

(19) 日本国特許庁 (JP)

(12) 特 許 公 報 (B2)

(11) 特許番号
特許第4630553号
(P4630553)

(45) 発行日 平成23年2月9日 (2011.2.9)

(24) 登録日 平成22年11月19日 (2010.11.19)

(51) Int. Cl.

F I

B 2 5 J 5/00 (2006.01)

B 2 5 J 13/00 (2006.01)

G O 6 N 3/00 (2006.01)

B 2 5 J 5/00 F

B 2 5 J 13/00 Z

G O 6 N 3/00 5 5 O E

請求項の数 5 (全 22 頁)

(21) 出願番号	特願2004-8130 (P2004-8130)	(73) 特許権者	000002185
(22) 出願日	平成16年1月15日 (2004.1.15)		ソニー株式会社
(65) 公開番号	特開2005-199383 (P2005-199383A)		東京都港区港南1丁目7番1号
(43) 公開日	平成17年7月28日 (2005.7.28)	(73) 特許権者	393031586
審査請求日	平成19年1月12日 (2007.1.12)		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	100064746
			弁理士 深見 久郎
		(74) 代理人	100085132
			弁理士 森田 俊雄
		(74) 代理人	100083703
			弁理士 仲村 義平
		(74) 代理人	100096781
			弁理士 堀井 豊

最終頁に続く

(54) 【発明の名称】 動的制御装置および動的制御装置を用いた2足歩行移動体

(57) 【特許請求の範囲】

【請求項1】

制御対象に対する制御信号を生成するための動的制御装置であって、
前記制御対象の状態情報を検知するためのセンサ群と、
周期的な時間発展を行なう内部状態を有し、前記制御対象の状態推定を行なって、前記
制御対象に対する前記制御信号を生成する制御手段を備え、
前記制御手段は、
前記センサ群から得られる前記状態情報のみで構成される状態空間を用いた方策勾配法
による強化学習を行い、前記制御対象から得られる前記状態情報と前記制御信号により規
定される報酬に基づく価値関数と前記強化学習中に得られる報酬系列とに基づいて、フィ
ードバックパラメータを更新し、更新された前記フィードバックパラメータにより規定さ
れる確率分布により出力値を決定するフィードバック制御器と、
前記周期的な時間発展を行なう前記内部状態を有し、前記フィードバック制御器からの
前記出力値に応じて変化する前記内部状態に基づき、前記内部状態に対応する目標値に対
するPDサーボ系の出力として、前記制御信号を生成するためのセントラルパターンジェ
ネレータとを含み、

前記制御信号に基づいて、前記制御対象を駆動するための駆動手段とを備える、動的制
御装置。

【請求項2】

前記確率分布は、前記フィードバックパラメータにより決定される平均値と分散とを有

する正規分布により表され、

前記価値関数は、正規化ガウスネットワークで表現される、請求項 1 記載の動的制御装置。

【請求項 3】

前記平均値は、前記フィードバックパラメータを重みとする正規化ガウス関数ネットワークで表現され、

前記分散は、前記フィードバックパラメータによるシグモイド関数として表現され、

前記フィードバックパラメータは、各前記フィードバックパラメータに対応する学習率とテンポラル・ディファレンス誤差とエリジビリティ・トレースとに応じて更新される、請求項 2 記載の動的制御装置。

10

【請求項 4】

前記フィードバック制御器は、前記出力値を所定のレベル以下に制限するための出力飽和手段を含む、請求項 2 記載の動的制御装置。

【請求項 5】

請求項 1 ~ 4 のいずれか 1 項に記載される動的制御装置により、2 足歩行制御が行なわれる、2 足歩行移動体。

【発明の詳細な説明】

【技術分野】

【0001】

本発明は、非線形の制御対象に関し、2 足歩行等の周期運動や、状態観測器を実現する上で有効となる動的制御装置と動的制御装置を用いた 2 足歩行移動体とに関する。

20

【背景技術】

【0002】

ロボットなどを制御しようとするとき、センサノイズや、そもそもセンサを配置することができないことによって、制御のために必要な状態変数を直接的には測定できない状況が考えられる。

【0003】

そのような場合、状態観測器（オブザーバ）（たとえば、非特許文献 1 を参照）やカルマンフィルタ（たとえば、非特許文献 2、非特許文献 3 を参照）を用いることが一般的である。しかしながら、対象のダイナミクスが非線形である場合、これらの手法では隠れ状態の推定は困難な場合がある。近年、非線形系に適用可能なオブザーバの提案がなされているものの、それぞれ特定の条件を満たさなければ適用できないなどの問題がある（たとえば、非特許文献 4、非特許文献 5 を参照）。

30

【0004】

また、拡張カルマンフィルタ（局所線形モデルを用いて状態分布の更新を行なう）（たとえば、非特許文献 3 を参照）やモンテカルロフィルタ（モンテカルロ法により生成した多数の粒子により状態分布を近似する）（たとえば、非特許文献 6 を参照）などの手法が非線形系での状態推定法として知られているが、いずれも状態分布を陽に求めなければならない。

【0005】

40

すなわち、従来の制御器は、基本的には、現在の状態を観測し、そこからの直接の写像によって制御出力を与える必要がある。

【非特許文献 1】D.G. Luenberger 著、“An introduction to observers”、IEEE Trans., AC, Vol.16, pp. 596--602, 1971.

【非特許文献 2】R.E. Kalman and R.S. Bucy 著、“New results in linear filtering and prediction theory”、Trans., ASME, Series D, J. of Basic Engineering, Vol.83, No.1, pp. 95--108, 1961.

【非特許文献 3】F.L. Lewis 著、“Optimal Estimation: with an Introduction to Stochastic Control Theory”、John Wiley & Sons, 1977.

【非特許文献 4】志水清孝，鈴木俊輔，田中哲史著、“こう配降下法による非線形オブザ

50

ーバ（非線形システムの状態観測器）”、電子情報通信学会論文誌 A, Vol. J83-A, No.8, pp. 956--964, 2000.

【非特許文献5】H.Nijmeijer and T.L. Fossen著、“Directions in Nonlinear Observer Design”、Springer-Verlag, London, 1999.

【非特許文献6】G.Kitagawa著、“Monte Carlo Filter and Smoother for Non-Gaussian Nonlinear State Models”、Journal of Computational and Graphical Statistics, Vol.5, pp. 1--25, 1996.

