

(19)日本国特許庁(JP)

(12)特許公報(B2)

(11)特許番号  
特許第7652276号  
(P7652276)

(45)発行日 令和7年3月27日(2025.3.27)

(24)登録日 令和7年3月18日(2025.3.18)

(51)国際特許分類

F I

G O 6 N 20/00 (2019.01)

G O 6 N 20/00 1 3 0

G O 6 N 3/0895(2023.01)

G O 6 N 3/0895

G O 6 F 18/214 (2023.01)

G O 6 F 18/214

請求項の数 5 (全23頁)

|             |                             |          |                                |
|-------------|-----------------------------|----------|--------------------------------|
| (21)出願番号    | 特願2023-550801(P2023-550801) | (73)特許権者 | 000005223                      |
| (86)(22)出願日 | 令和3年9月28日(2021.9.28)        |          | 富士通株式会社                        |
| (86)国際出願番号  | PCT/JP2021/035678           |          | 神奈川県川崎市中原区上小田中4丁目1番1号          |
| (87)国際公開番号  | WO2023/053216               | (74)代理人  | 110002147                      |
| (87)国際公開日   | 令和5年4月6日(2023.4.6)          |          | 弁理士法人酒井国際特許事務所                 |
| 審査請求日       | 令和6年1月11日(2024.1.11)        | (72)発明者  | 大川 佳寛                          |
|             |                             |          | 神奈川県川崎市中原区上小田中4丁目1番1号 富士通株式会社内 |
|             |                             | (72)発明者  | 横田 泰斗                          |
|             |                             |          | 神奈川県川崎市中原区上小田中4丁目1番1号 富士通株式会社内 |
|             |                             | 審査官      | 多賀 実                           |

最終頁に続く

(54)【発明の名称】 機械学習プログラム、機械学習方法および機械学習装置

(57)【特許請求の範囲】

【請求項1】

複数のデータを機械学習モデルに入力して、前記複数のデータの複数の予測結果を取得し、

前記複数のデータのうち前記複数の予測結果が第2のグループを示す第1の複数のデータから、予測結果が第1のグループを示す第1のデータに類似する第2のデータを選択し、前記第1のデータの特徴量と前記第2のデータの特徴量との間の特徴量に対応する一又は複数のデータを生成し、

前記機械学習モデルのパラメータに基づいて得られた、前記複数のデータと前記一又は複数のデータとのそれぞれの複数の特徴量に基づいて、前記複数のデータと前記一又は複数のデータとのクラスタリングを実行し、

前記クラスタリングの結果を正解ラベルとする前記複数のデータと前記一又は複数のデータとを含む訓練データに基づいて、前記機械学習モデルのパラメータを更新する、処理をコンピュータに実行させることを特徴とする機械学習プログラム。

【請求項2】

前記生成する処理は、前記第1のデータを複製することによって得られた第3のデータにノイズを付加することで、前記一又は複数のデータを生成する処理を含む、

ことを特徴とする請求項1に記載の機械学習プログラム。

【請求項3】

前記予測結果と、前記訓練データに含まれる正解ラベルとを基にして、前記機械学習モ

デルのパラメータを更新するか否かを判定する処理を更にコンピュータに実行させることを特徴とする請求項 1 に記載の機械学習プログラム。

【請求項 4】

複数のデータを機械学習モデルに入力して、前記複数のデータの複数の予測結果を取得し、

前記複数のデータのうち前記複数の予測結果が第 2 のグループを示す第 1 の複数のデータから、予測結果が第 1 のグループを示す第 1 のデータに類似する第 2 のデータを選択し、前記第 1 のデータの特徴量と前記第 2 のデータの特徴量との間の特徴量に対応する一又は複数のデータを生成し、

前記機械学習モデルのパラメータに基づいて得られた、前記複数のデータと前記一又は複数のデータとのそれぞれの複数の特徴量に基づいて、前記複数のデータと前記一又は複数のデータとのクラスタリングを実行し、

前記クラスタリングの結果を正解ラベルとする前記複数のデータと前記一又は複数のデータとを含む訓練データに基づいて、前記機械学習モデルのパラメータを更新する、

処理をコンピュータが実行することを特徴とする機械学習方法。

【請求項 5】

複数のデータを機械学習モデルに入力して、前記複数のデータの複数の予測結果を取得し、

前記複数のデータのうち前記複数の予測結果が第 2 のグループを示す第 1 の複数のデータから、予測結果が第 1 のグループを示す第 1 のデータに類似する第 2 のデータを選択し、前記第 1 のデータの特徴量と前記第 2 のデータの特徴量との間の特徴量に対応する一又は複数のデータを生成し、

前記機械学習モデルのパラメータに基づいて得られた、前記複数のデータと前記一又は複数のデータとのそれぞれの複数の特徴量に基づいて、前記複数のデータと前記一又は複数のデータとのクラスタリングを実行し、

前記クラスタリングの結果を正解ラベルとする前記複数のデータと前記一又は複数のデータとを含む訓練データに基づいて、前記機械学習モデルのパラメータを更新する、

処理を実行する制御部を有する機械学習装置。

【発明の詳細な説明】

【技術分野】

【0001】

本発明は、訓練データを用いた機械学習の技術に関する。

【背景技術】

【0002】

企業等の情報システムでは、データの判定や分類を行う場合に、機械学習モデルを利用する。機械学習モデルは、機械学習に利用した訓練データに基づいて判定や分類を行うため、運用中にデータの傾向が変化すると、機械学習モデルの性能が低下する。

【0003】

機械学習モデルの性能を維持するため、機械学習モデルの正解率等が低下した場合には、正解ラベル付けを行ったデータを生成し、機械学習モデルの機械学習を再度実行している。

【0004】

図 18 および図 19 は、自動で正解ラベルをデータに付与する従来技術を説明するための図である。従来技術を実行する装置を「従来装置」と表記する。

【0005】

運用を開始する前の処理を、図 18 について説明する。図 18 の縦軸は、特徴空間のデータの密度に対応する軸である。横軸は、特徴量（特徴空間の座標）に対応する軸である。線 1 は、特徴空間の座標と、座標に対応するデータの密度との関係を示す。従来装置は、傾向が変化する前のデータを特徴空間に写像して、写像した各データの密度を計算する。従来装置は、クラスタリングを実行し、クラスタ数と、各クラスタ中で密度が閾値  $D_{th}$

10

20

30

40

50

以上となる領域の中心座標を記録する。

【0006】

図18に示す例では、特徴空間のデータが、クラスAおよびクラスBに分類されている。クラスAについて、密度が閾値 $D_{th}$ 以上となる領域の中心座標を $X_A$ とする。クラスBについて、密度が閾値 $D_{th}$ 以上となる領域の中心座標を $X_B$ とする。この場合には、従来装置は、クラス数「2」と、クラスAの中心座標 $X_A$ と、クラスBの中心座標 $X_B$ とを記録する。

【0007】

運用を開始した後の処理を、図19について説明する。図19の縦軸は、特徴空間のデータの密度に対応する軸である。横軸は、特徴量（特徴空間の座標）に対応する軸である。線2は、特徴空間の座標と、座標に対応するデータの密度との関係を示す。従来装置は、運用を開始した後に、運用で使用しているデータを特徴空間に写像し、写像した各データの密度を計算する。従来装置は、密度の閾値を下げていき、クラス数が、運用を開始する前に記録したクラス数と同じになる最小の閾値を探索する。

【0008】

図18で説明した例を用いて、運用を開始する前のクラス数を「2」とする。従来装置は、密度の閾値を徐々に下げていき、閾値を $D$ に設定することで、運用で使用しているデータ（特徴空間に写像したデータ）のクラス数を「2」に調整する。従来装置は、領域2-1に含まれるデータと、領域2-2に含まれるデータとをそれぞれ抽出（クラスタリング）する。

【0009】

従来装置は、運用前に記憶しておいた中心座標と、運用を開始した後のクラスAの中心座標との移動距離の合計等に基づくマッチングを行うことで、データに正解ラベルを付与する。たとえば、かかるマッチングによって、領域2-1のクラスAが、クラスAに対応付けられ、領域2-2のクラスBが、クラスBに対応付けられる。この場合、従来装置は、領域2-1の各データに、正解ラベル「クラスA」を付与し、領域2-2の各データに、正解ラベル「クラスB」を付与する。

【先行技術文献】

【特許文献】

【0010】

【文献】国際公開第2021/079442号

【発明の概要】

【発明が解決しようとする課題】

【0011】

しかしながら、上述した従来技術では、あるクラスに属するデータの数が少ない場合には、自動的に正解ラベルを付与することができないという問題がある。

【0012】

図20は、従来技術の課題を説明するための図である。図20の縦軸は、特徴空間のデータの密度に対応する軸である。横軸は、特徴量（特徴空間の座標）に対応する軸である。線3は、特徴空間の座標と、座標に対応するデータの密度との関係を示す。図20に示す例では、データを機械学習モデルに入力した場合に、データが「正常データ」または「異常データ」のいずれかのクラスに分類されるものとする。

