

(19) 日本国特許庁(JP)

(12) 特 許 公 報(B2)

(11) 特許番号

特許第4811997号
(P4811997)

(45) 発行日 平成23年11月9日(2011.11.9)

(24) 登録日 平成23年9月2日(2011.9.2)

(51) Int.Cl.

F I

G O 5 B 13/02 (2006.01)

G O 5 B 13/02

L

G O 5 B 13/04 (2006.01)

G O 5 B 13/04

請求項の数 4 (全 17 頁)

(21) 出願番号 特願2005-320988 (P2005-320988)
 (22) 出願日 平成17年11月4日(2005.11.4)
 (65) 公開番号 特開2007-128318 (P2007-128318A)
 (43) 公開日 平成19年5月24日(2007.5.24)
 審査請求日 平成20年8月20日(2008.8.20)

(73) 特許権者 503360115
 独立行政法人科学技術振興機構
 埼玉県川口市本町四丁目1番8号
 (73) 特許権者 393031586
 株式会社国際電気通信基礎技術研究所
 京都府相楽郡精華町光台二丁目2番地2
 (74) 代理人 100078868
 弁理士 河野 登夫
 (74) 代理人 100114557
 弁理士 河野 英仁
 (72) 発明者 森本 淳
 埼玉県川口市本町四丁目1番8号 独立行政法人科学技術振興機構内

最終頁に続く

(54) 【発明の名称】 状態推定装置、状態推定システム及びコンピュータプログラム

(57) 【特許請求の範囲】

【請求項1】

観測対象の状態を観測した観測結果に基づいて前記観測対象の状態を推定する状態推定装置において、

前記観測対象の状態を推定する模擬モデルと、

該模擬モデルによる状態の推定結果に基づいて、観測結果を推定した推定観測結果を算出する手段と、

推定観測結果及び観測結果、並びに推定観測結果及び観測結果の差を用いて、強化学習モジュールにおける状態推定の方策に基づくフィードバック値を算出する強化学習モジュールと、

推定観測結果及び観測結果の差並びにフィードバック値に基づいて、報酬値を算出する手段と、

算出した報酬値を用いて強化学習モジュールの方策を更新する更新手段と、

前記強化学習モジュールにて算出されたフィードバック値に基づいて、前記模擬モデルの観測対象の状態を推定する手段と

を備え、

前記更新手段は、報酬値に応じて更新される学習パラメータに基づく平均値及び標準偏差にて示される正規分布に従って分布するように方策を更新する

ことを特徴とする状態推定装置。

【請求項2】

前記学習パラメータは、報酬値の移動平均を用いた関数に基づいて更新される様に構成してあることを特徴とする請求項 1 に記載の状態推定装置。

【請求項 3】

観測対象と、
該観測対象の状態を推定する請求項 1 又は請求項 2 に記載の状態推定装置と、
前記観測対象を制御する制御装置と
を備え、
前記状態推定装置の模擬モデルは、前記観測対象の状態の推定結果を前記制御装置へ出力する手段を更に備え、
前記制御装置は、
受け付けた推定結果に基づいて、観測対象を制御する制御命令を生成する手段と、
生成した制御命令を前記観測対象へ出力する手段と
を備え、
前記観測対象は、受け付けた制御命令に従って動作する手段を備える
ことを特徴とする状態推定システム。

10

【請求項 4】

観測対象の状態を観測した観測結果の入力を受け付けるコンピュータに、受け付けた観測結果に基づいて、前記制御対象の状態を推定させるコンピュータプログラムにおいて、
コンピュータに、前記観測対象の状態を模する模擬モデルを用いて、前記観測対象の状態を推定させる手順と、
コンピュータに、前記模擬モデルによる状態の推定結果に基づいて、観測結果を推定した推定観測結果を算出させる手順と、
コンピュータに、推定観測結果及び観測結果、並びに推定観測結果及び観測結果の差を用いて、強化学習モジュールにおける状態推定の方策に基づくフィードバック値を算出させる手順と、
コンピュータに、推定観測結果及び観測結果の差並びにフィードバック値に基づいて、報酬値を算出させる手順と、
コンピュータに、算出した報酬値を用いて強化学習モジュールの方策を更新させる手順と、
コンピュータに、前記強化学習モジュールにて算出されたフィードバック値に基づいて、前記模擬モデルの観測対象の状態を推定させる手順と
を実行させ、
前記更新させる手順は、報酬値に応じて更新される学習パラメータに基づく平均値及び標準偏差にて示される正規分布に従って分布するように方策を更新する
ことを特徴とするコンピュータプログラム。

20

30

【発明の詳細な説明】

【技術分野】

【0001】

本発明は、観測対象の状態を観測した観測値に基づいて前記観測対象の状態を推定する状態推定装置、該状態推定装置を用いた状態推定システム、及び前記状態推定装置を実現するためのコンピュータプログラムに関し、特に線形又は非線形な動きを示し、その動きを観測することで観測対象の状態を推定する状態推定装置、状態推定システム及びコンピュータプログラムに関する。

40

【背景技術】

【0002】

ロボット等の制御対象を制御する場合、制御対象に組み込まれた各種センサがノイズにより、又はセンサそのものを組み込むことができないことにより、必要な状態変数を直接測定することができない状態が生じうる。また例えば市場予測を行う場合においても、価格変動の要因となる株価等の変数は直接観測することができない。

【0003】

50

状態変数を直接観測することができない場合、状態観測器（オブザーバ）（非特許文献 1 参照）、カルマンフィルタ（非特許文献 2、3 参照）等の方法が一般に用いられている。

【非特許文献 1】ディー・ジー・ルエンバーガー（D.G.Luenberger）、An introduction to observers, IEEE Trans., AC, Vol. 16, P. 596-602, 1971

