

(19)日本国特許庁(JP)

(12)特許公報(B2)

(11)特許番号
特許第7537517号
(P7537517)

(45)発行日 令和6年8月21日(2024.8.21)

(24)登録日 令和6年8月13日(2024.8.13)

(51)国際特許分類

F I

G 0 6 N 20/00 (2019.01)

G 0 6 N 20/00

請求項の数 9 (全16頁)

(21)出願番号	特願2022-570960(P2022-570960)	(73)特許権者	000004237
(86)(22)出願日	令和2年12月25日(2020.12.25)		日本電気株式会社
(86)国際出願番号	PCT/JP2020/048791		東京都港区芝五丁目7番1号
(87)国際公開番号	WO2022/137520	(74)代理人	100103090
(87)国際公開日	令和4年6月30日(2022.6.30)		弁理士 岩壁 冬樹
審査請求日	令和5年6月7日(2023.6.7)	(74)代理人	100124501
			弁理士 塩川 誠人
		(72)発明者	江藤 力
			東京都港区芝五丁目7番1号 日本電気株式会社内
		審査官	渡辺 一帆

最終頁に続く

(54)【発明の名称】 学習装置、学習方法および学習プログラム

(57)【特許請求の範囲】

【請求項1】

リプシッツ連続条件を満たすように特徴量が設定された報酬関数の入力を受け付ける関数入力手段と、

熟練者の軌跡の確率分布と、前記報酬関数のパラメータに基づいて決定される軌跡の確率分布との距離を表わすワッサースタイン距離を最小にする軌跡を推定する推定手段と、

推定された軌跡に基づいて前記ワッサースタイン距離を最大にするように前記報酬関数のパラメータを更新する更新手段とを備えた

ことを特徴とする学習装置。

【請求項2】

更新手段は、非拡大写像による更新則である非拡大写像勾配法を用いて報酬関数のパラメータを更新する

請求項1記載の学習装置。

【請求項3】

更新手段は、パラメータ更新後のワッサースタイン距離が大きくなるように、一回前の更新時のワッサースタイン距離の勾配に対する今回の更新時のワッサースタイン距離の勾配の比の値と、一回前の更新時のステップ幅との積の値以下のステップ幅で報酬関数のパラメータを更新する

請求項1または請求項2記載の学習装置。

【請求項4】

ワッサースタイン距離が収束したか否か判定する判定手段を備え、

ワッサースタイン距離が収束していないと判定された場合、推定手段は、熟練者の軌跡の確率分布と、更新された報酬関数のパラメータに基づいて決定される軌跡の確率分布との距離を表わすワッサースタイン距離を最小にする軌跡を推定し、更新手段は、推定された軌跡に基づいてワッサースタイン距離を最大にするように報酬関数のパラメータを更新する

請求項 1 から請求項 3 のうちのいずれか 1 項に記載の学習装置。

【請求項 5】

関数入力手段は、線形関数になるように特徴量が設定された報酬関数の入力を受け付ける
請求項 1 から請求項 4 のうちのいずれか 1 項に記載の学習装置。

10

【請求項 6】

リップシッツ連続条件を満たすように特徴量が設定された報酬関数の入力を受け付け、
熟練者の軌跡の確率分布と、前記報酬関数のパラメータに基づいて決定される軌跡の確率分布との距離を表わすワッサースタイン距離を最小にする軌跡を推定し、
推定された軌跡に基づいて前記ワッサースタイン距離を最大にするように前記報酬関数のパラメータを更新する

ことを、コンピュータに実行させることを特徴とする学習方法。

【請求項 7】

非拡大写像による更新則である非拡大写像勾配法を用いて報酬関数のパラメータを更新する

20

請求項 6 記載の学習方法。

【請求項 8】

コンピュータに、

リップシッツ連続条件を満たすように特徴量が設定された報酬関数の入力を受け付ける関数入力処理、

熟練者の軌跡の確率分布と、前記報酬関数のパラメータに基づいて決定される軌跡の確率分布との距離を表わすワッサースタイン距離を最小にする軌跡を推定する推定処理、および、

推定された軌跡に基づいて前記ワッサースタイン距離を最大にするように前記報酬関数のパラメータを更新する更新処理

30

を実行させるための学習プログラム。

【請求項 9】

コンピュータに、

更新処理で、非拡大写像による更新則である非拡大写像勾配法を用いて報酬関数のパラメータを更新させる

請求項 8 記載の学習プログラム。

【発明の詳細な説明】

【技術分野】

【0001】

本発明は、逆強化学習を行う学習装置、学習方法および学習プログラムに関する。

40

【背景技術】

【0002】

機械学習の手法の一つに強化学習（RL：Reinforcement Learning）が知られている。強化学習は、様々な行動を試行錯誤しながら価値を最大化するような行動を学習する手法である。強化学習では、この価値を評価するための報酬関数が設定され、この報酬関数を最大にするような行動が探索される。しかし、報酬関数の設定は、一般には困難である。

【0003】

この報酬関数の設定を容易にする方法として、逆強化学習（IRL：Inverse Reinforcement Learning）が知られている。逆強化学習では、熟練者の意思決定履歴データを利用して、報酬関数を用いた最適化と、報酬関数のパラメータの更新とを繰り返すことで、