【発明の開示】

【発明が解決しようとする課題】

【0006】

本発明の目的は、制御器に内部状態とダイナミクスを持たせることで、周期運動に対し状態観測器の学習を容易とすることが可能な動的制御装置およびこのような動的制御装置を用いた2足歩行移動体を提供することである。

【課題を解決するための手段】

【0007】

本発明では、学習の目的は、ある瞬間での推定誤差を少なくするのではなく、タスクを行なっている期間を通じての推定誤差を少なくすることである。そこで、方策勾配法（強化学習）の枠組みを用いて非線形状態観測器の構築を行なっている。

【0008】

本発明の1つの局面では、状態観測器のダイナミクスへの入力を学習器の行動と考え、現在の観測可能な状態（観測出力）、状態観測器の状態、制御器の出力から、適切に状態観測器の状態を制御対象の状態へと導く行動則を方策勾配法によって獲得している。

【0009】

したがって、この発明の1つの局面に従うと、制御対象に対する制御信号を生成するための動的制御装置であって、制御対象の状態情報を検知するためのセンサ群と、周期的な時間発展を行なう内部状態を有し、制御対象の状態推定を行なって、制御対象に対する制御信号を生成する制御手段を備え、制御手段は、センサ群から得られる状態情報のみで構成される状態空間を用いた方策勾配法による強化学習を行い、制御対象から得られる状態情報と制御信号により規定される報酬に基づく価値関数と強化学習中に得られる報酬系列とに基づいて、フィードバックパラメータを更新し、更新されたフィードバックパラメータにより規定される確率分布により出力値を決定するフィードバック制御器と、周期的な時間発展を行なう内部状態を有し、フィードバック制御器からの出力値に応じて変化する内部状態に基づき、内部状態に対応する目標値に対するPDサーボ系の出力として、制御信号を生成するためのセントラルパターンジェネレータとを含み、制御信号に基づいて、制御対象を駆動するための駆動手段とを備える。

【0010】

好ましくは、確率分布は、フィードバックパラメータにより決定される平均値と分散とを有する正規分布により表され、価値関数は、正規化ガウスネットワークで表現される。

【0011】

好ましくは、平均値は、フィードバックパラメータを重みとする正規化ガウス関数ネットワークで表現され、分散は、フィードバックパラメータによるシグモイド関数として表現され、フィードバックパラメータは、各フィードバックパラメータに対応する学習率とテンポラル・ディファレンス誤差とエリジビリティ・トレースとに応じて更新される。

【0012】

好ましくは、フィードバック制御器は、出力値を所定のレベル以下に制限するための出力飽和手段を含む。

【0013】

好ましくは、この発明の2足歩行移動体は、上記動的制御装置により、2足歩行制御が行なわれる。

【発明の効果】

10

20

30

40

50

【 0 0 1 4 】

本発明では、行動則自体が内部変数とその微分方程式によって表わされるダイナミクスを持つため、行動則と物理系の引き込みの性質を利用することができ、周期運動に対して状態推定を行なう場合に、状態観測器の学習を容易とすることが可能である。さらに、センサ入力のノイズや時間遅れに対しても、ロバストな性質を持つ制御器を実現することができる。

【発明を実施するための最良の形態】

【 0 0 1 5 】

以下、図面を参照して本発明の実施の形態について説明する。

【 0 0 1 6 】

10

以下の説明では、一例として、本発明を2足歩行の制御に適用する場合を説明するが、本発明は、必ずしもこのような場合に限定されるものではなく、たとえば、より一般的に周期運動を行なう系に対して有効な制御システムを提供するものである。特に、本発明は、周期運動を行なう劣駆動（機械）系に適用するのに適した制御システムを提供する。

【 0 0 1 7 】

[実施の形態 1]

(本発明のシステム構成)

図1は、本発明の動的制御装置を用いた2足歩行移動システム1000の一例を示す概念図である。

【 0 0 1 8 】

20

図1を参照して、システム1000は、動的制御装置100と、動的制御装置100の上部に設けられる胴部40と、動的制御装置100により駆動制御される脚部とを備える。脚部は、右脚10rと左脚10lとを有し、各脚は、接地面近傍に設けられるセンサ20rおよび20lとを備える。一方、胴部40には、センサ30が設けられる。

【 0 0 1 9 】

センサ20rは、胴部40の中心線4に対する右脚の角度 θ_r 、角速度 $d\theta_r/dt$ という情報を検出し、また、センサ20lは、中心線4に対する左脚の角度 θ_l 、角速度 $d\theta_l/dt$ という情報を検出し、それぞれ、動的制御装置100に通知する。さらに、センサ30は、鉛直方向2に対する胴部40のピッチ角 θ_p 、角速度 $d\theta_p/dt$ という情報を検出し、それぞれ、動的制御装置100に通知する。

30

【 0 0 2 0 】

動的制御装置100は、センサ20r、センサ20l、センサ30からの情報に基づいて、右脚10rおよび左脚10lの動作を制御する。

【 0 0 2 1 】

図2は、図1に示した動的制御装置100の構成を示すブロック図である。

【 0 0 2 2 】

図2を参照して、動的制御装置100は、センサ20r、センサ20l、センサ30からの信号を受け取る通信インタフェース106と、後に説明する制御パラメータやセンサからの情報を格納しておくための記憶装置104と、センサ20r、センサ20l、センサ30からの情報を用いて学習して獲得した動的行動則に基づき、制御信号を生成する演算処理部102と、演算処理部102からの制御信号に基づいて、右脚10rおよび左脚10lの駆動制御を行なうための駆動部108とを備える。

40

【 0 0 2 3 】

以下では、動的制御装置100の制御動作のための準備の処理および制御動作について説明する。

【 0 0 2 4 】

(1 - 1 . 動的行動則)

まず、本発明の制御動作を説明する前提として、「動的行動則」について説明する。

【 0 0 2 5 】

図3は、このような動的行動則を説明するための概念図である。

50

【 0 0 2 6 】

「動的行動則」とは、図 3 (1) に示すような行動則自体が内部変数とその微分方程式によって表わされるダイナミクスを持つ枠組である。

【 0 0 2 7 】

これをより具体的に表現すると、図 3 (2) に示すように観測器が内部変数及びその微分方程式を持つようなものや、図 3 (3) に示すような制御器が内部変数及びその微分方程式を持つようなものである場合が考えられる。

【 0 0 2 8 】

このような「動的行動則」に基づいて、制御対象を制御するような制御装置を「動的制御装置」と呼ぶことにする。

10

【 0 0 2 9 】

以下では、まず、図 3 (3) の枠組を用いることで、動的制御装置 1 0 0 により、2 足歩行運動を実現する場合を考える。

【 0 0 3 0 】

行動則が内部状態を持つことによって、行動則と物理系の引き込みの性質を利用することが出来るため、周期運動や状態推定を行なう場合に有効であると考えられる。さらに、センサ入力ノイズや時間遅れに対しても、ある程度ロバストな性質を持つことが期待できる。

【 0 0 3 1 】

このような行動則を学習によって獲得する場合、行動則の内部状態が隠れ変数となることは問題となる。しかし、行動則の内部状態が物理系の状態に対して引き込むことで、それぞれの状態は一意に対応するようになる。ただし、過渡的な状態では隠れ状態を扱う必要がある。

20

【 0 0 3 2 】

そこで本発明では、動的行動則の獲得手法として、後に説明するように、隠れ変数が存在する環境においても適用可能な方策勾配法を用いる。

【 0 0 3 3 】

以下では、本発明の学習システムおよび、動的行動則を構成するセントラルパターンジェネレータ (central pattern generator : C P G) と C P G へのフィードバック制御器について説明を行なう。

30

【 0 0 3 4 】

(1 - 2 . 学習システム)

以下の説明では、周期運動の例として、3 リンク 2 足歩行ロボットモデルを用いた 2 足歩行運動に対して動的行動則を適用する。

【 0 0 3 5 】

図 4 は、図 2 で説明した演算処理部 1 0 2 の行なう処理を示す機能ブロック図である。以下に説明するとおり、演算処理部 1 0 2 は、学習システムとして機能する。

【 0 0 3 6 】

図 4 に示すとおり、この学習システムは、基本的に、C P G 処理部 1 0 2 6 とフィードバック制御器 1 0 2 2 によって動的行動則を構成する。学習に用いる状態 x は、以下の式で表わされる。

