から、現在時刻において最も環境をよく抽象化している一の状態を指定するインデックスを決定する状態決定ステップと、

前記状態決定ステップにより決定された一の離散状態において学習を行うことで取るべき一の行動を指定するインデックスを決定する行動決定ステップと、

下位報酬として連続関数である複数の下位報酬関数の中から、前記状態決定ステップにより決定されたインデックスと、前記行動決定ステップにより決定されたインデックスとを有する一の下位報酬関数を選択する下位報酬選択ステップと、

前記下位報酬選択ステップにより選択された下位報酬関数を最大にするように環境への制御出力を決定し、当該制御出力を用いて環境に対して働きかける制御出力決定ステップとを含み、

前記環境抽象化ステップでは、現在時刻における環境の状態と環境への制御出力とに基づいてモデルごとに異なる環境の状態変化の予測値を算出する複数の状態予測モデルにより環境を抽象化し、前記状態予測モデルにより算出された環境の状態変化の予測値と観測変数から推定された環境の状態変化の推定値との差に基づく状態予測誤差が小さいものほど大きな値を取る責任信号を算出し、当該責任信号を抽象化された環境ごとに出し、

前記状態決定ステップでは、前記責任信号を最大にするインデックスを、現在時刻において最も環境をよく抽象化している一の状態を指定するインデックスとすることを特徴とするエージェント学習方法。

【請求項 6】

環境に基づいて学習し且つ当該学習の結果に基づいて当該環境に対して働きかけるエージェント学習プログラムであって、

観測関数を通した観測変数として環境の状態を観測し、当該観測変数に基づいて環境を連続状態から離散状態へと抽象化する環境抽象化手段と、

前記環境抽象化手段により抽象化された後の離散状態の中から、現在時刻において最も環境をよく抽象化している一の状態を指定するインデックスを決定する状態決定手段と、

10

20

30

40

50

前記状態決定手段により決定された一の離散状態において学習を行うことで取るべき一の行動を指定するインデックスを決定する行動決定手段と、

下位報酬として連続関数である複数の下位報酬関数の中から、前記状態決定手段により決定されたインデックスと、前記行動決定手段により決定されたインデックスとを有する一の下位報酬関数を選択する下位報酬選択手段と、

前記下位報酬選択手段により選択された下位報酬関数を最大にするように環境への制御出力を決定し、当該制御出力を用いて環境に対して働きかける制御出力決定手段としてコンピュータを機能させ、

前記環境抽象化手段は、現在時刻における環境の状態と環境への制御出力とに基づいてモデルごとに異なる環境の状態変化の予測値を算出する複数の状態予測モデルにより環境を抽象化し、前記状態予測モデルにより算出された環境の状態変化の予測値と観測変数から推定された環境の状態変化の推定値との差に基づく状態予測誤差が小さいものほど大きな値を取る責任信号を算出し、当該責任信号を抽象化された環境ごとに出力し、

10

前記状態決定手段は、前記責任信号を最大にするインデックスを、現在時刻において最も環境をよく抽象化している一の状態を指定するインデックスとすることを特徴とするエージェント学習プログラム。

【発明の詳細な説明】

【技術分野】

【0001】

本発明は、主に強化学習に用いられ階層構造を有するエージェント学習装置、エージェント学習方法及びエージェント学習プログラムに関するものである。

20

【背景技術】

【0002】

強化学習は、エージェントが環境からの報酬信号をもとに、自律的に行動則を獲得できるという優れた特性を有する学習理論である。しかしながら、これまでの強化学習は学習速度が遅いため、対象となる環境が複雑な場合は実行不可能なほど膨大な学習時間を要し、実用的な手法ではないとの批判を受けてきた。このため、強化学習の学習速度を速める手段として階層構造の導入が注目されており、階層構造を用いたいくつかの強化学習手法が提案されている。

【0003】

30

最も簡単な階層構造を用いた強化学習装置は、上位層と下位層との二層から構成される。下位層は、上位層と環境との相互作用を取り持つ機能を有し、主に環境を抽象化して上位層にその情報を伝える。そして、上位層は、下位層によって抽象化された状態空間において学習を行う。この階層構造を用いることで、報酬の伝搬の速い上位層は下位層から環境の状態を受け取り強化学習を行う。この構成により、当該強化学習装置は高速に行動則を獲得することができる。

【0004】

階層構造を用いた強化学習手法として、例えば、解像度の粗い上位層を設けることで効率よく探索を行う手法（非特許文献1参照）や、マルコフ決定過程（Markov Decision Process, MDP）を分解し、それぞれに価値関数を割り当てる手法（非特許文献2参照）が提案されている。

40

【非特許文献1】ピー ダヤン（P. Dayan）他、「封建的な強化学習」（Feudal reinforcement learning）、ニューラルインフォメーションプロセスシステムの進歩（Advances in Neural Information Processing Systems）、1993年

【非特許文献2】ティー ジー ディエテリッチ（T. G. Dietterich）、「MAXQ価値関数の分解を用いた階層的な強化学習」（Hierarchical reinforcement learning with the MAXQ value function decomposition）、アーティフィシャルインテリジェンスリサーチジャーナル（Journal of Artificial Intelligence Research）、2000年、vol.13、p. 227 - p. 303

【発明の開示】

50

【発明が解決しようとする課題】

【0005】

しかしながら、上記の従来手法は、離散時間及び離散状態を対象としたものであるが、実世界においてヒトやロボットが環境と相互作用を行う場合に観測される環境の状態や環境からの報酬信号、ロボット等を制御するために出力する制御信号等は、一般に連続的な値を取るため、上記従来の手法をそのまま用いても、実世界における連続時間及び連続状態を取り扱うことができない。

【0006】

また、実験者が連続的な環境を予め離散化してロボットに入力したりすることにより、階層構造を作り込む手法も考えられるが、この手法では、強化学習において最も困難な課題である「環境の抽象化」を実験者が事前知識に基づいて解決しているため、一般性があるとは言えず、連続的である環境の状態を自律的に学習することはできない。

10

【0007】

本発明の目的は、連続時間及び連続状態を取り扱うことができるとともに、環境から自律的且つ高速に学習を行うことができるエージェント学習装置、エージェント学習方法及びエージェント学習プログラムを提供することである。

【課題を解決するための手段】

【0008】

本発明に係るエージェント学習装置は、環境に基づいて学習し且つ当該学習の結果に基づいて当該環境に対して働きかけるエージェント学習装置であって、観測関数を通した観測変数として環境の状態を観測し、当該観測変数に基づいて環境を連続状態から離散状態へと抽象化する環境抽象化手段と、環境抽象化手段により抽象化された後の離散状態の中から、現在時刻において最も環境をよく抽象化している一の状態を指定するインデックスを決定する状態決定手段と、状態決定手段により決定された一の離散状態において学習を行うことで取るべき一の行動を指定するインデックスを決定する行動決定手段と、下位報酬として連続関数である複数の下位報酬関数の中から、状態決定手段により決定されたインデックスと、行動決定手段により決定されたインデックスとを有する一の下位報酬関数を選択する下位報酬選択手段と、下位報酬選択手段により選択された下位報酬関数を最大にするように環境への制御出力を決定し、当該制御出力でもって環境に対して働きかける制御出力決定手段とを備えるものである。

20

30

【0009】

本発明に係るエージェント学習装置では、環境が連続状態から離散状態へと抽象化され、現在時刻において最も環境をよく抽象化している一の状態を指定するインデックスが決定されるとともに、取るべき一の行動を指定するインデックスが決定され、これら2つのインデックスにより指定される一の下位報酬関数が選択され、選択された下位報酬関数を最大にする制御出力を用いて環境に働きかけているので、一般に連続状態を有する環境を離散状態へと抽象化し、当該離散状態において強化学習を行うことができ、高速に学習を行うことができる。また、環境に働きかける際には、その離散状態を再度連続状態へと変換しているため、連続時間及び連続状態を取り扱うことができるとともに、環境から自律的且つ高速に学習を行うことができる。

40

【0010】

また、環境抽象化手段は、現在時刻における環境の状態と環境への制御出力とに基づいてモデルごとに異なる環境の状態変化の予測値を算出する複数の状態予測モデルにより環境を抽象化し、状態予測モデルにより算出された環境の状態変化の予測値と観測変数から推定された環境の状態変化の推定値との差に基づく状態予測誤差が小さいものほど大きな値を取る責任信号を算出し、当該責任信号を抽象化された環境ごとに出力する。したがって、複数の状態予測モデルが算出する環境の状態変化の予測値はそれぞれ異なるため、現在の環境を最もよく抽象化している一の状態を有効に決定できる。