50

熟練者の意図を反映する報酬関数を生成する。

【0004】

非特許文献1には、逆強化学習の一つである最大エントロピー逆強化学習(ME-IRL: Maximum Entropy-IRL)について記載されている。非特許文献1に記載された方法では、熟練者のデータ $D = \{\tau_1, \tau_2, \dots, \tau_N\}$ (ただし、 $\tau_i = ((s_1, a_1), (s_2, a_2), \dots, (s_N, a_N))$)からただ1つの報酬関数 $R(s, a) = \theta \cdot f(s, a)$ を推定する。この推定された θ を用いることで、熟練者の意思決定を再現できる。

【0005】

また、非特許文献2には、最大エントロピー逆強化学習を改良した逆強化学習の手法の一つであるGCL(Guided Cost Learning)について記載されている。非特許文献2に記載された手法では、重点サンプリングを用いて報酬関数の重みを更新する。

10

【0006】

また、報酬関数を学習する逆強化学習と、方策を直接学習する行動模倣とを合わせて、与えられた行動履歴を再現する模倣学習も知られている(例えば、非特許文献3参照)。

【先行技術文献】

【非特許文献】

【0007】

【文献】B. D. Ziebart, A. Maas, J. A. Bagnell, and A. K. Dey, "Maximum entropy inverse reinforcement learning", In AAAI, AAAI '08, 2008.

【文献】Chelsea Finn, Sergey Levine, Pieter Abbeel, "Guided Cost Learning: Deep Inverse Optimal Control via Policy Optimization", Proceedings of The 33rd International Conference on Machine Learning, PMLR 48, pp.49-58, 2016.

20

【文献】Jonathan Ho, Stefano Ermon, "Generative adversarial imitation learning", NIPS'16: Proceedings of the 30th International Conference on Neural Information Processing Systems, pp.4572-4580, December 2016

【発明の概要】

【発明が解決しようとする課題】

【0008】

逆強化学習や模倣学習では、再現したい熟練者の行動履歴と、最適化された実行結果との差異を小さくするように報酬関数が学習される。非特許文献1~3に記載された逆強化学習や模倣学習では、KL(Kullback-Leibler)ダイバージェンスや、JS(Jensen-Shannon)ダイバージェンスのような確率的な距離で上述する差異が定義される。

30

【0009】

ここで、報酬関数のパラメータを更新する際、一般には勾配法が用いられる。しかし、組み合わせ最適化問題では確率分布の設定が難しく、実問題の多くが属している組み合わせ最適化問題に上述するような逆強化学習を適用することが困難である。

【0010】

そこで、本発明は、組み合わせ最適化問題において、安定的に逆強化学習を実施できる学習装置、学習方法および学習プログラムを提供することを目的とする。

【課題を解決するための手段】

40

【0011】

本発明による学習装置は、リブシット連続条件を満たすように特徴量が設定された報酬関数の入力を受け付ける関数入力手段と、熟練者の軌跡の確率分布と、報酬関数のパラメータに基づいて決定される軌跡の確率分布との距離を表わすワッサースタイン距離を最小にする軌跡を推定する推定手段と、推定された軌跡に基づいてワッサースタイン距離を最大にするように報酬関数のパラメータを更新する更新手段とを備えたことを特徴とする。

【0012】

本発明による学習方法は、リブシット連続条件を満たすように特徴量が設定された報酬関数の入力を受け付け、熟練者の軌跡の確率分布と、報酬関数のパラメータに基づいて決定される軌跡の確率分布との距離を表わすワッサースタイン距離を最小にする軌跡を推定

50

し、推定された軌跡に基づいてワッサースタイン距離を最大にするように報酬関数のパラメータを更新することを、コンピュータに実行させることを特徴とする。

【 0 0 1 3 】

本発明による学習プログラムは、コンピュータに、リブシツ連続条件を満たすように特徴量が設定された報酬関数の入力を受け付ける関数入力処理、熟練者の軌跡の確率分布と、報酬関数のパラメータに基づいて決定される軌跡の確率分布との距離を表わすワッサースタイン距離を最小にする軌跡を推定する推定処理、および、推定された軌跡に基づいてワッサースタイン距離を最大にするように報酬関数のパラメータを更新する更新処理を実行させることを特徴とする。

【発明の効果】

10

【 0 0 1 4 】

本発明によれば、組み合わせ最適化問題において、安定的に逆強化学習を実施できる。

【図面の簡単な説明】

【 0 0 1 5 】

【図 1】本発明による学習装置の一実施形態の構成例を示すブロック図である。

【図 2】ワッサースタイン距離を用いた逆強化学習の例を示す説明図である。

【図 3】学習装置の動作例を示すフローチャートである。

【図 4】本発明による学習装置の概要を示すブロック図である。

【図 5】少なくとも 1 つの実施形態に係るコンピュータの構成を示す概略ブロック図である。

20

【発明を実施するための形態】

【 0 0 1 6 】

まず初めに、一般的な逆強化学習を、組み合わせ最適化問題に適用することが困難な理由を説明する。非特許文献 1 に記載された ME - IRL では、熟練者の軌跡（行動履歴）を再現する報酬関数が複数存在するという不定性を解決するため、最大エントロピー原理を用いて軌跡の分布を指定し、真の分布へ近づけること（すなわち、最尤推定）により報酬関数を学習する。

【 0 0 1 7 】

ME - IRL では、軌跡 τ は、以下に例示する式 1 で表わされ、軌跡の分布 $p_{\theta}(\tau)$ を表わす確率モデルは、以下に例示する式 2 で表わされる。式 2 における $c_{\theta}(\tau)$ は、コスト関数であり、符号を逆転させる（すなわち、 $-c_{\theta}(\tau)$ ）ことで報酬関数 $r_{\theta}(\tau)$ を表わす（式 3 参照）。また、 Z は、全ての軌跡に対する報酬の総和を表わす（式 4 参照）。

