

(19)日本国特許庁(JP)

(12)特許公報(B2)

(11)特許番号

特許第7299282号

(P7299282)

(45)発行日 令和5年6月27日(2023.6.27)

(24)登録日 令和5年6月19日(2023.6.19)

(51)国際特許分類

F I

G 0 6 T 7/00 (2017.01)

G 0 6 T 7/00 3 0 0 F

G 0 6 F 16/783 (2019.01)

G 0 6 T 7/00 3 5 0 C

G 0 6 F 16/732 (2019.01)

G 0 6 F 16/783

G 0 6 F 16/732

請求項の数 13 (全20頁)

(21)出願番号 特願2021-166004(P2021-166004)
 (22)出願日 令和3年10月8日(2021.10.8)
 (65)公開番号 特開2022-20647(P2022-20647A)
 (43)公開日 令和4年2月1日(2022.2.1)
 審査請求日 令和3年10月8日(2021.10.8)
 (31)優先権主張番号 202011358245.3
 (32)優先日 令和2年11月27日(2020.11.27)
 (33)優先権主張国・地域又は機関
 中国(CN)

(73)特許権者 514322098
 ベイジン バイドゥ ネットコム サイエ
 ンス テクノロジー カンパニー リミテ
 ッド
 Beijing Baidu Netco
 m Science Technolog
 y Co., Ltd.
 中華人民共和国 ペキン 100085,
 ハイディアン ディストリクト, シャン
 ディ テンス ストリート, 10番, バ
 イドゥ キャンパス 2階
 2/F Baidu Campus, N
 o.10, Shangdi 10th
 Street, Haidian Dis
 trict, Beijing 1000
 最終頁に続く

(54)【発明の名称】 ビデオ処理方法、装置、電子デバイス、記憶媒体、及びプログラム

(57)【特許請求の範囲】

【請求項1】

複数の第1ビデオフレームを取得し、前記複数の第1ビデオフレームに対してきめ細かい分割を行い、複数の第2ビデオフレームを取得することと、

前記複数の第2ビデオフレームに関するマルチモーダル情報に基づき、前記複数の第2ビデオフレームに対して特徴符号化を行い、前記マルチモーダル情報の融合を特徴付けるための特徴融合情報を取得することと、

前記特徴融合情報に基づき、前記複数の第2ビデオフレームに対して類似度のマッチングを行い、類似度のマッチング結果に基づいてターゲットビデオを取得することと、
予めトレーニングされた第1ニューラルネットワークモデルにより、前記複数の第2ビデオフレームから前記マルチモーダル情報を識別することと、

第2ニューラルネットワークモデルにより、前記マルチモーダル情報の中の各種類の情報に対して区別を行うことと、

第3ニューラルネットワークモデルにより、前記マルチモーダル情報に関する時系列情報に対して識別を行うことと、

前記第1ニューラルネットワークモデル、前記第2ニューラルネットワークモデル、前記第3ニューラルネットワークモデルの出力結果に対して融合を行い、前記特徴融合情報を取得することと、を含む

ことを特徴とするビデオ処理方法。

【請求項2】

10

20

複数の第 1 ビデオフレームを取得し、前記複数の第 1 ビデオフレームに対してきめ細かい分割を行い、複数の第 2 ビデオフレームを取得することは、

シーン及び色彩転換を特徴付けるためのパラメータに基づき、前記複数の第 1 ビデオフレームに対してきめ細かい分割を行い、複数の第 2 ビデオフレームを取得することを含むことを特徴とする請求項 1 に記載のビデオ処理方法。

【請求項 3】

前記複数の第 2 ビデオフレームに関するマルチモーダル情報に基づき、前記複数の第 2 ビデオフレームに対して特徴符号化を行い、前記マルチモーダル情報の融合を特徴付けるための特徴融合情報を取得することは、

前記マルチモーダル情報に基づき、前記複数の第 2 ビデオフレームに対して特徴抽出及び特徴融合の処理を行い、前記特徴融合情報を取得することを含む

ことを特徴とする請求項 1 に記載のビデオ処理方法。

【請求項 4】

前記特徴融合情報に基づき、前記複数の第 2 ビデオフレームに対して類似度のマッチングを行い、類似度のマッチング結果に基づいてターゲットビデオを取得することは、

前記特徴融合情報に基づき、前記複数の第 2 ビデオフレームの類似度に対して採点し、採点結果を前記類似度のマッチング結果とすることと、

前記類似度のマッチング結果として、同じイベントコンテンツに関する、隣接するビデオフレームが類似している場合、前記複数の第 2 ビデオフレームのそれぞれに対して、前記隣接するビデオフレームに対する結合がされるまで、前記隣接するビデオフレームに対してビデオ結合を行い、ビデオ結合結果に基づいて前記ターゲットビデオを取得することを含む

ことを特徴とする請求項 1 に記載のビデオ処理方法。

【請求項 5】

予めトレーニングされた第 1 ニューラルネットワークモデルにより、前記複数の第 2 ビデオフレームから前記マルチモーダル情報を識別することは、

前記第 1 ニューラルネットワークモデルの中のナレッジグラフ抽出器により、ナレッジグラフ情報を識別することと、

前記第 1 ニューラルネットワークモデルの中のテキスト抽出器により、テキスト情報を識別することと、

前記第 1 ニューラルネットワークモデルの中のオーディオ抽出器により、オーディオ情報を識別することと、

前記第 1 ニューラルネットワークモデルの中の色調抽出器により、色調情報を識別することと、

前記第 1 ニューラルネットワークモデルの中の物体抽出器により、物体情報を識別することと、

前記第 1 ニューラルネットワークモデルの中の動作抽出器により、動作情報を識別することを含む、

前記マルチモーダル情報は、前記ナレッジグラフ情報、前記テキスト情報、前記オーディオ情報、前記色調情報、前記物体情報、前記動作情報の中の少なくとも 1 つを含む

ことを特徴とする請求項 1 に記載のビデオ処理方法。

【請求項 6】

複数の第 1 ビデオフレームを取得し、前記複数の第 1 ビデオフレームに対してきめ細かい分割を行い、複数の第 2 ビデオフレームを取得するための分割モジュールと、

前記複数の第 2 ビデオフレームに関するマルチモーダル情報に基づき、前記複数の第 2 ビデオフレームに対して特徴符号化を行い、前記マルチモーダル情報の融合を特徴付けるための特徴融合情報を取得するための符号化モジュールと、

前記特徴融合情報に基づき、前記複数の第 2 ビデオフレームに対して類似度のマッチングを行い、類似度のマッチング結果に基づいてターゲットビデオを取得するためのビデオ処理モジュールと、

10

20

30

40

50

予めトレーニングされた第 1 ニューラルネットワークモデルにより、前記複数の第 2 ビデオフレームから前記マルチモーダル情報を識別するための識別モジュールと、
第 2 ニューラルネットワークモデルにより、前記マルチモーダル情報の中の各種類の情報に対して区別を行い、
第 3 ニューラルネットワークモデルにより、前記マルチモーダル情報に関する時系列情報に対して識別を行い、
前記第 1 ニューラルネットワークモデル、前記第 2 ニューラルネットワークモデル、前記第 3 ニューラルネットワークモデルの出力結果に対して融合を行い、前記特徴融合情報を取得するための融合モジュールと、を備える

ことを特徴とするビデオ処理装置。

10

【請求項 7】

前記分割モジュールは、

シーン及び色彩転換を特徴付けるためのパラメータに基づき、前記複数の第 1 ビデオフレームに対してきめ細かい分割を行い、複数の第 2 ビデオフレームを取得することに用いられる

ことを特徴とする請求項 6 に記載のビデオ処理装置。

【請求項 8】

前記符号化モジュールは、

前記マルチモーダル情報に基づき、前記複数の第 2 ビデオフレームに対して特徴抽出及び特徴融合の処理を行い、前記特徴融合情報を取得することに用いられる

20

ことを特徴とする請求項 6 に記載のビデオ処理装置。

【請求項 9】

前記ビデオ処理モジュールは、

前記特徴融合情報に基づき、前記複数の第 2 ビデオフレームの類似度に対して採点し、採点結果を前記類似度のマッチング結果とし、

前記類似度のマッチング結果として、同じイベントコンテンツに関する、隣接するビデオフレームが類似している場合、前記複数の第 2 ビデオフレームのそれぞれに対して、前記隣接するビデオフレームに対する結合がされるまで、前記隣接するビデオフレームに対してビデオ結合を行い、ビデオ結合の結果に基づいて前記ターゲットビデオを取得することに用いられる

