

(12) 特許協力条約に基づいて公開された国際出願

(19) 世界知的所有権機関
国際事務局

(43) 国際公開日
2019年3月7日(07.03.2019)



(10) 国際公開番号

WO 2019/045101 A1

- (51) 国際特許分類:
G06T 7/00 (2017.01) *G06T 7/11* (2017.01)
- (21) 国際出願番号: PCT/JP2018/032635
- (22) 国際出願日: 2018年9月3日(03.09.2018)
- (25) 国際出願の言語: 日本語
- (26) 国際公開の言語: 日本語
- (30) 優先権データ:
特願 2017-169632 2017年9月4日(04.09.2017) JP
- (71) 出願人: 国立大学法人 東京大学 (THE UNIVERSITY OF TOKYO) [JP/JP]; 〒1138654 東京都文京区本郷七丁目3番1号 Tokyo (JP).
- (72) 発明者: 相澤 清晴 (AIZAWA Kayoharu); 〒1138654 東京都文京区本郷七丁目3番1号 国立大学法人東京大学内 Tokyo (JP).

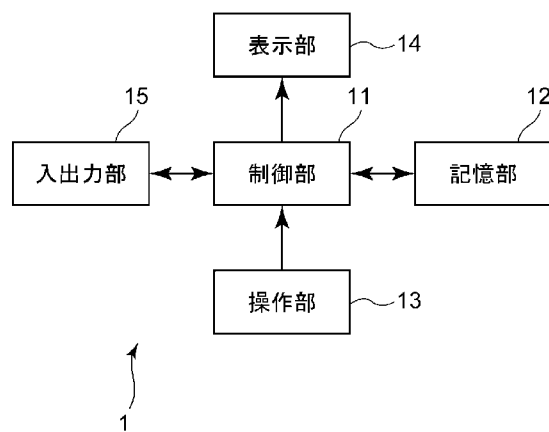
小川 徹 (OGAWA Toru); 〒1138654 東京都文京区本郷七丁目3番1号 国立大学法人東京大学内 Tokyo (JP).

(74) 代理人: 竹居 信利 (TAKEI Nobutoshi); 〒1600022 東京都新宿区新宿6丁目7-1 エルプリメント新宿308 Tokyo (JP).

(81) 指定国(表示のない限り、全ての種類の国内保護が可能): AE, AG, AL, AM, AO, AT, AU, AZ, BA, BB, BG, BH, BN, BR, BW, BY, BZ, CA, CH, CL, CN, CO, CR, CU, CZ, DE, DJ, DK, DM, DO, DZ, EC, EE, EG, ES, FI, GB, GD, GE, GH, GM, GT, HN, HR, HU, ID, IL, IN, IR, IS, JO, KE, KG, KH, KN, KP, KR, KW, KZ, LA, LC, LK, LR, LS, LU, LY, MA, MD, ME, MG, MK, MN, MW, MX, MY, MZ, NA, NG, NI, NO, NZ, OM, PA, PE, PG, PH, PL, PT, QA, RO, RS, RU, RW, SA, SC, SD, SE, SG, SK, SL,

(54) Title: IMAGE PROCESSING DEVICE AND PROGRAM

(54) 発明の名称: 画像処理装置及びプログラム



11 Control unit
12 Storage unit
13 Operation unit
14 Display unit
15 Input/output unit

(57) Abstract: An image processing device for accepting comics image data and generating, on the basis of the comics image data, information that specifies a frame portion, information that specifies a face portion, information that specifies a body portion, and information that specifies a text portion through use of a result obtained by machine learning so as to specify each of the frame portion, the face position, the body portion, and the text portion, respectively, the information being supplied to prescribed information processing.

(57) 要約: 漫画画像データを受け入れ、当該漫画画像データに基づいて、コマ部分、顔部分、身体部分、文字部分のそれぞれを特定するよう機械学習した結果を用い、コマ部分を特定する情報と、顔部分を特定する情報と、身体部分を特定する情報と、文字部分を特定する情報とをそれぞれ生成して、当該情報が所定の情報処理に供される画像処理装置である。

[続葉有]

SM, ST, SV, SY, TH, TJ, TM, TN, TR, TT, TZ, UA,
UG, US, UZ, VC, VN, ZA, ZM, ZW.

- (84) 指定国(表示のない限り、全ての種類の広域保護が可能): ARIPO (BW, GH, GM, KE, LR, LS, MW, MZ, NA, RW, SD, SL, ST, SZ, TZ, UG, ZM, ZW), ユーラシア (AM, AZ, BY, KG, KZ, RU, TJ, TM), ヨーロッパ (AL, AT, BE, BG, CH, CY, CZ, DE, DK, EE, ES, FI, FR, GB, GR, HR, HU, IE, IS, IT, LT, LU, LV, MC, MK, MT, NL, NO, PL, PT, RO, RS, SE, SI, SK, SM, TR), OAPI (BF, BJ, CF, CG, CI, CM, GA, GN, GQ, GW, KM, ML, MR, NE, SN, TD, TG).

添付公開書類:

- 一 国際調査報告 (条約第21条(3))

明 細 書

発明の名称： 画像処理装置及びプログラム

技術分野

[0001] 本発明は、画像処理装置及びプログラムに関する。

背景技術

[0002] 近年では、漫画の画像データを加工して、セリフ部分を抽出して他国語用に翻訳する技術や、コマごとに配列を変更して、スマートフォン等の画面に適した状態とする技術が考えられている。

[0003] このような処理を行うにあたり、従来から、色や文字認識処理の結果等を用いて、描かれた人物の部分やセリフ部分を特定する処理等が考えられている（非特許文献1）。

先行技術文献

非特許文献

[0004] 非特許文献1：Christophe Rigaud, et. al., Speech ballon and speaker as sociation for comics and manga understanding., Proceedings of the 13th International Conference on Document Analysis and Recognition, pp. 351-355, IEEE, 2015

発明の概要

発明が解決しようとする課題

[0005] 一方、近年では機械学習により画像中から物体を検出する技術が開発され、広く研究されている。しかしながら、従来の一般物体検出の処理では、画像中の特定の部分には一つの物体が含まれるとの前提で検出が行われるため、多数の検出対象が互いに重なりあっている場合については考慮されていない。

[0006] ところが漫画画像データにおいては、コマの内側に（ないしは複数のコマにまたがって）人物の身体や顔が描画され、また、これらの各部に重なり合わせてセリフの文字が配置されることが一般的である。従って、機械学習に

よる物体検出処理をそのまま適用したのでは、コマ、登場人物の身体、顔、セリフといった部分がそれぞれ十分な精度で検出できない。

[0007] 本発明は上記実情に鑑みて為されたもので、漫画画像データの処理において、機械学習処理を用いて、コマ、顔部分、身体部分、及び文字部分の認識精度を、従来のものに比べて向上できる画像処理装置及びプログラムを提供することを、その目的の一つとする。

課題を解決するための手段

[0008] 上記従来例の問題点を解決するための本発明は、画像処理装置であって、漫画画像データを受け入れる受入手段と、画像データから、当該画像データ内に描画された漫画のコマ部分を検出するよう機械学習された状態にあるフレーム検出手段と、画像データから、当該画像データ内に描画された顔部分を検出するよう機械学習された状態にある顔検出手段と、画像データから、当該画像データ内に描画された身体部分を検出するよう機械学習された状態にある身体検出手段と、画像データから、当該画像データ内に含まれる文字部分を検出するよう機械学習された状態にある文字検出手段と、前記受け入れた漫画画像データに基づいて、前記フレーム検出手段が検出したコマ部分を特定する情報と、前記顔検出手段が検出した、顔部分を特定する情報と、前記身体検出手段が検出した、身体部分を特定する情報と、前記文字検出手段が検出した、文字部分を特定する情報と、を生成する検出情報生成手段と、を含み、前記生成されたコマ部分を特定する情報と、顔部分を特定する情報と、身体部分を特定する情報と、文字部分を特定する情報とが所定の情報処理に供されることとしたものである。

発明の効果

[0009] 本発明によれば、機械学習処理を用いて、漫画画像データのうちからコマ、顔部分、身体部分、及び文字部分を認識する際の認識精度を、従来のものに比べて向上できる。

