

(19)日本国特許庁(JP)

(12)特 許 公 報(B1)

(11)特許番号
特許第7698855号
(P7698855)

(45)発行日 令和7年6月26日(2025. 6. 26)

(24)登録日 令和7年6月18日(2025. 6. 18)

(51)Int. Cl.

G 0 6 T 7/00 (2017. 01)

F I

G 0 6 T 7/00 3 5 0 C

請求項の数 6 (全 32 頁)

(21)出願番号	特願2025-14513(P2025-14513)	(73)特許権者	517169746
(22)出願日	令和7年1月31日(2025. 1. 31)		C T W株式会社
審査請求日	令和7年2月17日(2025. 2. 17)		東京都港区六本木一丁目9番10号
早期審査対象出願		(74)代理人	100215027
			弁理士 留場 恒光
		(72)発明者	廖 威
			東京都港区六本木一丁目9番10号 C T W株式会社内
		(72)発明者	姜 曉嵐
			東京都港区六本木一丁目9番10号 C T W株式会社内
		(72)発明者	楊 銘
			東京都港区六本木一丁目9番10号 C T W株式会社内
		最終頁に続く	

(54)【発明の名称】デフォルメキャラクタ生成プログラムおよびデフォルメキャラクタ生成システム

(57)【特許請求の範囲】

【請求項1】

コンピュータを、(1)元キャラクタ画像分析手段、(2)参照画像分析手段、および(3)デフォルメキャラクタ取得手段として機能させ、

前記(1)元キャラクタ画像分析手段は、
キャラクタの画像である元キャラクタ画像を取得する元キャラクタ画像取得手段、
前記元キャラクタ画像から、複数の特徴量を取得する複数特徴量取得手段、
前記複数の特徴量を、ポーズの特徴量とそれ以外の特徴量に分ける特徴量分離手段、
前記元キャラクタ画像から、当該元キャラクタ画像について説明している文字情報を取得するキャラクタプロンプト取得手段、および、

前記複数の特徴量のうち、前記ポーズの特徴量以外の特徴量のうち少なくとも一部を学習データとして学習している拡散モデルを取得するキャラクタモデル取得手段、を備え、

前記(2)参照画像分析手段は、
デフォルメキャラクタの参照画像を取得する参照画像取得手段、および、
少なくとも前記デフォルメキャラクタの一部を示す線画情報に基づくデータを学習データとして学習している制御用モデルを取得するコントロールモデル取得手段、を備え、

前記(3)デフォルメキャラクタ取得手段は、前記制御用モデルにより制御されている前記拡散モデルと前記文字情報とから、前記元キャラクタ画像をデフォルメ化しているデフォルメキャラクタ画像を取得することを特徴とする、デフォルメキャラクタ生成プログラム。

【請求項 2】

前記少なくともデフォルメキャラクタの一部を示す線画情報に基づくデータを学習データとして学習している制御用モデルが、

少なくともデフォルメキャラクタの一部を示す線画情報およびポーズ情報に基づくデータを学習データとして学習している制御用モデルであることを特徴とする、

請求項 1 に記載のデフォルメキャラクタ生成プログラム。

【請求項 3】

さらに、前記拡散モデルが備えるニューラルネットワークの少なくとも一部の層において、重みを変更して得る複数の出力画像の差分を取得して評価し、出力画像に影響の小さい層の重みの削減または削除を行う最適化手段を備えることを特徴とする、請求項 1 に記載のデフォルメキャラクタ生成プログラム。

10

【請求項 4】

前記キャラクタの画像が、キャラクタの顔画像であり、

前記デフォルメキャラクタ画像が、デフォルメキャラクタ顔画像であることを特徴とする、請求項 1 に記載のデフォルメキャラクタ生成プログラム。

【請求項 5】

(1) 元キャラクタ画像分析部と、(2) 参照画像分析部と、および(3) デフォルメキャラクタ取得部と、を備え、

前記(1) 元キャラクタ画像分析部は、

キャラクタの画像である元キャラクタ画像を取得する元キャラクタ画像取得部と、

20

前記元キャラクタ画像から、複数の特徴量を取得する複数特徴量取得部と、

前記複数の特徴量を、ポーズの特徴量とそれ以外の特徴量に分ける特徴量分離部と、

前記元キャラクタ画像から、当該元キャラクタ画像について説明している文字情報を取得するキャラクタプロンプト取得部と、

前記複数の特徴量のうち、前記ポーズの特徴量以外の特徴量のうち少なくとも一部を学習データとして学習している拡散モデルを取得するキャラクタモデル取得部と、を備え、

前記(2) 参照画像分析部は、

デフォルメキャラクタの参照画像を取得する参照画像取得部と、

少なくとも前記デフォルメキャラクタの一部を示す線画情報に基づくデータを学習データとして学習している制御用モデルを取得するコントロールモデル取得部と、を備え、

30

前記(3) デフォルメキャラクタ取得部は、前記制御用モデルにより制御されている前記拡散モデルと前記文字情報とから、前記元キャラクタ画像をデフォルメ化しているデフォルメキャラクタ画像を取得することを特徴とする、デフォルメキャラクタ生成システム。

【請求項 6】

(1) 元キャラクタ画像分析ステップ、(2) 参照画像分析ステップ、および(3) デフォルメキャラクタ取得ステップ、から構成されるデフォルメキャラクタ生成方法であって、

前記(1) 元キャラクタ画像分析ステップは、

プロセッサがキャラクタの画像である元キャラクタ画像を取得する元キャラクタ画像取得ステップ、

40

プロセッサが前記元キャラクタ画像から、複数の特徴量を取得する複数特徴量取得ステップ、

プロセッサが前記複数の特徴量を、ポーズの特徴量とそれ以外の特徴量に分ける特徴量分離ステップ、

プロセッサが前記元キャラクタ画像から、当該元キャラクタ画像について説明している文字情報を取得するキャラクタプロンプト取得ステップ、および、

プロセッサが、前記複数の特徴量のうち、前記ポーズの特徴量以外の特徴量のうち少なくとも一部を学習データとして学習している拡散モデルを取得するキャラクタモデル取得ステップ、を備え、

50

前記（２）参照画像分析ステップは、

プロセッサがデフォルメキャラクタの参照画像を取得する参照画像取得ステップ、および、

プロセッサが、少なくとも前記デフォルメキャラクタの一部を示す線画情報に基づくデータを学習データとして学習している制御用モデルを取得するコントロールモデル取得ステップ、を備え、

前記（３）デフォルメキャラクタ取得ステップは、プロセッサが、前記制御用モデルにより制御されている前記拡散モデルと前記文字情報とから、前記元キャラクタ画像をデフォルメ化しているデフォルメキャラクタ画像を取得することを特徴とする、デフォルメキャラクタ生成方法。

10

【発明の詳細な説明】

【技術分野】

【０００１】

本発明は、デフォルメキャラクタ生成プログラムおよびデフォルメキャラクタ生成システムに関するものである。

【背景技術】

【０００２】

ゲームやアニメーション、教育、報道など様々な分野のコンテンツにおいて、例えば２頭身キャラクタなどのいわゆるデフォルメ化されたキャラクタが使用されている。

これらは視聴者に対して可愛らしい印象を与え、コンテンツに対し親しみを覚えやすくさせる効果などを有する。

20

【０００３】

オリジナルのキャラクタとして２頭身キャラクタが作成されることもあるが、元々８頭身などで描かれているキャラクタをデフォルメ化して２頭身キャラクタにすることなども行われる。

【０００４】

しかしながら、このようなデフォルメキャラクタを作成する場合、表情や服飾などの特徴的部分を残しつつ、不要な部分は削除するといった工夫が必要であり多くの労力が生じる。

【０００５】

30

例えば特許文献１には、キャラクタ画像を二頭身や三頭身のデフォルメ化したキャラクタに変換可能な画像変換装置が開示されている。

【先行技術文献】

【特許文献】

【０００６】

【特許文献１】特開２０２３－１５７３３４号公報

【非特許文献】

【０００７】

【非特許文献１】Miao Hua, Jiawei Liu, Fei Ding, Wei Liu, Jie Wu, Qian Heomiita, arXiv, “DreamTuner: Single Image is Enough for Subject-Driven Generation”, [online]、[2025年1月10日検索]、インターネット<URL: <https://arxiv.org/html/2312.13691v1>>

40

【０００８】

近年、例えばオリジナルとなる人物イラストなどを入力すると、表情や服装などが異なる別のイラストを生成することができるようになるなど、生成ＡＩなどを道具として活用することでイラストレーターなどの労力を軽減できるようになっている。

【０００９】

非特許文献１には、キャラクタ画像やポーズを入力として、好みのキャラクタ画像を出力させる技術について開示されている。

【００１０】

50

しかしながらその一方で、例えば8頭身のキャラクタを入力として、2頭身のキャラクタを生成A Iの出力とするような場合、いくつかの課題がある。

例えば、デフォルメ化に際しては、残す部分（強調する部分）と残さない部分（強調しない部分）を取捨選択しないといけない点である。

【0011】

キャラクタの身体的特徴や服装などは多種多様である。例えば、髪型、目の色、服飾のデザイン、持ち物、装飾品のそれぞれに特徴のあるキャラクタが存在する。

デフォルメ化に際し、どの部分を残し、どの部分を削除するかはキャラクタにより異なる。

髪型に特徴のあるキャラクタをデフォルメ化したときにその髪型が変わっていたとすると、デフォルメ化としては適切でないと言える。

【0012】

しかしながら、望み通りのデフォルメキャラクタを画像生成A Iの出力として得ようとすると、プロンプトの工夫に頼らざるを得ないなど、多大な労力が生じる問題がある。

【発明の概要】

【発明が解決しようとする課題】

【0013】

解決しようとする問題点は、通常のキャラクタからデフォルメ化キャラクタを生成するにあたり、必要な特徴を残すデフォルメ化が困難な点である。

【課題を解決するための手段】

【0014】

本発明は、デフォルメ化キャラクタの生成に際して必要な特徴を残すために、入力画像（例えば8頭身キャラクタ画像）から複数の特徴量を取得し、そのなかで、ポーズの特徴量（とそれ以外などの特徴量と）を分離する（デカップリングする）こと（特徴量分離）を最も主要な特徴とする。

特徴量を分離することで、各特徴量が他の特徴量に依存せずに処理されるため、モデルの柔軟性と精度が著しく向上する。

【0015】

非特許文献1には、デフォルメキャラクタに変換することや、入力しているキャラクタ画像から得られる特徴量（feature）についてポーズの特徴量とそれ以外の特徴量とに分離すること（Feature Decoupling・後述）については記載されていない。

【0016】

本発明は、上記課題を鑑みてなされたものであり、例えば以下の手段を採用している。

すなわち、コンピュータを、（1）元キャラクタ画像分析手段、（2）参照画像分析手段、および（3）デフォルメキャラクタ取得手段として機能させ、

前記（1）元キャラクタ画像分析手段は、

キャラクタの画像である元キャラクタ画像を取得する元キャラクタ画像取得手段、

前記元キャラクタ画像から、複数の特徴量を取得する複数特徴量取得手段、

前記元キャラクタ画像から、当該元キャラクタ画像について説明している文字情報を取得するキャラクタプロンプト取得手段、および、

前記複数の特徴量のうち、ポーズの特徴量以外の特徴量のうち少なくとも一部を学習データとして学習している機械学習モデルを取得するキャラクタモデル取得手段、を備え、

前記（2）参照画像分析手段は、

デフォルメキャラクタの参照画像を取得する参照画像取得手段、および、

少なくとも前記デフォルメキャラクタの一部を示す線画情報に基づくデータを学習データとして学習している制御用モデルを取得するコントロールモデル取得手段、を備え、

前記（3）デフォルメキャラクタ取得手段は、前記制御用モデルにより制御されている前記機械学習モデルと前記文字情報とから、前記元キャラクタ画像をデフォルメ化しているデフォルメキャラクタ画像を取得することを特徴とする、デフォルメキャラクタ生成プログラムを提供する。

10

20

30

40

50

【発明の効果】

【0017】

本発明のデフォルメキャラクタ生成プログラムは、ポーズの特徴量と、ポーズの特徴量以外の特徴量とを分離する（デカップリングする）ため、元キャラクタ画像の主要な特徴量は残しつつ、デフォルメキャラクタを出力できるという利点がある。