30

ことを特徴とする請求項 6 に記載のビデオ処理装置。

【請求項 10】

前記識別モジュールは、

前記第 1 ニューラルネットワークモデルの中のナレッジグラフ抽出器により、ナレッジグラフ情報を識別し、

前記第 1 ニューラルネットワークモデルの中のテキスト抽出器により、テキスト情報を識別し、

前記第 1 ニューラルネットワークモデルの中のオーディオ抽出器により、オーディオ情報を識別し、

前記第 1 ニューラルネットワークモデルの中の色調抽出器により、色調情報を識別し、

40

前記第 1 ニューラルネットワークモデルの中の物体抽出器により、物体情報を識別し、

前記第 1 ニューラルネットワークモデルの中の動作抽出器により、動作情報を識別するために用いられ、

前記マルチモーダル情報は、前記ナレッジグラフ情報、前記テキスト情報、前記オーディオ情報、前記色調情報、前記物体情報、前記動作情報の中の少なくとも 1 つを含む

ことを特徴とする請求項 6 に記載のビデオ処理装置。

【請求項 11】

少なくとも 1 つのプロセッサと、

前記少なくとも 1 つのプロセッサに通信接続されたメモリと、を備え、

前記メモリは、前記少なくとも 1 つのプロセッサによって実行可能な命令を記憶し、前

50

記命令は、前記少なくとも 1 つのプロセッサにより実行されると、前記少なくとも 1 つのプロセッサに請求項 1 ~ 5 のいずれか 1 項に記載のビデオ処理方法を実行させることを特徴とする電子デバイス。

【請求項 1 2】

コンピュータに請求項 1 ~ 5 のいずれか 1 項に記載のビデオ処理方法を実行させるためのコンピュータ命令が記憶されている非一時的なコンピュータ可読記憶媒体。

【請求項 1 3】

コンピュータにおいて、プロセッサにより実行されると、請求項 1 ~ 5 のいずれか 1 項に記載のビデオ処理方法を実現することを特徴とするプログラム。

【発明の詳細な説明】

【技術分野】

【0001】

本開示は、人工知能分野に関し、特に、ディープラーニング、モデルトレーニング、ナレッジグラフ及びビデオ処理等の分野に関する。

【背景技術】

【0002】

ポータブル機器、携帯電話端末等の電子デバイスは、以前よりもっとインテリジェントになり、チップの分析能力がより強く、特に、ビデオ情報の分析、画面のレンダリング等は、以前より高速且つ鮮明になり、ビデオ品質に対するユーザの要求が以前より高く、特に、高適時性シナリオ（例えば、軍事パレードシナリオ、スポーツイベント、リアルタイムビデオ生中継等）の場合、各ビデオの瞬間の素晴らしい画面に対し、ユーザは、キャプチャしたいと希望しているので、より正確で且つ鮮明なビデオ画面が必要になっている。

【0003】

ビデオ処理においては、ビデオの分割を例とすると、手動の方式でビデオの分割を実現することができるが、多くの人件費を消費するだけではなく、上述した高適時性シナリオの要求を満たすことができない。一方、非手動での幾つかのビデオ分割方式では、ビデオフレームのコンテンツ情報（例えば、テキスト、ビデオ内の物体、動作等）を十分に理解することができず、ビデオイベントの連続性（例えば、シーン転換によるシナリオの切り替え等）も正しく制御することができないので、ビデオ画面に対する解釈の正確性が大幅に低下し、最終的なターゲットビデオにより提示されたビデオ品質の効果に影響を与えてしまう。

【発明の概要】

【0004】

本開示は、ビデオ処理方法、装置、電子デバイス及び記憶媒体を提供する。

【0005】

本開示の 1 つの側面では、

複数の第 1 ビデオフレームを取得し、前記複数の第 1 ビデオフレームに対してきめ細かい分割を行い、複数の第 2 ビデオフレームを取得することと、

前記複数の第 2 ビデオフレームに関するマルチモーダル情報に基づき、前記複数の第 2 ビデオフレームに対して特徴符号化を行い、前記マルチモーダル情報の融合を特徴付けるための特徴融合情報を取得することと、

前記特徴融合情報に基づき、前記複数の第 2 ビデオフレームに対して類似度のマッチングを行い、類似度のマッチング結果に基づいてターゲットビデオを取得することを含むビデオ処理方法が提供される。

【0006】

本開示のもう 1 つの側面では、

複数の第 1 ビデオフレームを取得し、前記複数の第 1 ビデオフレームに対してきめ細かい分割を行い、複数の第 2 ビデオフレームを取得するための分割モジュールと、

前記複数の第 2 ビデオフレームに関するマルチモーダル情報に基づき、前記複数の第 2 ビデオフレームに対して特徴符号化を行い、前記マルチモーダル情報の融合を特徴付ける

10

20

30

40

50

ための特徴融合情報を取得するための符号化モジュールと、

前記特徴融合情報に基づき、前記複数の第2ビデオフレームに対して類似度のマッチングを行い、類似度のマッチング結果に基づいてターゲットビデオを取得するためのビデオ処理モジュールとを備えるビデオ処理装置が提供される。

【0007】

本開示のもう1つの側面では、

少なくとも1つのプロセッサと、

前記少なくとも1つのプロセッサに通信接続されたメモリとを備える電子デバイスが提供される。

【0008】

ここで、前記メモリは、前記少なくとも1つのプロセッサによって実行可能な命令を記憶し、前記命令は、前記少なくとも1つのプロセッサにより実行されると、前記少なくとも1つのプロセッサに本開示の任意の1つの実施形態による方法を実行させる。

【0009】

本開示のもう1つの側面では、コンピュータ命令が記憶されている非一時的なコンピュータ可読記憶媒体が提供される。ここで、前記コンピュータ命令は、前記コンピュータに本開示の任意の1つの実施形態による方法を実行させることに用いられる。

【0010】

本開示によれば、複数の第1ビデオフレームを取得し、前記複数の第1ビデオフレームに対してきめ細かい分割を行い、複数の第2ビデオフレームを取得する。前記複数の第2ビデオフレームに関するマルチモーダル情報に基づき、前記複数の第2ビデオフレームに対して特徴符号化を行い、前記マルチモーダル情報の融合を特徴付けるための特徴融合情報を取得する。前記特徴融合情報に基づき、前記複数の第2ビデオフレームに対して類似度のマッチングを行い、類似度のマッチング結果に基づいてターゲットビデオを取得する。当該マルチモーダル情報に基づいて特徴符号化を行うことにより、より多くのビデオコンテンツの詳細を含む情報を取得し、類似度に基づいてマッチングされた後、得られたターゲットビデオの精度がより高いので、ビデオ分割の正確性を高めることができる。

【0011】

当該部分に記載の内容は、本開示の実施形態の肝心又は重要な特徴を示すことを意図するものではなく、本開示の範囲を制限しないことが理解されたい。本開示の他の特徴は、以下の説明により、より理解しやすくなる。

【0012】

添付の図面は、本実施形態をより良く理解するために用いられ、本開示に対する限定を構成するものではない。

【図面の簡単な説明】

【0013】

【図1】本開示の実施形態によるビデオ処理の高適時性シナリオの複数のビデオフレームの模式図である。

【図2】本開示の実施形態によるビデオ処理方法のフローチャート模式図である。

【図3】本開示の実施形態によるビデオ処理方法を実現するシステムモジュールのアーキテクチャ図である。

【図4】本開示の実施形態による、マルチモーダル情報に基づいて特徴符号化を実現する模式図である。

【図5】本開示の実施形態による類似度のマッチングの模式図である。

【図6】本開示の実施形態によるビデオ結合の模式図である。

【図7】本開示の実施形態によるビデオ処理装置の構成構造模式図である。

【図8】本開示の実施形態によるビデオ処理方法を実現するための電子デバイスのブロック図である。

【発明を実施するための形態】

【0014】

10

20

30

40

50

以下、図面を参照しながら、本開示の例示的な実施形態を説明し、理解を助けるために本開示の実施形態の様々な詳細を含んでいるが、これらは、単に例示的なものとみなされるべきである。よって、当業者は、本開示の範囲及び要旨から逸脱することなく、本明細書に記載の実施形態に様々な変更及び修正を加えることができることを識別すべきである。明瞭で且つ簡潔にするために、以下の説明では、周知の機能と構造の説明を省略している。

