



(51) International Patent Classification:

H04N 13/275 (2018.01) H04N 13/31 (2018.01)

H04N 13/359 (2018.01) G06N 3/08 (2006.01)

(21) International Application Number:

PCT/KR2019/006025

(22) International Filing Date:

20 May 2019 (20.05.2019)

(25) Filing Language:

English

(26) Publication Language:

English

(30) Priority Data:

201810884174.7 06 August 2018 (06.08.2018) CN

(71) Applicant: SAMSUNG ELECTRONICS CO., LTD.

[KR/KR]; 129, Samsung-ro, Yeongtong-gu, Suwon-si, Gyeonggi-do 16677 (KR).

(72) Inventor: WANG, Lei; 5~12F, 6 Building, Chuqiao Cheng,

57# Andemen Street, Nanjing, Jiangsu 210012 (CN).

(74) Agent: LEE, Keon-Joo et al.; Mihwa Bldg., 16 Daehak-ro

9-gil, Chongro-gu, Seoul 03079 (KR).

(81) Designated States (unless otherwise indicated, for every

kind of national protection available): 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,

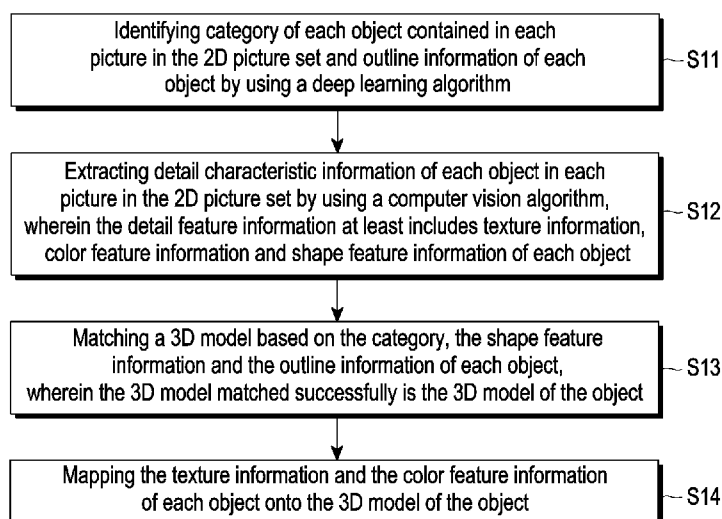
DZ, EC, EE, EG, ES, FI, GB, GD, GE, GH, GM, GT, HN, HR, HU, ID, IL, IN, IR, IS, JO, JP, KE, KG, KH, KN, KP, KR, KW, KZ, LA, LC, LK, LR, LS, LU, LY, MA, MD, ME, MG, MK, MN, MW, MX, MY, MZ, NA, NG, NI, NO, NZ, OM, PA, PE, PG, PH, PL, PT, QA, RO, RS, RU, RW, SA, SC, SD, SE, SG, SK, SL, SM, ST, SV, SY, TH, TJ, TM, TN, TR, TT, TZ, UA, UG, US, UZ, VC, VN, ZA, ZM, ZW.

(84) Designated States (unless otherwise indicated, for every kind of regional protection available): ARIPO (BW, GH, GM, KE, LR, LS, MW, MZ, NA, RW, SD, SL, ST, SZ, TZ, UG, ZM, ZW), Eurasian (AM, AZ, BY, KG, KZ, RU, TJ, TM), European (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).

Published:

— with international search report (Art. 21(3))

(54) Title: METHOD, STORAGE MEDIUM AND APPARATUS FOR CONVERTING 2D PICTURE SET TO 3D MODEL



(57) **Abstract:** Provided are a method for converting a 2D picture set to a 3D model. The method includes: identifying the category of each object contained in each picture in the 2D picture set and outline information of each object by using a deep learning algorithm; extracting detail characteristic information of each object by using a computer vision algorithm, wherein the detail feature information at least includes texture information, color feature information and shape feature information; matching the 3D model based on the category, the shape feature information and the outline information of each object, wherein the 3D model matched successfully is the 3D model of the object; and mapping the texture information and the color feature information of each object onto the 3D model of each object. According to the present disclosure, a realistic 3D model is constructed to overcome the shortcomings of generating a 3D model and a 3D video based on parallax.



Description

Title of Invention: METHOD, STORAGE MEDIUM AND APPARATUS FOR CONVERTING 2D PICTURE SET TO 3D MODEL

Technical Field

- [1] The present disclosure relates to the computer field, and in particular to a method, a storage medium and an apparatus for converting a 2D picture set to a 3D model.

Background Art

- [2] Currently, a 2D picture or video can be converted into a 3D model or a 3D video based on a parallax principle, which essentially generates two 2D pictures that are different for left and right eyes, not a substantial 3D model. The parallax is an illusion, therefore, a user may feel uncomfortable, distorted and easily fatigued when watching the 3D model or 3D video generated based on parallax, the user's experience can be poor, and entertainment and interestingness can be influenced.

Disclosure of Invention

Solution to Problem

- [3] Accordingly, the present disclosure provides a method, a storage medium and an apparatus for converting a 2D picture set to a 3D model, so as to solve the problem of how to construct a 3D model based on a 2D picture set.
- [4] The present disclosure provides a method for converting a 2D picture set to a 3D model, wherein, the 2D picture set includes at least one picture, and the method includes:
- [5] step 11, identifying the category of each object contained in each picture in the 2D picture set and outline information of each object by using a deep learning algorithm;
- [6] step 12, extracting detail characteristic information of each object in each picture in the 2D picture set by using a computer vision algorithm, wherein the detail feature information at least includes texture information, color feature information and shape feature information of each object;
- [7] step 13, matching a 3D model based on the category, the shape feature information and the outline information of each object, wherein the 3D model matched successfully is the 3D model of the object; and
- [8] step 14, mapping the texture information and the color feature information of each object matched successfully onto the 3D model of the object.
- [9] The present disclosure also provides a non-transitory computer-readable storage medium storing instructions that, when executed by a processor, cause the processor to

perform steps in the method for converting the 2D picture set to the 3D model described above.

[10] The present disclosure also provides an apparatus for converting a 2D picture set to a 3D model, including a processor and a non-transitory computer-readable storage medium described above.

[11] The present disclosure provides a method for converting a 2D picture set to a 3D model, image information of the object in the 2D picture is extracted, 3D model is matched, after the matching is successful, the texture information and the color feature information of the object extracted from the 2D picture are mapped onto the 3D model, so as to construct a realistic 3D model, which doesn't exist the defect of generating a 3D model and a 3D video based on parallax, therefore, the user's experience of the 3D video or 3D model can be improved, and the entertainment and interestingness can be enhanced.

Brief Description of Drawings

[12] FIG. 1 is a flowchart illustrating a method for converting a 2D picture set to a 3D model according to some examples of the present disclosure;

[13] FIG. 2 is an example of step 11 in FIG. 1;

[14] FIG. 3 is an example of step 12 in FIG. 1;

[15] FIG. 4 is an example of step 131 according to some examples of the present disclosure;

[16] FIG. 5 is an example of step 16 according to some examples of the present disclosure;

[17] FIG. 6 is a schematic diagram illustrating video decoding according to some examples of the present disclosure;

[18] FIG. 7 is a schematic diagram illustrating child extracting according to some examples of the present disclosure;

[19] FIG. 8 is a schematic diagram illustrating child posture synchronization according to some examples of the present disclosure; and

[20] FIG. 9 is a schematic diagram illustrating AR scene implementation according to some examples of the present disclosure.

Best Mode for Carrying out the Invention

[21] In order to make the object, technical solutions, and merits of the present disclosure clearer, the present disclosure will be illustrated in detail hereinafter with reference to the accompanying drawings and detailed examples.