【0011】

また、状態決定手段は、責任信号を最大にするインデックスを、現在時刻において最も

50

環境をよく抽象化している一の状態を指定するインデックスとする。したがって、現在の環境を最もよく抽象化している一の状態を確実に決定できる。

【 0 0 1 2 】

行動決定手段は、行動価値関数の状態を指定するインデックスとして責任信号を最大にするインデックスを有し且つ責任信号により重み付けされた当該行動価値関数の値を最大とする行動を指定するインデックスを、取るべき一の行動を指定するインデックスとすることが好ましい。この場合、現在時刻において最も環境をよく抽象化している一の状態を指定するインデックスを有する状態の中で、最適な行動を有効に決定することができる。

【 0 0 1 3 】

制御出力決定手段は、下位報酬関数を最大にする際に、状態予測モデルと下位報酬関数との組み合わせにより構成される下位価値関数の値を最大にするように制御出力を決定することが好ましい。この場合、環境に与えられる制御出力として、下位報酬関数を最大にするようなものが選ばれるので、強化学習を有効に行うことができる。

【 0 0 1 4 】

下位報酬関数は、環境の状態と制御出力とによって張られる関数空間において、状態決定手段により決定されたインデックスと、行動決定手段により決定されたインデックスとにより指定される離散的な領域内に頂点を持つ凸関数であることが好ましい。この場合、関数空間において、隣り合った下位報酬関数同士の区別が明確になるために、抽象化された状態間の遷移が確定的になるため、学習を効率的に行うことができる。

【 0 0 1 5 】

本発明に係るエージェント学習方法は、環境抽象化手段、状態決定手段、行動決定手段、下位報酬選択手段および制御出力決定手段を備えるエージェント学習装置を用いて、環境に基づいて学習し且つ当該学習の結果に基づいて当該環境に対して働きかけるエージェント学習方法であって、観測関数を通した観測変数として環境の状態を観測し、当該観測変数に基づいて環境を連続状態から離散状態へと抽象化する環境抽象化ステップと、環境抽象化ステップにより抽象化された後の離散状態の中から、現在時刻において最も環境をよく抽象化している一の状態を指定するインデックスを決定する状態決定ステップと、状態決定ステップにより決定された一の離散状態において学習を行うことで取るべき一の行動を指定するインデックスを決定する行動決定ステップと、下位報酬として連続関数である複数の下位報酬関数の中から、状態決定ステップにより決定されたインデックスと、行動決定ステップにより決定されたインデックスとを有する一の下位報酬関数を選択する下位報酬選択ステップと、下位報酬選択ステップにより選択された下位報酬関数を最大にするように環境への制御出力を決定し、当該制御出力を用いて環境に対して働きかける制御出力決定ステップとを含み、前記環境抽象化ステップでは、現在時刻における環境の状態と環境への制御出力とに基づいてモデルごとに異なる環境の状態変化の予測値を算出する複数の状態予測モデルにより環境を抽象化し、前記状態予測モデルにより算出された環境の状態変化の予測値と観測変数から推定された環境の状態変化の推定値との差に基づく状態予測誤差が小さいものほど大きな値を取る責任信号を算出し、当該責任信号を抽象化された環境ごとに出力し、前記状態決定ステップでは、前記責任信号を最大にするインデックスを、現在時刻において最も環境をよく抽象化している一の状態を指定するインデックスとするものである。

【 0 0 1 6 】

本発明に係るエージェント学習プログラムは、環境に基づいて学習し且つ当該学習の結果に基づいて当該環境に対して働きかけるエージェント学習プログラムであって、観測関数を通した観測変数として環境の状態を観測し、当該観測変数に基づいて環境を連続状態から離散状態へと抽象化する環境抽象化手段と、環境抽象化手段により抽象化された後の離散状態の中から、現在時刻において最も環境をよく抽象化している一の状態を指定するインデックスを決定する状態決定手段と、状態決定手段により決定された一の離散状態において学習を行うことで取るべき一の行動を指定するインデックスを決定する行動決定手段と、下位報酬として連続関数である複数の下位報酬関数の中から、状態決定手段により

決定されたインデックスと、行動決定手段により決定されたインデックスとを有する一の下位報酬関数を選択する下位報酬選択手段と、下位報酬選択手段により選択された下位報酬関数を最大にするように環境への制御出力を決定し、当該制御出力を用いて環境に対して働きかける制御出力決定手段としてコンピュータを機能させ、前記環境抽象化手段は、現在時刻における環境の状態と環境への制御出力とに基づいてモデルごとに異なる環境の状態変化の予測値を算出する複数の状態予測モデルにより環境を抽象化し、前記状態予測モデルにより算出された環境の状態変化の予測値と観測変数から推定された環境の状態変化の推定値との差に基づく状態予測誤差が小さいものほど大きな値を取る責任信号を算出し、当該責任信号を抽象化された環境ごとに出力し、前記状態決定手段は、前記責任信号を最大にするインデックスを、現在時刻において最も環境をよく抽象化している一の状態を指定するインデックスとするものである。

10

【発明の効果】

【0017】

本発明によれば、一般に連続状態を有する環境を離散状態へと抽象化し、当該離散状態において強化学習を行うので、高速に学習を行うことができ、さらに、環境に働きかける際には、その離散状態を再度連続状態へと変換しているので、連続時間及び連続状態を取り扱うことができるとともに、環境から自律的且つ高速に学習を行うことができる。

【発明を実施するための最良の形態】

【0018】

まず、本発明の基本原理について説明する。階層型強化学習の目的の一つは、抽象化された状態空間で行動選択を行う上位層を設けることにより学習時間を短縮することである。このような上位層を用いる場合、「どのように状態や行動を抽象化するか」が最も難しく、また学習装置の性能に大きく関わってくる。特に、非線形な連続システムでは、環境の特性を実験者が把握することは困難なため、この抽象化が難しい。したがって、実験者を介さず、自律的（自動的）に抽象化が行える手法が必要となる。このようなモデルのひとつとして、WolpertらによるMOSAIC (Module Switching And Identification for Control) モデルが挙げられる。

20

【0019】

強化学習におけるMOSAICモデルとは、環境に対して働きかけ、その結果得られる報酬を最大化するための行動出力を決定する強化学習システムと、環境の変化を予測する環境予測システムとの組によりなる学習モジュールが複数備えられ、各学習モジュールの環境予測システムの予測誤差が少ないものほど大きな値を取る責任信号が求められ、この責任信号に比例して強化学習システムによる制御出力（行動出力）が重み付けされて、環境に対する行動が与えられることを特徴とする学習モデルのことである（特許第3086206号参照）。

30

【0020】

上記のMOSAICモデルは、環境に対する制御出力と、環境を観測することで得られる観測変数との入出力特性に基づいて、環境の文節化を自動で行うことができる。そして、分けられた文節ごとに制御器を切り替えることにより、非線形又は非定常な環境への適応を可能としている。

40

【0021】

次に、上記のMOSAICモデルに基づいた環境の抽象化について説明する。本明細書で対象とする環境においては、ある時刻 t における環境の状態 $x(t)$ （次元を D_s とする）と、エージェント学習装置が環境に働きかける際の制御出力 $u(t)$ （次元を D_c とする）と、制御出力 $u(t)$ のもとでの環境の状態変化 $x'(t)$ と、エージェント学習装置が環境を観測することにより得られる観測変数 $y(t) = g(x(t)) + v(t)$ （ $v(t)$ は観測ノイズ）との同時分布が、以下のように分解できると仮定する。なお、本明細書においては、文字に付けられたダッシュ（プライム）記号「 $'$ 」は、文字の上に付けられた点（ドット）「 $\cdot$ 」を意味するものとする。ただし、文字としては、 i, j, S を除くものとし、 i, j に付けられたダッシュ記号はそのまま、そのものを意味し、 S

50

' は u による偏微分を表すものとする。

【 0 0 2 2 】

【 数 1 】

$$p(\dot{x}(t), y(t) \mid x(t), u(t)) \\ = \sum_{i=1}^{N^f} \lambda_i^f(t) p(y(t) \mid x(t), i) p(\dot{x}(t) \mid x(t), u(t), i) \\ \dots (1)$$

