

(51) International Patent Classification:
G06K 9/00 (2006.01)(21) International Application Number:
PCT/CN2016/073184(22) International Filing Date:
2 February 2016 (02.02.2016)

(25) Filing Language: English

(26) Publication Language: English

(72) Inventor; and

(71) Applicant : WANG, Xiaogang [CN/CN]; Department of Electrical Engineering, The Chinese University of Hong Kong, Shatin N.T., Hong Kong (CN).

(72) Inventors: WANG, Lijun; Room A442, Chuang Xin Yuan Building, Dalian University of Technology, Dalian, Liaoning (CN). OUYANG, Wanli; Department of Electrical Engineering, The Chinese University of Hong Kong, Shatin, N.T., Hong Kong (CN). LU, Huchuan; Room A442, Chuang Xin Yuan Building, Dalian University of Technology, Dalian, Liaoning (CN).

(74) Agent: INSIGHT INTELLECTUAL PROPERTY LIMITED; 19 A, Tower A, Indo Building, No. 48A Zhichun Road, Haidian District, Beijing 100098 (CN).

(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, DK, DM, DO, DZ, EC, EE, EG, ES, FI, GB, GD, GE, GH, GM, GT, HN, HR, HU, ID, IL, IN, IR, IS, JP, KE, KG, KN, KP, KR, 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: METHODS AND SYSTEMS FOR CNN NETWORK ADAPTION AND OBJECT ONLINE TRACKING

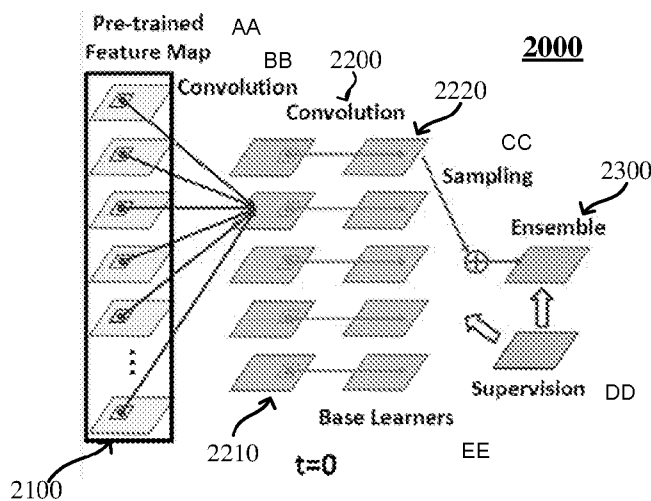


Fig. 2a

(57) Abstract: Disclosed are methods, apparatuses and systems for CNN network adaption and object online tracking. The CNN network adaption method comprises: transforming a first feature map into a plurality of sub-feature maps, wherein the first feature map is generated by the pre-trained CNN according to a frame of the target video; convolving each of the sub-feature maps with one of a plurality of adaptive convolution kernels, respectively, to output a plurality of second feature maps with improved adaptability; training, frame by frame, the adaptive convolution kernels.

METHODS AND SYSTEMS FOR CNN NETWORK ADAPTION AND OBJECT ONLINE TRACKING

Technical Field

[0001] The disclosures relate to methods, apparatuses and systems for Convolutional Neural Network (CNN) network adaption and object online tracking.

Background

[0002] Visual tracking is a fundamental problem in computer vision that has been received a rapidly growing attention. For a model-free object tracking problem, a category agnostic target is indicated by a bounding box in the first frame, and the tracker aims at locating the target in each of the following frames. Due to significant target appearance changes caused by abrupt motion, deformation, occlusion and illumination variation, visual tracking is still a challenging problem. Prior approaches rely on hand-crafted features to describe the target and have addressed the above challenging factors to a certain extend.

[0003] Recently, deep CNNs trained on large scale image classification data sets have demonstrated great success. These semantic representations discovered by the learning process are shown to be very effective at distinguishing objects of various categories. However, supervised training of deep CNNs with millions of parameters entails a large number of annotated training samples. To apply deep CNNs for tasks with a limited amount of training samples, previous approaches adopt a transfer learning method by first pre-training a deep CNN on a source task with a large scale training data set and then fine-tuning the learned feature on the target task. Due to the good generalization capability of CNN features across different data sets, this transfer learning approach is effective and has shown state-of-the-art performance in many applications.

[0004] However, for online visual tracking, the lack of training samples becomes even more severe, since the only training sample with ground truth label is provided in the first frame,

and the tracking results used for updating the tracker are also obtained in a sequential manner. Thus, directly fine-tuning a pre-trained deep CNN online is prone to over-fitting, which will degrade the tracker and gradually leads to tracking drift

Summary

[0005] The following presents a simplified summary of the disclosure in order to provide a basic understanding of some aspects of the disclosure. This summary is not an extensive overview of the disclosure. It is intended to neither identify key or critical elements of the disclosure nor delineate any scope of particular embodiments of the disclosure, or any scope of the claims. Its sole purpose is to present some concepts of the disclosure in a simplified form as a prelude to the more detailed description that is presented later.

[0006] In order to address, at least partially, one of the above issues, a CNN network adaption method is proposed for adapting a pre-trained CNN to a target video, in one aspect of the present application. The method comprises: transforming a first feature map into a plurality of sub-feature maps, wherein the first feature map is generated by the pre-trained CNN according to a frame of the target video; convolving each of the sub-feature maps with one of a plurality of adaptive convolution kernels, respectively, to output a plurality of second feature maps with improved adaptability; and training, frame by frame, the adaptive convolution kernels.

[0007] In one embodiment of the present application, the transforming and the convolving are implemented in an adaptive CNN comprising: a first convolution layer, linked to the pre-trained CNN and configured to transform the first feature map into the plurality of sub-feature maps; and a second convolution layer, linked to the first convolution layer and configured to convolve each of the sub-feature maps with one of the adaptive convolution kernels, respectively.

[0008] In one embodiment of the present application, the training comprises: feeding a first

training sample forward through the pre-trained CNN and the adaptive CNN to generate a first output image, wherein the first training sample is obtained according to a first frame of the target video; comparing the generated first output image with a first ground truth derived from the first frame to obtain a plurality of first training errors for the adaptive convolution kernels, respectively; back-propagating repeatedly the first training errors through the pre-trained CNN and the adaptive CNN to train the adaptive convolution kernels until the first training errors converge, wherein a plurality of parameters are obtained for the trained adaptive convolution kernels, respectively; grouping a parameter of the obtained parameters, which has a smallest first training error, and a rest of the obtained parameters into an ensemble set and a candidate set, respectively; and optimizing, according to a subsequent frame of the target video, the parameters grouped in the candidate set.

[0009] In one embodiment of the present application, the optimizing comprises: feeding the second training sample forward through the pre-trained CNN and the adaptive CNN to generate a second output image, wherein the second training sample is obtained according to a second frame of the target video and the second frame is subsequent to the first frame; comparing the second output image with a second ground truth derived from the second frame to obtain a plurality of second training errors for the plurality of adaptive convolution kernels; and if any of the second training errors is higher than a threshold, back-propagating the second training errors through the pre-trained CNN and the adaptive CNN to further refine the parameters in the candidate set and transferring at least one of the further refined parameters to the ensemble set.

[0010] In one embodiment of the present application, each of the adaptive convolution kernels is trained under a different loss criterion.

[0011] In one embodiment of the present application, the method further comprises further reducing, by a mask layer, a correlation among the sub-feature maps, wherein the mask layer is linked to the second convolution layer of the adaptive CNN.

[0012] In one embodiment of the present application, the mask layer comprises a plurality of binary masks, each of which is convolved with one of the sub-feature maps and has a same spatial size with the convolved sub-feature map.

[0013] In another aspect, a method is proposed for an object online tracking, comprising: determining a region of interest (ROI) in a first frame of a target video; feeding the determined ROI forward through a pre-trained CNN to extract an initial feature map thereof; initializing, with the initial feature map, an adaptive CNN used for detecting a location of the object and a scale estimation network used for defining a scale of the object; predicting, with the initialized adaptive CNN, a second location of the object in a second frame of the target video, wherein the second frame is subsequent to the first frame; estimating, with the initialized scale estimation network, a second scale of the object in the second frame of the target video; updating, with optimized network parameters acquired in the predicting and the estimating, the adaptive CNN and the scale estimation network, respectively; predicting, with the updated adaptive CNN, a third location of the object in a third frame of the target video, wherein the third frame is subsequent to the second frame; and estimating, with the updated scale estimation network, a third scale of the object in the third frame of the target video.