図面の簡単な説明

[0010] [図1]本発明の実施の形態に係る画像処理装置の構成例を表すブロック図であ

る。

[図2]本発明の実施の形態に係る画像処理装置の例を表す機能ブロック図である。

[図3]本発明の実施の形態に係る画像処理装置が処理の対象とする漫画画像データの概要例を表す説明図である。

[図4]本発明の実施の形態に係る画像処理装置の検出処理部の概要例を表す内部機能ブロック図である。

[図5]本発明の実施の形態に係る画像処理装置の検出処理部のもう一つの例を表す内部機能ブロック図である。

[図6]本発明の実施の形態に係る画像処理装置の検出処理部のもう一つの例による処理の概要例を表す説明図である。

[図7]本発明の実施の形態に係る画像処理装置の検出処理部の構成例を表す機能ブロック図である。

発明を実施するための形態

[0011] 本発明の実施の形態について図面を参照しながら説明する。本発明の実施の形態に係る画像処理装置1は、図1に例示するように、制御部11と、記憶部12と、操作部13と、表示部14と、入出力部15とを含んで構成されている。

[0012] 制御部11は、CPU等のプログラム制御デバイスであり、記憶部12に格納されたプログラムを実行して、漫画画像データを受け入れ、当該受け入れた漫画画像データに基づいて、コマ部分を特定する情報と、顔部分を特定する情報と、身体部分を特定する情報と、文字部分を特定する情報と、を生成する。本実施の形態の制御部11は、これらの各部分を特定する処理において、画像データから、当該画像データ内に描画された漫画のコマ部分を検出するよう機械学習された状態にあるフレーム検出器と、画像データから、当該画像データ内に描画された顔部分を検出するよう機械学習された状態にある顔検出器と、画像データから、当該画像データ内に描画された身体部分を検出するよう機械学習された状態にある身体検出器と、画像データから、

当該画像データ内に含まれる文字部分を検出するよう機械学習された状態にある文字検出器とを用いる。

[0013] またこの制御部 11 は、生成されたコマ部分を特定する情報と、顔部分を特定する情報と、身体部分を特定する情報と、文字部分を特定する情報とを用いて所定の情報処理を実行する。この情報処理としては、例えば、各部を表す画像を出力する処理や、文字部分を特定する情報により特定された範囲内の文字列に対する光学文字認識処理や、コマ部分ごとに画像データを分割する分割処理等がある。これらの制御部 11 の動作については、後に詳しく述べる。

[0014] 記憶部 12 は、メモリデバイス等であり、制御部 11 により実行されるプログラムを保持する。このプログラムは、コンピュータ可読かつ、非一時的な記録媒体に格納されて提供され、この記憶部 12 に格納されたものであってもよい。また、この記憶部 12 は、制御部 11 のワークメモリとしても動作する。

[0015] 操作部 13 は、マウスやキーボード等であり、利用者の指示操作を受け入れて制御部 11 に出力する。表示部 14 は、例えばディスプレイ等であり、制御部 11 から入力される指示に基づいて情報を表示出力する。

[0016] 入出力部 15 は、例えばネットワークインタフェース等であり、外部からデータ（画像データ等）を受信して、制御部 11 に出力する。またこの入出力部 15 は、制御部 11 から入力される指示に従って、データを外部の装置等に送出する。

[0017] 次に制御部 11 の動作について説明する。本実施の形態の制御部 11 は、記憶部 12 に格納されたプログラムを実行することで、機能的には、図 2 に例示するように、受入部 21 と、検出処理部 22 と、検出情報生成部 23 と、情報処理部 24 とを含んで構成される。また検出処理部 22 は、フレーム検出部 31 と、顔検出部 32 と、身体検出部 33 と、文字検出部 34 とを含む。

[0018] 受入部 21 は、漫画画像データを受け入れて検出処理部 22 に出力する。

ここで漫画画像データは、一般的には、顔部分（F）と身体部分（B）と文字部分（C）とが互いに重なり合って描画された画像データであり（図3）、少なくとも一つのコマ（M）を含む。また、この受入部21は、検出処理部22におけるニューラルネットワークを利用するため、漫画画像データを拡大または縮小して、ニューラルネットワークの入力に適したサイズにリサイズする。

[0019] 検出処理部22のフレーム検出部31は、画像データから、当該画像データ内に描画された漫画のコマ部分を検出するよう機械学習された状態にあるフレーム検出器を有する。具体的に、このフレーム検出部31が備えるフレーム検出器は、R-CNN (Regions with CNN features) (Girshick, Ross, et al. "Rich feature hierarchies for accurate object detection and semantic segmentation." Proceedings of the IEEE conference on computer vision and pattern recognition, 2014) や、Fast R-CNN (Girshick, Ross. "Fast r-cnn." Proceedings of the IEEE International Conference on Computer Vision, 2015)、Faster R-CNN (Ren, Shaoqing, et al. "Faster R-CNN: Towards real-time object detection with region proposal networks." Advances in neural information processing systems, 2015)、YOLO (You Only Look Once) (Redmon, Joseph, et al. "You only look once: Unified, real-time object detection." arXiv preprint arXiv:1506.02640 (2015))、あるいは、SSD (シングル・ショット・マルチボックス・ディテクタ; Single Shot MultiBox Detector) (Liu, Wei, et al. "SSD: Single Shot MultiBox Detector." arXiv preprint arXiv:1512.02325 (2015)) など、種々の方法で構成されたニューラルネットワークを採用して実現できる。

[0020] 図4にその概略を示すように、SSD等のニューラルネットワークを採用した検出器40は、ベースネットワーク部41と、分類器42とを含んで構成される。ここでベースネットワーク部41は、検出対象の候補が含まれる画像の範囲と、当該範囲内の画像の特徴量とを出力する。また分類器42は

、出力された画像の範囲に、検出対象（フレーム検出部 31 の場合、漫画画像データのコマを区分する枠線）が含まれるか否かを、出力された特徴量に基づいて判断する。

[0021] このような SSD 等を採用した検出器 40 は、検出対象の範囲（フレーム検出部 31 の場合、漫画画像データのコマを区分する枠線に外接する形状の範囲）を人為的に指定した画像データのサンプルを用いて機械学習させる。ここで機械学習の具体的方法や、検出器 40 の利用方法については、広く知られているので、ここでの詳しい説明を省略する。

[0022] 顔検出部 32 は、画像データから、当該画像データ内に描画されたキャラクターの顔部分を検出するよう機械学習された状態にある顔検出器を有する。この顔検出器も、フレーム検出部 31 が備えるフレーム検出器と同様、SSD 等、種々の方法で構成されたニューラルネットワークを採用して実現できる。この顔検出器は、検出対象の範囲である、漫画画像データに含まれるキャラクターの顔に外接する所定形状の範囲を人為的に指定した画像データのサンプルを用いて機械学習させる。

[0023] 身体検出部 33 は、画像データから、当該画像データ内に描画されたキャラクターの身体部分を検出するよう機械学習された状態にある身体検出器を有する。この身体検出器も、フレーム検出部 31 が備えるフレーム検出器と同様、SSD 等、種々の方法で構成されたニューラルネットワークを採用して実現できる。この身体検出器は、検出対象の範囲である、漫画画像データに含まれるキャラクターの身体に外接する所定形状の範囲を人為的に指定した画像データのサンプルを用いて機械学習させる。

[0024] 文字検出部 34 は、画像データから、当該画像データ内に描画された文字部分を検出するよう機械学習された状態にある文字検出器を有する。この文字検出器も、フレーム検出部 31 が備えるフレーム検出器と同様、SSD 等、種々の方法で構成されたニューラルネットワークを採用して実現できる。この文字検出器は、検出対象の範囲である、漫画画像データに含まれる文字部分に外接する所定形状の範囲を人為的に指定した画像データのサンプルを

用いて機械学習させる。

[0025] 検出情報生成部 23 は、受入部 21 が受け入れた漫画画像データについて、フレーム検出部 31 が検出したコマ部分を特定する情報と、顔検出部 32 が検出した、顔部分を特定する情報と、身体検出部 33 が検出した、身体部分を特定する情報と、文字検出部 34 が検出した、文字部分を特定する情報とを生成する。