【0013】

図20では、領域3-1に含まれるデータが「正常データ」のクラスに属し、領域3-2に含まれるデータが「異常データ」のクラスに属するものとする。領域3-2に含まれるデータが極端に少ないと、図19で説明したように、閾値を下げて、クラス数が、運用を開始する前に記録したクラス数と同じにならず、クラスタリングを正しく行うことができない。このため、あるクラスに属するデータ数が少ない場合には、自動的に正解ラベルを付与することができない。

【0014】

なお、運用中のデータを機械学習モデルに入力した際に分類されるクラス間のサンプル数が極端に異なる場合も、クラスタリングを正しく行えず、自動的に正解ラベルを付与することができない。

【0015】

1つの側面では、本発明は、あるクラスに属するデータの数が少ない場合でも、自動的に正解ラベルを付与することができる機械学習プログラム、機械学習方法および機械学習装置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】

【0016】

第1の案では、機械学習プログラムは、複数のデータを機械学習モデルに入力して、複数のデータの複数の予測結果を取得する処理をコンピュータに実行させる。機械学習プログラムは、複数のデータのうち予測結果が第1のグループを示す第1のデータに基づいて、一又は複数のデータを生成する処理をコンピュータに実行させる。機械学習プログラムは、機械学習モデルのパラメータに基づいて得られた、複数のデータと一又は複数のデータとのそれぞれの複数の特徴量に基づいて、複数のデータと一又は複数のデータとのクラスタリングを実行する処理をコンピュータに実行させる。機械学習プログラムは、クラスタリングの結果を正解ラベルとする複数のデータと一又は複数のデータとを含む訓練データに基づいて、機械学習モデルのパラメータを更新する処理をコンピュータに実行させる。

【発明の効果】

【0017】

あるクラスに属するデータの数が少ない場合でも、自動的に正解ラベルを付与することができる。

【図面の簡単な説明】

【0018】

【図1】図1は、疑似異常データを生成する際のアプローチと課題を説明するための図である。

【図2】図2は、疑似異常データを生成する処理を説明するための図である。

【図3】図3は、本実施例に係る機械学習装置の構成を示す機能ブロック図である。

【図4】図4は、訓練データのデータ構造の一例を示す図である。

【図5】図5は、ラベル付与部の処理を説明するための図(1)である。

【図6】図6は、ラベル付与部の処理を説明するための図(2)である。

【図7】図7は、ラベル付与部の処理を説明するための図(3)である。

【図8】図8は、ラベル付与部の処理を説明するための図(4)である。

【図9】図9は、ラベル付与部の処理を説明するための図(5)である。

【図10】図10は、劣化検出部の劣化判定を説明するための図である。

【図11】図11は、本実施例に係る機械学習装置の処理手順を示すフローチャートである。

【図12】図12は、外部環境の変化によるデータの傾向の変化を示す図である。

【図13】図13は、検証結果を示す図(1)である。

【図14】図14は、カメラのAUCスコアの推移の一例を示す図である。

【図15】図15は、異なる生成方法によって生成したデータの一例を示す図である。

【図16】図16は、検証結果を示す図(2)である。

【図17】図17は、実施例の機械学習装置と同様の機能を実現するコンピュータのハードウェア構成の一例を示す図である。

【図18】図18は、自動で正解ラベルをデータに付与する従来技術を説明するための図(1)である。

【図19】図19は、自動で正解ラベルをデータに付与する従来技術を説明するための図(2)である。

【図20】図20は、従来技術の課題を説明するための図である。

【発明を実施するための形態】

## 【0019】

以下に、本願の開示する機械学習プログラム、機械学習方法および機械学習装置の実施例を図面に基づいて詳細に説明する。なお、この実施例によりこの発明が限定されるものではない。

## 【実施例】

## 【0020】

本実施例に係る機械学習装置は、入力されたデータを異常クラスまたは正常クラスのうちのいずれかのクラスに分類する機械学習モデルを利用するものとする。たとえば、機械学習モデルに入力するデータは、画像データ等である。機械学習モデルは、DNN（Deep Neural Network）等である。正常クラスに分類されたデータを「正常データ」と表記する。異常クラスに分類されたデータを「異常データ」と表記する。

10

## 【0021】

機械学習装置は、運用時に分類した異常データと正常データとを用いて疑似的な異常データを生成し、疑似的な異常データを含めてクラスタリングを実行することで、自動的にデータに正解ラベルを付与する。以下の説明では、疑似的な異常データを「疑似異常データ」と表記する。

## 【0022】

図1は、疑似異常データを生成する際のアプローチと課題を説明するための図である。疑似異常データを生成する場合に、生成方法によっては、自動的に正解ラベルを付与することができない場合がある。

20

## 【0023】

図1において、グラフG1、G2、G3の縦軸はデータの特徴空間のデータの密度に対応する軸である。横軸は、特徴量（特徴空間の座標）に対応する軸である。たとえば、データを機械学習モデルに入力し、機械学習モデルの出力層よりも所定数前の層から出力されるベクトルが特徴量となる。特徴量に応じて、データの特徴空間上の座標が決まる。

## 【0024】

グラフG1において、分布dis1aは、「正常データの分布」を示す。特徴空間における正常データの図示を省略する。分布dis1bは、「真の異常データの分布」を示す。たとえば、特徴空間における異常データを、異常データ10、11、12、13、14とする。異常データの数が少ないと、異常データの分布は、分布dis1bのようにならず、図20で説明したように、自動的に正解ラベルを付与することができない。

30

## 【0025】

ここで、異常データの数を増やすために、単純に、異常データ10、11、12、13、14に対応するデータ（画像データ）と同一の画像データを複製すると、異常データの分布は、グラフG2に示す、分布dis2a、dis2b、dis2c、dis2d、dis2eとなる。分布dis2a、dis2b、dis2c、dis2d、dis2eは、真の異常データの分布dis1bとは異なるため、クラスタリングが失敗し、自動的にデータに正解ラベルを付与することができない。

## 【0026】

一方、機械学習装置は、図1のグラフG3に示すように、異常データの分布が、真の異常データの分布dis1bに近づくように、疑似異常データを生成する。たとえば、後述する図2で説明する処理を、機械学習装置が実行し、疑似異常データを生成することで、異常データの分布は、分布dis3となる。

40

## 【0027】

図2は、疑似異常データを生成する処理を説明するための図である。たとえば、機械学習装置は、疑似異常データを生成する場合に、ステップS1、S2の順に処理を実行する。

## 【0028】

機械学習装置が実行するステップS1の処理について説明する。機械学習装置は、運用データに含まれる複数のデータを、特徴空間Fに写像する。たとえば、機械学習装置は、データを機械学習モデルに入力し、機械学習モデルの出力層よりも所定数前の層から出力

50

される特徴量を、データを写像した値とする。特徴量により、特徴空間 F の座標が決まる。特徴空間 F に写像された異常データを、異常データ 20, 21 とする。特徴空間 F の写像された正常データを、正常データ 31, 32, 33, 34, 35, 36, 37, 38, 39 とする。異常データ 20, 21、正常データ 30~39 を用いて、機械学習装置の処理について説明する。

【0029】

機械学習装置は、特徴空間 F において、異常データと類似する正常データを選択する。特徴空間 F において、異常データとの距離が閾値未満となる正常データを、異常データに類似する正常データとする。

【0030】

機械学習装置は、異常データ 20 と、正常データ 30~39 とを比較して、異常データ 20 に類似する正常データ 30, 31, 32, 34 を選択する。機械学習装置は、異常データ 21 と、正常データ 30~39 とを比較して、異常データ 21 に類似する正常データ 30, 32, 33, 35 を選択する。

【0031】

機械学習装置が実行するステップ S2 の処理について説明する。機械学習装置は、ステップ S1 で選択した正常データそれぞれに対し、割合  $\alpha$  を一様乱数とし、異常データと正常データとの線形結合により合成して、疑似異常データを生成する。たとえば、機械学習装置は、 $\alpha$  ブレンディング等を用いて、疑似異常データを生成する。

【0032】

機械学習装置は、異常データ 20 と、正常データ 30 とを結ぶ線分を「 $1-\alpha:\alpha$ 」で分割した座標（特徴量）に対応する、疑似異常データ 51 を生成する。機械学習装置は、異常データ 20 と、正常データ 34 とを結ぶ線分を「 $1-\alpha:\alpha$ 」で分割した座標（特徴量）に対応する、疑似異常データ 52 を生成する。機械学習装置は、異常データ 20 と、正常データ 32 とを結ぶ線分を「 $1-\alpha:\alpha$ 」で分割した座標（特徴量）に対応する、疑似異常データ 53 を生成する。機械学習装置は、異常データ 20 と、正常データ 31 とを結ぶ線分を「 $1-\alpha:\alpha$ 」で分割した座標（特徴量）に対応する、疑似異常データ 54 を生成する。

【0033】

機械学習装置は、異常データ 21 と、正常データ 30 とを結ぶ線分を「 $1-\alpha:\alpha$ 」で分割した座標（特徴量）に対応する、疑似異常データ 55 を生成する。機械学習装置は、異常データ 21 と、正常データ 32 とを結ぶ線分を「 $1-\alpha:\alpha$ 」で分割した座標（特徴量）に対応する、疑似異常データ 56 を生成する。機械学習装置は、異常データ 21 と、正常データ 35 とを結ぶ線分を「 $1-\alpha:\alpha$ 」で分割した座標（特徴量）に対応する、疑似異常データ 57 を生成する。機械学習装置は、異常データ 21 と、正常データ 33 とを結ぶ線分を「 $1-\alpha:\alpha$ 」で分割した座標（特徴量）に対応する、疑似異常データ 58 を生成する。

