

(19)日本国特許庁(JP)

(12)特許公報(B2)

(11)特許番号

特許第7378172号

(P7378172)

(45)発行日 令和5年11月13日(2023.11.13)

(24)登録日 令和5年11月2日(2023.11.2)

(51)国際特許分類

F I

H 0 4 N 21/8549(2011.01)

H 0 4 N 21/8549

G 0 6 V 10/82 (2022.01)

G 0 6 V 10/82

G 0 6 N 3/08 (2023.01)

G 0 6 N 3/08

G 0 6 T 7/00 (2017.01)

G 0 6 T 7/00

3 5 0 C

G 0 6 N 3/045(2023.01)

G 0 6 N 3/045

請求項の数 8 (全24頁)

(21)出願番号 特願2022-72173(P2022-72173)

(22)出願日 令和4年4月26日(2022.4.26)

(65)公開番号 特開2023-129179(P2023-129179
A)

(43)公開日 令和5年9月14日(2023.9.14)

審査請求日 令和4年4月26日(2022.4.26)

(31)優先権主張番号 10-2022-0026671

(32)優先日 令和4年3月2日(2022.3.2)

(33)優先権主張国・地域又は機関
韓国(KR)

(73)特許権者 505224569

インハ インダストリー パートナーシッ
プ インスティテュートInha - Industry Part
nership Institute大韓民国 2 2 2 1 2 インチョン ミチ
ヨル - グ インハ - ロ 1 0 0

(74)代理人 110001519

弁理士法人太陽国際特許事務所

(72)発明者 チョ・グンシク

大韓民国 2 2 2 1 2 インチョン ミチ
ヨル - グ インハ - ロ 1 0 0

(72)発明者 ユン・ウィニョン

大韓民国 2 2 2 1 2 インチョン ミチ
ヨル - グ インハ - ロ 1 0 0

最終頁に続く

(54)【発明の名称】 効率的なキーフレーム選択報酬関数を備えた教師なし映像要約方法および装置

(57)【特許請求の範囲】

【請求項1】

フレームレベル映像特徴抽出モジュールによって入力映像からフレームレベル視覚的特徴を抽出する段階、

アテンション基盤の映像要約ネットワークモジュールによってアテンション加重値を計算し、前記アテンション加重値を利用してキーフレームを選択するためのフレーム追跡確率として重要度点数を示す段階、

評価モジュールによって前記キーフレームを選択するためのキーフレーム同士の視覚的類似性距離と時間的距離による時間的一貫性報酬関数と代表性報酬関数を求め、時間的一貫性報酬関数と代表性報酬関数を利用してアテンション基盤の映像要約ネットワークが映像要約のキーフレームを選択するための重要度点数を予測するように学習する段階、

前記予測された重要度点数によって該当のキーフレームを選択して映像要約を生成し、生成された映像要約の品質を評価し、方策勾配アルゴリズム基盤の学習モジュールによって前記アテンション基盤の映像要約ネットワークに対する方策勾配(Policy Gradient)学習を実行する段階、

前記選択されたキーフレームの重要度点数を用いて前記映像要約ネットワークモジュールによってキーフレームを選択する確率を制御するための正規化および再構成損失を計算する段階、および

前記計算された正規化および再構成損失に基づいて映像要約生成モジュールによって映像要約を生成する段階

10

20

を含む、アテンション基盤の映像要約方法。

【請求項 2】

前記アテンション基盤の映像要約ネットワークモジュールによってアテンション加重値を計算し、前記アテンション加重値を利用してキーフレームを選択するためのフレーム追跡確率として重要度点数を示す段階は、

エンコーダネットワーク、デコーダネットワーク、および前記エンコーダネットワークとデコーダネットワークの間のアテンション層において、拡張 R N N によってパラメータと計算を減少させて時間依存性を抽出し、

前記エンコーダネットワークは、キーフレーム同士のローカルおよびグローバルコンテキストと視覚的類似性をキャプチャし、

前記アテンション層では、エンコーダネットワークの出力と以前のデコーダネットワークの隠れ状態の両方を用いてアテンション加重値を計算し、

前記アテンション加重値は、ソフトマックス関数に基づいて各キーフレームの確率点数に正規化し、

前記アテンション加重値を用いてエンコーダネットワークの出力に掛けてコンテキストベクトルを求め、

前記デコーダネットワークの入力のためにコンテキストベクトルと初期化されたデコーダネットワークの以前の出力を連結して前記デコーダネットワークを学習し、前記デコーダネットワークおよび前記エンコーダネットワークの学習結果を利用して重要度点数を求める

請求項 1 に記載のアテンション基盤の映像要約方法。

【請求項 3】

前記評価モジュールによってキーフレームを選択するためのキーフレーム同士の視覚的類似性距離と時間的距離による時間的一貫性報酬関数と代表性報酬関数を求め、時間的一貫性報酬関数と代表性報酬関数を利用してアテンション基盤の映像要約ネットワークが映像要約のキーフレームを選択するための重要度点数を予測するように学習する段階は、

前記代表性報酬関数を利用して抽出された特徴を用いて選択したキーフレームと映像のすべてのキーフレームの類似性を計算し、前記代表性報酬関数を利用して映像要約のキーフレームを選択するための重要度点数を予測するように学習し、

前記時間的一貫性報酬関数を利用して代表的なショットレベルキーフレームを選択するために、すべてのキーフレームに対して選択されたキーフレームのうちから最も近い隣りを見つけ出す過程を繰り返し学習する

請求項 1 に記載のアテンション基盤の映像要約方法。

【請求項 4】

前記予測された重要度点数によって該当のキーフレームを選択して映像要約を生成し、生成された映像要約の品質を評価し、方策勾配アルゴリズム基盤の学習モジュールによって前記アテンション基盤の映像要約ネットワークに対する方策勾配 (P o l i c y G r a d i e n t) 学習を実行する段階は、

過小評価報酬 (U n d e r - a p p r e c i a t e d R e w a r d : U R E X) 方法を探索する探索戦略の目的関数を使用して、目標関数の近似値のためにソフトマックス関数を用いて各エピソードに対する報酬の正規化された重要度加重値集合を計算することによってパラメータ化された方策勾配学習を実行する

請求項 1 に記載のアテンション基盤の映像要約方法。

【請求項 5】

入力映像からフレームレベル視覚的特徴を抽出するフレームレベル映像特徴抽出モジュール、

アテンション基盤の映像要約ネットワークによってアテンション加重値を計算し、前記アテンション加重値を利用してキーフレームを選択するためのフレーム追跡確率として重要度点数を示す映像要約ネットワークモジュール、

前記キーフレームを選択するためのキーフレーム同士の視覚的類似性距離と時間的距離

10

20

30

40

50

による時間的一貫性報酬関数と代表性報酬関数を求め、時間的一貫性報酬関数と代表性報酬関数を利用してアテンション基盤の映像要約ネットワークが映像要約のキーフレームを選択するための重要度点数を予測するように学習する評価モジュール、

前記予測された重要度点数によって該当のキーフレームを選択して映像要約を生成し、生成された映像要約の品質を評価し、前記アテンション基盤の映像要約ネットワークに対する方策勾配 (Policy Gradient) 学習を実行する方策勾配アルゴリズム基盤の学習モジュール (前記選択されたキーフレームの重要度点数を用いて前記映像要約ネットワークモジュールによってキーフレームを選択する確率を制御するための正規化および再構成損失を計算する)、および

前記計算された正規化および再構成損失に基づいて映像要約を生成する映像要約生成モジュール

10

を含む、アテンション基盤の映像要約装置。

【請求項 6】

前記映像要約ネットワークモジュールは、

エンコーダネットワーク、デコーダネットワーク、および前記エンコーダネットワークとデコーダネットワークの間のアテンション層において、拡張 RNN によってパラメータと計算を減少させて時間依存性を抽出し、

前記エンコーダネットワークは、キーフレーム同士のローカルおよびグローバルコンテキストと視覚的類似性をキャプチャし、

前記アテンション層では、エンコーダネットワークの出力と以前のデコーダネットワークの隠れ状態の両方を用いてアテンション加重値を計算し、

20

前記アテンション加重値は、ソフトマックス関数に基づいて各キーフレームの確率点数に正規化し、

前記アテンション加重値を用いてエンコーダネットワークの出力に掛けてコンテキストベクトルを求め、

前記デコーダネットワークの入力のためにコンテキストベクトルと初期化されたデコーダネットワークの以前の出力を連結して前記デコーダネットワークを学習し、前記デコーダネットワークおよび前記エンコーダネットワークの学習結果を利用して重要度点数を求める

請求項 5 に記載のアテンション基盤の映像要約装置。

30

【請求項 7】

前記評価モジュールは、

前記代表性報酬関数を利用して抽出された特徴を用いて選択したキーフレームと映像のすべてのキーフレームの類似性を計算し、前記代表性報酬関数を利用して映像要約のキーフレームを選択するための重要度点数を予測するように学習し、