[0026] 情報処理部 24 は、生成されたコマ部分を特定する情報と、顔部分を特定する情報と、身体部分を特定する情報と、文字部分を特定する情報とを用いて所定の情報処理を実行する。この情報処理としては、例えば、特定された文字部分の画像に対して光学的文字認識（OCR）を行い、その結果を出力する処理等がある。また、情報処理部 24 は、光学的文字認識の結果、得られた文字列を、機械翻訳処理により他言語に翻訳して出力してもよい。

[0027] 本実施の形態の一例は以上の構成を備え、次のように動作する。なお、以下の説明では、制御部 11 によるフレーム検出部 31，顔検出部 32，身体検出部 33，及び文字検出部 34 は、SSD を採用し、それぞれ、予め画像データから、当該画像データ内に描画された漫画のコマ部分、顔部分、身体部分、及び文字部分を検出するよう機械学習した状態にあるものとする。

[0028] 画像処理装置 1 は、利用者から入力される漫画の画像データ（機械学習のサンプルに含まれないもの）を処理の対象として、当該処理対象の画像データに対して並列的に、フレーム検出器と、顔検出器と、身体検出器と、文字検出器とにより、コマ部分、キャラクタの顔部分、身体部分、及び文字部分をそれぞれ検出して、それぞれ検出した画像の範囲を特定する情報を得る。

[0029] そして画像処理装置 1 は、コマ部分を特定する情報と、顔部分を特定する情報と、身体部分を特定する情報と、文字部分を特定する情報とを用いて所定の情報処理、例えば特定された文字部分の画像に対して光学的文字認識（OCR）を行い、当該光学的文字認識の結果、得られた文字列を、機械翻訳処理により他言語に翻訳して出力する。

[0030] [ベースネットワークを共用する例]

またここまでの説明では、フレーム検出部 3 1，顔検出部 3 2，身体検出部 3 3，及び文字検出部 3 4 は、それぞれ独立したベースネットワークと、検出器を備えるものとしたが、本実施の形態はこの例に限られない。例えば一つのベースネットワークをフレーム検出部 3 1，顔検出部 3 2，身体検出部 3 3，文字検出部 3 4 が共用してもよい。

[0031] すなわちこの例では、フレーム検出部 3 1，顔検出部 3 2，身体検出部 3 3，及び文字検出部 3 4 は、図 5 に例示するように、それぞれに共通して、検出対象の候補となる画像の範囲と、当該範囲内の画像の特徴量とを機械学習した状態にあり、処理対象となった画像データに基づき、検出対象の候補となる画像の範囲と、当該範囲内の画像の特徴量とを出力するベースネットワーク部 4 1' と、フレーム検出部 3 1，顔検出部 3 2，身体検出部 3 3，及び文字検出部 3 4 のそれぞれに対応して、独立して設けられる分類器 4 2 a，4 2 b，4 2 c，4 2 d とを備える。

[0032] なお、この例でも、ベースネットワーク部 4 1' 及び分類器 4 2 a，b，c，d は、SSD に基づくニューラルネットワークとしてよいが、次の点で SSD を変形して用いる。すなわち一般的な SSD の出力段では、物体を検出する領域の候補（アンカーボックス）が予め複数定められており（複数のアンカーボックスの集合をアンカーセットと呼ぶ）、当該複数の領域の候補のうちから、対象となる物体が含まれる領域を特定する。

[0033] 本実施の形態のここでの例では、出力段より前のネットワーク（ベースネットワーク部 4 1'）は 1 つとするが、出力段において、アンカーセット（各アンカーセットには、例えば 8 7 3 2 個のアンカーボックスが含まれる）をフレーム検出部 3 1，顔検出部 3 2，身体検出部 3 3，及び文字検出部 3 4 のそれぞれに対応して、4 つ複製して、各分類器 4 2 a，b，c，d として用いる。すなわちこの例における SSD では、各アンカーボックスについて、当該アンカーボックス内で検出した物体の領域との位置特定誤差（左上座標の情報と幅及び高さの情報とからなる 4 次元の情報）と、物体が含まれ得るとされる確信度（ここではシグモイド関数により正規化しておく）との

合計 5 次元の情報を出力するが、アンカーボックスが複製した 4 つの同じアンカーセット中のアンカーボックスのうち、第 1 のアンカーセット A 1 中のアンカーボックスについては画像データのうちコマ部分を機械学習させた状態とする。また第 2 のアンカーセット A 2 中のアンカーボックスについては、画像データのうち顔部分を機械学習させた状態とし、第 3 のアンカーセット A 3 中のアンカーボックスについては、画像データのうち身体部分を機械学習させた状態とし、第 4 のアンカーセット A 4 中のアンカーボックスについては、画像データのうち文字部分を機械学習させた状態とする（図 6）。

[0034] 具体的には、出力段が出力する第 1 のアンカーセット中のアンカーボックスの情報については、学習用のサンプルを入力したときに、コマの枠線に外接する矩形が推定されることとなるよう、出力段から誤差を逆伝播して、分類器 4 2 a 及びベースネットワーク部 4 1' のパラメータを更新する。

[0035] 同様に、出力段が出力する第 2 のアンカーセット中のアンカーボックスの情報については、学習用のサンプルを入力したときに、キャラクタの顔部分に外接する矩形が推定されることとなるよう、出力段から誤差を逆伝播して、分類器 4 2 b 及びベースネットワーク部 4 1' のパラメータを更新する。また出力段が出力する第 3 のアンカーセット中のアンカーボックスの情報については、学習用のサンプルを入力したときに、キャラクタの身体に外接する矩形が推定されることとなるよう、出力段から誤差を逆伝播して、分類器 4 2 c 及びベースネットワーク部 4 1' のパラメータを更新する。さらに出力段が出力する第 4 のアンカーセット中のアンカーボックスの情報については、学習用のサンプルを入力したときに、文字に外接する矩形が推定されることとなるよう、出力段から誤差を逆伝播して、分類器 4 2 d 及びベースネットワーク部 4 1' のパラメータを更新する。

[0036] なお、第 i ($i = 1, 2, 3, 4$) のアンカーセット中の a 番目のアンカーボックス ($a = 1, 2, \dots, 8732$) に対する、 m 番目のサンプル（ミニバッチ学習を行うこととして、ミニバッチサイズを M とすると、 $m = 1, 2, \dots, M$) の割り当て $s(m, i, a)$ とその重なり $J(m, i, a)$ とを

次のように定義する。

[数1]

$$s(m, i, a) = \operatorname{argmax}_{g, c=t(m, g)} \operatorname{Jaccard}(B_a, B_{m, g})$$

$$J(m, i, a) = \operatorname{Jaccard}(B_a, B_{m, s(m, i, a)})$$

ここで、 g は 1 以上、 $G(m)$ 以下の整数であり、 $G(m)$ は、上記 m 番目のサンプルに含まれる正解の個数であり、 $t(m, g)$ 、及び $B(m, g)$ は上記 m 番目のサンプルの g 番目の正解のクラス（コマ、顔、身体、文字のいずれであるかを表す情報）と、外接矩形とを表す。

[0037] そして損失関数（Loss関数） $L(z)$ を、位置特定誤差 $L_{loc}(m, z)$ と、確信度 $L_{conf}(m, z)$ との和として次のように設定する。

[数2]

$$L(z) = \frac{1}{\sum_{m=1, \dots, M} |A(m, pos)|} \sum_{m=1, \dots, M} (L_{loc}(m, z) + L_{conf}(m, z))$$

ここで、 z はニューラルネットワークの出力を表し、 $A(m, pos)$ は、 m 番目のサンプルについてオブジェクトが割り当てられたアンカーボックスの添字集合であり、具体的には、

[数3]

$$A(m, pos) = \{(c, a) | J(m, i, a) \geq 0.5\}$$

などとしておく。

[0038] なお、 $L_{loc}(m, z)$ 及び、 $L_{conf}(m, z)$ は、次のように定義しておく。

[数4]

$$L_{loc}(m, z) = \sum_{A(m, pos)} \text{huber}(z(i, a, 2 \dots 5), \log \Delta B)$$