【0015】

本文における用語「及び／又は」は、関連対象の関連関係を説明するものに過ぎず、3つの関係があっても良いことを表し、例えば、A及び／又はBは、Aだけがあり、A及びBがあり、Bだけがあるという3つの場合を表すことができる。本文における用語「少なくとも1つ」は、複数の中の任意の1つ又は複数の中の少なくとも2つの任意の組み合わせを表し、例えば、A、B、Cの中の少なくとも1つを含むことは、A、B、Cからなる集合から選択された任意の1つ又は複数の元素を含むことを表すことができる。本文における用語「第1」、「第2」は、複数の類似する技術用語を指し、それらを区別するためのものであり、順序を制限する意味がなく、又は、2つだけに制限する意味がなく、例えば、第1特徴及び第2特徴は、2種類／2つの特徴を指し、第1特徴は1つ又は複数であっても良く、第2特徴も1つ又は複数であっても良い。

【0016】

また、本開示をより良く説明するために、後述における具体的な実施形態においては、沢山の具体的な詳細が記載されている。当業者は、幾つかの具体的な詳細がなくても、本開示は、同様に実施することができることを理解すべきである。幾つかの実施形態においては、本開示の主旨を強調するために、当業者に周知されている方法、手段、素子及び回路に対しては詳しく説明していない。

【0017】

ビデオ分割は、インターネットビデオ及びニューメディアの短いビデオコンテンツプラットフォームのニーズにより、従来のテレビメディア番組に対して行われた二次処理であり、即ち、元の完全な番組コンテンツを、ある論理的な思考又は特別のニーズに従って複数のビデオに分割するものである。インターネットビデオコンテンツの主なソースは、従来のテレビメディアの番組、様々な機関のビデオ製品、映画とテレビ会社の映画とテレビ作品を含み、これらのビデオを分割することにより、貴重な情報を深く掘り下げることができ、新たに編集された後、インターネットプロトコルテレビジョン（IPTV）、OTT、モバイルTV及びニューメディアの短いビデオプラットフォームに用いることができ、ニューメディアオーディオビジュアル番組の断片化の要求を満たすことができ、オーディオビジュアルの編集業界の1つの新しい試み及び探索となる。

【0018】

従来の手動によるビデオ分割技術は、人手でビデオを編集及び分割するので、処理時間が長く、ビデオの量が多い場合、生産性の向上が遅いため、高適時性等のシナリオに応用することができず、大量の時間及び経験を消費するだけではなく、経済的コスト及び適時性の要求を満たすことができない。非機械学習の従来のビデオ分割アルゴリズムのような、非手動によるビデオ分割技術の場合、色調及びブロックマッチングに基づいてビデオ分割を行うことができるが、ピクチャー、シーンの間の視覚情報しか考慮せず、ビデオフレームのコンテンツ情報を十分に理解することができない。また、機械学習に基づくビデオ分割技術の場合、クラスタリング方法に基づいてビデオ分割（キーフレームの抽出、画像特徴の説明、クラスタリング等を含む）を行うことができるが、ビデオイベントの連続性を考慮せず、幾つかのシーン切り替えが比較的頻繁なシナリオ（図1に示す体育イベントの中の複数のビデオフレームから構成される素晴らしい瞬間等）、又は、複雑なイベントのシナリオにおいて複数のシーンの連続的な転換があるビデオの場合、ビデオ分割の効果が良くなく、ビデオ分割の正確率が低い。

【0019】

本開示の実施形態では、ビデオ処理方法が提供される。図2は、本開示の実施形態によ

10

20

30

40

50

るビデオ処理方法のフローチャート模式図であり、当該方法は、ビデオ処理装置に用いることができ、例えば、当該装置は、端末、サーバ又は他の処理デバイスに展開でき、ビデオフレーム分割、ビデオフレーム特徴符号化、ビデオフレーム類似度のマッチングを実行することにより、最終的なターゲットビデオ等を取得することができる。ここで、端末は、ユーザ機器（UE、User Equipment）、モバイルデバイス、携帯電話、コードレス電話、パーソナルデジタルアシスタント（PDA、Personal Digital Assistant）、ハンドヘルドデバイス、コンピューティングデバイス、車載機器、ウェアラブル機器等であっても良い。幾つかの可能な実現方法においては、当該方法は、プロセッサがメモリに記憶されたコンピュータ可読命令を呼び出す方式により実現することができる。図2に示すように、ステップS101、ステップS102及びステップS103を含む。

10

【0020】

ステップS101においては、複数の第1ビデオフレームを取得し、前記複数の第1ビデオフレームに対してきめ細かい分割を行い、複数の第2ビデオフレームを取得する。

【0021】

ステップS102においては、前記複数の第2ビデオフレームに関するマルチモーダル情報に基づき、前記複数の第2ビデオフレームに対して特徴符号化を行い、前記マルチモーダル情報の融合を特徴付けるための特徴融合情報を取得する。

【0022】

ステップS103においては、前記特徴融合情報に基づき、前記複数の第2ビデオフレームに対して類似度のマッチングを行い、類似度のマッチング結果に基づいてターゲットビデオを取得する。

20

【0023】

ステップS101においては、シーン及び色彩転換を特徴付けるためのパラメータに基づき、前記複数の第1ビデオフレームに対してきめ細かい分割を行い、前記複数の第2ビデオフレームを取得する。シーン及び色彩転換を特徴付けるためのパラメータは、シーンの観点からは、ビデオエッジに対する分割、ビデオの中のブロックマッチングに基づくビデオ分割、統計的決定に基づくビデオ分割及び双閾値の比較（双閾値の比較を設定することにより、シーンの急変なのか、シーンの段階的な変化なのかを区別する）に基づくビデオ分割等を含むことができる。色彩転換の観点からは、色調に基づくビデオ分割を含むことができる。

30

【0024】

ステップS102においては、前記マルチモーダル情報に基づき、前記複数の第2ビデオフレームに対して特徴抽出及び特徴融合処理を行い、前記特徴融合情報を取得する。ここで、特徴融合処理は、複数のニューラルネットワークモデル、又は、多機能抽出に統合された1つのニューラルネットワークモデルをエキスパートモデルとして用いることにより、第2ビデオフレームに関するマルチモーダル情報に対してそれぞれ特徴抽出を行う。ここで、マルチモーダル情報は、ナレッジグラフ情報、テキスト情報、オーディオ情報、色調情報、物体情報、動作情報の中の少なくとも1つを含む。

【0025】

40

ステップS103においては、前記特徴融合情報に基づき、前記複数の第2ビデオフレームの類似度に対して採点し、採点結果を前記類似度のマッチング結果とし、前記類似度のマッチング結果として、同じイベントコンテンツに関する、隣接するビデオフレームが類似している場合、前記複数の第2ビデオフレームのそれぞれに対して、前記隣接するビデオフレームに対する結合がされるまで、前記隣接するビデオフレームに対してビデオ結合を行い、ビデオ結合の結果に基づいて前記ターゲットビデオを取得する。

【0026】

本開示によれば、複数の第1ビデオフレームを取得し、前記複数の第1ビデオフレームに対してきめ細かい分割を行い、複数の第2ビデオフレームを取得する。前記複数の第2ビデオフレームに関するマルチモーダル情報に基づき、前記複数の第2ビデオフレームに

50

対して特徴符号化を行い、前記マルチモーダル情報の融合を特徴付けるための特徴融合情報を取得することができる。前記特徴融合情報に基づき、前記複数の第2ビデオフレームに対して類似度のマッチングを行い、類似度のマッチング結果に基づいてターゲットビデオを取得することができる。当該マルチモーダル情報に基づいて特徴符号化を行うことにより、より多くのビデオコンテンツの詳細を含む情報を取得し、類似度に基づいてマッチングされた後、得られたターゲットビデオの精度がより高いので、ビデオ分割の正確性を高めることができる。