前記時間的一貫性報酬関数を利用して代表的なショットレベルキーフレームを選択するために、すべてのキーフレームに対して選択されたキーフレームのうちから最も近い隣りを見つけ出す過程を繰り返し学習する

請求項 5 に記載のアテンション基盤の映像要約装置。

【請求項 8】

40

前記方策勾配アルゴリズム基盤の学習モジュールは、

過小評価報酬 (Under-appreciated Reward: UREX) 方法を探索する探索戦略の目的関数を使用し、目標関数の近似値のためにソフトマックス関数を用いて各エピソードに対する報酬の正規化された重要度加重値集合を計算することによってパラメータ化された方策勾配学習を実行する

請求項 5 に記載のアテンション基盤の映像要約装置。

【発明の詳細な説明】

【技術分野】

【0001】

本発明は、効率的なキーフレーム選択報酬関数を備えた教師なし映像要約方法および装

50

置に関する。

【背景技術】

【0002】

Y o u T u b e (登録商標)のようなオンライン動画共有プラットフォームを利用して興味のある映像を検索するために時間を費やす人が多くなった。この時間を短縮するためにプレビューまたは要約映像を使用することで、全体映像コンテンツを効率的かつ迅速に把握することができる(非特許文献1)。ここ数年間にわたって映像要約が重要視されるようになっており、映像コンテンツを検索したり、長い映像から短い形式の要約映像を製作したりするための研究が積極的に行われている。しかし、映像要約では、映像のフレームレベルまたはショットレベルの重要度点数を予測するための困難な作業が問題となっている(非特許文献2)。映像要約は、明らかな視聴覚パターンや意味の規則がない、抽象的かつ主観的なマルチモード作業であるためである。映像のフレームが興味深いものであるか有益なものであれば、フレームの重要度点数が高くなければならない。また、このような高い点数のフレームが、映像要約を製作するために選択される。

10

【0003】

最近では多様な方法が提案されているが、このような方法はディープラーニング(非特許文献3、4、5)を利用することで高い性能を発揮する。ディープラーニング基盤の映像要約方法は、教師あり学習基盤の方法と教師なし学習基盤の方法とに分けられる。教師あり学習基盤の方法の場合は、ラベルが指定されたデータセットを生成するのが極めて難しい。さらに、多様なドメインや場면을包括した大規模なデータセットを生成することも極めて難しい。このような理由により、教師なし映像要約方法の開発に力が注がれてきた。

20

【0004】

強化学習(R e i n f o r c e m e n t L e a r n i n g : R L)基盤の映像要約方法が従来技術(非特許文献6)で提案され、優れた成果を示した。特に、R Lを用いて深層ニューラルネットワークを学習させるために、報酬関数であるキーフレームを選択するための効率的かつ明示的な評価方法がある。さらに、深層ニューラルネットワークは、評価方法を利用して、代表性、多様性、統一性などの映像の多様な特徴を効率的に学習する。R Lを使用しながら、部分線形補間法を使用する従来技術(非特許文献3)でI n t e r - S U Mを提案した。補間法を用いることでネットワークの出力を減らし、高い分散問題を緩和して性能を改善した。しかし、多くの映像において、キーフレームが特定の場面からしか選択されなかったり、興味深いキーフレームが適切に選択されなかったりした。

30

【0005】

さらに、以前のR L基盤の映像要約方法にはいくつかの短所がある。一つ目に、深層ニューラルネットワークによっては視覚的および時間的コンテキストの捕捉が難しい。二つ目に、多くの方法は、キーフレームの時間的分布を考慮せず、キーフレーム同士の視覚的な差を計算してネットワークを学習させるために報酬関数または損失関数を使用する。したがって、一律的に映像を容易に理解できるようにキーフレームを選択して演出者のストーリーラインを生かすことのできる映像要約方法が求められている。

【先行技術文献】

【非特許文献】

40

【0006】

【文献】E j a z , N . ; M e h m o o d , I . ; B a i k , S . W . E f f i c i e n t v i s u a l a t t e n t i o n b a s e d f r a m e w o r k f o r e x t r a c t i n g k e y f r a m e s f r o m v i d e o s . J . I m a g e C o m m u n . 2 0 1 3 , 2 8 , 3 4 - 4 4 .

【文献】G y g l i , M . ; G r a b n e r , H . ; R i e m e n s c h n e i d e r , H . ; G o o l , L . V . C r e a t i n g s u m m a r i e s f r o m u s e r v i d e o s . I n P r o c e e d i n g s o f t h e E u r o p e a n C o n f e r e n c e o n C o m p u t e r V i s i o n (E C C V) ; S p r i n g e r , 2 0 1 5 , p p . 5 0 5 - 5 2 0 .

50

【文献】Yoon, U. N. ; Hong, M. D. ; Jo, G. S. Interp - SUM: Unsupervised Video Summarization with Piecewise Linear Interpolation. Sensors 2021. vol. 21, no. 13, 4562.

【文献】Apostolidis, E. ; Adamantidou, E. ; Metsai, A. ; Mezaris, V. ; Patras, I. Unsupervised Video Summarization via Attention-Driven Adversarial Learning. In International Conference on Multimedia Modeling (MMM); Springer : Daejeon, Korea, 5 - 8 January 2020, pp. 492 - 504.

10

【文献】Jung, Y. J. ; Cho, D. Y. ; Kim, D. H. ; Woo, S. H. ; Kweon, I. S. Discriminative feature learning for unsupervised video summarization. AAAI Conference on Artificial Intelligence, Honolulu, Hawaii, USA, 27 January - 1 February 2019, pp. 8537 - 8544.

【文献】Zhou, K. ; Qiao, Y. ; Xiang, T. Deep Reinforcement Learning for Unsupervised Video Summarization with Diversity-Representativeness Reward. AAAI Conf. Artif. Intell. 2018, 32, 7582 - 7589.

20

【文献】Song, Y. ; Vallmitjana, J. ; Stent, A. ; Jaimes, A. Tvsum: Summarizing web videos using titles. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, Boston, MA, USA, 7 - 12 June 2015, pp. 5179 - 5187.

【文献】Feng, L. ; Li, Z. ; Kuang, Z. ; Zhang, W. Extractive Video Summarizer with Memory Augmented Neural Networks. MM, Seoul, Republic of Korea, 22 - 26 October 2018, pp. 976 - 983.

30

【文献】Zhang, K. ; Chao, W. L. ; Sha, F. ; Grauman, K. Video Summarization with Long Short-term Memory. In Proceedings of the European Conference on Computer Vision (ECCV); Springer : Amsterdam, Netherlands, 2016; pp. 766 - 782.

【文献】Zhang, Y. ; Kampffmeyer, M. ; Zhao, X. ; Tan, M. DTR-GAN: Dilated Temporal Relational Adversarial Network for Video Summarization. In Proceedings of the ACM Turing Celebration Conference (ACM TURC), Shanghai, China, 18 May 2018.

40

【文献】Ji, Z. ; Xiong, K. ; Pang, Y. ; Li, X. Video Summarization with Attention-Based Encoder-Decoder Networks. IEEE Transactions on circuits and systems for video technology, June 2020, vol. 30, no. 6. pp. 1709 - 1717.

【文献】Mahasseni, B. ; Lam, M. ; Todorovic, S. Unsupervised Video Summarization with Adversa

50

rial LSTM Networks. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, Honolulu, Hawaii, USA, 22 - 25 July 2017, pp. 202 - 211.

【文献】Yuan, L.; Tay, F. E.; Li, P.; Zhou, L.; Feng, F. Cycle-SUM: Cycle-consistent Adversarial LSTM Networks for Unsupervised Video Summarization. In Proceedings of the Thirty-Third AAAI Conference on Artificial Intelligence, Honolulu, Hawaii, USA, 27 Jan - 1 Feb 2019, Volume 33, pp. 9143 - 9150.

10

【文献】Kaufman, D.; Levi, G.; Hassner, T.; Wolf, L. Temporal Tessellation: A Unified Approach for Video Analysis. In Proceedings of the IEEE International Conference on Computer Vision (ICCV), 2017, pp. 94 - 104.

【文献】Rochan, M.; Ye, L.; Wang, Y. Video Summarization Using Fully Convolutional Sequence Networks. In Proceedings of the European Conference on Computer Vision (ECCV); Springer: Munich, Germany, 2018, pp. 347 - 363.

20

【文献】Silver, D.; Lever, G.; Heess, N.; Degris, T.; Wierstra, D.; Riedmiller, M. Deterministic Policy Gradient Algorithms. In Proceedings of the 31st International Conference on Machine Learning (ICML), Beijing, China, 21 - 26 June 2014, pp. 387 - 395.

【文献】Yu, Y. Towards Sample Efficient Reinforcement Learning. In Proceedings of the Twenty-Seventh International Joint Conference on Artificial Intelligence (IJCAI), Stockholm, Sweden, 13 - 19 July 2018, pp. 5739 - 5743.

