



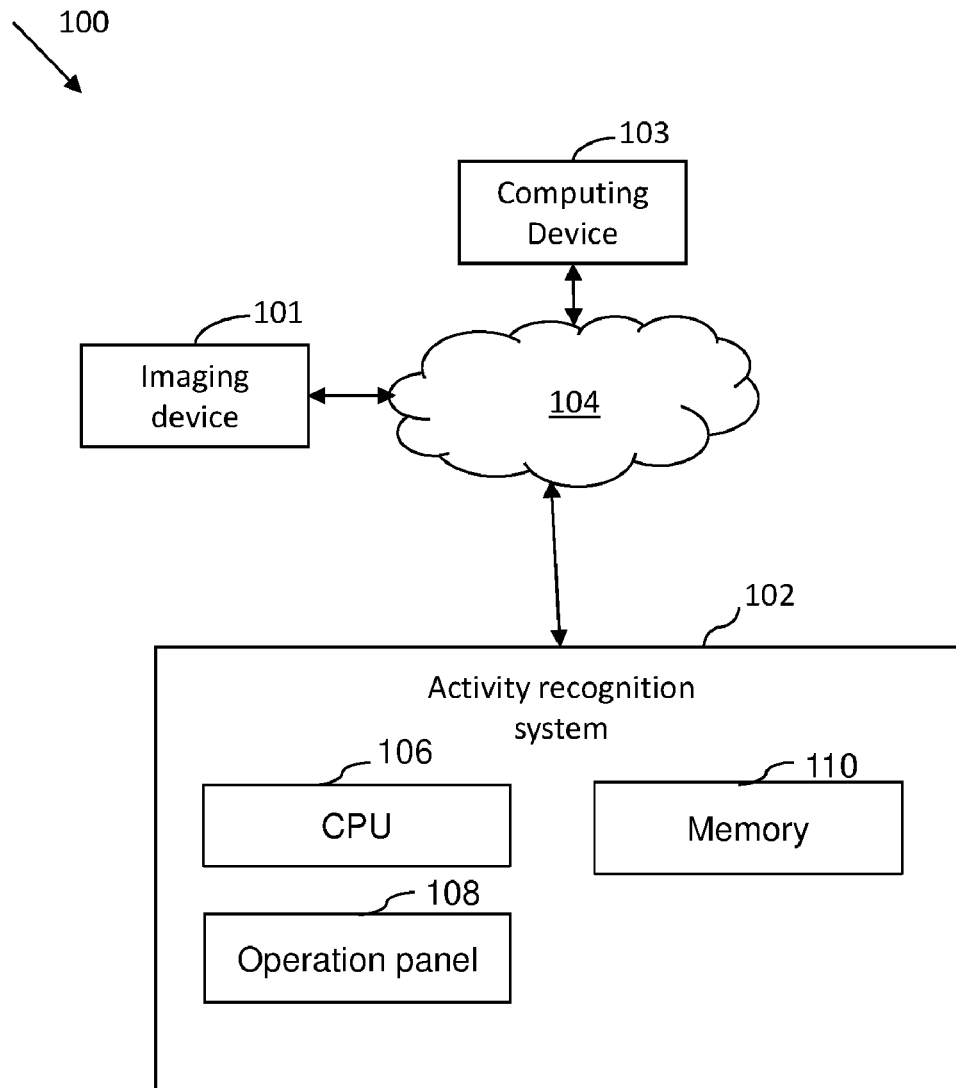
US 20210004575A1

(19) **United States**(12) **Patent Application Publication**
Pescaru et al.(10) **Pub. No.: US 2021/0004575 A1**(43) **Pub. Date: Jan. 7, 2021**(54) **QUANTIZED TRANSITION CHANGE
DETECTION FOR ACTIVITY
RECOGNITION**(52) **U.S. Cl.**CPC *G06K 9/00335* (2013.01); *G06K 9/00711*
(2013.01); *G06K 9/6262* (2013.01); *G06K*
9/6267 (2013.01); *G06K 9/00362* (2013.01)(71) Applicant: **Everseen Limited**, Blackpool (IE)

(57)

ABSTRACT(72) Inventors: **Dan Pescaru**, Timisoara (RO); **Cosmin
Cernazanu-glavan**, Timisoara Timis
(RO); **Vasile Gui**, Timisoara (RO)

A system for recognizing human activity from a video stream includes a classifier for classifying an image frame of the video stream in one or more classes and generating a class probability vector for the image frame based on the classification. The system further includes a data filtering and binarization module for filtering and binarizing each probability value of the class probability vector based on a pre-defined probability threshold value. The system furthermore includes a compressed word composition module for determining one or more transitions of one or more classes in consecutive image frames of the video stream and generating a sequence of compressed words based on the determined one or more transitions. The system furthermore includes a sequence dependent classifier for extracting one or more user actions by analyzing the sequence of compressed words to and recognizing human activity therefrom.

(21) Appl. No.: **16/458,288**(22) Filed: **Jul. 1, 2019****Publication Classification**(51) **Int. Cl.***G06K 9/00* (2006.01)
G06K 9/62 (2006.01)