【0027】

1つの例においては、ビデオ分割モジュール、マルチモーダル特徴符号化モジュール、類似度マッチングモジュール（主に、隣接するビデオセグメントに対する類似度のマッチング）、ビデオフレームスライシングモジュールにより、上述したステップS101～ステップS103を実現することができる。具体的には、当該ビデオ分割モジュールに入力された複数の第1ビデオフレームを取得した後、シーン及び色彩転換に従い、取得された複数の第1ビデオフレームに対してきめ細かい分割を行い、複数の第2ビデオフレームを取得することができる。当該マルチモーダル特徴符号化モジュールに入力された当該複数の第2ビデオフレームに対し、マルチモーダル情報に基づいて特徴符号化（例えば、マルチモーダル情報の特徴抽出及び特徴融合）を行い、マルチモーダル情報が融合された後の特徴情報を取得する。当該特徴情報を当該類似度マッチングモジュールに入力してビデオの類似度のマッチングを行い、類似度のマッチング結果（例えば、類似度の採点結果）を取得する。類似度のマッチング結果は、同じイベントコンテンツの2つの隣接するビデオフレームが類似している場合、同じイベント内のビデオセグメントに対して復元を行うという戦略に基づき、当該ビデオフレームスライシングモジュールにより同じイベントコンテンツの2つの隣接するビデオフレームに対してそれぞれビデオ結合を行い、ビデオ処理後の最終的なビデオ分割結果を取得する。同じイベントコンテンツの類似度に従って結合し、ビデオコンテンツの詳細の類似度をより注目するので、ビデオ分割がより正確になり、最終的なビデオ分割結果の正確性を大幅に高めることができる。

【0028】

1つの実施形態においては、予めトレーニングされた第1ニューラルネットワークモデルにより、前記複数の第2ビデオフレームから前記マルチモーダル情報を識別する。

【0029】

1つの例においては、第1ニューラルネットワークモデルは、複数のエキスパートモデルから構成されても良く、複数のエキスパートモデルのそれぞれの機能を、1つのニューラルネットワークに集積されて構成されても良い。複数のエキスパートモデルのそれぞれの機能を、1つのニューラルネットワークに集積して構成される第1ニューラルネットワークモデルを例とする場合、当該第1ニューラルネットワークモデルは、ナレッジグラフ抽出器、テキスト抽出器、オーディオ抽出器、色調抽出器、物体抽出器及び動作抽出器を含んでも良い。ここで、第1ニューラルネットワークモデルの中のナレッジグラフ抽出器（又は、ナレッジグラフに基づく構造化ラベルベクトル抽出器と呼ばれる）により、ナレッジグラフ情報（例えば、knowledge特徴等）を識別することができ、第1ニューラルネットワークの中のテキスト抽出器（又は、テキストに基づくテキストベクトル抽出器と呼ばれる）により、テキスト情報（例えば、text特徴）を識別することができ、第1ニューラルネットワークの中のオーディオ抽出器（又は、オーディオに基づくオーディオベクトル抽出器と呼ばれる）により、オーディオ情報（例えば、audio情報）を識別することができ、前記第1ニューラルネットワークの中の色調抽出器（又は、画像に基づくRGB抽出器と呼ばれる）により、色調情報（例えば、RGB特徴）を識別することができ、前記第1ニューラルネットワークの中の物体抽出器（又は、ターゲット検出に基づく物体特徴抽出器と呼ばれる）により、物体情報（例えば、object特徴）を識別することができ、前記第1ニューラルネットワークの中の動作抽出器（動作識別に基づく動作ベクトル抽出器）により、動作情報（例えば、action特徴）を識別することができる。ここで、前記マルチモーダル情報は、前記ナレッジグラフ情報、前記テキス

10

20

30

40

50

ト情報、前記オーディオ情報、前記色調情報、前記物体情報及び前記動作情報の中の少なくとも1つを含む。

【0030】

本実施形態によれば、本開示のインテリジェントビデオ分割技術と機械学習に基づく複数のエキスパートモデルを組み合わせることで、マルチモーダル情報の特徴識別、特徴抽出及び特徴融合を実現する。更に、融合された特徴情報（特徴融合情報と呼ばれる）に対して類似度の比較を実現する。よって、より詳細なビデオコンテンツ情報を取得し、ビデオコンテンツ及びイベント知識をより深く理解することができるので、最も正しいビデオ分割結果が得られ、最終的なビデオ分割結果の正確度を大幅に高めることができる。

【0031】

1つの実施形態においては、予めトレーニングされた、ビデオ特徴抽出モデル $F(u)$ のような第1ニューラルネットワークモデルにより、前記複数の第2ビデオフレームから前記マルチモーダル情報を識別及び抽出することができる。ビデオ特徴識別モデル $M(u)$ のような第2ニューラルネットワークモデルにより、前記マルチモーダル情報のそれぞれの情報を区別することができる。ビデオに対応する時系列情報の抽出モデル $T(u)$ のような、第3ニューラルネットワークモデルにより、前記マルチモーダル情報に関する時系列情報を識別及び抽出し、前記ビデオ特徴抽出の時間オフセット表現を記録し、前記第1ニューラルネットワークモデル、第2ニューラルネットワークモデル及び第3ニューラルネットワークモデルの出力結果を融合し、前記特徴融合情報を取得する。特徴融合情報は、より多くのビデオコンテンツの詳細を記述することができるので、後続において類似度の比較を行う際にマッチング速度及び精度を高め、同じイベントコンテンツの2つの隣接するビデオフレームに対して類似度のマッチングが行われた後にビデオ処理の最終的なビデオ分割結果を取得し、結果がより正確であり、ビデオ分割がより正確であり、最終的なビデオ分割結果の正確度を大幅に高めることができる。

【0032】

応用例

本開示の実施形態を応用する1つの処理流れは、以下の内容を含む。

【0033】

図3は、本開示の実施形態によるビデオ処理方法を実現するシステムモジュールのアーキテクチャ図であり、ビデオ分割モジュール（主に、シーン及び色彩転換に従い、ビデオに対してきめ細かい分割を行う）、マルチモーダル特徴符号化モジュール（主に、マルチモーダル情報を用いてビデオに対して特徴符号化を行う）、類似度マッチングモジュール（主に、隣接するビデオセグメントに対して類似度のマッチングを行い、更に、同じイベントコンテンツに従い、ビデオセグメントに対してビデオ結合を行うことにより、最終的なビデオ分割結果を取得することができる）から構成されるシステムにより、本開示のインテリジェントビデオ分割の流れを実現することができる。マルチモーダル情報の融合により、ビデオコンテンツ及びイベント知識を深く理解し、ディープラーニングを組み合わせることでビデオを分割する。図3に示すように、以下の内容を含む。

【0034】

ビデオ分割モジュール

ビデオ分割モジュールにより、ビデオセグメントのきめ細かい分割を行うことができる。きめ細かい分割の原則は、主に、次の内容を含む。1)色調に基づく分割であり、2つのフレームのグレースケール差を直接に計算することができ、合計するフレーム差が設定されたある閾値より大きい場合、シーンの急変がある。2)エッジに基づく分割であり、エッジ特徴がシーンの分割に用いられることができ、まず、フレーム間の全体的な変位を計算し、これに基づいて位置合わせを行い、次に、エッジの数及び位置を計算する。3)ブロックマッチングに基づく分割であり、非圧縮ビデオに用いられるブロックマッチングシーンを例とする場合、動きの滑らかさの度量を用いてシーンの変化を検出することができる。4)統計的決定に基づく分割であり、動き補償特徴、適応閾値方式、ビデオシーケンスの時系列シーン急変モード及びシーン長さ分布情報を用い、統計的決定モデルを確立

10

20

30

40

50

し、当該統計的決定モデルが推定した基準により、シーン検出エラー率を最小限に抑えることができる。５）双閾値の比較に基づく分割であり、双閾値（例えば、 T_b 、 T_s ）を設定することができる。フレーム差が T_b より大きい場合、シーンの急変があるが、フレーム差が T_b より小さく且つ T_s より大きい場合、シーンの段階的な変化がある。接続するフレームのフレーム差が T_s を超え始めると、このフレームは、シーンの段階的な変化の開始フレームと呼ばれ、これによって類推する。

【００３５】

二、マルチモーダル特徴符号化モジュール

図４は、本開示の実施形態による、マルチモーダル情報に基づいて特徴符号化を実現する模式図である。マルチモーダル特徴符号化は、主に、複数のエキスパートモデルにより
10 若干のエキスパートベクトル（ $experts\ embedding$ ）を取得し、これらの $experts\ embedding$ によりビデオ全体のマルチモーダル情報の描画及び特徴抽出を完了する。図４に示すように、以下の内容を含む。