ここで、 $\Delta B = B(m, s(m, i, a)) - B(a)$

$$L_{conf}(m, z) = \sum_{A(m, pos) \cup A(m, neg)} l(m, i, a, z)$$

ここで、 $l(m, i, a, z) = \text{sigmoid}(z(i, a, 1), J(m, i, a) \geq 0.5)$

なお、 $A(m, neg)$ は、ハードネガティブ (hard negative) の集合であって、オブジェクトに割り当てられていないアンカーボックスのうち、 $l(m, i, a, z)$ が大きい順に上位 $k | A(m, pos) |$ 個を選択して得られる (ハードネガティブマイニングと呼ばれる方法であるので、ここでの詳細な説明を省略する)。また、 $\text{huber}()$ は、ヒューバー関数 (huber関数) である。この関数についても広く知られているのでここでの詳細な説明を省略する。

[0039] 以上のようにフレーム検出部 3 1, 顔検出部 3 2, 身体検出部 3 3, 及び文字検出部 3 4 が構成された本実施の形態の画像処理装置 1 においても、フレーム検出部 3 1, 顔検出部 3 2, 身体検出部 3 3, 及び文字検出部 3 4 のそれぞれに対応する分類器 4 2 a, b, c, d が、それぞれコマ部分を特定する情報と、顔部分を特定する情報と、身体部分を特定する情報と、文字部分を特定する情報とを推定するので、これらを用いて所定の情報処理を実行する。

[0040] [画像データを分割する例]

また本実施の形態の制御部 1 1 は、受入部 2 1 の動作として、入力された漫画画像データの全体を、拡大または縮小して、ニューラルネットワークの入力に適したサイズにリサイズするのではなく、入力された漫画画像データを、所定の条件に基づいて複数の分割部分に分割し、当該分割して得られた分割部分 (部分的な漫画画像データ、以下、部分画像データと呼ぶ) を、ニューラルネットワークの入力に適したサイズにリサイズして、検出処理部 2

2に出力してもよい。

[0041] ここで上記所定の条件は、例えば、元の漫画画像データ（幅 w 、高さ h ）を、 2×2 個に分割（それぞれが幅 $w/2$ 、高さ $h/2$ となるような、重なり合わない4つの領域に分割）するとの条件であってもよい。またこの条件は、漫画画像データの内容に基づき、例えば、白色（背景色）が連続する部分で分割するとの条件であってもよい。さらに、この所定の条件は、コマごとに分割するとの条件であってもよい。

[0042] 本実施の形態のこの例では、検出処理部22のフレーム検出部31、顔検出部32、身体検出部33、及び文字検出部34は、部分画像データのそれぞれからコマ部分（コマごとに分割しない場合）、顔部分、身体部分及び文字部分を検出する。なお、この例では、機械学習の処理も、分割して得られた部分画像データを用いて行うこととしてもよい。

[0043] そして検出情報生成部23は、部分画像データごとにフレーム検出部31が検出したコマ部分を特定する情報と、顔検出部32が検出した、顔部分を特定する情報と、身体検出部33が検出した、身体部分を特定する情報と、文字検出部34が検出した、文字部分を特定する情報とを生成し、これらをまとめて元の漫画画像データにおける、コマ部分、顔部分、身体部分、及び文字部分のそれぞれを特定する情報を生成する。

[0044] 情報処理部24は、検出情報生成部23により生成されたコマ部分を特定する情報と、顔部分を特定する情報と、身体部分を特定する情報と、文字部分を特定する情報とを用いて所定の情報処理を実行する。

[0045] また、上述の通り、部分画像データに分割する所定の条件として、コマごとに分割するとの条件であってもよい。この場合、制御部11は、検出処理部22のフレーム検出部31としての動作によりコマ部分を検出し、当該検出したコマ部分ごとに分割して部分画像データを生成することとしてもよい。

[0046] すなわち本実施の形態のこの例では、図7に例示するように、フレーム検出部31が出力する、コマ部分（コマを区分する枠線）に外接する多角形（

または円弧等の曲線を少なくとも一部に含んでもよい)を特定する情報を、顔検出部32, 身体検出部33, 文字検出部34に出力する。そして、顔検出部32, 身体検出部33, 文字検出部34のそれぞれが、フレーム検出部31が出力する情報で特定されるコマ部分ごとに、各コマ部分をそれぞれ部分画像データとして、部分画像データのそれぞれから顔部分、身体部分及び文字部分を検出する。なお、この例でも、顔検出部32, 身体検出部33, 文字検出部34に係る機械学習の処理は、分割して得られた部分画像データを用いて行うこととしてもよい。

[0047] そしてこの例でも、検出情報生成部23は、フレーム検出部31が検出したコマ部分ごとに、顔検出部32が検出した、顔部分を特定する情報と、身体検出部33が検出した、身体部分を特定する情報と、文字検出部34が検出した、文字部分を特定する情報とを生成し、これらをまとめて元の漫画画像データにおける、コマ部分、顔部分、身体部分、及び文字部分のそれぞれを特定する情報を生成する。

[0048] 情報処理部24は、検出情報生成部23により生成されたコマ部分を特定する情報と、顔部分を特定する情報と、身体部分を特定する情報と、文字部分を特定する情報とを用いて所定の情報処理を実行する。

[0049] [検出結果の合成]

また処理の対象とする画像データを分割する場合、制御部11は、分割前の画像データについても、ニューラルネットワークの入力に適したサイズにリサイズして、検出処理部22としての動作を行ってもよい。すなわち、制御部11は、フレーム検出部31, 顔検出部32, 身体検出部33, 及び文字検出部34の動作として、分割前の画像データのそれぞれからコマ部分、顔部分、身体部分及び文字部分を検出する。

[0050] そして制御部11は、ここで検出したコマ部分、顔部分、身体部分及び文字部分を特定する情報を記憶しておき、さらに処理の対象とする画像データを、所定の条件に基づいて複数の分割部分に分割し、当該分割して得られた部分画像データごとに、ニューラルネットワークの入力に適したサイズにリ

サイズして、検出処理部 22 としての動作を行う。

[0051] この例によると、分割前の画像データについて検出されたコマ部分、顔部分、身体部分及び文字部分を特定する情報と、分割後に得られた部分画像データごとの顔部分、身体部分及び文字部分を特定する情報とが得られることとなる。

[0052] そして制御部 11 は、分割前の画像データから検出されたコマ部分、顔部分、身体部分及び文字部分を特定する情報と、分割後の部分画像データのそれぞれから検出された顔部分、身体部分及び文字部分を特定する情報とを用い、分割前、または分割後のいずれか少なくとも一方から顔部分、身体部分及び文字部分が検出されたならば、検出情報生成部 23 は、当該少なくとも一方から検出した顔部分、身体部分及び文字部分を特定する情報を生成して出力する（各部分の検出結果をそれぞれ統合して出力する）。

[0053] この例では、いわば、分割前の画像データから検出した顔部分、身体部分、文字部分のそれぞれと、分割後の画像データから検出した顔部分、身体部分、文字部分のそれぞれとの論理和が、処理対象となった画像データから検出した顔部分、身体部分、文字部分として、当該処理対象となった画像データから検出した顔部分、身体部分、文字部分を特定する情報が出力される。

[0054] なお、コマ部分は、顔部分、身体部分、または文字部分よりも一般的に高い精度で検出できるため、分割前の画像データ（または分割後の画像データであってもよい）のいずれか一方のみから検出すれば十分と考えられるが、制御部 11 は、コマ部分についても、分割前の画像データまたは分割後の画像データの少なくともいずれかから検出した場合に、当該コマ部分を特定する情報を出力するようにしてもよい。

[0055] また、このように、いずれかから検出された各部分（コマ部分、顔部分、身体部分、文字部分のそれぞれ）の情報を出力する場合は、重複している部分の情報については、重複を除いて出力する。

[0056] またここでは分割前の画像データと、分割後の画像データとのいずれかから検出されたコマ部分、顔部分、身体部分及び文字部分を特定する情報を出

力することとした。つまり、例えば処理対象の画像データが漫画の 1 ページ分の画像データである場合、ページ全体で検出したものと、部分ごとに分割した分割部分ごとに検出したものとの「OR（論理和）」をとることとした。しかしながら、本実施の形態のこの例は、これに限られず、分割前の画像データと、分割後の画像データとの双方から共通して検出されたコマ部分、顔部分、身体部分及び文字部分を特定する情報を出力してもよい（つまり、例えば処理対象の画像データが漫画の 1 ページ分の画像データである場合、ページ全体で検出したものと、部分ごとに分割した分割部分ごとに検出したものとの「AND（論理積）」をとってもよい）。