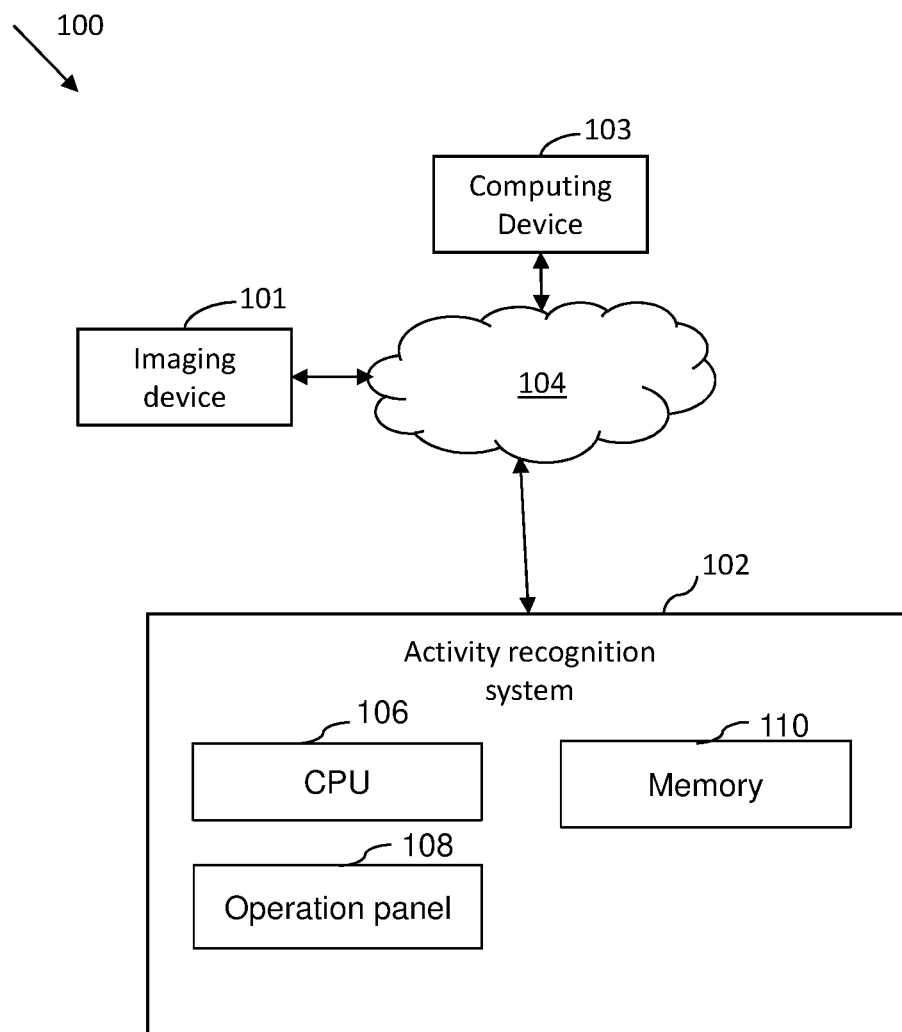


FIG. 1

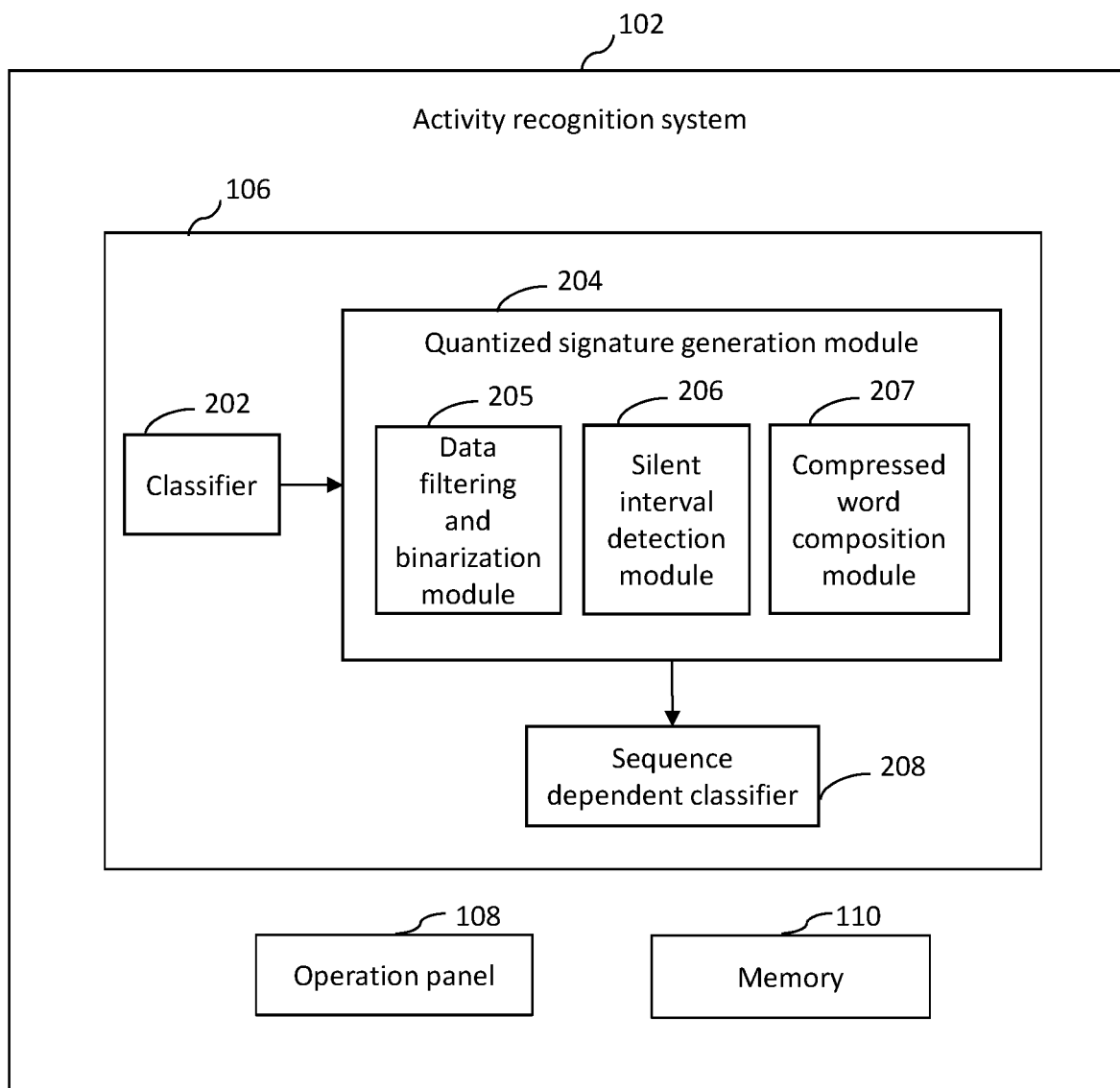


FIG. 2

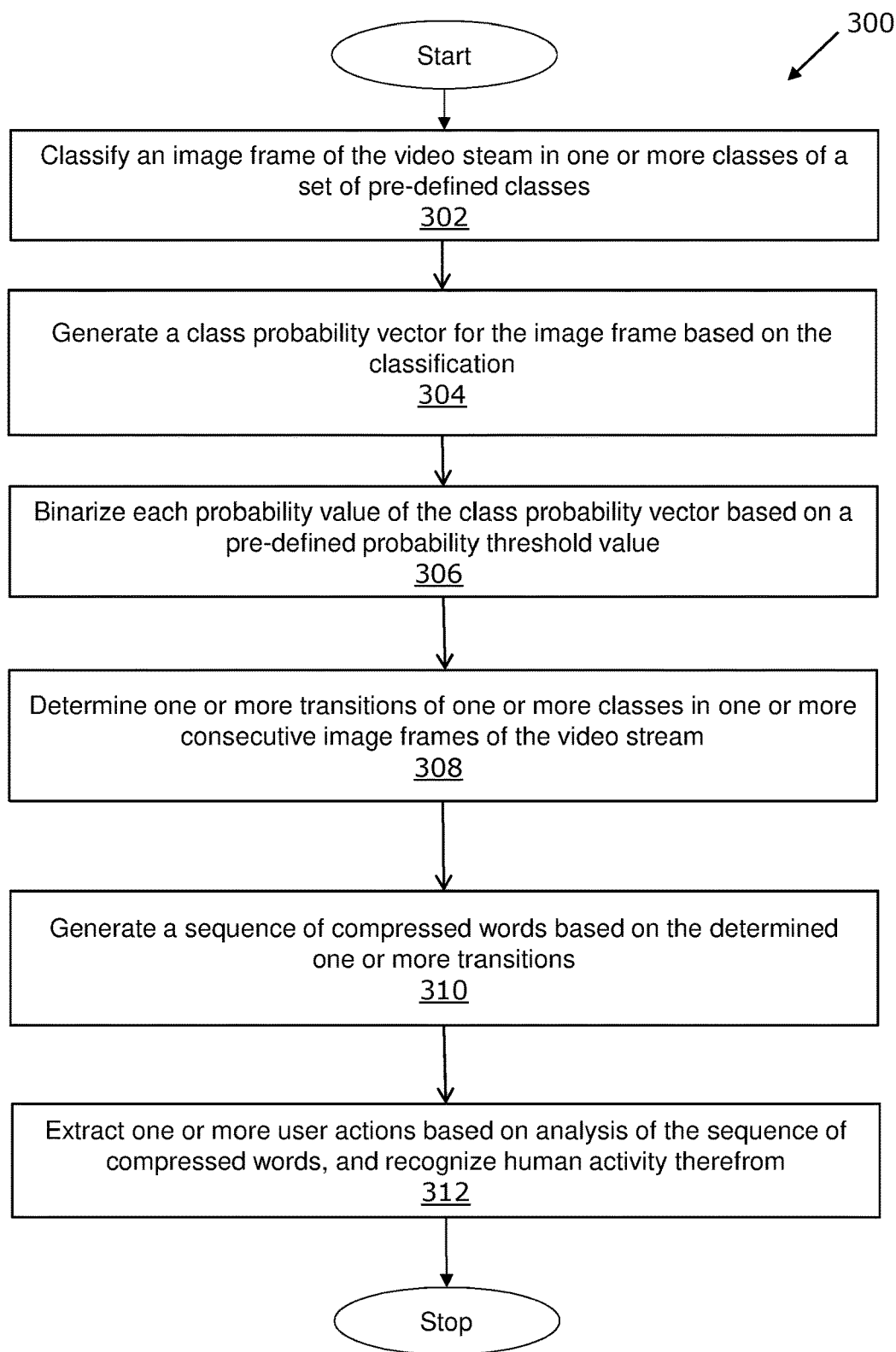


FIG. 3

QUANTIZED TRANSITION CHANGE DETECTION FOR ACTIVITY RECOGNITION

TECHNICAL FIELD

[0001] The present disclosure relates generally to artificial intelligence, and more specifically, to human activity recognition from a video stream and symbolic processing.

BACKGROUND

[0002] With advancement in technology, recognition of human physical activities is gaining tremendous importance. The recognition of human physical activities contributes towards various applications such as surveillance of a retail store check-out process involving a self-check out (SCO) system. Such a system allows buyers to complete a process of purchasing by themselves. Another example of application of recognition of human physical activities is providing assistance in video surveillance by detecting unfair activities done by shop lifters such as theft and thereby alerting a personnel employed in the shop to prevent the theft. Moreover, recognition of human physical activities is employed in intelligent driver assisting systems, assisted living systems for humans in need, video games, physiotherapy, and so forth. Furthermore, recognition of human physical activities is actively used in the field of sports, military, medical, robotics and so forth.

[0003] Human physical activities represent the building blocks of most process modelling. However, as human behaviour is unpredictable, the recognition of such human physical activities in a diverse environment is a difficult task. The human physical activity is typically decomposable into a set of basic actions involving various human body parts, such as hands, feet, face, and so forth. Moreover, the set of basic actions associated with the human physical activity are spanned over a plurality of time intervals. Recognition tasks of such activities face the problem of summarizing the overall sequence of actions over a variable time interval.