40

【 0 0 3 7 】

【数 1】

$$x = (x_1, x_2, x_3, x_4)^T = (\theta_l + \theta_p, \theta_r + \theta_p, \dot{\theta}_l + \dot{\theta}_p, \dot{\theta}_r + \dot{\theta}_p)^T$$

【 0 0 3 8 】

ただし、上述のとおり、 θ_r 、 θ_l は、それぞれロボットの鉛直方向からの左右の脚 1 0 r、1 0 l の角度であり、 θ_p は胴体 4 0 のピッチ角である。

【 0 0 3 9 】

50

つまり、学習システムはロボットから直接得られる信号のみで状態空間を構成しており、C P Gの内部状態を用いていないという特徴を有している。

【 0 0 4 0 】

また、ここでは全状態観測を仮定し、 $y = x$ であるとする。

【 0 0 4 1 】

(1 - 3 . セントラルパターンジェネレータ (C P G))

演算処理部 1 0 2 により実現される学習システムで、動的行動則を構成するC P G 処理部 1 0 2 2 の構成として、以下の式で表わされる神経振動子モデルを用いる。なお、このような神経振動子モデルについては、たとえば、文献：Kiyoshi Matsuoka著、“Sustained oscillations generated by mutually inhibiting neurons with adaptation.”、Biological Cybernetics, Vol.52, pp. 367-376, 1985に開示がある。

10

【 0 0 4 2 】

【数 2】

$$\tau \dot{z}_1 = -z_1 - \omega_{12}q_2 - \omega_{13}q_3 - \beta p_1 + z_0 + a_1 \quad (1)$$

$$\tau' \dot{p}_1 = -p_1 + q_1 \quad (2)$$

$$q_1 = \max(0, z_1) \quad (3)$$

$$\tau \dot{z}_2 = -z_2 - \omega_{21}q_1 - \omega_{24}q_4 - \beta p_2 + z_0 + a_2 \quad (4)$$

20

$$\tau' \dot{p}_2 = -p_2 + q_2 \quad (5)$$

$$q_2 = \max(0, z_2) \quad (6)$$

$$\tau \dot{z}_3 = -z_3 - \omega_{34}q_4 - \omega_{31}q_1 - \beta p_3 + z_0 + a_3 \quad (7)$$

$$\tau' \dot{p}_3 = -p_3 + q_3 \quad (8)$$

$$q_3 = \max(0, z_3) \quad (9)$$

30

$$\tau \dot{z}_4 = -z_4 - \omega_{43}q_3 - \omega_{42}q_2 - \beta p_4 + z_0 + a_4 \quad (10)$$

$$\tau' \dot{p}_4 = -p_4 + q_4 \quad (11)$$

$$q_4 = \max(0, z_4) \quad (12)$$

【 0 0 4 3 】

ここで、変数： z 、 p はニューロン内の状態、 q はニューロンの出力、 z_0 は持続入力、定数 β はニューロンの疲労係数、 τ 、 τ' は、 z 、 p の時定数、 ω は拮抗ニューロン間の結合係数である。また、 a は、後に説明するフィードバック制御器からの出力項である。

40

【 0 0 4 4 】

図5は、式(1)～(12)で表わされる神経振動子モデルによるC P Gを示す概念図である。図5においては、ニューロン内の状態 z 、 p の間で、相互に正の結合を行なうものは白丸で、負の結合を行なうものは黒丸で示している。

【 0 0 4 5 】

図6は、このようなC P Gの出力 q を構成する変数 z_1 、 z_2 の波形を示す図である。なお、この計算では、例として、 $\tau = 0.05$ 、 $\tau' = 0.6$ 、 $\beta = 2.5$ 、 $\omega = 2.0$ 、 $z_0 = 0.1$ 、 $a = 0$ を用いた。

【 0 0 4 6 】

50

式(1)～(12)で表わされるモデルにしたがって、CPGの内部状態変数 z_1 、 z_2 が周期的に変化していることがわかる。

【0047】

さらに、図4の学習システムのPDサーボ処理部1028では、以下に示すとおり、各ニューロンの出力の差を両脚のサーボ系の目標関節角 θ^d とした。

【0048】

【数3】

$$\theta_l^d = (q_1 - q_2) \quad (13)$$

$$\theta_r^d = (q_3 - q_4) \quad (14)$$

10

【0049】

ただし、 θ_l^d は左脚の目標関節角、 θ_r^d は右脚の目標関節角である。

【0050】

PDサーボ処理部1028の結果出力されるロボットへのトルク入力 u は、次に示すPDサーボ系の出力を用いる。

【0051】

【数4】

$$u_l = K_p(\theta_l^d - \theta_l) - K_d\dot{\theta}_l \quad (15)$$

20

$$u_r = K_p(\theta_r^d - \theta_r) - K_d\dot{\theta}_r \quad (16)$$

【0052】

ただし、 u_l は左脚に対するトルク入力、 u_r は右脚に対するトルク入力である。また、 K_p は位置ゲイン、 K_d は速度ゲインである。

【0053】

(1-4. フィードバック制御器1022)

上述のCPGへのフィードバック制御器1022は、次の確率分布(17)によって表わされる。

【0054】

30

【数5】

$$\pi(v_j, \mathbf{x}; \mathbf{w}) = \frac{1}{\sigma_j \sqrt{2\pi}} \exp \left(-\frac{(v_j - \mu_j)^2}{2\sigma_j^2} \right) \quad (17)$$

【0055】

ただし、 $\mathbf{x}$ は制御対象の状態ベクトル、 $\mathbf{w}$ はパラメータベクトルである。従って j 番目の出力の実現値 v_j は、以下の式(18)によって与えられる。

【0056】

【数6】

40

$$v_j(t) = \mu_j(t) + \sigma_j n_j(t) \quad (18)$$

【0057】

ただし、 $n_j(t) \sim N(0, 1)$ であり、 $N(0, 1)$ は平均0、分散1の正規分布を表わす。

【0058】

ここでは出力を飽和させるために、出力飽和处理部1024において、関数 $d(\cdot)$ を用いて以下の式(19)のように、最終的な制御器の出力 $a_j(t)$ を決定する。

【0059】

50

【数 7】

$$a_j(t) = a_j^{max} d(v_j(t)) \quad (19)$$

【0060】

ただし、以下の説明では、一例として、 $d()$ としては、以下の式を用いる。

【0061】

【数 8】

$$d(x) = \frac{2}{\pi} \arctan\left(\frac{\pi}{2} x\right) \quad 10$$

【0062】

ここでの a_j ($j=1 \sim 4$) は、式 (1) の左右の脚の神経振動子の伸筋、屈筋にそれぞれ対応する。

【0063】

(2. 方策勾配法)

「方策勾配法」とは、パラメータ化された確率的方策に従って行動選択を行ない、方策を改善する方向に方策のパラメータを少しずつ更新する強化学習手法の1種である。以下に方策勾配法を用いた行動則の学習方法について述べる。

【0064】

20

(2-1. 連続時間・状態系でのテンポラル・ディファレンス (Temporal Difference) 誤差)