30

【 0 0 1 8 】

【数 1】

$$\tau = \{(s_t, a_t) | t = 0, \dots, T\} \quad (\text{式1})$$

$$p_{\theta}(\tau) := \frac{1}{Z} \exp(-c_{\theta}(\tau)) \quad (\text{式2})$$

40

ただし、

$$-c_{\theta}(\tau) = r_{\theta}(\tau) = \sum_{t=0}^T \gamma^t r_{\theta}(s_t, a_t) \quad (\text{式3})$$

$$Z = \sum_{\tau} \exp(-c_{\theta}(\tau)) \quad (\text{式4})$$

【 0 0 1 9 】

50

そして、最尤推定による報酬関数の重みの更新則（具体的には、勾配上昇法）は、以下に例示する式 5 および式 6 で表わされる。式 5 における α はステップ幅であり、 $L_{ME}(\theta)$ は、ME - IRL で用いられる分布間の距離尺度である。

【 0 0 2 0 】

【数 2】

$$\theta \leftarrow \theta + \alpha \nabla_{\theta} L_{ME}(\theta) \quad (\text{式5})$$

$$L_{ME}(\theta) := \frac{1}{N} \sum_{i=1}^N \log p_{\theta}(\tau) = \frac{1}{N} \sum_{i=1}^N (-c_{\theta}(\tau^{(i)})) - \log \sum_{\tau} \exp(-c_{\theta}(\tau)) \quad (\text{式6})$$

【 0 0 2 1 】

上述するように、式 6 における第 2 項は、全ての軌跡に対する報酬の総和である。ME - IRL では、この第 2 項の値を厳密に計算できることを前提としている。しかし、現実的には、全ての軌跡に対する報酬の総和を計算することは困難であるため、非特許文献 2 に記載された GCL では、重点サンプリングで近似的にこの値を算出する。

【 0 0 2 2 】

しかし、組み合わせ最適化問題は、離散的な値（言い換えると、連続的ではない値）をとるため、ある値を入力したときに、その値に対応する確率を返す確率分布の設定が難しい。組み合わせ最適化問題では、目的関数における値が少しでも変化すると、結果も大きく変化する可能性があるためである。

【 0 0 2 3 】

例えば、典型的な組み合わせ最適化問題の例として、経路問題やスケジューリング問題、切り出し・詰め込み問題や割り当て・マッチング問題などが挙げられる。具体的には、経路問題は、例えば、運搬経路問題や巡回セールスマン問題などであり、スケジューリング問題は、例えば、ジョブショップ問題や勤務スケジュール問題などである。また、切り出し・詰め込み問題は、例えば、ナップサック問題やピンパッキング問題などであり、割り当て・マッチング問題は、最大マッチング問題や一般化割当問題などである。

【 0 0 2 4 】

本開示の学習装置を用いることで、これらの組み合わせ最適化問題において、安定的に逆強化学習を実施することが可能になる。以下、本発明の実施形態を図面を参照して説明する。

【 0 0 2 5 】

図 1 は、本発明による学習装置の一実施形態の構成例を示すブロック図である。本実施形態の学習装置 100 は、機械学習により、対象者（熟練者）の行動から報酬関数を推定する逆強化学習を行う装置であり、熟練者の行動特性に基づく情報処理を具体的に行う装置である。学習装置 100 は、記憶部 10 と、入力部 20 と、特徴量設定部 30 と、重み初期値設定部 40 と、数理最適化実行部 50 と、重み更新部 60 と、収束判定部 70 と、出力部 80 とを備えている。

【 0 0 2 6 】

なお、数理最適化実行部 50、重み更新部 60 および収束判定部 70 により、後述する逆強化学習が行われることから、数理最適化実行部 50、重み更新部 60 および収束判定部 70 を含む装置を、逆強化学習装置とすることができる。

【 0 0 2 7 】

記憶部 10 は、学習装置 100 が各種処理を行うために必要な情報を記憶する。記憶部 10 は、後述する入力部 20 が受け付けた熟練者の意思決定履歴データ（軌跡）を記憶し

10

20

30

40

50

てもよい。また、記憶部 10 は、後述する数理最適化実行部 50 および重み更新部 60 が学習に用いる報酬関数の特徴量の候補を記憶していてもよい。ただし、特徴量の候補は、必ずしも目的関数に使用される特徴量である必要はない。

【0028】

また、記憶部 10 は、後述する数理最適化実行部 50 を実現するための数理最適化ソルバを記憶していてもよい。なお、数理最適化ソルバの内容は任意であり、実行する環境や装置に応じて決定されればよい。

【0029】

入力部 20 は、学習装置 100 が各種処理を行うために必要な情報の入力を受け付ける。入力部 20 は、例えば、上述する熟練者の意思決定履歴データ（具体的には、状態と行動のペア）の入力を受け付けてもよい。また、入力部 20 は、後述する逆強化学習装置が逆強化学習を行う際に用いる初期状態の制約 z の入力を受け付けてもよい。