30

【文献】Lehnert, L.; Larocche, R.; Seijen, H. V. On Value Function Representation of Long Horizon Problems. In Proceedings of the Thirty-Second AAAI Conference on Artificial Intelligence, New Orleans, Louisiana, USA, 2 - 7 Feb 2018, pp. 3457 - 3465.

【文献】Szegedy, C.; Liu, W.; Jia, Y.; Sermanet, P.; Reed, S.; Anguelov, D.; Erhan, D.; Vanhoucke, B.; Rabinovich, A. Going deeper with convolutions. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR), Boston, MA, USA, 7 - 12 June 2015, pp. 1 - 9.

40

【文献】Luong, T.; Pham, H.; Manning, C. D.; Effective Approaches to Attention-based Neural Machine Translation. Proceedings of the 2015 Conference on Empirical Methods in Natu

50

ral Language Processing (EMNLP), Lisbon, Portugal, 17 - 21 September 2015, pp. 1412 - 1421.

【文献】Nachum, O.; Norouzi, M.; Schuurmans, D. Improving Policy Gradient by Exploring Under-Appreciated Rewards. arXiv 2016, arXiv:1611.09321.

【文献】Potapov, D.; Douze, M.; Harchaoui, Z.; Schmid, C.; Category-specific video summarization. European Conference on Computer Vision (ECCV), Zurich, Switzerland, Sep 2014, pp. 540 - 555.

10

【文献】Rochan, M.; Wang, Y. Video Summarization by Learning from Unpaired Data. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR) 2019, pp. 7902 - 7911.

【発明の概要】

【発明が解決しようとする課題】

【0007】

本発明が達成しようとする技術的課題は、新たな時間的一貫性報酬関数および代表性報酬関数 (TR-SUM) を備えた強化学習基盤の映像要約フレームワーク、および重要度点数を正確に予測するためのアテンション基盤の映像要約方法および装置を提案することにある。より詳細には、長い映像のコンテキストを効率的にキャプチャするために、アテンション基盤のエンコーダ-デコーダアーキテクチャをもつ映像要約ネットワークおよび関心のあるキーフレームを効率的かつ均一に選択するための時間的一貫性報酬関数および代表性報酬関数である新たな報酬関数を提案する。

20

【課題を解決するための手段】

【0008】

一側面において、本発明で提案する効率的なキーフレーム選択報酬関数を備えた教師なし映像要約方法は、フレームレベル映像特徴抽出モジュールによって入力映像からフレームレベル視覚的特徴を抽出する段階、アテンション基盤の映像要約ネットワークモジュールによってアテンション加重値を計算し、前記アテンション加重値を利用してキーフレームを選択するためのフレーム追跡確率として重要度点数を示す段階、評価モジュールによって前記キーフレームを選択するためのキーフレーム同士の視覚的類似性距離と時間的距離による時間的一貫性報酬関数と代表性報酬関数を求め、時間的一貫性報酬関数と代表性報酬関数を利用してアテンション基盤の映像要約ネットワークが映像要約のキーフレームを選択するための重要度点数を予測するように学習する段階、前記予測された重要度点数によって該当のキーフレームを選択して映像要約を生成し、生成された映像要約の品質を評価し、方策勾配アルゴリズム基盤の学習モジュールによって前記アテンション基盤の映像要約ネットワークに対する方策勾配 (Policy Gradient) 学習を実行する段階、前記選択されたキーフレームの重要度点数を利用して前記映像要約ネットワークモジュールによってキーフレームを選択する確率を制御するための正規化および再構成損失を計算する段階、および前記計算された正規化および再構成損失に基づいて映像要約生成モジュールによって映像要約を生成する段階を含む。

30

40

【0009】

前記アテンション基盤の映像要約ネットワークモジュールによってアテンション加重値を計算し、前記アテンション加重値を利用してキーフレームを選択するためのフレーム追跡確率として重要度点数を示す段階は、エンコーダネットワーク、デコーダネットワーク、および前記エンコーダネットワークとデコーダネットワークの間のアテンション層において、拡張RNNによってパラメータと計算を減少させて時間依存性を抽出し、前記エン

50

コードネットワークは、キーフレーム同士のローカルおよびグローバルコンテキストと視覚的類似性をキャプチャし、前記アテンション層ではエンコードネットワークの出力と以前のデコードネットワークの隠れ状態の両方を用いてアテンション加重値を計算し、前記アテンション加重値は、ソフトマックス関数に基づいて各キーフレームの確率点数に正規化し、前記アテンション加重値を用いてエンコードネットワークの出力に掛けてコンテキストベクトルを求め、前記デコードネットワークの入力のためにコンテキストベクトルと初期化されたデコードネットワークの以前の出力を連結して前記デコードネットワークを学習し、前記デコードネットワークおよび前記前記エンコードネットワークの学習結果を利用して重要度点数を求める。

【0010】

前記評価モジュールによってキーフレームを選択するためのキーフレーム同士の視覚的類似性距離と時間的距離による時間的一貫性報酬関数と代表性報酬関数を求め、時間的一貫性報酬関数と代表性報酬関数を利用してアテンション基盤の映像要約ネットワークが映像要約のキーフレームを選択するための重要度点数を予測するように学習する段階は、前記代表性報酬関数を利用して抽出された特徴を用いて選択したキーフレームと映像のすべてのキーフレーム同士の類似性を計算し、前記代表性報酬関数に基づいて映像要約のキーフレームを選択するための重要度点数を予測するように学習し、前記時間的一貫性報酬関数を利用して代表的なショットレベルキーフレームを選択するために、すべてのキーフレームに対して選択されたキーフレームのうちから最も近い隣りを見つけ出す過程を繰り返し学習する。

【0011】

前記予測された重要度点数によって該当のキーフレームを選択して映像要約を生成し、生成された映像要約の品質を評価し、方策勾配アルゴリズム基盤の学習モジュールによって前記アテンション基盤の映像要約ネットワークに対する方策勾配 (Policy Gradient) 学習を実行する段階は、過小評価報酬 (Under-appreciated Reward: UREX) 方法を探る探索戦略の目的関数を使用し、目標関数の近似値のためにソフトマックス関数を用いて各エピソードに対する報酬の正規化された重要度加重値集合を計算することによってパラメータ化された方策勾配学習を実行する。

【0012】

また他の側面において、本発明で提案する効率的なキーフレーム選択報酬関数を備えた教師なし映像要約装置は、入力映像からフレームレベル視覚的特徴を抽出するフレームレベル映像特徴抽出モジュール、アテンション基盤の映像要約ネットワークによってアテンション加重値を計算し、前記アテンション加重値を利用してキーフレームを選択するためのフレーム追跡確率として重要度点数を示す映像要約ネットワークモジュール、前記キーフレームを選択するためのキーフレーム同士の視覚的類似性距離と時間的距離による時間的一貫性報酬関数と代表性報酬関数を求め、時間的一貫性報酬関数と代表性報酬関数を利用してアテンション基盤の映像要約ネットワークが映像要約のキーフレームを選択するための重要度点数を予測するように学習する評価モジュール、前記予測された重要度点数によって該当のキーフレームを選択して映像要約を生成し、生成された映像要約の品質を評価し、前記アテンション基盤の映像要約ネットワークに対する方策勾配 (Policy Gradient) 学習を実行する方策勾配アルゴリズム基盤の学習モジュール (前記選択されたキーフレームの重要度点数を使用して、前記映像要約ネットワークモジュールによってキーフレームを選択する確率を制御するための正規化および再構成損失を計算する)、前記計算された正規化および再構成損失に基づいて映像要約を生成する映像要約生成モジュール映像要約生成モジュールを含む。

【発明の効果】

【0013】

本発明の実施形態によると、新たな時間的一貫性報酬関数および代表性報酬関数 (TR-SUM) を備えた強化学習基盤の映像要約フレームワークおよび重要度点数を正確に予測するためのアテンション基盤の映像要約方法および装置を提案する。提案する教師なし

10

20

30

40

50

映像要約方法は、アテンション基盤のエンコーダ - デコーダアーキテクチャをもつ映像要約ネットワークによって長い映像のコンテキストを効率的にキャプチャすることができ、時間的一貫性報酬関数および代表性報酬関数である新たな報酬関数を利用することで関心のあるキーフレームを効率的かつ均一に選択することができる。