[0004] The conventional human physical activity recognition techniques are inefficient in recognizing the human physical activities, due to a different body structure, a different body shape, a different skin colour and so forth of each human body. Also, the time frame for a human activity pose important variation in time depending on the subject, and maybe other environment conditions. Moreover, not all the basic body parts movements are related with the purpose of the considered activity. Therefore, the activity recognition process face two major problems related with actions time variation and physical trajectory variation of human body parts involved in the activity.

[0005] Therefore, in light of the foregoing discussion, there exists a need to overcome the aforementioned drawbacks associated with the recognition of human physical activities, and provide a system and method that aims to reduce the influence of time variation and the variety of body parts movements in activity recognition using a recurrent neural network.

SUMMARY

[0006] The present disclosure seeks to provide a system for recognizing human activity from a video stream and a method thereof.

[0007] According to an aspect of the present disclosure, there is provided a system for recognizing human activity from a video stream captured by an imaging device. The system includes a memory to store one or more instructions, and a processor communicatively coupled to the memory. The system includes a classifier communicatively coupled to the imaging device, and configured to classify an image frame of the video stream in one or more classes of a set of pre-defined classes, wherein the image frame is classified based on user action in a region of interest of the image frame, and generate a class probability vector for the image frame based on the classification, wherein the class probability vector includes a set of probabilities of classification of the image frame in each pre-defined class. The system further includes a data filtering and binarization module configured to filter and binarize each probability value of the class probability vector based on a pre-defined probability threshold value. The system further includes a compressed word composition module configured to determine one or more transitions of one or more classes in one or more consecutive image frames of the video stream, based on corresponding binarized probability vectors, and generate a sequence of compressed words based on the determined one or more transitions in the one or more consecutive image frames. The system further includes a sequence dependent classifier configured to extract one or more user actions by analyzing the sequence of compressed words, and recognize human activity therefrom.

[0008] According to another aspect of the present disclosure, there is provided a method for recognizing human activity from a video stream. The method includes classifying by a classifier, an image frame of the video stream in one or more classes of a set of pre-defined classes, wherein the image frame is classified based on user action in a region of interest of the image frame. The method further includes generating a class probability vector for the image frame based on the classification, wherein the class probability vector includes a set of probabilities of classification of the image frame in each pre-defined class. The method furthermore includes binarizing each probability value of the class probability vector based on a pre-defined probability threshold value. The method furthermore includes determining one or more transitions of one or more classes in one or more consecutive image frames of the video stream, based on corresponding binarized probability vectors. The method furthermore includes generating a sequence of compressed words based on the determined one or more transitions in the one or more consecutive image frames. The method furthermore includes extracting one or more user actions by analyzing the sequence of compressed words to, and recognize human activity therefrom.

[0009] According to yet another aspect of the present disclosure, there is provided a computer programmable product for recognizing human activity from a video stream, the computer programmable product comprising a set of instructions. The set of instructions when executed by a processor causes the processor to classify an image frame of the video stream in one or more classes of a set of pre-defined classes, wherein the image frame is classified based on user action in a region of interest of the image frame, generate a class probability vector for the image frame based on the classification, wherein the class probability vector includes a set of probabilities of classification of the image frame in each pre-defined class, binarize each probability value of the

class probability vector based on a pre-defined probability threshold value, determine one or more transitions of one or more classes in one or more consecutive image frames of the video stream, based on corresponding binarized probability vectors, generate a sequence of compressed words based on the determined one or more transitions in the one or more consecutive image frames, and extract one or more user actions by analyzing the sequence of compressed words to extract one or more user actions, and recognize human activity therefrom.

[0010] The present disclosure seeks to provide a system for recognizing human activity from a video stream. Such a system enables efficient and reliable recognition of human activities from the video stream.

[0011] It will be appreciated that features of the present disclosure are susceptible to being combined in various combinations without departing from the scope of the present disclosure as defined by the appended claims.

DESCRIPTION OF THE DRAWINGS

[0012] The summary above, as well as the following detailed description of illustrative embodiments, is better understood when read in conjunction with the appended drawings. For the purpose of illustrating the present disclosure, exemplary constructions of the disclosure are shown in the drawings. However, the present disclosure is not limited to specific methods and instrumentalities disclosed herein. Moreover, those in the art will understand that the drawings are not to scale. Wherever possible, like elements have been indicated by identical numbers.

[0013] Embodiments of the present disclosure will now be described, by way of example only, with reference to the following diagrams wherein:

[0014] FIG. 1 illustrates an environment, wherein various embodiments of the present disclosure can be practiced;

[0015] FIG. 2 illustrates the activity recognition system for recognizing one or more human actions and activity in the video stream captured by the imaging device of FIG. 1, in accordance with an embodiment of the present disclosure; and

[0016] FIG. 3 is a flowchart illustrating a method for recognizing human activity from a video stream, in accordance with an embodiment of the present disclosure.

[0017] In the accompanying drawings, an underlined number is employed to represent an item over which the underlined number is positioned or an item to which the underlined number is adjacent. A non-underlined number relates to an item identified by a line linking the non-underlined number to the item. When a number is non-underlined and accompanied by an associated arrow, the non-underlined number is used to identify a general item at which the arrow is pointing.

DESCRIPTION OF EMBODIMENTS

[0018] The following detailed description illustrates embodiments of the present disclosure and ways in which they can be implemented. Although some modes of carrying out the present disclosure have been disclosed, those skilled in the art would recognize that other embodiments for carrying out or practicing the present disclosure are also possible.