10

【 0 0 2 3 】

ここで、 $p(\dot{x}'(t) \mid x(t), u(t), i)$ は局所的なダイナミクスを表し、 N^f はその個数、 $\lambda_i^f(t)$ は前述の責任信号と呼ばれる重み付けのための信号の中の i 番目の信号である。以下同様に、 i 等の下付の添字は、その添字が付けられた変数や状態等の i 番目のものを表すものとする。また、 g は観測関数と呼ばれ、エージェント学習装置はこの観測関数 g を通じて環境の状態 x を獲得することができる。

【 0 0 2 4 】

一般に、複数個の局所ダイナミクスが切り替わる環境のモデルは、“Switching State-Space Models (SSM)” と呼ばれる。本発明においては、MOSAICモデルに基づいて、局所ダイナミクスが切り替わるSSMとして環境をモデル化し、各時刻において有効になっている局所ダイナミクスのインデックスを上位層の状態として採用する。

20

【 0 0 2 5 】

MOSAICモデルに基づいた環境の文節化は、観測変数 $y(t)$ の時系列から現在の状態 $x(t)$ を推定し、状態予測モデルの予測誤差に基づいて責任信号 $\lambda_i^f(t)$ を求めることにより実現される。ここで、状態予測モデルとは、状態変化 $\dot{x}'(t)$ の予測値を出力するモデルであり、 $\omega(t)$ をシステムノイズとして以下の式で表される。

【 0 0 2 6 】

【 数 2 】

$$p(\hat{\dot{x}}(t) \mid x(t), u(t), i) = f_i(x(t), u(t)) + \omega(t) \dots (2)$$

30

【 0 0 2 7 】

図1(a)は後述するセミマルコフ決定過程(Semi-Markov Decision Processes、以下、SMDPという)を模式的に示す図であり、図1(b)は環境の状態変化を求めるための状態予測モデルを模式的に示す図である。図1(b)は、連続状態空間において、環境の状態 $x(t)$ と制御出力 $u(t)$ とに基づいて、状態変化 $\dot{x}'(t)$ を算出するダイナミクスを表している。この連続状態空間を複数の局所領域に分割し、それぞれにインデックス $i = 1, 2, \dots, N^f$ を割り当てると、図に示したように、そのインデックス i により指定される離散状態が形成される。図1(a)は、この離散状態間のダイナミクスの移り変わりを示しており、例えば、矢印の始点から終点へと離散状態が変化の様子を表している。

40

【 0 0 2 8 】

また、上記した状態 $x(t)$ の推定にはカルマンフィルタやモンテカルロ法などの手法を利用することができる。以下では、それらの手法により、状態 $x(t)$ の推定値 $\hat{x}^-(t)$ が得られたものとして責任信号の説明を行う。責任信号 $\lambda_i^f(t)$ は、時刻 t までの状態(の推定値) $\hat{x}^-(t)$ とその変化 $\hat{x}'^-(t)$ 、及び制御出力 $u(t)$ の時系列の事後分布として以下の式に基づいて推定される。なお、本明細書においては、文字に付けられたティルダ「 $\sim$ 」及びハット「 $\wedge$ 」は、文字の上に付けられていることを意味するものとする。そして、上記の x のように「 $\sim$ 」及び「 $\wedge$ 」、又は「 $\sim$ 」及び「 $\wedge$ 」が2つ

50

重ねてある場合には、文字の上にドットが付けられ、さらにその上にティルダ又はハットが付けられているものとする。

【 0 0 2 9 】

【 数 3 】

$$\lambda_i^f(t) \equiv P(i | \tilde{x}[0:t], \tilde{x}[0:t], u[0:t]) \quad \dots (3)$$

【 0 0 3 0 】

責任信号 $\lambda_i^f(t)$ は、各状態予測モデル $f_i(x(t), u(t))$ の尤度を正規化することで求められるが、その尤度を状態予測誤差 $\|x_i'(t) - \tilde{x}(t)\|^2$ に基づいて計算するのが責任信号の特徴である。そのような文節化によって非線形・非定常な環境を適応的に制御できることがこれまでの研究で示されている（例えば、D. M. Wolpert and M. Kawato: "Multiple paired forward and inverse models for motor control", Neural Networks, 11, pp. 1317-1369 (1998). 参照）。

10

【 0 0 3 1 】

したがって、各時刻において、最も高い責任信号を持つ状態予測モデルのインデックス i は、環境を良く抽象化していると言え、そのインデックス i を状態変数とする上位層は M O S A I C モデルの利点を有する状態空間で行動選択を行える。抽象化状態のインデックスを i 、選択可能な行動（サブゴール）のインデックスを j とすると、上位層が扱う環境は以下のような S M D P として定義できる。

20

【 0 0 3 2 】

1) 状態 i で行動 j を選択したとき、次状態 i' は、ある状態遷移確率 $P(i' | i, j)$ に従って決定される。

【 0 0 3 3 】

2) 状態 i から次状態 i' に遷移するまでの時間 l は、一定ではなく、ある確率分布 $P(l | i', i, j)$ を持つ。

【 0 0 3 4 】

3) 次状態 i' に遷移したときに与えられる報酬は、状態 i にいたときの報酬 $R(t)$ を時間で積分した値となる。

【 0 0 3 5 】

30

この S M D P は、状態間の遷移や値の更新といった処理の時間間隔が一定ではなく、イベントドリブンすなわちそれら処理の時間間隔が任意な場合に対応した数理モデルである。また、上記したような S M D P 環境において、どの抽象化状態が現時刻における環境の状態を最もよく表しているかを示す最適状態価値関数 V^* は、以下に示す B e l l m a n 方程式を満たさなければならない。

【 0 0 3 6 】

【 数 4 】

$$V^*(i) = \max_j \left[\sum_{i'=1}^{N^f} P(i'|i, j) \int_0^\infty P(l|i', i, j) \left\{ \int_t^{t+l} e^{-\frac{s-t}{\tau}} R(s) ds + e^{-\frac{l}{\tau}} V^*(i') \right\} dl \right] \quad \dots (4)$$

40

【 0 0 3 7 】

ここで、右辺中括弧内の第一項は、時刻 t から $t + l$ まで状態 i に滞在したときに得る

50

報酬の重み付き累積和であり、 τ はその時定数である。また、右辺中括弧内の第二項中の指数部分は、次状態 i' に遷移するまでの時間 l に応じた割引率である。そして、これら中括弧内の量を、状態遷移確率 $P(i' | i, j)$ と確率分布 $P(l | i', i, j)$ とで平均化している。この $SMDP$ における定式化は、時間が連続で且つ遷移する時間が分布を持つことを除けば MDP における定義と同様である。また、どの行動を取れば環境から得られる報酬を最も大きくすることができるかを示す最適行動価値関数 $Q^*(i, j)$ は以下の式を満たす。

【0038】

【数5】

$$Q^*(i, j) = \sum_{i'=1}^{N_f} P(i' | i, j) \int_0^{\infty} P(l | i', i, j) \left\{ \int_t^{t+l} e^{-\frac{s-t}{\tau}} R(s) ds + e^{-\frac{l}{\tau}} \max_j Q^*(i', j') \right\} dl$$

... (5)

10

【0039】

以上のように、 $MOSAIC$ モデルに基づいた SSM は、セミマルコフ的に局所ダイナミクスが切り替わることを特徴とする。

20

【0040】

次に、本発明の一実施の形態によるエージェント学習装置について説明する。図2は、本発明の一実施の形態によるエージェント学習装置の構成を示すブロック図である。

【0041】

図2に示すエージェント学習装置1は、上位層10及び下位層20の二層構造からなる。エージェント学習装置1は、ROM（リードオンリメモリ）、CPU（中央演算処理装置）、RAM（ランダムアクセスメモリ）等を備える通常のマイクロコンピュータ、A/D（アナログ/デジタル）変換器、D/A（デジタル/アナログ）変換器等から構成され、ROMに記憶されたエージェント学習プログラムをCPUにおいて実行することにより、上位層10及び下位層20として機能し、環境50の状態を通じて学習を行う。

30

【0042】

本実施の形態においては、 $MOSAIC$ モデルが階層型強化学習装置に適用され、 $MOSAIC$ モデルに基づいて切り分けられた文節を、行動選択を行う上位層に与える抽象化状態とすると、下位層20は、環境状態の抽象化を行い、上位層10は、各抽象化状態においてどの行動を目標とするべきかを強化学習により獲得する。そして、下位層20は、上位層10が選択した行動の達成度に応じて与えられる下位報酬を最大にするような制御則の学習を行う。