【図面の簡単な説明】

【0018】

【図1】デフォルメキャラクタ生成プログラムP1による処理の概要を示す図である。

【図2】基本設定画面を示す図である。

【図3】元キャラクタ画像確認画面を示す図である。

10

【図4】デフォルメキャラクタ設定欄を示す図（拡大図）である。

【図5】生成画像表示画面を示す図である。

【図6】変換画像取得処理を示すフローチャートである。

【図7】複数の領域を認識するイメージを表す図である。

【図8】最適化処理を示すフローチャートである。

【図9】重みの削減等を説明するための図である。

【図10】最適化処理前後のデフォルメキャラクタ画像を示す図である。

【図11】デフォルメキャラクタ生成処理を（図1とは別の形で）示す図である。

【図12】デフォルメキャラクタ生成システム1の概要を示すネットワーク構成図である。

20

【図13】サーバ10のハードウェア構成図である。

【図14】第二の実施形態の概要を示す図である。

【図15】デフォルメキャラクタ生成プログラムP1によるデフォルメ化の応用例を示す図である。

【発明を実施するための形態】

【0019】

本発明の実施形態を図面に基づいて説明する。以下の各実施形態では、同一又は対応する部分については同一の符号を付して説明を適宜省略する場合がある。

また、以下に用いる図面は本実施形態を説明するために用いるものであり、実際の装置の構成やユーザーインターフェース（UI）、データ構成などとは異なる場合がある。

30

【0020】

（実施形態の概要）

本実施形態の概要について、図1を用いて説明する。

図1は、本実施形態のデフォルメキャラクタ生成プログラムP1による処理の概要を示す図である。

【0021】

デフォルメキャラクタ生成プログラムP1は、高頭身キャラクタ（図1では左上の5～6頭身キャラクタ）の画像を入力として、低頭身キャラクタ（図1では右下の2～3頭身キャラクタ）の画像を出力するプログラムである。

【0022】

40

デフォルメキャラクタ生成プログラムP1は、（1）元キャラクタ画像を取得して各種情報（キャラクタプロンプト・キャラクタモデルなど）を取得する元キャラクタ画像分析手段（図1右上）、（2）参照画像を取得してコントロールモデルを取得する参照画像分析手段（図1左下）、および（3）（1）および（2）で得られている情報をもとにデフォルメキャラクタを出力するデフォルメキャラクタ取得手段（図1右下）を備える。

【0023】

また、デフォルメキャラクタ生成プログラムP1による処理を行うシステムを以下「デフォルメキャラクタ生成システム1」とする。

【0024】

（実施形態の詳細）

50

以下、本実施形態に係るデフォルメキャラクタ生成システム１について、詳細を説明する。デフォルメキャラクタ生成システム１は、デフォルメキャラクタ生成プログラムＰ１を備えるコンピュータ（サーバ１０）を含み、ユーザに対し、入力に係るキャラクタ画像をデフォルメ化して出力するシステムをオンラインで提供する。

すなわちデフォルメキャラクタ生成システム１は、デフォルメキャラクタ生成プログラムＰ１による情報処理が、ハードウェア資源を用いて具体的実現されるものである。

以下、デフォルメキャラクタ生成システム１を構成する、１．ユーザーインターフェース、２．プログラム処理、および３．ハードウェア構成、について順に説明する。

【００２５】

（用語の定義）

10

ここで、いくつか言葉の定義を行う。

「キャラクタ」とは、目視可能な外観を有するものであり、生物・非生物や実在・非実在を問わない。ここでは特に、全身のうち頭部とそれ以外を分けられるものを意味する。例えば創作アニメの人物キャラクタである（生物・非実在）。ここで非生物キャラクタとは、例えばロボットキャラクタなどを意味する。さらに、実在キャラクタとは、例えば芸能人をアニメ調に描いたものなどである。なお、例えばシーツを頭からすっぽり被った人物のように、首の場所が隠れるなどし、頭部が明確でない場合であっても、およそその位置が推認できるものであればよい。また例えば、顔（目や口）のある部位を頭部としてもよい。

「キャラクタ画像」とは、キャラクタの画像である。イラスト画像か写真画像かを問わない。

20

「ポーズ」とは、キャラクタがとる姿勢を意味する。

「デフォルメ」とは、対象の形態を意図的に変えることをいう。例えば、キャラクタの頭身比率を変化させ、高頭身キャラクタから低頭身キャラクタにすることを指す。

「高頭身キャラクタ」「低頭身キャラクタ」はそれぞれ一般的に、高い頭身（例えば６～９頭身）のキャラクタおよび低い頭身（例えば２～５頭身）のキャラクタをいう。ただし、ここでいう高い・低いとは絶対的なものではなく相対的なものである。デフォルメと言う場合、元キャラクタ画像（以下「入力画像」と称する場合がある。）が高い頭身、元キャラクタ画像をデフォルメ化している画像（以下「出力画像」と称する場合がある。）が低い頭身となる。例えば、４頭身キャラクタを２頭身キャラクタに変化させることもデフォルメと称する。なお一般に、低頭身キャラクタはちびキャラクタ（chibi character）などと呼ばれる。

30

「特徴量分離」（Feature Decoupling）とは、機械学習モデル、特に画像生成モデルにおいて、異なる特徴量（例えばポーズ（姿勢）、衣服、表情に関する特徴量など。）を独立して扱う手法である。特徴量は変数などともいう。

「ユーザ」は、デフォルメキャラクタ生成システム１を使用する者をいう。端末２０を操作し、個人と法人とを問わない。デフォルメキャラクタ生成システム１において、ユーザは高頭身キャラクタの画像を入力し、変換後の低頭身キャラクタを出力として取得する。

「取得」は、プロセッサがデータまたはプログラム等を入力部または外部端末などから取得することのほか、プロセッサがデータまたはプログラム等を作成・更新等することを含む。下記実施形態のキャラクタモデルなどを参照されたい。

40

【００２６】

以下に図面等で使用されている用語を説明する。

ただし以下は発明の理解を助けるためのものであり、本明細書の内容を限定するものではない。

「Stable Diffusion」は、テキスト等を画像に変換する機械学習モデルであり、いわゆる画像生成ＡＩの１つである。Stable Diffusionは拡散モデル（Diffusion Model）（特に、潜在拡散モデル（Latent Diffusion Model））を用いていることを特徴とする。また、Stable Diffusionは条件付けを行うテキストを入力するためのテキストエンコーダ（Te

50

xt Encoder)を有する。テキストエンコーダはテキストをベクトルに変換するものであり、例えばCLIPという手法(大量の画像とテキストの組み合わせを学習し、画像とテキストの類似度を算出)で学習済みのTransformer(機械学習モデル)である。Stable Diffusionにおける処理は、VAE(次項)により画像情報から潜在変数を得る処理、および、VAEにより潜在変数から画像情報を得る処理を含む。VAEを用いることで計算量を抑えていることが特徴である。拡散モデルの具体的態様については下記の実施形態において説明する。

「VAE(Variational Auto-Encoder)」は、変分オートエンコーダ(変分自己符号化器)のことである。VAEはエンコーダー・デコーダ構造を有する生成モデルである。ただしVAEの語は、このような構造を有するニューラルネットワーク、このようなネットワーク構造、機械学習モデル、またはこれらを含む技術を指すこともある。VAEは、入力データを平均と標準偏差(または分散など)で表現するように学習する(平均ベクトル、分散ベクトルなど)。例えば、エンコーダは入力データを(統計分布のある一点となる)潜在変数に変換し、デコーダは潜在変数(統計分布から(ランダムに)サンプリングした一点)を復元することで新しいデータを生成する。また、学習時において誤差逆伝播(Backpropagation)のための工夫(Reparametrization trick)がなされている。

「Unet」は、CNN(Convolutional Neural Network)をベースとするニューラルネットワークである。例えば畳み込み層を含むモジュールと、アテンション層を含むモジュールを含む。一般的には画像のセグメンテーションに使用されるが、ここでは拡散モデルの逆拡散過程、例えば潜在空間におけるノイズ除去に使用される。Unetは例えば、ノイズを入力として、除去すべきノイズを出力する。入力されているノイズから除去すべきノイズを除去することで、当該ノイズが除去されている画像が得られる。これを繰り返すことで、徐々に画像からノイズが除去されていく。

「ControlNet」は、Stable Diffusionの拡張機能であり、ユーザが指定したポーズや構図に基づいて(高品質な)画像を生成することを可能にする。ControlNetは、画像の構造的特徴(ポーズ、線画、エッジなど)を条件(Condition)として与え、それを反映して画像を生成するための追加的なニューラルネットワークである。Unetのニューラルネットワークブロックを保持しつつ、ControlNet上のニューラルネットワークでユーザから与えられた追加条件を付加してUnetに反映させる。つまり、条件(Condition)(例えば構図データ)などが入力となり、Unet上のブロックを出力先とする。詳細は割愛するが、ゼロ畳み込みやスキップコネクションといった技術が組み込まれている。

「LoRA(Low-Rank Adaptation)」は、(事前)学習済み(機械学習)モデルのファインチューニングに関する技術である。なお、低消費電力長距離無線通信技術のLoRa(登録商標)とは異なる。

【0027】

以下において、「○○」処理と記載している場合、コンピュータのプロセッサは、プログラム格納部に記憶されている「○○」プログラムに基づく処理を実行することを意味する。本段落において、「○○」の箇所には同じ語が入る。

すなわち、「○○」プログラムは、「○○」処理の実行により、コンピュータを「○○」手段として機能させるプログラムである。またこの際、当該プロセッサを備える制御部は、「○○」部(または「○○」装置)としても機能することを意味する。

この場合において、「○○」部は、「○○」プログラムに基づく「○○」処理を実行することを意味する。

また、「○○」処理の)処理手順を示す場合は「○○」ステップと記載する。

【0028】

例えば、デフォルメキャラクタ生成プログラムP1は、デフォルメキャラクタ生成処理の実行により、コンピュータをデフォルメキャラクタ生成手段として機能させるプログラムである。またこの際、プロセッサ122を備えるコンピュータの制御部12は、デフォルメキャラクタ生成部(またはデフォルメキャラクタ生成装置)として機能する。

【0029】

デフォルメキャラクタ生成システム1において、端末20などの各端末(コンピュータ

10

20

30

40

50

）はそれぞれプロセッサを備えるが、単にプロセッサという場合は、デフォルメキャラクタ生成プログラム P 1 により処理を行うプロセッサ、本実施形態ではサーバ 1 0 のプロセッサ 1 2 2、を指すものとする。

【 0 0 3 0 】

なお以下において、簡単のため、「サーバ 1 0 のプロセッサ 1 2 2 が、端末からのリクエストを受けて、当該端末のブラウザに表示するためのデータを返す」ことを、「プロセッサ 1 2 2 が端末のブラウザに表示する（させる）」または「プロセッサ 1 2 2 が表示する（させる）」などと記載する場合がある。

また同様に、「サーバ 1 0 のプロセッサ 1 2 2 が、記憶部 1 4 のデータ記憶部 1 4 b にデータを保存させる」ことを、「プロセッサが（データを）保存する（させる）」などと記載する場合がある。

【 0 0 3 1 】

（第一の実施形態）

【 0 0 3 2 】

1 . ユーザーインターフェース (U s e r I n t e r f a c e (U I))

まず、本実施形態のデフォルメキャラクタ生成システム 1 が端末 2 0 に表示させるインターフェースについて、図を用いて説明する。

以降で説明するインターフェースは、プロセッサ 1 2 2 が端末 2 0 のブラウザに表示させるものを簡略化したものである。

【 0 0 3 3 】

また、説明に必要な機能に関わるアイコン等のみ表示することとし、それ以外の公知のアイコンなどは省略する。例えば、直前に表示されていたページに戻るための戻るボタンなどは省略している。

【 0 0 3 4 】

図 2 は、基本設定画面を示す図である。

基本設定画面左側の画像アップロード欄 U I - 1 1 は、ユーザが入力画像である元キャラクタ画像をアップロードするための U I であるとともに、アップロードされた元キャラクタ画像を表示する欄である。