【図面の簡単な説明】

【 0 0 1 4 】

【図 1】本発明の一実施形態における、アテンション基盤の映像要約ネットワークの概念を説明するための図である。

【図 2】本発明の一実施形態における、時間一貫性報酬および代表性報酬関数を利用した強化学習基盤の映像要約フレームワークを示した図である。

10

【図 3】本発明の一実施形態における、効率的なキーフレーム選択報酬関数を備えた教師なし映像要約装置の構成を示した図である。

【図 4】本発明の一実施形態における、効率的なキーフレーム選択報酬関数を備えた教師なし映像要約方法を説明するためのフローチャートである。

【図 5】本発明の一実施形態における、アテンション加重値を計算するためのアテンション基盤の映像要約ネットワークの構成を示した図である。

【図 6】本発明の一実施形態における、代表性報酬関数を説明するための図である。

【図 7】本発明の一実施形態における、時間的一貫性報酬関数を説明するための図である。

【発明を実施するための形態】

【 0 0 1 5 】

20

本発明は、効率的なキーフレーム選択報酬関数を備えた教師なし映像要約方法を提案する。効率的なキーフレームの選択のために、新たな時間的一貫性報酬および代表性報酬関数 (TR - SUM) を備えた強化学習基盤の映像要約フレームワークを設計する。映像のキーフレームレベル重要度点数を予測するために、アテンション基盤のエンコーダ - デコーダアーキテクチャとして映像要約方法および装置を提案する。本発明の実施形態によると、重要度点数を使用して、関心のあるキーフレームを効率的かつ均一に選択するために役立つ報酬関数に基づいて報酬を計算する。以下、本発明の実施例について、添付の図面を参照しながら詳しく説明する。

【 0 0 1 6 】

図 1 は、本発明の一実施形態における、アテンション基盤の映像要約ネットワークの概念を説明するための図である。

30

【 0 0 1 7 】

本発明では、新たな時間的一貫性報酬および代表性報酬関数 (TR - SUM (video summarization framework with the new temporal consistency reward and representativeness reward functions)) を備えた強化学習基盤の映像要約フレームワークを提案し、図 1 に示すように、重要度点数を正確に予測するためのアテンション基盤の映像要約ネットワークを提案する。

【 0 0 1 8 】

提案する教師なし映像要約方法は、アテンション基盤のエンコーダ - デコーダアーキテクチャをもつ映像要約ネットワークによって長い映像のコンテキストを効率的にキャプチャすることができ、時間的一貫性報酬関数および代表性報酬関数である新たな報酬関数を利用することで関心のあるキーフレームを効率的かつ均一に選択することができる。

40

【 0 0 1 9 】

映像要約方法は、教師あり学習基盤の方法と教師なし学習基盤の方法とに分けられる。2つの方法すべて、多数の使用者が注釈をつけた映像のフレームレベルまたはショットレベルの重要度点数を含む映像要約データセットを使用する (非特許文献 2)。教師あり学習基盤の方法は、重要度点数を予測するための入力として映像のフレームレベルまたはショットレベル特徴を用いてモデルを学習する。この方法は、データセットを使用して、予測された重要度点数と注釈がついた重要度点数の差によって費用を計算する。また、この

50

ような方法は、最適のモデルを探索するための費用を最小化する。

【0020】

非特許文献8では、メモリ増強映像要約方法が提案された。メモリネットワークは、全体映像から抽出した支援知識を効率的に提供する。現在は、ショットの重要度点数を予測するために、グローバルアテンションメカニズムを用いて元の映像に対する全体的な理解によって点数を調整する。非特許文献9では、学習の代表性と多様性のために、フレームをサンプリングする確率をエンコードする決定的点過程(Determinantal Point Process: DPP)のあるLSTM基盤のネットワークを提示した。非特許文献10では、映像フレーム同士の時間的コンテキストの表現を強化するために、生成器の拡張時間関係(Dilated Temporal Relational: DTR)装置を提示した。映像の最上の要約を予測するためのネットワークを学習させるために、敵対的学習方法を3-プレイヤー(three-player)損失関数とともに使用する。非特許文献11では、キーショット選択のための重要度点数を予測するために、アテンション基盤のエンコーダ-デコーダネットワークが提案された。このネットワークは、双方向LSTMネットワークがあるエンコーダとアテンションメカニズムがあるデコーダを用いて映像表現を学習する。

10

【0021】

しかし、教師あり学習基盤の方法は、多様なドメインの映像を含んでいるため、人間がラベリングした映像要約データセットを生成することが極めて難しいという問題を抱えている。この反面、教師なし学習基盤の方法は、人間がラベリングしたデータセットを必要としない。

20

【0022】

非特許文献4で、アテンションオートエンコーダ(Attention Autoencoder: AAE)ネットワークは、SUM-GANで提案された敵対的オートエンコーダ(Adversarial Autoencoder: AAE)の学習の効率性と性能を高めるためにSUM-GANでVAE(Variational Autoencoder)に替わる。ネットワークを学習させる間、映像を要約するための興味深いフレームに加重値を付与する。非特許文献5では、映像の要約を適切に予測するための映像のローカルおよびグローバルコンテキストを効率的に学習するために、VAEとGAN(Generative Adversarial Networks)アーキテクチャを基盤としたCSNet(Chunk and Stride Network)を提案した。非特許文献12では、敵対的オートエンコーダ(AAE)基盤の映像要約モデルが提案された。選択器LSTMは、映像の入力フレームレベル特徴からフレームを選択する。この後、VAEは、選択したフレームを用いて再構成された映像を生成する。全体ネットワークを学習させるために、判別器は、再構成された映像と原本入力映像とを区別する。モデル学習には4つの損失関数を使用される。非特許文献13ではSUM-GANの変形であるCycle-SUMを提案したが、要約映像で原本映像の情報を保存するために2つのVAE基盤の生成器と2つの判別器がある周期敵対的生成網を採択した。非特許文献14で提案された、テッセレーション接近方式を使用する映像要約方法は、視覚的に類似するクリップを探索し、グラフ基盤の方法であるViterbialgorithmを使用して時間的一貫性を維持するクリップを選択する。非特許文献15では教師なし学習基盤のSUM-FCNを提案した。このような方法は、映像シーケンスを処理するために、空間畳み込みで変換された時間畳み込みをもつ新たなFCNアーキテクチャを提示する。この方法は、デコーダの出力点数を使用してフレームを選択し、要約映像でフレームの多様性を適用するためにリペリング(repelling)正規化器によって損失関数を計算する。

30

40

【0023】

深層強化学習は、深層ニューラルネットワークを強化学習方法と結合する(非特許文献16)。方策勾配(Policy Gradient)方法は、モデルがない強化学習方法のうちの1つである。方策勾配方法は、方策を深層ニューラルネットワークモデルにパラメータ化し、確率的勾配降下法(Stochastic Gradient Desce

50

n t : S G D) のような勾配降下方法を用いて、方策によって定義された状態分布に対する報酬を極大化してモデルを最適化する。モデルを学習させるために、この方法は、目的関数として費用を計算して最小化する。しかし、方策勾配方法には、低いサンプル効率性問題（非特許文献 17）と高い分散のようないくつかの問題を抱えている。特に、低いサンプル効率性の問題は、エージェントが人間ほど知能的でないため、環境（状態）において学習行動のための人間の経験のようなサンプルを人間よりも遥かに多く必要とするためである。他の問題は、推定された勾配の高分散である。この問題は、長い水平問題と高い次元の作用空間によって発生する（非特許文献 18）。長い水平の問題は、目標を達成するための長い一連の決定に対する極めて遅延した報酬から始まる。提案する方法では、分散を減らすために基準とともに方策勾配を使用し、サンプル効率性の問題を緩和するためにエピソード数を増やす。

10

【0024】

図 2 は、本発明の一実施形態における、時間一貫性報酬および代表性報酬関数を利用した強化学習基盤の映像要約フレームワークを示した図である。

【0025】

本発明では、映像要約ネットワークが予測した重要度点数（Importance Score）を用いて映像要約の問題をフレーム選択の問題として公式化する。特に、拡張された（dilated）GRU エンコーダとアテンションメカニズム（attention mechanism）がある GRU デコーダネットワークを用いてネットワークを開発する。このネットワークは、映像表現を学習し、重要度点数をフレーム選択確率によって効率的に予測する。重要度点数は、図 2 に示すようにベルヌーイ分布（Bernoulli Distribution）を用いることで、要約からキーフレームを選択するためのフレーム選択動作に変換される。

20

【0026】

図 3 は、本発明の一実施形態における、効率的なキーフレーム選択報酬関数を備えた教師なし映像要約装置の構成を示した図である。

【0027】

提案する効率的なキーフレーム選択報酬関数を備えた教師なし映像要約装置は、フレームレベル映像特徴抽出モジュール 310、映像要約ネットワークモジュール 320、評価モジュール 330、方策勾配アルゴリズム基盤の学習モジュール 340、および映像要約生成モジュール 350 を含む。

30

【0028】

本発明の実施形態に係るフレームレベル映像特徴抽出モジュール 310 は、入力映像からフレームレベル視覚的特徴を抽出し、アテンション基盤の映像要約ネットワークモジュール 320 によってアテンション加重値を計算し、前記アテンション加重値を利用してキーフレームを選択するためのフレーム追跡確率として重要度点数を示す。

【0029】

本発明の実施形態に係る映像要約ネットワークモジュール 320 は、エンコーダネットワーク、デコーダネットワーク、および前記エンコーダネットワークとデコーダネットワークの間のアテンション層において、拡張 RNN によってパラメータと計算を減少させて時間依存性を抽出する。