【0034】

機械学習装置が、図 2 で説明した処理を実行して疑似異常データを生成すると、疑似異常データを含む異常データの分布が、図 1 で説明した分布  $d_{is3}$  となる。このため、機械学習装置は、正常データ、異常データ、疑似異常データの各特徴量を基にして、クラスタリングを実行すると、訓練データの特徴量に基づくクラスタリング結果と対応付けることができる。したがって、あるクラスに属するデータ（たとえば、異常データ）の数が少ない場合でも、自動的に正解ラベルを付与することができる。

【0035】

次に、本実施例に係る機械学習装置の構成の一例について説明する。図 3 は、本実施例に係る機械学習装置の構成を示す機能ブロック図である。図 3 に示すように、この機械学習装置 100 は、通信部 110 と、入力部 120 と、表示部 130 と、記憶部 140 と、制御部 150 を有する。

【0036】

10

20

30

40

50

通信部 110 は、ネットワークを介して、外部装置との間でデータ通信を行う。通信部 110 は、外部装置から、訓練データ 141、運用データ 143 等を受信する。機械学習装置 100 は、訓練データ 141、運用データ 143 を、後述する入力部 120 から受け付けてもよい。

【0037】

入力部 120 は、データを入力するためのインタフェースである。入力部 120 は、マウス、およびキーボードなどの入力装置を介してデータの入力を受け付ける。

【0038】

表示部 130 は、データを出力するためのインタフェースである。たとえば、表示部 130 は、ディスプレイなどの出力装置にデータを出力する。

【0039】

記憶部 140 は、訓練データ 141、機械学習モデル 142、運用データ 143、再訓練データ 144、クラスタ関連データ 145 を有する。記憶部 140 は、メモリ等の記憶装置の一例である。

【0040】

訓練データ 141 は、機械学習モデル 142 の機械学習を実行する場合に用いられる。図 4 は、訓練データのデータ構造の一例を示す図である。図 4 に示すように、訓練データは、項番と、データと、正解ラベルとを対応付ける。項番は、訓練データ 141 のレコードを識別する番号である。データは、画像データである。正解ラベルは、データが正常であるか、異常であるかを示すラベルである。

【0041】

たとえば、項番「1」のデータの正解ラベルは「正常」であるため、項番「1」のデータは、正常データである。項番「3」のデータの正解ラベルは「異常」であるため、項番「3」のデータは、異常データである。

【0042】

機械学習モデル 142 は、DNN 等であり、入力層、隠れ層、出力層を有する。機械学習モデル 142 は、誤差逆伝播法等に基づいて機械学習が実行される。

【0043】

機械学習モデル 142 に、データを入力すると、入力されたデータが正常であるか異常であるかの分類結果が出力される。

【0044】

運用データ 143 は、運用時に利用する複数のデータを含むデータセットである。

【0045】

再訓練データ 144 は、機械学習モデル 142 の機械学習を再度実行する場合に用いられる訓練データである。

【0046】

クラスタ関連データ 145 は、訓練データ 141 に含まれる各データを特徴空間に写像した場合における、クラスタ数と、各クラスタ中で密度が閾値以上となる領域の中心座標とを有する。また、クラスタ関連データ 145 は、後述するラベル付与部 156 のクラスタリング結果に基づく各クラスタの中心座標を有する。

【0047】

制御部 150 は、取得部 151、機械学習部 152、事前処理部 153、推論部 154、生成部 155、ラベル付与部 156、劣化検出部 157 を有する。

【0048】

取得部 151 は、外部装置または入力部 120 から、訓練データ 141 を取得し、訓練データ 141 を記憶部 140 に格納する。取得部 151 は、外部装置または入力部 120 から、運用データ 143 を取得し、運用データ 143 を記憶部 140 に格納する。

【0049】

機械学習部 152 は、訓練データ 141 を用いて、誤差逆伝播法により、機械学習モデル 142 の機械学習を実行する。機械学習部 152 は、訓練データ 141 の各データを、

10

20

30

40

50

機械学習モデル 1 4 2 の入力層に入力した場合に、出力層から出力される出力結果が、入力したデータの正解ラベルに近づくように、機械学習モデル 1 4 2 を訓練する。機械学習部 1 5 2 は、検証データを用いて、機械学習モデル 1 4 2 の検証を行う。

【 0 0 5 0 】

事前処理部 1 5 3 は、訓練データ 1 4 1 のデータを特徴空間に写像して、クラスタリングを実行することで、運用を開始する前のデータのクラスタ数と、クラスタ中で密度が閾値以上となる領域の中心座標を特定する。事前処理部 1 5 3 は、クラスタ数と、各クラスタの中心座標を、クラスタ関連データ 1 4 5 に記録する。

【 0 0 5 1 】

事前処理部 1 5 3 は、訓練データ 1 4 1 に含まれる各データを、特徴空間に写像する。たとえば、事前処理部 1 5 3 は、訓練データ 1 4 1 の各データを、機械学習モデル 1 4 2 に入力し、機械学習モデル 1 4 2 の出力層よりも所定数前の層から出力される特徴量を、データを写像した値とする。この特徴量は、訓練された機械学習モデル 1 4 2 のパラメータに基づいて得られる値である。特徴量により、特徴空間 F の座標が決まる。

【 0 0 5 2 】

事前処理部 1 5 3 は、式 ( 1 ) を用いて、特徴空間におけるデータの密度を算出する。式 ( 1 ) において、N はデータの総数を示し、 $\sigma$  は標準偏差を示す。x は、データの特徴量の期待値 ( 平均値 ) であり、 $x_j$  は、j 番目のデータの特徴量を示す。

【 0 0 5 3 】

【 数 1 】

$$\text{Gauss\_density}(x) = \frac{1}{N(\sqrt{2\pi}\sigma)^d} \sum_{j=0}^{N-1} \exp\left(-\frac{\|x - x_j\|^2}{2\sigma^2}\right) \quad \dots (1)$$

【 0 0 5 4 】

ここでは、事前処理部 1 5 3 は、データの密度として、ガウス密度を算出する場合について説明したが、これに限定されるものではなく、eccentricity や、K N N 距離 ( K-Nearest Neighbor Algorithm ) 等を用いて密度を計算してもよい。

【 0 0 5 5 】

事前処理部 1 5 3 は、縦軸を密度、横軸を特徴量とするグラフを生成する。事前処理部 1 5 3 によって生成されるグラフは、図 1 8 で説明したグラフに対応する。事前処理部 1 5 3 は、クラスタリングを実行し、クラスタ数と、各クラスタ中で密度が閾値  $D_{th}$  以上となる領域の中心座標を記録する。

【 0 0 5 6 】

図 1 8 に示す例では、特徴空間のデータが、クラスタ A およびクラスタ B に分類されている。たとえば、クラスタ A は、正常データの属するクラスタである。クラスタ B は、異常データの属するクラスタである。クラスタ A について、密度が閾値  $D_{th}$  以上となる領域の中心座標を  $X_A$  とする。クラスタ B について、密度が閾値  $D_{th}$  以上となる領域の中心座標を  $X_B$  とする。この場合には、事前処理部 1 5 3 は、クラスタ数「2」と、クラスタ A の中心座標  $X_A$  と、クラスタ B の中心座標  $X_B$  とを、クラスタ関連データ 1 4 5 に記録する。

【 0 0 5 7 】

ここでは、事前処理部 1 5 3 が、クラスタ数と、各クラスタの中心座標を特定する場合について説明したが、クラスタ数および各クラスタの中心座標を、外部装置から事前に取得しておいてもよい。

【 0 0 5 8 】

推論部 1 5 4 は、運用データ 1 4 3 からデータを取得し、取得したデータを機械学習モデル 1 4 2 に入力することで、入力したデータが、正常データであるか、異常データであるかを推論する。推論部 1 5 4 は、運用データ 1 4 3 に含まれる各データについて、上記処

10

20

30

40

50



理を繰り返し実行する。推論部 154 は、運用データ 143 の各データについて、データが正常データであるか異常データであるかの推定結果を設定し、生成部 155 に出力する。推論部 154 は、推論結果を、表示部 130 に出力して、推論結果を表示させてもよい。

【0059】

生成部 155 は、図 2 で説明した処理を実行することで、疑似異常データを生成する。以下において、生成部 155 の処理の一例について説明する。

【0060】

生成部 155 は、運用データ 143 に含まれる複数のデータを、特徴空間 F に写像する。たとえば、生成部 155 は、データを機械学習モデル 142 に入力し、機械学習モデル 142 の出力層よりも所定数前の層から出力される特徴量を、データを写像した値とする。この特徴量は、訓練された機械学習モデル 142 のパラメータに基づいて得られる値である。たとえば、特徴空間に写像された異常データ、正常データは、図 2 に示す、異常データ 20、21、正常データ 30～39 となる。生成部 155 は、データが異常データであるか、正常データであるかを、推論部 154 の推論結果を基にして特定する。