ユーザはドラッグアンドドロップで元キャラクタ画像をアップロードする。

なお、アップロードの方法はこれに限らない。画像が格納されているフォルダを選択し、画像ファイルを選択できるようにしてもよい（不図示）。

【 0 0 3 5 】

デフォルメキャラクタ設定欄 U I - 1 2 は、ユーザが出力するデフォルメキャラクタ画像に関する設定を行うための U I である。

デフォルメキャラクタ設定欄 U I - 1 2 は等身比率設定ボタン U I - 1 2 1、参照デフォルメキャラクタ画像 U I - 1 2 2、およびデフォルメキャラクタ生成開始ボタン U I - 1 2 3 を表示する。詳細は後述する。

【 0 0 3 6 】

図 3 は、元キャラクタ画像確認画面を示す図である。

ユーザが元キャラクタ画像を画像アップロード欄 U I - 1 1 にアップロードすると、プロセッサ 1 2 2 はその画像から後述するキャラクタプロンプト（文字情報）およびキャラクタモデル（機械学習モデル）を取得・作成し、元キャラクタ画像確認画面（ウィンドウ） U I - 1 3 を表示する。

なお図中において、キャラクタモデルの作成については「トレーニング」という用語を用いている。

【 0 0 3 7 】

なお本実施形態において、プロセッサ 1 2 2 はアップロード後にトレーニング（学習）を行うかについての確認画面を表示し、ユーザがそこで表示されている O K ボタンアイコン等を押下することによりパラメータ更新などによる学習を開始するが、ここでは省略している。

10

20

30

40

50

【0038】

図3に示すように、プロセッサ122は元キャラクタ画像確認画面（ウィンドウ）UI-13に画像情報UI-131（ファイル名、ユーザID、作成者、および作成日時など）、ユーザがアップロードしている元キャラクタ画像UI-132、OKボタンアイコンUI-133、再トレーニングアイコンUI-134、および削除ボタンアイコンUI-135を表示する。

【0039】

ユーザがOKボタンアイコンUI-133を押下すると、プロセッサ122は基本設定画面の画像アップロード欄UI-11に元キャラクタ画像を表示する（不図示）。

なおユーザが基本設定画面で当該元キャラクタ画像を選択（押下）すると、プロセッサ122は再度元キャラクタ画像確認画面（ウィンドウ）UI-13を表示する。

10

【0040】

元キャラクタ画像確認画面（ウィンドウ）UI-13でユーザが再トレーニングアイコンUI-134を押下すると、プロセッサ122は元キャラクタ画像を選択する画面を表示する（不図示）。ユーザがそこで元キャラクタ画像を選択して再トレーニングを実行させるボタンを押下すると、プロセッサ122はユーザの選択による当該元キャラクタ画像（入力画像）からキャラクタプロンプトおよびキャラクタモデルを（再度）作成する。

この際ユーザの任意により、または自動的に、パラメータが調整される。このほか、トレーニングの方法はユーザが調整することができる（不図示）。

【0041】

ユーザが削除ボタンアイコンUI-135を押下すると、プロセッサ122はアップロードされている元キャラクタ画像を削除する。ユーザは新しい元キャラクタ画像をアップロードすることができる。

20

【0042】

図4は、デフォルメキャラクタ設定欄（UI-12）を示す図（拡大図）である。

【0043】

本実施形態において、参照デフォルメキャラクタ画像（参照されるデフォルメキャラクタの画像）はあらかじめ記憶部14（データ格納部14b）に多数格納されており、ユーザはデフォルメキャラクタの向きやポーズなどを自由に選択することができる。

【0044】

また、ユーザは自分の好みのデフォルメキャラクタをアップロードすることもできる。

30

例えば本実施形態において、デフォルメキャラクタ設定欄UI-12に画像をドラッグアンドドロップすることで好みのキャラクタを参照画像とすることもできる。

【0045】

等身比率設定ボタンUI-121は、出力画像の頭身比率を設定するボタンアイコンである。例えば、1:4のボタンを選択すると、プロセッサ122はデフォルメキャラクタ設定欄UI-12に頭の高さ:全身の高さが1:4となる画像を出力させる（図2参照。また、図2・図4における参照デフォルメキャラクタ画像UI-122の影の部分参照）。

【0046】

図4に示すように、等身比率設定ボタンUI-121のA11ボタンは、示されている（図4であれば1:2.5~1:5）比率の画像について一通り表示するものである。

40

【0047】

ユーザが参照デフォルメキャラクタ画像UI-122を選択して押下すると、プロセッサ122はチェック（レ点）を表示してその画像が選択されたことを表示する。

ユーザは複数の参照デフォルメキャラクタ画像を選択することができ、その場合プロセッサ122はその選択に応じて複数のデフォルメキャラクタを生成する。

つまりプロセッサ122は、一度に複数のデフォルメキャラクタ画像を出力することができる。

【0048】

50

参照デフォルメキャラクター画像UI - 122は、出力する予定のデフォルメキャラクターのポーズと頭身比率の目安を表示するものである。この画像は多くのストックのなかからプロセッサ122が自動で選択する。

その一方でプロセッサ122はユーザの選択により、（立ち絵や座り絵といった）ポーズを選択するためのポーズ表示ウィンドウを表示し（不図示）、ユーザは当該ポーズ表示ウィンドウに一覧表示されるポーズ群から好みのポーズを選択することもできる。

【0049】

つまりユーザは、等身比率設定ボタンUI - 121を押下することで、または参照画像をアップロードすることで、簡単に頭身比率を入力することができる。

言い換えると、ユーザはGUIによる簡単な操作で頭身比率を入力することができる。

【0050】

ユーザがデフォルメキャラクター生成開始ボタンUI - 123を押下すると、プロセッサ122はデフォルメキャラクター生成に関する処理を開始する。処理については後述する。

【0051】

図5は、生成画像表示画面を示す図である。

プロセッサ122がデフォルメキャラクターの画像（出力画像）を取得すると、生成画像表示画面の出力画像表示欄UI - 14に出力画像UI - 141を表示する。またプロセッサはこれ以外にも、出力画像加工関連ボタンUI - 142、出力画像情報UI - 143、および出力画像保存関連ボタンUI - 144を表示する。

【0052】

出力画像加工関連ボタンUI - 142は、出力されている画像に対してさらに加工を行うか、画像を再生成させるためのボタンアイコンである。

【0053】

図5に示す拡大率ボタンは表示されている拡大率を変化させるためのボタンアイコンである。

色ボタンはキャラクターの肌の色や髪の色など、色の加工を行うためのボタンアイコンである。

パーツ分けボタンはデフォルメキャラクターのパーツ（例えば髪型、服装など）を他のパーツと分離して表示させるためのボタンアイコンである。ユーザはそのパーツについて別個独立に加工等を行うことができる。

再生成ボタンはデフォルメキャラクター画像を再生成するためのボタンアイコンである。機械学習モデルにおいて異なるパラメータが設定され（拡散モデルにおいてランダムで潜在変数がサンプリングされ）、パラメータに応じたデフォルメキャラクター画像（通常異なるデフォルメキャラクター画像）を生成することができる。つまり、ユーザはボタン一つで簡単に異なる画像を取得することができるという利点がある。

【0054】

出力画像情報UI - 143は画像に関する情報を表示する欄である。

スタイルは、参照画像のファイル名である。

バリエーション強度は、バリエーションを加えるためのパラメータである。バリエーション強度が大きいほど、元キャラクター画像からの変化が大きい画像が出力画像として出力される。

本実施形態において、バリエーション強度はユーザによる設定が可能である（UIは不図示）。

【0055】

出力画像保存関連ボタンUI - 144は、出力している画像の保存等に関する操作ボタンアイコンである。DLは出力している画像を端末（端末20）にダウンロードするものであり、削除は出力している画像を削除するものであり、保存ボタンはサーバ（サーバ10）上に出力している画像を（サーバ上などに）保存するものである。

【0056】

以上のような構成により、ユーザはGUI（Graphical User Interface）による簡単な

10

20

30

40

50

操作でデフォルメキャラクタ画像を生成することができる。

特に、ユーザは参考例としての参照デフォルメキャラクタ画像UI - 122を見ながら頭身比率を選択できるため、出力される画像が予想しやすく、希望するデフォルメキャラクタ画像をより確実に得ることができる。

またこの際、プロンプトを細かく調整するなどの入力を行わなくてもGUIで処理を実行できるという著しい効果がある。

【0057】

2. プログラム処理

< 2 - 1. デフォルメキャラクタ生成処理 >

本実施形態のデフォルメキャラクタ生成システム1において行われるプログラム処理について説明する。

【0058】

本実施形態において、プロセッサ122は、デフォルメキャラクタ生成プログラムP1に基づき、デフォルメキャラクタ生成処理を行う。

デフォルメキャラクタ生成プログラムP1は、少なくとも(1)元キャラクタ画像分析プログラムP11、(2)参照画像分析プログラムP12、(3)デフォルメキャラクタ取得プログラムP13、および(4)最適化プログラムP14を含み、プロセッサ122はこれらの各プログラムに基づいて、元キャラクタ画像分析処理、参照画像分析処理、デフォルメキャラクタ取得処理、および最適化処理をそれぞれ実行する。

また上記のほか、デフォルメキャラクタ生成処理は、ユーザに上述したUIをオンラインで提供するオンラインUI提供処理を含む。提供するUIについては説明済みであるためここでは省略する。

【0059】

(1)元キャラクタ画像分析処理は、元キャラクタ画像を取得して分析する処理であり、

(加工対象となる)キャラクタの画像である元キャラクタ画像を取得する元キャラクタ画像取得処理、

1枚の前記元キャラクタ画像から、複数の領域の画像情報を取得して複数の特徴量を取得する複数特徴量取得処理、

前記領域(の画像)および/または特徴量ごとにタグを付して管理するタグ管理処理、

前記複数の特徴量を、ポーズの特徴量とそれ以外の特徴量に分ける(Feature Decoupling(特徴量分離)する)特徴量分離処理、

前記元キャラクタ画像から、当該元キャラクタ画像について説明している文字情報を取得する、(特に、前記複数の特徴量のうち、ポーズの特徴量以外の特徴量を含むコンセプトの特徴量について説明する文字情報を取得する)キャラクタプロンプト取得処理、および、

前記複数の特徴量のうち、ポーズの特徴量以外の特徴量のうち少なくとも一部(ポーズの特徴量以外の特徴量を含むコンセプトの特徴量)を学習データとして学習している(学習する)機械学習モデル(拡散モデル)を取得するキャラクタモデル取得処理、を含む。

【0060】

ここで、上記文字情報を「キャラクタプロンプト」、上記機械学習モデル(拡散モデル)を「キャラクタモデル」と称する。

【0061】

(2)参照画像分析処理は、デフォルメキャラクタの参照画像を取得して分析する処理であり、

デフォルメキャラクタの参照画像を取得する参照画像取得処理、および、

前記(参照画像における)デフォルメキャラクタの(形状の)一部を示す線画情報を少なくとも1つの学習データとして学習している制御用(機械学習)モデルを取得するコントロールモデル取得処理、を含む。

【0062】

10

20

30

40

50

ここで、上記制御用（機械学習）モデルを「コントロールモデル」と称する。

なお、上述した機械学習モデル（キャラクタモデル）と区別するために制御用モデルと表記している。

例えばキャラクタモデルを第一の機械学習モデル、コントロールモデルを第二の機械学習モデルと称してもよい。

【 0 0 6 3 】

（ 3 ）デフォルメキャラクタ取得処理は、

前記制御用モデルにより制御されている前記機械学習モデル（拡散モデル）と前記文字情報とから、前記元キャラクタ画像をデフォルメ化しているデフォルメキャラクタ画像を取得する処理である。

【 0 0 6 4 】

より具体的には、前記制御用モデル（＝コントロールモデル）により制御されている前記機械学習モデル（＝キャラクタモデル）に対して前記文字情報（＝キャラクタプロンプト）を入力し、前記元キャラクタ画像をデフォルメ化しているデフォルメキャラクタ画像を出力として取得する処理である。