【非特許文献 2】アール・イー・カルマン、アール・エス・ビューシー（R.E.Kalman and R.S.Bucy）、New results in linear filtering and prediction theory, Trans., ASME, Series D, J. of Basic Engineering, Vol. 83, No. 1, P. 95-108, 1961

【非特許文献 3】エフ・エル・ルイス（F.L.Lewis）、Optimal Estimation: with an Introduction to Stochastic Control Theory, John Wiley & Sons, 1977

10

【発明の開示】

【発明が解決しようとする課題】

【0004】

しかしながら非特許文献 1 の状態観測器及び非特許文献 2、3 のカルマンフィルタでは、対象のダイナミクスが非線形である場合に、隠れ状態の推定が困難になるという問題がある。

【0005】

本発明は斯かる事情に鑑みてなされたものであり、観測対象の状態を模する模擬モデル、模擬モデルによる観測対象の状態の推定結果等に基づき状態推定の方策を示すフィードバック値を算出する強化学習モジュール、強化学習モジュールが算出したフィードバック値及び観測対象の観測結果等に基づき報酬値を算出する報酬関数等のプログラムモジュール及び関数を用い、強化学習モジュールでは、フィードバック値が適正な値になる様に報酬値に基づく強化学習を行うことで、線形のダイナミクスだけでなく、非線形のダイナミクスへも容易に適用することが可能な状態推定装置、該状態推定装置を用いた状態推定システム、及び前記状態推定装置を実現するためのコンピュータプログラムの提供を目的とする。

20

【課題を解決するための手段】

【0008】

第 1 発明に係る状態推定装置は、観測対象の状態を観測した観測結果に基づいて前記観測対象の状態を推定する状態推定装置において、前記観測対象の状態を推定する模擬モデルと、該模擬モデルによる状態の推定結果に基づいて、観測結果を推定した推定観測結果を算出する手段と、推定観測結果及び観測結果、並びに推定観測結果及び観測結果の差を用いて、強化学習モジュールにおける状態推定の方策に基づくフィードバック値を算出する強化学習モジュールと、推定観測結果及び観測結果の差並びにフィードバック値に基づいて、報酬値を算出する手段と、算出した報酬値を用いて強化学習モジュールの方策を更新する更新手段と、前記強化学習モジュールにて算出されたフィードバック値に基づいて、前記模擬モデルの観測対象の状態を推定する手段とを備え、前記更新手段は、報酬値に応じて更新される学習パラメータに基づく平均値及び標準偏差にて示される正規分布に従って分布するように方策を更新することを特徴とする。

30

【0009】

本発明では、強化学習モジュールが、フィードバック値が適正な値になる様に報酬値に基づく強化学習を行うことで、線形のダイナミクスだけでなく、非線形のダイナミクスへも容易に適用することが可能である。

40

【0011】

本発明では、状態推定の方策を示すフィードバック値を、学習パラメータに基づく分布関数に従って様々な方策をとるべく展開させることにより、模擬モデルの推定観測結果が観測対象の観測結果に近づく様に、分布関数の平均値及び標準偏差が更新されながら強化学習が繰り返されるので、観測対象の実際の状態を正確に推定することができる様に、フィードバック値が収束していき、非線形のダイナミクスへも容易に展開することが可能である。

50

【 0 0 1 2 】

第2発明に係る状態推定装置は、第1発明において、前記学習パラメータは、報酬値の移動平均を用いた関数に基づいて更新される様に構成してあることを特徴とする。

【 0 0 1 3 】

本発明では、強化学習を用いて学習パラメータを更新することにより、ある瞬間での推定誤差を小さくするのではなく、タスクを行っている期間を通じての推定誤差を小さくし、模擬モデルと強化学習モジュールにより推定される状態を適切に観測対象の状態に近付けることが可能である。

【 0 0 1 4 】

第3発明に係る状態推定システムは、観測対象と、該観測対象の状態を推定する第1発明又は第2発明に記載の状態推定装置と、前記観測対象を制御する制御装置とを備え、前記状態推定装置の模擬モデルは、前記観測対象の状態の推定結果を前記制御装置へ出力する手段を更に備え、前記制御装置は、受け付けた推定結果に基づいて、観測対象を制御する制御命令を生成する手段と、生成した制御命令を前記観測対象へ出力する手段とを備え、前記観測対象は、受け付けた制御命令に従って動作する手段を備えることを特徴とする。

10

【 0 0 1 5 】

本発明では、強化学習モジュールが、フィードバック値が適正な値になる様に報酬値に基づく強化学習を行うことで、線形のダイナミクスだけでなく、非線形のダイナミクスへも容易に適用することが可能であり、しかも強化学習モジュールの学習結果に基づき模擬モデルにて推定された推定結果に基づいて、観測対象を制御することにより、ロボットの制御等の様々な分野への展開が可能である。

20

【 0 0 1 6 】

第4発明に係るコンピュータプログラムは、観測対象の状態を観測した観測結果の入力を受け付けるコンピュータに、受け付けた観測結果に基づいて、前記制御対象の状態を推定させるコンピュータプログラムにおいて、コンピュータに、前記観測対象の状態を模する模擬モデルを用いて、前記観測対象の状態を推定させる手順と、コンピュータに、前記模擬モデルによる状態の推定結果に基づいて、観測結果を推定した推定観測結果を算出させる手順と、コンピュータに、推定観測結果及び観測結果、並びに推定観測結果及び観測結果の差を用いて、強化学習モジュールにおける状態推定の方策に基づくフィードバック値を算出させる手順と、コンピュータに、推定観測結果及び観測結果の差並びにフィードバック値に基づいて、報酬値を算出させる手順と、コンピュータに、算出した報酬値を用いて強化学習モジュールの方策を更新させる手順と、コンピュータに、前記強化学習モジュールにて算出されたフィードバック値に基づいて、前記模擬モデルの観測対象の状態を推定させる手順とを実行させ、前記更新させる手順は、報酬値に応じて更新される学習パラメータに基づく平均値及び標準偏差にて示される正規分布に従って分布するように方策を更新することを特徴とする。