【0061】

生成部 155 は、特徴空間 F において、異常データと類似する正常データを選択する。特徴空間 F において、異常データとの距離が閾値未満となる正常データを、異常データに類似する正常データとする。たとえば、生成部 155 は、図 2 において、異常データ 20 に類似する正常データとして、正常データ 30、31、32、34 を選択する。生成部 155 は、異常データ 21 に類似する正常データとして、正常データ 30、32、33、35 を選択する。

【0062】

生成部 155 は、上記処理によって選択した正常データそれぞれに対し、割合  $\alpha$  を一様乱数とし、異常データと正常データとの線形結合により合成して、疑似異常データを生成する。たとえば、生成部 155 は、 $\alpha$  ブレンディング等を用いて、疑似異常データを生成する。生成部 155 は、図 2 で説明した処理を実行することで、疑似異常データ 51～58 を生成する。

【0063】

生成部 155 は、異常データの特徴量、正常データの特徴量、疑似異常データの特徴量を、ラベル付与部 156 に出力する。

【0064】

ラベル付与部 156 は、異常データの特徴量、正常データの特徴量、疑似異常データの特徴量に基づいて、クラスタリングを実行し、クラスタリング結果に応じて、データに正解ラベルを付与する。ラベル付与部 156 は、正解ラベルを付与した各データを、再訓練データ 144 として、記憶部 140 に登録する。以下において、ラベル付与部 156 の処理の一例について説明する。ラベル付与部 156 は、 $\alpha$  ブレンディングによって生成した疑似異常データについても、正解ラベルを付与して、再訓練データ 144 に登録する。

【0065】

ラベル付与部 156 が実行するクラスタリング処理を実行する。図 5 は、ラベル付与部の処理を説明するための図 (1) である。ラベル付与部 156 は、異常データの特徴量、正常データの特徴量、疑似異常データの特徴量に基づいて、縦軸を密度、横軸を特徴量とするグラフ G10 を生成する (ステップ S10)。ラベル付与部 156 は、事前処理部 153 と同様にして、式 (1) を基にして、データ (正常データ、異常データおよび疑似異常データ) の密度を算出する。

【0066】

ラベル付与部 156 は、密度に対応する閾値を所定値ごとに下げていき、クラスタ関連データ 145 に記録された事前のクラスタ数と同じになる最小の閾値を探索する (ステップ S11)。ここでは、クラスタ関連データ 145 に記録された事前のクラスタ数を「2」とする。

【0067】

ラベル付与部 156 は、閾値以上であるデータの特徴量に対してパーシステントホモロジ変換（PH変換）を実行して、0次元の連結成分を参照する。ラベル付与部 156 は、予め定めた閾値以上の半径を有するバー（bar）の数が事前に設定したクラスタ数と一致するか否かにより、クラスタの計算および特定を実行する（ステップ S12）。

【0068】

ラベル付与部 156 は、閾値を超えるバーの数が事前のクラスタ数と一致しない場合は、閾値を所定値下げて、処理を繰り返す（ステップ S13）。

【0069】

上記のように、ラベル付与部 156 は、密度の閾値を下げて密度が閾値以上のデータを抽出する処理と、抽出されたデータに対する PH 変換処理によりクラスタ数を計算する処理とを、事前のクラスタ数と一致するまで繰り返す。ラベル付与部 156 は、クラスタ数が一致した場合に、その時の閾値（密度）以上の密度を有するデータ領域の中心座標 C1、C2 を特定し、クラスタ関連データ 145 に記録する。ラベル付与部 156 は、クラスタリング処理を行うたびに、中心座標を、クラスタ関連データ 145 に記録する。

【0070】

ラベル付与部 156 が実行する PH 変換は、たとえば、特許文献 1（国際公開第 2021/079442 号）に記載された PH 変換である。

【0071】

ラベル付与部 156 は、上記のクラスタリング処理の結果を基にして、運用データ 143 に含まれる各データに正解ラベルを付与する。ラベル付与部 156 は、クラスタリング処理によって決定された密度が閾値以上のデータに対して、それぞれが属するクラスタに基づく正解ラベル付けを行うことで、再訓練データ 144 を生成する。

【0072】

図 6 は、ラベル付与部の処理を説明するための図（2）である。図 6 のグラフ G10 に関する説明は、図 5 のグラフ G10 に関する説明と同様である。ラベル付与部 156 は、上記のクラスタリング処理を実行することで、クラスタ数が 2 の状態で最小となった閾値以上となったデータと、2 つの中心座標 C1、C2 を特定する。ラベル付与部 156 は、クラスタ関連データ 145 に記録された中心座標の履歴と、マッチング処理に基づき、2 つの中心座標それぞれが属するクラスタを決定する。

【0073】

図 7 は、ラベル付与部の処理を説明するための図（3）である。図 7 を用いて、マッチング処理の一例について説明する。ラベル付与部 156 は、機械学習モデル 142 の訓練が完了してから、現在に至るまでに特定された各クラスタの中心座標を特徴空間にマッピングし、進行方向を推定して、現在抽出された 2 つの中心座標（C1、C2）それぞれのクラスタを決定する。

【0074】

単に一番近い中心座標のマッチングでは、中心座標の変動を加味すると妥当でないことがある。図 7 の（a）のように、中心座標が変動していて、新しい 2 点を新たにマッチングする場合、近い点でマッチングすると図 7 の（b）のようになるが、これは変動の方向からは不自然な動きである。図 7 の（c）のように、変動する方が自然である。

【0075】

このため、ラベル付与部 156 は、補正距離を導入する。たとえば、進行方向に進む場合はより近い点と判定する仕組みを導入し、前回の座標からの進行方向ベクトルと、前回の座標から今回の座標を結ぶベクトルとの内積を計算することで、進行方向を特定する。ラベル付与部 156 は、内積の値を  $c$  として、 $(\tan(c) + 1) / 2$  を重みとして 2 点間の距離に乘算した値を補正距離として、最近傍点を選択する。たとえば、中心座標 Cb1 および中心座標 C1 の距離に、ベクトル  $v1$  およびベクトル  $v2$  の内積  $c$  に基づく重み  $((\tan(c) + 1) / 2)$  を乗算した値が、補正距離となる。

【0076】

ラベル付与部 156 は、クラスタの中心座標が特定される度に、中心座標間の補正距離

10

20

30

40

50

を算出し、補正距離の近い中心座標同士をマッチングする処理を繰り返し実行する。

【0077】

たとえば、事前処理部153のクラスタリング結果により特定されたクラスAの中心座標をCb3-1とし、クラスBの中心座標をCb3-2とする。ラベル付与部156によって、中心座標Cb3-1とCb2-1とがマッチングされ、中心座標Cb2-1とCb1-1とがマッチングされ、中心座標Cb1-1とC1とがマッチングされたとすると、中心座標C1は、クラスAに対応付けられる。本実施例では、クラスAに対応するクラスを「正常クラス」とする。

【0078】

ラベル付与部156によって、中心座標Cb3-2とCb2-2とがマッチングされ、中心座標Cb2-2とCb1-2とがマッチングされ、中心座標Cb1-2とC2とがマッチングされたとすると、中心座標C2は、クラスBに対応付けられる。本実施例では、クラスBに対応するクラスを「異常クラス」とする。

【0079】

図6の説明に戻る。図6に示す例では、中心座標C1にクラスA（正常クラス）が対応付けられる。中心座標C2にクラスB（異常クラス）が対応付けられる。この場合、ラベル付与部156は、運用データ143に含まれるデータのうち、密度が閾値以上かつ、中心座標C1と同じクラスAに属するデータに、正解ラベル「正常」を設定する。一方、ラベル付与部156は、運用データ143に含まれるデータのうち、密度が閾値以上かつ、中心座標C2と同じクラスBに属するデータに、正解ラベル「異常」を設定する。

【0080】

続いて、ラベル付与部156は、クラスタリング処理によって抽出されなかった閾値未満のデータそれぞれに正解ラベルを付与する。図8は、ラベル付与部の処理を説明するための図(4)である。ラベル付与部156は、抽出されなかった各データについて、各クラスの中心座標C1との距離およびC2との距離をそれぞれ計測し、2番目に近い距離が各クラスの中心間の距離の最大値より大きい場合は、一番近いクラスAに属するデータと決定する。

【0081】

図8の例の場合、ラベル付与部156は、上記手法によりクラスAが決定された領域X（クラスA）と領域Y（クラスB）以外の領域のうち、領域Xよりも外側の領域Pのデータについては、クラスAと決定する。ラベル付与部156は、領域Yよりも外側の領域Qのデータについては、クラスBと決定する。

【0082】

ラベル付与部156は、2番目に近い距離が各クラスの中心間の距離の最大値より小さい（複数のクラスAの中間にある）領域Zのデータについては、近くにある複数のクラスAのデータが混在していると判定する。この場合、ラベル付与部156は、各データに関して各クラスの確率を測定して付与する。たとえば、ラベル付与部156は、k近傍法、一様確率法、分布比率保持法などを用いて、領域Zに属する各データについて、各クラスAに属する確率を算出し、確率的なラベル（正常クラスAの確率、異常クラスBの確率、他のクラスCの確率）を生成して付与する。

【0083】

ラベル付与部156は、領域Zに属する各入力データに対して、そのデータに近傍に位置するすでにラベル付けされたデータをk個抽出し、その割合が正常クラス=0.6、異常クラス=0.4、他のクラス=0であれば、その割合をラベルとして付与する。