【 0 0 6 5 】

大きく分けると、画像（元キャラクタ画像・参照画像）を取得して分析等する処理（上記（ 1 ）および（ 2 ））と、画像生成によりデフォルメキャラクタ画像を得る処理（上記（ 3 ））である。

言い換えると、画像生成によりデフォルメ画像を得る処理（上記（ 3 ））は、（上記（ 1 ）および（ 2 ）の）から得られるデータ・モデル等からデフォルメキャラクタを実際に生成する処理である。

【 0 0 6 6 】

図 6 は、デフォルメキャラクタ生成処理を示すフローチャートである。

プロセッサ 1 2 2 は、デフォルメキャラクタを生成する処理の開始の指示（例えば上記デフォルメキャラクタ設定欄 UI - 1 2 のデフォルメキャラクタ生成開始ボタン UI - 1 2 3 の押下（図 4 参照））を受け付けることにより、デフォルメキャラクタ生成処理を開始する。

【 0 0 6 7 】

まず、（ 1 ）元キャラクタ画像分析処理について説明する。

プロセッサ 1 2 2 は、ユーザにより入力されている、デフォルメ化加工対象の元キャラクタ画像を取得する（ステップ 1 ・元キャラクタ画像取得処理）。

【 0 0 6 8 】

プロセッサ 1 2 2 は取得した画像から複数の特徴量を抽出し、取得する（ステップ 2 ・複数特徴量取得処理）。

本実施形態において、プロセッサ 1 2 2 は、 1 枚の元キャラクタ画像から、例えば頭部、上半身、または全身に係る部分など、複数の領域を認識してその画像情報を取得して複数の特徴量を取得する。

【 0 0 6 9 】

図 7 は複数の領域を認識するイメージを表す図である。

図 7 中の、スケールの異なる複数の矩形領域（四角形の線で囲まれた部分・バウンディングボックス）が複数の領域を示す。

【 0 0 7 0 】

本実施形態において、プロセッサ 1 2 2 は顔を含む画像を異なるスケールで取得（クロップ）している。

本実施形態において特に、顔を認識する顔認識（Face Detection）の処理を行っている。顔認識の方法は公知の方法が適宜用いられる。

なお、このように画像を（切り取って）取得する処理は一般的に「クロップ（crop）」または「クロッピング（cropping）」と呼ばれる。

【 0 0 7 1 】

10

20

30

40

50

図7に示す例の場合、複数の領域は、図7の矩形領域の面積が小さいものから順に、顔、頭部、上半身、頭頂より腰まで、膝より上（帽子を含む、以下同じ。）、ふくらはぎより上、くるぶしより上、踵より上、および爪先より上（全身）、に係る部分を含む。

【0072】

スケールの異なる複数の領域の画像を取得（クロップ）することは、特徴量を取得すべき部位（位置）を明確にし、部位（位置）に応じた複数の特徴量を得ることを可能にする。例えば、顔とその他の部位とに関するデータを取得することなどができる。また、少数（1枚）の画像からデータ量を増やすという利点がある。

【0073】

1枚の元キャラクタ画像から複数の領域を認識して画像を取得する方法として公知の物体検出技術を用いることができる。

【0074】

なお物体検出技術のうち、大まかな物体の位置を特定したあとにその物体のクラスを識別するものは2段階モデルと呼ばれ、例えばR-CNN（Region Based Convolutional Neural Networks）やFPN（Feature Pyramid Networks）が挙げられる。また、物体の位置の特定とクラスの識別を同時に行うものは1段階モデルと呼ばれ、例えばYOLO（You Only Look Once）やSSD（Single Shot multibox Detector）が挙げられる。このほか、カスケード分類器（Open CV）などを用いてもよい。

【0075】

「複数の領域」は、元キャラクタ画像における認識単位（切り出し（クロッピング）単位）を示す。プロセッサ122は1つの領域について1つのクロッピング画像を得る。

「複数の特徴量」は、元キャラクタ画像から得られる複数の特徴量である。

本実施形態では、クロッピングした画像をVAEの入力として得られる潜在変数である。

【0076】

プロセッサ122は、領域（の画像）および／または特徴量ごとにタグを付して管理する（タグ管理処理）。

本実施形態において、クロッピングにより得ている画像とタグとが1対1で対応する。

本実施形態においてタグは、頭部、衣装、全身、顔、上半身、髪、装飾、およびポーズに関するものである（このあたりについては図11も参照されたい）。

【0077】

本実施形態においてプロセッサ122は、画像（ファイル）と、タグを含むテキスト（ファイル）とをセットで記憶部に記憶させることで管理する。

またプロセッサ122は、このデータ（テキスト）の一部を後述するキャラクタプロンプトに用いる。

【0078】

以下にこのテキストについてタグごとに例示する。各タグは領域（の画像）および／または特徴量を説明する文字情報を含む。

なお、この領域（の画像）および／または特徴量を説明する文字情報を便宜上プロンプト要素またはタグ情報と称する場合がある。

- ・頭部（Head）：羽根付き帽子（hat feather）、海賊帽（pirate hat）、海軍帽（navy hat）
- ・衣装（Cloth）：ミニスカート（miniskirt）、制服（uniform）、ニーハイブーツ（knee boots）、ジャケット（jacket）
- ・全身（Full body）：全身（full body）、立っている（standing）、寝ている（lying）、座っている（sitting）
- ・顔：笑顔（smile）、閉じた口（closed mouth）、髭のある（beard）、赤面した（blush）
- ・上半身（Upper body）：上半身（upper body）
- ・髪（Hair）：長髪（long hair）、短髪（short hair）、1つの三つ編み（single braid）、ツインポニーテール（twin ponytails）

10

20

30

40

50

・装飾 (Decoration) : ネックレス (necklace)、イヤリング (earrings)、包帯 (band age)、宝石 (jewelry)

・ポーズ (Pose) : 手を腰に回している (hand on hip)、片足を挙げている (leg lift)、膝を立てている (knee up)

【 0 0 7 9 】

この例ではタグやテキストは単語または熟語である。ただしこれに限るものではなく、文章などであってもよい。

【 0 0 8 0 】

なお一般に、画像データの入力を受け付け、その内容を要約する (説明する) テキストを出力する技術を Image Captioning という。Image Captioning を実行するための (マルチ

10

【 0 0 8 1 】

このような技術はイメージ to プロンプト (Image to Prompt) と呼ばれ、画像からテキストを生成する技術であるが、これは画像中の文字情報を取得する OCR (Optical Character Recognition) とは異なる。

【 0 0 8 2 】

なおここでは画像からテキストに取得する方法について説明しているが、画像に基づく特徴量をテキストに変換する方法を用いてもよい。

【 0 0 8 3 】

またプロセッサ 1 2 2 は、タグに応じた特徴量を取得する。

20

例えばタグごと (つまりはクロッピングしている画像ごと) に特徴量を取得するようにしてもよいし、取得している複数の特徴量をタグに応じて分類してもよい。

【 0 0 8 4 】

なお、一つの領域 (クロッピングにより得ている 1 枚の画像) から複数の特徴量を得てもよいし、複数の領域から 1 つの特徴量を得てもよい。

例えばプロセッサ 1 2 2 は、頭部の画像から頭部に係る特徴量と髪に係る特徴量を取得してもよい。

【 0 0 8 5 】

また図 7 に示すように、爪先より上の部分の画像はキャラクタの全身を含むため、ポーズの特徴量を含んでいる。

30

ただし、ポーズの特徴量は全身の画像がなければ取得できないというものではない。例えば上半身を含む画像は、キャラクタが手を腰に当てている、というポーズの特徴量を含み得る。

【 0 0 8 6 】

また、例えば顔の領域の画像があるからといって、プロセッサ 1 2 2 は上半身の領域の画像から顔の部分に係る特徴量を得ないというものではない。

つまりプロセッサ 1 2 2 が、ある領域 (画像) からどのような情報を抽出するかについては制限されない。

【 0 0 8 7 】

本実施形態において、プロセッサ 1 2 2 が元キャラクタ画像の顔の領域の画像を取得し、顔の特徴量を取得することは必須である。

40

例えばデフォルメキャラクタ画像において特に顔部分が大きく表示されるように、顔の特徴量はデフォルメキャラクタの特徴量を決定づけるためである。

【 0 0 8 8 】

図 6 に戻り、プロセッサ 1 2 2 は、ステップ 2 で得ている複数の特徴量のうち、ポーズの特徴量と、ポーズの特徴量以外の特徴量を含むコンセプトの特徴量とを分離する (ステップ 3 ・特徴量分離処理)。この分離を特徴量分離 (Feature Decoupling) と称する。

【 0 0 8 9 】

特徴量分離を実現する方法として、生成モデルの潜在空間 (Latent Space) を設計し、異なる特徴量を異なる次元で制御するように学習させることが挙げられる。

50

本実施形態において特に、ポーズの特徴量をそれ以外の特徴量と分けている。

特徴量分離により、各特徴量が他の特徴量に依存せずに処理されるため、モデルの柔軟性と精度が向上する。

【0090】

本実施形態において、ポーズ以外の特徴・要素をコンセプトと表現する。例えばポーズの特徴量以外の特徴量をコンセプトの特徴量としている。

上記の例であれば、複数の特徴量のうち、頭、衣装、全身、装飾、顔、上半身、および髪に関する特徴量はコンセプトの特徴量となる。

つまり本実施形態において、複数の特徴量 = ポーズの特徴量 + コンセプトの特徴量である。

10

【0091】

プロセッサ122は元キャラクタ画像から、当該元キャラクタ画像について説明している文字情報を取得する（ステップ4・キャラクタプロンプト取得処理）。

この元キャラクタ画像について説明している文字情報（プロンプト）を便宜上「キャラクタプロンプト」と称する。

【0092】

プロセッサ122は特に、ポーズ以外の要素（画像および/または特徴量）について説明する文字情報を取得する。

本実施形態においてプロセッサ122は、画像のうちポーズ以外に係る部分、つまり頭部、衣装、全身、顔、上半身、髪、または装飾など（コンセプト）を説明する文字情報を取得する。

20

【0093】

キャラクタプロンプトを例示すると以下のようになる。

「1girl, blue eyes, purple hair, long hair, police uniform, transparent background, open clothes, black skirt, mini skirt, high-heeled shoes, simple background, ...（略）」

【0094】

キャラクタプロンプトを取得して後述するデフォルメキャラクタ取得処理の入力に使用することにより、キャラクタの詳細な特徴を制御し、デフォルメキャラクタに反映することができるという利点がある。

30

【0095】

また一方で、プロセッサ122は、上述した複数の特徴量のうち、ポーズの特徴量以外の特徴量のうち少なくとも一部（ポーズの特徴量以外の特徴量を含むコンセプトの特徴量）を学習データとして学習している（学習する）機械学習モデル（拡散モデル）を取得する（ステップ5・キャラクタモデル取得処理）。

この機械学習モデル（拡散モデル）を便宜上「キャラクタモデル」と称する。

【0096】

このキャラクタモデルの訓練構造は、拡散モデル（Diffusionモデル）と当該拡散モデルを制御するネットワーク（ControlNet）を含む。

キャラクタモデルは、ControlNetの重みは凍結して、Diffusionモデルのパラメータのみを更新して（学習して）得られるものである。

40

「（プロセッサ122は）キャラクタモデルを都度作成する」と表現する場合がある。

【0097】

ここで「拡散モデル（Diffusion Model）」の一般的な内容について説明する。

拡散モデルは、本実施形態でも用いているStable Diffusionに関する技術の中核となる機械学習モデルである。

拡散モデルは、ノイズ画像などからノイズを取り除いてノイズなし画像などを生成可能なモデルである。拡散モデルは学習段階において、元情報（画素、潜在空間など）にノイズを加えていく拡散過程と、ノイズあり情報からノイズを除去していく逆拡散過程を含む。

50

【0098】

拡散モデルは、画像などの元情報、加えていくノイズ、元情報からノイズを加えていくことで得るノイズあり情報（例えばノイズ画像など）、ノイズあり情報から除去すべきノイズ、および、ノイズあり情報からノイズを除去していくことで得る元情報、などを学習データとする。