[0014] In one embodiment of the present application, the adaptive CNN comprises: a first convolution layer, linked to the pre-trained CNN and configured to transform a first feature map into a plurality of sub-feature maps, wherein the first feature map is generated by the pre-trained CNN according to any frame of the target video; and a second convolution layer, linked to the first convolution layer and configured to convolve each of the sub-feature maps with one of a plurality of adaptive convolution kernels, respectively, to output a plurality of second feature maps with improved adaptability.

[0015] In one embodiment of the present application, the adaptive CNN is initialized by: feeding a first training sample forward through the pre-trained CNN and the adaptive CNN to generate a first output image, wherein the first training sample is obtained according to a first frame of the target video; comparing the generated first output image with a first ground truth

derived from the first frame to obtain a plurality of first training errors for the adaptive convolution kernels, respectively; back-propagating repeatedly the first training errors through the pre-trained CNN and the adaptive CNN to train the adaptive convolution kernels until the first training errors converges, wherein a plurality of parameters are obtained for the trained adaptive convolution kernels, respectively; grouping a parameter of the obtained parameters, which has a smallest first training error, and a rest of the obtained parameters into an ensemble set and a candidate set, respectively.

[0016] In one embodiment of the present application, the adaptive CNN is updated by: feeding a second training sample forward through the pre-trained CNN and the adaptive CNN to generate a second output image, wherein the second training sample is obtained according to a second frame of the target video and the second frame is subsequent to the first frame; comparing the second output image with a second ground truth derived from the second frame to obtain a plurality of second training errors for the plurality of adaptive convolution kernels, respectively; and if any of the second training errors is higher than a threshold, back-propagating the second training errors through the pre-trained CNN and the adaptive CNN to further refine the parameters in the candidate set and transferring at least one of the further refined parameters to the ensemble set.

[0017] In one embodiment of the present application, each of the adaptive convolution kernels is trained under a different loss criterion.

[0018] In one embodiment of the present application, the adaptive CNN further comprises a mask layer linked to the second convolution layer to further reduce a correlation among the sub-feature maps.

[0019] In one embodiment of the present application, the mask layer comprises a plurality of binary masks, each of which is convolved with one of the sub-feature maps and has a same spatial size with the convolved sub-feature map.

[0020] In one embodiment of the present application, the location of the object is predicted by a heat map generated by the adaptive CNN, wherein a location with a maximum value is predicted to be the location of the object and the maximum value is sampled as a confidence.

[0021] In one embodiment of the present application, the updating is performed only if the confidence is higher than a pre-defined threshold.

[0022] In one embodiment of the present application, the ROI is centered at an object to be tracked.

[0023] In another aspect, a system is proposed for adapting a pre-trained CNN to a target video, comprising: a memory that stores executable components; and a processor electrically coupled to the memory to execute the executable components. The executable components are executed for: transforming a first feature map into a plurality of sub-feature maps, wherein the first feature map is generated by the pre-trained CNN according to a frame of the target video; convolving each of the sub-feature maps with one of a plurality of adaptive convolution kernels, respectively, to output a plurality of second feature maps with improved adaptability; training, frame by frame, the adaptive convolution kernels.

[0024] In another aspect, a system is proposed for an object online tracking, comprising: a memory that stores executable components; and a processor electrically coupled to the memory to execute the executable components. The executable components are executed for: determining a region of interest (ROI) in a first frame of a target video; feeding the determined ROI forward through a pre-trained CNN to extract an initial feature map thereof; initializing, with the initial feature map, an adaptive CNN used for detecting a location of the object and a scale estimation network used for defining a scale of the object; predicting, with the initialized adaptive CNN, a second location of the object in a second frame of the target video, wherein the second frame is subsequent to the first frame; estimating, with the initialized scale estimation network, a second scale of the object in the second frame of the target video; updating, with optimized network parameters acquired in the predicting and the

estimating, the adaptive CNN and the scale estimation network, respectively; predicting, with the updated adaptive CNN, a third location of the object in a third frame of the target video, wherein the third frame is subsequent to the second frame; and estimating, with the updated scale estimation network, a third scale of the object in the third frame of the target video.

[0025] In another aspect, an apparatus is proposed for adapting a pre-trained CNN to a target video, comprising: means for transforming a first feature map into a plurality of sub-feature maps, wherein the first feature map is generated by the pre-trained CNN according to a frame of the target video; means for convolving each of the sub-feature maps with one of a plurality of adaptive convolution kernels, respectively, to output a plurality of second feature maps with improved adaptability; and means for training, frame by frame, the adaptive convolution kernels.

[0026] In another aspect, an apparatus is proposed for an object online tracking. The apparatus comprises a feature extraction unit, configured for: determining a region of interest (ROI) in a first frame of a target video; and feeding the determined ROI forward through a pre-trained CNN to extract an initial feature map thereof. The apparatus further comprises: an initialization and update unit, configured for initializing, with the initial feature map, an adaptive CNN used for detecting a location of the object and a scale estimation network used for defining a scale of the object; a location prediction unit, configured for predicting, with the initialized adaptive CNN, a second location of the object in a second frame of the target video, wherein the second frame is subsequent to the first frame; and a scale estimation unit, configured for estimating, with the initialized scale estimation network, a second scale of the object in the second frame of the target video. In addition, the initialization and update unit is further configured for updating, with optimized network parameters acquired in the predicting and the estimating, the adaptive CNN and the scale estimation network, respectively; the location prediction unit is further configured for predicting, with the updated adaptive CNN, a third location of the object in a third frame of the target video, wherein the third frame is subsequent to the second frame; and the scale estimation unit is further configured for

estimating, with the updated scale estimation network, a third scale of the object in the third frame of the target video.

[0027] Based on the proposed CNN adaption method and system, pre-trained deep features can be effectively transferred for online application with a reduced over-fitting. The proposed object online tracking method, apparatus and system are constructed based on the proposed CNN adaption method, apparatus and system. Due to the reduced over-fitting, the proposed object online tracking method and system can perform an improved object online tracking.

Brief Description of the Drawing

[0028] Exemplary non-limiting embodiments of the present application are described below with reference to the attached drawings. The drawings are illustrative and generally not to an exact scale. The same or similar elements on different figures are referenced with the same reference numbers.

[0029] Fig. 1 illustrates a conventional method for adapting a pre-trained CNN to a target image.

[0030] Figs. 2a-2c illustrate an adaptive CNN for adapting a pre-trained CNN to a target video, according to an embodiment of the present application.

[0031] Fig. 3 is a flowchart illustrating the initialization of the adaptive CNN, according to an embodiment of the present application.

[0032] Fig. 4 is a flowchart illustrating the process of object online tracking, according to an embodiment of the present application.

[0033] Fig. 5 is a schematic diagram illustrating a system architecture for the object online tracking, according to an embodiment of the present application.

Detailed Description

[0034] Reference will now be made in detail to some specific embodiments of the invention including the best modes contemplated by the inventors for carrying out the invention. Examples of these specific embodiments are illustrated in the accompanying drawings. While the invention is described in conjunction with these specific embodiments, it will be appreciated by one skilled in the art that it is not intended to limit the invention to the described embodiments. On the contrary, it is intended to cover alternatives, modifications, and equivalents as may be included within the spirit and scope of the invention as defined by the appended claims. In the following description, numerous specific details are set forth in order to provide a thorough understanding of the present application. The present application may be practiced without some or all of these specific details. In other instances, well-known process operations have not been described in detail in order not to unnecessarily obscure the present application.

[0035] The terminology used herein is for the purpose of describing particular embodiments only and is not intended to be limiting of the invention. As used herein, the singular forms "a", "an" and "the" are intended to include the plural forms as well, unless the context clearly indicates otherwise. It will be further understood that the terms "comprises" and/or "comprising" when used in this specification, specify the presence of stated features, integers, steps, operations, elements, and/or components, but do not preclude the presence or addition of one or more other features, integers, steps, operations, elements, components, and/or groups thereof.