【0084】

ラベル付与部156は、領域Zに属する各データに対して、各クラスAにすべて同じ確率を付与する。例えば、ラベル付与部156は、2クラス分類の場合には、正常クラス=0.5、異常クラス=0.5をラベルとして付与し、3クラス分類の場合には、正常クラス=0.3、異常クラス=0.3、他のクラス=0.3などをラベルとして付与する。

【0085】

上述した手法により推定して、ラベル付与部 156 が、各データに付与する正解ラベルの情報が図 9 である。図 9 は、ラベル付与部の処理を説明するための図 (5) である。推定された正解ラベルは、各クラスに属する確率 (正常クラスに属する確率, 異常クラスに属する確率, 他のクラスに属する確率) で付与される。図 9 に示すように、領域 X と領域 P の各データには、推定ラベル (正解ラベル) [1, 0, 0] が付与され、領域 Y と領域 Q の各入力データには、推定ラベル [0, 1, 0] が付与され、領域 Z の各入力データには、推定ラベル [a, b, c] が付与される。なお、a, b, c は、k 近傍法などの手法により算出される確率である。そして、ラベル付与部 156 は、各データと推定ラベルとの対応付けた再訓練データ 144 を、記憶部 140 に格納する。

【0086】

図 3 の説明に戻る。劣化検出部 157 は、機械学習モデル 142 の精度劣化を検出する。たとえば、劣化検出部 157 は、機械学習モデル 142 の判定結果と、ラベル付与部 156 により生成された推定結果 (再訓練データ 144) とを比較して、機械学習モデル 142 の精度劣化を検出する。

【0087】

図 10 は、劣化検出部の劣化判定を説明するための図である。図 10 に示すように、劣化検出部 157 は、データ (運用データ 143 のデータ) を機械学習モデル 142 に入力した場合の出力結果 (正常クラス) に基づき、判定結果 [1, 0, 0] を生成する。一方で、劣化検出部 157 は、データに対する上記推定処理により、領域 X または領域 P に属した場合の推定結果 [1, 0, 0]、領域 Y または領域 Q に属した場合の推定結果 [0, 1, 0]、または、領域 Z に属した場合の推定結果 [a, b, c] を取得する。

【0088】

劣化検出部 157 は、各入力データについて、判定結果と推定結果とを取得し、これらの比較により劣化判定を実行する。例えば、劣化検出部 157 は、各推定結果で示される各データ (各点) の確率ベクトルに対し、機械学習モデル 142 による判定結果のベクトル表示の成分積の和 (内積) をその点のスコアとし、そのスコアの合計をデータ数で割った値と閾値との比較により、劣化判定を実行する。

【0089】

なお、劣化検出部 157 は、次の処理を実行して、機械学習モデル 142 の精度劣化を検出してもよい。劣化検出部 157 は、クラスタ関連データ 145 を参照し、訓練データ 141 のクラスターリング処理によって特定されるクラスタ A の中心座標と、現在の運用データ 143 のクラスターリング処理によって特定されるクラスタ A の中心座標との距離の逆数を、スコアとして算出する。劣化検出部 157 は、かかるスコアが閾値未満である場合に、機械学習モデル 142 の精度が劣化したと判定する。

【0090】

劣化検出部 157 は、機械学習モデル 142 の精度劣化を検出した場合に、機械学習部 152 に対して、機械学習の再実行依頼を出力する。機械学習部 152 は、劣化検出部 157 から、機械学習の再実行依頼を受け付けた場合、再訓練データ 144 を用いて、機械学習モデル 142 の機械学習を再度実行する。

【0091】

次に、本実施例に係る機械学習装置 100 の処理手順について説明する。図 11 は、本実施例に係る機械学習装置の処理手順を示すフローチャートである。図 11 に示すように、機械学習装置 100 の機械学習部 152 は、訓練データ 141 を用いて、機械学習モデル 142 の機械学習を実行する (ステップ S101)。

【0092】

機械学習装置 100 の事前処理部 153 は、訓練データ 141 を基にして、クラスタ数、各クラスタの中心座標を特定し、クラスタ関連データ 145 に記録する (ステップ S102)。機械学習装置 100 の取得部 151 は、運用データ 143 を取得し、記憶部 140 に格納する (ステップ S103)。

【0093】

10

20

30

40

50

機械学習装置１００の推論部１５４は、運用データ１４３のデータを、機械学習モデル１４２に入力し、データのクラスを推定する（ステップＳ１０４）。機械学習装置１００の生成部１５５は、正常データの特徴量、異常データの特徴量を基にして、異常疑似データを生成する（ステップＳ１０５）。

【００９４】

機械学習装置１００のラベル付与部１５６は、正常データ、異常データ、疑似異常データの各特徴量を基にして、クラスタリング処理を実行する（ステップＳ１０６）。ラベル付与部１５６は、クラスタリング処理の結果を基にして、データに正解ラベルを付与し、再訓練データ１４４を生成する（ステップＳ１０７）。

【００９５】

機械学習装置１００の劣化検出部１５７は、機械学習モデル１４２の性能に関するスコアを算出する（ステップＳ１０８）。機械学習装置１００は、スコアが閾値未満でない場合には（ステップＳ１０９，Ｎｏ）、ステップＳ１０３に移行する。一方、機械学習装置１００は、スコアが閾値未満である場合には（ステップＳ１０９，Ｙｅｓ）、ステップＳ１１０に移行する。

【００９６】

機械学習部１５２は、再訓練データ１４４を基にして、機械学習モデル１４２の機械学習を再度実行し（ステップＳ１１０）、ステップＳ１０３に移行する。

【００９７】

次に、本実施例に係る機械学習装置１００の効果について説明する。機械学習装置１００は、訓練済みの機械学習モデル１４２に、運用データ１４３のデータを入力することで、正常データおよび異常データの特徴量を特定する。機械学習装置１００は、正常データおよび異常データの特徴量を基にして、疑似異常データを生成し、正常データ、異常データ、疑似異常データの各特徴量を基にして、クラスタリングを実行する。機械学習装置１００は、クラスタリング結果に基づく正解ラベルを、運用データおよび疑似異常データの各データに付与することで、再訓練データ１４４を生成し、再訓練データ１４４を基にして、機械学習モデルのパラメータを更新する。上記のように、正常データおよび異常データの特徴量を基にして、疑似異常データを生成することで、あるクラスに属するデータの数が少ない場合には、自動的に正解ラベルを付与することができる。

【００９８】

機械学習装置１００は、上記のように、自動的に正解ラベルを付与することで、自動的に再訓練データ１４４を生成でき、再訓練データ１４４を用いて、機械学習モデル１４２の機械学習を再度実行して、機械学習モデル１４２の精度劣化を抑止することができる。

【００９９】

機械学習装置１００は、特徴空間において、異常データに類似する正常データを選択し、異常データと選択した正常データとの間に、疑似異常データを生成する。これによって、特徴空間のデータの分布を、自動的に正解ラベルを付与することが可能な分布とすることができる。

【０１００】

次に、機械学習装置１００が実行するその他の処理（１）、（２）について説明する。

【０１０１】

その他の処理（１）について説明する。上述した機械学習装置１００は、特徴空間において、正常データと異常データの特徴量を基にして、疑似異常データを生成していたが、これに限定されるものではない。たとえば、機械学習装置１００の生成部１５５は、運用データ１４３に含まれるデータのうち、異常データを複製し、複製した異常データに、ガウシアンノイズ等のノイズを付与した異常データを生成してもよい。以下の説明では、ノイズを付与した異常データを、ノイズデータと表記する。

【０１０２】

機械学習装置１００のラベル付与部１５６は、異常データの特徴量、ノイズデータの特徴量、正常データの特徴量に基づいて、クラスタリング処理を実行し、クラスタリング結

10

20

30

40

50

果に応じて、データに正解ラベルを付与する。ノイズデータの特徴量は、ノイズデータを、訓練済みの機械学習モデル142に入力した場合に、機械学習モデル142の出力層よりも所定数前の層から出力される特徴量である。

#### 【0103】

その他の処理(2)について説明する。上述した機械学習装置100は、運用データ143について、異常データの数と、正常データの数とに差異がある場合に、疑似異常データを生成していたが、これに限定されるものではない。機械学習装置100の生成部155は、訓練データ141について、異常データの数と、正常データの数とに差異がある場合でも、訓練データ141の異常データの特徴量と、正常データの特徴量とを用いて、疑似異常データを生成し、機械学習モデル142の機械学習に利用してもよい。

10

#### 【0104】

次に、本実施例に係る機械学習装置100を、ある工場における異常検知AI(Artificial Intelligence)に適用して性能を検証した結果について説明する。検証条件として、機械学習モデルをDNNとし、データを異常データまたは正常データに分類するように、機械学習モデル142を事前に訓練する。

#### 【0105】

運用時想定シナリオとして、証明器具の寿命により、徐々に暗くなることを想定する。バッチ(batch)毎に10%ずつ明度が低下する。各バッチにおいて、運用データとして、正常データを80枚、異常データを5枚取得する。