【0099】

そして拡散モデル（に含まれるUnet）は、上記ノイズあり情報などを入力として、ノイズあり情報から取り除くべきノイズを学習する（学習段階）。

そして拡散モデル（に含まれるUnet）は、ノイズあり情報（ノイズ画像など）を入力として、当該ノイズあり情報から除去すべきノイズを出力する。この工程を繰り返す結果として、ノイズあり情報（ノイズ画像など）からノイズが除去されている情報（ノイズなし画像など）を出力することができる（推論段階）。

10

【0100】

本実施形態では、特徴量分離を行うことで得ているコンセプトの特徴量により、機械学習モデル（拡散モデル）のパラメータを更新する。

【0101】

なお、ポーズをVAEにより潜在変数に変換したものはレイアウトの維持などに使用している（ControlNet・図1中の「Freezed」の部分）。

【0102】

例えば本実施形態において、特徴量分離を行うことで得ているコンセプトの特徴量が上記元情報となる。

20

つまりプロセッサ122は、コンセプトの特徴量（潜在変数）にノイズを加えていく拡散工程およびノイズを除去する逆拡散工程を経る学習により機械学習モデル（拡散モデル）を作成し取得する。

【0103】

なおデフォルメキャラクタ画像を得られるのであれば、元キャラクタ画像から複数の画像を取得しない態様（態様1）や、複数の画像を取得した上でポーズの特徴量とそれ以外の特徴量とにデカップリング（Decoupling）をしない方法（態様2）であってもよいが、本実施態様の方法によると元キャラクタ画像の特徴を精度高く反映しているデフォルメキャラクタ画像を得られるという顕著な効果がある。

30

【0104】

小括すると、プロセッサ122は一枚の元キャラクタ画像から複数の特徴量を取得してポーズの特徴量とそれ以外の特徴量（コンセプト）に分離（デカップリング）する。またこのコンセプトを説明する文字情報（キャラクタプロンプト）を取得する。さらにコンセプトの特徴量により学習する機械学習モデル（拡散モデル（キャラクタモデル））を作成する。

【0105】

本実施形態のキャラクタモデルは、1枚の人物（キャラクタ）画像のみを使用して訓練される。この訓練を完了したモデルは、異なるポーズやスタイル（例えば、SD（スーパーデフォルメ）スタイルなど）を持つこの人物（キャラクタ）の特徴を生成するために使用することができる。

40

【0106】

具体的な訓練方法について以下にまとめる。プロセッサは、顔が含まれる画像であって異なるスケールでクロップ（crop, cropping）した画像を使用する。プロセッサはこの異なるスケールの画像からプロンプト（タグ情報）を抽出する。機械学習モデルの訓練構造はDiffusionモデルとControlNetを含み、ControlNetの重みは凍結して、プロセッサはDiffusionモデルのパラメータのみを更新する。パラメータ更新後のDiffusionモデルが、本実施形態のキャラクタモデルである。

【0107】

なお、既存の機械学習モデル（拡散モデル）を基礎として、組み込んで、または改良し

50

て使用してもよい。

例えば、（デフォルメキャラクタ画像を含む）キャラクタ画像にノイズを加えていくことなどによりすでに学習している機械学習モデル（拡散モデル）などを使用してもよい。

【0108】

つぎに、（2）参照画像分析処理について説明する。

プロセッサ122は、デフォルメキャラクタの参照画像を取得する（ステップ6・参照画像取得処理）。

【0109】

そして本実施形態においてプロセッサ122は、まず前処理として参照画像のコンセプト情報を削除するとともに、デフォルメキャラクタの形状の一部を示す線画情報（Scribbleと呼ばれる線画情報）と、ポーズ情報（Pose）と、を取得する。

10

【0110】

図1左下に示すように、プロセッサ122は線画情報として少なくとも参照画像の頭部、手先部分、および足先部分の線画を取得する。

これにより出力するデフォルメキャラクタのポーズ（向きなど）や頭身比率などを制御することができる（後述）。

【0111】

また、プロセッサ122はポーズ情報として参照画像の関節位置情報を取得する。

【0112】

図1左下において、白黒反転している画像の上段が線画情報（Scribble）、下段がポーズ情報（Pose）のイメージ図である。

20

【0113】

本実施形態においてプロセッサ122は、取得している線画情報（Scribble）およびポーズ情報（Pose）から、制御用（機械学習）モデルを作成する。

この制御用（機械学習）モデルを便宜上「コントロールモデル」と称する。

【0114】

このコントロールモデルは、拡散モデルを制御する追加的なニューラルネットワークである。

特に、入力画像のキャラクタなどを、意図したポーズなどに制御できるという優れた効果を奏する。

30

本実施形態のコントロールモデルは、デフォルメキャラクタ画像のポーズおよび/または（頭身数に影響する）外形などを制御する。

【0115】

上述した処理は、Stable Diffusionのプラグインや拡張機能で実行することができる。例えばStable Diffusion Web UI・Automatic1111版のWeb UIに追加できる「ControlNet Extension」などを好適に用いることができる。

【0116】

本実施形態において、線画情報の取得には（ControlNet）ControlNet Scribble（ControlNetのモデルの一つ）を用いている。

また、ポーズ情報の取得には（ControlNet）Open Pose（ControlNetのモデルの一つ）を用いている。

40

【0117】

なお、Open Poseは姿勢推定手法である。キャラクタの画像にOpen Poseを適用することで、そのキャラクタのポーズを推定する。Open PoseはParts Affinity Fieldsと呼ばれる骨格間の位置関係を考慮した処理を導入している。

【0118】

ControlNetの各種「preprocessor」が入力画像を解析・変換（例：ポーズ検出、エッジ検出、線画化など）し、それをモデルの条件として使用する。つまりControlNetの機能として、画像を線画化（edge detection / scribble detection）する仕組みが用意されている場合もあり、ユーザはそれを活用して「Scribble（線画）」を生成することができる。

50

【0119】

小括するとプロセッサ122は、少なくとも（参照画像における）デフォルメキャラクタの（形状の）一部を示す線画情報に基づくデータを学習データとして学習している制御用（機械学習）モデルを作成し取得する（ステップ7・コントロールモデル取得処理）。

【0120】

より具体的には、（参照画像における）デフォルメキャラクタの（形状の）一部を示す線画情報およびポーズ情報に基づくデータ（具体的には線画情報およびポーズ情報をVAEにより特徴量（潜在変数）化したもの）を学習データとして学習している、（上述した拡散モデルを制御するための追加的なニューラルネットワークである）制御用（機械学習）モデル（＝コントロールモデル）を作成し取得する。

10

つまり、この線画情報およびポーズ情報は制御用モデルの入力（条件（Condition））になる。

【0121】

なお本実施形態において、このコントロールモデルは都度作成しているが、参照画像ごとにあらかじめ作成してサーバ10の記憶部14に記憶しているコントロールモデルを取得するものであってもよい。この場合、処理が早くなるという利点がある。

ただし、ユーザが新たな画像を参照画像としてアップロードしているときは、その新たな画像についてプロセッサ122はコントロールモデルを作成する。

【0122】

図6に戻り、最後に（3）デフォルメキャラクタ取得処理について説明する。

20

プロセッサ122は、制御用モデルにより制御されている機械学習モデル（拡散モデル）と文字情報とから、元キャラクタ画像をデフォルメ化しているデフォルメキャラクタ画像を取得する（ステップ8・デフォルメキャラクタ取得処理）。

【0123】

より具体的には、上述の制御用モデル（＝コントロールモデル）により制御されている拡散モデル（＝キャラクタモデル）に対して文字情報（＝キャラクタプロンプト）を入力することで、元キャラクタ画像をデフォルメ化しているデフォルメキャラクタ画像を生成して出力として取得する。

【0124】

（不慣れなユーザなどが）プロンプトにより条件入力することなどに比べ、格段に精度が高いデフォルメキャラクタ画像が得られるという顕著な効果がある。

30

【0125】

そしてプロセッサ122は、取得しているデフォルメキャラクタ画像を端末20の表示部284aに表示させる。

ユーザにより複数の頭身比率が選択されているときは、それぞれの頭身比率に対応するデフォルメキャラクタ画像を出力する。

またUIの項で示したように、プロセッサ122は、ユーザによる入力を受け付けてデフォルメキャラクタ画像を再生成することなどができる。

【0126】

なお、本実施形態におけるデフォルメキャラクタ生成処理の効果を発揮できるものであれば、処理の順序は図6のものに限られない。例えば図6において、ステップ4、ステップ5、およびステップ6～7を逐次的に描画しているが、これらは並列的に処理してもよい。

40

または、ステップ4、ステップ2・3・5、およびステップ6～7を並列的に処理してもよい。

【0127】

<2-2.最適化処理>

最適化処理において、プロセッサ122は、上述した拡散モデル（キャラクタモデル）が備えるニューラルネットワーク（Unet）の少なくとも一部の層において、重みを変更して得る複数の出力画像の差分を取得して評価し、出力画像に影響の小さい層の重みの削減

50

または削除を行う。

【0128】

プロセッサ122は、最適化プログラムP14に基づき、最適化処理を行う。

すなわち、最適化プログラムP14は、プロセッサ122による最適化処理の実行により、コンピュータを最適化手段として機能させる。

【0129】

図8は、最適化処理を示すフローチャートである。

プロセッサ122は、ユーザによるデフォルメ画像再生成の指示を受け付けることにより、またはあらかじめ設置している条件に基づき自動的に、最適化処理を開始する（図5・生成画像表示画面参照）。

本実施形態において、ユーザは画像の少なくとも一部を変化させている画像を再生成することができる。

【0130】

まずプロセッサ122は、（拡散モデルの）ニューラルネットワークにおける複数の層のそれぞれにおいて重みを1および0に設定した場合における、デフォルメキャラクタ画像の差分を取得する（ステップ21）。

プロセッサ122は、影響が大きい層の重みを保持する一方で、影響が大きい層の重みを削減または削除（以下「削減等」とする。）する（ステップ22）。

【0131】

図9は、重みの削減等を説明するための図である。

図9において、右上の角丸四角で囲まれている部分はニューラルネットワークの同じ層であることを示す。また、同じ角丸四角で囲まれている2体のキャラクタはそれぞれ、当該層において重みを0としている場合と1としている場合におけるデフォルメキャラクタ画像である。

図9中のDiffは、重みが0のときのデフォルメキャラクタ画像と1のときのデフォルメキャラクタ画像との差分を示す。白抜きの部分が差分である。

【0132】

ある層において重みを1および0に設定して得たそれぞれの画像の差分があまり大きくない場合、その層は全体の結果に対して大きな影響を与えていないことを意味する。

つまり各層でこの差分を評価することで、当該各層の全体に対する影響度を把握することができる。そしてプロセッサ122は、この結果を受けて影響の少ない層の重みを削除等する。

【0133】

例えば、図9の（4つの角丸四角で表す層の）一番右の層では、重みの違い（重み1および0）により生成されるそれぞれの画像の差が大きい（差分が大きい）。このことから、この層の重みの変化により生成される画像の変化が大きいこと、つまりこの層の重みが画像生成に大きく影響することがわかる。

一方で、左から2番目の層は差分が小さいことから、この層が生成画像に与える影響は小さいと考えられる。

【0134】

この設定は学習済みモデルのファインチューニング方法（LoRA）におけるパラメータ層を最適化し、最終的に高品質な結果を生成するという効果を奏する。

【0135】

図8のフローチャートに戻る。

この影響が大きい層、つまり鍵となる特徴に関する層を（Unetの）ニューラルネットワークの出力層に配置し、また、誤差逆伝播によりニューラルネットワークの層の重みを更新する（ファインチューニング）（ステップ23）。

ファインチューニングにより、元キャラクタ画像の特徴を適切に再現しているデフォルメキャラクタ画像が得られる。

【0136】

10

20

30

40

50

最適化処理により得られた最適な画像の例を示す。

図10は、最適化処理前後のデフォルメキャラクタ画像を示す図である。

図10に示すように、最適化後の画像は最適化前の画像より元キャラクタ画像を忠実に再現している（特に腰のあたりの四角で囲っている部分）。

【0137】

ここで、本実施形態の具体的な方法について説明する。

上記の処理を行う方法として例えばLoRA Block Weightを用いることができる。LoRA Block Weightにより、LoRAの適用度を調整することができる。