30

【 0 0 1 7 】

本発明では、汎用コンピュータ等のコンピュータにて実行することにより、コンピュータが状態推定装置として動作し、強化学習モジュールが、フィードバック値が適正な値になる様に報酬値に基づく強化学習を行うことで、線形のダイナミクスだけでなく、非線形のダイナミクスへも容易に適用することが可能である。

40

【 発明の効果 】

【 0 0 1 8 】

本発明に係る状態推定装置、状態推定システム及びコンピュータプログラムは、様々な装置、自然現象、更には経済現象等の線形又は非線形な動きを示し、その状態の一部又は全部を観測することが可能な観測対象と、該観測対象の状態を観測した観測結果に基づいて前記観測対象の状態を推定する状態推定装置とを備え、特に制御可能な様々な装置等の観測対象を観測する場合に、該観測対象を制御する制御装置を更に備える。そして前記観測対象の状態を模する模擬モデルを用いて、前記観測対象の状態を推定し、前記模擬モデ

50

ルによる状態の推定結果に基づいて、観測結果を推定した推定観測結果を算出し、推定観測結果及び観測結果、並びに推定観測結果及び観測結果の差を用いて推定結果を評価した報酬値に基づいて、状態推定の方策を示すフィードバック値を算出し、推定観測結果及び観測結果の差並びにフィードバック値に基づいて、報酬値を算出し、フィードバック値に基づいて、前記模擬モデルの観測対象の状態推定方法を更新する。

【0019】

この構成により、本発明では、強化学習モジュールが、フィードバック値が適正な値になる様に報酬値に基づく強化学習を行うことで、線形のダイナミクスだけでなく、非線形のダイナミクスへも容易に適用することが可能である等、優れた効果を奏する。

【0020】

さらに前記強化学習モジュールは、報酬値に応じて更新される学習パラメータに基づく平均値及び標準偏差にて示される正規分布に従って分布するフィードバック値を算出する様に構成してあり、報酬値の移動平均を用いた関数に基づいて学習パラメータを更新する様に構成してある。

【0021】

この構成により、本発明では、強化学習に基づいて学習パラメータを更新するため、ある瞬間での推定誤差を小さくするのではなく、タスクを行っている期間を通じての推定誤差を小さくし、模擬モデルと強化学習モジュールにより推定される状態を適切に観測対象の状態に近付けることが可能である等、優れた効果を奏する。

【発明を実施するための最良の形態】

【0022】

以下、本発明をその実施の形態を示す図面に基づいて詳述する。図1は、本発明の状態推定方法の構成例を概念的に示すブロック図である。図1中1は、観測対象(Plant)であり、本発明の状態推定方法は、観測対象1の状態を推定することを目的とする。観測対象1は、様々な装置、自然現象、更には経済現象等の線形又は非線形な動きを示し、その状態の一部又は全部を観測することが可能である。なお本実施の形態では、観測対象1は、制御装置(Controller)2により制御することが可能な装置として以降の説明を行う。また観測対象1の状態は、強化学習状態推定装置(RLSE: Reinforcement Learning State Estimator)3により推定される。

【0023】

強化学習状態推定装置3は、例えば汎用コンピュータにて構成されており、観測対象1の状態を推定する模擬モデル(Plant Model)3aとして機能するモジュールと、模擬モデル3aに正確な状態を推定させるべく状態推定のための方策を出力する強化学習モジュール(Reinforcement Learning Module)3bとを有している。強化学習モジュール3bは本発明の状態推定方法において主要な処理を行うプログラムモジュールであり、後述する様々な関数、定数、その他設定値を用いて強化学習を行い、模擬モデル3aの動きを観測対象1に近付ける方策を展開する。

【0024】

さらに強化学習状態推定装置3は、模擬モデル3aが推定した状態から観測結果を推定した推定観測結果を算出する出力関数3c、模擬モデル3aの推定結果を評価する報酬値を算出する報酬関数3d等の様々な関数を有している。

【0025】

この様に構成される本発明の状態推定方法において、観測対象1は、制御装置2からの制御命令に基づく制御により動作し、また観測対象1の状態を観測した観測結果は、強化学習状態推定装置3へ入力される。

【0026】

強化学習状態推定装置3の模擬モデル3aは、設定されている推定のためのモデルに従って観測対象1の状態を推定し、観測対象1の推定結果として制御装置2へ出力する。なお模擬モデル3aは、強化学習モジュール3bにおける状態推定の方策から出力されたフ

10

20

30

40

50

ィードバック値の入力に基づいて、状態の推定を行う。そして制御装置 2 は、入力を受け付けた推定結果に基づいて観測対象 1 を制御する制御命令を生成し、生成した制御命令を観測対象 1 及び強化学習状態推定装置 3 へ出力する。

【 0 0 2 7 】

なお模擬モデル 3 a から出力される推定結果は、強化学習状態推定装置 3 内の出力関数 3 c 及び強化学習モジュール 3 b でも用いられる。強化学習状態推定装置 3 では、模擬モデル 3 a により推定した観測対象の状態を示す推定結果から、出力関数 3 c により、観測対象 1 の観測結果を推定した推定観測結果を算出し、算出した推定観測結果を報酬関数 3 d へ渡す。

【 0 0 2 8 】

報酬関数 3 d は、観測対象 1 の観測結果、及び出力関数 3 c にて算出された推定観測結果、並びに強化学習モジュール 3 b における状態推定の方策から出力されたフィードバック値に基づいて推定を行い、その結果を評価した報酬値を算出し、算出した報酬値を強化学習モジュール 3 b へ渡す。

【 0 0 2 9 】

強化学習モジュール 3 b は、観測対象 1 の観測結果及び制御装置 2 から出力された制御命令、並びに模擬モデル 3 a による推定結果及び報酬関数 3 d にて算出した報酬値に基づいて、状態推定の方策を更新し、更新した方策から出力されるフィードバック値を模擬モデル 3 a 及び報酬関数 3 d へ渡す。