連続時間・状態系のダイナミクスを以下の式 (20) で表わす。

【0065】

【数 9】

$$\frac{d\mathbf{x}(t)}{dt} = f(\mathbf{x}(t), \mathbf{u}(t)) \quad (20)$$

【0066】

ただし、 $\mathbf{x} \in \mathbb{R}^n$ は状態、 $\mathbf{u} \in \mathbb{R}^m$ は制御入力を表わす。

30

【0067】

報酬は状態と制御入力の関数として、以下の式 (21) で与えられるとする。

【0068】

【数 10】

$$r(t) = r(\mathbf{x}(t), \mathbf{u}(t)) \quad (21)$$

【0069】

ある制御則 $\pi(\mathbf{u}(t) | \mathbf{x}(t))$ のもとで、状態 $\mathbf{x}(t)$ の価値関数を以下の式 (22) で定義する。

【0070】

40

【数 11】

$$V^\pi(\mathbf{x}(t)) = E \left\{ \int_t^\infty e^{-\frac{s-t}{\tau}} r(\mathbf{x}(s), \mathbf{u}(s)) ds \mid \pi \right\} \quad (22)$$

【0071】

ただし、 τ は価値関数の時定数である。また、式 (22) の両辺の時間微分から、以下の式 (23) という拘束条件が与えられる。

【0072】

【数 1 2】

$$\frac{dV^\pi(\mathbf{x}(t))}{dt} = \frac{1}{\tau} V^\pi(\mathbf{x}(t)) - r(t) \quad (23)$$

【0 0 7 3】

$V(\mathbf{x}(t)) = V(\mathbf{x}(t); w)$ を価値関数の予測値とする。ただし、 w は評価値の予測値のパラメータである。

【0 0 7 4】

予測が正しければ、式(23)を満たす。予測が正しくない場合、下式(24)に示した予測誤差を減らすように学習を行なう。 10

【0 0 7 5】

【数 1 3】

$$\delta(t) = r(t) - \frac{1}{\tau} V(t) + \dot{V}(t) \quad (24)$$

【0 0 7 6】

上式は連続時間系でのTD誤差である。

【0 0 7 7】

20

(2-2. 方策勾配法の一般論)

動的計画法やグリーディ方策(greedy policy)などの価値関数の評価を基に学習を行なう場合には、環境がマルコフ決定過程である必要があるが、実問題に適用する場合には、ノイズやセンサの能力によってマルコフ決定過程を保証することは困難である。しかし、方策勾配法は、価値関数と共に、試行中に得られた累積報酬系列を考慮することで、環境が非マルコフ決定過程(POMDP)でも適用することが出来る。

【0 0 7 8】

ここで、パラメータ w を持つ方策 π_w を用いた場合、以下の式(25)が成り立つ。

【0 0 7 9】

【数 1 4】

30

$$\frac{\partial}{\partial w_i} E\{V(t) | \pi_w\} = E\{\delta(t) e_i(t)\} \quad (25)$$

【0 0 8 0】

ただし、以下の式(26)が成り立つ。

【0 0 8 1】

【数 1 5】

$$\dot{e}_i(t) = -\frac{1}{\kappa} e_i(t) + \frac{\partial \ln \pi_w}{\partial w_i} \quad (26)$$

40

【0 0 8 2】

ここで、 κ はエリジビリティ・トレース(eligibility trace)の時定数である。テンポラル・ディファレンス誤差 δ と方策のエリジビリティ・トレース $e(t)$ により、価値関数の方策パラメータ w に関する勾配の不偏推定量を求めることが出来ることが与えられている。このような方策勾配法については、たとえば、文献：木村元、小林重信、“Actorに適正度の履歴を用いたactor-criticアルゴリズム-不完全なvalue-functionのもとでの強化学習”、人工知能学会誌, Vol.15, No.2, pp. 267-275, 2000に記載がある。

【0 0 8 3】

50

よって、パラメータの更新則は次の式(27)のようになる。

【0084】

【数16】

$$\dot{w}_i = \eta \delta(t) e(t) \quad (27)$$

【0085】

ただし、 η は学習率である。

【0086】

(2-3. 動的行動則の学習)

10

以下では、上述の方策勾配法を用いて、動的行動則の獲得を行なう。

【0087】

ここでは、式(17)に示したフィードバック制御器の学習を行なうことで望みの動的行動則を獲得することを考える。

【0088】

(2-3-1. 価値関数の更新)

まず、価値関数処理部1032において演算される、連続状態における価値関数の表現方法として、以下の式(28)による正規化ガウス関数ネットワーク(normalized Gaussian network: NGnet)を用いる。なお、正規化ガウス関数ネットワークについては、

20

【0089】

【数17】

$$V(\mathbf{x}(t)) = \sum_i w_i^c b_i^c(\mathbf{x}(t)) \quad (28)$$

【0090】

ただし、 $b_i^c(\cdot)$ は、正規化処理部1030において $\mathbf{x}$ に施される基底関数であり、 w_i^c は価値関数のパラメータである。

【0091】

パラメータ w_i^c に対するエリジビリティ・トレース e_i^c と、TD誤差を用いたパラメータ w_i^c の更新式は、それぞれ以下の式(29)および(30)のようになる。

30

【0092】

【数18】

$$\dot{e}_i^c(t) = -\frac{1}{\kappa^c} e_i^c(t) + b_i^c(\mathbf{x}(t)) \quad (29)$$

$$\dot{w}_i^c(t) = \alpha \delta(t) e_i^c(t) \quad (30)$$

40

【0093】

ただし、 α は価値関数の学習率、 κ^c はエリジビリティ・トレースの時定数である。

【0094】

(2-3-2. フィードバック制御器の更新)

式(17)に示した確率的なフィードバック制御器1022を用いる場合、その j 番目の出力の平均 μ_j と標準偏差 σ_j に関するエリジビリティは式(26)右辺第2項と同様、それぞれ以下のように与えられる。

【0095】

【数 1 9】

$$\frac{\partial \ln \pi}{\partial \mu_j} = \frac{a_t - \mu_j}{\sigma_j^2} \quad (31)$$

$$\frac{\partial \ln \pi}{\partial \sigma_j} = \frac{(a_t - \mu_j)^2 - \sigma_j^2}{\sigma_j^3} \quad (32)$$

【0 0 9 6】

10

ここではさらに、以下の式(33)および(34)のように、平均 μ を正規化ガウス関数ネットワークによって表わし、標準偏差 σ をシグモイド関数によって表わす。

【0 0 9 7】

【数 2 0】

$$\mu_j = \sum_i w_{ij}^\mu b_i^\mu(\mathbf{x}(t)) \quad (33)$$

$$\sigma_j = \frac{1}{1 + \exp(-w_j^\sigma)} \quad (34)$$

20

【0 0 9 8】

ただし、ノーターションとしては以下のとおりである。

【0 0 9 9】

【数 2 1】

b_i^μ : 基底関数、

$\omega_{ij}^\mu, \omega_j^\sigma$: 式(17)に示したフィードバック制御器のパラメータ

【0 1 0 0】

30

これらのパラメータに対応するエリジビリティは以下の式(35)および(36)のように求められる。

【0 1 0 1】

【数 2 2】

$$\frac{\partial \ln \pi}{\partial w_{ij}^\mu} = \frac{\partial \ln \pi}{\partial \mu_j} \frac{\partial \mu_j}{\partial w_{ij}^\mu} = \frac{(a_t - \mu_j) b_i^\mu(\mathbf{x}(t))}{\sigma_j^2} \quad (35)$$

$$\frac{\partial \ln \pi}{\partial w_j^\sigma} = \frac{\partial \ln \pi}{\partial \sigma_j} \frac{\partial \sigma_j}{\partial w_j^\sigma} = \frac{((a_t - \mu_j)^2 - \sigma_j^2)(1 - \sigma_j)}{\sigma_j^2} \quad (36)$$