【0138】

なおLoRA (Low-Rank Adaptation) は、学習済みモデルに対してファインチューニングを行う方法である。

そのなかで、コンセプトを対象にチューニングを行うもの、生成したい画像のコンセプトを反映させるもの、またはキャラクタに特定のコンセプトを適用させるものをConcept LoRAと称する。Concept LoRAを用いることで、上記のような調整を行うことができる。

【0139】

以上まとめると、元キャラクタ画像分析処理により、1枚の元キャラクタ画像の複数の領域から複数の特徴量を取得し、デフォルメキャラクタに元キャラクタ画像の特徴を多く反映できるようにしている。ここではキャラクタプロンプトやキャラクタモデルを作成（取得）する。

また、参照画像分析処理により、参照画像から線図やポーズを抽出する処理などを経て、キャラクタモデルを制御するためのコントロールモデルを取得する。

そしてデフォルメキャラクタ取得処理において、コントロールモデルにコントロール（制御）されているキャラクタモデルに対してキャラクタプロンプトを入力することで、元キャラクタ画像に忠実なデフォルメキャラクタを生成することができる。その精度の高さは図面で示すとおりである。

そして最適化処理により、さらに高品質なデフォルメ画像を得ることができる。

【0140】

また、上述する拡散モデル（キャラクタモデル）は、キャラクタの画像である元キャラクタ画像から当該元キャラクタ画像をデフォルメ化しているデフォルメキャラクタ画像を取得する学習済みモデルであって、

前記元キャラクタ画像から複数の特徴量を取得し、当該複数の特徴量のうちポーズの特徴量以外の特徴量を含むコンセプトの特徴量、前記コンセプトの特徴量に加えるノイズ、およびノイズを加えることで得るノイズあり情報、を少なくとも含むデータを学習データとして、ノイズあり情報から取り除くべきノイズを学習し、

ノイズあり情報を入力として、当該ノイズあり情報から除去すべきノイズを出力することを特徴とし、

さらに、参照画像におけるデフォルメキャラクタの（形状の）一部を示す線画情報を少なくとも1つの学習データとして学習している制御用モデル（＝コントロールモデル）により、出力するデフォルメキャラクタ画像におけるポーズが制御されることを特徴とする、学習済みモデルである。

【0141】

図11は、上述したデフォルメキャラクタ生成処理を（図1とは別の形で）示す図である。

キャラクタの画像（元キャラクタ画像）からデフォルメキャラクタ画像を得るほか、デフォルメキャラクタ画像を最適化する処理を示す。

【0142】

3．ハードウェア構成

図12はデフォルメキャラクタ生成システム1の概要を示すネットワーク構成図である。

図12に示すように、本実施形態におけるデフォルメキャラクタ生成システム1は、シ

10

20

30

40

50

ステムサーバ１０（サーバ１０）、および（ユーザ）端末２０を備える。また、これらの各装置は、ネットワークＮを介して接続されている。ネットワークＮは例えばインターネットなどである。

サーバ１０には、デフォルメキャラクタ生成プログラムＰ１を含み、本実施形態に係るデフォルメキャラクタ生成システム１を動作させるためのソフトウェア（アプリケーションソフトウェア）がインストールされており、当該ソフトウェアの機能により、各種処理が実行される。

以下、各ハードウェアについて説明する。

【０１４３】

<サーバ１０>

サーバ１０は、デフォルメキャラクタ生成プログラムＰ１を実行するためのコンピュータである。またサーバ１０は、ユーザの用に供するユーザインターフェースを提供するウェブサイトＷを表示する。

【０１４４】

図１２においてサーバ１０は１台のみ図示しているが、数は１台に限られるものではなく、複数のサーバにより実現してもよい。例えばウェブサーバと機械学習用サーバを別にするなどが挙げられる。

複数のサーバは負荷分散や可用性の観点から用いることも考えられる。

また、画像生成に関する処理を、他のサーバに任せるものまたは他のサーバと協力して行うものであってもよい。

このほか、サーバ１０はクラウドサービス事業者のコンピュータを利用するものであってもよいし、ユーザがコンピュータを用意してもよい。

【０１４５】

図１３は、サーバ１０のハードウェア構成図である。

図１３に示すように、サーバ１０は、制御部１２、記憶部１４、および通信制御部１６を備える。また制御部１２は、プロセッサ１２２、ＲＯＭ１２４、ＲＡＭ１２６、計時部１２８を備える。それぞれの基本的な機能については後でまとめて説明する。

【０１４６】

プロセッサ１２２を備える制御部１２は、サーバ１０においてデフォルメキャラクタ生成部としても機能する（不図示）。デフォルメキャラクタ生成部は、デフォルメキャラクタ生成プログラムＰ１を実行してデフォルメキャラクタ生成処理を行う。本実施形態において、プロセッサ１２２はＣＰＵ（Central Processing Unit）である。

【０１４７】

また、一のプログラムは、別のプログラムを含んでいてもよい。例えば本実施形態において、デフォルメキャラクタ生成プログラムＰ１は、元キャラクタ画像分析プログラムＰ１１や参照画像分析プログラムＰ１２などを含む。

【０１４８】

図１３に示すように、記憶部１４は、プログラム格納部１４ａとデータ格納部１４ｂを備え、各種処理に必要なプログラムやデータを備える。

例えば、プログラム格納部１４ａには、本実施形態に係るデフォルメキャラクタ生成プログラムＰ１のほか、サーバ１０に接続されている機器を制御するための制御プログラム、例えば通信制御部１６を制御する通信制御プログラムなどが格納されている。

【０１４９】

図１３に示すように、通信制御部１６は、サーバ１０と、外部にある端末、例えば、後述する端末２０などとの間で通信を行う装置である。通信制御部１６は図１２に示すように、サーバ１０をネットワークＮに接続する。

【０１５０】

上記のほか、サーバ１０は、命令やデータの入力を行うための入力部（例えばキーボード）や、情報を何らかの形で出力するための出力部（例えば音声出力装置）などを備えていてもよい（不図示）。また、本実施形態の用途のために追加的に必要な装置や、本実

10

20

30

40

50

施形態の用途について利便性を向上させるための装置を備えていてもよい。

【0151】

< 端末20 >

端末20は、ユーザがデフォルメキャラクタ生成システム1を利用するための情報処理装置である。ユーザは、端末20を用いてサーバ10にアクセスすることにより、デフォルメキャラクタ生成システム1を利用する。

端末20は、制御部22、記憶部24、通信制御部26、および入出力部28を備える。また制御部は、プロセッサ、ROM、RAM、および計時部を備える。

入出力部28は入力部282および出力部284を備え、出力部284は表示部284aを備える。

説明済みの内容と重複するものや、後述する基本的な機能に関するものは説明を省略する。

【0152】

本実施形態において、端末20はデスクトップPCである。ただし、端末20はこれに限られるものではなく、スマートフォンやタブレットなどの携帯型端末であってもよい。

【0153】

なお本実施形態において、端末20は特定のアプリケーションプログラムをインストールする必要はなく、ウェブサイトWにアクセスすることで各種処理が実行可能であり、デフォルメキャラクタ生成システム1を利用することができる。

【0154】

(コンピュータの基本的機能に係る説明)

以下、制御部(プロセッサ、ROM、RAM、計時部)、記憶部、通信制御部、入力部、および出力部について説明する。

なお、本実施形態のいずれの端末においても、機能部間の接続態様(ネットワークトポロジ)は特に限定されない。例えばバス型であってもよいし、スター型、メッシュ型などであってもよい。

【0155】

プロセッサは、ROMや記憶部などに記憶されたプログラムに従って、情報処理や各種装置の制御を行う。本実施形態において、プロセッサはCPU(Central Processing Unit)である。

【0156】

ただし、プロセッサはCPUに限られない。CPU、DSP(Digital Signal Unit)、GPU(Graphics Processing Unit)、GPGPU(General Purpose computing on GPU)、TPU(Tensor Processing Unit)、またはASIC(Application Specific Integrated Circuit)など、各種プロセッサを単独で、あるいは組み合わせて用いてもよい。

例えば、CPUとGPUを統合したプロセッサはAPU(Accelerated Processing Unit)などと呼ばれるが、このようなプロセッサを用いてもよい。

【0157】

ROMは、プロセッサが各種制御や演算を行うための各種プログラムやデータがあらかじめ格納された、リードオンリーメモリである。

【0158】

RAMは、プロセッサにワーキングメモリとして使用されるランダムアクセスメモリである。このRAMには、本実施形態の各種処理を行うための各種エリアが確保可能になっている。

【0159】

計時部は、時間情報の取得などに係る計時処理を行う。コンピュータが通信制御部を備える場合は、NTP(ネットワーク・タイム・プロトコル)により外部から時間情報を取得してもよい。

【0160】

10

20

30

40

50

記憶部は、プログラムやデータなどの情報を記憶するための装置である。記憶部はストレージとも称する。記憶部は内蔵型か、外付型かを問わない。

【0161】

記憶部は、データの読み書きが可能な記憶媒体と、当該記憶媒体に読み書きするドライブとを含む。

記憶媒体は例えば、内蔵型や外付型があり、H D（ハードディスク）、C D - R O M、フラッシュメモリなどが挙げられる。

ドライブは例えば、H D D（ハードディスクドライブ）、S S D（ソリッドステートドライブ）などが挙げられる。

【0162】

記憶部は、機能部としてプログラム格納部とデータ格納部を備える。

プログラム格納部には、各種機器を制御するための制御プログラム、例えば通信を制御する通信制御プログラムなどが格納されている。

【0163】

通信制御部は、端末などの間で通信を行うための装置である。通信制御部は、当該通信制御部を備える端末をネットワークNに接続する。

【0164】

通信制御部の通信方式は公知の方式であり、機器に応じて有線による方式や無線による方式が適用される。

例えば、端末がデスクトップP Cであれば有線、無線の両方の場合が考えられ、また、端末がスマートフォンであれば、無線による通信方式が考えられる。

【0165】

有線であれば、例えばI E E E 8 0 2 . 3（例えばバス型やスター型の有線L A N）で規定される通信方式を好適に用いることができるが、それ以外にも、I E E E 8 0 2 . 5（例えばリング型の有線L A N）で規定される通信方式などを用いてもよい。

【0166】

無線であれば、例えばI E E E 8 0 2 . 1 1（例えばW i - F i）で規定される通信方式を好適に用いることができるが、それ以外にも、I E E E 8 0 2 . 1 5（例えばブルートゥース（登録商標）、B L E（ブルートゥース（登録商標）ローエナジー）など）、I E E E 8 0 2 . 1 6（例えばW i M A X）、Z i g B e e（登録商標）、9 2 0 M H z 帯無線（W i - S U Nなど）、または赤外線通信などの光通信で規定される通信方式などを用いてもよい。

【0167】

入力部および出力部は、それぞれ端末に対する入力と出力を担う装置である。入力部および出力部をあわせて入出力部と称する場合がある。

入力部はユーザからの入力を受け付ける装置である。このような入力部として例えば、キーボード、ポインティングデバイスとしてのマウス、トラックパッド、タブレット、またはタッチパネルなどが挙げられる。

【0168】

端末がタブレットやスマートフォンなどであって、入力部がタッチパネルの場合、入力部はタッチスクリーンなど、画像などを表示する表示部の表面に配置される。この場合、入力部は、表示部に表示される各種操作アイコンに対応したユーザのタッチ位置を特定し、ユーザによる入力を受け付ける。

【0169】

出力部は例えば、画像や音声、帳票などを出力するための装置である。

出力部として例えば、タッチスクリーンやディスプレイ（液晶ディスプレイや有機E Lディスプレイ）などの表示装置（表示部）や、スピーカなどの音声出力装置、プリンタなどの帳票出力装置が挙げられる。

【0170】

以上のような構成により、ユーザはオンラインで用意に所望のデフォルメキャラクタを

10

20

30

40

50

取得することができる。

【0171】

(第二の実施形態)

第二の実施形態は、デフォルメキャラクタの全身ではなく顔部分のみを生成する実施形態である。

【0172】

図14は第二の実施形態の概要を示す図である。

本実施形態において、デフォルメキャラクタ生成プログラムP1は図14左上の元キャラクタ画像と左下の参照画像(顔部分)を入力として、右下のデフォルメキャラクタの顔部分を出力する。