[22] The present disclosure provides a method for converting a 2D picture set to a 3D model, as shown in FIG. 1, the 2D picture set includes at least one picture, and the method includes:

- [23] step 11 (S11), category of each object contained in each picture in the 2D picture set and outline information of each object is identified by using a deep learning algorithm.
- [24] The outline information includes not only periphery of each object, but also position information of the periphery, center point coordinates of the periphery, width and height of the periphery and the like.
- [25] The deep learning algorithm includes an unsupervised pre-trained network, a convolutional neural network, a recurrent neural network, a recursive neural network, and the like, and any one or combination of networks capable of identifying the category and outline information of the object from 2D pictures is applicable to the present disclosure.
- [26] For example, FIG. 2 is an implementation of step 11, after performing the method as shown in FIG. 2, each picture can obtain the category information and outline information of each object in the picture, and the method includes:
- [27] step 111, any picture in the 2D picture set is inputted into a convolutional neural network, to obtain an n-level feature map $P_1 \dots P_n$ of the picture, and $n \geq 2$.
- [28] The Convolutional Neural Network CNN (Convolutional Neural Network) model is generally used to do feature extracting work. The backbone network of the CNN includes residual network ResNeXt-101 and feature pyramid network FPN.
- [29] The residual network ResNeXt-101 is a simple, highly modular network structure for image classification, and is to extract features in CNN. The present disclosure also improves the network structure of ResNeXt-101, by using an acceleration strategy, and replacing 3x3 convolution of the ResNeXt-101 with a depth wise separable convolution, so that model miniaturization is achieved and n-level features $C_0 \dots C_{n-1}$ is output.
- [30] The feature pyramid network FPN, as an extension of ResNeXt-101, enables the entire CNN network to better characterize the target on multiple scales. And the performance of the pyramid extracting standard features is improved by adding a second pyramid, wherein, P_{n-1} is obtained by 1X1 convolution of C_{n-1} , P_{n-2} is obtained by 1X1 convolution of C_{n-2} plus sampling on P_{n-1} . P_{n-3} is obtained by 1X1 convolution of C_{n-3} plus sampling on P_{n-2} ; P_{n-4} is obtained by 1X1 convolution of C_{n-4} plus sampling on P_{n-3} . P_n is obtained by 1x1 max-pooling of P_{n-1} .
- [31] The output of a first pyramid from bottom layer to top layer is sequentially input to top layer to bottom layer of a second pyramid, for example, the second pyramid can select high-level features from the first pyramid and transfer them to the bottom layer. Based on this process, features at each level are allowed be combined with high-level and low-level features.
- [32] As a backbone network, ResNeXt101+FPN is used to extract features, and finally outputs feature map $P_1 \dots P_n$, in the FPN, the second pyramid has a feature map

containing features of each level, rather than a single backbone feature map in the standard backbone (i.e., a highest layer Cn-1 in the first pyramid), and the strategy for selecting features are as follows: selecting which level of features is dynamically determined by the target size.

[33] step 112, target proposal region in the P1...Pn is located by using a Region Proposal Network, wherein each proposal region includes at least one anchor box.

[34] The Region Proposal Network (RPN) takes any convolution feature map as an input, and outputs a proposal region of the convolution feature map, each proposal region contains a plurality of anchor boxes, which is similar to the step of Selective Search in the target detection.

[35] The RPN scans the feature map with a sliding window and looks for the presence of a target region, the region scanned is referred to as an anchor (also referred to as an anchor box), five specifications (32, 64, 128, 256, 512) of anchor are defined, and each specification has 3 ratios (1:1, 1:2, 2:1).

[36] RPN produces 2 outputs for each region: class cls (object or no object) of the anchor and bounding-box precision reg (change percentages in x, y, width and height). The sliding window adopts a special fully-connected layer with two branches for outputting the class and the precision of the anchor.

[37] The specific implementation is as follows: if a 512-dimensional fc feature is generated, mapping from a feature map to a first fully-connected feature is implemented by a convolutional layer Conv2D with Num_out = 512, kernel_size = 3X3, stride = 1, padding is same. Then, mapping from feature at a previous layer to features of two branches cls and reg is implemented by two convolutional layers Conv2D with Num_out = 2X15 (15 is the class of anchor 5X3) = 30 and Num_out = 5X15 = 75, respectively, kernel_size = 1X1, stride = 1, padding is valid.

[38] Anchor containing the target can be located with prediction of the RPN, and the position and the size of the anchor containing the target are finely adjusted.

[39] Step 113, when any of the proposal regions includes at least two anchor boxes, the anchor boxes of each proposal region may be screened by adopting a non-maximum suppression algorithm, to reserve the anchor box with the highest foreground score, and discard other anchor boxes.

[40] If the plurality of anchors selected by the PRN overlap with each other, non-maximum suppression can be adopted, to reserve the anchor with the highest foreground score, and discard the rest.

[41] Step 114, for each anchor box in the P1...Pn, the anchor box is divided into a pooling unit with a first preset size, extracted characteristic value of each subunit by using max-pooling, and then output pooled P1...Pn;

[42] step 115, the pooled P1...Pn is mapped to a fully-connected feature, the object

category of each anchor box is identified on the fully-connected feature and the size of the anchor box is reduced;

[43] In step 115, an object-based anchor box needs a fully-connected layer to identify the object category, and the fully-connected layer can only process input with a fixed size, however, the anchor boxes obtained in step 113 have different sizes. Step 114 is required to normalize the anchor box confirmed in step 113 into the first preset size, and the specific implementation includes:

[44] a. each anchor box may be traversed, to keep the boundary of a floating-point number performing no quantization.

[45] b. proposal region may be divided into $m \times m$ units, and boundary of each unit also does not perform quantization.

[46] c. 4 fixed coordinate positions may be calculated in each unit, and values of these 4 positions may be calculated by using a bilinear interpolation method, and then max-pooling operation is performed. The fixed position refers to a position determined by a fixed rule in each rectangular unit. For example, if the number of sampling points is 4, the unit is averaged into four small blocks, and then their respective center points are determined. It is apparent that the coordinates of these sample points are typically floating-point numbers, therefore, the pixel values thereof are obtained with an interpolated method.

[47] Step 115 mainly relates to a classification analysis algorithm and a regression analysis algorithm to obtain a classification of the anchor box, and a regression of an anchor bounding-box. As with the RPN, the classification analysis algorithm and regression analysis algorithm generates two outputs for each anchor box: category (specifically, category of object) and fine adjustment of bounding-box (further finely adjust the position and size of the anchor bounding-box).

[48] The specific implementation for the classification analysis algorithm and regression analysis algorithm is as follows:

[49] a. if a 1024-dimensional fc feature is generated, mapping from $P_1 \dots P_n$ to a first fully-connected feature is implemented by a convolutional layer with $\text{Num_out} = 1024$, $\text{kernel_size} = 3 \times 3$, $\text{stride} = 1$, padding is valid.

[50] b. The first fully-connected feature is followed by a BatchNormal, then relu is activated and then dropout, wherein, rate of dropout is 0.5.

[51] c. Then, another 1024-dimensional fc feature is output, mapping to a second fully-connected feature is implemented by a convolutional layer Conv2D with $\text{Num_out} = 1024$, $\text{kernel_size} = 1 \times 1$, $\text{stride} = 1$, padding is valid, a BatchNormal is followed, and then relu is activated.

[52] d. Finally, mapping from feature at a previous layer to features of two branches softmax classification (the region is classified into specific categories, such as person,

car, chair, and etc.) and linear regression (further finely adjust the position and size of the bounding-box) is implemented by two fully-connected layers with Num_out is 80 (category of class of the object identified) and $4 \times 80 = 320$ (position information multiplied by category of class), respectively.

[53] Step 116, outline information of objects in the reduced anchor box region is identified by using a full convolutional network.

[54] The Full Convolutional Network FCN may achieve pixel-wise classification (i.e., end to end, pixel-wise). The Full Convolutional Network FCN takes the positive region selected by the anchor box classification as an input, generates a mask thereof, segments pixels of different objects based on the mask and determines outline information of the object.

[55] For example, the FCN may be composed of four same convolutional layers Conv2D with num_out = 256, kernel_size = 3X3, stride = 1, padding is valid, and one deconvolutional layer (num_out = 256, kernel_size = 2X2, stride = 2), and then be mapped to a mask binarizing layer sigmoid with an output dimension of 80, Conv2D(80,(1,1),strides = 1,activation = "sigmoid").