[0036] For online applications, one simple approach to transfer offline pre-trained CNN features is to add one or more randomly initialized CNN layers, i.e., adaptive CNN, sequence to the pre-trained CNN model. Then the parameters, i.e. convolutional kernels and bias, of the pre-trained CNN is fixed while only the parameters of the adaptative CNN is trained online to fit the current task, for example, a target video or a target image. Fig. 1 illustrates such a conventional method 1000 for adapting a pre-trained CNN 1100 to a target image. As can be

seen from Fig. 1, an adaptive CNN 1200 is disposed sequence to the pre-trained CNN 1100 and configured for refining the pre-trained features to a final adapted feature 1300 for supervision. However, as the parameters of the adaptive CNN 1200 are jointly learned during a training process, this transfer learning method suffers from severe over-fitting for an online application, where training sample with ground truth label is only provided in the first frame. The online learned parameters mainly focus on recent training samples and are less likely to well generalize to both historical and future samples. This phenomenon can be fatal to online visual tracking where the target often undergoes significant appearance changes or heavy occlusion.

[0037] To tackle the above issue, an adaptive CNN 2200 is proposed for better transferring pre-trained deep features, as shown in Fig. 2a to Fig. 2c. The pre-trained CNN 2100 is denoted as CNN-E, and an RGB image is taken as input and a convolution feature map X is output. The online adaptive CNN 2200, denoted as CNN-A, is randomly initialized and consists of two convolution layers interleaved with a ReLU layer as the nonlinear activation unit. The first convolution layer is linked to the pre-trained CNN 2100 and configured to transform the first feature map into the plurality of sub-feature maps. The second convolution layer is linked to the first convolution layer and configured to convolve each of the sub-feature maps with one of the adaptive convolution kernels, respectively. The online adaptive CNN 2200 takes the feature map X as input and generates the final feature map 2300, which is formulated by:

$$\{F_2^c(X) | c = 1, 2, \dots, C_2\} \quad (1)$$

where $F_2^c(X) \in \mathbb{R}^{m \times n}$ indicates the c -th channel of the sub-feature map generated by the second convolution layer with spatial size of $m \times n$. The sub-feature map in the second layer is obtained by convolving the kernel with the sub-feature map in the first layer as:

$$F_2^c(X) = \sum_{k=1}^{C_1} w_k^c * F_1^k(X) + b_c, \quad (2)$$

where C_1 denotes the number of channels of the sub-feature map output by the first convolution layer; w_k^c denotes the convolution kernel connecting the k -th channel of the first convolution layer sub-feature map with the c -th channel of the second convolution layer sub-feature map; b_c is the bias and the symbol "*" denotes a convolution operation. The summation is conducted over all the channels.

[0038] In order to introduce randomness into the parameter learning process, the output sub-feature map is regarded as an set of base learners, formulated by:

$$F_2^c(X) = \sum_{k=1}^{C_1} f(X; \gamma_k^c) \quad (3)$$

where each base learner is defined as:

$$f(X; \gamma_k^c) = w_k^c * F_1^k(X) + b_c^k \quad (4)$$

and the parameter γ_k^c indicates the corresponding kernel weights and bias in both the first and second convolution layers of CNN-A.

[0039] The online training of the CNN-A network is then equivalent to online training each base learner and to sequentially sample a well-trained parameter for each of base learners into an ensemble set. Since the proposed online training method is conducted independently in each channel of the output sub-feature map, in the following discussion, only one output channel will be discussed as an example to describe the training method. For notification simplicity, the superscript channel number is omitted and the notation $\{\gamma_k | k = 1, 2, \dots, C_1\}$ is used to denote the parameters of the base learners for any one output sub-feature map channel.

[0040] Fig. 3 shows how the online training process is initialized. At the beginning of the online training process, a first training sample is prepared according to a first frame of the target video at step S301. The strategy for preparing a training sample from a given image, such as a stochastic gradient descent (SGD), is well-known to one skilled in the art and thus will not be discussed in detail hereinafter. The first training sample is then fed forward, at step S302, through the pre-trained CNN and the adaptive CNN to generate a first output image. Then at step S303, the first output image is compared with the first ground truth, which is derived from the first frame, to obtain a plurality of first training errors for the plurality of adaptive convolution kernels, respectively. The first training errors are back-propagated, through the pre-trained CNN and the adaptive CNN to train the adaptive convolution kernels in an iterative manner until the first training errors converge, as shown at step S304.

[0041] In one implement, a plurality of parameters are trained for each of the adaptive convolution kernels, respectively, after the initialization. The parameter γ^* with the smallest training error is selected and grouped into an ensemble set $\mathcal{E}$ and the rest of trained parameters are grouped into a candidate set $\mathcal{C}$. The parameters in the candidate set will be optimized in following frames of the target video. However, in an alternative implement, two or more parameters with smallest training errors may be selected into the ensemble set.

[0042] In the following training process, i.e., an optimizing process, the parameters in the candidate set are sequentially added to the ensemble set in a similar manner. Since the optimizing process is similar to that of the initialization, such as the preparing, the feeding, the comparing and the back-propagating, only the difference will be discussed hereinafter.

[0043] In the optimizing process, all the parameters in the ensemble set are used to form an ensemble with output $F(X; \mathcal{E}) = \frac{1}{|\mathcal{E}|} \sum_{\gamma_i \in \mathcal{E}} f(X; \gamma_i)$ for online testing. At the t-th step,

a newly training sample X_t with target output Y_t is obtained. The parameters in the ensemble set $\mathcal{E}$ is jointly refined, for example by SGD with the loss function $L_{\mathcal{E}} = L(Y_t, F(X_t; \mathcal{E}))$. Meanwhile, each parameter $\gamma_j \in \mathcal{C}$ is refined independently, for example by SGD using with following loss function:

$$L_{\mathcal{C}}(Y_t, f(X_t; \gamma_j)) = L(Y_t, f(X_t; \gamma_j) + \eta F(X_t; \mathcal{E})) \quad (5)$$

where $F(X_t; \mathcal{E})$ is fixed and the parameter η is used to balance the impact of the ensemble on the candidate parameters, such that the refining of the parameter $\gamma_j \in \mathcal{C}$ considers both the target output Y_t and the output of the ensemble $F(X_t; \mathcal{E})$. If the training error $L_{\mathcal{E}}$ is higher than a predefined threshold and the candidate set $\mathcal{C}$ is not empty, a refined parameter is sampled from the candidate set $\mathcal{C}$, for example, according to the following sampling probability density

$$p(\gamma) = q(L_{\mathcal{C}}(Y_t, f(X_t; \gamma))), \gamma \in \mathcal{C} \quad (6)$$

where $q(\cdot)$ is a monotonically decreasing function. Thus the sampled parameter is removed from the candidate set $\mathcal{C}$ and added into the ensemble set $\mathcal{E}$. The above online training approach is conducted sequentially in each time step. When all the parameters are sampled from the candidate set to the ensemble set, the ensemble $F(X; \mathcal{E})$ evolves into a well-trained CNN model. In an alternate implement, the parameters incorporated in the well-trained CNN model would still be jointly updated with a further training process during the period of the subsequent frames. The proposed adaptive CNN demonstrate a moderate diversity since the parameters thereof are trained independently, especially by using different loss criterions.

[0044] In one embodiment, a mask layer may be contained in the adaptive CNN and linked

to the second convolution layer to further reduce a correlation among the sub-feature maps. Specifically, each channel of the output sub-feature map from the second convolution layer is associated with an individual binary mask which has the same spatial size with the sub-feature map. All the masks are initialized in a random manner and then fixed throughout the online training process. The forward propagation of the convolution layer at the training stage is then conducted as

$$F^c = \sum_{k=1}^K w_k^c * (M^c \odot X^k) + b_c \quad (7)$$

where X^k indicates the k-th channel of the sub-feature map; M^c denotes the binary mask associated with the c-th channel of the output feature map F^c ; and the symbol " $\odot$ " denotes a *Hadamard* product. Accordingly, the backward propagation is also conducted by considering the binary masks. Trained in this way, the learned convolution kernels are enforced to focus on different part of the input feature maps through the binary masks.