10

本実施形態において、「顔」は頭部全体、例えば帽子などの装飾を含んでよい。

【0173】

このとき、UI、処理(デフォルメキャラクタ生成処理および最適化処理)、およびハードウェア構成は原則として第一の実施形態と同じである。

【0174】

ただし本実施形態の処理やデータにおいて、元キャラクタ画像からは頭部(顔)の特徴量が抽出できればよい。

例えば、上述した複数特徴量取得処理において、顔を含む領域の画像が取得できればよい。

キャラクタプロンプトは顔部分の特徴量を説明する文字情報であり、キャラクタモデルは顔部分のデフォルメキャラクタ画像(デフォルメキャラクタ顔画像)を得るためのモデルである。

20

【0175】

また、コントロールモデル取得処理において、上述した実施形態ではプロセッサ122がポーズ情報(Pose)とデフォルメキャラクタの形状の一部を示す線画情報(Scribbleと呼ばれる線画情報)とを取得していたが、本実施形態ではデフォルメキャラクタの形状の一部を示す線画情報だけでよい。

【0176】

図14に示すように、本実施形態において、プロセッサ122はデフォルメキャラクタの形状の一部を示す線画情報として目、口、顎部、および首部の線画情報を取得する。

30

【0177】

ただし、線画情報以外の情報から特徴量を取得し、当該特徴量をコントロールモデル(の入力)に用いてもよい。

このような特徴量として例えば、元キャラクタ画像における右目と左目の間隔の情報、または両目と鼻筋の関係を示す情報(デッサンにおける十字線、いわゆる「あたり」と呼ばれるものなど)に基づく特徴量、つまり顔の部位(目、鼻、または口など)の位置関係から得られる特徴量が挙げられる。

追加の情報を加えることにより、顔の特徴をよりよく捉えたデフォルメキャラクタ画像を得ることができる。

【0178】

小括すると、第一の実施形態におけるキャラクタの画像がキャラクタの顔画像であり、デフォルメキャラクタ画像が、デフォルメキャラクタ顔画像であってもよい。

40

【0179】

このように、デフォルメキャラクタ生成プログラムP1が生成できるのはデフォルメキャラクタの全身像だけではなく、全身の一部分である頭部(顔)部分のデフォルメも可能とする。

キャラクタの顔アイコンがSNSのプロフィール画像に使用されるなど、デフォルメキャラクタの顔画像の生成は高いニーズがある。

デフォルメキャラクタ生成プログラムP1は、このようなニーズに対応することもできる。

50

【 0 1 8 0 】

図 1 5 は、本実施形態（第一の実施形態と第二の実施形態）のデフォルメキャラクタ生成プログラム P 1 によるデフォルメ化の応用例を示す図である。

1 つの元キャラクタ画像から、顔または全身についての複数のデフォルメ画像を生成することができる。

【 0 1 8 1 】

（変形例）

本発明は上述の実施形態に限られるものではなく、本発明の趣旨を逸脱しない範囲において、上述の実施形態に種々の変更を加えたものを含む。

【 0 1 8 2 】

例えば上述の実施形態では、パラメータの更新により機械学習モデル（拡散モデル）を都度作成していたが、機械学習モデルの学習や作成はこれに限るものではない。例えば、パラメータの更新に限らず、モデル構造を変更するようなものや、モデル構造から作成するようなものであってもよい。

ただし、パラメータの更新は処理が早くより好ましい。

【 0 1 8 3 】

また上述したように、すでに作成されている機械学習モデル（拡散モデル）を取得するようにしてもよい。

また、作成したキャラクタモデルを、（既存の・他の）Stable Diffusionの拡散モデルに統合する（組み込む）ようにしてもよい。

ただし、元キャラクタ画像を取得して分析し、拡散モデルのパラメータを更新して新たな拡散モデルとすることで、デフォルメキャラクタ画像における元キャラクタ画像の再現度は著しく向上する。

【 0 1 8 4 】

上述の実施形態において、機械学習モデル（キャラクタモデル）は拡散モデルとコントロールモデルを含むものであったが、これに限られない。まず、拡散モデルのみで構成されていてもよい。また、本発明の目的を達するものであれば、拡散モデル以外の機械学習モデルで構成してもよい。

ただし、拡散モデルとコントロールモデルを含む態様は、元キャラクタ画像の特徴量を精度高く再現しているデフォルメキャラクタ画像を生成することができる。

【 0 1 8 5 】

特徴量分離（Feature Decoupling）を実現するために、以下のアプローチが用いられ得る。上記の実施形態では次の（１）を用いている。

（１）潜在空間の分離

生成モデルの潜在空間（Latent Space）を設計し、異なる次元が異なる特徴量を制御するように学習させる。例えば、各次元が「髪型」、「顔（表情）」、または「ポーズ」などをそれぞれ担当する。

（２）アーキテクチャ設計

生成ネットワーク（GANやDiffusion Modelなど）の層やモジュールを設計し、特定の層が特定の特徴量を処理するようにする。例えばスタイル変換のための「マルチスケール特徴量抽出」や「レイヤーレベルでのスタイル適用」などである。

（３）学習の正則化（Regularization）

特徴量間の独立性を保つために、学習プロセス中に特徴量間の相互作用を抑制するペナルティを導入する。

（４）モジュール間の条件付け

特定の条件（テキスト、属性、スケッチ）に基づいて生成プロセスを調整することで、特定の特徴量を強調する。

【 0 1 8 6 】

本実施形態を含む本発明の態様は、換言すると以下の特徴量を備える。下記は本願出願時における特許請求の範囲と対応する。ただし、出願後における特許請求の範囲の補正に

10

20

30

40

50

より、当該補正後の特許請求の範囲の記載とは異なる場合がある。

(1) 第1の態様では、コンピュータを、(1)元キャラクタ画像分析手段、(2)参照画像分析手段、および(3)デフォルメキャラクタ取得手段として機能させ、前記(1)元キャラクタ画像分析手段は、キャラクタの画像である元キャラクタ画像を取得する元キャラクタ画像取得手段、前記元キャラクタ画像から、複数の特徴量を取得する複数特徴量取得手段、前記元キャラクタ画像から、当該元キャラクタ画像について説明している文字情報を取得するキャラクタプロンプト取得手段、および、前記複数の特徴量のうち、ポーズの特徴量以外の特徴量のうち少なくとも一部を学習データとして学習している機械学習モデルを取得するキャラクタモデル取得手段、を備え、前記(2)参照画像分析手段は、デフォルメキャラクタの参照画像を取得する参照画像取得手段、および、少なくとも前記デフォルメキャラクタの一部を示す線画情報に基づくデータを学習データとして学習している制御用モデルを取得するコントロールモデル取得手段、を備え、前記(3)デフォルメキャラクタ取得手段は、前記制御用モデルにより制御されている前記機械学習モデルと前記文字情報とから、前記元キャラクタ画像をデフォルメ化しているデフォルメキャラクタ画像を取得することを特徴とする、デフォルメキャラクタ生成プログラムを提供する。

10

(2) 第2の態様では、前記少なくともデフォルメキャラクタの一部を示す線画情報に基づくデータを学習データとして学習している制御用モデルが、少なくともデフォルメキャラクタの一部を示す線画情報およびポーズ情報に基づくデータを学習データとして学習している制御用モデルであることを特徴とする、第1の態様に記載のデフォルメキャラクタ生成プログラムを提供する。

20

この場合、デフォルメキャラクタに関する情報量が多くなることから、入力している参照画像により近いポーズ等のデフォルメキャラクタ画像が得られやすくなる。

(3) 第3の態様では、さらに、前記機械学習モデルが備えるニューラルネットワークの少なくとも一部の層において、重みを変更して得る複数の出力画像の差分を取得して評価し、出力画像に影響の小さい層の重みの削減または削除を行う最適化手段を備えることを特徴とする、第1の態様に記載のデフォルメキャラクタ生成プログラムを提供する。

この場合、元キャラクタ画像の再現度がさらに高いデフォルメキャラクタ画像を得ることができる。

(4) 第4の態様では、前記キャラクタの画像が、キャラクタの顔画像であり、前記デフォルメキャラクタ画像が、デフォルメキャラクタ顔画像であることを特徴とする、第1の態様に記載のデフォルメキャラクタ生成プログラムを提供する。

30

この場合、顔の(アップの)画像についてのデフォルメキャラクタ画像が得られる。

(5) 第5の態様では、(1)元キャラクタ画像分析部と、(2)参照画像分析部と、および(3)デフォルメキャラクタ取得部と、を備え、前記(1)元キャラクタ画像分析部は、キャラクタの画像である元キャラクタ画像を取得する元キャラクタ画像取得部と、前記元キャラクタ画像から、複数の特徴量を取得する複数特徴量取得部と、前記元キャラクタ画像から、当該元キャラクタ画像について説明している文字情報を取得するキャラクタプロンプト取得部と、前記複数の特徴量のうち、ポーズの特徴量以外の特徴量のうち少なくとも一部を学習データとして学習している機械学習モデルを取得するキャラクタモデル取得部と、を備え、前記(2)参照画像分析部は、デフォルメキャラクタの参照画像を取得する参照画像取得部と、少なくとも前記デフォルメキャラクタの一部を示す線画情報に基づくデータを学習データとして学習している制御用モデルを取得するコントロールモデル取得部と、を備え、前記(3)デフォルメキャラクタ取得部は、前記制御用モデルにより制御されている前記機械学習モデルと前記文字情報とから、前記元キャラクタ画像をデフォルメ化しているデフォルメキャラクタ画像を取得することを特徴とする、デフォルメキャラクタ生成システムを提供する。

40

(6) 第6の態様では、(1)元キャラクタ画像分析ステップ、(2)参照画像分析ステップ、および(3)デフォルメキャラクタ取得ステップ、から構成されるデフォルメキャラクタ生成方法であって、前記(1)元キャラクタ画像分析ステップは、プロセッサが

50

キャラクタの画像である元キャラクタ画像を取得する元キャラクタ画像取得ステップ、プロセッサが前記元キャラクタ画像から、複数の特徴量を取得する複数特徴量取得ステップ、プロセッサが前記元キャラクタ画像から、当該元キャラクタ画像について説明している文字情報を取得するキャラクタプロンプト取得ステップ、および、プロセッサが、前記複数の特徴量のうち、ポーズの特徴量以外の特徴量のうち少なくとも一部を学習データとして学習している機械学習モデルを取得するキャラクタモデル取得ステップ、を備え、前記(2)参照画像分析ステップは、プロセッサがデフォルメキャラクタの参照画像を取得する参照画像取得ステップ、および、プロセッサが、少なくとも前記デフォルメキャラクタの一部を示す線画情報に基づくデータを学習データとして学習している制御用モデルを取得するコントロールモデル取得ステップ、を備え、前記(3)デフォルメキャラクタ取得ステップは、プロセッサが、前記制御用モデルにより制御されている前記機械学習モデルと前記文字情報とから、前記元キャラクタ画像をデフォルメ化しているデフォルメキャラクタ画像を取得することを特徴とする、デフォルメキャラクタ生成方法を提供する。

10

(7)第7の態様では、キャラクタの画像である元キャラクタ画像から当該元キャラクタ画像をデフォルメ化しているデフォルメキャラクタ画像を取得する学習済みモデルであって、前記元キャラクタ画像から複数の特徴量を取得し、当該複数の特徴量のうちポーズの特徴量以外の特徴量を含むコンセプトの特徴量、前記コンセプトの特徴量に加えるノイズ、およびノイズを加えることで得るノイズあり情報、を少なくとも含むデータを学習データとして、ノイズあり情報から取り除くべきノイズを学習し、ノイズあり情報を入力として、当該ノイズあり情報から除去すべきノイズを出力することを特徴とし、さらに、参照画像におけるデフォルメキャラクタの一部を示す線画情報を少なくとも1つの学習データとして学習している制御用モデルにより出力するデフォルメキャラクタ画像におけるポーズが制御されることを特徴とする、学習済みモデルを提供する。