[0057] さらに、ここでは分割の態様を一種類としたが、複数種類の分割態様で分割して得た複数種類の部分画像データを生成してもよい。例えば、コマごとに分割して得た部分画像データと、 2×2 の 4 分割した部分画像データと…といったように複数種類の態様で分割して得られた部分画像データ（さらに分割前の画像データを加えてもよい）のいずれか少なくとも一つから（あるいはそれぞれから共通して）検出されたコマ部分、顔部分、身体部分及び文字部分を特定する情報を出力することとしてもよい。この場合も、重複が生じる場合は、重複を除いて出力する。

[0058] [実施形態の効果]

このように本実施の形態によれば、漫画画像データ内で互いに重なり合い、または包含関係となるコマ部分、キャラクタの顔部分、身体部分、及び文字部分の各部に対応してそれぞれ独立した検出器（または分類器）により、それぞれ検出を行うので、従来の機械学習を利用した検出に比べ、検出の精度を向上できる。

符号の説明

[0059] 1 画像処理装置、11 制御部、12 記憶部、13 操作部、14 表示部、15 入出力部、21 受入部、22 検出処理部、23 検出情報生成部、24 情報処理部、31 フレーム検出部、32 顔検出部、33 身体検出部、34 文字検出部、40 検出器、41, 41' ベース

ネットワーク部、42 分類器。

請求の範囲

[請求項1]

漫画画像データを受け入れる受入手段と、
画像データから、当該画像データ内に描画された漫画のコマ部分を
検出するよう機械学習された状態にあるフレーム検出手段と、
画像データから、当該画像データ内に描画された顔部分を検出する
よう機械学習された状態にある顔検出手段と、
画像データから、当該画像データ内に描画された身体部分を検出す
るよう機械学習された状態にある身体検出手段と、
画像データから、当該画像データ内に含まれる文字部分を検出する
よう機械学習された状態にある文字検出手段と、
前記受け入れた漫画画像データに基づいて、前記フレーム検出手段
が検出したコマ部分を特定する情報と、前記顔検出手段が検出した、
顔部分を特定する情報と、前記身体検出手段が検出した、身体部分を
特定する情報と、前記文字検出手段が検出した、文字部分を特定する
情報と、を生成する検出情報生成手段と、
を含み、
前記生成されたコマ部分を特定する情報と、顔部分を特定する情報
と、身体部分を特定する情報と、文字部分を特定する情報とが所定の
情報処理に供される画像処理装置。

[請求項2]

請求項1記載の画像処理装置であって、
前記フレーム検出手段と、顔検出手段と、身体検出手段と、文字検
出手段とは、それぞれに共通して、検出対象の候補となる画像の範囲
と、当該範囲内の画像の特徴量とを機械学習した状態にあり、処理対
象となった画像データに基づき、検出対象の候補となる画像の範囲と
、当該範囲内の画像の特徴量とを出力するベースネットワーク部と、
前記フレーム検出手段と、顔検出手段と、身体検出手段と、文字検
出手段とのそれぞれに対応して設けられる分類器であって、前記画像
の特徴量に基づき、対応する画像の範囲内に含まれる画像が、それぞ

れコマ部分、顔部分、身体部分、文字部分であるか否かを分類する分類器とを含む画像処理装置。

[請求項3] 請求項 1 または 2 に記載の画像処理装置であって、

前記検出情報生成手段は、前記受け入れた漫画画像データを、所定の条件に基づいて複数の分割部分に分割して得られた分割部分のそれぞれのうちから、前記顔検出手段と、身体検出手段と、文字検出手段とにより、分割部分ごとに、顔部分を特定する情報と、身体部分を特定する情報と、文字部分を特定する情報とを生成する画像処理装置。

[請求項4] 請求項 3 に記載の画像処理装置であって、

前記検出情報生成手段は、フレーム検出手段が検出したコマ部分を特定する情報を用いて、当該特定されたコマ部分のそれぞれを前記分割部分として、当該分割部分のそれぞれのうちから、前記顔検出手段と、身体検出手段と、文字検出手段とにより、分割部分であるコマ部分ごとに、顔部分を特定する情報と、身体部分を特定する情報と、文字部分を特定する情報とを生成する画像処理装置。

[請求項5] 請求項 3 または 4 に記載の画像処理装置であって、

前記検出情報生成手段は、前記受け入れた漫画画像データを、分割する前の画像データから前記顔検出手段と、身体検出手段と、文字検出手段とにより検出した顔部分を特定する情報と、身体部分を特定する情報と、文字部分を特定する情報とを生成し、

当該生成した情報と、前記分割部分のそれぞれのうちから検出した、顔部分を特定する情報と、身体部分を特定する情報と、文字部分を特定する情報とを統合して出力する画像処理装置。

[請求項6] 請求項 1 から 5 のいずれか一項に記載の画像処理装置であって、

前記フレーム検出手段と、顔検出手段と、身体検出手段と、文字検出手段とは、シングル・ショット・マルチボックス・ディテクタ（SSD）を用いて構成される画像処理装置。

[請求項7] コンピュータを、

漫画画像データを受け入れる受入手段と、

画像データから、当該画像データ内に描画された漫画のコマ部分を
検出するよう機械学習された状態にあるフレーム検出手段と、

画像データから、当該画像データ内に描画された顔部分を検出する
よう機械学習された状態にある顔検出手段と、

画像データから、当該画像データ内に描画された身体部分を検出する
よう機械学習された状態にある身体検出手段と、

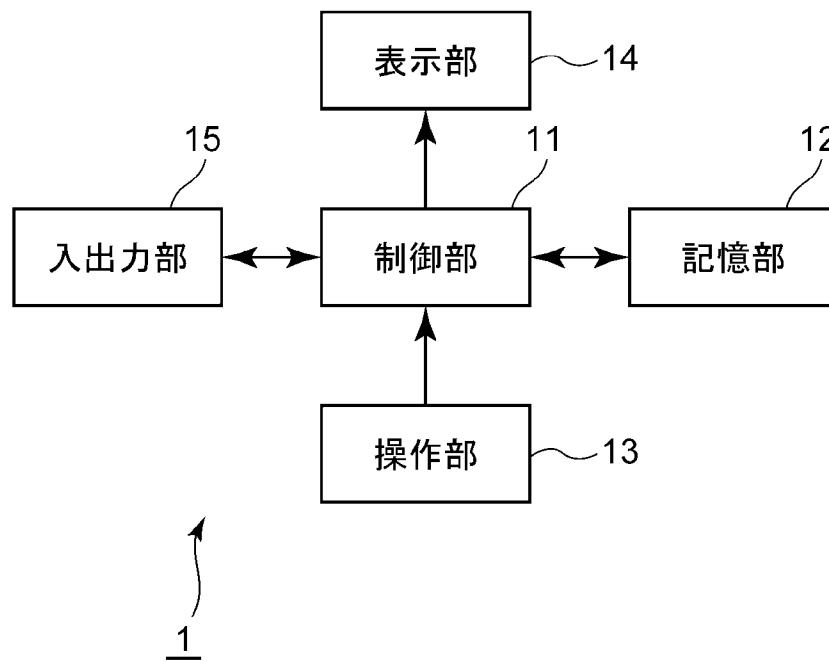
画像データから、当該画像データ内に含まれる文字部分を検出する
よう機械学習された状態にある文字検出手段と、

前記受け入れた漫画画像データに基づいて、前記フレーム検出手段
が検出したコマ部分を特定する情報と、前記顔検出手段が検出した、
顔部分を特定する情報と、前記身体検出手段が検出した、身体部分を
特定する情報と、前記文字検出手段が検出した、文字部分を特定する
情報と、を生成する検出情報生成手段と、