[0019] FIG. 1 illustrates an environment 100, wherein various embodiments of the present disclosure can be prac-

ticed. The environment 100 includes an imaging device 101, an activity recognition system 102, and a computing device 103, communicatively coupled to each other through a communication network 104. The communication network 104 may be any suitable wired network, wireless network, a combination of these or any other conventional network, without limiting the scope of the present disclosure. Few examples may include a Local Area Network (LAN), wireless LAN connection, an Internet connection, a point-to-point connection, or other network connection and combinations thereof.

[0020] The imaging device 101 is configured to capture a video stream. In an embodiment of the present disclosure, the imaging device 101 is configured to capture one or videos of a retail check out process including a Selfcheck out system (SCO). Optionally, the imaging device 101 includes, but not limited to, an Internet protocol (IP) camera, a Pan-Tilt-Zoom (PTZ) camera, a thermal image camera or an Infrared camera.

[0021] The activity recognition system 102 is configured to recognize human actions and human activities in the video stream captured by the imaging device 101.

[0022] The activity recognition system 102 includes a central processing unit (CPU) 106, an operation panel 108, and a memory 110. The CPU 106 is a processor, computer, microcontroller, or other circuitry that controls the operations of various components such as the operation panel 108, and the memory 110. The CPU 106 may execute software, firmware, and/or other instructions, for example, that are stored on a volatile or non-volatile memory, such as the memory 110, or otherwise provided to the CPU 106. The CPU 106 may be connected to the operation panel 108, and the memory 110, through wired or wireless connections, such as one or more system buses, cables, or other interfaces. In an embodiment of the present disclosure, the CPU 106 may include a custom Graphic processing unit (GPU) server software to provide realtime object detection and prediction, for all cameras on a local network.

[0023] The operation panel 108 may be a user interface for the image forming apparatus 100 and may take the form of a physical keypad or touchscreen. The operation panel 108 may receive inputs from one or more users relating to selected functions, preferences, and/or authentication, and may provide and/or receive inputs visually and/or audibly.

[0024] The memory 110, in addition to storing instructions and/or data for use by the CPU 106 in managing operation of the image forming apparatus 100, may also include user information associated with one or more users of the image forming apparatus 100. For example, the user information may include authentication information (e.g. username/password pairs), user preferences, and other user-specific information. The CPU 106 may access this data to assist in providing control functions (e.g. transmitting and/or receiving one or more control signals) related to operation of the operation panel 108, and the memory 110.

[0025] The imaging device 101 and the activity recognition system 102 may be controlled/operated by the computing device 103. Examples of the computing device 103 include a smartphone, a personal computer, a laptop, and the like. The computing device 103 enables the user/operator to view and save the videos captured by the imaging device 101, and access the videos/images processed by the activity recognition system 102. The computing device 103 may execute a mobile application of the activity recognition

system **102** so as to enable a user to access and process the video stream captured by the imaging device **101**.

[0026] In an embodiment, the camera **101**, the activity recognition system **102**, and the computing device **103** may be integrated in a single device, where the single device is a portable smartphone having a built-in camera and a display.

[0027] FIG. 2 illustrates the activity recognition system **102** for recognizing one or more human actions and activity in the video stream captured by the imaging device **101**, in accordance with an embodiment of the present disclosure.

[0028] The activity recognition system **102** includes the CPU **106** that includes a classifier **202** that is operable to analyze each frame of the video stream to determine at least one action region of interest, wherein the at least one region of interest comprise at least one object. The action region of interest refers to a rectangular area in each frame of the video stream, where in the at least one object is seen and one or more actions take place. In an example, the at least one object may be a person, objects such as clothing items, groceries, wallet and so forth, and one or more actions may include a person taking out wallet from its pocket, the person walking in a queue, the person swiping a credit card, and the like. Each action can be used as a building block for process model extraction, wherein a process can be expressed as a chain of actions.

[0029] In an embodiment of the present disclosure, the classifier **202** may be an algorithm-based classifier such as a convolutional neural network (CNN) trained to classify an image frame of the video of the SCO scan area (scanning action region of interest) in classes such as hand, object in hand, object, body part, empty scanner. The criteria for classification of an image frame in each class has been mentioned below:

[0030] Hand—The image frame shows human hand(s).

[0031] Object in hand—The image frame shows an object in a hand of the user.

[0032] Object—The image frame shows only object

[0033] Body part—The image frame shows a human body part

[0034] Empty scanner—The image frame shows only the empty scanner

[0035] The CNN as referred herein is defined as trained deep artificial neural networks that is used primarily to classify the at least one object in the at least one region of interest. Notably, they are algorithms that can identify faces, individuals, street signs, and the like. The term “neural network” as used herein can include a highly interconnected network of processing elements, each optionally associated with a local memory. In an example, the neural network may be a Kohonen map, a multi-layer perceptron, and so forth. Furthermore, the processing elements of the neural networks can be “artificial neural units”, “artificial neurons,” “neural units,” “neurons,” “nodes,” and the like. Moreover, the neuron can receive data from an input or one or more other neurons, process the data, and send processed data to an output or yet one or more other neurons. The neural network or one or more neurons thereof can be generated in either hardware, software, or a combination of hardware and software, and the neural network can be subsequently trained. It will be appreciated that the convolutional neural network (CNN) consists of an input layer, a plurality of hidden layers and an output layer. Moreover, the plurality of hidden layers of the convolutional neural network typically

consist of convolutional layers, pooling layers, fully connected layers and normalization layers. Optionally, a Visual Geometry Group 19 (VGG 19) model is used as a convolutional neural network architecture. The VGG 19 model is configured to classify the at least one object in the frame of the video stream into classes. It will be appreciated that hidden layers comprise a plurality of sets of convolution layers.