[56] The mask generated is low in resolution: 14x14 pixels, a predicting mask is enlarged to the size of the anchor bounding-box to give a final mask result, and each object has one mask.

[57] Step 12, detail characteristic information of each object in each picture in the 2D picture set is extracted by using a computer vision algorithm; the detail feature information at least includes texture information, color feature information and shape feature information of each object.

[58] In addition to the texture information, the color feature information and the shaped characteristic information of the object, the detailed characteristic information of the present disclosure may also include: whether the 2D picture is a separate target frame of the object.

[59] In particular, step 12 is accomplished by step 121 and step 122, as shown in FIG. 3.

[60] step 121, locating objects in any picture in the 2D picture set by using a super-pixel and/or threshold segmentation method; if the picture includes only one object, then the picture is a separate target frame of the object;

[61] step 122, based on the locating information of each object in the picture, the texture of each object is extracted by using Tamura texture feature algorithm and wavelet transform, color feature information of each object is extracted by using a color histogram matching, and shape feature information of each object is extracted by using a geometric parameter method.

[62] Shape features of an object are calculated, for example, by using a shape parameter method (shape factor) relating to shape quantitative measurement (e.g., length, width,

moment, area, perimeter, etc.). If a plurality of objects are included, then the shape ratio between objects is also calculated.

[63] Step 13, matching the 3D model is performed based on the category, the shape feature information and the outline information of each object, and the 3D model matched successfully is the 3D model of the object.

[64] Matching the existing models in the 3D model library (3Dmax) may be performed according to the category, the shape feature information (such as, length, width, height), and the outline information of each object acquired by identification. The matching rule may be as follows, first match the category, match the shape feature information in the same category, after the matching of the shape feature information is completed, match the outline, to match models progressively and orderly.

[65] step 14, mapping the texture information and the color feature information of each object matched successfully onto the 3D model of the object is performed.

[66] The present disclosure provides a method for converting a 2D picture set to a 3D model, image information of the object in the 2D picture is extracted, 3D model is matched, after the matching is successful, the texture information and the color feature information of the object extracted from the 2D picture are mapped onto the 3D model, so as to construct a realistic 3D model, which doesn't exist the defect of generating a 3D model and a 3D video based on parallax, therefore, the user's experience of the 3D video or 3D model can be improved, and the entertainment and interestingness can be enhanced.

[67] In addition, the step 13 further includes: if the matching fails, then step 131 is performed;

[68] Step 131, constructing a 3D model of the object based on the separate target frame of the object matched unsuccessfully is performed.

[69] Specifically, as shown in FIG. 4, step 131 includes:

[70] Step 131-1, feature points of an object in the separate target frame may be extracted.

[71] Step 131-2, matching the feature points of the object in the separate target frame is performed to obtain registration points of the feature points.

[72] The separate target frame only includes information of one object, wherein, step 131 and step 132 may be implemented by using a Scale Invariant Feature Transform (Scale Invariant Feature Transform) algorithm, which is an excellent image matching algorithm with better robustness to rotation, scale and perspective. In some examples, other feature extraction algorithms may be considered, such as SURF, ORB, and the like.

[73] Registration points of a feature point also needs to be screened, for example, by using a Ratio Test method, 2 registration points that best match the feature point are sought by using the KNN algorithm, if the ratio between the matching distance of the first reg-

istration point and that of the second registration point is less than a certain threshold, then this match is accepted, otherwise, it is deemed to be as mismatching.

- [74] In some examples, a Cross Test method may be also used to screen the feature points and the registration points.
- [75] Step 131-3, an essential matrix of the separate target frame may be extracted based on the feature points and the registration points.
- [76] Step 131-4, intrinsic and extrinsic parameters of a camera may be solved based on the essential matrix.
- [77] Step 131-5, converting two-dimensional coordinates of the feature points and the registration points into three-dimensional coordinates may be performed based on the intrinsic and extrinsic parameters of the camera.
- [78] Step 131-6, judging whether there are other separate target frames of the object un-analyzed may be performed, if yes, returning to step 131-1, otherwise, step 131-7 is performed.
- [79] Step 131-7, a 3-dimensional point cloud may be generated based on three-dimensional coordinates of the feature points and the registration points, and geometric modeling of the 3D model of the object may be completed by using a Poisson curved surface reconstruction method based on the 3D point cloud.
- [80] Once the registration points are obtained, the essential matrix may be resolved by using the newly added function `findEssentialMat ()` in OpenCV3.0. After the essential matrix is obtained, another function `recoverPose` is used to decompose the essential matrix and return the relative transformation R and T , namely the intrinsic and extrinsic parameters of the camera, between two cameras, and complete calibration of the camera.
- [81] After the intrinsic and extrinsic parameters of the camera are solved, the two-dimensional coordinates of the feature points and the registration points are converted into three-dimensional coordinates, a sparse three-dimensional point cloud is generated, and then a dense point cloud is obtained by using PMVS2. There are a lot of the point cloud processing methods, and PMVS2 is only one of them.
- [82] Further, based on the 3D point cloud, geometric modeling of the 3D model of the object is completed by using the Poisson curved surface reconstruction method;
- [83] Step 131-8, refining the texture information, the color feature information and the shape feature information of the object may be performed based on the separate target frame, and the refined information may be mapped onto the 3D model of the object.
- [84] Finally, feature parameters of object-related feature information are refined, for example, body proportion, head feature, eyes, mouth, nose, eyebrows, facial contour in a person object, these parameters are acquired and then synchronously mapped onto a 3D model, and a real target model is reconstructed.

- [85] The method in FIG. 1 of the present disclosure may be also applied to a 2D video in addition to a 2D picture, in the case of 2D video, before step 11, the method further includes:
- [86] Step 10, key frames may be extracted in a 2D video as pictures in the 2D picture set.
- [87] Specifically, this step may be as follows:
- [88] Step 101, decoding the 2D video may be performed to acquire all static frames of the 2D video;
- [89] Step 102, clustering analysis may be performed on all the static frames to extract a static frame with the largest entropy in each cluster as a key frame of the cluster.
- [90] For example, one 1-minute video (1800 frame data) may obtain 30 key frames after performing steps 101 and 102 described above.
- [91] The detailed flow is as follows:
- [92] a. a video file is opened, data is read into a buffer from a hard disk, video file information is acquired, and a video file in the buffer is sent to a decoder for decoding.
- [93] b. the 2D video is decoded to acquire all static frames of the 2D video.
- [94] Decoders, including the FFMPEG, the MediaCodec of the Android platform, the AVFoundation of the IOS platform and the like, all can decode the 2D video to obtain a sequence of static frames of the 2D video.
- [95] c. The static frames are aggregated into n clusters by performing clustering analysis, the static frames in each cluster are similar, while static frames between different clusters are dissimilar.
- [96] If frame number in one cluster is too small, the frame is directly merged with the adjacent frame. For each cluster, a centroid is maintained, for each frame, the similarity of its cluster centroid is calculated. If the similarity is less than a certain threshold, then it is classified into a new cluster, otherwise it is added to the previous cluster.
- [97] d. One representative is extracted from each cluster as a key frame, for example, an image with the largest entropy in each cluster may be calculated and used as a key frame.
- [98] In some examples, after step 14 or step 131, the method further includes:
- [99] Step 15, a posture of any object in any picture in the 2D picture set may be identified, and the posture of the 3D model of the object may be adjusted to be consistent with the posture of the object.
- [100] Step 16, the adjusted 3D model may be rendered into an AR scene.
- [101] Assuming that a plurality of pictures are included in the 2D picture set, step 15 and step 16 may be performed one by one according to time information (e.g., a generation time) of the picture so that dynamic AR contents may be formed.
- [102] Furthermore, as shown in FIG. 5, step 16 includes:

- [103] Step 161, real scene information may be acquired.
- [104] Step 162, analyzing the real scene information and camera position information may be performed to obtain an affine transformation matrix of the 3D model projected on a camera view plane.
- [105] Step 163, merging the 3D model with a real scene video may be performed, the 3D model of the object is imported with the affine transformation matrix, and the 3D model and the real scene video may be displayed together on an AR presenting device or other device.
- [106] Merging video or displaying directly, that is, the graphics system first calculates an affine transformation matrix between the virtual object (3D model of the object) coordinates to the camera view plane according to the camera position information and the positioning marks in the real scene, and then draws the virtual object on the view plane according to the affine transformation matrix, and finally displays the 3D model and the real scene video together on an AR presenting device or other displays after directly merging with the real scene video.
- [107] When the key frames in the 2D video sequentially generate corresponding 3D models one by one, the corresponding 3D models are put into the VR environment one by one, and corresponding AR contents based on the 2D video may be generated.
- [108] What described above is description of a method for converting a 2D picture set to a 3D model of the present disclosure, and examples of applying the method of the present disclosure are set forth below.
- [109] Example 1
- [110] In some examples, a user takes a 2D video of a child playing with a mobile phone, the video may be converted to a piece of AR content based on the method of the present disclosure, and the user may directly watch the AR content to experience a sense of being personally on the scene. Specific operations are carried out as follows:
- [111] Step 201: parsing the 2D video is performed; the video is opened to acquire all static frames of the video, the static frames are analyzed to find and save a key frame, as shown in FIG. 6.
- [112] Step 202, identifying and extracting child and related feature information in the key frame may be performed by using a deep learning algorithm and a computer vision algorithm, as shown in FIG. 7.
- [113] Upon extracting, not setting the target object can be selected and a default object is adopted, and the default object includes person, car, chair, cup, bird, cow, cat, dog, sheep and etc. A specific extracting object may be also selected, for example, only target object of the person and the relevant characteristics are extracted. Alternatively, a target object framed manually and related characteristics may be also extracted.
- [114] Step 203: identifying the category and the feature information of the object may be

performed according to step 202, and retrieving and matching corresponding model category in the 3D model library may be performed, for example, in this example, the extracting object is a child (person category), first, 3D model matching the child is retrieved in the 3D model library.

[115] Then, according to the extracted child detail features such as eyes, mouth, nose, eyebrows, facial contour, textures and the like, these parameters are then synchronized to corresponding model to make the 3D model more vivid and realistic, and a 3D model of the child consistent with the key frame information is built.

[116] Step 204: according to the key frame obtained in the step 201 and the 3D model generated in the step 203, the posture of the model is adjusted. The posture of the 3D model may be adjusted to the posture in the key frame, and actions of the child in the video may be synchronized to the model, as shown in FIG. 8.

[117] Step 205: the action behaviors corresponding to the child model and the model are rendered into an AR scene, and displayed on the AR presenting device or other display, as shown in FIG. 9.

[118] Step 206: the AR content of the child playing is created successfully.

[119] Example 2

[120] In some examples, sometimes, a user may not watch the car exhibition at the scene for some reason, but only can watch ordinary videos of the car on the exhibition taken by his friend. The video may be converted into a piece of AR content based on the method of the present disclosure, so that the user may watch the car in a sense of being personally on the scene.

[121] Step 301: parsing the 2D video is performed; the video is opened to acquire all static frames of the video, the static frames are analyzed to find and save a key frame.

[122] Step 302, identifying and extracting car and related feature information in the key frame may be performed by using a deep learning algorithm and a computer vision algorithm.

[123] Upon extracting, not setting the target object can be selected and a default object is adopted, and the default object includes person, car, chair, cup, bird, cow, cat, dog, sheep and etc. A specific extracting object may be also selected, for example, only target object of the person and the relevant features are extracted. Alternatively, a target object framed manually and related features may be also extracted.

[124] Step 303: identifying the category and the feature information of the object may be performed according to step 202, and retrieving and matching corresponding model category in the 3D model library may be performed, for example, in this example, the extracting object is a car, first, 3D model matching the car is retrieved in the 3D model library.

[125] Then, according to the extracted car detail features such as shape, color, textures and

the like, these parameters are then synchronized to corresponding model to make the 3D model more vivid and realistic.

[126] Step 304: according to the key frame obtained in the step 301 and the 3D model generated in the step 303, the posture of the model is adjusted. The posture of the 3D model may be adjusted to the posture in the key frame, and various angles of watching the car in the video may be synchronized to the model.

[127] Step 305: the direction corresponding to the car model and the model are rendered into an AR scene, and displayed on the AR presenting device or other display.

[128] Step 306: the AR content of the car exhibition is created successfully.

[129] Example 3

[130] In some examples, a user often watches some 2D performance videos. The video may be converted into a piece of AR content based on the method of the present disclosure, so that the user or others may experience of watching on the scene in a sense of being personally on the scene.

[131] Step 401: parsing the 2D video is performed; the video is opened to acquire all static frames of the video, the static frames are analyzed to find and save a key frame.

[132] Step 402, identifying and extracting stage and related feature information in the key frame may be performed by using a deep learning algorithm and a computer vision algorithm.

[133] Upon extracting, not setting the target object can be selected and a default object is adopted, and the default object includes person, car, chair, cup, bird, cow, cat, dog, sheep and etc. A specific extracting object can also be selected to be set as a stage, for example, only target object of the stage and the relevant features are extracted. Alternatively, a target object framed manually and related features may be also extracted.

[134] Step 403: identifying the category and the feature information of the object may be performed according to step 202, and retrieving and matching corresponding model category in the 3D model library may be performed, for example, in this example, the extracting object is a stage, first, 3D model matching the stage is retrieved in the 3D model library.

[135] Then, according to the extracted stage detail features such person, seat, performance property and the like, these parameters are then synchronized to corresponding model to make the 3D model more vivid and realistic.

[136] Step 404: according to the key frame obtained in the step 401 and the 3D model generated in the step 403, the posture of the model is adjusted. The posture of the 3D model may be adjusted to the posture in the key frame, and actions in the video may be synchronized to the model.

[137] Step 405: the stage model and the direction corresponding to the stage model and the model are rendered into the AR scene, and displayed on the AR presenting device or

other display.

[138] Step 406: the AR content of the performance is created successfully.

[139] The present disclosure also provides a non-transitory computer-readable storage medium storing instructions that, when executed by a processor, cause the processor to perform any of the steps in the method for converting the 2D picture set to the 3D model described above.

[140] The present disclosure also provides an apparatus for converting a 2D picture set to a 3D model, including a processor and a non-transitory computer-readable storage medium described above.

[141] Specifically, an apparatus for converting a 2D picture set to a 3D model, the 2D picture set includes at least one picture, and the apparatus includes:

[142] an object category and outlineidentifying module, configured to identify the category of each object contained in each picture in the 2D picture set and outlineinformation of each object by using a deep learning algorithm;

[143] an object detail feature extracting module, configured to extract detail characteristic information of each object in each picture in the 2D picture set by using a computer vision algorithm, wherein the detail feature information at least includes texture information, color feature information and shape feature information of each object;

[144] a model matching module, configured to match a 3D model based on the category, the shape feature information and the outlineinformation of each object, wherein the 3D model matched successfullyis the 3D model of the object; and

[145] a 3D object refining module, configured to map the texture information and the color feature information of each object matched successfullyonto the 3D model of the object.

[146] In some examples, before the object category and outlineidentifying module, the apparatus furtherincludes:

[147] a key frame extracting module, configured to extract key frames in a 2D video as pictures in the 2D picture set.

[148] In some examples, the key frame extracting module includes:

[149] a video decoding module, configured to decode the 2D video to acquire all static frames of the 2D video;

[150] a clustering analyzing module, configured to perform clustering analysis on all the static frames to extract a static frame with the largest entropy in each cluster as a key frame of the cluster.