10

【0030】

特徴量設定部 30 は、状態および行動を含むデータから、報酬関数の特徴量を設定する。具体的には、特徴量設定部 30 は、後述する逆強化学習装置が分布間の距離尺度としてワッサーズタイン (Wasserstein) 距離を利用できるように、関数全体で接線の勾配が有限になるように報酬関数の特徴量を設定する。特徴量設定部 30 は、例えば、リップシツ連続条件を満たすように報酬関数の特徴量を設定してもよい。

【0031】

例えば、 f_τ を軌跡 τ の特徴量ベクトルとする。コスト関数 $c_\theta(\tau) = \theta^\top f_\tau$ と線形に限定した場合、写像 $F: \tau \rightarrow f_\tau$ がリップシツ連続であれば、 $c_\theta(\tau)$ もリップシツ連続である。そのため、特徴量設定部 30 は、報酬関数が線形関数になるように特徴量を設定してもよい。

20

【0032】

なお、例えば、以下に例示する式 7 は、 a_0 において勾配が無限大になってしまうため、本開示において不適切な報酬関数と言える。

【0033】

【数 3】

$$f_\tau = \begin{cases} 1 & (a_0 \geq 0) \\ 0 & (otherwise) \end{cases} \quad (\text{式7})$$

30

【0034】

特徴量設定部 30 は、例えば、ユーザの指示に応じて特徴量が設定された報酬関数を決定してもよく、記憶部 10 からリップシツ連続条件を満たす報酬関数を取得してもよい。

【0035】

重み初期値設定部 40 は、報酬関数の重みを初期化する。具体的には、重み初期値設定部 40 は、報酬関数に含まれる個々の特徴量の重みを設定する。なお、重みを初期化する方法は特に限定されず、ユーザ等に応じて予め定められた任意の方法に基づいて重みが初期化されればよい。

40

【0036】

数理最適化実行部 50 は、熟練者の軌跡（行動履歴）の確率分布と、最適化された（報酬関数の）パラメータに基づいて決定される軌跡の確率分布との間の距離を最小にする軌跡 $\tau^\wedge$ （ $\tau^\wedge$ は、 τ の上付き $\wedge$ ）を導出する。具体的には、数理最適化実行部 50 は、分布間の距離尺度として、KL / JS ダイバージェンスの代わりにワッサーズタイン距離を利用して、そのワッサーズタイン距離を最小にするよう数理最適化を実行することにより、熟練者の軌跡 $\tau^\wedge$ を推定する。

【0037】

50

ワッサーズタイン距離は、以下に例示する式 8 で定義される。なお、ワッサーズタイン距離の制約から、コスト関数 $c_{\theta}(\tau)$ は、リブシッツ連続条件を満たす関数である必要がある。一方、本実施形態では、特徴量設定部 30 によってリブシッツ連続条件を満たすように報酬関数の特徴量が設定されているため、数理最適化実行部 50 は、以下に例示するようなワッサーズタイン距離を利用することが可能になる。

【0038】

【数 4】

$$W(\theta) := \frac{1}{N} \sum_{i=1}^N (-c_{\theta}(\tau^{(i)})) - \frac{1}{N} \sum_{i=1}^N (-c_{\theta}(\hat{\tau}(\theta, z^{(i)}))) \quad (式8)$$

【0039】

上記に例示する式 8 で定義されるワッサーズタイン距離は 0 以下の値をとり、この値を大きくすることが、分布同士を近づけることに対応する。また、式 8 の第 2 項において、コスト関数 c_{θ} の引数（すなわち、 $\tau^{\wedge}(\theta, z^{(i)})$ ）は、パラメータ θ で最適化した i 番目の軌跡を表わす。なお、 z は、軌跡パラメータである。式 8 の第 2 項は、組み合わせ最適化問題でも算出可能な項である。そのため、式 8 に例示するワッサーズタイン距離を分布間の距離尺度として用いることで、組み合わせ最適化問題において、安定的に逆強化学習を実施することが可能になる。

【0040】

重み更新部 60 は、推定された熟練者の軌跡 $\tau^{\wedge}$ に基づいて分布間の距離尺度を最大にするように報酬関数のパラメータ θ を更新する。具体的には、重み更新部 60 は、上述するワッサーズタイン距離を最大にするように報酬関数のパラメータを更新する。重み更新部 60 は、例えば、推定された軌跡 $\tau^{\wedge}$ を固定して、勾配上昇法によりパラメータを更新してもよい。

【0041】

また、本実施形態では、重み更新部 60 は、報酬関数のパラメータを更新する際、ワッサーズタイン距離を単調増加させるために、非拡大写像による更新則（以下、非拡大写像勾配法と記すこともある。）を用いてもよい。以下、非拡大写像勾配法について詳述する。

【0042】

ここでは、報酬関数として線形関数が用いられる場合を例示する。上記のように軌跡 τ の特徴量ベクトルを f_{τ} とすると、報酬関数は、以下に例示する式 9 のように表される。

【0043】

【数 5】

$$-c_{\theta}(\tau) = r_{\theta}(\tau) = \theta^{\top} f_{\tau} \quad (式9)$$

【0044】

このとき、ワッサーズタイン距離の単調増加性を保証するため、任意の軌跡 τ_a および軌跡 τ_b 、並びに、各軌跡に対する特徴量ベクトル f_{τ_a} および特徴量ベクトル f_{τ_b} について、以下の式 10 に例示する関係を満たす定数 K が存在する必要がある。

【0045】

【数 6】

10

20

30

40

50

$$\|\theta^\top f_{\tau_a} - \theta^\top f_{\tau_b}\| \leq K \|\tau_a - \tau_b\| \quad (\text{式10})$$