[0036] In operation, the classifier **202** receives and classifies an image frame of the video stream of the SCO scan area (scanning action region of interest) in classes such as hand, object in hand, object, body part, empty scanner based on content of the image frame. In an embodiment of the present disclosure, the classifier **202** analyses each image frame statically and for each image frame, outputs a class probability vector P_v having one component for each considered class, such that, $P_v = \{P_{Hand}, P_{HandObject}, P_{Object}, P_{BodyPart}, P_{EmptyScanner}\}$

Where P_{Hand} = Probability of the image frame to be classified in class ‘hand’

$P_{HandObject}$ = Probability of the image frame to be classified in class ‘object in hand’

P_{Object} = Probability of the image frame to be classified in class ‘object’

$P_{BodyPart}$ = Probability of the image frame to be classified in class ‘body part’

$P_{EmptyScanner}$ = Probability of the image frame to be classified in class ‘empty scanner’

[0037] In an example, the classifier **202** generates six probability vectors P_{v1} till P_{v6} for six consecutive image frames in five classes, in a format given below.

$$P_{v1} = \{0.0, 0.0, 0.0, 0.0, 1.0\}$$

$$P_{v2} = \{0.0, 0.0, 0.28, 0.0, 0.72\}$$

$$P_{v3} = \{0.0, 0.0, 0.26, 0.0, 0.74\}$$

$$P_{v4} = \{0.0, 0.0, 0.19, 0.0, 0.81\}$$

$$P_{v5} = \{0.0, 0.0, 0.29, 0.0, 0.71\} \quad P_{v6} = \{0.0, 0.45, 0.14, 0.0, 0.41\}$$

[0038] The CPU **106** further includes a quantized signature generation module **204** for generating a quantized signature for each scan action determined by the classifier **202**. A scan action is a user action performed for scanning an item in a scanning zone of a self-check out (SCO) terminal.

[0039] The quantized signature generation module **204** includes a data filtering and binarization module **205**, a silent interval detection module **206**, and a compressed word composition module **207**.

[0040] The data filtering and binarization module **205** is configured to apply a filter on the class probability vectors generated by the classifier **202** to minimize errors by the classifier **202**. A classifier error appears if the classifier **202** classifies a continuous movement on the scanner using a single class for the entire sequence except one isolated frame. In such case, the isolated frame may be wrongly classified.

[0041] Below is an example output of probability vectors from the classifier **202** for six consecutive image frames of the video stream, wherein the six consecutive image frames cover a continuous movement over the scanner. For an image frame in, each probability vector P_{vn} includes prob-

abilities of classification of the image frame in each of the five classes “hand”, “object in hand”, “object”, “body part”, and “empty scanner”.

$$P_{v1}=\{0.0,0.0,0.28,0.0,0.72\}$$

$$P_{v2}=\{0.0,0.0,0.28,0.0,0.72\}$$

$$P_{v3}=\{0.0,0.0,0.01,0.27,0.72\}$$

$$P_{v4}=\{0.0,0.0,0.28,0.0,0.72\}$$

$$P_{v5}=\{0.0,0.0,0.28,0.0,0.72\}$$

$$P_{v6}=\{0.0,0.0,0.28,0.0,0.72\}$$

[0042] It can be clearly seen that the probability vector P_{v3} of the third image frame of the video sequence is different, which means that there is an error in the classification of the third image frame by the classifier **202**. The data filtering and binarization module **205** rectifies the error in the classification of the third image frame based on the information that the six frames cover substantially similar information. In an embodiment of the present disclosure, the data filtering and binarization module **205** rectifies the error by removing the erroneous frame.

[0043] The data filtering and binarization module **205** is then configured to binarize the filtered values of probability vectors using a heuristic threshold value, such that each component of a probability vector is assigned a value “1” if it is equal to or greater than the heuristic threshold value, else “0”.

[0044] In an example, when heuristic threshold value is 0.2, exemplary filtered probability vectors P_{vf} for five consecutive image frames may be represented as below:

$$P_{vf1}=\{0.0,0.0,1.0\}$$

$$P_{vf2}=\{0.0,0.0,0.28,0.0,0.72\}$$

$$P_{vf3}=\{0.0,0.0,0.26,0.0,0.74\}$$

$$P_{vf4}=\{0.0,0.0,0.39,0.0,0.71\}$$

$$P_{vf5}=\{0.0,0.45,0.14,0.0,0.41\}$$

and corresponding binarized probability vectors P_{vb} may be represented as below:

$$P_{vb1}=\{0\ 0\ 0\ 0\ 1\}$$

$$P_{vb2}=\{0\ 0\ 1\ 0\ 1\}$$

$$P_{vb3}=\{0\ 0\ 1\ 0\ 1\}$$

$$P_{vb4}=\{0\ 0\ 1\ 0\ 1\}$$

$$P_{vb5}=\{0\ 1\ 0\ 0\ 1\}$$

[0045] Each binarized probability vector P_{vb} is thus a binarized string of a series of binary numbers, that can be used to determine transitions of classes in consecutive frames. For example, in the first image frame, the binary value corresponding to class ‘object’ is ‘0’, and in the second image frame, the binary value corresponding to class ‘object’ is ‘1’, which means that there is clearly a transition of class from the first to second image frame. Similarly, in the fourth image frame, the binary value corresponding to class ‘object in hand’ is ‘0’, and the binary value corresponding to class ‘object’ is ‘1’. In the fifth frame, the binary

value for ‘object in hand’ changes to ‘1’, and the binary value for ‘object’ changes to ‘0’. This clearly indicates that the user has kept the object in their hand during transition from fourth to fifth frame. Thus, the binarized/quantized probability vectors provide information about transition of classes in consecutive image frames.