[151] In some examples, the object category and outlineidentifying module includes:

[152] a convolutional neural network, configured to input any picture in the 2D picture set into the convolutional neural network, wherein the convolutional neural network outputsan n-level feature map $P_1 \dots P_n$ of the picture, and $n \geq 2$;

- [153] a region construction network, configured to locate target proposal region in the $P1...Pn$ by using the Region Proposal Network, wherein each proposal region includes at least one anchor box;
- [154] an anchor box screening module, when any of the proposal regions includes at least two anchor boxes, configured to screen the anchor box of each proposal region by adopting a non-maximum suppression algorithm, reserve the anchor box with the highest foreground score, and discard other anchor boxes;
- [155] a pooling module, for each anchor box in the $P1...Pn$, configured to divide the anchor box into a pooling unit with a first preset size, extract the feature value of each subunit by using max-pooling, and then output the pooled $P1...Pn$;
- [156] a classification and regression module, configured to map the pooled $P1...Pn$ to a fully-connected feature, identify the object category of each anchor box on the fully-connected features and reduce the size of the anchor box; and
- [157] a full convolutional network, configured to identify outline information of objects in region of each reduced anchor box by using the full convolutional network.
- [158] In some examples, the convolutional neural network includes a residual network and a feature pyramid network, wherein 3×3 convolution of the residual network is replaced with a depth wise separable convolution.
- [159] In some examples, the feature pyramid network includes a first pyramid and a second pyramid, and the output of the first pyramid from bottom layer to top layer is sequentially input to top layer to bottom layer of the second pyramid.
- [160] In some examples, in the object detail feature extracting module, the detail feature also includes: whether the 2D picture is a separate target frame of an object.
- [161] Furthermore, the object detail feature extracting module includes:
- [162] an object locating module, configured to locate objects in any picture in the 2D picture set by using a superpixel and/or threshold segmentation method; if the picture includes only one object, then the picture is a separate target frame of the object; if the picture includes only one object, then the picture is a separate target frame of the object;
- [163] a detail feature parsing module, based on the locating information of each object in the picture, configured to extract the texture of each object by using Tamura texture feature algorithm and wavelet transform, extract color feature information of each object by using a color histogram matching, and extract shape feature information of each object by using a geometric parameter method.
- [164] In some examples, the model matching module further includes: if the matching fails, then a model constructing model is performed;
- [165] the model constructing model, configured to construct a 3D model of the object based on the separate target frame of the object matched unsuccessfully.

- [166] Furthermore, the model constructing model includes:
- [167] a feature point module, configured to extract feature points of the object in the separate target frame;
- [168] a registration point module, configured to match the feature points of the object in the separate target frame to obtain registration points of the feature points;
- [169] an essential matrix generating module, configured to extract an essential matrix of the separate target frame based on the feature points and the registration points;
- [170] a camera parameter parsing module, configured to solve intrinsic and extrinsic parameters of a camera based on the essential matrix;
- [171] a coordinate converting module, configured to convert two-dimensional coordinates of the feature points and the registration points into three-dimensional coordinates based on the intrinsic and extrinsic parameters of the camera;
- [172] a separate target frame remaining judgment module, configured to judge whether there are other separate target frames of the object unanalyzed, if yes, the feature point module is returned, otherwise, a 3D model constructing module is performed;
- [173] the 3D model constructing module, configured to generate a three-dimensional point cloud based on three-dimensional coordinates of the feature points and the registration points, and complete geometric modeling of the 3D model of the object by using a Poisson curved surface construction method based on the 3D point cloud; and
- [174] a 3D model refining module, configured to refine the texture information, the color feature information and the shape feature information of the object based on the separate target frame, and mapping the refined information to the 3D model of the object.
- [175] In some examples, after the 3D object refining module or the model constructing module, the apparatus further includes:
- [176] a posture synchronizing module, configured to identify a posture of any object in any picture in the 2D picture set, and adjust the posture of the 3D model of the object to be consistent with the posture of the object;
- [177] an AR projecting module, configured to render the adjusted 3D model into an AR scene.
- [178] Furthermore, the AR projecting module includes:
- [179] an information acquiring module, configured to acquire real scene information;
- [180] an affine transformation matrix resolving module, configured to analyze the real scene information and camera position information to obtain an affine transformation matrix of the 3D model of the object projected on a camera view plane; and
- [181] a projecting module, configured to merge the 3D model with a real scene video, wherein the 3D model of the object is imported with the affine transformation matrix, and display the 3D model and the real scene video together on an AR presenting

device or other device.

[182] It should be noted that the examples of the apparatus for converting a 2D picture set to a 3D model of the present disclosure are the same as the principles of examples of the method for converting a 2D picture set to a 3D model, and relevant points may be cross-referenced.

[183] The examples described above are merely preferred examples of the present disclosure and they do not limit the present disclosure. Any modification, equivalent replacement, and improvement made within the spirit and principle of technical solutions of the present disclosure shall fall within the scope of the present disclosure.

Claims

- [Claim 1] A method for converting a 2D picture set to a 3D model, characterized in that, the 2D picture set comprises at least one picture, and the method comprises:
identifying category of each object contained in each picture in the 2D picture set and outline information of each object by using a deep learning algorithm;
extracting detail characteristic information of each object in each picture in the 2D picture set by using a computer vision algorithm, wherein the detail feature information at least includes texture information, color feature information and shape feature information of each object;
matching a 3D model based on the category, the shape feature information and the outline information of each object, wherein the 3D model matched successfully is the 3D model of the object; and
mapping the texture information and the color feature information of each object onto the 3D model of the object.
- [Claim 2] The method according to claim 1, characterized in that, before the identifying category of each object contained in each picture in the 2D picture set and outline information of each object, the method further comprises:
extracting key frames in a 2D video as pictures in the 2D picture set.
- [Claim 3] The method according to claim 2, characterized in that, the extracting key frames in a 2D video as pictures in the 2D picture set comprises:
decoding the 2D video to acquire all static frames of the 2D video; and
performing clustering analysis on all the static frames to extract a static frame with the largest entropy in each cluster as a key frame of the cluster.
- [Claim 4] The method according to claim 1, characterized in that, the identifying category of each object contained in each picture in the 2D picture set and outline information of each object comprises:
inputting any picture in the 2D picture set into a convolutional neural network, wherein the convolutional neural network outputs an n-level feature map $P_1 \dots P_n$ of the picture, and $n \geq 2$;
locating target proposal region in the $P_1 \dots P_n$ by using a Region Proposal Network, wherein each proposal region comprises at least one anchor box;

when any of the proposal regions comprises at least two anchor boxes, screening the anchor boxes of each proposal region by adopting a non-maximum suppression algorithm, to reserve the anchor box with a highest foreground score, and discard other anchor boxes;
 for each anchor box in the $P_1 \dots P_n$, dividing the anchor box into a pooling unit with a first preset size, extracting the characteristic value of each subunit by using max-pooling, and then outputting the pooled $P_1 \dots P_n$;
 mapping the pooled $P_1 \dots P_n$ to a fully-connected feature, identifying the object category of each anchor box on the fully-connected feature and reducing the size of the anchor box; and
 identifying outline information of objects in region of each reduced anchor box by using a full convolutional network.

[Claim 5] The method according to claim 4, characterized in that, the convolutional neural network comprises a residual network and a feature pyramid network, wherein 3x3 convolution of the residual network is replaced with a depth wise separable convolution.

[Claim 6] The method according to claim 5, characterized in that, the feature pyramid network comprises a first pyramid and a second pyramid, and the output of the first pyramid from bottom layer to top layer is sequentially input to top layer to bottom layer of the second pyramid.

[Claim 7] The method according to claim 1, characterized in that, the detail feature information further comprises: whether the 2D picture is a separate target frame of an object.

[Claim 8] The method according to claim 1, characterized in that, the extracting detail characteristic information of each object in each picture in the 2D picture set comprises:

locating objects in any picture in the 2D picture set by using a super-pixel and/or threshold segmentation method; if the picture comprises only one object, then the picture is a separate target frame of the object; and

based on the locating information of each object in the picture, extracting the texture of each object by using Tamura texture feature algorithm and wavelet transform, extracting color feature information of each object by using a color histogram matching, and extracting shape feature information of each object by using a geometric parameter method.