【 0 0 3 0 】

図 2 は、本発明の強化学習状態推定装置 3 の構成例を示すブロック図である。汎用コンピュータ等のコンピュータを用いた強化学習状態推定装置 3 は、装置全体を制御する CPU 等の制御手段 3 1 と、本発明の強化学習状態推定装置用のコンピュータプログラム 3 0 0 及びデータ等の各種情報を記録した CD - ROM 等の記録媒体 3 0 1 から各種情報を読み取る CD - ROM ドライブ等の補助記憶手段 3 2 と、補助記憶手段 3 2 により読み取った各種情報を記録するハードディスク等の記録手段 3 3 と、制御手段 3 1 の制御により一時的に発生するデータを記憶する RAM 等の記憶手段 3 4 とを備えている。そして記録手段 3 3 に記録したコンピュータプログラム 3 0 0 を記憶手段 3 4 に記憶して制御手段 3 1 の制御により実行することで、汎用コンピュータは、本発明の強化学習状態推定装置 3 として動作する。さらに強化学習状態推定装置 3 は、観測対象 1 の観測結果及び制御装置 2 の制御命令の入力を受け付ける入力手段 3 5 並びに制御装置 2 へ推定結果を出力する出力手段 3 6 を備えている。

【 0 0 3 1 】

さらに強化学習状態推定装置 3 が備える記録手段 3 3 には、模擬モデル 3 a、強化学習モジュール 3 b、出力関数 3 c、報酬関数 3 d 等の様々なプログラムモジュール及び関数が記録されている。

【 0 0 3 2 】

次に本発明の強化学習状態推定装置 3 を用いた状態推定方法について説明する。本発明の状態推定方法は、観測対象 1 の推定すべき真の状態 x 及び観測対象 1 の観測結果（出力値） y を用いた下記の式 1 ～ 式 3 にて示される。

【 0 0 3 3 】

10

20

30

40

【数 1】

$$\dot{x} = f(x, u) + n(t) \quad \dots \text{式 1}$$

但し、 $\dot{x}$: 観測対象 1 の真の状態の時間微分(= dx/dt)

$f(x, u)$: 観測対象 1 のダイナミクス

x : 観測対象 1 の真の状態

u : 観測対象 1 のダイナミクスへの制御入力

$n(t)$: 時刻 t の関数として示されるノイズ入力

10

$$y = h(x) + v(t) \quad \dots \text{式 2}$$

但し、 y : 観測対象 1 の観測結果 (出力値)

$h(x)$: 状態 x に基づき観測結果 y を決定する出力関数

$v(t)$: 時刻 t の関数として示される観測ノイズ

$$\dot{\hat{x}} = f(\hat{x}, u) + a(\hat{x}, y) \quad \dots \text{式 3}$$

但し、 $\dot{\hat{x}}$: 観測対象 1 の推定状態 $\hat{x}$ の時間微分(= $d\hat{x}/dt$)

$f(\hat{x}, u)$: 観測対象 1 の推定状態 $\hat{x}$ を用いて示されるダイナミクス

$a(\hat{x}, y)$: 観測対象 1 の推定状態 $\hat{x}$ 及び観測結果 y に基づき決定される
状態推定の方策に基づくフィードバック値

20

【0034】

なお上記式 1 において、ノイズ入力 $n(t)$ は、下記の式 4 となる性質を有している。

【0035】

【数 2】

$$E[n(t)n^T(\tau)] = U\delta(t-\tau) \quad \dots \text{式 4}$$

30

但し、 $E[n(t)n^T(\tau)]$: プロセスノイズの共分散行列

Q : プロセスノイズの大きさを示す行列

$\delta()$: ディラックのデルタ関数

τ : 時刻 (t とは異なる)

【0036】

また上記式 2 において、観測ノイズ $v(t)$ は、下記の式 5 となる性質を有している。

40

【0037】

【数 3】

$$E[v(t)v^T(\tau)] = S\delta(t-\tau) \quad \dots \text{式 5}$$

但し、 $E[v(t)v^T(\tau)]$: 観測ノイズの共分散行列
 S : 観測ノイズの大きさを示す行列

10

【0038】

上記式 1 において、観測対象 1 のダイナミクス $f(x, u)$ は、観測対象 1 の観測結果 y に基づいて、観測対象 1 の推定状態 $\hat{x}$ (以降の文章中において、観測対象 1 の推定状態 (推定結果) を示す記号を $\hat{x}$ と表記する。) と真の状態 x との差異を小さくすべく強化学習することにより取得する関数である。但し、予め既知のダイナミクス $f(x, u)$ を設定しておく様にしても良い。なお本発明の状態推定方法では、模擬モデル 3 a により推定した観測対象の状態 x の推定結果 $\hat{x}$ から算出した推定観測結果 $\hat{y}$ (以降の文章中において、推定観測結果を示す記号を $\hat{y}$ と表記する。) と、観測対象 1 の実際の観測結果 y との差異が小さい程、高い評価となる下記の式 6 の報酬関数を用いることにより、推定結果 $\hat{x}$ を実際の状態 x に近付けるべく強化学習を行う。

20

【0039】

【数 4】

$$r(\hat{x}, y, a) = e(|y - h(\hat{x})|) - c(a) \quad \dots \text{式 6}$$

但し、 $r(\hat{x}, y, a)$: 推定結果 $\hat{x}$ 、観測結果 y 及びフィードバックパラメータ a に
 基づく報酬関数 3 d

$e(|y - h(\hat{x})|)$: 推定結果の正確さを評価する関数

$h(\hat{x})$: 推定状態 $\hat{x}$ から推定観測結果 $\hat{y}$ を算出する出力関数 3 c

$c(a)$: 方策の出力 a に対するコスト関数

a : 模擬モデル 3 a の状態推定の方策に基づくフィードバック値

30

【0040】

式 6 中の推定結果の正確さを評価する関数 $e()$ は、実際の観測結果 y と、出力関数 3 c との差異が小さい程、大きい値をとり、関数 $c(a)$ は、強化学習モジュール 3 b における方策の出力の大きさが大きい程、大きい値をとるように設定される。