#### 【0106】

図12は、外部環境の変化によるデータの傾向の変化を示す図である。正常データを、Im1-0~Im1-8とする。異常データを、Im2-0~Im2-8とする。正常データIm1-0は、0バッチ目の正常なデータ(元画像データ)である。異常データIm2-0は、0バッチ目の異常なデータ(元画像データ)である。

20

#### 【0107】

正常データIm1-1は、1バッチ目の正常なデータ(明度90%の画像)である。異常データIm2-1は、1バッチ目の異常なデータ(明度90%の画像データ)である。正常データIm1-2は、2バッチ目の正常なデータ(明度80%の画像)である。異常データIm2-2は、2バッチ目の異常なデータ(明度80%の画像データ)である。

#### 【0108】

正常データIm1-3は、3バッチ目の正常なデータ(明度70%の画像)である。異常データIm2-3は、3バッチ目の異常なデータ(明度70%の画像データ)である。正常データIm1-4は、4バッチ目の正常なデータ(明度60%の画像)である。異常データIm2-4は、4バッチ目の異常なデータ(明度60%の画像データ)である。

30

#### 【0109】

図示を省略するが、正常データIm1-5は、5バッチ目の正常なデータ(明度50%の画像)である。異常データIm2-5は、5バッチ目の異常なデータ(明度50%の画像データ)である。図示を省略するが、正常データIm1-6は、6バッチ目の正常なデータ(明度40%の画像)である。異常データIm2-6は、6バッチ目の異常なデータ(明度40%の画像データ)である。

40

#### 【0110】

正常データIm1-7は、7バッチ目の正常なデータ(明度30%の画像)である。異常データIm2-7は、7バッチ目の異常なデータ(明度30%の画像データ)である。正常データIm1-8は、8バッチ目の正常なデータ(明度20%の画像)である。異常データIm2-8は、8バッチ目の異常なデータ(明度20%の画像データ)である。

#### 【0111】

評価指標として、データを機械学習モデルに入力し、入力したデータが正常データであるか異常データであるかを判定し、各バッチにおけるAUC(Area Under Curve)スコアを算出する。AUCスコアが高いほど、機械学習モデルの検知性能が維持されていることを示す。工場の7か所のカメラ(カメラID1~7)の異常検知データセット(運用

50

データ) に対して、本実施例の機械学習装置 100 を適用した異常検知 AI と、再訓練を行わない異常検知 AI との AUC スコアは、図 13 に示す検証結果となった。

【0112】

図 13 は、検証結果を示す図 (1) である。図 13 の検証結果は、最終バッチ (8 バッチ目) における AUC スコアを示す。ベースラインは、再訓練を行わない異常検知 AI を示す。提案手法は、本実施例の機械学習装置 100 を適用した異常検知 AI を示す。図 13 に示すように、全てのカメラにおいて、提案手法の AUC スコアは、ベースラインの AUC スコアを上回っており、暗い状態 (データの傾向が変化) でも、検知性能を維持している。

【0113】

図 14 は、カメラの AUC スコアの推移の一例を示す図である。図 14 のグラフ G20 は、カメラ ID「3」の AUC スコアの推移を表す。グラフ G20 の縦軸は、AUC スコアに対応する軸であり、横軸はバッチ (batch number) に対応する軸である。線分 20a は、ベースラインの各バッチにおける AUC スコアの推移を示す。線分 20b は、提案手法の各バッチにおける AUC スコアの推移を示す。

【0114】

図 14 のグラフ G21 は、カメラ ID「6」の AUC スコアの推移を表す。グラフ G21 の縦軸は、AUC スコアに対応する軸であり、横軸はバッチ (batch number) に対応する軸である。線分 21a は、ベースラインの各バッチにおける AUC スコアの推移を示す。線分 21b は、提案手法の各バッチにおける AUC スコアの推移を示す。

【0115】

図 14 のグラフ G20、G21 に示すように、本実施例の機械学習装置 100 を適用した異常検知 AI は、暗い状態でも検知性能を維持している。

【0116】

続いて、生成方法 (1) ~ (5) によって、疑似異常データを生成し、機械学習モデル 142 の性能を検証した結果について説明する。前提条件として、データは、異常データまたは正常データに分類され、異常データの数が、正常データの数と比較して少ないものとする。

【0117】

- (1) 同一の異常データを複製する。
- (2) (1) で複製した異常データにガウシアンノイズ (ノイズ強度: 弱 < 標準偏差  $\sigma = 0.01$  >) を付加したノイズデータを生成する。
- (3) (1) で複製した異常データにガウシアンノイズ (ノイズ強度: 中 < 標準偏差  $\sigma = 0.1$  >) を付加したノイズデータを生成する。
- (4) (1) で複製した異常データにガウシアンノイズ (ノイズ強度: 強 < 標準偏差  $\sigma = 1$  >) を付加したノイズデータを生成する。
- (5) 機械学習装置 100 の  $\alpha$  ブレンディングにより、異常データと、異常データに類似の正常データとを合成した疑似異常データを生成する。

【0118】

図 15 は、異なる生成方法によって生成したデータの一例を示す図である。データ D1-1 は、正常データである。データ D1-2 は、異常データである。データ D (1) は、生成方法 (1) により生成したデータである。データ D (2) は、生成方法 (2) により生成したデータである。データ D (3) は、生成方法 (3) により生成したデータである。データ D (4) は、生成方法 (4) により生成したデータである。データ D (5) は、生成方法 (5) により生成したデータである。

【0119】

図 16 は、検証結果を示す図 (2) である。図 16 の検証結果は、生成方法 (1) ~ (5) を用いた場合において、全バッチのカメラ ID 別の平均 AUC スコアを示す。図 16 に示すように、生成方法 (5) は、他の生成方法 (1) ~ (4) と比較して、AUC スコアが、最高値または次点の性能維持を達成している。

10

20

30

40

50

## 【0120】

なお、本実施例では一例として、データが異常データまたは正常データに分類され、異常データの数、正常データの数と比較して少ない場合について説明したが、正常データの数、異常データの数と比較して少ない場合も同様に適用可能である。また、本実施例では、データが、異常データまたは正常データに分類される場合について説明したが、これに限定されるものではなく、他のクラスに分類されてもよい。

## 【0121】

次に、上記実施例に示した機械学習装置100と同様の機能を実現するコンピュータのハードウェア構成の一例について説明する。図17は、実施例の機械学習装置と同様の機能を実現するコンピュータのハードウェア構成の一例を示す図である。

10

## 【0122】

図17に示すように、コンピュータ200は、各種演算処理を実行するCPU201と、ユーザからのデータの入力を受け付ける入力装置202と、ディスプレイ203とを有する。また、コンピュータ200は、有線または無線ネットワークを介して、外部装置等との間でデータの授受を行う通信装置204と、インタフェース装置205とを有する。また、コンピュータ200は、各種情報を一時記憶するRAM206と、ハードディスク装置207とを有する。そして、各装置201～207は、バス208に接続される。

## 【0123】

ハードディスク装置207は、取得プログラム207a、機械学習プログラム207b、事前処理プログラム207c、推論プログラム207d、生成プログラム207e、ラベル付与プログラム207f、劣化検出プログラム207gを有する。また、CPU201は、各プログラム207a～207gを読み出してRAM206に展開する。

20

## 【0124】

取得プログラム207aは、取得プロセス206aとして機能する。機械学習プログラム207bは、機械学習プロセス206bとして機能する。事前処理プログラム207cは、事前処理プロセス206cとして機能する。推論プログラム207dは、推論プロセス206dとして機能する。生成プログラム207eは、生成プロセス206eとして機能する。ラベル付与プログラム207fは、ラベル付与プロセス206fとして機能する。劣化検出プログラム207gは、劣化検出プロセス206gとして機能する。

## 【0125】

取得プロセス206aの処理は、取得部151の処理に対応する。機械学習プロセス206bの処理は、機械学習部152の処理に対応する。事前処理プロセス206cの処理は、事前処理部153の処理に対応する。推論プロセス206dの処理は、推論部154の処理に対応する。生成プロセス206eの処理は、生成部155の処理に対応する。ラベル付与プロセス206fの処理は、ラベル付与部156の処理に対応する。劣化検出プロセス206gの処理は、劣化検出部157の処理に対応する。

30

## 【0126】

なお、各プログラム207a～207gについては、必ずしも最初からハードディスク装置207に記憶させておかなくても良い。例えば、コンピュータ200に挿入されるフレキシブルディスク(FD)、CD-ROM、DVD、光磁気ディスク、ICカードなどの「可搬用の物理媒体」に各プログラムを記憶させておく。そして、コンピュータ200が各プログラム207a～207gを読み出して実行するようにしてもよい。