【 0 0 4 6 】

ここで、上記に示す式 1 0 は、以下に例示する式 1 1 のように書き換えることができる。

【 0 0 4 7 】

【数 7 】

$$\|f_{\tau_a} - f_{\tau_b}\| \leq \tilde{K} \|\tau_a - \tau_b\| \quad (\text{式11})$$

10

【 0 0 4 8 】

t 回目に更新される報酬関数のパラメータを θ_t 、ワッサースタイン距離を $W(\theta_t)$ 、および、ステップ幅を α_t とする。このとき、報酬関数のパラメータの更新則は、以下に例示する式 1 2 のように表すことができる。

【 0 0 4 9 】

【数 8 】

$$\theta_{t+1} = \theta_t + \alpha_t \nabla W(\theta_t) \quad (\text{式12})$$

20

【 0 0 5 0 】

重み更新部 6 0 は、報酬関数のパラメータの更新則（すなわち、 $\theta(t) \rightarrow \theta(t+1)$ ）が非拡大写像になるという制約下で、ワッサースタイン距離を大きくする勾配のステップ幅を探索し、そのステップ幅で報酬関数のパラメータを更新する。具体的には、重み更新部 6 0 は、以下の式 1 3 および式 1 4 に示す条件を満たすステップ幅 α_t で報酬関数のパラメータを更新する。

30

【 0 0 5 1 】

【数 9 】

$$0 < \alpha_t \leq \alpha_{t-1} \frac{\|\nabla W(\theta_{t-1})\|}{\|\nabla W(\theta_t)\|} \quad (\text{式13})$$

$$W(\theta_{t+1}) > W(\theta_t) \quad (\text{式14})$$

40

【 0 0 5 2 】

式 1 3 および式 1 4 は、パラメータ更新後のワッサースタイン距離が大きくなる（ $W(\theta_{t+1}) > W(\theta_t)$ ）ように、一回前の更新時 $t-1$ のワッサースタイン距離 $W(\theta_{t-1})$ の勾配 $W(\theta_{t-1})$ に対する今回の更新時 t のワッサースタイン距離 $W(\theta_t)$ の勾配 $W(\theta_t)$ の比 $(\|\nabla W(\theta_{t-1})\| / \|\nabla W(\theta_t)\|)$ の値と一回前の更新時 $t-1$ のステップ幅 α_{t-1} との積以下になるような正のステップ幅 α_t の値を探索することを示す。

【 0 0 5 3 】

例えば、組み合わせ最適化問題の場合、数理最適化実行部 5 0 による推定結果が報酬関

50

数の変化に対して不連続になる場合がある。具体的には、ある値の最大化と最小化とを交互に実施する更新では、多くの場合、その値が振動し、収束するまでに時間を要することがある。一方、本実施形態では、数理最適化実行部 50 が、上述する非拡大写像勾配法を用いることで、ワッサーズタイン距離の単調増加性を保証しながらパラメータを更新することが可能になる。

【0054】

以降、後述する収束判定部 70 によって、ワッサーズタイン距離が収束されたと判定されるまで、数理最適化実行部 50 による軌跡の推定処理、および、重み更新部 60 によるパラメータの更新処理が繰り返される。

【0055】

収束判定部 70 は、分布間の距離尺度が収束したか否かを判定する。具体的には、収束判定部 70 は、ワッサーズタイン距離が収束したか否かを判定する。判定方法は任意であり、収束判定部 70 は、例えば、分布間のワッサーズタイン距離の絶対値が予め定めた閾値より小さくなったときに、分布間の距離尺度が収束したと判定してもよい。

【0056】

収束判定部 70 は、距離が収束していないと判断した場合、数理最適化実行部 50 および重み更新部 60 による処理を継続させる。一方、収束判定部 70 は、距離が収束したと判断した場合、数理最適化実行部 50 および重み更新部 60 による処理を終了させる。

【0057】

出力部 80 は、学習された報酬関数を出力する。

【0058】

図 2 は、ワッサーズタイン距離を用いた逆強化学習の例を示す説明図である。なお、本開示で示すワッサーズタイン距離を用いた逆強化学習のことを、Wasserstein IRL (WIRL) と記すこともある。

【0059】

まず、初期状態の制約 z および初期値が設定されたパラメータ θ の報酬関数に基づき、最適化ソルバを用いて、ワッサーズタイン距離を最小化するように数理最適化を行うことで軌跡 τ^* を推定する。なお、図 2 に例示する最適化ソルバは、数理最適化実行部 50 に対応する。

【0060】

一方、推定された軌跡 τ^* と入力された熟練者の軌跡 τ に基づいて、ワッサーズタイン距離を最大化するように数理最適化を行うことで報酬関数（コスト関数）のパラメータを更新する。この処理は、重み更新部 60 の処理に対応する。

【0061】

以降、ワッサーズタイン距離が収束したと判定されるまで、図 2 に例示する処理が繰り返される。

【0062】

入力部 20 と、特徴量設定部 30 と、重み初期値設定部 40 と、数理最適化実行部 50 と、重み更新部 60 と、収束判定部 70 と、出力部 80 とは、プログラム（学習プログラム）に従って動作するコンピュータのプロセッサ（例えば、CPU (Central Processing Unit)）によって実現される。