[Claim 9] The method according to claim 7, characterized in that, the matching a

3D model based on the category, the shape feature information and the outline information of each object further comprises: if the matching fails, constructing a 3D model of the object based on the separate target frame of the object matched unsuccessfully.

[Claim 10]

The method according to claim 9, characterized in that, the constructing a 3D model of the object based on the separate target frame of the object matched unsuccessfully comprises:

extracting feature points of the object in the separate target frame;

matching the feature points of the object in the separate target frame to obtain registration points of the feature points;

extracting an essential matrix of the separate target frame based on the feature points and the registration points;

solving intrinsic and extrinsic parameters of a camera based on the essential matrix;

converting two-dimensional coordinates of the feature points and the registration points into three-dimensional coordinates based on the intrinsic and extrinsic parameters of the camera;

judging whether there are other separate target frames of the object un-analyzed, if yes, the extracting feature points of the object in the separate target frame is returned, otherwise, generating a 3-dimensional point cloud based on three-dimensional coordinates of the feature points and the registration points, and completing geometric modeling of the 3D model of the object by using a Poisson curved surface reconstruction method based on the 3D point cloud; and
refining the texture information, the color feature information and the shape feature information of the object based on the separate target frame, and mapping the refined information to the 3D model of the object.

[Claim 11]

The method according to claim 9, characterized in that, after the mapping the texture information and the color feature information of each object onto the 3D model of the object or the constructing a 3D model of the object based on the separate target frame of the object matched unsuccessfully, the method further comprises:

identifying a posture of any object in any picture in the 2D picture set, and adjusting the posture of the 3D model of the object to be consistent with the posture of the object;

rendering the adjusted 3D model into an AR scene.

[Claim 12]

The method according to claim 11, characterized in that, the rendering

the adjusted 3D model into an AR scene comprises:

acquiring real scene information;

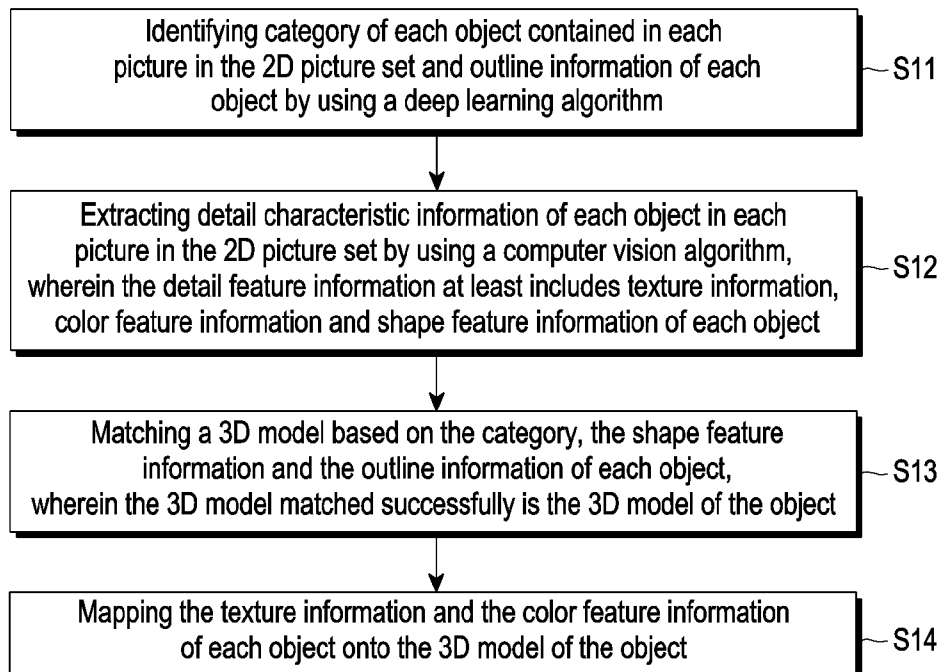
analyzing the real scene information and camera position information to obtain an affine transformation matrix of the 3D model of the object projected on a camera view plane; and

merging the 3D model with a real scene video, wherein the 3D model of the object is imported with the affine transformation matrix, and displaying the 3D model and the real scene video together on an AR presenting device or other device.

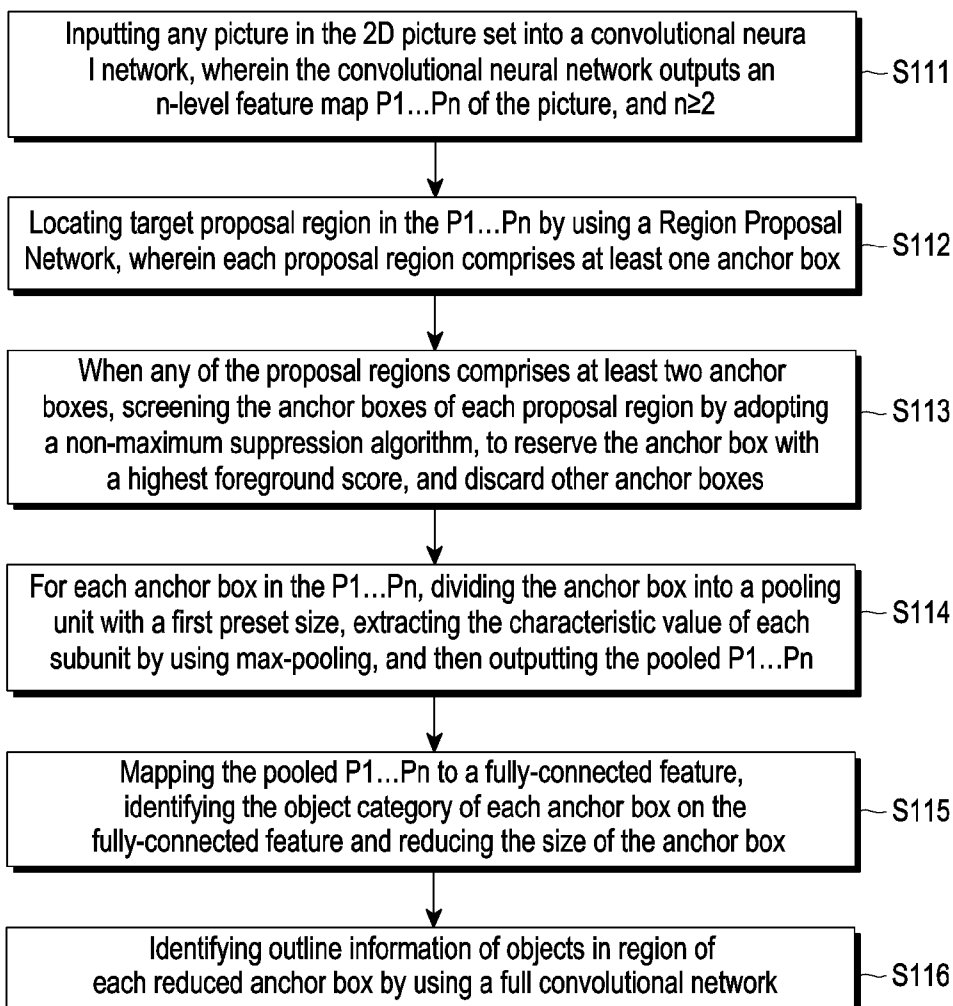
[Claim 13] A non-transitory computer-readable storage medium storing instructions that, when executed by a processor, cause the processor to perform steps in the method for converting the 2D picture set to the 3D model according to any one of claims 1 to 12.

[Claim 14] An apparatus for converting a 2D picture set to a 3D model, characterized by comprising a processor and a non-transitory computer-readable storage medium according to claim 13.