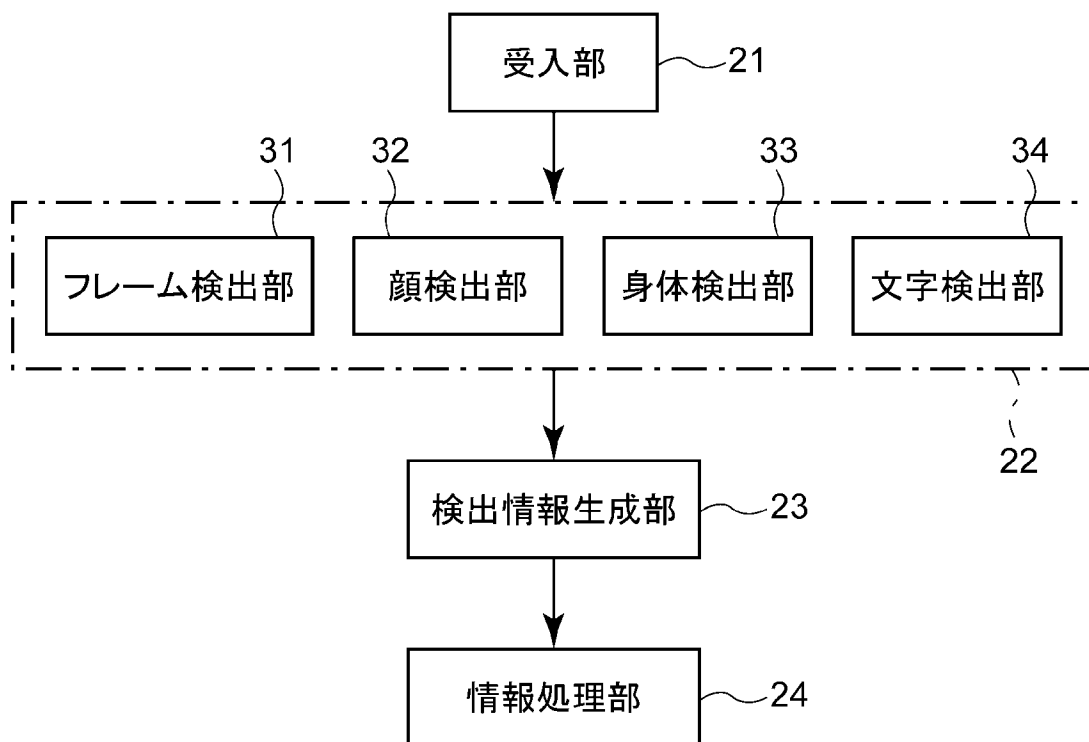
として機能させ、

前記生成されたコマ部分を特定する情報と、顔部分を特定する情報
と、身体部分を特定する情報と、文字部分を特定する情報とが所定の
情報処理に供されるプログラム。

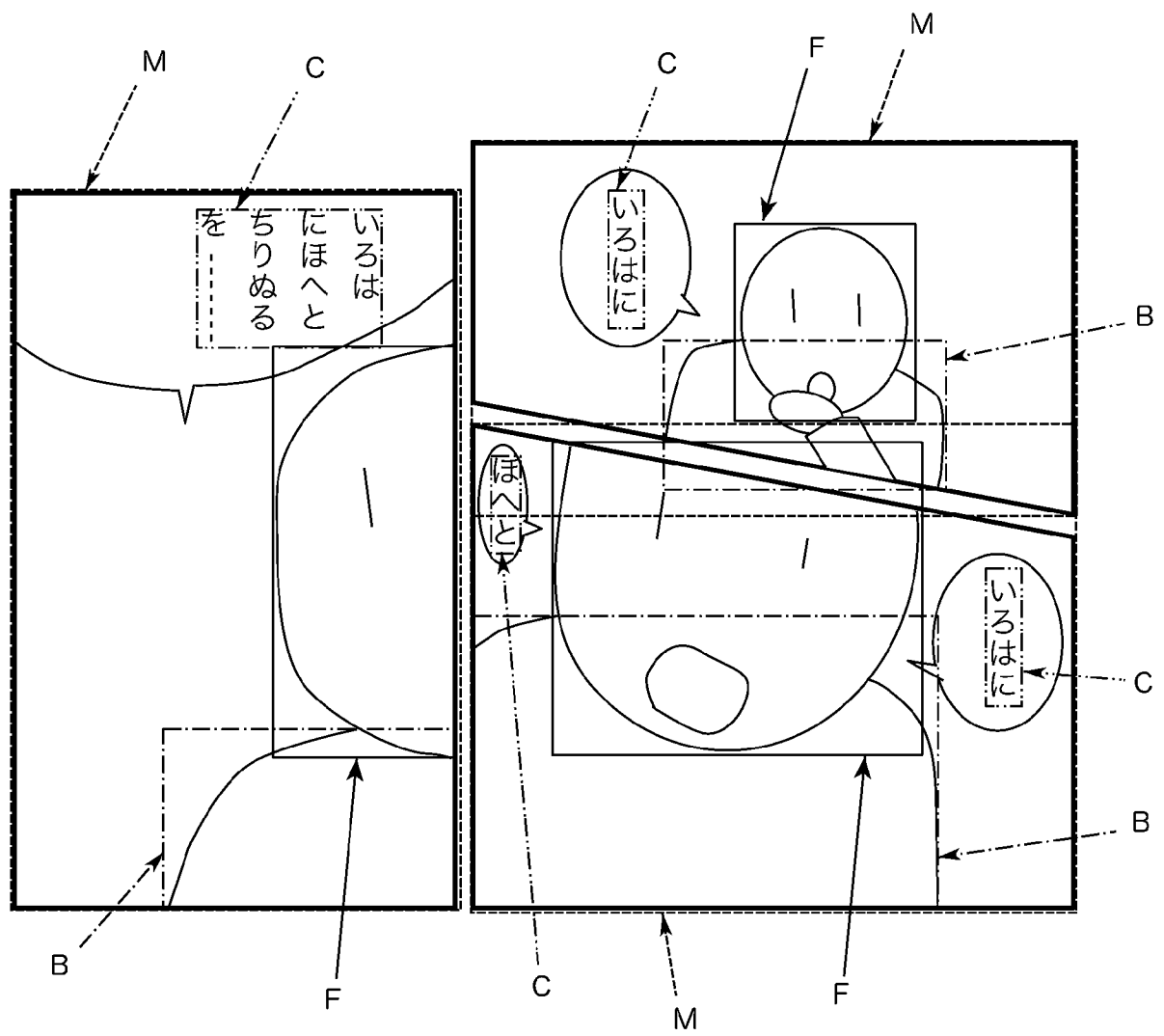
[図1]



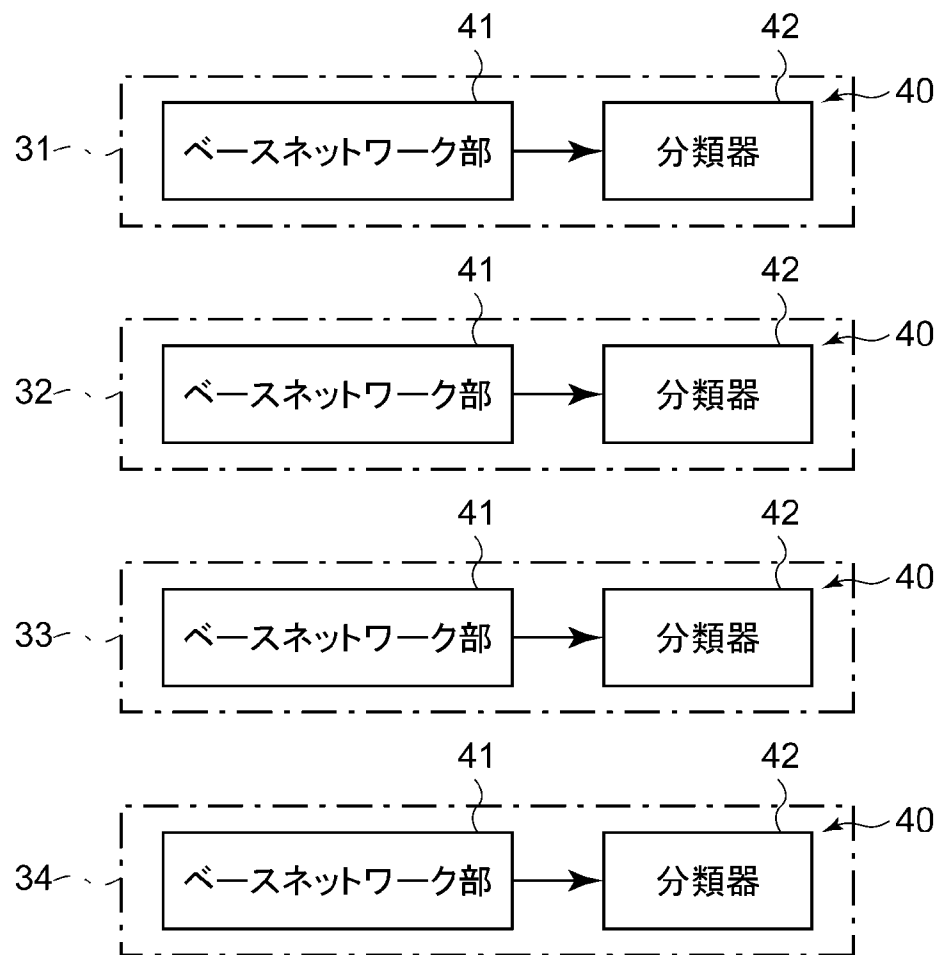
[図2]