20

【産業上の利用可能性】

【0187】

可愛いキャラクタはオンライン・オフラインを問わず産業（特に娯楽産業）上の至るところでニーズがあることから、このようなキャラクタを生成可能なプログラムはニーズに即したキャラクタの提供を可能にする。

【符号の説明】

【0188】

30

1 デフォルメキャラクタ生成システム

- 10 サーバ
- 12 制御部
- 122 プロセッサ
- 124 ROM
- 126 RAM
- 128 計時部
- 14 記憶部
- 14a プログラム格納部
- 14b データ格納部
- 16 通信制御部
- 18 入出力部
- 20 端末
- 22 制御部
- 24 記憶部
- 26 通信制御部
- 28 入出力部
- 282 入力部
- 284 出力部
- 284a 表示部

40

50

U I - 1 1 アップロード欄
U I - 1 2 デフォルメキャラクター設定欄
U I - 1 2 1 等身比率設定ボタン
U I - 1 2 2 参照デフォルメキャラクター画像
U I - 1 2 3 デフォルメキャラクター生成開始ボタン
U I - 1 3 元キャラクター画像確認画面（ウィンドウ）
U I - 1 3 1 画像情報
U I - 1 3 2 元キャラクター画像
U I - 1 3 3 O K ボタンアイコン
U I - 1 4 出力画像表示欄
U I - 1 4 1 出力画像
U I - 1 4 2 出力画像加工関連ボタン
U I - 1 4 3 出力画像情報
U I - 1 4 4 出力画像保存関連ボタン
P 1 デフォルメキャラクター生成プログラム
P 1 1 元キャラクター画像分析プログラム
P 1 2 参照画像分析プログラム
P 1 3 デフォルメキャラクター取得プログラム
P 1 4 最適化プログラム

10

【要約】

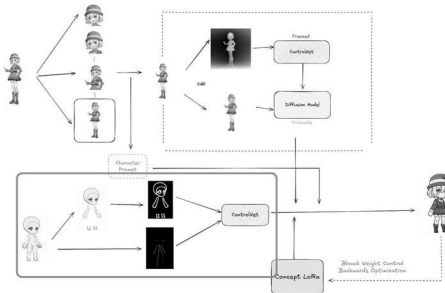
20

【課題】 必要な特徴を残しながら、通常のキャラクターからデフォルメ化キャラクターを生成する。

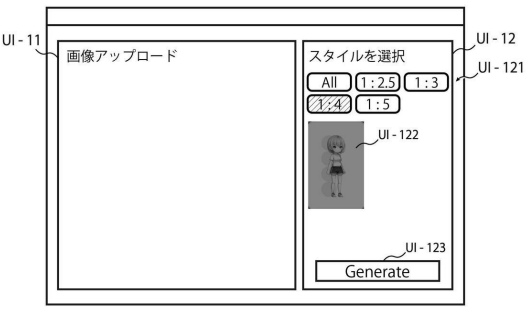
【解決手段】 デフォルメキャラクター生成プログラム P 1 は、（ 1 ）元キャラクター画像分析処理、（ 2 ）参照画像分析処理、および（ 3 ）デフォルメキャラクター取得処理により、元キャラクター画像からデフォルメ化キャラクターを生成する。デフォルメ化キャラクターの生成に際して必要な特徴を残すために入力画像（例えば 8 頭身キャラクター画像）から複数の特徴量を取得し、そのなかで、ポーズの特徴量と、ポーズの特徴量以外の特徴量とを分離する（デカップリングする）ことを最も主要な特徴とする。

【選択図】 図 1

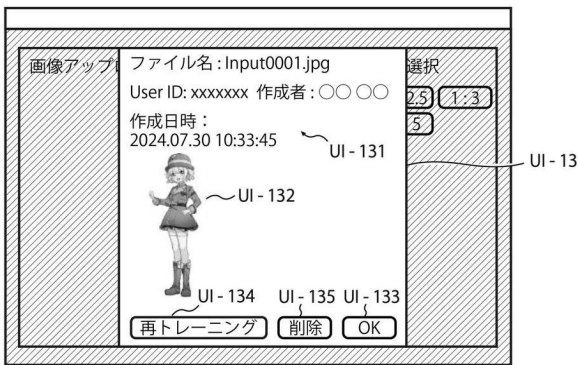
【図 1】



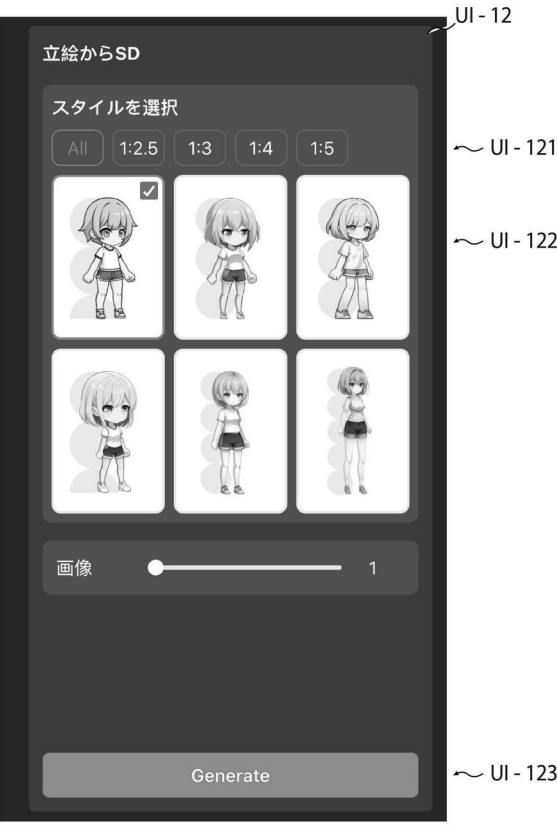
【図 2】



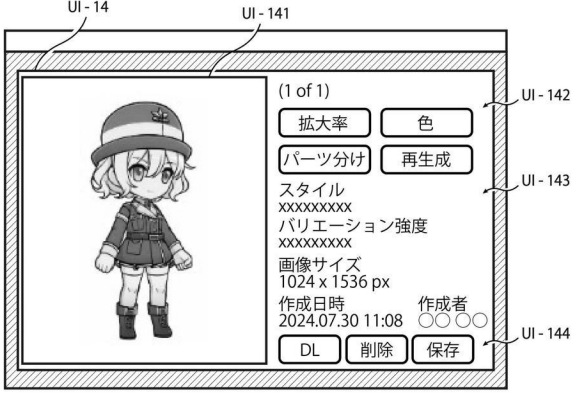
【図 3】



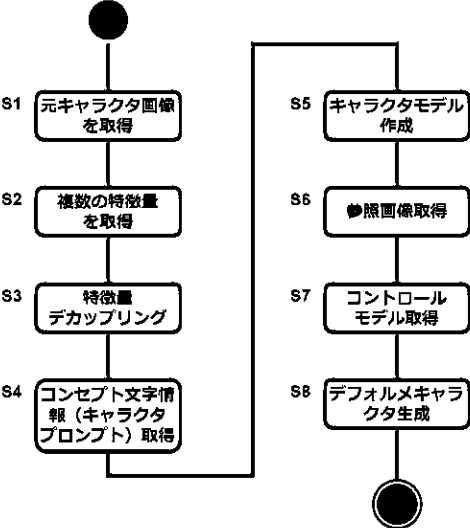
【図 4】



【図 5】



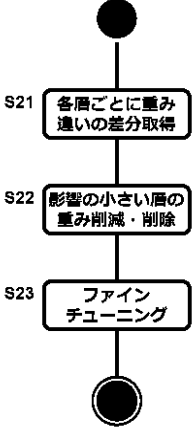
【図 6】



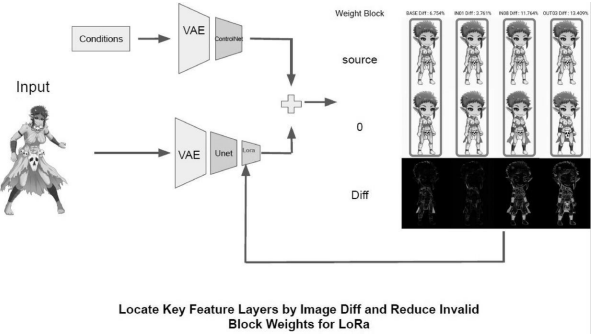
【図 7】



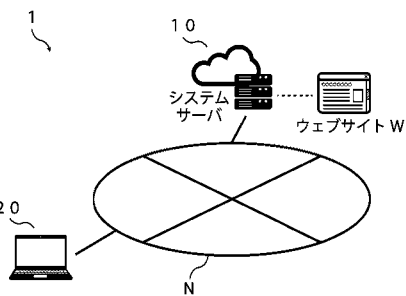
【図 8】



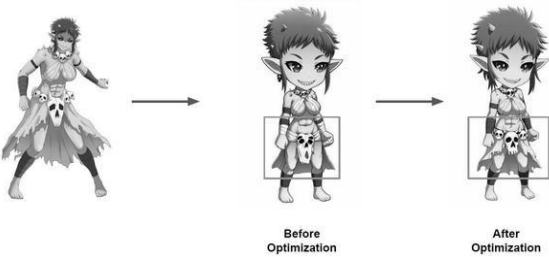
【図 9】



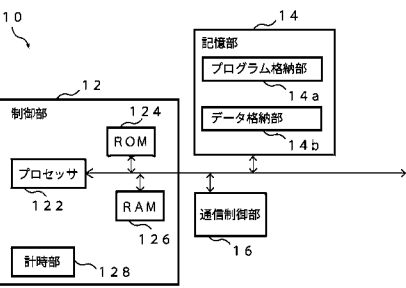
【図 1 2】



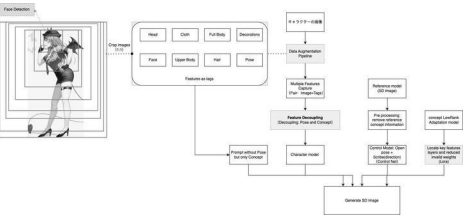
【図 1 0】



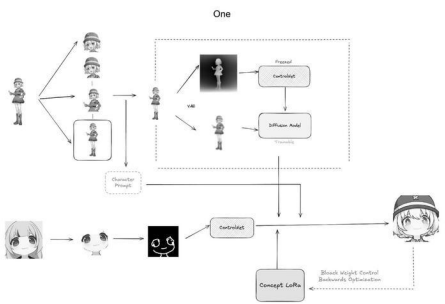
【図 1 3】



【図 1 1】



【図 14】



【図 15】



フロントページの続き

審査官 鈴木 肇

- (56)参考文献 特開 2 0 2 3 - 1 5 7 3 3 4 (J P , A)
国際公開第 2 0 2 4 / 1 8 5 0 5 4 (W O , A 1)
特開 2 0 2 1 - 1 0 5 9 5 0 (J P , A)
特開 2 0 2 1 - 0 4 7 7 1 1 (J P , A)
特開 2 0 2 3 - 1 6 9 0 4 8 (J P , A)
国際公開第 2 0 2 4 / 1 6 1 8 3 9 (W O , A 1)
Miao Hua, Jiawei Liu, Fei Ding, Wei Liu, Jie Wu, Qian He , DreamTuner: Single Image is
Enough for Subject-Driven Generation , arXiv, [online] , 2023年12月21日 , [2025年3月25日
検索], インターネット<URL: <https://arxiv.org/pdf/2312.13691v1>> , DOI: 10.48550/arXiv.23
12.13691
Stable Diffusion 『ControlNet』の使い方ガイド！導入・更新方法も , romptn Magazine, [onl
ine] , [2025年3月25日検索], インターネット<URL: <https://romptn.com/article/7868>>
澁谷美樹, 鈴木優 , 描いたイラストのデフォルメ具合を自動で調節するツール , インタラクショ
ン 2 0 1 8 論文集 [o n l i n e] , 日本 , 情報処理学会 , 2018年 , p994-997 , URL:[http://www
w.interaction-ipsj.org/proceedings/2018/data/pdf/3B33.pdf](http://www.w.interaction-ipsj.org/proceedings/2018/data/pdf/3B33.pdf)

(58)調査した分野(Int.Cl. , D B 名)

G 0 6 T 1 / 0 0 - 1 9 / 2 0