【００３６】

ビデオレベルベクトル（ $embedding$ ）表現は、マルチモーダルトランスフォーマー（MMT、 $Multi-modal\ Transformer$ ）のようなマルチモーダルコーディングモジュールの１つの例により得られる。MMTは、 $Transformer$ エンコーダーのアーキテクチャに従うことができ、 $Transformer$ エンコーダーは、スタックされた自己注意機構（ $Self-Attention$ ）層及び完全接続層で構成される。MMTの入力 $\Omega(u)$ は、ワンセットの $embedding$ 表現であり、
20 全ての次元が同じであり、 d_{model} 次元として定義され、その中のそれぞれの $embedding$ は、何れもワンセットの特徴表現を表し、この入力は、式（１）に示されている。

【００３７】

【数１】

$$\Omega(u) = F(u) + M(u) + T(u) \quad (1)$$

30

【００３８】

ここで、式（１）の中のそれぞれのパラメータの意味は、次の通りである。

$\Omega(u)$ は、ビデオフレーム特徴符号化の後のベクトル出力を表す。

【００３９】

$F(u)$ は、ビデオ特徴抽出モデルを表し、ビデオの中のマルチモーダル情報を抽出することに用いられる。ここで、 F_{agg}^{know} 、 F_1^{know} 、 $\dots$ 、 F_k^{know} は、ナレッジグラフ情報（例えば、 $knowledge$ ）を抽出することを表し、「 $know$ 」は、 $knowledge$ 特徴の略称であり、「 k 」は、ベクトルの次元を表し、「 agg 」は、平均
40 ベクトルを表し、これによって類推し、テキスト情報（例えば、 $text$ 特徴）、オーディオ情報（例えば、 $audio$ 特徴）、色調情報（例えば、 RGB 特徴）、物体情報（例えば、 $object$ 特徴）及び動作情報（例えば、 $action$ 特徴）に対してそれぞれ抽出する。

【００４０】

$M(u)$ は、ビデオ特徴識別モデルを表し、マルチモーダル情報の中の異なるタイプの情報を区別することに用いられる。

【００４１】

$T(u)$ は、ビデオに対応する時系列情報の抽出モデルを表し、時系列情報（即ち、時系列ベクトル）を抽出及び記録することにより、特徴抽出の時間オフセット表現を記録す
50

るために用いられ、ここで、 T_{agg} 、 T_1 、 $\dots$ 、 T_D は、抽出する時系列情報を表し、「 D 」は、何秒を表す。

【0042】

$F(u)$ により抽出された前記マルチモーダル情報の場合、ビデオデータの固有の異なる形式から効果的な表現を学習するために、様々なエキスパートモデルをビデオ特徴抽出器として用いることにより、前記マルチモーダル情報を抽出することができる。様々なエキスパートモデルは、ナレッジグラフに基づく構造化ラベル $embedding$ 抽出器、テキストに基づく $text\ embedding$ 抽出器、オーディオに基づく $audio\ embedding$ 抽出器、画像に基づく RGB 抽出器、ターゲット検出に基づく $object$ 特徴抽出器及び動作識別に基づく $action\ embedding$ 抽出器を主に 10
含み、様々なエキスパートモデルがそれぞれ抽出された前記マルチモーダル情報に対して特徴融合を行うことにより、様々なコンテンツ次元でビデオ情報を完全に特徴付けることができる。特徴融合により、学習された、異なるエキスパートモデルにより抽出された前記マルチモーダル情報の間の関係に基づき、クロスモードとロングシーケンスの時間関係を利用して共同表現を行うことで、予めトレーニングされた異なるエキスパートモデル $\{F^n\}_{n=1}^N$ を用いてより正確なビデオコンテンツの詳細を取得することができる。

【0043】

それぞれのエキスパートモデルは、特別のタスクによるトレーニングによって取得されるものであり、その後、マルチモーダル情報の特徴抽出に用いられる。ビデオ u に対し、それぞれのエキスパートモデルは、 K 個の特徴 ($features$) を含むシーケンスを抽出することができ、 $F^n(u)=[F_1^n, \dots, F_K^n]$ として表す。 20

【0044】

様々なエキスパートモデルにより抽出されたビデオの $feature$ 特徴付けは、異なるエキスパートモデルを用いて特徴抽出を行うので、抽出された異なるエキスパートベクトルの特徴 (又は、特徴ベクトルと呼ばれる) を共通の d_{model} 次元にマッピングするために、 N 個の $linear\ layers$ (それぞれのエキスパートごとに1つ) を用いて全ての特徴を $R^{d_{model}}$ に投影することができる。

【0045】

$Transformer$ エンコーダーは、それぞれの特徴入力ごとに1つの $embedding$ を生成するので、複数の特徴に複数の $embedding$ 表現を提供する。それぞれの特徴の唯一の $embedding$ 表現を得るために、1つの纏め $embedding$ F_{agg}^n を定義することができ、当該 $embedding$ は、収集された特徴をコンテキスト化 (現在の表現を集合して) し、最大プーリングにより当該埋め込みを初期化する: $F_{agg}^n = \maxpool(\{F_k^n\}_{k=1}^N)$ 、よって、入力 $feature$ シーケンス全体の形式は、式 (2) に示されている。 30

【0046】

【数2】

$$F(u) = [F_{agg}^1, F_1^1, \dots, F_K^1, \dots, F_{agg}^N, F_1^N, \dots, F_K^N] \quad (2)$$

40

【0047】

式 (2) においては、 N は、エキスパートモデルの数 (N は、1 より大きい正整数である) を表し、 K は、ベクトル次元 (K は、1 より大きい正整数である) を表す。

【0048】

$M(u)$ によりマルチモーダル情報の中の異なるタイプの情報を区別する場合、マルチモーダル情報をより良く処理及び区別するために、 MMT は、現在処理している $embedding$ 埋め込みがどのエキスパートモデルからの入力であるかを区別する必要があり、 N 個の d_{model} 次元の $embedding$ 埋め込み $\{E_1, \dots, E_N\}$ を学習することにより、異 50

なるエキスパートの $embedding$ 表現を区別することができる。エキスパートモデルの $embedding$ からビデオエンコーダへのシーケンスは、次の式 (3) に示す形式を採用することができる。

【0049】

【数3】

$$E(v) = [E^1, E^1, \dots, E^1, \dots, E^N, E^N, \dots, E^N] \quad (3)$$

【0050】

式 (3) においては、 N は、エキスパートモデルの数 (N は、1 より大きい正整数である) を表す。

【0051】

$T(v)$ により時系列情報を記録する場合、ビデオのそれぞれの特徴が抽出した MMT からの時間情報を提供する。1つのビデオの最長継続時間は、 t_{max} 秒であっても良く、「秒」を測定パラメータとし、 $\{T_1, \dots, T_D\}$ のように、秒ごとに1つの d_{model} 次元の $D = |t_{max}|$ を学習する。時間範囲 $[t, t+1)$ から抽出されたそれぞれのエキスパートモード $feature$ を T_{t+1} として表す。例えば、ビデオにおいては、2.2秒で抽出された特徴は、時間埋め込み T_3 により時間符号化される。2つの追加する時間埋め込み T_{agg} 及び T_{unk} を学習し、集約特徴及び未知の時間情報特徴に対してそれぞれ符号化する。最後に、 $Temporal\ embeddings\ T$ は、次の式 (4) に示す形式を採用する。

【0052】

【数4】

$$T(v) = [T_{agg}, T_1, \dots, T_D, \dots, T_{agg}, T_1, \dots, T_D] \quad (4)$$

【0053】

式 (4) においては、 T_{agg} は、時間情報の平均ベクトルを表し、 T_D は、第 D 秒 (D は、1秒より大きい数値である) の時間情報を表す。

【0054】

MMT の実現方式は、式 (5) に示されている。

【0055】

【数5】

$$\psi_{agg}(v) = MMT(\Omega(v)) = [\psi_{agg}^1, \dots, \psi_{agg}^N] \quad (5)$$

【0056】

式 (5) においては、 N は、エキスパートモデルの数を表し、同時に、 $\psi_{agg}(v)$ は、ビデオ纏めの情報を表す。 $\Omega(v)$ は、 MMT の入力を表す。