【0043】

上位層10は、行動選択器101、下位報酬関数選択器102、行動価値関数学習器103及び行動価値関数記憶器104を備え、行動選択を行う層であり、標準的な離散系強化学習の枠組みで記述される。そのため、 $MOSAIC$ モデルも含め、これまでに離散系を対象として提案されている手法を容易に適用できる。

40

【0044】

下位層20は、内部状態推定器201、責任信号推定器202、制御信号出力器203、状態予測モデル記憶器204、下位状態価値関数学習器205及び下位状態価値関数記憶器206を備え、上位層10と連続な状態をとる環境50との間を取り持つ層である。なお、以下の説明では、行動選択器101が行う行動選択と行動価値関数学習器103が行う行動価値関数の更新とは、上位層10に与えられる抽象化状態が変化したとき、つま

50

り最大責任信号を持つ状態予測モデルのインデックス i が変化したときに行われるとする。また、上位層 10 の状態が時刻 t にて I だったものが、時刻 $t + 1$ にて I' へ変化したものとして説明する。

【0045】

内部状態推定器 201 は、時刻 0 から時刻 t までの間に環境 50 から受け取った観測変数 $y[0:t]$ と、時刻 0 から時刻 t までの間に制御信号出力器 203 から受け取った制御出力 $u[0:t]$ とから、現時刻 t における環境 50 の状態 $x(t)$ 及びその変化（状態変化） $x'(t)$ の推定値 $\hat{x}(t)$ 、 $\hat{x}'(t)$ を算出し、責任信号推定器 202 に送出する（例えば、R. E. Kalman: “A new approach to linear filtering and prediction problems”, Transaction of the ASME-Journal of Basic Engineering, 82, Series D, pp. 35-45 (1960). または A. Doucet, N. de Freitas and N. Gordon: “Sequential Monte Carlo Methods in Practice”, Springer-Verlag (2001). 参照）。

10

【0046】

ここで、 $va[0:t]$ （ va は任意の変数を表す）は、時刻 0 から時刻 t の間に観測された変数 va の閉区間における時系列である。また、 $va(0:t)$ は、時刻 0 から時刻 t の間に観測された変数 va の开区間における時系列である。

【0047】

責任信号推定器 202 は、内部状態推定器 201 から受け取った環境 50 の状態 $x(t)$ を抽象化する。つまり、責任信号推定器 202 は、複数の状態予測モデル $\{f_i(x(t), u(t))\}$ を有する状態予測モデル記憶器 204 を参照し、制御の各時刻において責任信号 $\lambda_i^f(t)$ を上位層 10 に伝える。責任信号 $\lambda_i^f(t)$ は、現時刻の情報から求められる各状態予測モデル $f_i(x(t), u(t))$ の尤度と、それ以前の時系列から求められる事前予測値 $\hat{\lambda}_i^f(t)$ との積に比例する。この責任信号は、時間が連続である場合、以下の式で定義される。

20

【0048】

【数 6】

$$\lambda_i^f(t) \equiv \frac{\hat{\lambda}_i^f(t) p(\tilde{x}(t) | \tilde{x}(t), u(t), i)}{\sum_{i'=1}^{N^f} \hat{\lambda}_{i'}^f(t) p(\tilde{x}(t) | \tilde{x}(t), u(t), i')} \cdots (6)$$

30

$$\hat{\lambda}_i^f(t) \equiv \frac{P(i | \tilde{x}(t), u(t)) P(i | \tilde{x}[0:t], \tilde{x}[0:t], u[0:t])}{\sum_{i'=1}^{N^f} P(i' | \tilde{x}(t), u(t)) P(i' | \tilde{x}[0:t], \tilde{x}[0:t], u[0:t])} \cdots (7)$$

40

【0049】

ここで、確率分布 $p(\hat{x}'(t) | \hat{x}(t), u(t), i)$ は状態予測モデル $f_i(x(t), u(t))$ の尤度であり、予測誤差が分散 σ^{f^2} に従う正規分布だとした場合は、以下の式で定義される。

【0050】

【数 7】

$$p(\hat{\tilde{x}}(t)|\tilde{x}(t),u(t),i)=$$

$$\frac{1}{(2\pi)^{D_s/2}\sigma_f^{D_s}}\exp\left[-\frac{1}{2\sigma_f^2}\|\hat{\tilde{x}}_i(t)-\tilde{x}(t)\|^2\right]$$

$$\dots(8)$$

10

【0051】

ここで、 $\hat{x}^i(t)$ は i 番目の状態予測モデル $f_i(x(t), u(t))$ の予測出力値である。責任信号の事前予測値 $\lambda_i^f(t)$ は、空間的局所性と時間的連続性との2つの要素から得られ、式(7)においては、 $P(i|x^{\sim}(t), u(t))$ が空間的局所性、 $P(i|x^{\sim}[0:t], x^{\sim}[0:t], u[0:t])$ が時間的連続性に基づいた要素となっている。

【0052】

以上のようにして推定される責任信号を用いて各状態予測モデルの出力や学習係数を重み付けすることにより、非線形且つ非定常なダイナミクスの各局所領域を各状態予測モデルが適応的に予測するような分化が行われる。特に、入力に責任信号で重み付けした値の分散を用いて空間的局所性を更新することで分化が促進される。

20

【0053】

行動選択器101は、責任信号推定器202から責任信号 $\lambda_i^f(t)$ の推定値を受け取り、当該責任信号を最大にするインデックス i を決定(それを I とする)し、当該インデックス I を、現在時刻において最も環境をよく抽象化している一の状態を指定するインデックスとして選択する。上位層10は、この抽象化状態 I の中で学習を行う。

【0054】

また、行動選択器101は、行動価値関数記憶器104に記憶されている一つの行動価値関数 $Q(i, j)$ を参照し、当該行動価値関数 $Q(i, j)$ を責任信号 $\lambda_i^f(t)$ で重み付けした値を用いて行動選択を行う。これは、観測にノイズや隠れ状態が存在するとき、責任信号 $\lambda_i^f(t)$ は分布を持ち、信念状態の一種とみなされるからである。そして、行動選択器101は、新たに選択された(インデックス I で指定される)抽象化状態が、前回選択された抽象化状態と同じか否か、つまり環境50の抽象化状態が変化したか否かを判定する。その結果、抽象化状態が変化したと判定された場合には、行動選択器101は、行動価値関数学習器103に対して通知信号を送出する。

30

【0055】

ある行動を選択する確率(行動選択確率)は、例えば Boltzmann (ボルツマン) 分布に従うとすると、逆温度パラメータ β を用いて以下の式で記述される。

【0056】

40

【数 8】

$$P(j|\lambda_i^f(t), i=1, \dots, N^f) = \frac{\exp\left[\beta \sum_{i=1}^{N^f} \lambda_i^f(t) Q(i, j)\right]}{\sum_{j'} \exp\left[\beta \sum_{i=1}^{N^f} \lambda_i^f(t) Q(i, j')\right]} \dots (9)$$

10

【0057】

行動選択器 101 は、式 (9) の確率分布に従って確率的に j の値を決定し、その値をエージェント学習装置 1 が取るべき一の行動を指定するインデックス $J (= j)$ として選択する。そして、行動選択器 101 は、抽象化状態 I と、選択された行動を指定するインデックス J とを下位報酬関数選択器 102 に送出する。

【0058】

下位報酬関数選択器 102 は、抽象化状態を指定するインデックス i と、行動を表すインデックス j によって指定される複数の下位報酬関数 $\{r_{ij}\}$ を有している。そして、下位報酬関数選択器 102 は、行動選択器 101 によって決定された抽象化状態を指定するインデックス I と、行動選択器 101 によって選択された行動を表すインデックス J によって指定される下位報酬関数 r_{IJ} を、下位状態価値関数学習器 205 に送出する。

20

【0059】

上位層 10 は、ある状態遷移確率 $P(i' | i, j)$ に従って遷移する抽象化状態を環境として学習を行う。その状態遷移確率は、各抽象化状態 i に対する下位報酬関数 $r_{ij}(x(t), u(t))$ によって決定されるが、その状態遷移が確定的であるほど上位層 10 は学習を行い易い。そこで、下位報酬関数 $r_{ij}(x(t), u(t))$ は、関数空間において、抽象化状態と行動とを表すインデックスにより指定される離散的な領域内に頂点を持つ凸関数を採用することが好ましい。このような下位報酬関数を下位層 20 に与えれば、任意の抽象化状態への遷移確率をより確定的にすることができる。