40

【0030】

本発明の実施形態に係る映像要約ネットワークモジュール 320 は、前記エンコーダネットワークはキーフレーム同士のローカルおよびグローバルコンテキストと視覚的類似性をキャプチャし、前記アテンション層ではエンコーダネットワークの出力と以前のデコーダネットワークの隠れ状態の両方を用いてアテンション加重値を計算する。

【0031】

本発明の実施形態に係る映像要約ネットワークモジュール 320 は、前記アテンション加重値はソフトマックス関数によって各キーフレームの確率点数に正規化し、前記アテンション加重値を用いてエンコーダネットワークの出力に掛けてコンテキストベクトルを求

50

める。

【 0 0 3 2 】

本発明の実施形態に係る映像要約ネットワークモジュール 3 2 0 は、前記デコーダネットワークの入力のためにコンテキストベクトルと初期化されたデコーダネットワークの以前の出力を連結して前記デコーダネットワークを学習し、前記デコーダネットワークおよび前記前記エンコーダネットワークの学習結果を利用して重要度点数を求める。

【 0 0 3 3 】

本発明の実施形態に係る評価モジュール 3 3 0 は、前記キーフレームを選択するためのキーフレーム同士の視覚的類似性距離と時間的距離による時間的一貫性報酬関数と代表性報酬関数を求め、時間的一貫性報酬関数と代表性報酬関数を利用してアテンション基盤の映像要約ネットワークが映像要約のキーフレームを選択するための重要度点数を予測するように学習する。

10

【 0 0 3 4 】

本発明の実施形態に係る評価モジュール 3 3 0 は、前記代表性報酬関数を利用して抽出された特徴を用いて選択したキーフレームと映像のすべてのキーフレーム同士の類似性を計算し、前記代表性報酬関数によって映像要約のキーフレームを選択するための重要度点数を予測するように学習する。

【 0 0 3 5 】

本発明の実施形態に係る評価モジュール 3 3 0 は、前記時間的一貫性報酬関数を利用して代表的なショットレベルキーフレームを選択するために、すべてのキーフレームに対して選択されたキーフレームのうちから最も近い隣りを見つけ出す過程を繰り返し学習する。

20

【 0 0 3 6 】

本発明の実施形態に係る方策勾配アルゴリズム基盤の学習モジュール 3 4 0 は、前記予測された重要度点数によって該当のキーフレームを選択して映像要約を生成し、生成された映像要約の品質を評価し、前記アテンション基盤の映像要約ネットワークに対する方策勾配 (Policy Gradient) 学習を実行する。

【 0 0 3 7 】

本発明の実施形態に係る方策勾配アルゴリズム基盤の学習モジュール 3 4 0 は、過小評価報酬 (Under - appreciated Reward : U R E X) 方法を探索する探索戦略の目的関数を使用し、目標関数の近似値のためにソフトマックス関数を用いて各エピソードに対する報酬の正規化された重要度加重値集合を計算することによってパラメータ化された方策勾配学習を実行する。

30

【 0 0 3 8 】

本発明の実施形態に係る映像要約ネットワークモジュール 3 2 0 は、前記選択されたキーフレームの重要度点数を使用して、キーフレームを選択する確率を制御するための正規化および再構成損失を計算する。本発明の実施形態に係る映像要約生成モジュール 3 5 0 は、前記計算された正規化および再構成損失に基づいて映像要約を生成する。

【 0 0 3 9 】

図 4 は、本発明の一実施形態における、効率的なキーフレーム選択報酬関数を備えた教師なし映像要約方法を説明するためのフローチャートである。

40

【 0 0 4 0 】

提案する、効率的なキーフレーム選択報酬関数を備えた教師なし映像要約方法は、フレームレベル映像特徴抽出モジュールによって入力映像からフレームレベル視覚的特徴を抽出する段階 4 1 0、アテンション基盤の映像要約ネットワークモジュールによってアテンション加重値を計算し、前記アテンション加重値を利用してキーフレームを選択するためのフレーム追跡確率として重要度点数を示す段階 4 2 0、評価モジュールによって前記キーフレームを選択するためのキーフレーム同士の視覚的類似性距離と時間的距離による時間的一貫性報酬関数と代表性報酬関数を求め、時間的一貫性報酬関数と代表性報酬関数を利用してアテンション基盤の映像要約ネットワークが映像要約のキーフレームを選択するための重要度点数を予測するように学習する段階 4 3 0、前記予測された重要度点数によ

50

って該当のキーフレームを選択して映像要約を生成し、生成された映像要約の品質を評価し、方策勾配アルゴリズム基盤の学習モジュールによって前記アテンション基盤の映像要約ネットワークに対する方策勾配 (Policy Gradient) 学習を実行する段階 440、前記選択されたキーフレームの重要度点数を用いて前記映像要約ネットワークモジュールによってキーフレームを選択する確率を制御するための正規化および再構成損失を計算する段階 420、および前記計算された正規化および再構成損失に基づいて映像要約生成モジュールによって映像要約を生成する段階 450 を含む。

【0041】

段階 410 で、フレームレベル映像特徴抽出モジュールによって入力映像からフレームレベル視覚的特徴を抽出する。

【0042】

段階 420 で、アテンション基盤の映像要約ネットワークモジュールによってアテンション加重値を計算し、前記アテンション加重値を利用してキーフレームを選択するためのフレーム追跡確率として重要度点数を示す。

【0043】

本発明の実施形態によると、エンコーダネットワーク、デコーダネットワーク、および前記エンコーダネットワークとデコーダネットワークの間のアテンション層において、拡張 RNN によってパラメータと計算を減少させて時間依存性を抽出する。

【0044】

前記エンコーダネットワークは、キーフレーム同士のローカルおよびグローバルコンテキストと視覚的類似性をキャプチャし、前記アテンション層ではエンコーダネットワークの出力と以前のデコーダネットワークの隠れ状態の両方を用いてアテンション加重値を計算する。

【0045】

前記アテンション加重値は、ソフトマックス関数に基づいて各キーフレームの確率点数に正規化し、前記アテンション加重値を用いてエンコーダネットワークの出力に掛けてコンテキストベクトルを求める。

【0046】

前記デコーダネットワークの入力のためにコンテキストベクトルと初期化されたデコーダネットワークの以前の出力を連結して前記デコーダネットワークを学習し、前記デコーダネットワークおよび前記前記エンコーダネットワークの学習結果を利用して重要度点数を求める。

【0047】

段階 430 で、評価モジュールによってキーフレームを選択するためのキーフレーム同士の視覚的類似性距離と時間的距離による時間的一貫性報酬関数と代表性報酬関数を求め、時間的一貫性報酬関数と代表性報酬関数を利用してアテンション基盤の映像要約ネットワークが映像要約のキーフレームを選択するための重要度点数を予測するように学習する。

【0048】

前記代表性報酬関数を利用して抽出された特徴を用いて選択したキーフレームと映像のすべてのキーフレームの類似性を計算し、前記代表性報酬関数によって映像要約のキーフレームを選択するための重要度点数を予測するように学習する。

【0049】

前記時間的一貫性報酬関数を利用して代表的なショットレベルキーフレームを選択するために、すべてのキーフレームに対して選択されたキーフレームのうちから最も近い隣りを見つけ出す過程を繰り返し学習する。

【0050】

段階 440 で、前記予測された重要度点数によって該当のキーフレームを選択して映像要約を生成し、生成された映像要約の品質を評価し、方策勾配アルゴリズム基盤の学習モジュールによって前記アテンション基盤の映像要約ネットワークに対する方策勾配 (Policy Gradient) 学習を実行する。

10

20

30

40

50

【 0 0 5 1 】

本発明の実施形態によると、過小評価報酬 (Under-appreciated Reward: UREX) 方法を探索する探索戦略の目的関数を使用し、目標関数の近似値のためにソフトマックス関数を用いて各エピソードに対する報酬の正規化された重要度加重値集合を計算することによってパラメータ化された方策勾配学習を実行する。

【 0 0 5 2 】

再び段階 4 2 0 で、前記選択されたキーフレームの重要度点数を用いて映像要約ネットワークモジュールによってキーフレームを選択する確率を制御するための正規化および再構成損失を計算し、段階 4 5 0 で、前記計算された正規化および再構成損失に基づいて映像要約生成モジュールによって映像要約を生成する。

10

【 0 0 5 3 】

図 5 ~ 7 を参照しながら、効率的なキーフレーム選択報酬関数を備えた教師なし映像要約方法についてさらに詳しく説明する。

【 0 0 5 4 】

図 5 は、本発明の一実施形態における、アテンション加重値を計算するためのアテンション基盤の映像要約ネットワークの構成を示した図である。

【 0 0 5 5 】

先ず、ImageNet データセットによって学習された強力な深層畳み込みニューラルネットワークである GoogleNet (非特許文献 19) を使用して、入力映像からフレームレベル視覚的特徴