40

## 【符号の説明】

## 【0127】

- 100 機械学習装置
- 110 通信部
- 120 入力部
- 130 表示部
- 140 記憶部
- 141 訓練データ

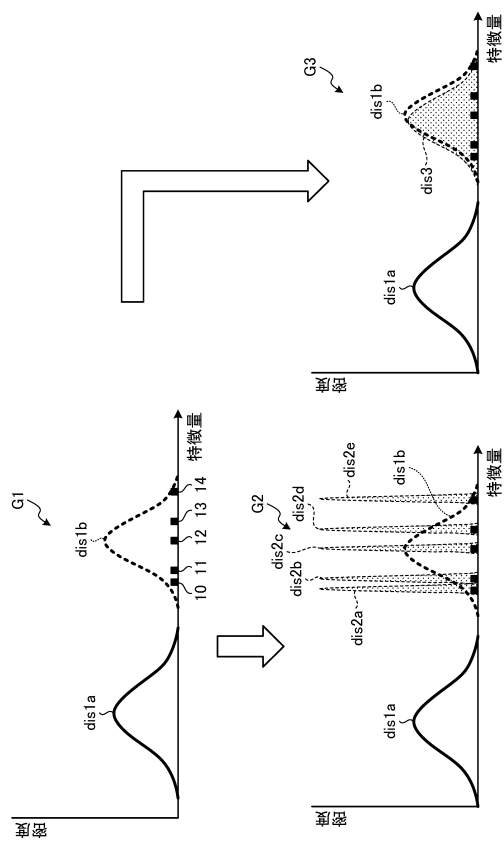
50



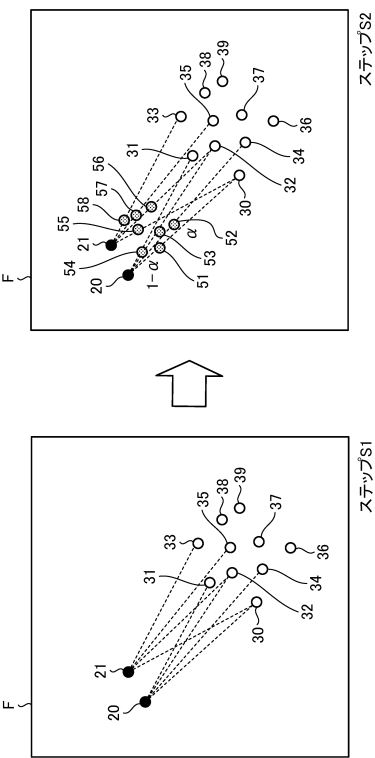
- 1 4 2 機械学習モデル
- 1 4 3 運用データ
- 1 4 4 再訓練データ
- 1 4 5 クラスタ関連データ
- 1 5 0 制御部
- 1 5 1 取得部
- 1 5 2 機械学習部
- 1 5 3 事前処理部
- 1 5 4 推論部
- 1 5 5 生成部
- 1 5 6 ラベル付与部
- 1 5 7 劣化検出部

【図面】

【図 1】



【図 2】



10

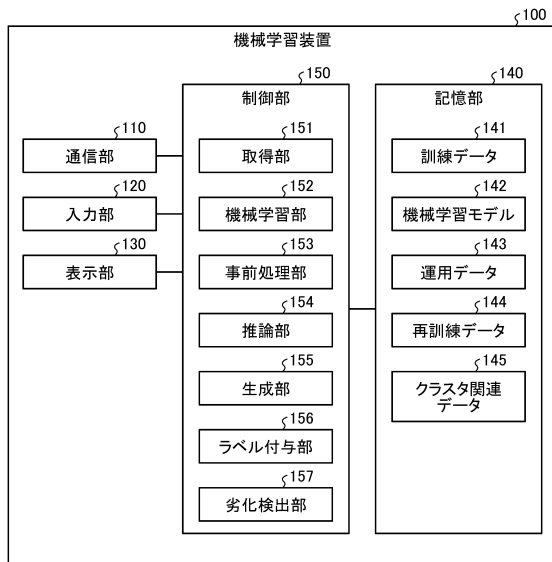
20

30

40

50

【圖 3】



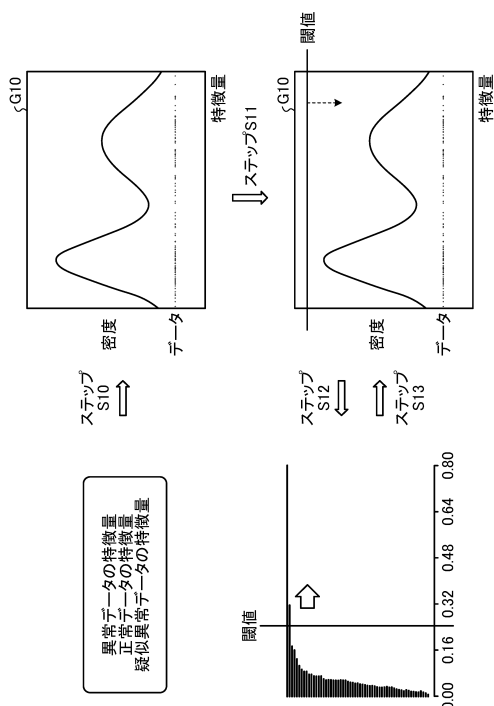
【図 4】

| §141 |             |       |
|------|-------------|-------|
| 項番   | データ(画像データ)  | 正解ラベル |
| 1    | 項番「1」の画像データ | 正常    |
| 2    | 項番「2」の画像データ | 正常    |
| 3    | 項番「3」の画像データ | 異常    |
| 4    | 項番「4」の画像データ | 正常    |
| ...  | ...         | ...   |