30

【0060】

状態予測モデル記憶器 204 は、状態予測モデル $f_i(x(t), u(t))$ を記憶している。また、状態予測モデル記憶器 204 は、制御信号出力器 203 に対して、状態予測モデル f_i を制御出力 u で偏微分した値を出力する。さらに、状態予測モデル記憶器 204 は、下位状態価値関数学習器 205 に対して、状態予測モデル f_i から算出された予測出力値 $x'^\wedge(t)$ を出力する。

【0061】

以下、状態予測モデル $f_i(x(t), u(t))$ と、下位報酬関数 $r_{ij}(x(t), u(t))$ との組み合わせからなる関数を下位状態価値関数 $V_{ij}(x(t))$ という。この下位状態価値関数 $V_{ij}(x(t))$ の最適値である最適下位状態価値関数 $V_{ij}^*(x(t))$ は、 ξ を $0 < \xi \leq \tau$ を満たす定数として、以下の Bellman 方程式で記述される。

40

【0062】

【数 9】

$$\frac{1}{\xi} V_{ij}^*(x(t)) = \max_u \left[r_{ij}(x(t), u(t)) + \frac{\partial V_{ij}^*(x(t))}{\partial x} f_i(x(t), u(t)) \right] \quad \dots (10)$$

10

【0063】

下位状態価値関数記憶器206は、下位状態価値関数 $V_{ij}(x(t))$ を記憶している。下位状態価値関数学習器205は、下位状態価値関数記憶器206に記憶されている下位状態価値関数 V_{ij} を参照し、当該下位状態価値関数 V_{ij} を前述の式(10)で与えられる最適下位状態価値関数 V_{ij}^* に近づけるための学習を行う。下位状態価値関数学習器205が行う学習は、以下の式に示す誤差を小さくするように行われる。

【0064】

【数10】

$$\delta_{ij}(t) = r_{ij}(x(t), u(t)) - \frac{1}{\tau} V_{ij}(x(t)) + \frac{\partial V_{ij}(x(t))}{\partial x} \hat{x}_i(t) \quad \dots (11)$$

20

【0065】

ここで δ_{ij} はTD誤差と呼ばれ、当該学習はこのTD誤差の自乗値が減少する方向へ式(11)のパラメータを更新することで実現される(例えば、K. Doya: "Reinforcement learning in continuous time and space", Neural Computation, 12, pp. 219-245 (2000). 参照)。そして、下位状態価値関数学習器205は、更新後の新たな下位状態価値関数 V_{ij} を下位状態価値関数記憶器206に記憶させるとともに、 V_{ij} を状態 x で偏微分した値を制御信号出力器203に対して送出する。

30

【0066】

また、最適下位状態価値関数 $V_{ij}^*(x(t))$ に従った最適制御出力 $u_{ij}^*(x(t))$ は、学習が収束したときに得られるものであり、以下の式で表される。

【0067】

【数11】

$$u_{ij}^*(x(t)) = S'^{-1} \left(\frac{\partial f_i(x(t), u(t))}{\partial u} \frac{\partial V_{ij}^*(x(t))}{\partial x} \right) \quad \dots (12)$$

40

【0068】

ここで、 $S'(\cdot)$ は出力のコスト関数を制御出力 u で偏微分した関数であり、 $S'^{-1}(\cdot)$ はその逆関数を表す。制御信号出力器203は、式(12)によって算出された制御出力 $u_{ij}(x(t))$ を内部状態推定器201に送出するとともに、当該制御出力 $u_{ij}(x(t))$ を用いて環境50に働きかける。つまり、下位層20は、上位層10

50

の状態と行動とに応じた下位価値関数を学習しながら制御出力を行う。

【 0 0 6 9 】

行動価値関数 $Q(i, j)$ は、各抽象化状態に対する行動の評価値を表すものである。行動価値関数学習器 103 は、制御信号出力器 203 が与えた制御出力 $u_{I, j}$ に応じて環境 50 から受け取った報酬 $R(t)$ を観測し、行動価値関数 $Q(i, j)$ が持つ評価値の誤差を算出する。そして、行動価値関数学習器 103 は、抽象化状態が変化したことを知らせる通知信号を行動選択器 101 から受け取ると、行動価値関数記憶器 104 に記憶されている行動価値関数 $Q(i, j)$ の誤差を修正する。行動価値関数 $Q(i, j)$ は、報酬 R の将来に渡った重み付累積和によって行動を評価するものであり、エージェントはその評価が最大となるような行動則を学習する。行動価値関数 $Q(i, j)$ は、式 (5) で表される最適に近づくように学習が行われる。抽象化状態 I に対して選択された行動を J 、変化した抽象化状態を I' とすると、学習は学習係数 α を用いて以下の式で行われる。

【 0 0 7 0 】

【数 12】

$$Q(I, J) := Q(I, J) + \alpha \left[\int_t^{t+1} e^{-\frac{s-t}{\tau}} R(x(s), u(s)) ds + e^{-\frac{1}{\tau}} \max_j Q(I', j) - Q(I, J) \right] \dots (13)$$

【 0 0 7 1 】

ここで、 τ は $0 < \tau < \infty$ を満たす割引の時定数である。また、式 (13) における符号「 $:=$ 」は、当該符号の右辺の値を新たに左辺の値とする、更新処理を行うことを意味する。

【 0 0 7 2 】

本実施の形態においては、内部状態推定器 201、責任信号推定器 202 及び行動選択器 101 が環境抽象化手段の一例に相当し、行動選択器 101 が状態決定手段及び行動決定手段の一例に相当し、下位報酬関数選択器 102 が下位報酬選択手段の一例に相当し、制御信号出力器 203 が制御出力決定手段の一例に相当する。

【 0 0 7 3 】

次に、上記のように構成されたエージェント学習装置の学習処理について説明する。図 3 は、図 2 に示すエージェント学習装置の学習処理を説明するためのフローチャートである。

【 0 0 7 4 】

まず、ステップ S1 において、内部状態推定器 201 は、環境 50 から受け取った観測変数 y の時系列と、制御信号出力器 203 から受け取った制御出力 u とから、現時刻 t における環境 50 の状態 $x(t)$ 及びその変化 $x'(t)$ の推定値 $\hat{x}(t)$ 、 $\hat{x}'(t)$ を算出する。内部状態推定器 201 から推定値 $\hat{x}(t)$ 、 $\hat{x}'(t)$ を受け取った責任信号推定器 202 は、ステップ S2 において、それらの値に基づき、時刻 t における責任信号 $\lambda_i^f(t)$ を推定する。

【 0 0 7 5 】

次に、ステップ S3 において、行動選択器 101 は、まず、責任信号推定器 202 から受け取った責任信号の推定値 $\lambda_i^f(t)$ を最大にするインデックス i を決定し、当該インデックス i を環境 50 の抽象化状態 I として選択する。そして、ステップ S4 において、行動選択器 101 は、新たに選択されたインデックス I が、前回選択されたインデックス I と同じか否か、つまり環境 50 の抽象化状態が変化したか否かを判定する。その結

果、抽象化状態が変化していないと判定された場合（ステップS4でNO）には、引き続きステップS7以降の処理が行われる。それに対して、抽象化状態が変化すると判定された場合（ステップS4でYES）には、ステップS5以降の処理として新たな選択処理が行われる。

【0076】

具体的には、ステップS5において、行動選択器101は、責任信号で重み付けされた行動価値関数Qを参照して、前述の式（9）の行動選択確率に従って確率的にjの値を決定し、その値をエージェント学習装置1が取るべき行動Jとして選択する。そして、行動選択器101は、抽象化状態Iと、行動を指定するインデックスJとを下位報酬関数選択器102に送出する。次に、下位報酬関数選択器102は、ステップS6において、抽象化状態Iと行動Jとによって指定される下位報酬関数 $r_{I,J}$ を下位状態価値関数学習器205に送出する。

10

【0077】

次に、ステップS7において、下位報酬関数 $r_{I,J}$ を受け取った下位状態価値関数学習器205は、式（11）で与えられるTD誤差の自乗値を小さくするように下位状態価値関数を更新する。制御信号出力器203は、ステップS8において、下位状態価値関数 $V_{i,j}(x(t))$ に従った制御出力 $u_{I,J}$ を決定し、内部状態推定器201へ送出するとともに、環境50に対して出力することで環境50に働きかける。