【0057】

三、類似度マッチングモジュール

図5は、本開示の実施形態による類似度のマッチングの模式図である。類似度マッチングモジュールにより、2つの隣接するビデオセグメントの類似度の計算を主に完成し、類似度マッチングは、上下対称するモジュールの設計を採用する。複数のエキスパートモデル $embedding$ の類似性を計算し、重み計算 (重みは、注意機構により自動的に学習することができる) を用いて類似度の採点を取得することにより、類似度のマッチング結果を取得する。また、損失関数は、双方向の最大結合ランキング損失関数 ($bi-dir$

10

20

30

40

50

ectional max-margin ranking loss)を採用でき、式(6)に示されている。

【0058】

【数6】

$$\mathcal{L} = \frac{1}{B} \sum_{i=1}^B \sum_{j \neq i} [\max(0, s_{ij} - s_{ii} + m) + \max(0, s_{ji} - s_{ii} + m)] \quad (6)$$

【0059】

式(6)においては、Lは、前記損失関数を表し、Bは、サンプルバッチ処理に用いられるハイパーパラメータ(batch size)を表し、 s_{ij} =similarity(v_i, v_j)であり、 s_{ij} は、2つのビデオセグメントの類似度を表し、mは、marginであり、値(0, 1)を取ることができる。

10

【0060】

四、類似度マッチングモジュール又は上述した図3を元に、類似度マッチングモジュールの後に、ビデオスライシング処理に特別に用いられるビデオフレームスライシングモジュールを追加する。

【0061】

図6は、本開示の実施形態によるビデオ結合の模式図であり、当該ビデオスライシング処理を統合する類似度マッチングモジュールを例とする。図6に示すように、類似度マッチングモジュールにより、隣接するビデオセグメントの結合とスライシングを実現することができる。主に、同じイベント内の細かいビデオセグメントを復元する。2つの隣接するビデオセグメントが類似していると判断された場合、2つのビデオを結合し、順に比較し、最終的なビデオ分割結果を取得する。

20

【0062】

本応用例によれば、マルチモーダル情報を抽出するための複数のエキスパートモデルの情報を融合し、マルチモーダル情報をキャプチャ及び融合することができるので、ビデオコンテンツ全体を完全に描画し、ビデオ画面の再現効果を高めることができる。ディープラーニングの方式により、大規模な且つ大量のビデオ分割を行うことができ、プロセス全体がより効率であり、コストがより低く、高適時性のビデオ要求を満たすことができる。KGに基づく構造化ラベル技術(例えば、実体、主題等である)、テキストに基づく表現技術、ビジョンに基づく(RGB、Object Action)等を結合し、ビデオコンテンツの角度からビデオを分割し、複数のシーンの頻繁な切り替えによる分割効果の悪い問題を解決することができる。また、スケーラビリティが良く、使用シナリオは、ビデオ技術に限らず、ビデオ指紋識別、ビデオ短帯域長、同じビデオマッチング等のような、任意のビデオの類似度のマッチングシナリオに適用できる。

30

【0063】

本開示の実施形態によれば、ビデオ処理装置が提供される。図7は、本開示の実施形態によるビデオ処理装置の構成構造模式図である。図7に示すように、当該ビデオ処理装置は、分割モジュール41、符号化モジュール42及びビデオ処理モジュール43を備える。

40

【0064】

分割モジュール41は、複数の第1ビデオフレームを取得し、前記複数の第1ビデオフレームに対してきめ細かい分割を行い、複数の第2ビデオフレームを取得することに用いられる。

【0065】

符号化モジュール42は、前記複数の第2ビデオフレームに関するマルチモーダル情報に基づき、前記複数の第2ビデオフレームに対して特徴符号化を行い、前記マルチモーダル情報の融合を特徴付けるための特徴融合情報を取得することに用いられる。

【0066】

ビデオ処理モジュール43は、前記特徴融合情報に基づき、前記複数の第2ビデオフレ

50

ームに対して類似度のマッチングを行い、類似度のマッチング結果に基づいてターゲットビデオを取得することに用いられる。

【0067】

1つの実施形態においては、前記分割モジュールは、シーン及び色彩転換を特徴付けるためのパラメータに基づき、前記複数の第1ビデオフレームに対してきめ細かい分割を行い、複数の第2ビデオフレームを取得することに用いられる。

【0068】

1つの実施形態においては、前記符号化モジュールは、前記マルチモーダル情報に基づき、前記複数の第2ビデオフレームに対して特徴抽出及び特徴融合の処理を行い、前記特徴融合情報を取得することに用いられる。

10

【0069】

1つの実施形態においては、前記特徴融合情報に基づき、前記複数の第2ビデオフレームの類似度に対して採点し、採点結果を前記類似度のマッチング結果とし、前記類似度のマッチング結果として、同じイベントコンテンツに関する、隣接するビデオフレームが類似している場合、前記複数の第2ビデオフレームのそれぞれに対して、隣接するビデオフレームに対する結合がされるまで、前記隣接するビデオフレームに対してビデオ結合を行い、ビデオ結合の結果に基づいて前記ターゲットビデオを取得することに用いられる。

【0070】

1つの実施形態においては、予めトレーニングされた第1ニューラルネットワークモデルにより、前記複数の第2ビデオフレームから前記マルチモーダル情報を識別するための識別モジュールを更に備える。

20

【0071】

1つの実施形態においては、前記識別モジュールは、前記第1ニューラルネットワークモデルの中のナレッジグラフ抽出器により、ナレッジグラフ情報を識別し、前記第1ニューラルネットワークモデルの中のテキスト抽出器により、テキスト情報を識別し、前記第1ニューラルネットワークモデルの中のオーディオ抽出器により、オーディオ情報を識別し、前記第1ニューラルネットワークモデルの中の色調抽出器により、色調情報を識別し、前記第1ニューラルネットワークモデルの中の物体抽出器により、物体情報を識別し、前記第1ニューラルネットワークモデルの中の動作抽出器により、動作情報を識別することに用いられる。前記マルチモーダル情報は、前記ナレッジグラフ情報、前記テキスト情報、前記オーディオ情報、前記色調情報、前記物体情報、前記動作情報の中の少なくとも1つを含む。

30

【0072】

1つの実施形態においては、第2ニューラルネットワークモデルにより、前記マルチモーダル情報の中の各種類の情報に対して区別を行い、第3ニューラルネットワークモデルにより、前記マルチモーダル情報に関する時系列情報に対して識別を行い、前記第1ニューラルネットワークモデル、前記第2ニューラルネットワークモデル、前記第3ニューラルネットワークモデルの出力結果に対して融合を行い、前記特徴融合情報を取得するための融合モジュールを更に備える。

【0073】

40

本開示の実施形態におけるそれぞれの装置の中の各モジュールの機能は、上述した方法の対応する記載を参照することができ、ここでは、繰り返して説明しない。

【0074】

本開示の実施形態によれば、本開示は、電子デバイス及び可読記憶媒体が更に提供される。

図8は、本開示の実施形態による例示するビデオ処理方法を実現するための電子デバイスのブロック図である。当該デバイスは、上述した展開デバイス又はエージェントデバイスであっても良い。電子デバイスは、ラップトップコンピュータ、デスクトップコンピュータ、ワークステーション、パーソナルデジタルアシスタント、サーバ、ブレードサーバ、メインフレームコンピュータのような様々な形態のデジタルコンピュータ及び他の好適

50

なコンピュータを表すことを目的としている。また、電子デバイスは、また、様々な形態のモバイルデバイス、例えば、パーソナルデジタルアシスタント、携帯電話、スマートフォン、ウェアラブルデバイス、及び他の類似のコンピューティングデバイスを表すことができる。本明細書に記載のコンポーネント、それらの接続及び関係、ならびにそれらの機能は、例としてのみ意図されており、本明細書に記載及び／又は要求される本開示の実現を限定することを意図するものではない。

【 0 0 7 5 】

当該電子デバイスは、１つ以上のプロセッサ 8 0 1、メモリ 8 0 2、及び各コンポーネントを接続するための、高速インターフェース及び低速インターフェースを含むインターフェースを有する。様々なコンポーネントは、異なるバスを用いて相互に接続されており、共通のマザーボード上に実装されてもよいし、必要に応じて他の方式で実装されてもよい。プロセッサは、電子デバイス内で実行される命令を処理してもよく、当該命令は、メモリに又はメモリ上に記憶されることによって、外部入出力装置（例えば、インターフェースに結合されたディスプレイ装置）に GUI のグラフィカル情報を表示させるための命令を含む。他の実施形態では、必要に応じて、複数のプロセッサ及び／又は複数のバスを複数のメモリと一緒に使用してもよい。同様に、複数の電子デバイスが接続されていてもよく、個々のデバイスが必要な操作の一部を提供する（例えば、サーバアレイ、ブレードサーバのグループ又はマルチプロセッサシステムとして）。図 8 は、一つのプロセッサ 8 0 1 を例としている。