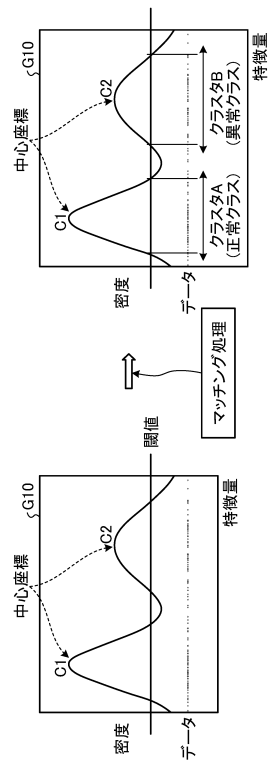
10

20

【図 5】



【図 6】

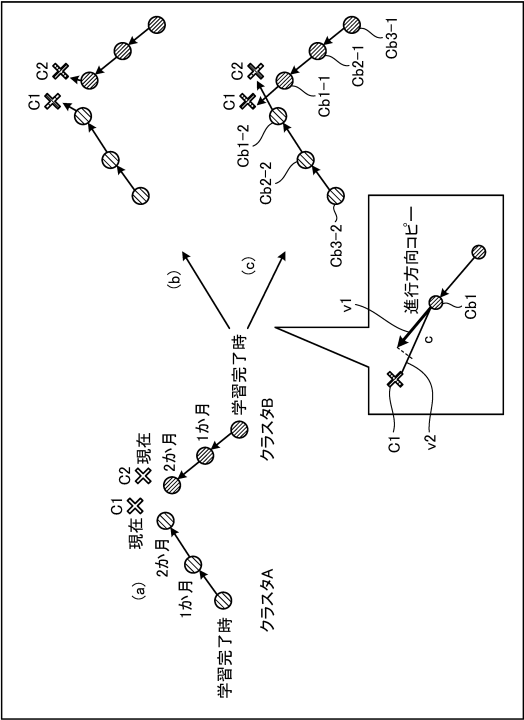


30

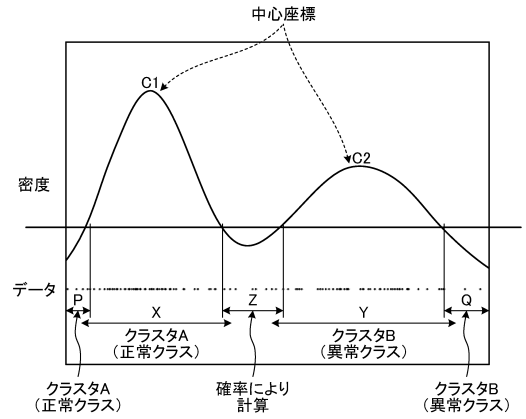
40

50

【図 7】



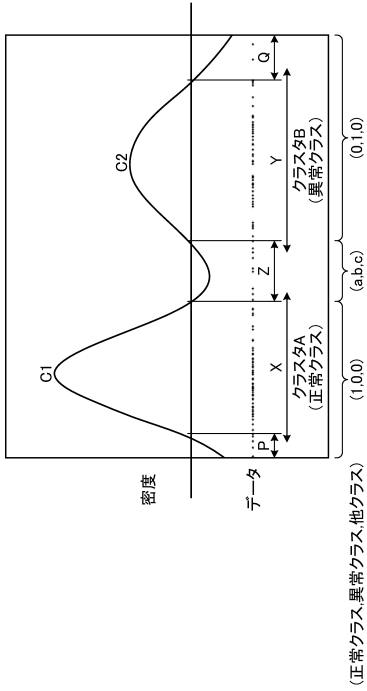
【図 8】



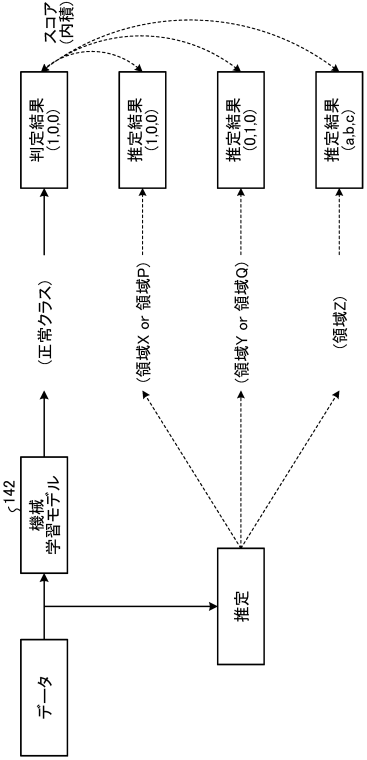
10

20

【図 9】



【図 10】

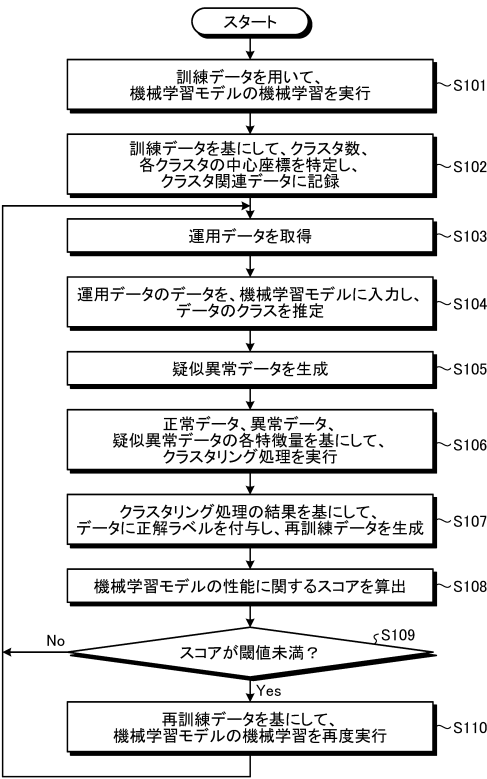


30

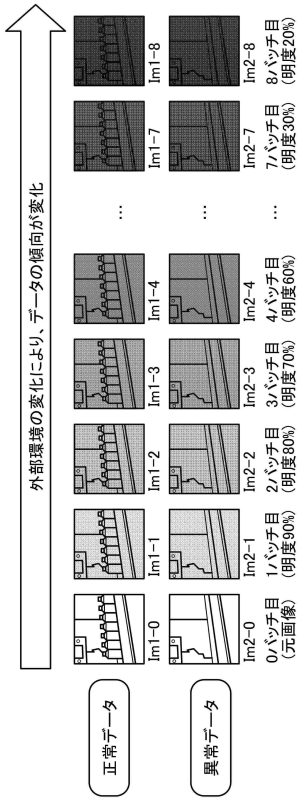
40

50

【図 1 1】



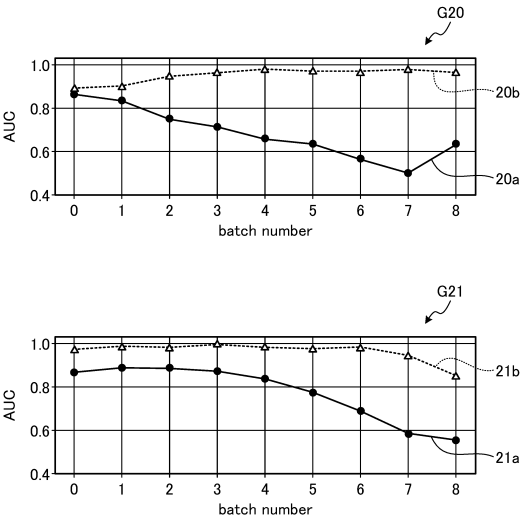
【図 1 2】



【図 1 3】

| カメラID | ベースライン<br>(再学習無し) | 提案手法 |
|-------|-------------------|------|
| 1     | 0.91              | 0.99 |
| 2     | 0.84              | 1.00 |
| 3     | 0.63              | 0.97 |
| 4     | 0.78              | 0.98 |
| 5     | 0.84              | 1.00 |
| 6     | 0.55              | 0.85 |
| 7     | 0.75              | 0.99 |

【図 1 4】



10

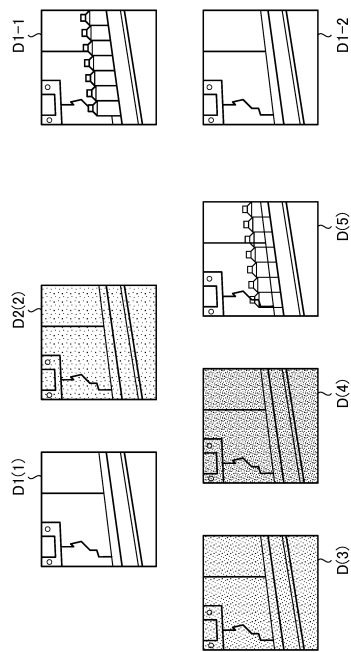
20

30

40

50

【図 1 5】



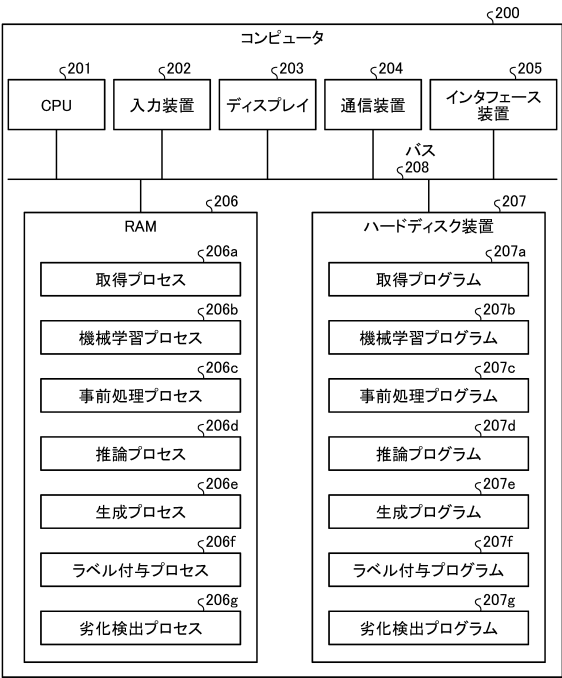
【図 1 6】

| カメラID | 生成方法(1) | 生成方法(2) | 生成方法(3) | 生成方法(4) | 生成方法(5) |
|-------|---------|---------|---------|---------|---------|
| 1     | 0.991   | 0.995   | 0.991   | 0.991   | 0.995   |
| 2     | -       | 0.992   | 0.906   | 0.953   | 0.992   |
| 3     | 0.947   | 0.912   | 0.767   | 0.878   | 0.965   |
| 4     | 0.999   | 0.997   | 0.995   | 0.918   | 0.997   |
| 5     | 1.000   | 1.000   | 1.000   | 1.000   | 1.000   |
| 6     | 0.873   | 0.982   | 0.702   | 0.571   | 0.875   |
| 7     | 0.977   | 0.992   | 0.818   | 0.908   | 0.994   |

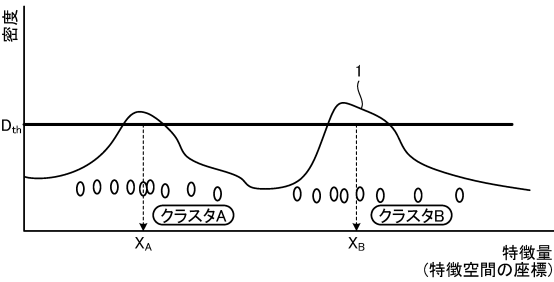
10

20

【図 1 7】



【図 1 8】

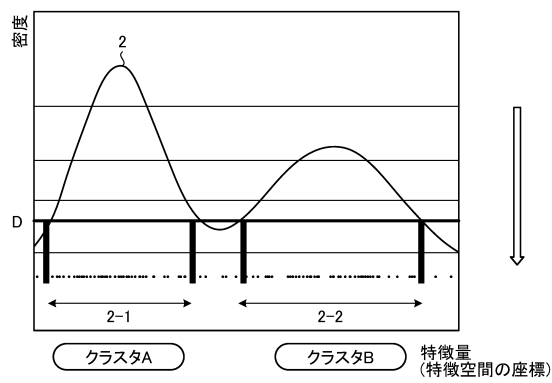


30

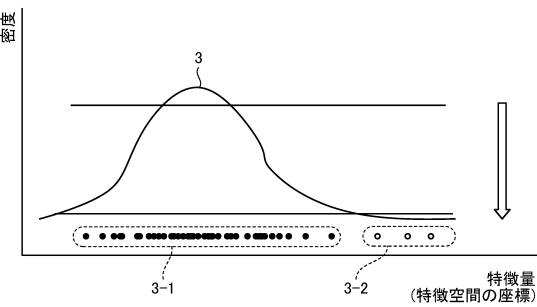
40

50

【図 1 9】



【図 2 0】



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2021/079442 (WO, A1)  
特開2021-057042 (JP, A)  
国際公開第2021/059509 (WO, A1)  
国際公開第2018/134952 (WO, A1)  
MORE, Ajinkya, "Survey of resampling techniques for improving classification performance in unbalanced datasets", arXiv.org [online], arXiv:1608.06048v1, Cornell University, 2016年, [検索日 2021.12.02], インターネット: <URL: <https://arxiv.org/pdf/1608.06048v1>>  
(58)調査した分野 (Int.Cl., DB名)  
G06N 3/00-99/00  
G06F 18/00-18/40