[0045] In another aspect, a method is proposed for an object online tracking. Fig. 4 schematically illustrates a general flowchart for an object online tracking. At the beginning of the tracking process, as shown at step S401, a target object is manually selected with a target boundary box in the first frame of the video, and accordingly, a region of interest (ROI) is determined. In one implement, the ROI is centered at the target object. In a further implement, the ROI may have twice the size of the target boundary box, as an example. At step S402, the ROI is fed forward a pre-trained CNN to extract an initial feature map for initialization. The initial feature map comprises the information of the location and the scale of the target object in the first frame. At step S403, an adaptive CNN and a scale estimation network is initialized with the initial feature map, wherein adaptive CNN is utilized to predict the location of the target object and the scale estimation network is utilized to estimate the scale of the target object. The scale estimation process will be further described in the following paragraphs. With the initialized adaptive CNN and the scale estimation network, the location and scale of the target object is predicted and estimated in a following

frame, for example, a second frame of the target video, as shown at step S404 and S405. Besides, the process of location prediction and scale estimation is also a training process, which gives a plurality of optimized parameters for the adaptive CNN and the scale estimation network, respectively. With the optimized parameters, the adaptive CNN and the scale estimation network are jointly updated at step S406. After updating process, both the adaptive CNN and the scale estimation network have better adaptabilities with respect to the relating frame. The location and scale of target object in the subsequent frames of the target video will be predicted and estimated based on the updated adaptive CNN (step S407 and S408). The result of the prediction and estimation further updates, in turn, the adaptive CNN and the scale estimation network.

[0046] The above-mentioned object online tracking process may be implemented in a system illustrated in Fig. 5. A feature extraction unit 501 comprising a pre-trained CNN, i.e., CNN-E, is configured to determine a ROI and to extract the feature of any frame. The extracted feature is transmitted to a location prediction unit 502 comprising the adaptive CNN and a scale estimation unit 503 comprising the scale estimation network. In addition, the extracted feature is also transmitted to an initialization and update unit 504 for initializing at the first frame of the target video. In subsequent frames, the extracted feature map of the current frame and the output from the location prediction unit 502 and scale estimation unit 503 are transmitted to the initialization and update unit 504 to update the adaptive CNN and the scale estimation network.

[0047] The structure, the initializing and the optimizing of the adaptive CNN is previously discussed and will not be further detailed hereinafter. In the case of object online tracking, the adaptive CNN is specifically used to transform, by performing a heat map regression, the feature map extracted from the pre-trained CNN into a target heat map. The location of the target object is then determined by the location on the heat map with the maximum value and the corresponding maximum heat map value serves as the confidence of this prediction. In one embodiment, the updating of the adaptive CNN and the scale estimation network are

conducted only if the confidence is higher than a pre-defined threshold, in order to avoid updating using contaminated training sample.

[0048] The scale estimation network has a conventional deep learning network structure, such as a CNN or a fully-connected network, and thus would not be described in detail herein. In the case of object online tracking, the scale prediction unit receives the feature map extracted from the pre-trained network and applies a set of pre-defined scale transformations to obtain the corresponding scale-transformed feature maps. The scale-transformed feature maps are fed forward through the scale estimation network, which assigns a score for each scale transformation. The scale with the highest score is then predicted as the current scale of the target. With the location and the scale resulted from the tracking system, the target object will be tracked at an improved precision.

[0049] As will be appreciated by one skilled in the art, the present application may be embodied as a system, a method or a computer program product. Accordingly, the present application may take the form of an entirely hardware embodiment and hardware aspects that may all generally be referred to herein as a "unit", "circuit", "module", or "system". Much of the inventive functionality and many of the inventive principles when implemented, are best supported with or integrated circuits (ICs), such as a digital signal processor and software therefore or application specific ICs. It is expected that one of ordinary skill, notwithstanding possibly significant effort and many design choices motivated by, for example, available time, current technology, and economic considerations, when guided by the concepts and principles disclosed herein will be readily capable of generating ICs with minimal experimentation. Therefore, in the interest of brevity and minimization of any risk of obscuring the principles and concepts according to the present application, further discussion of such software and ICs, if any, will be limited to the essentials with respect to the principles and concepts used by the preferred embodiments. In addition, the present invention may take the form of an entirely software embodiment (including firmware, resident software, micro-code, etc.) or an embodiment combining software. For example, the system may comprise a memory that

stores executable components and a processor, electrically coupled to the memory to execute the executable components to perform operations of the system, as discussed in reference to Figs. 1-6. Furthermore, the present invention may take the form of a computer program product embodied in any tangible medium of expression having computer-usable program code embodied in the medium.

[0050] Although the preferred examples of the present application have been described, those skilled in the art can make variations or modifications to these examples upon knowing the basic inventive concept. The appended claims are intended to be considered as comprising the preferred examples and all the variations or modifications fell into the scope of the present application.

[0051] Obviously, those skilled in the art can make variations or modifications to the present application without departing the spirit and scope of the present application. As such, if these variations or modifications belong to the scope of the claims and equivalent technique, they may also fall into the scope of the present application.

What is claimed is:

1. A method for adapting a pre-trained CNN to a target video, comprising:
transforming a first feature map into a plurality of sub-feature maps, wherein the first feature map is generated by the pre-trained CNN according to a frame of the target video;
convolving each of the sub-feature maps with one of a plurality of adaptive convolution kernels, respectively, to output a plurality of second feature maps with improved adaptability;
and
training, frame by frame, the adaptive convolution kernels.
2. The method of claim 1, wherein the transforming and the convolving are implemented in an adaptive CNN comprising:
a first convolution layer, linked to the pre-trained CNN and configured to transform the first feature map into the plurality of sub-feature maps; and
a second convolution layer, linked to the first convolution layer and configured to convolve each of the sub-feature maps with one of the adaptive convolution kernels, respectively.
3. The method of claim 2, wherein the training comprises:
feeding a first training sample forward through the pre-trained CNN and the adaptive CNN to generate a first output image, wherein the first training sample is obtained according to a first frame of the target video;
comparing the generated first output image with a first ground truth derived from the first frame to obtain a plurality of first training errors for the adaptive convolution kernels, respectively;
back-propagating repeatedly the first training errors through the pre-trained CNN and the adaptive CNN to train the adaptive convolution kernels until the first training errors converge,

wherein a plurality of parameters are obtained for the trained adaptive convolution kernels, respectively;

grouping a parameter of the obtained parameters, which has a smallest first training error, and a rest of the obtained parameters into an ensemble set and a candidate set, respectively; and

optimizing, according to a subsequent frame of the target video, the parameters grouped in the candidate set.

4. The method of claim 3, wherein the optimizing comprises:

feeding a second training sample forward through the pre-trained CNN and the adaptive CNN to generate a second output image, wherein the second training sample is obtained according to a second frame of the target video and the second frame is subsequent to the first frame;

comparing the second output image with a second ground truth derived from the second frame to obtain a plurality of second training errors for the plurality of adaptive convolution kernels; and

if any of the second training errors is higher than a threshold,

back-propagating the second training errors through the pre-trained CNN and the adaptive CNN to further refine the parameters in the candidate set; and
transferring at least one of the further refined parameters to the ensemble set.

5. The method of claim 1, wherein each of the adaptive convolution kernels is trained under a different loss criterion.

6. The method of claim 2, wherein the method further comprises:

reducing, by a mask layer, a correlation among the sub-feature maps, wherein the mask layer is linked to the second convolution layer of the adaptive CNN.

7. The method of claim 6, wherein the mask layer comprises a plurality of binary masks, each of which is convolved with one of the sub-feature maps and has a same spatial size with the convolved sub-feature map.