【0063】

例えば、プログラムは、学習装置 100 が備える記憶部 10 に記憶され、プロセッサは、そのプログラムを読み込み、プログラムに従って、入力部 20、特徴量設定部 30、重み初期値設定部 40、数理最適化実行部 50、重み更新部 60、収束判定部 70 および出力部 80 として動作してもよい。また、学習装置 100 の機能が SaaS (Software as a Service) 形式で提供されてもよい。

【0064】

また、入力部 20 と、特徴量設定部 30 と、重み初期値設定部 40 と、数理最適化実行部 50 と、重み更新部 60 と、収束判定部 70 と、出力部 80 とは、それぞれが専用のハ

10

20

30

40

50

ードウェアで実現されていてもよい。また、各装置の各構成要素の一部又は全部は、汎用または専用の回路（circuitry）、プロセッサ等やこれらの組合せによって実現されもよい。これらは、単一のチップによって構成されてもよいし、バスを介して接続される複数のチップによって構成されてもよい。各装置の各構成要素の一部又は全部は、上述した回路等とプログラムとの組合せによって実現されてもよい。

【0065】

また、学習装置100の各構成要素の一部又は全部が複数の情報処理装置や回路等により実現される場合には、複数の情報処理装置や回路等は、集中配置されてもよいし、分散配置されてもよい。例えば、情報処理装置や回路等は、クライアントサーバシステム、クラウドコンピューティングシステム等、各々が通信ネットワークを介して接続される形態として実現されてもよい。

10

【0066】

次に、本実施形態の学習装置100の動作を説明する。図3は、本実施形態の学習装置100の動作例を示すフローチャートである。入力部20は、エキスパートデータ（すなわち、熟練者の軌跡／意思決定履歴データ）の入力を受け付ける（ステップS11）。特徴量設定部30は、状態および行動を含むデータから、リブシット連続条件を満たすように報酬関数の特徴量を設定する（ステップS12）。また、重み初期値設定部40は、報酬関数の重み（パラメータ）を初期化する（ステップS13）。

【0067】

数理最適化実行部50は、リブシット連続条件を満たすように特徴量が設定された報酬関数の入力を受け付ける（ステップS14）。そして、数理最適化実行部50は、ワッサースタイン距離を最小にするように数理最適化を実行する（ステップS15）。具体的には、数理最適化実行部50は、熟練者の軌跡の確率分布と、報酬関数のパラメータに基づいて決定される軌跡の確率分布との距離を表わすワッサースタイン距離を最小にする軌跡を推定する。

20

【0068】

重み更新部60は、推定された軌跡に基づいてワッサースタイン距離を最大にするように報酬関数のパラメータを更新する（ステップS16）。重み更新部60は、例えば、非拡大写像勾配法を用いて報酬関数のパラメータを更新してもよい。

【0069】

収束判定部70は、ワッサースタイン距離が収束したか否か判定する（ステップS17）。ワッサースタイン距離が収束していないと判定された場合（ステップS17におけるNo）、更新された軌跡を用いてステップS15以降の処理が繰り返される。一方、ワッサースタイン距離が収束したと判定された場合（ステップS17におけるYes）、出力部80は、学習された報酬関数を出力する（ステップS18）。

30

【0070】

以上のように、本実施形態では、数理最適化実行部50が、リブシット連続条件を満たすように特徴量が設定された報酬関数の入力を受け付け、熟練者の軌跡の確率分布と、報酬関数のパラメータに基づいて決定される軌跡の確率分布との距離を表わすワッサースタイン距離を最小にする軌跡を推定する。そして、重み更新部60が、推定された軌跡に基づいてワッサースタイン距離を最大にするように報酬関数のパラメータを更新する。よって、組み合わせ最適化問題において、安定的に逆強化学習を実施できる。

40

【0071】

次に、本発明の概要を説明する。図4は、本発明による学習装置の概要を示すブロック図である。本発明による学習装置90（例えば、学習装置100）は、リブシット連続条件を満たすように特徴量が設定された報酬関数の入力を受け付ける関数入力手段91（例えば、数理最適化実行部50）と、熟練者の軌跡の確率分布と、報酬関数のパラメータに基づいて決定される軌跡の確率分布との距離を表わすワッサースタイン距離を最小にする軌跡を推定する推定手段92（例えば、数理最適化実行部50）と、推定された軌跡に基づいてワッサースタイン距離を最大にするように報酬関数のパラメータを更新する更新手

50

段 9 3 (例えば、重み更新部 6 0) とを備えている。

【 0 0 7 2 】

そのような構成により、組み合わせ最適化問題において、安定的に逆強化学習を実施できる。

【 0 0 7 3 】

更新手段 9 3 は、非拡大写像による更新則である非拡大写像勾配法を用いて報酬関数のパラメータを更新してもよい。

【 0 0 7 4 】

具体的には、更新手段 9 3 は、パラメータ更新後のワッサースタイン距離 (例えば、 $W(\theta)$) が大きくなるように (すなわち、 $W(\theta_{t+1}) > W(\theta_t)$)、一回前の更新時 ($t - 1$ 回目) のワッサースタイン距離の勾配 (例えば、 $W(\theta_{t-1})$) に対する今回の更新時 (例えば、 t 回目) のワッサースタイン距離の勾配 (例えば、 $W(\theta_t)$) の比の値と、一回前の更新時のステップ幅 (例えば、 α_{t-1}) との積の値以下のステップ幅 (例えば、 α_t) で報酬関数のパラメータを更新してもよい (例えば、式 1 3 および式 1 4 参照)。