【0078】

そして、行動価値関数学習器103は、与えられた制御出力 $u_{I,J}$ に応じた報酬Rを環境50から受け取ると、ステップS9において、報酬Rの積分を行う。このとき、ステップS10において、行動価値関数学習器103が行動選択器101から通知信号を受け取っていない場合、つまり抽象化状態が変化していないと判定された場合（ステップS10でNO）には、引き続きステップS12以降の処理が行われる。それに対して、抽象化状態が変化すると判定された場合（ステップS10でYES）には、行動価値関数学習器103は当該行動価値関数による更新処理を行う。

20

【0079】

つまり、上位層10における学習の結果を反映させ、抽象化状態iが選択されたときにどの行動jを選択すれば最も価値があるか、つまり最も大きな報酬が得られるかを行動価値関数Qの更新という形で取り込む。具体的には、ステップS11において、行動価値関数学習器103は、環境50から与えられた報酬Rに基づいて算出された新たな行動価値関数Qにより、行動価値関数記憶器104に記憶された行動価値関数Qの値を更新する。そして、以上のステップS1からS11までの処理を、学習が終了する（ステップS12でYES）まで繰り返し行う。

30

【0080】

ここで、図3においては、ステップS4及びS10で抽象化状態Iが変化した場合にのみ、ステップS5、S6及びS11の処理を行うとして説明した。つまり、例えば、行動価値関数Qの更新は、一定の時間間隔ではなく、責任信号を最大にする状態が変わったときに（セミマルコフ的に）行われるとした。しかしながら、本実施の形態はそれに限られることなく、ステップS5、S6及びS11の処理を、一定の時間間隔ごとに行う形態であってもよい。

40

【0081】

上記の一連の処理により、本発明に係る実施の形態では、一般に連続状態を有する環境を離散状態へと抽象化し、当該離散状態において強化学習を行うので、高速に学習を行うことができる。また、環境に働きかける際には、その離散状態を再度連続状態へ変換しているので、連続時間及び連続状態を有する環境を取り扱うことが可能となり、当該環境から自律的且つ高速に学習を行うことができる。

【0082】

次に、上記のエージェント学習装置の学習効果について、具体例を挙げて説明する。図4は、倒立振子の構成を示す図である。

50

【 0 0 8 3 】

図 4 に示すように、振子 6 0 は、台車 7 0 の一側面の中心に回動自在に取り付けられている。振子 6 0 の位置は、当該振子 6 0 が取り付けられた中心を通る鉛直線の上方の部分から測った角度 θ によって指定されるものとする。台車 7 0 は、図 4 における右向きを z 軸の正の向きとした場合に、 z 軸に沿って加えられる力 F によって一軸方向に移動可能なように構成されている。

【 0 0 8 4 】

力 F の大きさ (F の絶対値) $|F|$ は、 $|F| < 20 [N]$ であるとする。また、台車 7 0 の位置は、図の点線で示される位置を基準 ($z = 0$ とする) とし、当該基準位置から振子 6 0 が取り付けられた中心までの z 座標で表すものとする。さらに、台車 7 0 が移動可能な幅に制限はないものとする。

10

【 0 0 8 5 】

このような構成において、状態 $x(t)$ は、台車 7 0 の座標 z 、その座標 z の変化 z' 、角度 θ 、そしてその角度の変化 θ' によって指定され、 $x(t) = [z \ \theta \ z' \ \theta']^T$ (T は行列の転置を表す) と表すことができる。また、制御出力 $u(t)$ は合力 F であるため、 $u(t) = F$ である。各観測関数は、 $g_i(x(t)) = x(t) + N(0, 0.1 I_{D_s})$ とした (ただし、 I_{D_s} は D_s 次元の単位行列を表し、 N はガウスノイズ (ホワイトノイズ) を表すものとする。)。物理パラメータは、Barto ら (A. G. Barto, R. S. Sutton and C. W. Anderson: "Neuronlike adaptive elements that can solve difficult learning control problems", IEEE Transactions on Systems Man and Cybernetics, 13, pp.834-846 (1983). 参照) が用いたものと同じ設定にし、制御の間隔は $0.01 [sec]$ とした。また、 $N(0, 0.1 I_{D_s})$ はシステムノイズであり、状態変数の次元 D_s は 4、制御入力次元 D_c は 1 である。

20

【 0 0 8 6 】

1 試行の長さは $10 [sec]$ とし、各試行の初期状態は振子 6 0 がぶら下がった状態 ($\theta = \pi [rad]$) から $\pm 15\pi / 180 [rad]$ の範囲でランダムに選択した。報酬関数 $R(x(t), u(t))$ は、振子 6 0 が倒立した状態 ($\theta = 0 [rad]$) の周りで大きく与えられる正の報酬と出力のコストの和で、次式に示すように与えた。

【 0 0 8 7 】

【数 1 3】

30

$$R(x(t), u(t)) = \exp\left[-\frac{1}{2}\theta(t)^2\right] - 0.0001F(t)^2$$

・・・ (14)

【 0 0 8 8 】

上記の倒立振子を用い、図 1 に示すエージェント学習装置 1 による階層型強化学習アルゴリズム (以下、「SSRL」という) と、標準的な強化学習の 2 つの手法を用いて実験を行った。これら 2 つの手法の設定を、それぞれ以下に示す。

40

【 0 0 8 9 】

図 1 に示すエージェント学習装置 1 では、下位層の状態予測モデル f_i として、以下のような線形モデルを計 100 個用意した。

【 0 0 9 0 】

【数 1 4】

$$f_i(x, u) = A_i(x - x_i^f) + B_i u + c_i \quad \dots (15)$$

$$P(i|\dot{x}[0:t], x[0:t], u[0:t]) = \frac{1}{N^f} \quad \dots (16)$$

$$P(i|x(t), u(t)) = \frac{P(i, x(t), u(t))}{\sum_{i'=1}^{N^f} P(i', x(t), u(t))} \quad \dots (17)$$

10

$$P(i, x(t), u(t)) = \frac{1}{(2\pi)^{D_s/2} |\sum_i^f|^{D_s/2}} \exp \left[-(x(t) - x_i^f)^T \sum_i^{f-1} (x(t) - x_i^f) \right] \quad \dots (18)$$

20

【0091】

また、責任信号 $\lambda_i^f(t)$ の事前予測値 $\lambda_i^{\wedge f}(t)$ は、空間的局所性に基づくもののみを用いた。予測誤差の分散 σ^{f^2} は 0.1^2 で固定とし、その他のパラメータ A_i 、 B_i 、 c_i 、 x_i^f 、 Σ_i^f 等は以下に示すEMアルゴリズムを用いて予め学習したものを実験に用いた。

【0092】

まず、学習に用いたデータは $[0 - \pi \sim \pi - 10 \sim 10 - 2\pi \sim 2\pi]^T$ の範囲で一様にランダムな初期状態 $x(0)$ からガウスノイズを制御入力 ($u(t) = [10 N(0, 1)]$) として与える事で生成した。そして、 $0.1 [sec]$ の試行を計 1000 回行った。

30

【0093】

学習の対象としたパラメータとその初期値は以下のようにして設定した。式(15)のような線形モデルを状態予測モデルとした場合、 A_i の上 2 行が $[O_{D_s/2} \quad I_{D_s/2}]$ (ただし、 $O_{D_s/2}$ は $D_s/2$ 次元の正方なゼロ行列を表す。)、 B_i の上 2 行と c_i の上半分の要素が 0 となることは明らかであるため、その他の要素のみ学習の対象とし、初期値は $-0.1 \sim 0.1$ の間における一様な乱数として与えた。入力の分散の初期値は $0.1 I_{D_s}$ で与え、台車の横方向 (z 方向) に関係する要素のみを固定し、出力の分散 σ^{f^2} は 0.1^2 で固定した。状態予測モデルの中心 x_i^f は、用意したデータの範囲で一様な乱数とした。

40

【0094】

そして、以下の処理 (Eステップ及びMステップ) を、下記の式(19) ~ (21) を用いて各パラメータが十分に収束するまで繰り返した。

・Eステップ：各データに対する i 番目の状態予測モデルの責任信号 $\lambda_i^f(t)$ を求める。

・Mステップ：各パラメータを更新する。

【0095】

【数 15】

$$x_i^f = \langle x \rangle_i / \langle 1 \rangle_i \quad \dots (19)$$

$$\Sigma_i^f = \langle (x - x_i^f)(x - x_i^f)^T \rangle / \langle 1 \rangle_i \quad \dots (20)$$

$$W_i = \langle \dot{x} [x^T u^T 1] \rangle \langle [x^T u^T 1]^T [x^T u^T 1] \rangle^{-1} \quad \dots (21)$$