【 0 0 7 6 】

メモリ 8 0 2 は、本開示によって提供される非一時的なコンピュータ可読記憶媒体である。ここで、前記メモリは、本開示により提供されるビデオ処理方法を前記少なくとも 1 つのプロセッサに実行させるために、前記少なくとも 1 つのプロセッサにより実行可能な命令を記憶している。本開示の非一時的なコンピュータ可読記憶媒体は、本開示によって提供されるタッチコマンドの処理方法をコンピュータに実行させるために使用されるコンピュータ命令を記憶している。

【 0 0 7 7 】

メモリ 8 0 2 は、非一時的なコンピュータ可読記憶媒体として、非一時的なソフトウェアプログラム、非一時的なコンピュータ実行可能プログラム及びモジュール、例えば、本開示の実施形態におけるビデオ処理方法に対応するプログラム命令／モジュール（例えば、図 7 に示す分割モジュール、符号化モジュール、ビデオ処理モジュール等のモジュール）を格納するために使用することができる。プロセッサ 8 0 1 は、メモリ 8 0 2 に記憶された非一時的なソフトウェアプログラム、命令及びモジュールを実行することにより、サーバの各種機能アプリケーション及びデータ処理を実行し、上述した方法の実施形態におけるビデオ処理方法を実現する。

【 0 0 7 8 】

メモリ 8 0 2 は、プログラム記憶領域とデータ記憶領域とを含んでもよく、プログラム記憶領域は、オペレーティングシステム、少なくとも 1 つの機能に必要なアプリケーションプログラムを格納してもよく、データ記憶領域は、タッチコマンドの処理方法の電子デバイスの使用により作成されたデータなどを格納してもよい。また、メモリ 8 0 2 は、高速ランダムアクセスメモリを含んでもよく、少なくとも 1 つのディスクメモリ装置、フラッシュメモリ装置又は他の非一時的なソリッドステートメモリ装置などの非一時的なメモリを更に含んでもよい。幾つかの実施形態では、メモリ 8 0 2 は、プロセッサ 8 0 1 に対して相対的に遠隔に配置されたメモリを含むことが好ましく、これらの遠隔メモリは、ネットワークを介して、タッチコマンドの処理方法の電子デバイスに接続されてもよい。前記ネットワークの例としては、インターネット、企業のイントラネット、ローカルエリアネットワーク、移動体通信ネットワーク及びそれらの組合せが挙げられるが、これらに限定されない。

【 0 0 7 9 】

ビデオ処理方法の電子デバイスは、入力装置 8 0 3 と出力装置 8 0 4 を更に含んでもよ

10

20

30

40

50

い。プロセッサ 801、メモリ 802、入力装置 803 及び出力装置 804 は、バスを介して接続されていてもよく、他の方式で接続されていてもよく、図 8 ではバスを介した接続を例に挙げている。

【0080】

入力装置 803 は、入力された数値情報又は文字情報を受信するとともにタッチコマンドの処理方法の電子デバイスのユーザ設定及び機能制御に関連するキー信号入力を生成してもよく、例えば、タッチスクリーン、キーパッド、マウス、トラックパッド、タッチパッド、インジケータスティック、1つ以上のマウスボタン、トラックボール、ジョイスティック、その他の入力装置などが挙げられる。出力装置 804 は、表示装置、補助照明装置（例えば、LED）、触覚フィードバック装置（例えば、振動モータ）等を含んでもよい。当該表示装置としては、液晶ディスプレイ（Liquid Crystal Display、LCD）、発光ダイオード（Light Emitting Diode、LED）、プラズマディスプレイ等が挙げられるが、これらに限定されない。幾つかの実施形態では、表示装置は、タッチスクリーンであってもよい。

10

【0081】

本明細書に記載のシステム及び技術の様々な実施形態は、デジタル電子回路システム、集積回路システム、専用集積回路（Application Specific Integrated Circuits、ASIC）、コンピュータハードウェア、ファームウェア、ソフトウェア、及び/又はそれらの組合せで実現することができる。これらの様々な実施形態は、以下を含み得る。1つ以上のコンピュータプログラムで実施し、当該1つ以上のコンピュータプログラムは、少なくとも1つのプログラマブルプロセッサを含むプログラマブルシステム上で実行及び/又は解釈され、当該プログラマブルプロセッサは、記憶システム、少なくとも1つの入力装置、及び少なくとも1つの出力装置からデータ及び命令を受信し、且つデータ及び命令を当該記憶システム、当該少なくとも1つの入力装置、及び当該少なくとも1つの出力装置へ転送することができる専用又は汎用のプログラマブルプロセッサであってもよい。

20

【0082】

これらのコンピュータプログラムは、プログラム、ソフトウェア、ソフトウェアアプリケーション、又はコードとも呼ばれ、プログラマブルプロセッサのための機械命令を含み、高レベル手順及び/又はオブジェクト指向のプログラミング言語、及び/又はアセンブリ/機械語を使用してこれらのコンピュータプログラムを実装することができる。本明細書で使用されるように、「機械可読媒体」及び「コンピュータ可読媒体」という用語は、機械命令及び/又はデータをプログラマブルプロセッサに提供するために使用される任意のコンピュータプログラム製品、デバイス、及び/又は装置、例えば、磁気ディスク、光ディスク、メモリ、プログラマブルロジックデバイス（programmable logic device、PLD）を指し、機械可読信号である機械命令を受け取る機械可読媒体を含む。「機械可読信号」という用語は、機械命令及び/又はデータをプログラマブルプロセッサに提供するために使用される任意の信号を指す。

30

【0083】

ユーザとのインタラクティブを提供するために、本明細書に記載されているシステム及び技術は、ユーザに情報を表示するための表示装置（例えば、CRT（Cathode Ray Tube、陰極線管）又はLCD（液晶ディスプレイ）モニター）と、ユーザがコンピュータに入力を提供するためのキーボード及びポインティング装置（例えば、マウス又はトラックボール）とを有するコンピュータ上に実装されてもよい。他の種類の装置もユーザとのインタラクティブを提供するためにも使用されてもよく、例えば、ユーザに提供されるフィードバックは、任意の形態の感覚フィードバック（例えば、視覚フィードバック、聴覚フィードバック、又は触覚フィードバック）であってもよく、ユーザからの入力、任意の形態（音響入力、音声入力、又は触覚入力を含む）で受信されてもよい。

40

【0084】

本明細書に記載されているシステム及び技術は、バックエンドコンポーネントを含むコ

50

ンピューティングシステム（例えば、データサーバー）、ミドルウェアコンポーネントを含むコンピューティングシステム（例えば、アプリケーションサーバー）、又はフロントエンドコンポーネントを含むコンピューティングシステム（例えば、グラフィカルユーザインターフェース又はウェブブラウザを備えたユーザコンピューター。当該グラフィカルユーザインターフェース又は当該ウェブブラウザを介して、ユーザはここで説明するシステム及び技術の実装とインタラクティブできる）、又はそのようなバックエンドコンポーネント、ミドルウェアコンポーネント、又はフロントエンドコンポーネントの任意の組合せを含むコンピューティングシステムで実装されてもよい。システムのコンポーネントは、任意の形態又は媒体のデジタルデータ通信（例えば、通信ネットワーク）を介して相互に接続されていてもよい。通信ネットワークの例としては、ローカルエリアネットワーク（Local Area Network、LAN）、ワイドエリアネットワーク（Wide Area Network、WAN）及びインターネット等がある。

10

【0085】

コンピュータシステムは、クライアントとサーバを含むことができる。クライアントとサーバは一般的に互いに遠隔地にあり、通常は、通信ネットワークを介してインタラクティブする。クライアント-サーバ関係は、対応するコンピュータ上で実行され、互いにクライアント-サーバ関係を有するコンピュータプログラムによって生成される。サーバは、クラウドコンピューティングサーバ又はクラウドホストとも呼ばれるクラウドサーバであっても良く、クラウドコンピューティングサービスシステムのホスト製品であり、従来の物理ホスト及び仮想プライベートサーバー（VPS）サービスにおける管理困難の問題及び業務拡大性が弱いという欠陥を解決する。サーバは、分散システムのサーバであっても良く、ブロックチェーンと組み合わせたサーバであっても良い。