40

【0 1 0 2】

上式(35)(36)と式(26)(27)を考慮すると、以下の式(37)および(38)のようなフィードバックパラメータの更新則が得られる。

【0 1 0 3】

【数 2 3】

$$\dot{w}_{ij}^{\mu} = \beta^{\mu} \delta(t) e_{ij}^{\mu}(t) \quad (37)$$

$$\dot{w}_j^{\sigma} = \beta^{\sigma} \delta(t) e_j^{\sigma}(t) \quad (38)$$

【0 1 0 4】

ただし、ノーターションとしては以下のとおりである。

【0 1 0 5】

10

【数 2 4】

 $\beta^{\mu}, \beta^{\sigma}$: 学習率 $e_{ij}^{\mu}(t), e_j^{\sigma}(t)$: それぞれのパラメータのエリジビリティ・トレース

【0 1 0 6】

また、式(35)(36)において、パラメータ σ が分母となっていることにより、 σ が0へと近付くとエリジビリティが発散することが問題となる。そこでエリジビリティ・トレースの更新には式(26)の代わりに次式を用いる。

【0 1 0 7】

20

【数 2 5】

$$\dot{e}_{ij}^{\mu}(t) = -\frac{1}{\kappa^{\mu}} e_{ij}(t) + \sigma_j^2 \frac{\partial \pi_{\mathbf{w}}}{\partial w_{ij}} \quad (39)$$

$$\dot{e}_j^{\sigma}(t) = -\frac{1}{\kappa^{\sigma}} e_j(t) + \sigma_j^2 \frac{\partial \pi_{\mathbf{w}}}{\partial w_j} \quad (40)$$

【0 1 0 8】

ただし、ノーターションとしては以下のとおりである。

30

【0 1 0 9】

【数 2 6】

 $\kappa^{\mu}, \kappa^{\sigma}$: それぞれのパラメータのエリジビリティ・トレースの時定数

【0 1 1 0】

(2 - 4 . 具体例)

図4に示した学習システムにおいて、数値シミュレーションを行なった結果について以下説明する。

【0 1 1 1】

このシミュレーションにおいて、図1に示した2足歩行移動システム1000(2足歩行ロボット)は、脚長が0.2m、両脚の質量がそれぞれ0.5kgとし、胴体が0.1kgであるものとした。さらに、膝関節がないことを考慮して、遊脚を振り出す場合は足先が地面を通過出来るように設定した。

40

【0 1 1 2】

それぞれの学習パラメータは、以下のとおりである。

【0 1 1 3】

【数 2 7】

学習率 : $\alpha = 0.04$
 $\beta^\mu = 0.2$
 $\beta^\sigma = 0.01$
 各エリジビリティの時定数 : $\kappa = 2.0$
 サーボ系のゲイン : $K_p = 5.0$
 $K_d = 0.1$

10

【0 1 1 4】

また、N G n e t の基底関数は、実際にロボットが歩行運動を行なう際に必要であると予想される状態空間に格子状に均等に配置することを考え、以下のようにする。

【0 1 1 5】

【数 2 8】

$$\begin{aligned} x &= (x_1, x_2, x_3, x_4)^T \\ &= (\theta_l + \theta_p, \theta_r + \theta_p, \dot{\theta}_l + \dot{\theta}_p, \dot{\theta}_r + \dot{\theta}_p)^T : (12, 6, 12, 6) \end{aligned}$$

【0 1 1 6】

20

この結果、計 5 1 8 4 (= 1 2 × 6 × 1 2 × 6) 個をそれぞれ、以下の範囲に均等に配置した。

【0 1 1 7】

【数 2 9】

$$\left(-\frac{\pi}{2} \leq x_1 \leq \frac{\pi}{2}, -3\pi \leq x_2 \leq 3\pi, -\frac{\pi}{2} \leq x_3 \leq \frac{\pi}{2}, -3\pi \leq x_4 \leq 3\pi\right)$$

【0 1 1 8】

報酬関数は以下の式で表わす。

【0 1 1 9】

30

【数 3 0】

$$r(x) = k_H r_H(x) + k_S r_S(x) \quad (41)$$

【0 1 2 0】

ただし、それぞれがロボットの腰の高さに関する項 $r_H(t)$ 、歩行速度に関する項 $r_S(t)$ は、以下の式で表わされる。

【0 1 2 1】

【数 3 1】

40

$$\begin{aligned} r_H(x) &= h_1 - h' - \min(f_l, f_r), \\ r_S &= \max(-1, \min(\dot{x}_1, 1)) \end{aligned}$$

【0 1 2 2】

ここで、 h_1 はロボットの腰の高さ、 h' は腰の高さのオフセット、 f_l, f_r は左右の脚の高さである。したがって、式 (4 1) の右辺第 1 項は、ロボットの位置エネルギーに関連する量であり、右辺第 2 項はロボットの運動エネルギーに関連する量である。

【0 1 2 3】

以下に説明するシミュレーションでは各パラメータは、 $k_S = 0.06$ 、 $k_H = 0.5$ 、 $h' = 0.15$ とした。また、C P G のパラメータは (1 - 2 . 学習システム) で述べた

50

ものを用いている。

【0124】

計算機上でのロボット及びCPGのダイナミクスの時間刻みは1 msec、学習システムの時間刻みは10 msecとした。

【0125】

また、シミュレーションにおいて、1学習試行の終了条件は以下のようにした。

【0126】

i) 17700 msec 経過(約100歩の歩行終了後)

ii) 転倒時(ただし、同時に $r = -1$ の報酬を与える)

(2-5. 平地歩行の獲得及び、環境変化に対するロバスト性)

10

図7は、1試行で獲得した報酬の総和を、試行回数ごとに取った学習曲線を示す図である。図7においては、地面の傾斜0°のときの学習曲線を示している。

【0127】

図7より、学習は約350回で収束しており、定常歩行運動を獲得出来ていることが分かる。

【0128】

図8は、図7の学習曲線に対応する歩行の軌跡を示す図である。

【0129】

図8において、(1)は学習前の歩行軌跡、(2)は600回学習後の歩行軌跡を示す。600回の学習後では、歩幅が大きくなり歩行速度も向上して、良好な歩行軌跡が得られていることがわかる。

20

【0130】

また、600回学習試行を行なうことによって学習した各学習パラメータを用い、数度の

傾斜を付けることによって環境を変化させた場合でも、ある程度歩行動作を維持することが可能である。さらに、数回の学習試行を行なうことによって、新しい環境に適応することが出来る。これは、行動則の内部状態(ここではCPGの内部状態)と、ロボットの状態が引き込みを行なうことによって、ロバストなリミットサイクルを構成しているからであると考えられる。

【0131】

30

図9は、図7で獲得した歩行において、CPGの内部状態とロボットの状態の間のリミットサイクルを、CPGの内部状態 z_1 と脚角度の時間変化として示す図である。

【0132】

また、図10は、図7で獲得した歩行において、CPGの内部状態とロボットの状態の間のリミットサイクルを、脚角度、脚の角速度、CPGの内部状態 z_1 の関係として示す図である。

【0133】

外部からの擾乱に対しても、本発明の制御システムは、周期運動を継続させることが可能なことがわかる。

【0134】

40

(2-6. 報酬と獲得した運動の関係)