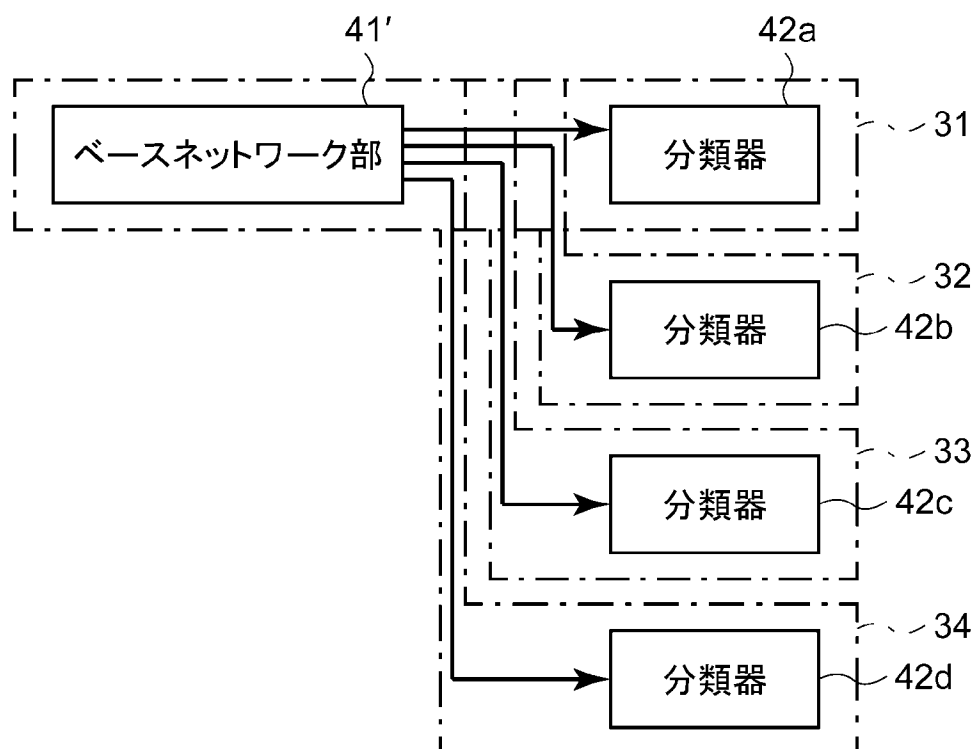
[図3]



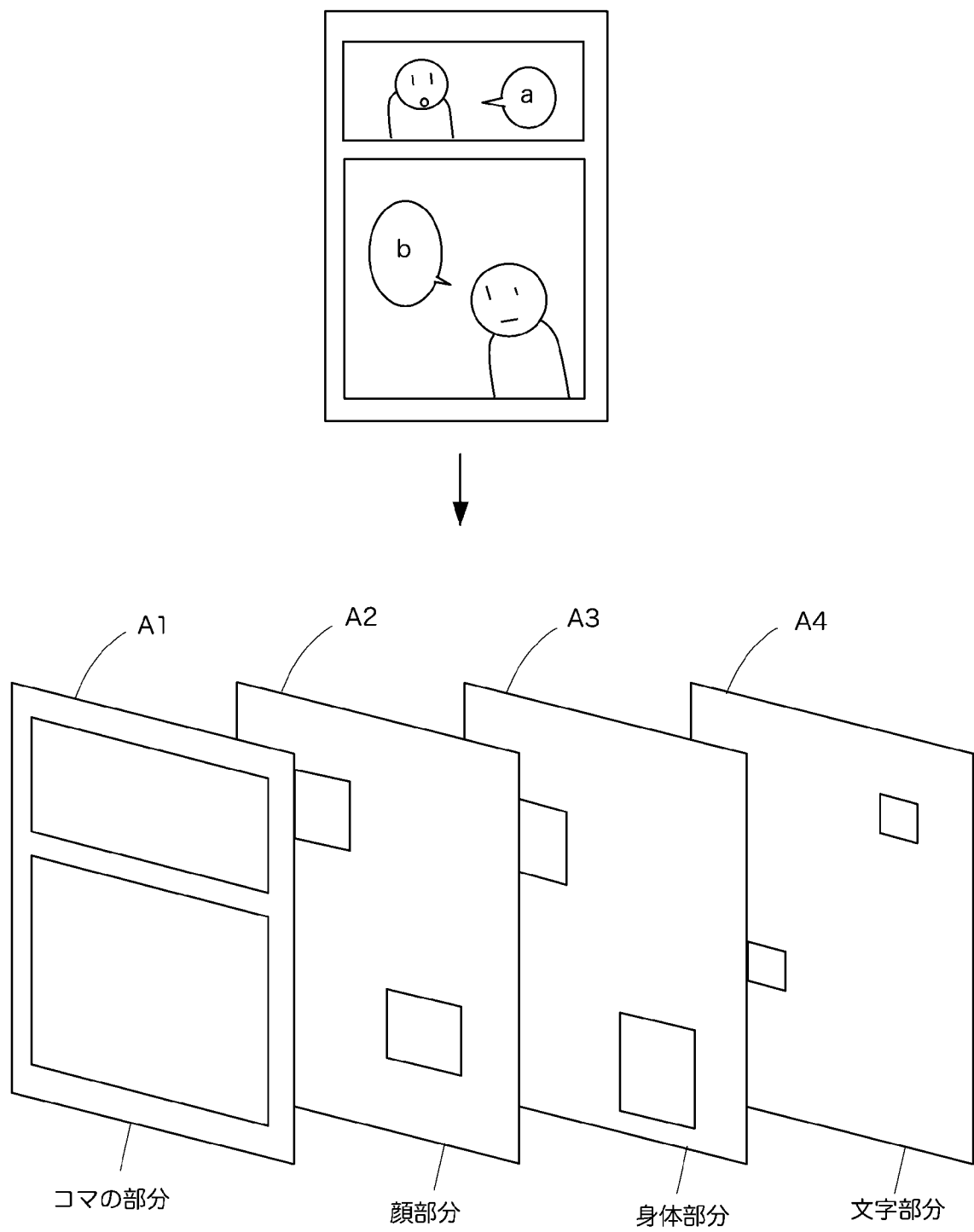
[図4]



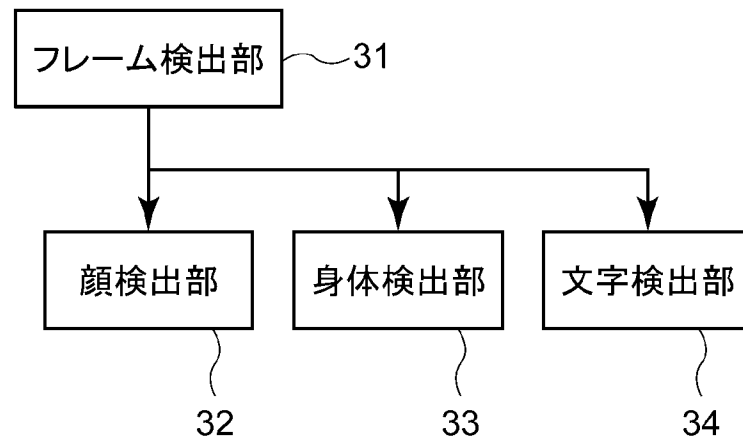
[図5]



[図6]



[図7]



INTERNATIONAL SEARCH REPORT

International application No.