【0041】

強化学習モジュール 3 b では、式 6 に示す報酬関数 r に基づく学習処理を繰り返すことで、模擬モデル 3 a の状態推定の方策を更新し、更新した状態推定の方策からフィードバック値 a を算出し、模擬モデル 3 a へ出力する。なお状態推定の方策を示すフィードバック値 a は、不明な状態から下記の式 7 に示す確率分布に従って様々な方策をとるべく展開され、模擬モデル 3 a の推定観測結果 $\hat{y}$ が観測対象 1 の観測結果 y に近づく様に、強化学習の繰り返しのて収束する。

40

【0042】

【数 5】

$$\pi_w(a_j | \hat{x}, y) = \frac{1}{\sqrt{2\pi}\sigma_j} \exp \frac{-(a_j - \mu_j(\hat{x}, y))^2}{2\sigma_j^2} \quad \dots \text{式 7}$$

但し、 $\pi_w(\cdot)$: 方策を表す確率分布

j : 状態推定の方策の出力回数を示す自然数

μ : 状態推定の方策を示すフィードバックパラメータ a の平均値

σ^2 : 分散

10

【0043】

式 7 にて示される様にフィードバック値 a は、平均値 μ 、分散 σ^2 の正規分布を示し、分散 σ^2 が小さい程、方策のバラツキが小さくなる。

【0044】

式 7 中の平均値 μ は、下記の式 8 ~ 式 10 にて示され、学習すべきパラメータを更新していくことにより学習が行われる。

【0045】

【数 6】

20

$$\mu_j(\hat{x}, y) = v_j(\tilde{x})(y - h(\hat{x})) \quad \dots \text{式 8}$$

但し、 $v_j(\tilde{x})$: 下記の式 9 にて示されるゲイン関数

$\tilde{x}$: 下記の式 10 にて示される強化学習モジュール 3 b の状態

$$v_j(\tilde{x}) = \sum_{i=1}^N w_i^{\mu_j} b_i(\tilde{x}) \quad \dots \text{式 9}$$

但し、 $b_i(\tilde{x})$: 基底関数

$w_i^{\mu_j}$: 平均値 μ に関する学習すべきパラメータ

$$\tilde{x} = (\hat{x}^T, (y - h(\hat{x}))^T)^T \quad \dots \text{式 10}$$

30

【0046】

式 7 中の分散 σ^2 の平方根である標準偏差 σ は、下記の式 11 にて示され、学習すべきパラメータを更新していくことにより学習が行われる。

【0047】

【数 7】

$$\sigma_j = \frac{1}{1 + \exp(-w_j^\sigma)} \quad \dots \text{式 11}$$

40

但し、 w_j^σ : 標準偏差 σ に関する学習すべきパラメータ

【0048】

一方、強化学習においては、式 7 及び式 8 のパラメータを更新するために下記の式 12 にて示される価値関数 Q が用いられる。

【0049】

【数 8】

$$Q_w^v(\hat{x}, y, a) = \sum_{i,j} v_i^{\mu_j} \phi_i^{\mu_j}(\hat{x}, y, a) + \sum_j v^{\sigma_j} \phi^{\sigma_j}(\hat{x}, y, a) \quad \dots \text{式 1 2}$$

但し、Q: 価値関数

 $v_i^{\mu_j}$: 価値関数を表現するパラメータ $\phi_i^{\mu_j}$: 下記の式 1 3 にて示される関数 v^{σ_j} : 価値関数を表現するパラメータ ϕ^{σ_j} : 下記の式 1 4 にて示される関数

$$\phi_i^{\mu_j}(\hat{x}, y, a) = \frac{\partial \ln \pi_w}{\partial w_i^{\mu_j}} = \frac{a_j - \mu_j}{\sigma_j^2} (y - h(\hat{x})) b_i(\tilde{x}) \quad \dots \text{式 1 3}$$

$$\phi^{\sigma_j}(\hat{x}, y, a) = \frac{\partial \ln \pi_w}{\partial w^{\sigma_j}} = \frac{(a_j - \mu_j)^2 - \sigma_j^2}{\sigma_j^2} (1 - \sigma_j) \quad \dots \text{式 1 4}$$

10

【0050】

次に学習すべきパラメータの更新方法について説明する。強化学習状態推定装置 3 の強化学習モジュール 3 b では、報酬関数 3 d にて算出された報酬値 $R(t)$ に基づき、予め設定されている時定数 τ を用いた直近の移動平均値を下記の式 1 5 にて計算する。

【0051】

【数 9】

$$\tau \dot{R}(t) = -R(t) + r(\hat{x}(t), y(t), a(t)) \quad \dots \text{式 1 5}$$

但し、 τ : 移動平均値を算出する期間を示す時定数 r : 式 6 から計算される報酬

30

【0052】

また平均報酬の近似誤差 $\Delta(t)$ を下記の式 1 6 にて計算する。

【0053】

【数 10】

$$\Delta(t) = (r(\hat{x}(t), y(t), a(t)) - R(t)) + \dot{Q}(t) \quad \dots \text{式 1 6}$$

40

【0054】

式 1 6 の右辺の第 1 項は、報酬 r を式 1 5 にて示した報酬値 $R(t)$ の平均値の差であり、右辺の第 2 項に示した価値関数 $Q(t)$ の時間微分 (Q の上方にドットを付与した記号にパラメータ (t) を付して表記) が大きく近似誤差 $\Delta(t)$ が大きいほど、現在の状態推定の方策が正しいことになる。更に近似誤差 $\Delta(t)$ を修正するための価値関数 Q のパラメータの更新方法を下記の式 1 7 及び式 1 8 に示す。