8. A method for an object online tracking, comprising:

determining a region of interest (ROI) in a first frame of a target video;

feeding the determined ROI forward through a pre-trained CNN to extract an initial feature map thereof;

initializing, with the initial feature map, an adaptive CNN used for detecting a location of the object and a scale estimation network used for defining a scale of the object;

predicting, with the initialized adaptive CNN, a second location of the object in a second frame of the target video, wherein the second frame is subsequent to the first frame;

estimating, with the initialized scale estimation network, a second scale of the object in the second frame of the target video;

updating, with optimized network parameters acquired in the predicting and the estimating, the adaptive CNN and the scale estimation network, respectively;

predicting, with the updated adaptive CNN, a third location of the object in a third frame of the target video, wherein the third frame is subsequent to the second frame; and

estimating, with the updated scale estimation network, a third scale of the object in the third frame of the target video.

9. The method of claim 8, wherein the adaptive CNN comprises:

a first convolution layer, linked to the pre-trained CNN and configured to transform a first feature map into a plurality of sub-feature maps, wherein the first feature map is generated by the pre-trained CNN according to any frame of the target video; and

a second convolution layer, linked to the first convolution layer and configured to convolve each of the sub-feature maps with one of a plurality of adaptive convolution kernels, respectively, to output a plurality of second feature maps with improved adaptability.

10. The method of claim 9, wherein the adaptive CNN is initialized by:

feeding a first training sample forward through the pre-trained CNN and the adaptive CNN to generate a first output image, wherein the first training sample is obtained according to a first frame of the target video;

comparing the generated first output image with a first ground truth derived from the first frame to obtain a plurality of first training errors for the adaptive convolution kernels, respectively;

back-propagating repeatedly the first training errors through the pre-trained CNN and the adaptive CNN to train the adaptive convolution kernels until the first training errors converges, wherein a plurality of parameters are obtained for the trained adaptive convolution kernels, respectively;

grouping a parameter of the obtained parameters, which has a smallest first training error, and a rest of the obtained parameters into an ensemble set and a candidate set, respectively.

11. The method of claim 9, wherein the adaptive CNN is updated by:

feeding the second training sample forward through the pre-trained CNN and the adaptive CNN to generate a second output image, wherein the second training sample is obtained according to a second frame of the target video and the second frame is subsequent to the first frame;

comparing the second output image with a second ground truth derived from the second frame to obtain a plurality of second training errors for the plurality of adaptive convolution kernels, respectively; and

if any of the second training errors is higher than a threshold,

back-propagating the second training errors through the pre-trained CNN and the adaptive CNN to further refine the parameters in the candidate set; and

transferring at least one of the further refined parameters to the ensemble set.

12. The method of claim 9, wherein each of the adaptive convolution kernels is trained

under a different loss criterion.

13. The method of claim 9, wherein the adaptive CNN further comprises a mask layer linked to the second convolution layer to further reduce a correlation among the sub-feature maps.

14. The method of claim 13, wherein the mask layer comprises a plurality of binary masks, each of which is convolved with one of the sub-feature maps and has a same spatial size with the convolved sub-feature map.

15. The method of claim 9, wherein the location of the object is predicted by a heat map generated by the adaptive CNN, wherein a location with a maximum value is predicted to be the location of the object and the maximum value is sampled as a confidence.

16. The method of claim 15, wherein the updating is performed only if the confidence is higher than a pre-defined threshold.

17. The method of claim 8, wherein the ROI is centered at an object to be tracked.

18. A system for adapting a pre-trained CNN to a target video, comprising:
a memory that stores executable components; and
a processor electrically coupled to the memory to execute the executable components
for:

transforming a first feature map into a plurality of sub-feature maps, wherein the first feature map is generated by the pre-trained CNN according to a frame of the target video;

convolving each of the sub-feature maps with one of a plurality of adaptive convolution kernels, respectively, to output a plurality of second feature maps with improved adaptability; and

training, frame by frame, the adaptive convolution kernels.

19. The system of claim 18, wherein the executable components comprises an adaptive CNN comprising:

a first convolution layer, linked to the pre-trained CNN and configured to transform the first feature map into the plurality of sub-feature maps; and

a second convolution layer, linked to the first convolution layer and configured to convolve each of the sub-feature maps with one of the adaptive convolution kernels, respectively.

20. The system of claim 19, wherein the training comprises:

feeding a first training sample forward through the pre-trained CNN and the adaptive CNN to generate a first output image, wherein the first training sample is obtained according to a first frame of the target video;

comparing the generated first output image with a first ground truth derived from the first frame to obtain a plurality of first training errors for the adaptive convolution kernels, respectively;

back-propagating repeatedly the first training errors through the pre-trained CNN and the adaptive CNN to train the adaptive convolution kernels until the first training errors converge, wherein a plurality of parameters are obtained for the trained adaptive convolution kernels, respectively;

grouping a parameter of the obtained parameters, which has a smallest first training error, and a rest of the obtained parameters into an ensemble set and a candidate set, respectively; and

optimizing, according to a subsequent frame of the target video, the parameters grouped in the candidate set.

21. The system of claim 20, wherein the optimizing comprises:

feeding a second training sample forward through the pre-trained CNN and the adaptive

CNN to generate a second output image, wherein the second training sample is obtained according to a second frame of the target video and the second frame is subsequent to the first frame;

comparing the second output image with a second ground truth derived from the second frame to obtain a plurality of second training errors for the plurality of adaptive convolution kernels; and

if any of the second training errors is higher than a threshold,

back-propagating the second training errors through the pre-trained CNN and the adaptive CNN to further refine the parameters in the candidate set; and

transferring at least one of the further refined parameters to the ensemble set.

22. The system of claim 18, wherein each of the adaptive convolution kernels is trained under a different loss criterion.

23. The system of claim 19, wherein the adaptive CNN further comprises a mask layer linked to the second convolution layer to further reduce a correlation among the sub-feature maps.

24. The system of claim 23, wherein the mask layer comprises a plurality of binary masks, each of which is convolved with one of the sub-feature maps and has a same spatial size with the convolved sub-feature map.

25. A system for an object online tracking, comprising:
a memory that stores executable components; and
a processor electrically coupled to the memory to execute the executable components
for:

determining a region of interest (ROI) in a first frame of a target video;

feeding the determined ROI forward through a pre-trained CNN to extract an initial feature map thereof;

initializing, with the initial feature map, an adaptive CNN used for detecting a location of the object and a scale estimation network used for defining a scale of the object;

predicting, with the initialized adaptive CNN, a second location of the object in a second frame of the target video, wherein the second frame is subsequent to the first frame;

estimating, with the initialized scale estimation network, a second scale of the object in the second frame of the target video;

updating, with optimized network parameters acquired in the predicting and the estimating, the adaptive CNN and the scale estimation network, respectively;

predicting, with the updated adaptive CNN, a third location of the object in a third frame of the target video, wherein the third frame is subsequent to the second frame; and

estimating, with the updated scale estimation network, a third scale of the object in the third frame of the target video.

26. The system of claim 25, wherein the adaptive CNN comprises:

a first convolution layer, linked to the pre-trained CNN and configured to transform a first feature map into a plurality of sub-feature maps, wherein the first feature map is generated by the pre-trained CNN according to any frame of the target video; and

a second convolution layer, linked to the first convolution layer and configured to convolve each of the sub-feature maps with one of a plurality of adaptive convolution kernels, respectively, to output a plurality of second feature maps with improved adaptability.

27. The system of claim 26, wherein the adaptive CNN is initialized by:

feeding a first training sample forward through the pre-trained CNN and the adaptive CNN to generate a first output image, wherein the first training sample is obtained according to a first frame of the target video;

comparing the first output image with a first ground truth derived from the first frame to

obtain a plurality of first training errors for the plurality of adaptive convolution kernels, respectively;

back-propagating repeatedly the first training errors through the pre-trained CNN and the adaptive CNN to train the adaptive convolution kernels until the first training errors converges, wherein a plurality of parameters are obtained for trained adaptive convolution kernels, respectively;

grouping a parameter of the obtained parameters, which has a smallest first training error, and a rest of the obtained parameters into an ensemble set and a candidate set, respectively.

28. The system of claim 26, wherein the adaptive CNN is updated by:

feeding the second training sample forward through the pre-trained CNN and the adaptive CNN to generate a second output image, wherein the second training sample is obtained according to a second frame of the target video and the second frame is subsequent to the first frame;

comparing the second output image with a second ground truth derived from the second frame to obtain a plurality of second training errors for the plurality of adaptive convolution kernels, respectively; and

if any of the second training errors is higher than a threshold,

back-propagating the second training errors through the pre-trained CNN and the adaptive CNN to further refine the parameters in the candidate set; and

transferring at least one of the further refined parameters to the ensemble set.