[0046] The silent interval detection module **206** is configured to detect one or more silent intervals in the video stream. In an embodiment of the present disclosure, during silent interval, no activity is detected in the scanning zone for a threshold time duration. In an example, the threshold time duration may be set as ‘0.5 s’, and a time interval of more than 0.5 s is marked as ‘silent interval’ when the binary value of class “empty scanner” of corresponding image frames remains ‘1’ during the entire time interval.

[0047] The compressed word composition module **207** is configured to generate a sequence of compressed words based on the binarized strings generated by the data filtering and binarization module **205**. The compressed words are generated based on the transition of classes from ‘1’ to ‘0’ and ‘0’ to ‘1’ in consecutive image frames.

[0048] In an embodiment of the present disclosure, each word is composed from letters of an alphabet containing $2*N$ letters correlated with the process actions semantics, where N represents the number of classes. In an example, if the number of classes is 5, then each word is composed from total 10 letters. For each class a “0->1” transition generates a specific “beginning” letter (e.g. ‘O’ for the class Object), while a “1->0” transition generates an “ending” letter (e.g. ‘o’ for the class Object).

[0049] Thus, the alphabet for five classes: ‘hand’, ‘object in hand’, ‘object’, ‘body part’, and ‘empty scanner’, contains the following letters:

```
classHand up:H down:h
classHandObject up:Q down:q
classObject up:O down:o
classBodyPart up: B down: b
classEmptyScanner up: E down: e
```

[0050] In an embodiment of the present disclosure, two adjacent words are separated by at least one frame classified as “empty scanner”. This could represent or not a silent interval depending on the length of consecutive ‘1’ ‘empty scanner’ values.

[0051] An example of quantized output generated by the compressed word composition module **207** is represented below:

[0052] Silence

[0053] OoE

[0054] Silence

[0055] OQoOqBobE

[0056] Silence

[0057] The sequence dependent classifier **208** is configured to receive the quantized output from the compressed word composition module **207**, and extract one or more scan actions from the continuous sequence of transitions represented as alphabet letters. The sequence dependent classifier **208** includes a machine learning based engine, as used herein relates to an engine that is capable of studying of algorithms and statistical models and use them to effectively perform a specific task without using explicit instructions, relying on patterns and inference. Examples of the sequence dependent classifier **208** include a recurrent neural network (RNN), a K nearest neighbor algorithm (KNN), and a support vector machine (SVM) algorithm, and so forth.

[0058] The sequence dependent classifier **208** analyzes the sequence of compressed words to recognize the human activity from the video stream. The sequence of compressed words is analyzed in order to determine various transitions of the classes in the region of interest. Such determination of the transitions of the classes leads to the recognition of the human activity from the video stream. The sequence dependent classifier **208** recognize transitions of the binarized input signal which suggest basic actions.

[0059] Thus, the quantized signature generation module **204** provides a quantization process for input signals coming from the classifier **202** observing a region of interest where an activity take place. The method for transitions quantization aims to reduce the influence of time variation and the variety of body parts movements in activity recognition using the sequence dependent classifier **208**.

[0060] FIG. 3 is a flowchart illustrating a method **300** for recognizing human activity from a video stream, in accordance with an embodiment of the present disclosure. Some steps may be discussed with respect to the system as shown in FIG. 2.

[0061] At step **302**, an image frame of the video stream in one or more classes of a set of pre-defined classes is classified by a classifier, wherein the image frame is classified based on user action in a region of interest of the image frame. In an embodiment of the present disclosure, the classifier is a convolutional neural network. In another embodiment of the present disclosure, the set of predefined classes for a Self-check out (SCO) scanning zone, include classes such as hand, object in hand, object, body part, and empty scanner.

[0062] At step **304**, a class probability vector is generated for the image frame based on the classification, wherein the class probability vector includes a set of probabilities of classification of the image frame in each pre-defined class. In an example, a class probability vector P_v is represented by:

$$P_v = \{P_{Hand}, P_{HandObject}, P_{Object}, P_{BodyPart}, P_{EmptyScanner}\}$$

Where P_{Hand} = Probability of the image frame to be classified in class 'hand'

$P_{HandObject}$ = Probability of the image frame to be classified in class 'object in hand'

P_{Object} = Probability of the image frame to be classified in class 'object'

$P_{BodyPart}$ = Probability of the image frame to be classified in class 'body part'

$P_{EmptyScanner}$ = Probability of the image frame to be classified in class 'empty scanner'

[0063] At step **306**, each probability value of the class probability vector is binarized based on a pre-defined probability threshold value. In an example, each component of a probability vector is assigned a value "1" if it is equal to or greater than the heuristic threshold value, else "0".

[0064] At step **308**, one or more transitions of one or more classes are determined in one or more consecutive image frames of the video stream, based on corresponding binarized probability vectors. For example, if in the first image frame, the binary value corresponding to class 'object' is '0', and in the second image frame, the binary value corresponding to class 'object' is '1', which means that there is clearly a transition of class from the first to second image frame.

[0065] At step **310**, a sequence of compressed words is generated based on the determined one or more transitions

in the one or more consecutive image frames. The compressed words are generated based on the transition of classes from '1' to '0' and '0' to '1' in consecutive image frames. In an embodiment of the present disclosure, a compressed word is formed from letters of an alphabet containing number of letters equivalent to twice the number of pre-defined classes. Further, each of the compressed word of the sequence of compressed words comprise at least one frame of non-activity therebetween. In an example, if the number of classes is 5, then each word is composed from total 10 letters. For each class a "0->1" transition generates a specific "beginning" letter (e.g. 'O' for the class Object), while a "1->0" transition generates an "ending" letter (e.g. 'o' for the class Object).

[0066] At step **312**, one or more user actions are extracted based on analysis of the sequence of compressed words by a sequence dependent classifier. The one or more user actions may be used to recognize human activity in the SCO scan area (scanning action region of interest), and transmits the recognition results to a user computing device. In some embodiments, the user computing device may be configured to store or display the recognition results. In an embodiment of the present disclosure, the sequence dependent classifier is a recurrent neural network.