【0055】

【数 1 1】

$$\dot{v}_i^{\mu j}(t) = \alpha^{\mu} \Delta(t) z_i^{\mu j}(t) \quad \dots \text{式 1 7}$$

但し、 α^{μ} :学習率 z :式 1 9 により算出される値(*eligibility trace*)

$$\dot{v}_i^{\sigma j}(t) = \alpha^{\sigma} \Delta(t) z_i^{\sigma j}(t) \quad \dots \text{式 1 8}$$

但し、 α^{σ} :学習率 z :式 2 0 により算出される値(*eligibility trace*)

10

【0 0 5 6】

また式 1 7 及び式 1 8 に示した関数の出力の時間に関する加重平均値の更新は下記の式 1 9 及び式 2 0 を用いて行われる。

【0 0 5 7】

【数 1 2】

$$\dot{z}_i^{\mu j}(t) = -\frac{1}{\kappa} z_i^{\mu j}(t) + \phi_i^{\mu j}(t) \quad \dots \text{式 1 9}$$

20

$$\dot{z}_i^{\sigma j}(t) = -\frac{1}{\kappa} z_i^{\sigma j}(t) + \phi_i^{\sigma j}(t) \quad \dots \text{式 2 0}$$

但し、 κ :*eligibility trace* の時定数

【0 0 5 8】

そして式 1 3 及び式 1 4 にて示した関数の出力を用い下記の式 2 1 及び式 2 2 により学習すべきパラメータを更新する。

【0 0 5 9】

【数 1 3】

30

$$\dot{w}_i^{\mu j} = \beta^{\mu} Q(\hat{x}(t), y(t), a_j(t)) \phi_i^{\mu j}(t) \quad \dots \text{式 2 1}$$

$$\dot{w}_i^{\sigma j} = \beta^{\sigma} Q(\hat{x}(t), y(t), a_j(t)) \phi_i^{\sigma j}(t) \quad \dots \text{式 2 2}$$

但し、 β^{μ} :学習率 β^{σ} :学習率

【0 0 6 0】

図 3 は、本発明の強化学習状態推定装置 3 の処理を示すフローチャートである。上述した様に本発明の強化学習状態推定装置 3 は、コンピュータプログラム 3 0 0 を実行する制御手段 3 1 の制御により、強化学習モジュール 3 b により、上述した様々な関数及び数式にて、状態推定の方策に基づくフィードバック値を算出し(ステップ S 1)、模擬モデル 3 a にて式 1 を用いて観測対象の状態を推定し(ステップ S 2)、模擬モデル 3 a による状態の推定結果に基づいて、出力関数 3 c により、観測結果を推定した推定観測結果を算出し(ステップ S 3)、式 6 を用いて推定観測結果及び観測結果の差並びに強化学習モジュール 3 b にて算出されるフィードバック値に基づいて報酬値を算出し(ステップ S 4)、算出した報酬値を用いて強化学習モジュール 3 b の方策を更新する(ステップ S 5)。そして強化学習状態推定装置 3 は、制御手段 3 1 の制御に基づいて、強化学習モジュール 3 b により、算出したフィードバック値に基づいて模擬モデル 3 a による観測対象の状態

40

50

を推定する（ステップ S 6）。

【実施例】

【 0 0 6 1 】

次に本発明の状態推定方法を用いて行った実験結果について説明する。図 4 は、本発明の状態推定方法の実験に用いた観測対象である単振り子を示す模式図である。観測対象である単振り子の重りの質量は m 、振り子長は l であり、外力 T を加えることにより支点 o を中心とする揺動を開始する。なお図 4 に示した単振り子の力学モデルは、下記の式 2 3 にて示される。

【 0 0 6 2 】

【数 1 4】

$$ml^2\ddot{\omega} = -\mu\omega - mgl\sin\theta + T \quad \cdots \text{式 2 3}$$

但し、 m : 重りの質量

l : 振り子長

ω : 角速度

g : 重力加速度

θ : 鉛直線と振り子とがなす角度

μ : 摩擦係数

T : 外力

【 0 0 6 3 】

なお当実験において、 $\mu = 0.01$ 、 $m = 1.0 \text{ kg}$ 、 $l = 1.0 \text{ m}$ 、そして $g = 9.8 \text{ m/s}^2$ である。

【 0 0 6 4 】

ここで観測対象である振り子の状態を示す状態ベクトルを $x = (\theta, \omega)^T$ とし、 $u = T$ とすると、上述した式 1 及び式 2 は、下記の式 2 4 及び式 2 5 として示される。

【 0 0 6 5 】

【数 1 5】

$$\dot{x} = \begin{pmatrix} \omega \\ -\frac{g}{l}\sin\theta - \frac{\mu}{ml^2}\omega \end{pmatrix} + \begin{pmatrix} 0 \\ \frac{1}{ml^2} \end{pmatrix} u + n(t) \quad \cdots \text{式 2 4}$$

$$y = Cx + v(t) \quad \cdots \text{式 2 5}$$

【 0 0 6 6 】

なお式 2 4 及び式 2 5 において、ノイズ $n(t)$ 及びノイズ $v(t)$ は、上述した式 4 及び式 5 に対し、 $U = \text{diag}\{0.01, 0.01\}$ 、 $S = 1.0$ を代入することにより表現することができる。

【 0 0 6 7 】

また上述した式 3 の観測対象の推定状態は、当実験において下記の式 2 6 にて示される。

【 0 0 6 8 】

10

20

30

40

【数 1 6】

$$\dot{\hat{x}} = \begin{pmatrix} \hat{\omega} \\ -\frac{g}{l} \sin \hat{\theta} - \frac{\mu}{ml^2} \hat{\omega} \end{pmatrix} + \begin{pmatrix} 0 \\ \frac{1}{ml^2} \end{pmatrix} u + a(\hat{x}, y) \quad \dots \text{式 2 6}$$