20

【0086】

本開示によれば、複数の第1ビデオフレームを取得し、前記複数の第1ビデオフレームに対してきめ細かい分割を行い、複数の第2ビデオフレームを取得し、前記複数の第2ビデオフレームに関するマルチモーダル情報に基づき、前記複数の第2ビデオフレームに対して特徴符号化を行い、前記マルチモーダル情報の融合を特徴付けるための特徴融合情報を取得し、前記特徴融合情報に基づき、前記複数の第2ビデオフレームに対して類似度のマッチングを行い、類似度のマッチング結果に基づいてターゲットビデオを取得する。本開示によれば、当該マルチモーダル情報に基づいて特徴符号化を行うことができるので、より多くのビデオコンテンツの詳細を含む情報を取得することができ、類似度のマッチングがされた後、取得されたターゲットビデオがより正確であり、ビデオ分割の正確性を高めることができる。

30

【0087】

上述した処理の様々なプロセスを用い、順序を変えたり、ステップを追加又は削除したりすることができることが理解されるべきである。例えば、本開示に記載の各ステップは、並行して実行されてもよく、順次実行されてもよく、異なる順序で実行されてもよく、本開示に開示された技術案の所望の結果が達成される限り、限定されない。

【0088】

上記の具体的な実施形態は、本開示の保護範囲の制限を構成するものではない。設計要件及び他の要因に応じて、様々な変更、組み合わせ、サブ組み合わせ及び置換えが行われ得ることは、当業者によって理解されるべきである。本開示の要旨及び原則の範囲内で行われる如何なる修正、同等の代替、改良等は、すべて本開示の保護範囲に含まれる。

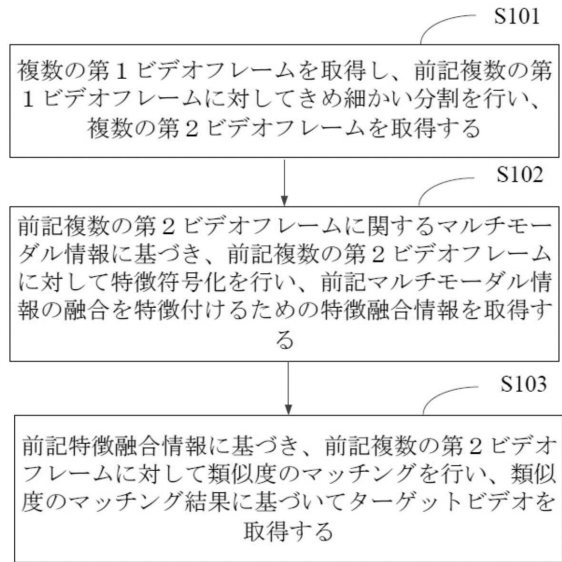
40

【図面】

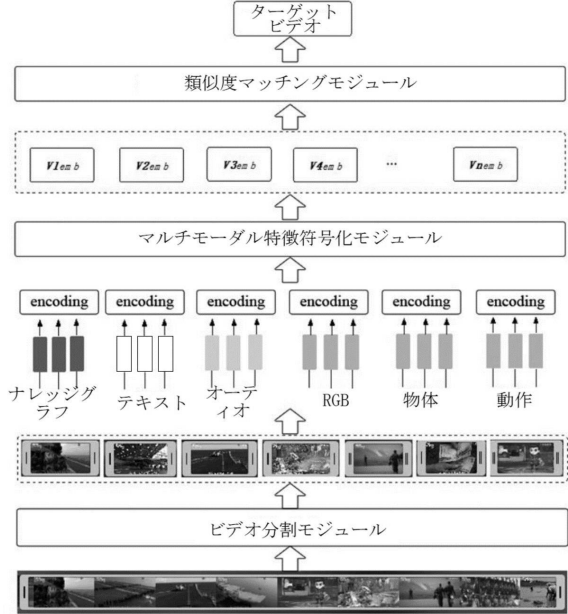
【図 1】



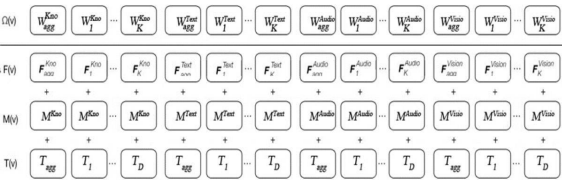
【図 2】



【図 3】



【図 4】



10

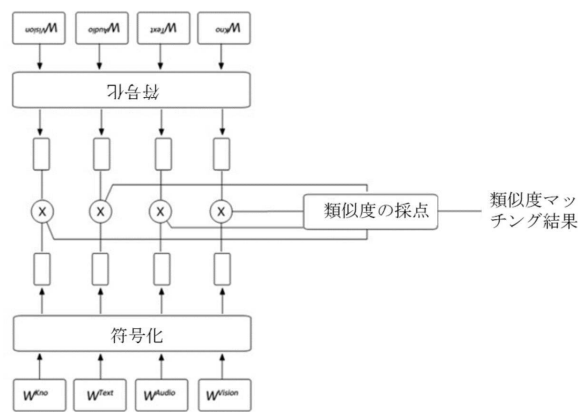
20

30

40

50

【図 5】

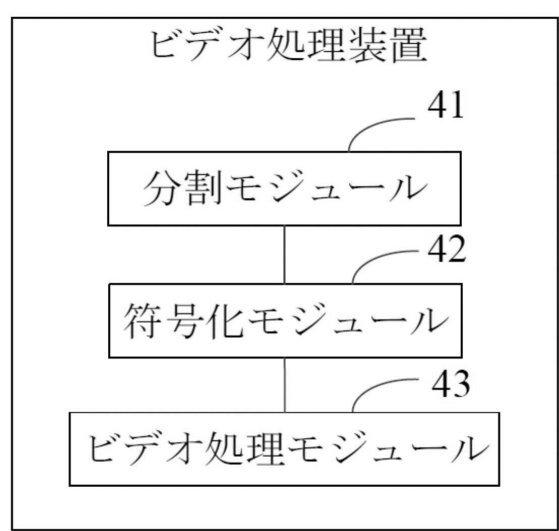


【図 6】

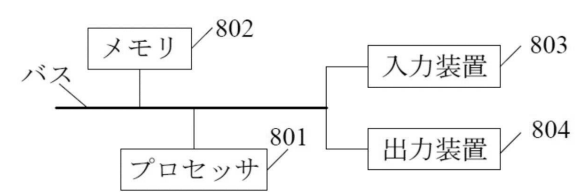


10

【図 7】



【図 8】



20

30

40

50

フロントページの続き

- 85, China
- (74)代理人 110001519
弁理士法人太陽国際特許事務所
- (72)発明者 チー ワン
中華人民共和国 ペキン 100085, ハイディアン ディストリクト, シャンディ テンス ス
トリート, 10番, バイドウ キャンパス 2階
- (72)発明者 チーファン フォン
中華人民共和国 ペキン 100085, ハイディアン ディストリクト, シャンディ テンス ス
トリート, 10番, バイドウ キャンパス 2階
- (72)発明者 フー ヤン
中華人民共和国 ペキン 100085, ハイディアン ディストリクト, シャンディ テンス ス
トリート, 10番, バイドウ キャンパス 2階
- (72)発明者 チュンコアン チャイ
中華人民共和国 ペキン 100085, ハイディアン ディストリクト, シャンディ テンス ス
トリート, 10番, バイドウ キャンパス 2階
- 審査官 山田 辰美
- (56)参考文献 特開平09-044639(JP, A)
特開2011-124681(JP, A)
国際公開第2007/052395(WO, A1)
丹野 良介, マルチモーダル深層学習によるドライブレコーダーデータの分類, 映像情報メ
ディア学会誌 第74巻 第1号, 日本, 一般社団法人映像情報メディア学会, 2020年, 第
74巻 第1号, p.44 - p.48
- (58)調査した分野 (Int.Cl., DB名)
G06T 7/00
G06F 16/783
G06F 16/732