PCT/JP2018/032635

A. CLASSIFICATION OF SUBJECT MATTER Int. Cl. G06T7/00 (2017.01) i, G06T7/11 (2017.01) i According to International Patent Classification (IPC) or to both national classification and IPC		
B. FIELDS SEARCHED Minimum documentation searched (classification system followed by classification symbols) Int. Cl. G06T7/00, G06T7/11 Documentation searched other than minimum documentation to the extent that such documents are included in the fields searched Published examined utility model applications of Japan 1922-1996 Published unexamined utility model applications of Japan 1971-2018 Registered utility model specifications of Japan 1996-2018 Published registered utility model applications of Japan 1994-2018 Electronic data base consulted during the international search (name of data base and, where practicable, search terms used)		
C. DOCUMENTS CONSIDERED TO BE RELEVANT		
Category*	Citation of document, with indication, where appropriate, of the relevant passages	Relevant to claim No.
Y A	柳澤秀彰, 渡辺裕, Faster R-CNN を用いたマンガの構造解析, 第 31 回 画像符号化シンポジウム 第 21 回 映像メディア処理シンポジウム, 30 November 2016 (intake date), pp. 80-81 (YANAGISAWA, Hideaki, WATANABE, Hiroshi. Structural Analysis of Comic Images using Faster R-CNN. PCSJ/IMPS 2016)	1-4, 6, 7 5
Y A	藤本東, 小川徹, 山本和慶, 松井勇佑, 山崎俊彦, 相澤清晴, 学術用アノテーション付き漫画画像データセットの構築と解析, 電子情報通信学会技術研究報告, 15 March 2017 (intake date), vol. 116, no. 64, pp. 35-40 (FUJIMOTO, Azuma, OGAWA, Toru, YAMAMOTO, Kazuyoshi, MATSUI, Yusuke, YAMASAKI, Toshihiko, AIZAWA, Kiyoharu. Creation and Analysis of Academic Manga Image Dataset with Annotations. IEICE Technical Report)	1-4, 6, 7 5
☒ Further documents are listed in the continuation of Box C. ☐ See patent family annex.		
* Special categories of cited documents: "A" document defining the general state of the art which is not considered to be of particular relevance "E" earlier application or patent but published on or after the international filing date "L" document which may throw doubts on priority claim(s) or which is cited to establish the publication date of another citation or other special reason (as specified) "O" document referring to an oral disclosure, use, exhibition or other means "P" document published prior to the international filing date but later than the priority date claimed "T" later document published after the international filing date or priority date and not in conflict with the application but cited to understand the principle or theory underlying the invention "X" document of particular relevance; the claimed invention cannot be considered novel or cannot be considered to involve an inventive step when the document is taken alone "Y" document of particular relevance; the claimed invention cannot be considered to involve an inventive step when the document is combined with one or more other such documents, such combination being obvious to a person skilled in the art "&" document member of the same patent family		
Date of the actual completion of the international search 10.10.2018		Date of mailing of the international search report 23.10.2018
Name and mailing address of the ISA/ Japan Patent Office 3-4-3, Kasumigaseki, Chiyoda-ku, Tokyo 100-8915, Japan		Authorized officer Telephone No.

INTERNATIONAL SEARCH REPORT

International application No.

PCT/JP2018/032635

C (Continuation). DOCUMENTS CONSIDERED TO BE RELEVANT

Category*	Citation of document, with indication, where appropriate, of the relevant passages	Relevant to claim No.
Y A	JP 2011-238043 A (KDDI CORP.) 24 November 2011, paragraphs [0031], [0035] (Family: none)	3, 4, 6 1, 2, 5, 7
Y A	JP 2002-229766 A (EASTMAN KODAK CO.) 16 August 2002, paragraph [0023] & US 2002/0065088 A1, paragraph [0031] & EP 1211631 A1 & DE 60126896 T2 & FR 2817429 A1 & AT 355567 T	3, 4, 6 1, 2, 5, 7
Y A	福井宏, 山下隆義, 山内悠嗣, 藤吉弘亘, Deep Learning を用いた歩行者検出の研究動向, 電子情報通信学会技術研究報告, 17 January 2017 (intake date), vol. 116, no. 366, pp. 37-46, non-official translation (FUKUI, Hiroshi, YAMASHITA, Takayoshi, YAMAUCHI, Yuji, FUJIYOSHI, Hironobu. Research Trends in Pedestrian Detection Using Deep Learning. IEICE Technical Report)	6 1-5, 7

A. 発明の属する分野の分類（国際特許分類（I P C））

Int.Cl. G06T7/00(2017.01)i, G06T7/11(2017.01)i

B. 調査を行った分野

調査を行った最小限資料（国際特許分類（I P C））

Int.Cl. G06T7/00, G06T7/11

最小限資料以外の資料で調査を行った分野に含まれるもの

日本国実用新案公報	1922-1996年
日本国公開実用新案公報	1971-2018年
日本国実用新案登録公報	1996-2018年
日本国登録実用新案公報	1994-2018年

国際調査で利用した電子データベース（データベースの名称、調査に使用した用語）

C. 関連すると認められる文献

引用文献の カテゴリー*	引用文献名 及び一部の箇所が関連するときは、その関連する箇所の表示	関連する 請求項の番号
Y A	柳澤 秀彰、渡辺 裕, F a s t e r R-CNN を用いたマンガの構造解析, 第31回 画像符号化シンポジウム 第21回 映像メディア処理 シンポジウム, 2016.11.30 (受入日), p. 80-81	1-4, 6, 7 5

☒ C欄の続きにも文献が列挙されている。☐ パテントファミリーに関する別紙を参照。

* 引用文献のカテゴリー

「A」特に関連のある文献ではなく、一般的技術水準を示すもの
「E」国際出願日前の出願または特許であるが、国際出願日以後に公表されたもの
「L」優先権主張に疑義を提起する文献又は他の文献の発行日若しくは他の特別な理由を確立するために引用する文献（理由を付す）
「O」口頭による開示、使用、展示等に言及する文献
「P」国際出願日前で、かつ優先権の主張の基礎となる出願

の日の後に公表された文献

「T」国際出願日又は優先日後に公表された文献であって出願と矛盾するものではなく、発明の原理又は理論の理解のために引用するもの
「X」特に関連のある文献であって、当該文献のみで発明の新規性又は進歩性がないと考えられるもの
「Y」特に関連のある文献であって、当該文献と他の1以上の文献との、当業者にとって自明である組合せによって進歩性がないと考えられるもの
「&」同一パテントファミリー文献

国際調査を完了した日

10.10.2018

国際調査報告の発送日

23.10.2018

国際調査機関の名称及びあて先

日本国特許庁（I S A/J P）

郵便番号100-8915

東京都千代田区霞が関三丁目4番3号

特許庁審査官（権限のある職員）

川▲崎▼ 博章

電話番号 03-3581-1101 内線 3531

5H

5284

C (続き) . 関連すると認められる文献		
引用文献の カテゴリ*	引用文献名 及び一部の箇所が関連するときは、その関連する箇所の表示	関連する 請求項の番号
Y A	藤本 東、小川 徹、山本 和慶、松井 勇佑、山崎 俊彦、 相澤 清晴、 学術用アノテーション付き漫画画像データセットの構築と解析、 電子情報通信学会技術研究報告, 2017.03.15 (受入日) , V o l . 1 1 6、N o . 4 6 4, p. 3 5－4 0	1－4, 6, 7 5
Y A	JP 2011-238043 A (KDD I 株式会社) 2011.11.24, 段落 [0031]、[0035] (ファミリーなし)	3, 4, 6 1, 2, 5, 7
Y A	JP 2002-229766 A (イーストマン コダック カンパニー) 2002.08.16, 段落 [0023] & US 2002/0065088 A1, 段落 [0031] & EP 1211631 A1 & DE 60126896 T2 & FR 2817429 A1 & AT 355567 T	3, 4, 6 1, 2, 5, 7
Y A	福井 宏、山下 隆義、山内 悠嗣、藤吉 弘亘、 D e e p L e a r n i n g を用いた歩行者検出の研究動向、 電子情報通信学会技術研究報告, 2017.01.17 (受入日) , V o l . 1 1 6、N o . 3 6 6, p. 3 7－4 6	6 1－5, 7