【 0 0 7 5 】

また、学習装置 9 0 は、ワッサースタイン距離が収束したか否か判定する判定手段 (例えば、収束判定部 7 0) を備えていてもよい。そして、ワッサースタイン距離が収束していないと判定された場合、推定手段 9 2 は、熟練者の軌跡の確率分布と、更新された報酬関数のパラメータに基づいて決定される軌跡の確率分布との距離を表わすワッサースタイン距離を最小にする軌跡を推定し、更新手段 9 3 は、推定された軌跡に基づいてワッサースタイン距離を最大にするように報酬関数のパラメータを更新してもよい。

【 0 0 7 6 】

また、関数入力手段 9 1 は、線形関数になるように特徴量が設定された報酬関数の入力を受け付けてもよい。

【 0 0 7 7 】

図 5 は、少なくとも 1 つの実施形態に係るコンピュータの構成を示す概略ブロック図である。コンピュータ 1 0 0 0 は、プロセッサ 1 0 0 1、主記憶装置 1 0 0 2、補助記憶装置 1 0 0 3、インタフェース 1 0 0 4 を備える。

【 0 0 7 8 】

上述の学習装置 9 0 は、コンピュータ 1 0 0 0 に実装される。そして、上述した各処理部の動作は、プログラム (学習プログラム) の形式で補助記憶装置 1 0 0 3 に記憶されている。プロセッサ 1 0 0 1 は、プログラムを補助記憶装置 1 0 0 3 から読み出して主記憶装置 1 0 0 2 に展開し、当該プログラムに従って上記処理を実行する。

【 0 0 7 9 】

なお、少なくとも 1 つの実施形態において、補助記憶装置 1 0 0 3 は、一時的でない有形の媒体の一例である。一時的でない有形の媒体の他の例としては、インタフェース 1 0 0 4 を介して接続される磁気ディスク、光磁気ディスク、CD-ROM (Compact Disc Read-only memory)、DVD-ROM (Read-only memory)、半導体メモリ等が挙げられる。また、このプログラムが通信回線によってコンピュータ 1 0 0 0 に配信される場合、配信を受けたコンピュータ 1 0 0 0 が当該プログラムを主記憶装置 1 0 0 2 に展開し、上記処理を実行してもよい。

【 0 0 8 0 】

また、当該プログラムは、前述した機能の一部を実現するためのものであっても良い。さらに、当該プログラムは、前述した機能を補助記憶装置 1 0 0 3 に既に記憶されている他のプログラムとの組み合わせで実現するもの、いわゆる差分ファイル (差分プログラム) であってもよい。

【 0 0 8 1 】

上記の実施形態の一部又は全部は、以下の付記のようにも記載されうるが、以下には限られない。

10

20

30

40

50

【 0 0 8 2 】

(付記 1) リブシツ連続条件を満たすように特徴量が設定された報酬関数の入力を受け付ける関数入力手段と、

熟練者の軌跡の確率分布と、前記報酬関数のパラメータに基づいて決定される軌跡の確率分布との距離を表わすワッサースタイン距離を最小にする軌跡を推定する推定手段と、

推定された軌跡に基づいて前記ワッサースタイン距離を最大にするように前記報酬関数のパラメータを更新する更新手段とを備えた

ことを特徴とする学習装置。

【 0 0 8 3 】

(付記 2) 更新手段は、非拡大写像による更新則である非拡大写像勾配法を用いて報酬関数のパラメータを更新する

付記 1 記載の学習装置。

【 0 0 8 4 】

(付記 3) 更新手段は、パラメータ更新後のワッサースタイン距離が大きくなるように、一回前の更新時のワッサースタイン距離の勾配に対する今回の更新時のワッサースタイン距離の勾配の比の値と、一回前の更新時のステップ幅との積の値以下のステップ幅で報酬関数のパラメータを更新する

付記 1 または付記 2 記載の学習装置。

【 0 0 8 5 】

(付記 4) ワッサースタイン距離が収束したか否か判定する判定手段を備え、

ワッサースタイン距離が収束していないと判定された場合、推定手段は、熟練者の軌跡の確率分布と、更新された報酬関数のパラメータに基づいて決定される軌跡の確率分布との距離を表わすワッサースタイン距離を最小にする軌跡を推定し、更新手段は、推定された軌跡に基づいてワッサースタイン距離を最大にするように報酬関数のパラメータを更新する

付記 1 から付記 3 のうちのいずれか 1 つに記載の学習装置。

【 0 0 8 6 】

(付記 5) 関数入力手段は、線形関数になるように特徴量が設定された報酬関数の入力を受け付ける

付記 1 から付記 4 のうちのいずれか 1 つに記載の学習装置。

【 0 0 8 7 】

(付記 6) リブシツ連続条件を満たすように特徴量が設定された報酬関数の入力を受け付け、

熟練者の軌跡の確率分布と、前記報酬関数のパラメータに基づいて決定される軌跡の確率分布との距離を表わすワッサースタイン距離を最小にする軌跡を推定し、

推定された軌跡に基づいて前記ワッサースタイン距離を最大にするように前記報酬関数のパラメータを更新する

ことを特徴とする学習方法。

【 0 0 8 8 】

(付記 7) 非拡大写像による更新則である非拡大写像勾配法を用いて報酬関数のパラメータを更新する

付記 6 記載の学習方法。

【 0 0 8 9 】

(付記 8) コンピュータに、

リブシツ連続条件を満たすように特徴量が設定された報酬関数の入力を受け付ける関数入力処理、

熟練者の軌跡の確率分布と、前記報酬関数のパラメータに基づいて決定される軌跡の確率分布との距離を表わすワッサースタイン距離を最小にする軌跡を推定する推定処理、および、