10

【0096】

ここで、Eステップにおいて、 W_i はパラメータ A_i , B_i , c_i を横に並べたものである。つまり、 $W_i = [A_i, B_i, c_i]$ である。また、 $\langle \rangle_i$ は責任信号による重み付き平均を表し、例えば、 $\langle h(t) \rangle_i = (1/T) (\int_{t=0}^T h(t) \lambda_i^f(t) dt)$ ($\int_{t=0}^T$ は $t=0$ から $t=T$ に亘る定積分を求めることを意味する積分記号であり、 T はデータの個数を表す。) である。

【0097】

上位層が下位層に与える下位報酬関数は、獲得された状態予測モデルをもとにし、次のように自動的に設定した。ある抽象化状態 i の中心 x_i^f と、他の抽象化状態の中心との間の距離を調べ、その距離が短かった順に自分自身を含めた5つを選択する。そして、選ばれた各抽象化状態の中心 x_i^f に頂点を持ち、その領域の逆共分散 $(\Sigma_i^f)^{-1}$ を傾きとした2次関数： $-(x(t) - x_i^f)^T (\Sigma_i^f)^{-1} (x(t) - x_i^f)$ を下位報酬関数として設定した。

20

【0098】

また、行動価値関数 Q は、*Adaptive RTDP* (Adaptive Real-Time Dynamic Programming; 例えば、A. G. Barto, S. J. Bradtke and S. P. Singh: "Learning to act using real-time dynamic programming", *Artificial Intelligence*, 72, 1-2, pp.81-138 (1995). 参照) を用いて行い、行動選択は ϵ -Greedy (例えば、R. S. Sutton and A. G. Barto: "Reinforcement Learning", MIT Press (1998). 参照) を用い、 $1/10$ の確率でランダムに行うとした。さらに、行動価値関数の時定数 τ は2とした。

30

【0099】

下位価値関数 V_{ij} は、すべての状態予測モデル f_i と、下位報酬関数 r_{ij} との組み合わせに対して用意し、*NGNet* (例えば、J. Moody and C. J. Darken: "Fast learning in networks of locally-tuned processing units", *Neural Computation*, 1, 2, pp. 281-294 (1989). 参照) を用いて学習した。下位価値関数 V_{ij} は、 $1 \times 5 \times 5 \times 5$ 個の基底を状態予測モデル f_i の中心 x_i^f の周りに配置し、*exponential eligibility trace* (例えば、K. Doya: "Reinforcement learning in continuous time and space", *Neural Computation*, 12, pp. 219-245 (2000). 参照) によってパラメータを更新した。その際、学習係数 α は0.1、下位価値関数の時定数 ξ は1.0、適格度の係数は0.5とした。下位層の出力 u は、各時刻における上位層の状態と行動とに対応した状態予測モデル f_I と下位価値関数 V_{IJ} とを用いて次のように定義した。また、コスト関数は、 $S(F(t)) = 0.0001 F(t)^2$ とした。

40

【0100】

【数 16】

$$u(t) = S'^{-1} \left(\frac{\partial f_I(x(t), u(t))^T}{\partial u} \frac{\partial V_{IJ}(x(t))^T}{\partial x} \right) \dots (22)$$

【0101】

10

一方、標準的な強化学習として、上記のSSRLとは異なり、階層構造を持たない標準的な強化学習手法で行動則の獲得を行った。価値関数の近似には、NGNetを用いた。NGNetの基底は、 $1 \times 20 \times 25 \times 25$ 個を状態空間へ格子上に配置し、報酬 $R(t)$ をもとに学習を行った。行動出力 $u(t)$ は、SSRLの下位層と同じように、状態予測モデルと価値関数の勾配をもとに生成した。状態予測モデルは、SSRLの場合と同じ条件で事前に学習したものを用いた。学習に用いた価値関数の学習係数、時定数とその適格度の係数及び基底の個数と分散は、SSRLの下位層と同じものを用いて公平な条件になるように設定した。

【0102】

20

図5は、長さが1000試行の実験を10回行ったときの結果を示す。図5において、実線がSSRL、破線が標準的な強化学習による結果を表し、図の横方向に実線又は破線で繋がれた各点は、それぞれSSRL、標準的な強化学習による1000試行毎の平均を表し、上下方向に描かれた実線及び破線の線分は、それぞれSSRL、標準的な強化学習による標準偏差を表すものである。

【0103】

30

図5からわかるように、価値関数の学習に用いた関数近似器、状態予測モデルの性能は、本実施の形態のSSRLと標準的な強化学習とにおいて同じであるにもかかわらず、本実施の形態のSSRLが勝る結果となった。つまり、標準的な強化学習では、連続な状態空間全体に報酬が伝播するのを待たなければ振子60が振り上げられないが、SSRLでは、抽象化された状態空間で報酬を伝播させることができるため、標準的な強化学習に比べはるかに高速に学習を行えた。また、SSRLでは早い段階で報酬が広く伝播していたが、標準的な強化学習の場合は初期状態の付近に報酬が伝播する前に実験が終了してしまい、ほとんど振子60を振り上げることができなかった。その結果、最後の100試行における振り上げの成功率は、SSRLが39.875[%]、標準的な強化学習が2.125[%]であった。

【0104】

40

公平な比較を行うために、SSRLの実験においてNGNetを用いて下位価値関数 V_{ij} の学習を行った。下位価値関数 V_{ij} は、組になっている状態予測モデル f_i と下位報酬関数 r_{ij} とを環境の代替とすることで、直接制御則 $u_{ij}(x)$ を導くことができる。特に、今回の実験のように状態予測モデル f_i が線形であり、下位報酬関数 r_{ij} が2次形式の場合はLinear Quadratic Controller (LQC; 例えば、D. P. Bertsekas: “Dynamic programming and optimal control”, Athena Scientific (1995). 参照)として容易に求められる。本実施の形態の実用性を高めるためには、下位の行動則をLQCとして与えた方が良く、今回の実験と同じ条件で下位の行動則をLQCとした場合の成功率は、最初の100試行の段階で91.5[%]となった。

【0105】

以上説明したように、抽象化された空間での状態遷移は、SMDPの枠組みで記述されるため、本実施の形態では、SMDP環境に対する既存の手法をそのまま上位層の学習則として適用可能である。この利点を生んだ大きな要因は、SMDP環境と連続システムである環境との受け渡しをしている下位層の存在である。つまり、下位層は、連続システム

50

を SMDP 環境に抽象化する役割を果たしている。

【 0 1 0 6 】

また、連続変数で記述される事象と離散変数で記述される事象とが混在するシステムは、Hybrid Dynamical System と呼ばれ、非線形で大規模な制御対象をこの Hybrid Dynamical System としてモデル化することにより、制御器の性能が向上することが知られている。本実施の形態においては、この Hybrid Dynamical System を階層型強化学習により制御する枠組みを示したものであるとも言える。