【0 0 6 9】

なお式 2 6 において、右辺の第 3 項は、状態推定のフィードバック値である。

10

【0 0 7 0】

そして上述した式 6 に示す報酬関数 r は、当実験において下記の式 2 7 にて示される。

【0 0 7 1】

【数 1 7】

$$r = \exp\left(-\frac{(y - C\hat{x})^2}{\sigma_r^2}\right) - \sum_{j=1}^2 c(a_j) \quad \dots \text{式 2 7}$$

20

【0 0 7 2】

なお当実験において $\sigma_r = 0.5$ である。また式 2 7 中において、知識の正確さに依存した関数 $c(a_j)$ は、フィードバック値 a の最大値である $a^{\max} = (a_1^{\max}, a_2^{\max}) = (5.0, 5.0)$ を用いた下記の式 2 8 にて示される。

【0 0 7 3】

【数 1 8】

$$c(a_j) = d \int_0^{a_j} s^{-1} \left(\frac{u}{a_j^{\max}} \right) du \quad \dots \text{式 2 8}$$

30

但し、 $d = 0.1$

【0 0 7 4】

当実験に際し、学習率は、下記の式 2 9 及び式 3 0 にて示される値を用いた。

【0 0 7 5】

【数 1 9】

$$\alpha^\mu = \alpha^\sigma = 2.5 \times 10^2 \quad \dots \text{式 2 9}$$

40

$$\beta^\mu = \beta^\sigma = 0.5 \quad \dots \text{式 3 0}$$

【0 0 7 6】

さらに当実験に際し、時定数 $\tau = 0.2 \text{ sec}$ 及び eligibility trace の時定数 $\kappa = 0.2 \text{ sec}$ との条件設定を行った。

【0 0 7 7】

また強化学習によるフィードバック値の更新周期は、 0.02 sec である。

50

【 0 0 7 8 】

図 5 乃至図 7 は、本発明の状態推定方法の実験結果を示すグラフである。なお図 5 乃至図 6 において (a) は、角度 θ のみが観測される場合、(b) は、角速度 ω のみが観測される場合、そして (c) は、角度及び角速度 ω の線形和 $\theta + \omega$ のみが観測される場合の結果を示している。図 5 は、上述した条件に基づく実験において、時間と観測値との関係を示しており、図 5 (a) は、角度 θ の観測値の経時変化を示しており、図 5 (b) は、角速度 ω の観測値の経時変化を示しており、そして図 5 (c) は、観測値が $\theta + \omega$ である場合の経時変化を示している。図 5 は、いずれも横軸に時間を取り、縦軸に観測値をとって、その関係を示している。図 5 に示す様に観測値はいずれも非線形に変動する。

【 0 0 7 9 】

図 6 は、横軸に時間を取り、縦軸に角度 θ をとって、実際の角度及び推定した角度の経時変化を示している。図 6 (a) は、図 5 (a) に示した角度の観測値に基づいて推定した角度と、実際の角度との関係を示している。この実験では、図 6 (a) に示す様に 1 秒程度で推定した角度が実際の角度に一致している。図 6 (b) は、図 5 (b) に示した角速度の観測値に基づいて推定した角度と、実際の角度との関係を示している。この実験では、図 6 (b) に示す様に 3 秒程度で推定した角度が実際の角度に一致している。図 6 (c) は、図 5 (c) に示した角度及び角速度の和の観測値に基づいて推定した角度と、実際の角度との関係を示している。この実験では、図 6 (c) に示す様に 2 秒程度で推定した角度が実際の角度に一致している。

【 0 0 8 0 】

図 7 は、横軸に時間を取り、縦軸に角速度 ω をとって、実際の角速度及び推定した角速度の経時変化を示している。図 7 (a) は、図 5 (a) に示した角度の観測値に基づいて推定した角速度と、実際の角速度との関係を示している。この実験では、図 7 (a) に示す様に 2 秒程度で推定した角速度が実際の角速度に一致している。図 7 (b) は、図 5 (b) に示した角速度の観測値に基づいて推定した角速度と、実際の角速度との関係を示している。この実験では、図 7 (b) に示す様に 2 秒程度で推定した角速度が実際の角速度に一致している。図 7 (c) は、図 5 (c) に示した角度及び角速度の和の観測値に基づいて推定した角速度と、実際の角速度との関係を示している。この実験では、図 7 (c) に示す様に 2 秒程度で推定した角度が実際の角度に一致している。

【 0 0 8 1 】

この様に本発明では、非線形のダイナミクスの状態を容易に推定することが可能である。

【 0 0 8 2 】

前記実施の形態では、単振り子を観測対象とする実験を示したが、本発明はこれに限らず、様々な装置、自然現象、更には経済現象等の線形又は非線形な動きを示す観測対象の状態推定に適用することが可能である。