推定された軌跡に基づいて前記ワッサースタイン距離を最大にするように前記報酬関数

10

20

30

40

50

のパラメータを更新する更新処理
を実行させるための学習プログラムを記憶するプログラム記憶媒体。

【 0 0 9 0 】

(付記 9) コンピュータに、
更新処理で、非拡大写像による更新則である非拡大写像勾配法を用いて報酬関数のパラメータを更新させる
ための学習プログラムを記憶する付記 8 記載のプログラム記憶媒体。

【 0 0 9 1 】

(付記 1 0) コンピュータに、
リブシット連続条件を満たすように特徴量が設定された報酬関数の入力を受け付ける関数入力処理、
熟練者の軌跡の確率分布と、前記報酬関数のパラメータに基づいて決定される軌跡の確率分布との距離を表わすワッサーズタイン距離を最小にする軌跡を推定する推定処理、および、
推定された軌跡に基づいて前記ワッサーズタイン距離を最大にするように前記報酬関数のパラメータを更新する更新処理
を実行させるための学習プログラム。

【 0 0 9 2 】

(付記 1 1) コンピュータに、
更新処理で、非拡大写像による更新則である非拡大写像勾配法を用いて報酬関数のパラメータを更新させる
付記 1 0 記載の学習プログラム。

【符号の説明】

【 0 0 9 3 】

1 0 記憶部
2 0 入力部
3 0 特徴量設定部
4 0 重み初期値設定部
5 0 数理最適化実行部
6 0 重み更新部
7 0 収束判定部
1 0 0 学習装置

10

20

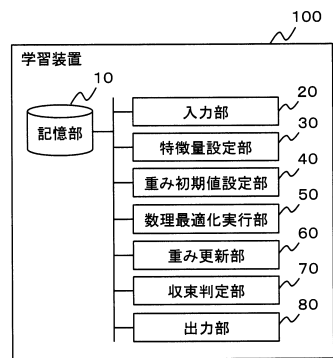
30

40

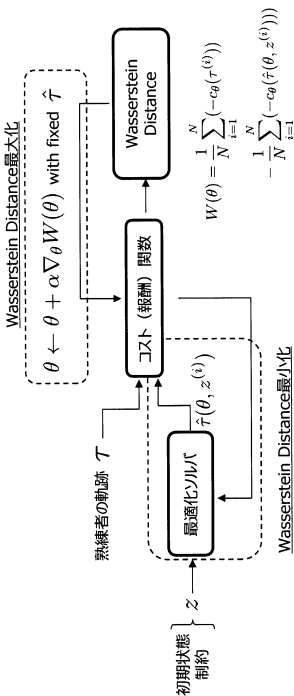
50

【 図 面 】

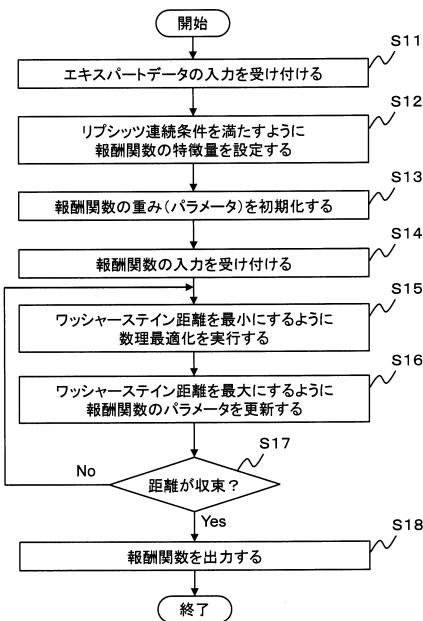
【 図 1 】



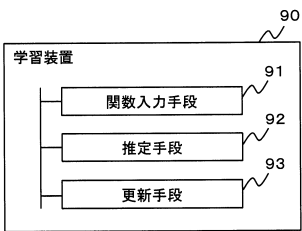
【 図 2 】



【 図 3 】



【 図 4 】



10

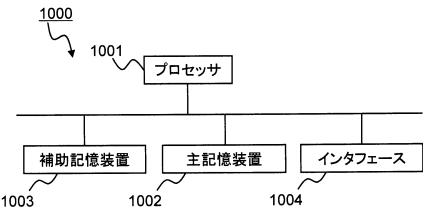
20

30

40

50

【 図 5 】



10

20

30

40

50

フロントページの続き

(56)参考文献 国際公開第 2 0 1 9 / 1 5 5 0 5 2 (W O , A 1)
 特開 2 0 2 0 - 1 7 7 0 1 6 (J P , A)
 国際公開第 2 0 1 8 / 1 3 1 2 1 4 (W O , A 1)
 ラシュカ セバスチャン ほか, "17.3 畳み込み層とワッサーズタイン距離を使って合成画像
 の品質を改善する", [第3版] Python機械学習プログラミング 達人データサイエンティス
 トによる理論と実践, 第1版, 株式会社インプレス, 2020年10月, pp. 548-567, ISBN 97
 8-4-295-01007-4
(58)調査した分野 (Int.Cl., D B 名)
 G 0 6 N 2 0 / 0 0 - 2 0 / 2 0