20

$$\{x_t\}_{t=1}^N$$

を抽出する。特徴の抽出は、フレームイメージの視覚的特性を低次元特徴ベクトルで捕捉するのに重要となる。また、抽出された特徴は、映像のフレーム同士の視覚的な差を効率的に計算するのに役立つ。

【 0 0 5 6 】

本発明では、図 5 に示すように、キーフレームレベルの重要度点数を予測するためのアテンション基盤の映像要約ネットワークを提案する。アテンション基盤要約ネットワークは、エンコーダネットワーク 5 1 0、デコーダネットワーク 5 2 0、および 2 つのネットワークの間のアテンション層 (Attention Layer) で構成される。このネットワークは、SUM-GANAE 方法に対して提案されたアテンションオートエンコーダ (非特許文献 4) で LSTM ネットワークを拡張繰り返しニューラルネットワーク (Dilated RNN) に替え、GRU (Gated Recurrent Unit) ネットワークをデコーダに替えることでアテンションオートエンコーダを改善する。拡張 RNN は、より少ないパラメータと計算効率性の向上のために拡張スキップ連結とともに実現される。特に、ネットワークは、拡張された繰り返し層を積み、層全体にわたって拡張を幾何級数的に増加させることで複雑な時間依存性を抽出する。また、拡張 RNN の層として GRU セルを使用する。エンコーダネットワークは、キーフレーム同士のローカルおよびグローバルコンテキストと視覚的類似性をキャプチャする。アテンション層では、エンコーダネットワークの出力と以前のデコーダネットワークの隠れ状態の両方を用いてアテンション加重値を計算する。アテンション加重値を計算するために、非特許文献 20 で提案されたアテンションメカニズムを利用する。特に、数式 (1) で説明するように、エンコーダネットワーク E_{out} の出力とデコーダネットワーク h_{t-1} の以前の隠れ状態でアテンション加重値を計算するためにコンテンツ基盤の点数関数を使用する。 $t = 1$ であるときに、デコーダネットワーク h の出力を 0 に設定する。 W_a は、学習可能なパラメータであるアテンション加重値マトリックスである。

30

40

【 0 0 5 7 】

【 数 1 】

$$Score(E_{out}, h_{t-1}) = E_{out}^T W_a h_{t-1}$$

50

・・・(1)

【0058】

次に、加重値は、ソフトマックス関数に基づいて各キーフレームの確率点数に正規化される。アテンション加重値を用いてエンコーダネットワークの出力に掛けてコンテキストベクトルCを求める。

【0059】

デコーダネットワークD_{in}の入力のために、コンテキストベクトルと0に初期化されたデコーダネットワークの以前の出力を連結し、デコーダネットワークを学習させる。

【0060】

この後、出力D_{out}を取得し、エンコーダネットワークによって学習された豊富な情報を再使用するためにE_{out}と連結することで、長いシーケンスの映像に対する性能を高める。そして、線形関数とともに、次の段階t+1でデコーダネットワークに伝達するための特徴の大きさの次元を減らす。最後に、全結合層およびシグモイド関数を用いて出力の次元を減らし、重要度点数

$$S = \{s_t\}_{t=1}^N$$

で出力を生成する。すなわち、重要度点数は、映像の要約としてキーフレームを選択するためのフレーム追跡確率(0 to 1)である。

【0061】

図6は、本発明の一実施形態における、代表性報酬関数を説明するための図である。

【0062】

本発明の実施形態に係る効率的な映像要約のために非特許文献6の多様性報酬関数を採択しながら、効率的なキーフレーム選択のための2つ新たな報酬を提案する。提案された報酬関数は、キーフレーム同士の視覚的類似性距離と時間的距離を考慮した時間的一貫性報酬関数と代表性報酬関数である。

【0063】

多様性報酬関数R_{div}数式(2)は、抽出された特徴とともにフレーム選択作業によって選択されたキーフレーム同士の差を計算する。この報酬関数に基づき、ネットワークは、要約のキーフレームから多様なフレームを選択するための重要度点数を予測するように学習される。このようなキーフレームで構成された要約は、映像の内容を容易に把握できるようにする。映像のストーリーラインを維持しながら演算の複雑性を減らすために、選択されたキーフレーム同士の差を計算するための時間距離を20に制限する。この制限がなければ、フラッシュバック場面や類似する場面が選択したキーフレームと遠く離れていても、多様なフレームを選択するときこのような場面を無視することができる。

【0064】

選択されたキーフレームの指数を

$$J = \{i_k | a_{i_k} = 1, k = 1, 2, \dots, |J|\}$$

とする。多様性報酬関数は次のとおりとなる。

【0065】

【数2】

$$R_{div} = \frac{1}{J(J-1)} \sum_{t \in J} \sum_{\substack{t' \in J \\ t \neq t'}} \left(1 - \frac{x_t^T x_{t'}}{\|x_t\|_2 \|x_{t'}\|_2} \right)$$

・・・(2)

【0066】

代表性報酬関数R_{rep}数式(3)は、抽出された特徴を用いて選択したキーフレームと

映像のすべてのキーフレームの類似性を計算する。この報酬関数に基づいて、ネットワークは、映像を示す要約のキーフレームを選択するための重要度点数を予測するように学習される。キーフレームで構成された要約は、映像の主題を容易に把握できるようにする。本発明では、映像の要約を適切に生成して性能を高めるために、数式(3)で示す $D \times s_t$ のように、代表性報酬関数として重要度点数 S を適用する新たな技法を提示する。ネットワークを学習させるためには、代表性報酬関数を増やし、報酬関数で距離 D 数式(4)を最小化しなければならない。

【0067】

【数3】

$$R_{rep} = \exp\left(-\frac{1}{N} \sum_{t=1}^N D_t \times s_t\right)$$

10

・・・(3)

【0068】

【数4】

$$D_t = \min_{t' \in J} \|x_t - x_{t'}\|_2$$

・・・(4)

【0069】

20

図6を参照すると、Aは、選択されたキーフレームとすべてのキーフレームとの距離が長い例であり、Bは、選択されたキーフレームとすべてのキーフレームとの距離が短い例である。キーフレームの重要度点数 S が高ければ、要約としてキーフレームを選択することができる。この反面、キーフレームの重要度点数が低ければ、要約としてキーフレームが選択される確率は低くなる。AとBはすべて $D \times S$ の平均値であり、図6に示すように $\bar{A}$

は
 $\bar{B}$

30

よりも高いため、提案された報酬関数を基準にBよりもAに対する報酬がより低い。

【0070】

重要度点数を使用するトリックの効果を説明する3つの事例がある。

【0071】

1. キーフレームの距離 D が短い場合、重要度点数が変更しても報酬関数が高い報酬を返還する。

【0072】

2. キーフレームの距離 D が長くてキーフレームの重要度点数 S が高ければ、報酬関数が低い報酬を返還する。また、低い報酬でネットワークを学習させた後、選択を阻むために、ネットワークはキーフレームの重要度が低い点数を予測する。

40

【0073】

3. キーフレームの距離 D が長くてキーフレームの重要度点数 S が低ければ、報酬関数が中間報酬を返還する。しかし、キーフレームのほとんどは重要度点数が低いため、要約に選定されないであろう。

【0074】

図7は、本発明の一実施形態における、時間的一貫性報酬関数を説明するための図である。

【0075】

代表的なショットレベルキーフレームを効率的かつ均一に選択するために、時間的一貫性報酬関数 R_{con} 数式(6)が提案される。

50

【 0 0 7 6 】

図 7 に示すように、先ずは

$$V_i^{summary}$$

の

$$V_i^{3'}$$

のような選択されたキーフレームと

$$V_i^{all}$$

10

の他のフレームとの類似性を計算した後、自身を除いて最も類似するフレーム

$$V_i^3$$

を選択する。

【 0 0 7 7 】

時間的一貫性報酬について説明するために、図 7 の B に含まれたキーフレームのような
代表的なキーフレームが周辺フレームの周りに類似する場面をもつと定義する。この反面
、図 7 の A に含まれたキーフレームは、選択したキーフレームと一時的に距離が遠い類似
の場面をもっている。演出者が意図したストーリーラインにより、A に含まれたキーフレ
ームと類似する場面は少し後に登場することもあるし、一回だけ示されることもある。し
かし、A に含まれたキーフレームの問題は、これらで生成された要約が映像コンテンツに
対する使用者の理解を妨げるという点にある。したがって、要約からこのキーフレームを
取り除く。このようなキーフレームを取り除くことによる他の長所は、一側だけで過度に
選択される問題を防ぐことができるという点にある。すなわち、キーフレームを均一に選
択することができるのである。本発明では、すべてのキーフレームの数

20

$$|J|$$

30

まで、すべてのキーフレーム

$$\{x_t\}_{t=1}^N$$

から選択したキーフレーム

$$j_k$$

のうちから最も近い隣りを見つけ出す過程を繰り返す。本発明では、

$$i_k$$

40

と

$$j_k$$

との距離を最小化することによって要約の時間的一貫性を学習することができる。距離を
最小化するために、報酬関数を次のように計算する。

【 0 0 7 8 】

【数 5】

$$j_k = (\argmin_{k \in |J|} \|x_{i_k} - x\|)$$

50

・・・ (5)