[0067] The present disclosure also relates to software products recorded on machine-readable non-transient data storage media, wherein the software products are executable upon computing hardware to implement methods of recognizing human activity from a video stream.

[0068] Modifications to embodiments of the invention described in the foregoing are possible without departing from the scope of the invention as defined by the accompanying claims. Expressions such as "including", "comprising", "incorporating", "consisting of", "have", "is" used to describe and claim the present invention are intended to be construed in a non-exclusive manner, namely allowing for items, components or elements not explicitly described also to be present. Reference to the singular is also to be construed to relate to the plural. Numerals included within parentheses in the accompanying claims are intended to assist understanding of the claims and should not be construed in any way to limit subject matter claimed by these claims.

1. A system for recognizing human activity from a video stream captured by an imaging device, the system comprising:

- a memory to store one or more instructions; and
- a processor communicatively coupled to the memory to execute the one or more instructions, wherein the processor comprises:
 - a classifier communicatively coupled to the imaging device, and configured to:

- classify an image frame of the video stream in one or more classes of a set of pre-defined classes, wherein the image frame is classified based on user action in a region of interest of the image frame; and

- generate a class probability vector for the image frame based on the classification, wherein the class probability vector includes a set of probabilities of classification of the image frame in each pre-defined class;

- a data filtering and binarization module configured to filter and binarize each probability value of the class probability vector based on a pre-defined probability threshold value;
- a compressed word composition module configured to: determine one or more transitions of one or more classes in one or more consecutive image frames of the video stream, based on corresponding binarized probability vectors; and
- generate a sequence of compressed words based on the determined one or more transitions in the one or more consecutive image frames; and
- a sequence dependent classifier configured to extract one or more user actions by analyzing the sequence of compressed words to, and recognize human activity therefrom.
2. The system as claimed in claim 1, wherein the classifier is a convolutional neural network.
3. The system as claimed in claim 1, wherein the set of predefined classes for a Self-check out (SCO) scanning zone, include classes such as hand, object in hand, object, body part, and empty scanner.
4. The system as claimed in claim 1, wherein the data filtering and binarization module is further operable to eliminate classifier errors in the class probability vectors of one or more consecutive image frames.
5. The system as claimed in claim 1, wherein the processor further comprises a silent interval detection module, wherein the silent interval detection module is configured to detect one or more silent intervals in the video stream based on no activity detection in the region of interest for a predefined threshold duration.
7. The system as claimed in claim 1, wherein a compressed word is formed from letters of an alphabet containing number of letters equivalent to twice the number of pre-defined classes.
8. The system as claimed in claim 1, wherein each of the compressed word of the sequence of compressed words comprise at least one frame of non-activity therebetween.
9. The system as claimed in claim 1, wherein the sequence dependent classifier is a recurrent neural network.
10. A method for recognizing human activity from a video stream, the method comprising
- classifying by a classifier, an image frame of the video stream in one or more classes of a set of pre-defined classes, wherein the image frame is classified based on user action in a region of interest of the image frame;
 - generating a class probability vector for the image frame based on the classification, wherein the class probability vector includes a set of probabilities of classification of the image frame in each pre-defined class;
 - binarizing each probability value of the class probability vector based on a pre-defined probability threshold value;
 - determining one or more transitions of one or more classes in one or more consecutive image frames of the video stream, based on corresponding binarized probability vectors;
 - generating a sequence of compressed words based on the determined one or more transitions in the one or more consecutive image frames; and
 - extracting one or more user actions by analyzing the sequence of compressed words by a sequence dependent classifier and recognizing human activity therefrom.
11. The method as claimed in claim 10, wherein the classifier is a convolutional neural network.
12. The method as claimed in claim 10, wherein the set of predefined classes for a Self-check out (SCO) scanning zone, include classes such as hand, object in hand, object, body part, and empty scanner.
13. The method as claimed in claim 10 further comprising eliminating classifier errors in the class probability vectors of one or more consecutive image frames.
14. The method as claimed in claim 10, further comprising detecting one or more silent intervals in the video stream based on no activity detection in the region of interest for a predefined threshold duration.
15. The method as claimed in claim 10, wherein a compressed word is formed from letters of an alphabet containing number of letters equivalent to twice the number of pre-defined classes.
16. The method as claimed in claim 10, wherein each of the compressed word of the sequence of compressed words comprise at least one frame of non-activity therebetween.
17. The method as claimed in claim 10, wherein the sequence dependent classifier is a recurrent neural network.
18. A computer programmable product for recognizing human activity from a video stream, the computer programmable product comprising a set of instructions, the set of instructions when executed by a processor causes the processor to:
- classify an image frame of the video stream in one or more classes of a set of pre-defined classes, wherein the image frame is classified based on user action in a region of interest of the image frame;
 - generate a class probability vector for the image frame based on the classification, wherein the class probability vector includes a set of probabilities of classification of the image frame in each pre-defined class;
 - binarize each probability value of the class probability vector based on a pre-defined probability threshold value;
 - determine one or more transitions of one or more classes in one or more consecutive image frames of the video stream, based on corresponding binarized probability vectors;
 - generate a sequence of compressed words based on the determined one or more transitions in the one or more consecutive image frames; and
 - extract one or more user actions by analyzing the sequence of compressed words to, and recognizing human activity therefrom.
19. The computer programmable product as claimed in claim 18, wherein a compressed word is formed from letters of an alphabet containing number of letters equivalent to twice the number of pre-defined classes.
20. The computer programmable product as claimed in claim 18, wherein each of the compressed word of the sequence of compressed words comprise at least one frame of non-activity therebetween.