式(41)の報酬関数中の、速度項係数 k_s を変化させた場合の、ロボットの歩行速度の関係を表1に示す。

【0135】

【表 1】

表 1: 報酬の速度項係数と歩行速度との関係

報酬の速度項係数 k_v	ロボットの歩行速度 (m/sec)
0.06	0.345
0.1	0.388
0.6	0.419

【0136】

10

ここで、ロボットの腰の高さに関する項の係数は、前節と同様 $k_H = 0.5$ とした。

【0137】

表 1 より、速度項を増加させるとロボットの歩行速度も増加することが分かり、よってロボットのダイナミクスから構成するようなコントローラを陽に用いることなく、学習の報酬を変化させることによって、ロボットを制御出来ることが確認出来る。

【0138】

(2-7. センサノイズ・時間遅れに対するロバスト性)

図 7 で獲得された歩行を教師信号として学習した各パラメータを初期値として用い、さらに図 4 の学習システムから CPG を取り除いたものを用いて、150 回学習試行を行なうことによって、内部状態を持たない行動則によって 2 足歩行運動を獲得した。これと、図 7 の学習によって獲得した歩行運動を用いて、コントローラのセンサノイズ及び時間遅れに対するロバスト性について比較を行なった。

20

【0139】

センサノイズは x_1 、 x_3 に対しては、 $N(0, 0.01)$ 、 x_2 、 x_4 に対しては、 $N(0, 0.09)$ を用い、時間遅れは 20 msec としてシミュレーションを行なった。

【0140】

図 11 は、センサノイズ・時間遅れに対するシミュレーション結果を示す図である。図 11 において、(1) は CPG 有り、(2) は CPG 無しのコントローラで構成された歩行を示す。また、図 11 において、(a) は通常の条件での歩行、(b) はセンサノイズのある状態での歩行、(c) は時間遅れがある場合の歩行であり、また、図 11 中で、“→” はロボットの進行方向を表わしている。

30

【0141】

CPG を持たない行動則で構成された歩行は、ノイズ及び時間遅れのどちらの場合についても歩行動作を保つ事は出来なかったが、図 4 に示した学習システムでは、ノイズ及び時間遅れがある場合でも歩行が可能であることがわかる。

【0142】

よって内部状態を持つ行動則を構成することによって、センサノイズや時間遅れに対してロバストなコントローラを構成出来ることが分かる。

【0143】

(3. 正規化ガウス関数ネットワークによる関数近似)

40

2-3-1 で述べた価値関数、フィードバック制御器を表現するために用いた、正規化ガウス関数ネットワークについて、以下説明する。

【0144】

NGnet は 3 層のネットワークで構成されており、中間素子は正規化ガウス関数である。

入力ベクトル $x = (x_1, \dots, x_n)^T$ に対して、 k 番目のユニットの活性化関数は、以下の式のようになる。

【0145】

【数 3 2】

$$\phi_k(\mathbf{x}) = e^{-\frac{1}{2}\|M_k(\mathbf{x}-\mathbf{c}_k)\|^2} \quad (42)$$

【0 1 4 6】

ただし、 $\mathbf{c}_k$ は活性化関数の中心であり、 M_k は活性化関数の形状を決定する行列である。ここで、活性化関数 $\phi_k(\mathbf{x})$ を各点で総和が1になるように以下の式(43)のように正規化したものを、基底関数 $b_k(\mathbf{x})$ とする。

【0 1 4 7】

【数 3 3】

10

$$b_k(\mathbf{x}) = \frac{\phi_k(\mathbf{x})}{\sum_{l=1}^K \phi_l(\mathbf{x})} \quad (43)$$

【0 1 4 8】

ただし、 K は基底関数の個数である。

【0 1 4 9】

このような正規化を行なうことによって、中心点 $\mathbf{c}_k$ が密に配置されている部分では、 $b_k(\mathbf{x})$ は局所的な基底関数となり、 $\mathbf{c}_k$ の分布の端の部分では $b_k(\mathbf{x})$ はシグモイド関数のような大域的な基底関数になる。

20

【0 1 5 0】

ネットワークの出力は、基底関数と重みの内積によって以下の式(44)ようになる。

【0 1 5 1】

【数 3 4】

$$y_k(\mathbf{x}) = \sum_{k=1}^K w_k b_k(\mathbf{x}) \quad (44)$$

【0 1 5 2】

30

この出力が、図4の正規化処理部1030の出力となる。

【0 1 5 3】

[実施の形態1の変形例]

以上の説明では、図3(3)の構成による制御について説明した。以下では、実施の形態1の変形例として、図3(2)の構成による制御について説明する。

【0 1 5 4】

図12は、図3(2)の構成に相当するシステムであって、制御器と制御対象を含めたシステム全体の構成を示す図である。

【0 1 5 5】

図12において、状態観測器2002は、状態観測器のダイナミクス2004と、方策勾配法(強化学習)に基づいた強化学習器2006によって構成される。

40

【0 1 5 6】

状態観測器2002中の強化学習器2006は、以下に説明するとおり、制御対象の観測出力 y と、状態観測器のダイナミクス2004に基づく出力と、制御器2010の制御出力 u とに基づいて、学習器出力 U を出力する。出力関数処理部2030は、状態観測器2002からの推定状態に基づいて、状態観測器2002の出力を報酬演算部2020に与える。報酬演算部2020は、状態観測器2002の出力と観測対象からの観測出力 y と学習器出力 U とに基づいて、報酬を計算し、強化学習器2006に与える。

【0 1 5 7】

(方策勾配法を用いた状態観測器の学習)

50

以下では、実施の形態 1 の変形例の状態観測器 2 0 0 2 の構造について説明する。

【 0 1 5 8 】

状態の推定値を、 $\hat{x}$ の頭部に “ ^ ” を付加して表現 ($= x \hat{i}$) (以下、本文中では 「 x ハット 」 と呼ぶ) したとき、つぎのような状態観測器を考える。

【 0 1 5 9 】

【 数 3 5 】

$$\dot{\hat{x}} = f(\hat{x}, u(\hat{x})) + U(\hat{x}, u(\hat{x}), y)$$

$$= g(\hat{x}, u(\hat{x}), U)$$

10

【 0 1 6 0 】

ここではまず、通常のオブザーバやカルマンフィルタ同様、対象のダイナミクス $f(x, u)$ は既知または学習によって獲得可能であるとし、対象システムの観測出力 y を基にして、現在の推定状態 x ハットと制御出力 u から、推定状態を真の状態にどのように近づけるべきかを方策勾配法を用いて学習する。

【 0 1 6 1 】

ここでは学習器の目的を、状態観測器の出力 y ハット (y の頭部に “ ^ ” を付加したもの) と対象システムの出力 y との誤差を最小にすることとする。

20

【 0 1 6 2 】

よって、報酬演算部 2 0 2 0 により演算される報酬関数は次のようになる。

【 0 1 6 3 】

【 数 3 6 】

$$r(t) = -(y(t) - \hat{y}(t))^T Q (y(t) - \hat{y}(t)) - U^T(t) R U(t)$$

【 0 1 6 4 】

ただし、 Q , R は報酬関数の形を決めるパラメータである。この結果、学習器は状態観測器のダイナミクス 2 0 0 4 への以下のようなノーテーションのフィードバック入力 U を獲得することになる。