【図面の簡単な説明】

【 0 1 0 7 】

【図 1】セマルコフ決定過程を模式的に示す図および環境の状態変化を求めるための状態予測モデルを模式的に示す図である。 10

【図 2】本発明の一実施の形態によるエージェント学習装置の構成を示すブロック図である。

【図 3】図 2 に示すエージェント学習装置の学習処理を説明するためのフローチャートである。

【図 4】倒立振子の構成を示す図である。

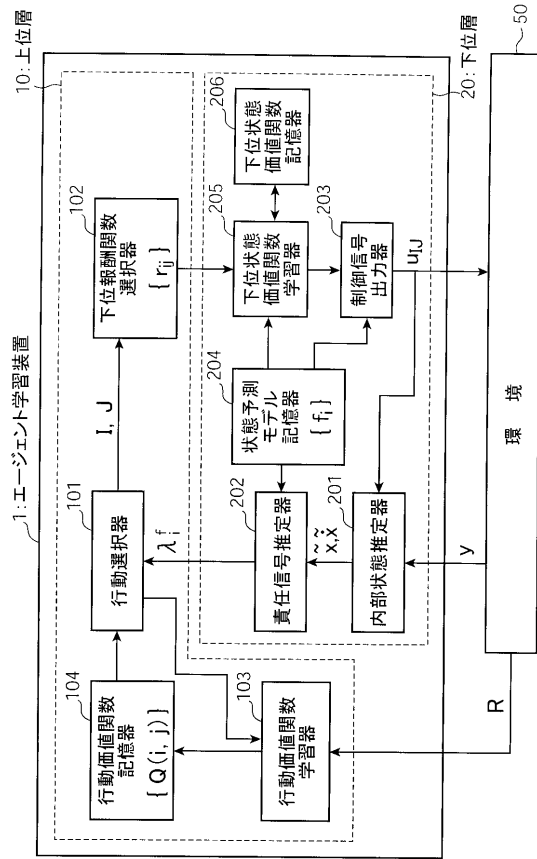
【図 5】図 4 に示す倒立振子を用いた実験の結果を示す図である。

【符号の説明】

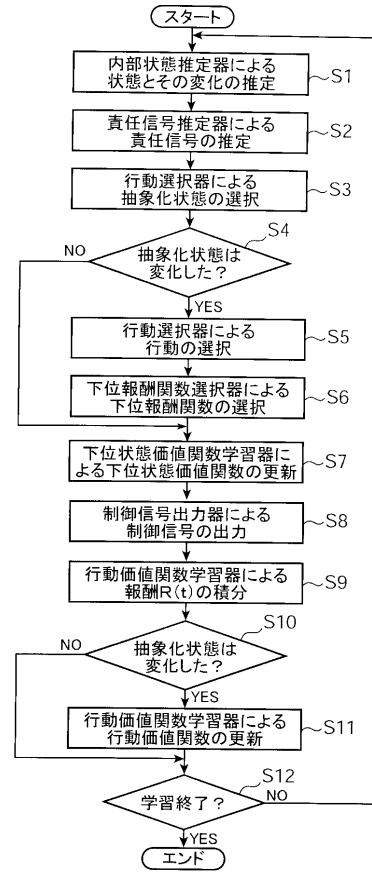
【 0 1 0 8 】

- 1 エージェント学習装置 20
- 1 0 上位層
- 2 0 下位層
- 1 0 1 行動選択器
- 1 0 2 下位報酬関数選択器
- 1 0 3 行動価値関数学習器
- 1 0 4 行動価値関数記憶器
- 2 0 1 内部状態推定器
- 2 0 2 責任信号推定器
- 2 0 3 制御信号出力器
- 2 0 4 状態予測モデル記憶器 30
- 2 0 5 下位状態価値関数学習器
- 2 0 6 下位状態価値関数記憶器

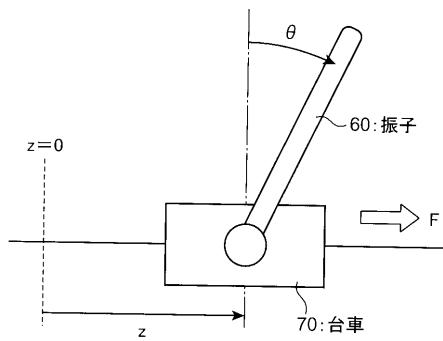
【図 2】



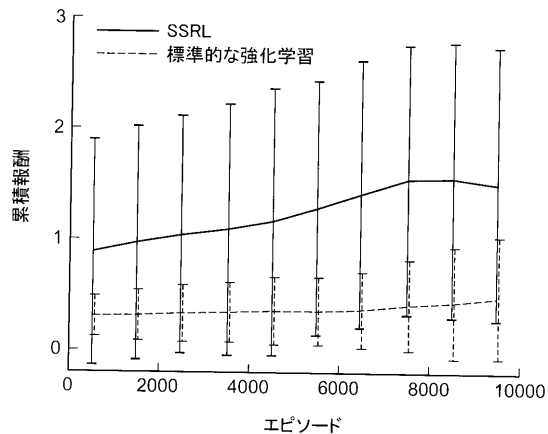
【図 3】



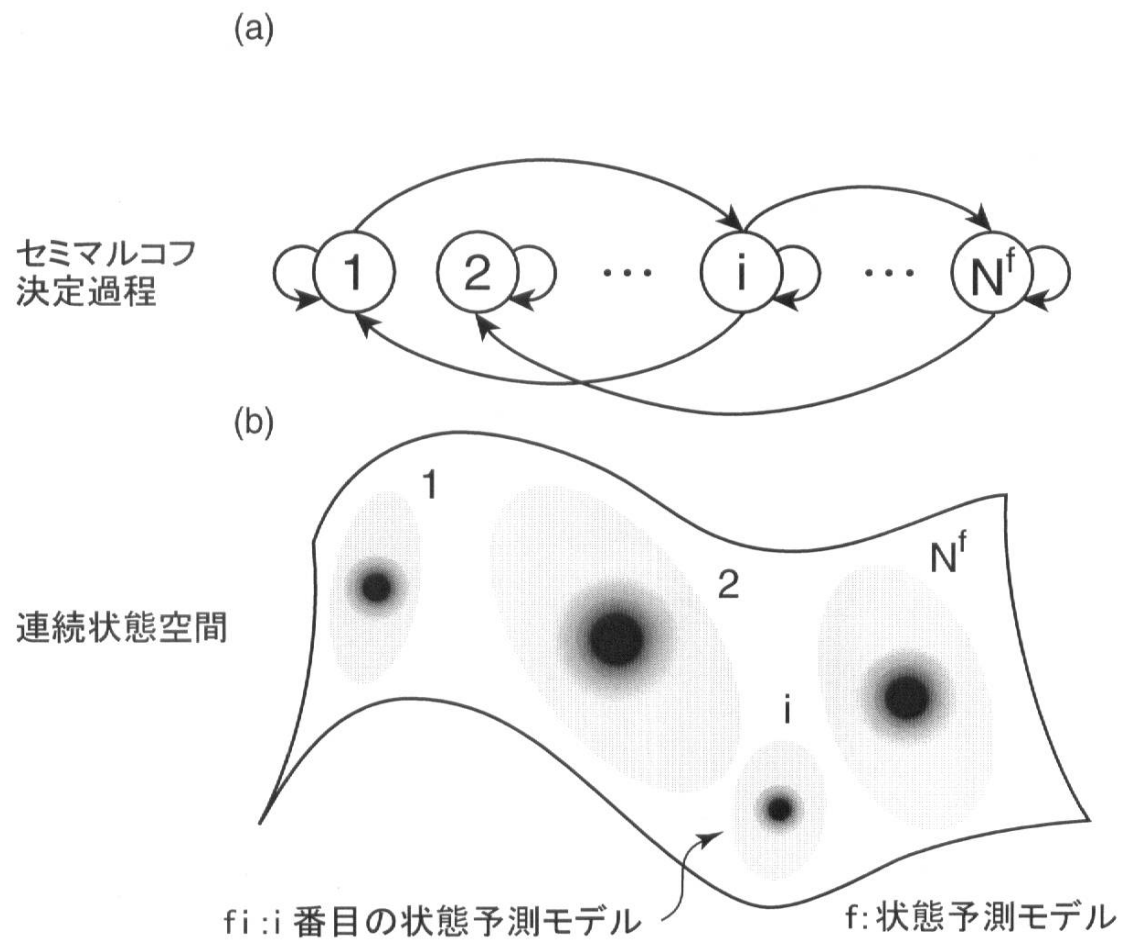
【図 4】



【図 5】



【図 1】



フロントページの続き

(72)発明者 銅谷 賢治

京都府相楽郡精華町光台二丁目2番地2 株式会社国際電気通信基礎技術研究所内

(72)発明者 川人 光男

京都府相楽郡精華町光台二丁目2番地2 株式会社国際電気通信基礎技術研究所内

審査官 長谷川 篤男

(56)参考文献 森本 淳、銅谷 賢治、階層型強化学習を用いた3リンク2関節ロボットによる起立運動の獲得
，日本ロボット学会誌，2001年 7月15日，第19巻 第5号，第32-37頁

鮫島 和行、片桐 憲一、銅谷 賢治、川人 光男、複数の予測モデルを用いた強化学習による
非線形制御，電子情報通信学会論文誌，2001年 9月20日，第J84-D-II巻 第9号，第2092
-2106頁

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，2
003年 3月15日，第21巻 第2号，第38-45頁

鮫島 和行、銅谷 賢治、川人 光男、強化学習MOSAIC：予測性によるシンボル化と見ま
ね学習，日本ロボット学会誌，2001年 7月15日，第19巻 第5号，第9-14頁

(58)調査した分野(Int.Cl.，DB名)

G06N 3/00

JSTPlus(JDreamII)