29. The system of claim 26, wherein each of the adaptive convolution kernels is trained under a different loss criterion.

30. The system of claim 26, wherein the adaptive CNN further comprises a mask layer linked to the second convolution layer to further reduce a correlation among the sub-feature maps.

31. The system of claim 30, wherein the mask layer comprises a plurality of binary masks, each of which is convolved with one of the sub-feature maps and has a same spatial size with the convolved sub-feature map.

32. The system of claim 26, wherein the location of the object is predicted by a heat map generated by the adaptive CNN, wherein a location with a maximum value is predicted to be the location of the object and the maximum value is sampled as a confidence.

33. The system of claim 32, wherein the updating is performed only if the confidence is higher than a pre-defined threshold.

34. The system of claim 25, wherein the ROI is centered at an object to be tracked.

35. An apparatus for adapting a pre-trained CNN to a target video, comprising:
means for transforming a first feature map into a plurality of sub-feature maps, wherein the first feature map is generated by the pre-trained CNN according to a frame of the target video;

means for convolving each of the sub-feature maps with one of a plurality of adaptive convolution kernels, respectively, to output a plurality of second feature maps with improved adaptability; and

means for training, frame by frame, the adaptive convolution kernels.

36. The apparatus of claim 35, wherein the means for transforming and the means for convolving are organized in an adaptive CNN comprising:

a first convolution layer, linked to the pre-trained CNN and configured to transform the first feature map into the plurality of sub-feature maps; and

a second convolution layer, linked to the first convolution layer and configured to convolve each of the sub-feature maps with one of the adaptive convolution kernels,

respectively.

37. The apparatus of claim 36, wherein the training comprises:

feeding a first training sample forward through the pre-trained CNN and the adaptive CNN to generate a first output image, wherein the first training sample is obtained according to a first frame of the target video;

comparing the generated first output image with a first ground truth derived from the first frame to obtain a plurality of first training errors for the adaptive convolution kernels, respectively;

back-propagating repeatedly the first training errors through the pre-trained CNN and the adaptive CNN to train the adaptive convolution kernels until the first training errors converge, wherein a plurality of parameters are obtained for the trained adaptive convolution kernels, respectively;

grouping a parameter of the obtained parameters, which has a smallest first training error, and a rest of the obtained parameters into an ensemble set and a candidate set, respectively; and

optimizing, according to a subsequent frame of the target video, the parameters grouped in the candidate set.

38. The apparatus of claim 37, wherein the optimizing comprises:

feeding a second training sample forward through the pre-trained CNN and the adaptive CNN to generate a second output image, wherein the second training sample is obtained according to a second frame of the target video and the second frame is subsequent to the first frame;

comparing the second output image with a second ground truth derived from the second frame to obtain a plurality of second training errors for the plurality of adaptive convolution kernels; and

if any of the second training errors is higher than a threshold,

back-propagating the second training errors through the pre-trained CNN

and the adaptive CNN to further refine the parameters in the candidate set; and
transferring at least one of the further refined parameters to the ensemble
set.

39. The apparatus of claim 35, wherein each of the adaptive convolution kernels is trained under a different loss criterion.

40. The apparatus of claim 36, wherein the adaptive CNN further comprises a mask layer linked to the second convolution layer to further reduce a correlation among the sub-feature maps.

41. The apparatus of claim 40, wherein the mask layer comprises a plurality of binary masks, each of which is convolved with one of the sub-feature maps and has a same spatial size with the convolved sub-feature map.

42. An apparatus for an object online tracking, comprising:
a feature extraction unit configured for:
determining a region of interest (ROI) in a first frame of a target video; and
feeding the determined ROI forward through a pre-trained CNN to extract an initial feature map thereof;

an initialization and update unit configured for initializing, with the initial feature map, an adaptive CNN used for detecting a location of the object and a scale estimation network used for defining a scale of the object;

a location prediction unit configured for predicting, with the initialized adaptive CNN, a second location of the object in a second frame of the target video, wherein the second frame is subsequent to the first frame; and

a scale estimation unit configured for estimating, with the initialized scale estimation network, a second scale of the object in the second frame of the target video;

wherein, the initialization and update unit is further configured for updating, with

optimized network parameters acquired in the predicting and the estimating, the adaptive CNN and the scale estimation network, respectively;

the location prediction unit is further configured for predicting, with the updated adaptive CNN, a third location of the object in a third frame of the target video, wherein the third frame is subsequent to the second frame; and

the scale estimation unit is further configured for estimating, with the updated scale estimation network, a third scale of the object in the third frame of the target video.

43. The apparatus of claim 42, wherein the adaptive CNN comprises:

a first convolution layer linked to the pre-trained CNN and configured to transform a first feature map into a plurality of sub-feature maps, wherein the first feature map is generated by the pre-trained CNN according to any frame of the target video; and

a second convolution layer linked to the first convolution layer and configured to convolve each of the sub-feature maps with one of a plurality of adaptive convolution kernels, respectively, to output a plurality of second feature maps with improved adaptability.

44. The apparatus of claim 43, wherein the adaptive CNN is initialized by:

feeding a first training sample forward through the pre-trained CNN and the adaptive CNN to generate a first output image, wherein the first training sample is obtained according to a first frame of the target video;

comparing the generated first output image with a first ground truth derived from the first frame to obtain a plurality of first training errors for the adaptive convolution kernels, respectively;

back-propagating repeatedly the first training errors through the pre-trained CNN and the adaptive CNN to train the adaptive convolution kernels until the first training errors converges, wherein a plurality of parameters are obtained for the trained adaptive convolution kernels, respectively;

grouping a parameter of the obtained parameters, which has a smallest first training error, and a rest of the obtained parameters into an ensemble set and a candidate set, respectively.

45. The apparatus of claim 43, wherein the adaptive CNN is updated by:

feeding a second training sample forward through the pre-trained CNN and the adaptive CNN to generate a second output image, wherein the second training sample is obtained according to a second frame of the target video and the second frame is subsequent to the first frame;

comparing the second output image with a second ground truth derived from the second frame to obtain a plurality of second training errors for the plurality of adaptive convolution kernels, respectively; and

if any of the second training errors is higher than a threshold,

back-propagating the second training errors through the pre-trained CNN and the adaptive CNN to further refine the parameters in the candidate set; and

transferring at least one of the further refined parameters to the ensemble set.

46. The apparatus of claim 42, wherein each of the adaptive convolution kernels is trained under a different loss criterion.

47. The apparatus of claim 43, wherein the adaptive CNN further comprises a mask layer linked to the second convolution layer to further reduce a correlation among the sub-feature maps.

48. The apparatus of claim 47, wherein the mask layer comprises a plurality of binary masks, each of which is convolved with one of the sub-feature maps and has a same spatial size with the convolved sub-feature map.

49. The apparatus of claim 43, wherein the location of the object is predicted by a heat map generated by the adaptive CNN, wherein a location with a maximum value is predicted to be the location of the object and the maximum value is sampled as a confidence.

50. The apparatus of claim 49, wherein the updating is performed only if the confidence is higher than a pre-defined threshold.