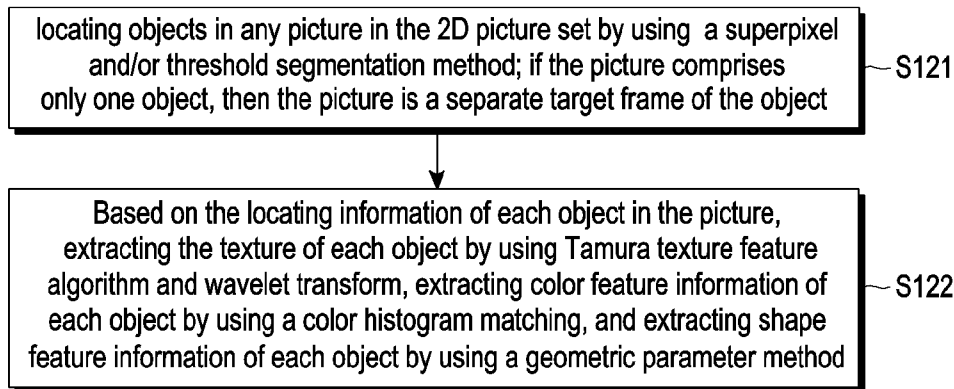
[Fig. 1]



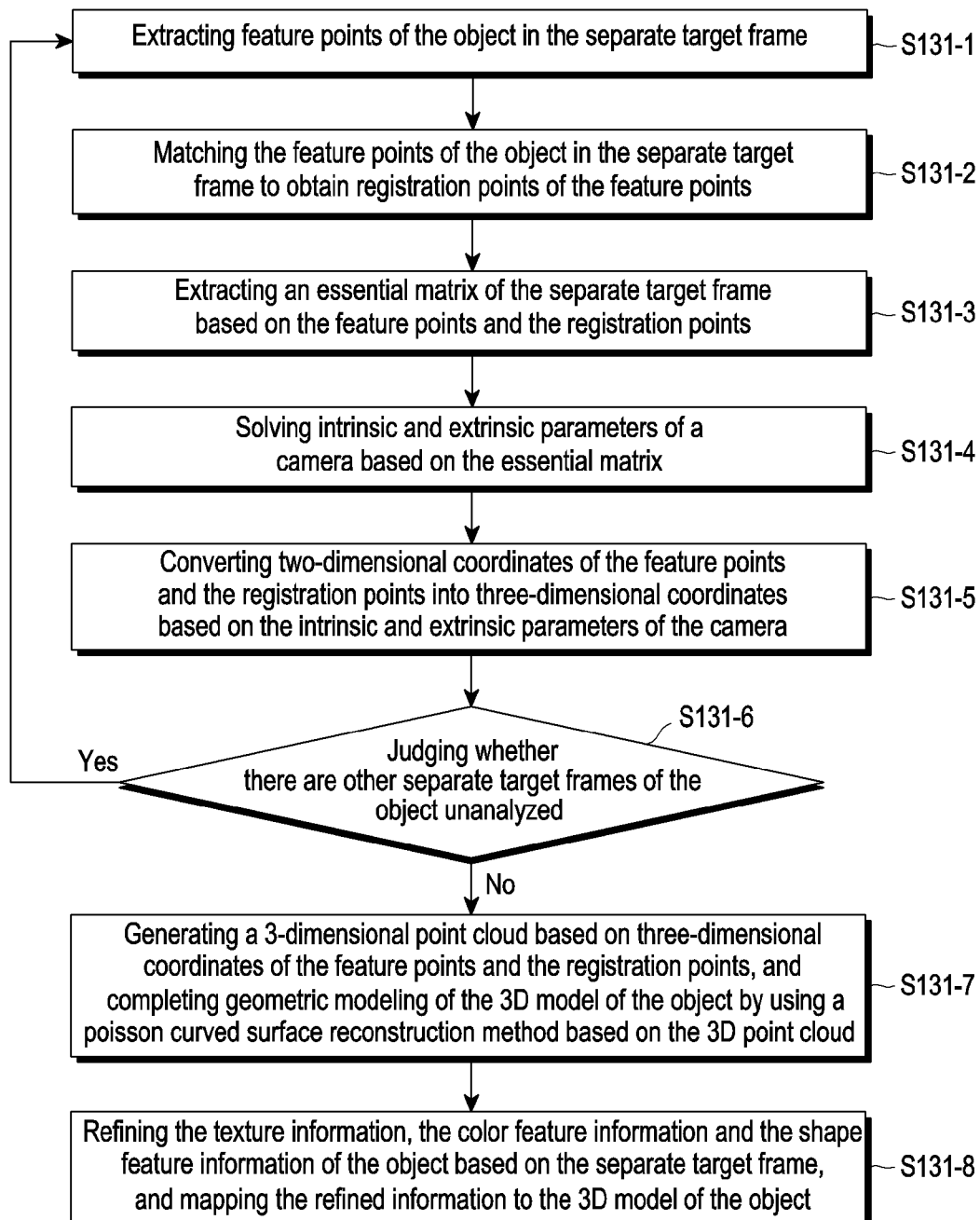
[Fig. 2]



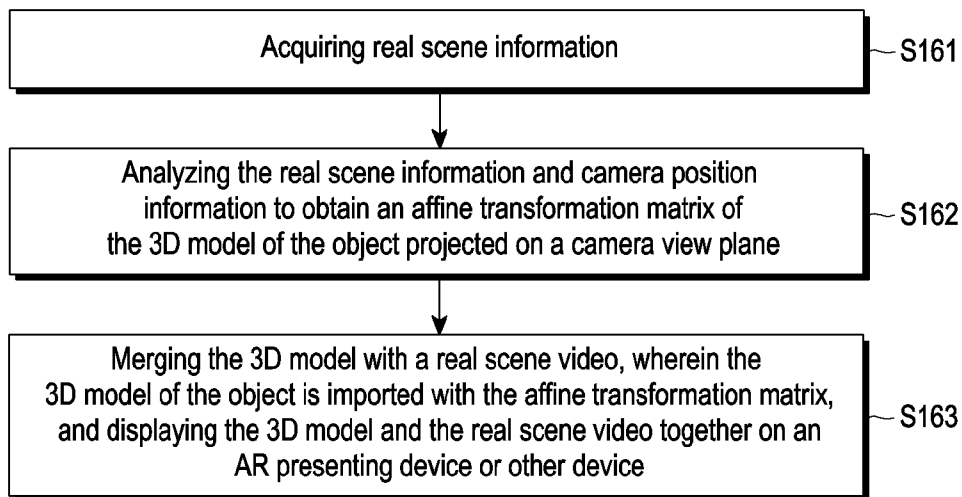
[Fig. 3]



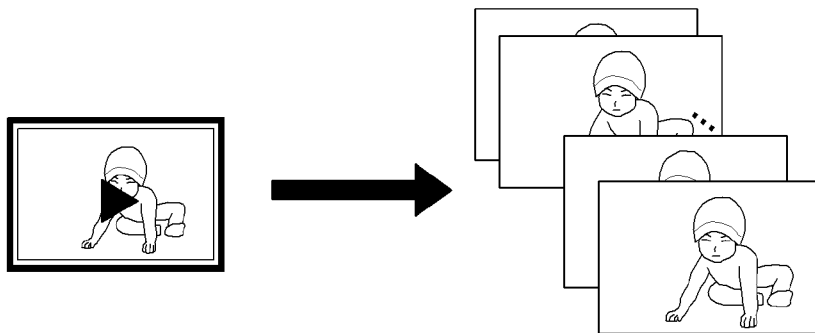
[Fig. 4]



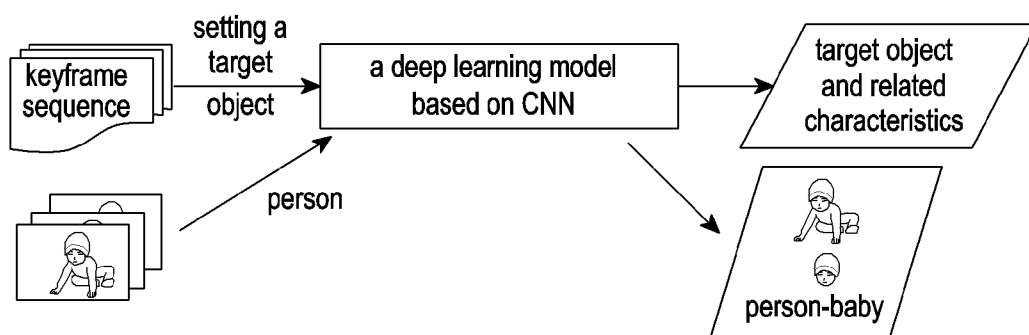
[Fig. 5]