30

【 0 1 6 5 】

【 数 3 7 】

$$U(\hat{x}, u, y)$$

【 0 1 6 6 】

ここで、フィードバック入力 U は次の確率分布により表現される。

【 0 1 6 7 】

【 数 3 8 】

$$\pi(U_j, \hat{x}, y; w) = \frac{1}{\sigma_j \sqrt{2\pi}} \exp \left(-\frac{(U_j - \mu_j)^2}{2\sigma_j^2} \right)$$

40

【 0 1 6 8 】

したがって、 j 番目の出力の実現値 U_j は、以下の式により与えられる。

【 0 1 6 9 】

【 数 3 9 】

$$U_j(t) = \mu_j(t) + \sigma_j n_j(t)$$

50

【 0 1 7 0 】

ただし、 $n_j(t) \sim N(0, 1)$ であり、 $N(0, 1)$ は、上述のとおり、平均0、分散1の正規分布を表わす。フィードバック入力 U を生成する確率分布 π の更新は、二足歩行運動の学習の場合と同様に行なわれる。

【 0 1 7 1 】

このような構成によっても、周期運動に対し状態観測器の学習を容易とすることが可能な動的制御装置およびこのような動的制御装置を用いた2足歩行移動体を提供することができる。

【 0 1 7 2 】

今回開示された実施の形態はすべての点で例示であって制限的なものではないと考えられるべきである。本発明の範囲は上記した説明ではなくて特許請求の範囲によって示され、特許請求の範囲と均等の意味および範囲内でのすべての変更が含まれることが意図される。

【図面の簡単な説明】

【 0 1 7 3 】

【図1】本発明の動的制御装置を用いた2足歩行移動システム1000の一例を示す概念図である。

【図2】図1に示した動的制御装置100の構成を示すブロック図である。

【図3】動的行動則を説明するための概念図である。

【図4】演算処理部102の行なう処理を示す機能ブロック図である。

【図5】神経振動子モデルによるCPGを示す概念図である。

【図6】CPGの出力 q を構成する変数 z_1 、 z_2 の波形を示す図である。

【図7】1試行で獲得した報酬の総和を、試行回数ごとに取った学習曲線を示す図である。

【図8】図7の学習曲線に対応する歩行の軌跡を示す図である。

【図9】CPGの内部状態とロボットの状態の間のリミットサイクルを、CPGの内部状態 z_1 と脚角度の時間変化として示す図である。

【図10】CPGの内部状態とロボットの状態の間のリミットサイクルを、脚角度、脚の角速度、CPGの内部状態 z_1 の関係として示す図である。

【図11】センサノイズ・時間遅れに対するシミュレーション結果を示す図である。

【図12】制御器と制御対象を含めたシステム全体の構成を示す図である。

【符号の説明】

【 0 1 7 4 】

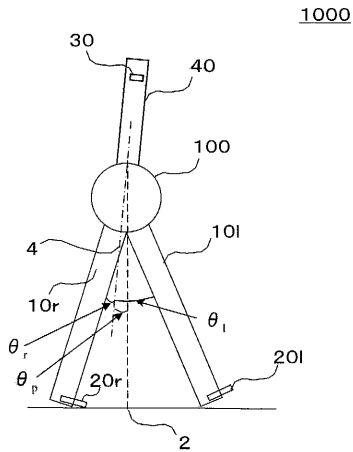
10r, 10l 脚部、20r, 20l, 30 センサ、40 胴部、100 動的制御装置、102 演算処理部、104 記憶装置、106 通信インタフェース、108 駆動部、1000 2足歩行移動システム。