51. The apparatus of claim 42, wherein the ROI is centered at an object to be tracked.

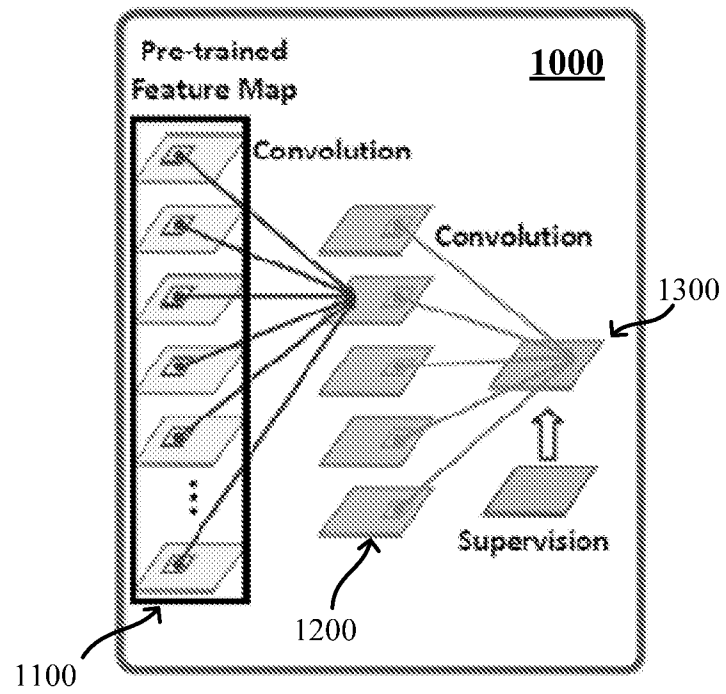
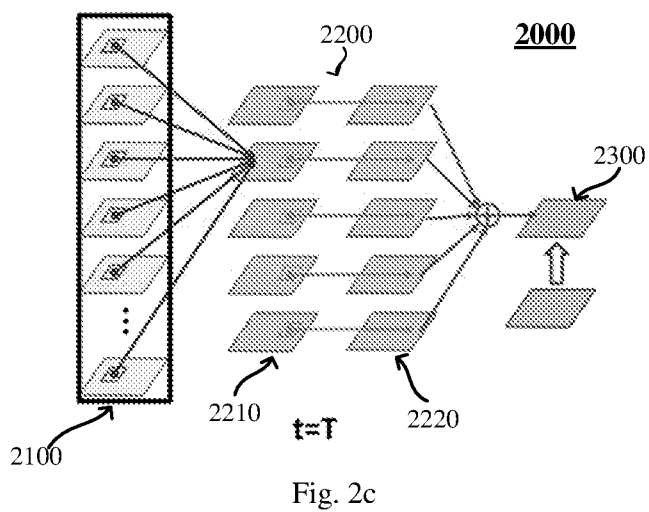
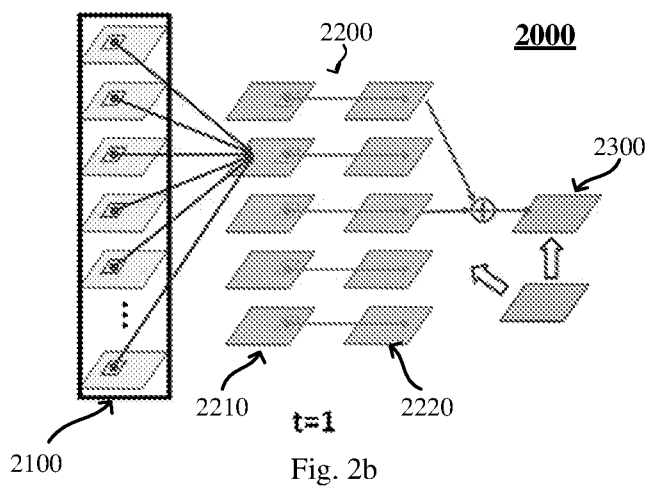
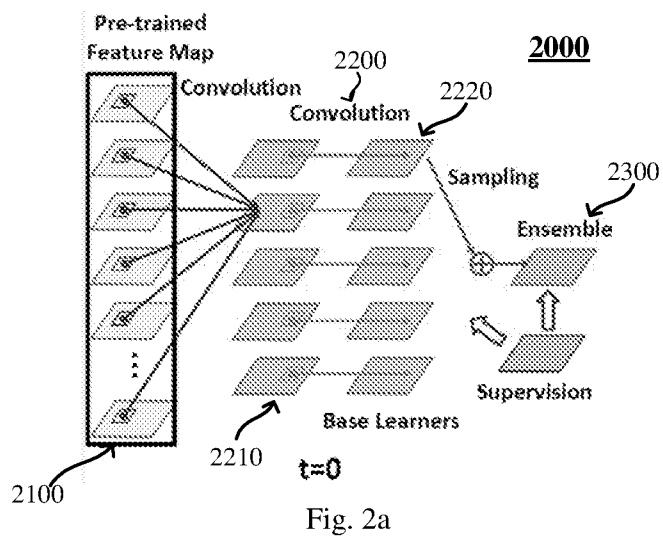