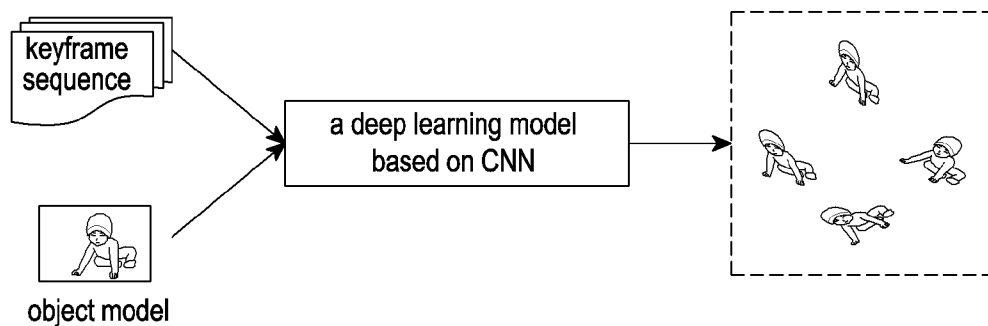
[Fig. 6]



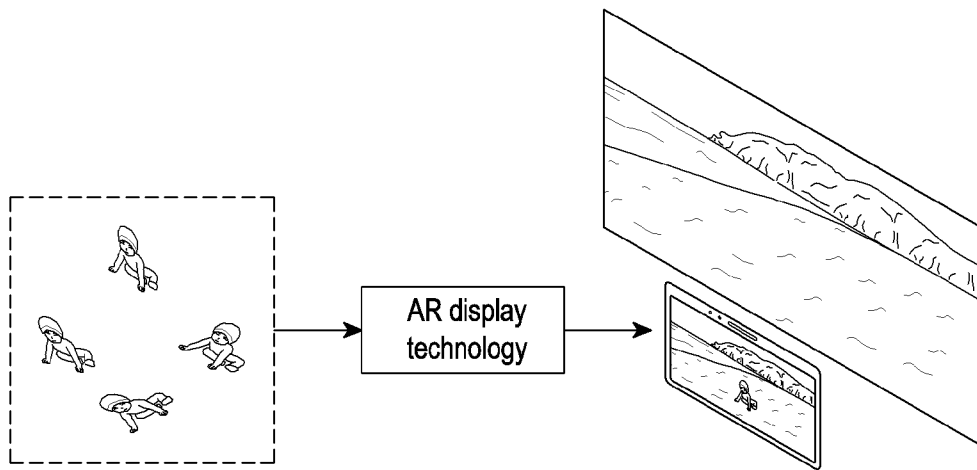
[Fig. 7]



[Fig. 8]



[Fig. 9]



A. CLASSIFICATION OF SUBJECT MATTER**H04N 13/275(2018.01)i, H04N 13/359(2018.01)i, H04N 13/31(2018.01)i, G06N 3/08(2006.01)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)

H04N 13/275; G06E 1/00; G06K 9/00; G06K 9/62; G06N 99/00; G06T 15/00; G06T 19/00; H04N 13/00; H04N 13/02; H04N 13/359; H04N 13/31; G06N 3/08

Documentation searched other than minimum documentation to the extent that such documents are included in the fields searched

Korean utility models and applications for utility models

Japanese utility models and applications for utility models

Electronic data base consulted during the international search (name of data base and, where practicable, search terms used)

eKOMPASS(KIPO internal) & Keywords: 2D picture set, 3D model, deep learning, category, outline information, texture, color, shape, key frame, extract, clustering analysis**C. DOCUMENTS CONSIDERED TO BE RELEVANT**

Category*	Citation of document, with indication, where appropriate, of the relevant passages	Relevant to claim No.
Y	US 2017-0085863 A1 (LEGEND3D, INC.) 23 March 2017 See paragraphs [0100], [0103]; claim 1; and figure 2.	1-3, 7-8, 13-14
A		4-6, 9-12
Y	US 2018-0189611 A1 (AQUIFI, INC.) 05 July 2018 See paragraph [0097]; claims 1, 4, 6; and figure 2A.	1-3, 7-8, 13-14
A	WO 2018-126270 A1 (NOVUMIND, INC.) 05 July 2018 See paragraphs [0027], [0039]; and figure 15.	1-14
A	US 2018-0190033 A1 (FACEBOOK, INC.) 05 July 2018 See paragraphs [0043]-[0044]; and figure 2.	1-14
A	US 6466205 B2 (TODD SIMPSON et al.) 15 October 2002 See column 7, lines 38-52.	1-14



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

27 August 2019 (27.08.2019)

Date of mailing of the international search report

27 August 2019 (27.08.2019)

Name and mailing address of the ISA/KR

International Application Division

Korean Intellectual Property Office

189 Cheongsu-ro, Seo-gu, Daejeon, 35208, Republic of Korea



Facsimile No. +82-42-481-8578

Authorized officer

AHN, Jeong Hwan

Telephone No. +82-42-481-8633



INTERNATIONAL SEARCH REPORT

Information on patent family members

International application No.

PCT/KR2019/006025

Patent document cited in search report	Publication date	Patent family member(s)	Publication date
US 2017-0085863 A1	23/03/2017	AU 2002-305387 B2 AU 2011-348121 A1 AU 2013-216732 A1 AU 2015-213286 A1 CA 2446150 A1 CA 2866672 A1 CN 104272377 A EP 1405269 A2 EP 2656315 A2 EP 2656315 A4 EP 2812894 A4 MX PA03010039 A US 03630698 A US 2004-0131249 A1 US 2011-0164109 A1 US 2012-0063681 A1 US 2013-0183023 A1 US 2013-0266292 A1 US 2015-0358612 A1 US 7181081 B2 US 8385684 B2 US 8396328 B2 US 9241147 B2 US 9270965 B2 US 9407904 B2 US 9438878 B2 US 9443555 B2 US 9595296 B2 US 9609307 B1 US 9615082 B2 WO 02-091302 A2 WO 2002-091302 A3 WO 2012-058490 A2 WO 2012-058490 A3 WO 2012-088477 A2 WO 2012-088477 A3 WO 2013-120115 A2 WO 2013-120115 A3 WO 2017-031113 A1 WO 2017-031117 A1	03/04/2008 11/07/2013 25/09/2014 03/09/2015 14/11/2002 15/08/2013 07/01/2015 07/04/2004 30/10/2013 31/12/2014 06/04/2016 06/12/2004 28/12/1971 08/07/2004 07/07/2011 15/03/2012 18/07/2013 10/10/2013 10/12/2015 20/02/2007 26/02/2013 12/03/2013 19/01/2016 23/02/2016 02/08/2016 06/09/2016 13/09/2016 14/03/2017 28/03/2017 04/04/2017 14/11/2002 22/01/2004 03/05/2012 28/06/2012 28/06/2012 26/10/2012 15/08/2013 24/10/2013 23/02/2017 23/02/2017
US 2018-0189611 A1	05/07/2018	WO 2018-129201 A1	12/07/2018
WO 2018-126270 A1	05/07/2018	None	
US 2018-0190033 A1	05/07/2018	EP 3343490 A1 EP 3343491 A1 US 2018-0189840 A1 WO 2018-125764 A1	04/07/2018 04/07/2018 05/07/2018 05/07/2018

INTERNATIONAL SEARCH REPORT

Information on patent family members

International application No.

PCT/KR2019/006025Patent document
cited in search reportPublication
datePatent family
member(s)Publication
date

US 6466205 B2

15/10/2002

WO 2018-125766 A1

05/07/2018

AU 2000-19180 A1

05/06/2000

AU 2000-19180 B2

15/05/2003

AU 760594 B2

15/05/2003

CA 2350657 A1

25/05/2000

CN 1331822 A

16/01/2002

US 2001-0052899 A1

20/12/2001

WO 00-30039 A1

25/05/2000