【 0 0 7 9 】

【 数 6 】

$$R_{con} = \frac{1}{\log(\sum_{k=1}^J (i_k - j_k)^2 / |J|)}$$

・・・ (6)

【 0 0 8 0 】

このような報酬を正規化するために、距離を
|J|

10

で割ってログ確率を使用する。

【 0 0 8 1 】

キーフレームを要約として選択するために、重要度点数 S をフレーム追跡動作

$$A = \{a_t | a_t \in \{0, 1\}, t = 1, \dots, N\}$$

に変換するための離散確率分布であるベルヌーイ分布を使用する。フレームのフレーム選択動作が 1 であれば、このキーフレームを要約として選択する。ベルヌーイ分布は、フレーム選択動作の変形をランダムに生成するため、映像の多様な要約に対する探索を促進する。

20

【 0 0 8 2 】

【 数 7 】

$$A \sim \text{Bernoulli}(a_t; s_t) = \begin{cases} s_t, & \text{for } a_t = 1 \\ 1 - s_t, & \text{for } a_t = 0 \end{cases}$$

・・・ (7)

【 0 0 8 3 】

また、生成された要約の品質を報酬の合計で評価する。この報酬により、アテンション基盤の映像要約ネットワークを、方策勾配方法を用いてパラメータ化された方策で学習する。方策勾配は、勾配降下アルゴリズムを用いてより効率的な要約を得るために行動戦略を探索するための強化学習方法のうちの 1 つである。行動戦略の不足を防ぐために、過小評価報酬 (Under - appreciated Reward : UREX) 方法を探る探索戦略の目的関数を使用する (非特許文献 21)。方策下で動作

30

$$\pi_{\theta}(a_t | h_t)$$

のログ確率

$$\log \pi_{\theta}(a_t | h_t)$$

40

が報酬

$$r(a_t | h_t) = \frac{R_{rep} + R_{div}}{2} + R_{con}$$

を過小評価する場合、該当の動作は探索戦略によってより探求されるであろう。

【 0 0 8 4 】

$$\text{UREX} \propto_{\text{UREX}}$$

50

の目的関数を計算するために、エピソード

j

に動作と報酬のログ確率を維持する。

$\mathcal{O}_{UREX}$

は、RAML (Reward Augmented Maximum Likelihood) 目的関数の合計である報酬

$R(a_t | h_t)$

10

の期待値である。本発明では、RAML目標関数の近似値のために、ソフトマックス関数を用いて各エピソード j に対する報酬の正規化された重要度加重値集合を計算する。

【0085】

【数8】

$$\mathcal{O}_{UREX}(\theta; \tau) = \mathbb{E}_{h \sim p(h_t)} \left\{ \sum_{a \in \mathcal{A}} R(a_t | h_t) \right\}$$

・・・(8)

20

【0086】

本発明では、勾配推定値の分散を減らして計算効率性を高めるために、方策勾配に重要な技術である基準を用いる。基準は、今まで経験した報酬の移動平均で計算される。多様な映像の移動平均を用いて多様性を改善するために、以下のように各映像 b_1 に対する基準とすべての映像 b_2 に対する基準を合算して基準を計算する。最後に、ネットワークを学習させるための費用として L_{rwd} を極大化する。

【0087】

【数9】

$$\mathcal{B} = 0.7 \times b_1 + 0.3 \times b_2$$

30

・・・(9)

【0088】

【数10】

$$L_{rwd} = \mathcal{O}_{UREX}(\theta; \tau) - \mathcal{B}$$

・・・(10)

【0089】

本発明では、重要度点数を用いてキーフレームを選択する確率を制御するために、非特許文献6で提案された正規化ターム L_{reg} を使用する。重要度点数のほとんどが1か0に近ければ、誤ったキーフレームを要約して選択する確率が高まる。したがって、学習中に重要度点数を0.5に近くするとき L_{reg} を用いる。重要度点数が0.5に早く収斂されることを防ぐために、これに0.01を掛ける。

40

【0090】

【数11】

$$L_{reg} = 0.01 \times \left(\frac{1}{N} \times \sum_1^N s_t - 0.5 \right)^2$$

50

・・・ (1 1)

【 0 0 9 1 】

すべての損失関数を計算した後、映像要約 $L_{summary}$ に対する最終損失を計算して逆伝播を実行する。

【 0 0 9 2 】

【 数 1 2 】

$$L_{summary} = L_{reg} - L_{rwd}$$

・・・ (1 2)

【 0 0 9 3 】

アルゴリズム 1 は、方策勾配方法を用いた教師なし映像要約ネットワークの学習手順に関するものである。

【 0 0 9 4 】

Algorithm 1 Training the Network.

1: Input: x_t frame-level features of the video

2: Output: proposed network's parameters (θ)

3:

4: for the number of iterations do

5: $S \leftarrow \text{Network}(x_t)$ % Predict the importance score S

6: $A \leftarrow \text{Bernoulli Distribution}(S)$ % Action A from the score S

7: % Calculate the reward functions and the loss using A and S

8: $\{\theta\} \leftarrow -\nabla(L_{reg} - L_{rwd})$ % Minimization

9: % Update the network using the policy gradient method:

10: end for

【 0 0 9 5 】

本発明の実施形態によると、アテンション基盤の映像要約ネットワークをテストするために、ショット内のフレームレベルの重要度点数を平均化してショットレベルの重要度点数を計算する。他の方法の性能と比べるためには、映像のショットレベルの重要度点数が必要となる。ショットを感知するために、ショットの境界のような変更地点を感知する KTS (Kernel Temporal Segmentation) 方法を使用する (非特許文献 22)。要約を生成するために、映像の長さの上位 15% に該当する主要ショットを点数で整列して選択する。この段階は、非特許文献 6 で説明した要約映像の重要性を最大化するための '0 - 1 Knapsack 問題' と同じ概念である。

【 0 0 9 6 】

上述した装置は、ハードウェア構成要素、ソフトウェア構成要素、および/またはハードウェア構成要素とソフトウェア構成要素との組み合わせによって実現されてよい。例えば、実施形態で説明された装置および構成要素は、例えば、プロセッサ、コントローラ、ALU (arithmetic logic unit)、デジタル信号プロセッサ、マイクロコンピュータ、FPA (field programmable array)、PLU (programmable logic unit)、マイクロプロセッサ、または命令を実行して応答することができる様々な装置のように、1 つ以上の汎用コンピュータまたは特殊目的コンピュータを利用して実現されてよい。処理装置は、オペレーティングシステム (OS) および OS 上で実行される 1 つ以上のソフトウェアアプリケーションを実行してよい。また、処理装置は、ソフトウェアの実行にตอบสนองし、データにアクセスし、データを記録、操作、処理、および生成してもよい。理解の便宜のために、1 つの処理装置が使用されるとして説明される場合もあるが、当業者であれば、処理装置が複数の処理要素および/または複数種類の処理要素を含んでもよいことが理解できるであろう。例え

10

20

30

40

50

ば、処理装置は、複数個のプロセッサまたは1つのプロセッサおよび1つのコントローラを含んでよい。また、並列プロセッサのような、他の処理構成も可能である。

【0097】

ソフトウェアは、コンピュータプログラム、コード、命令、またはこれらのうちの1つ以上の組み合わせを含んでもよく、思うままに動作するように処理装置を構成したり、独立的または集散的に処理装置に命令したりしてよい。ソフトウェアおよび/またはデータは、処理装置に基づいて解釈されたり、処理装置に命令またはデータを提供したりするために、いかなる種類の機械、コンポーネント、物理装置、仮想装置(virtual equipment)、コンピュータ記録媒体または装置に具現化されてよい。ソフトウェアは、ネットワークによって接続されたコンピュータシステム上に分散され、分散された状態で記録されても実行されてもよい。ソフトウェアおよびデータは、1つ以上のコンピュータ読み取り可能な記録媒体に記録されてよい。

10

【0098】

実施形態に係る方法は、多様なコンピュータ手段によって実行可能なプログラム命令の形態で実現されてコンピュータ読み取り可能な媒体に記録されてよい。前記コンピュータ読み取り可能な媒体は、プログラム命令、データファイル、データ構造などを単独でまたは組み合わせて含んでよい。前記媒体に記録されるプログラム命令は、実施形態のために特別に設計されて構成されたものであってもよいし、コンピュータソフトウェアの当業者に公知な使用可能なものであってもよい。コンピュータ読み取り可能な記録媒体の例としては、ハードディスク、フロッピー(登録商標)ディスク、および磁気テープのような磁気媒体、CD-ROM、DVDのような光媒体、フロプティカルディスク(optical disk)のような光磁気媒体、およびROM、RAM、フラッシュメモリなどのようなプログラム命令を格納して実行するように特別に構成されたハードウェア装置が含まれる。プログラム命令の例は、コンパイラによって生成されるもののような機械語コードだけではなく、インタプリタなどを使用してコンピュータによって実行される高級言語コードを含む。