10

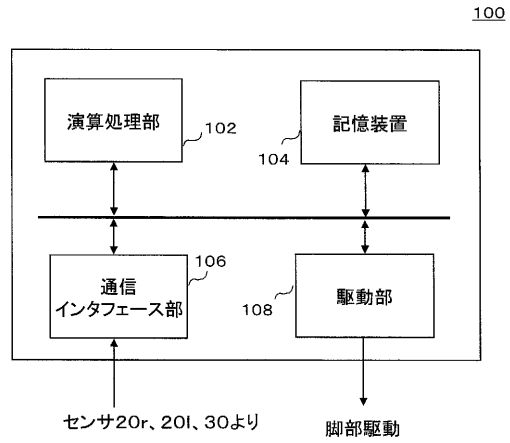
20

30

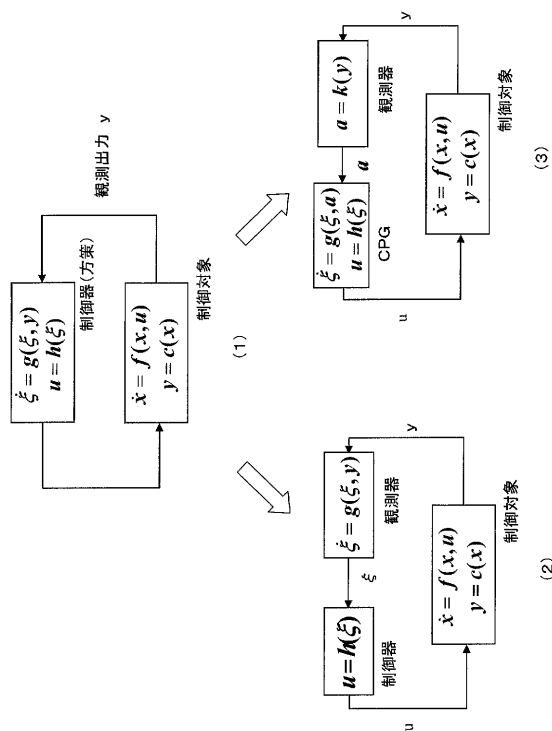
【図 1】



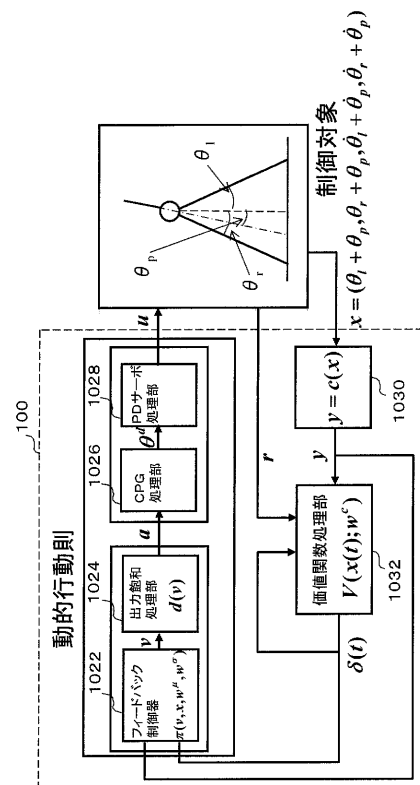
【図 2】



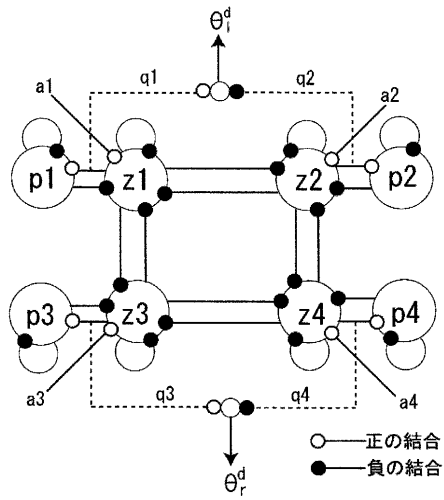
【図 3】



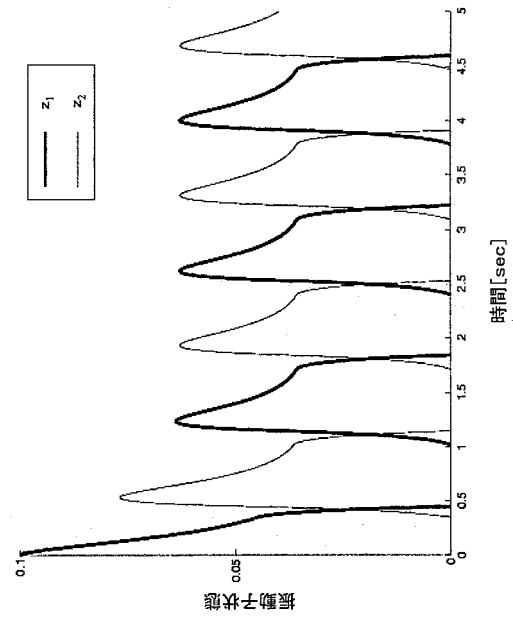
【図 4】



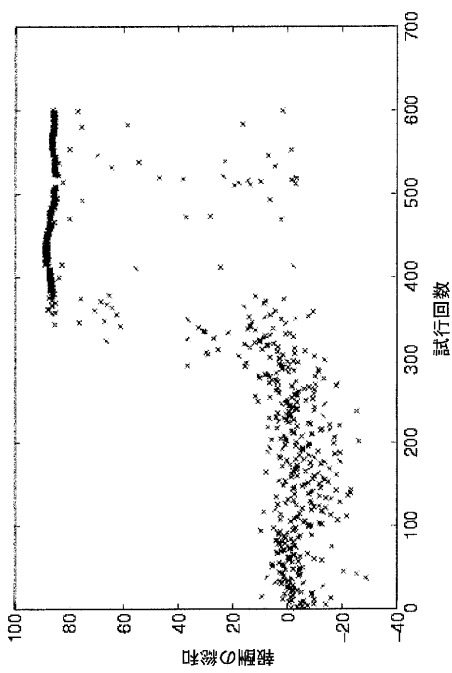
【図 5】



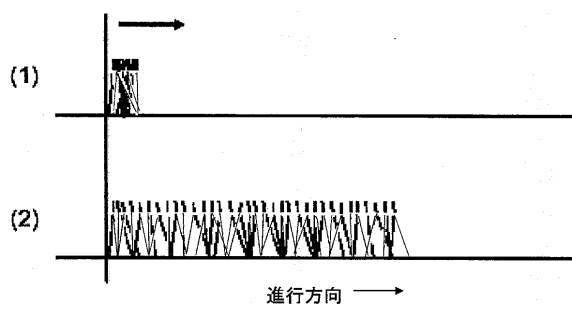
【図 6】



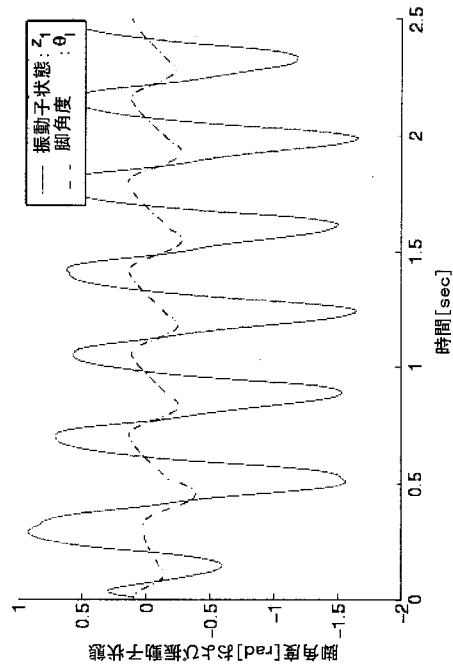
【図 7】



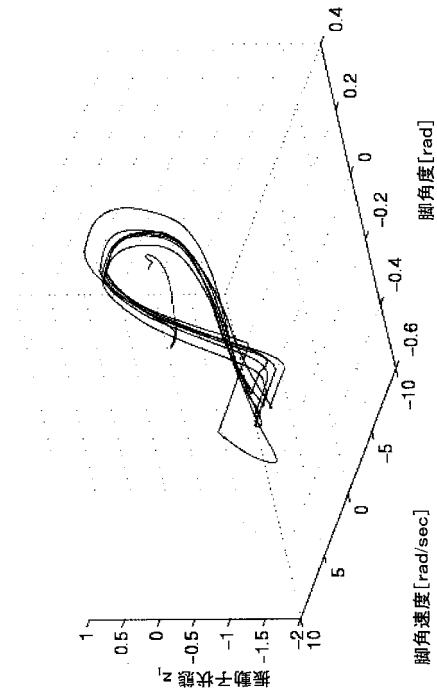
【図 8】



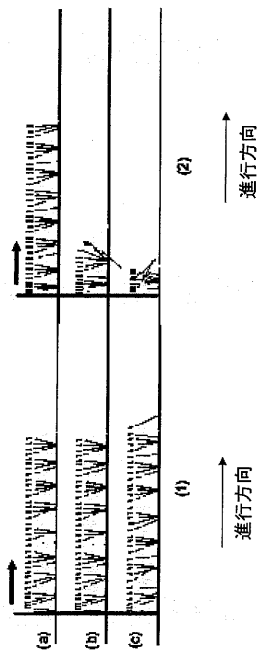
【図 9】



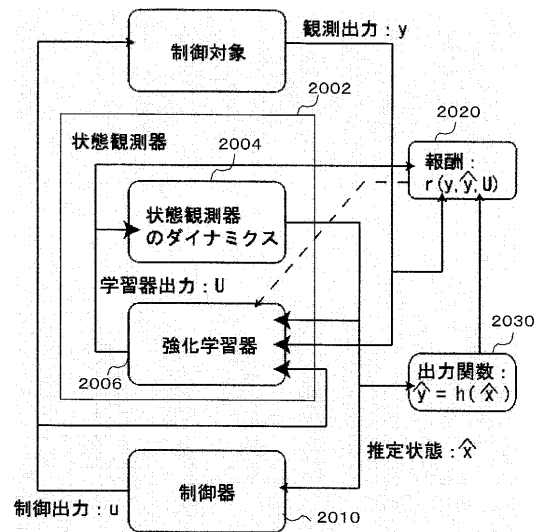
【図 10】



【図 11】



【図 12】



フロントページの続き

- (74)代理人 100098316
弁理士 野田 久登
- (74)代理人 100109162
弁理士 酒井 將行
- (72)発明者 森本 淳
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 松原 崇充
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 佐藤 雅昭
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 植村 森平

- (56)参考文献 中村泰、石井信、佐藤正樹、神経振動子ネットワークを用いた強化学習法による歩行運動の獲得，電子情報通信学会技術研究報告，日本，社団法人 電子情報通信学会，2002年 3月18日，Vol. 101, No. 735, pp. 183-190
森健、吉本潤一郎、石井信、確率の方策勾配法に基づくactor-critic法と連続システムの制御への応用，電子情報通信学会技術研究報告，日本，社団法人 電子情報通信学会，2003年 3月19日，Vol. 102, No. 731, pp. 137-142

(58)調査した分野(Int.Cl. , DB名)

B25J 1/00 - 21/02
G06N 3/00
Cini
JSTPlus(JDreamII)