Fig. 1



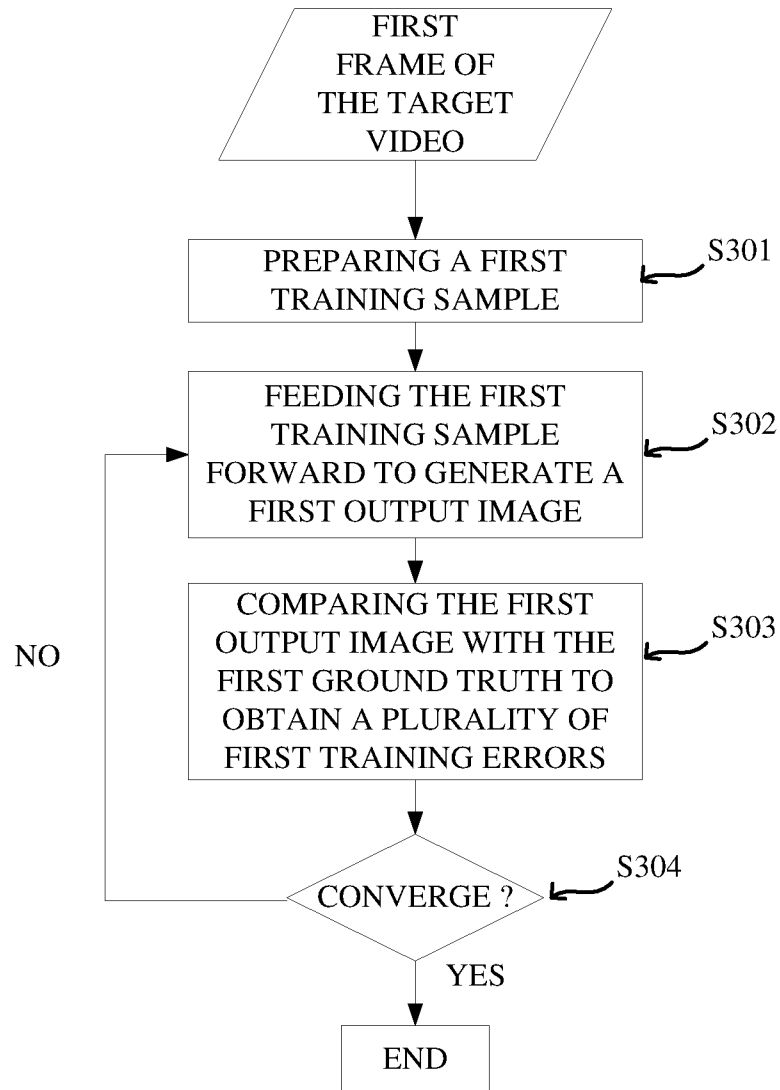


Fig. 3

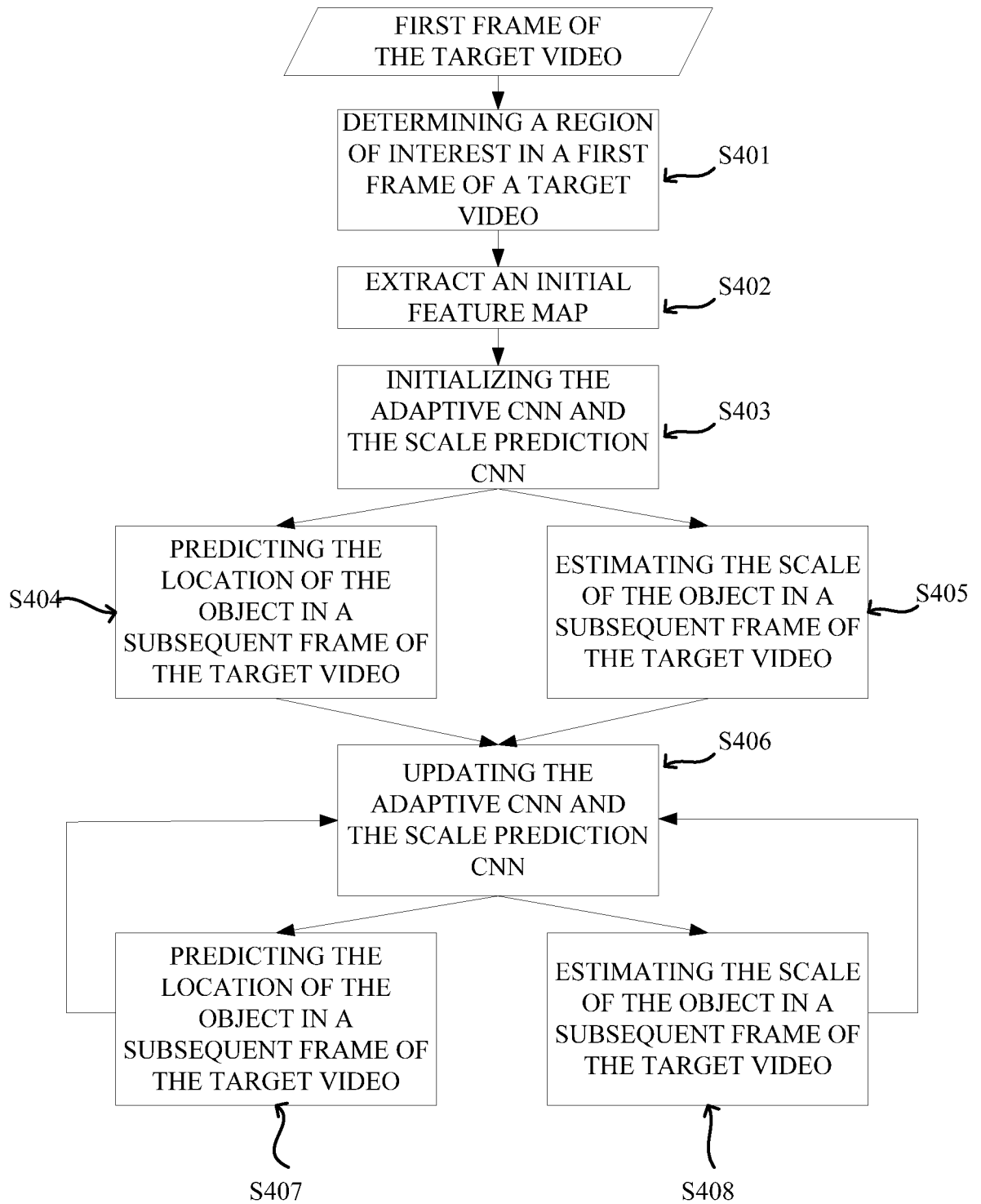


Fig. 4

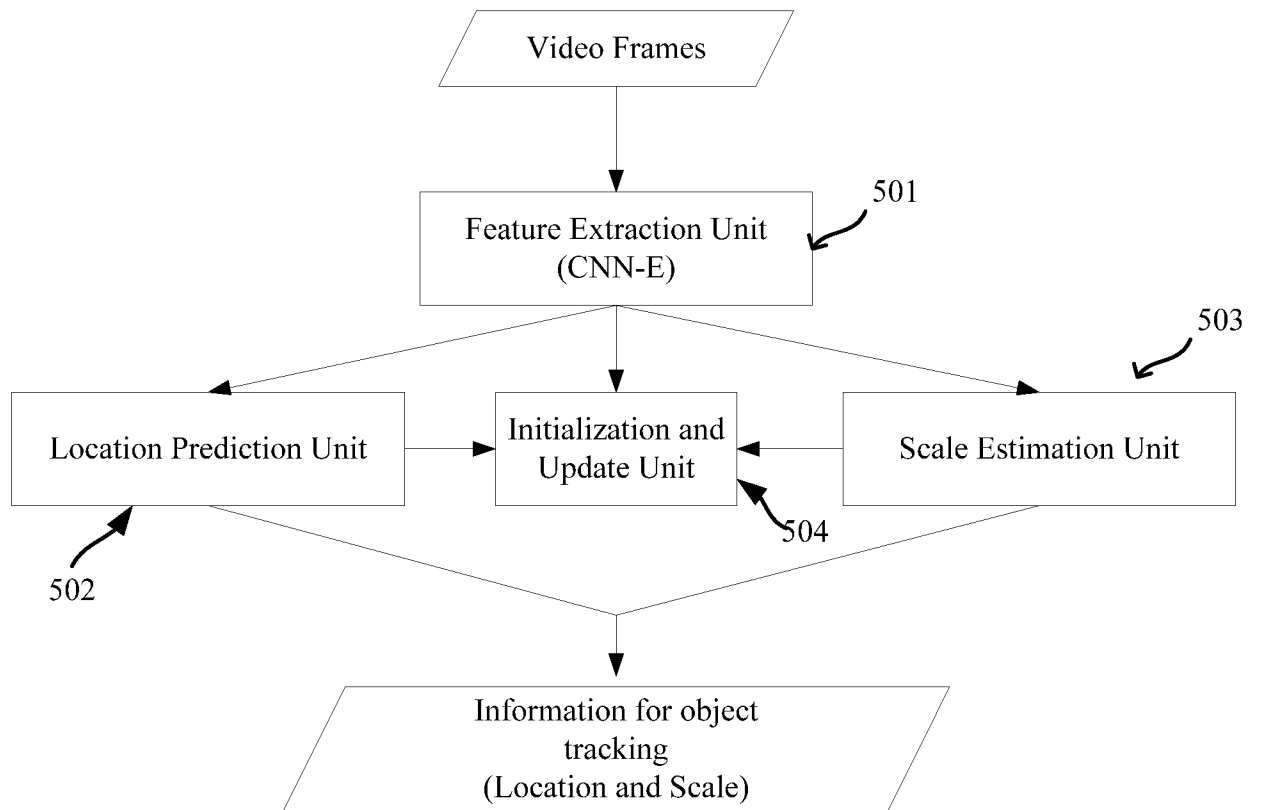


Fig. 5

INTERNATIONAL SEARCH REPORT

International application No.

PCT/CN2016/073184

A. CLASSIFICATION OF SUBJECT MATTER

G06K 9/00(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)

G06K, G06T

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

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

CNPAT,CNKI,WPI,EPODOC;GOOGLE;IEEE:visual, convolution, two, layer?, online, adapt+, kernel?, cnn, over, fitting, network, second, tracking, neural, convolv+, visual, updat+, sub+, feature, map

C. DOCUMENTS CONSIDERED TO BE RELEVANT

Category*	Citation of document, with indication, where appropriate, of the relevant passages	Relevant to claim No.
A	LI Huanyu et al. "Research on Visual Tracking Algorithm Based on Deep Feature Expression and Learning" <i>Journal of Electronics & Information Technology</i> , Vol. 37, No. 9, 30 September 2015 (2015-09-30), the abstract, chapters1-3	1-51
A	WANG Lijun et al. "Visual Tracking with Fully Convolutional Networks" <i>2015 IEEE International Conference on Computer Vision</i> , 31 December 2015 (2015-12-31), the whole document	1-51
A	US 8582807 B2 (NEC LABORATORIES AMERICA, INC.) 12 November 2013 (2013-11-12) the whole document	1-51
A	CN 102054170 A (INSTITUTE OF AUTOMATION, CHINESE ACADEMY OF SCIENCES) 11 May 2011 (2011-05-11) the whole document	1-51

☐ 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

20 October 2016

Date of mailing of the international search report

09 November 2016

Name and mailing address of the ISA/CN

STATE INTELLECTUAL PROPERTY OFFICE OF THE
P.R.CHINA
6, Xitucheng Rd., Jimen Bridge, Haidian District, Beijing
100088
China

Authorized officer

YAN,Jie

Facsimile No. (86-10)62019451

Telephone No. (86-10)62413393

INTERNATIONAL SEARCH REPORT
Information on patent family members

International application No.

PCT/CN2016/073184

Patent document cited in search report			Publication date (day/month/year)	Patent family member(s)			Publication date (day/month/year)
US	8582807	B2	12 November 2013	US	2011222724	A1	15 September 2011
CN	102054170	A	11 May 2011	CN	102054170	B	31 July 2013