20

【0099】

以上のように、実施形態を、限定された実施形態および図面に基づいて説明したが、当業者であれば、上述した記載から多様な修正および変形が可能であろう。例えば、説明された技術が、説明された方法とは異なる順序で実行されたり、かつ/あるいは、説明されたシステム、構造、装置、回路などの構成要素が、説明された方法とは異なる形態で結合されたりまたは組み合わされたり、他の構成要素または均等物によって対置されたり置換されたとしても、適切な結果を達成することができる。

30

【0100】

したがって、異なる実施形態であっても、特許請求の範囲と均等なものであれば、添付される特許請求の範囲に属する。

【符号の説明】

【0101】

310：フレームレベル映像特徴抽出モジュール

320：映像要約ネットワークモジュール

40

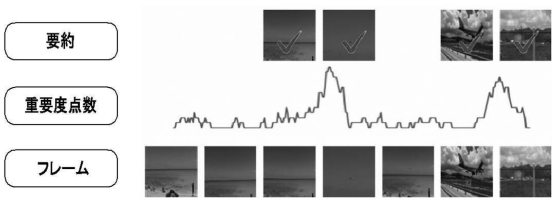
330：評価モジュール

340：方策勾配アルゴリズム基盤の学習モジュール

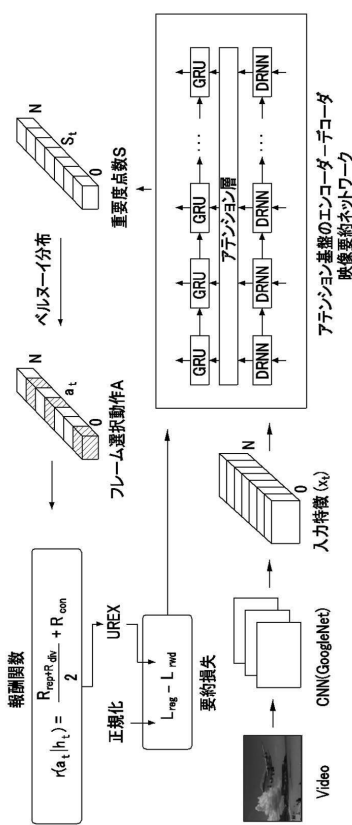
350：映像要約生成モジュール

【図面】

【図 1】



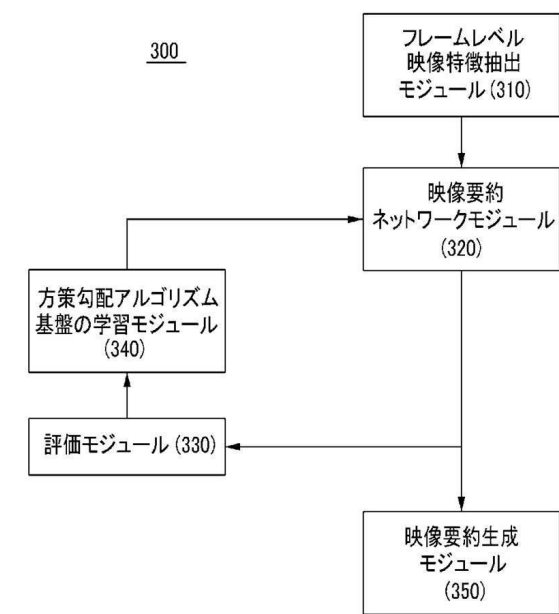
【図 2】



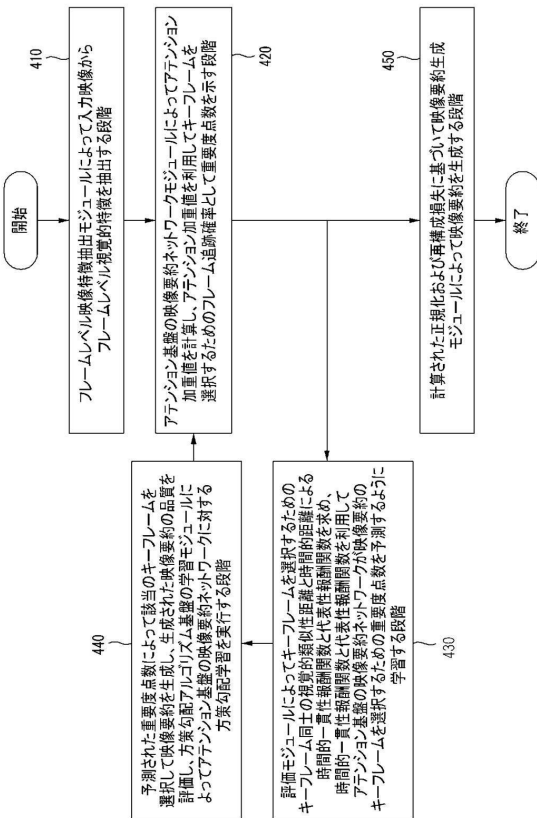
10

20

【図 3】



【図 4】

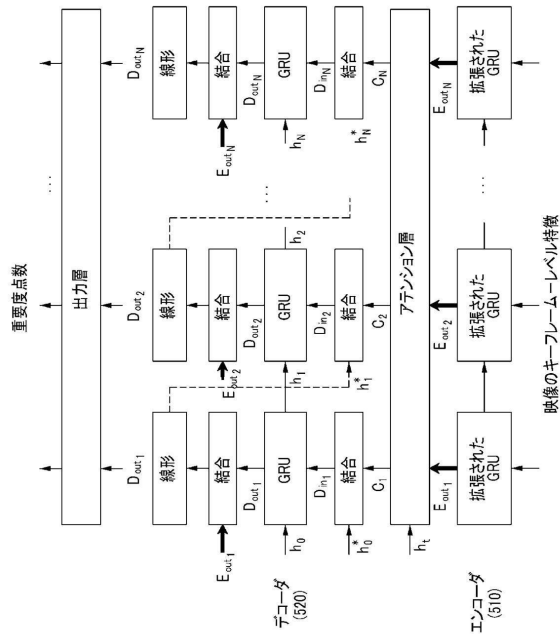


30

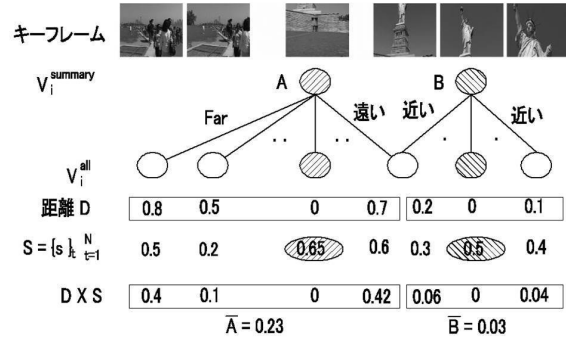
40

50

【図 5】



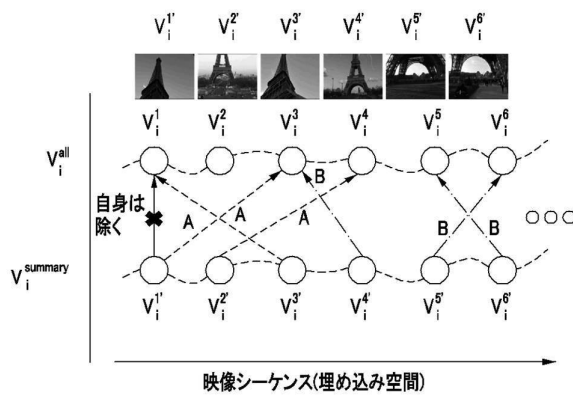
【図 6】



10

20

【図 7】



30

40

50

フロントページの続き

(72)発明者 ホン・ミョンドク

大韓民国 2 2 2 1 2 インチョン ミチョル - グ インハ - ロ 1 0 0

審査官 鈴木 隆夫

(56)参考文献 特開 2 0 2 2 - 1 6 9 0 0 9 (J P , A)

特開 2 0 2 1 - 0 6 0 8 7 4 (J P , A)

中国特許出願公開第 1 0 9 8 8 5 7 2 8 (C N , A)

Zhong Ji et al. , “ Deep Attentive Video Summarization With Distribution Consistency Learning ” , IEEE Transactions on Neural Networks and Learning Systems , Vol. 32 , No. 4 , 2020 年05月11日 , p.1765-1775 , DOI: 10.1109/TNNLS.2020.2991083

(58)調査した分野 (Int.Cl. , D B 名)

H 0 4 N 2 1 / 0 0 - 2 1 / 8 5 8

G 0 6 V 1 0 / 8 2

G 0 6 N 3 / 0 8

G 0 6 T 7 / 0 0

G 0 6 N 3 / 0 4 5