【 図面の簡単な説明 】

【 0 0 8 3 】

【 図 1 】 本発明の状態推定方法の構成例を概念的に示すブロック図である。

【 図 2 】 本発明の強化学習状態推定装置の構成例を示すブロック図である。

【 図 3 】 本発明の強化学習状態推定装置の処理を示すフローチャートである。

【 図 4 】 本発明の状態推定方法の実験に用いた観測対象である単振り子を示す模式図である。

【 図 5 】 本発明の状態推定方法の実験結果を示すグラフである。

【 図 6 】 本発明の状態推定方法の実験結果を示すグラフである。

【 図 7 】 本発明の状態推定方法の実験結果を示すグラフである。

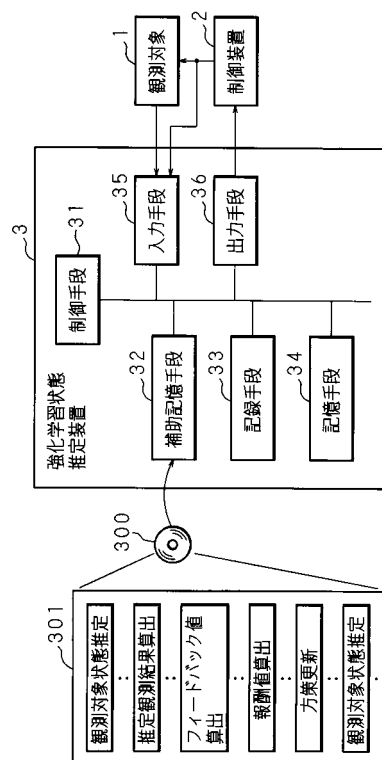
【 符号の説明 】

【 0 0 8 4 】

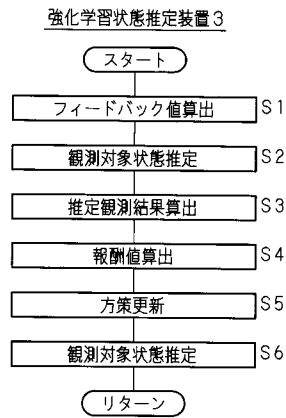
- 1 観測対象
- 2 制御装置

- 3 強化学習状態推定装置
- 3 a 模擬モデル
- 3 b 強化学習モジュール
- 3 c 出力関数
- 3 d 報酬関数
- 3 0 0 コンピュータプログラム
- 3 0 1 記録媒体

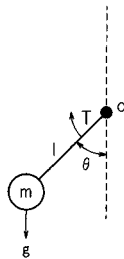
【圖 2】



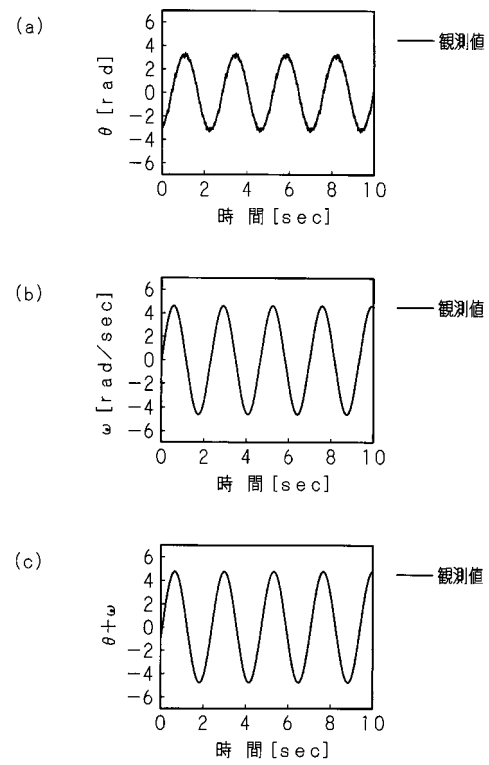
【図 3】



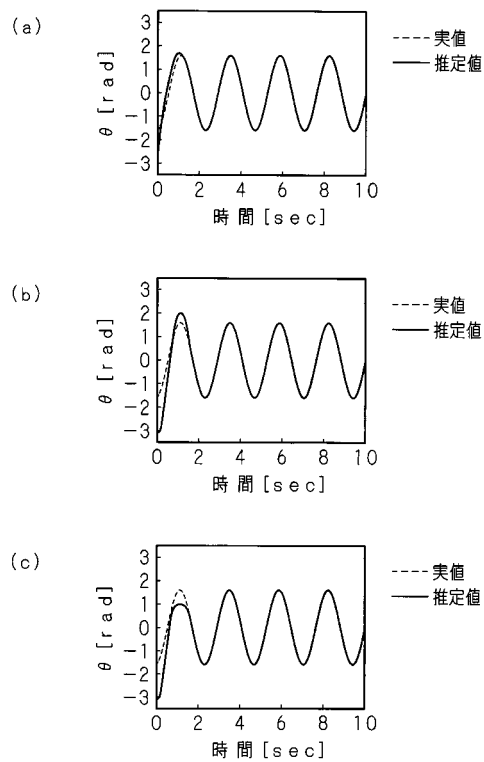
【図 4】



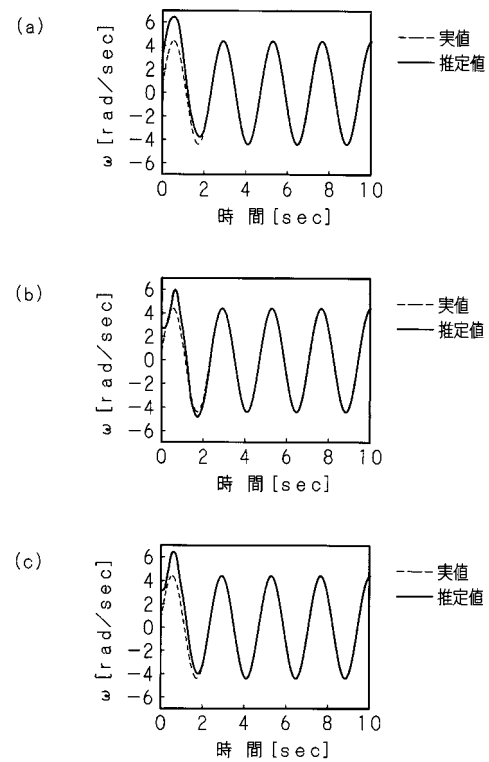
【図 5】



【図 6】



【図 7】



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

審査官 星名 真幸

(56)参考文献 森本 淳 Morimoto Jun, 銅谷 賢治 Doya Kenji, 強化学習を用いた状態観測器の構築 Development of An Observer by using Reinforcement Learning, 日本神経回路学会第12回全国大会講演論文集, 日本, 日本神経回路学会, 2002年 9月19日, pp.275-278

(58)調査した分野(Int.Cl., DB名)

G 0 5 B 1 3 / 0 2

G 0 5 B 1 3 / 0 4