



(22) Date de dépôt/Filing Date: 2018/05/01

(41) Mise à la disp. pub./Open to Public Insp.: 2019/11/01

(45) Date de délivrance/Issue Date: 2021/10/05

(51) Cl.Int./Int.Cl. *G06T 13/00* (2011.01),
G06F 3/14 (2006.01), *G06T 13/40* (2011.01),
G10L 15/187 (2013.01), *G10L 21/0272* (2013.01),
H04N 21/81 (2011.01)

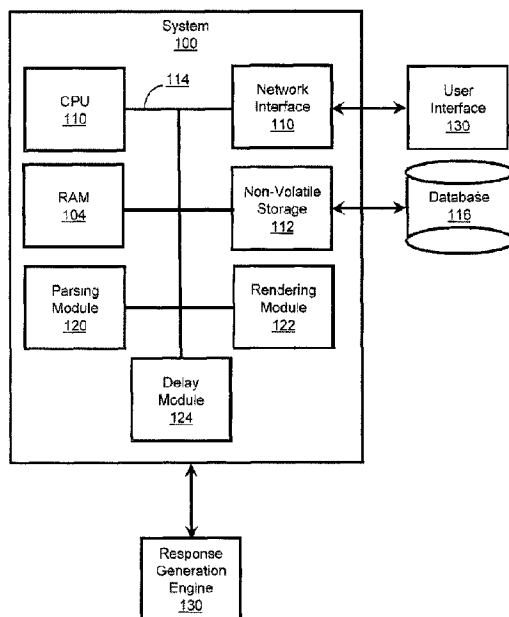
(72) Inventeurs/Inventors:
JENKIN, MICHAEL, CA;
TARAWNEH, ENAS, CA

(73) Propriétaires/Owners:
JENKIN, MICHAEL, CA;
TARAWNEH, ENAS, CA

(74) Agent: Bhole IP Law

(54) Titre : SYSTEME ET PROCEDE DE RENDU D'UN AVATAR ANIME

(54) Title: SYSTEM AND METHOD FOR RENDERING OF AN ANIMATED AVATAR



(57) Abrégé/Abstract:

There are provided systems and methods for rendering of an animated avatar. An embodiment of the method includes: determining a first rendering time of a first clip as approximately equivalent to a predetermined acceptable rendering latency, a first playing time of the first clip determined as approximately the first rendering time multiplied by a multiplicative factor; rendering the first clip; determining a subsequent rendering time for each of one or more subsequent clips, each subsequent rendering time is determined to be approximately equivalent to the predetermined acceptable rendering latency plus the total playing time of the preceding clips, each subsequent playing time is determined to be approximately the rendering time of the respective subsequent clip multiplied by the multiplicative factor; and rendering the one or more subsequent clips.

ABSTRACT

There are provided systems and methods for rendering of an animated avatar. An embodiment of the method includes: determining a first rendering time of a first clip as approximately equivalent to a predetermined acceptable rendering latency, a first playing time of the first clip determined as approximately the first rendering time multiplied by a multiplicative factor; rendering the first clip; determining a subsequent rendering time for each of one or more subsequent clips, each subsequent rendering time is determined to be approximately equivalent to the predetermined acceptable rendering latency plus the total playing time of the preceding clips, each subsequent playing time is determined to be approximately the rendering time of the respective subsequent clip multiplied by the multiplicative factor; and rendering the one or more subsequent clips.

1 SYSTEM AND METHOD FOR RENDERING OF AN ANIMATED AVATAR

2 TECHNICAL FIELD

3 [0001] The following is related generally to computer animation and more specifically
4 to a system and method for rendering of an animated avatar.

5 BACKGROUND

6 [0002] As robotics and internet-of-things (IOT) applications grow and become more
7 pervasive, human-machine interaction necessarily grows as well. Increasingly, this
8 interaction involves audio or oral interactions between a human user and an artificially
9 intelligent device; for example, oral interaction with an intelligent personal assistant
10 located in a smart speaker device. Generally, this interaction involves capturing the
11 audio signal of the user locally, sending this audio signal to a cloud computing resource,
12 utilizing a machine learning technique to digitally parse and identify words and phrases
13 in the audio signal, using a machine learning technique to build a response to the
14 sequence of words, and transmitting this to the human user and rendering it. In some
15 cases, in order to allow users to add their own concepts to the response system, hooks
16 can be programmed for application specific responses.

17 [0003] The above determined response can, in some cases, take the form of a
18 sequence of words or actions to be sent back to the local environment. Actions can be,
19 for example, to control IOT devices or to control an autonomous system. Where the
20 response is a sequence of words, a response can be delivered to the user, often via
21 computer-generated speech. In this case, the cloud computing resource can be used to
22 convert the words to an audio file via a computer-generated speech technique, the
23 audio file can be sent to the device local to the user, and the audio file can be played for
24 the user.

25 [0004] These applications are generally limited in that they only involve audio or text
26 interactions or interfaces, or IOT action responses.

27 SUMMARY

1 [0005] In an aspect, there is provided a method for rendering of an animated avatar
2 with a response on one or more computing devices, the method comprising: receiving
3 the response, the response comprising a plurality of pieces; determining a first
4 rendering time of a first clip of an animation of the avatar as approximately equivalent to
5 a predetermined acceptable rendering latency, the first clip comprising one or more
6 sequential pieces of the response, a first playing time of the first clip determined as
7 approximately the first rendering time multiplied by a multiplicative factor; rendering the
8 first clip of the animation of the avatar; determining a subsequent rendering time for
9 each of one or more subsequent clips of the animation of the avatar, each of the
10 subsequent clips comprising one or more sequential pieces of the response that
11 succeed the preceding clip of the animation of the avatar, each subsequent rendering
12 time is determined to be approximately equivalent to the predetermined acceptable
13 rendering latency plus the total playing time of the preceding clips, each subsequent
14 playing time is determined to be approximately the rendering time of the respective
15 subsequent clip multiplied by the multiplicative factor; and rendering the one or more
16 subsequent clips of the animation of the avatar.

17 [0006] In a particular case, the multiplicative factor being an approximation of the
18 ratio between a playing time of a representative clip and a rendering time of that
19 respective clip.

20 [0007] In another case, the first clip and at least one of the one or more subsequent
21 clips are rendered approximately contemporaneously.

22 [0008] In yet another case, the playing time of the first clip is reduced such that the
23 first clip ends at a natural break in speech in the response, and wherein the playing time
24 of each of the subsequent clips are reduced such that and each of the subsequent clips
25 end at a natural break in speech in the response.

26 [0009] In yet another case, the natural breaks in speech comprise a break between
27 words or at a punctuation.

1 [0010] In yet another case, the natural breaks in speech comprise a break that is
2 closest in time to the total respective rendering time of the respective clip.

3 [0011] In yet another case, each of the pieces are phonemes of the response.

4 [0012] In another aspect, there is provided a system for rendering of an animated
5 avatar displayable on a user interface with a response received from a response
6 generation engine, the system comprising one or more processors and a data storage
7 device, the one or more processors configured to execute: a parsing module to receive
8 the response, the response comprising a plurality of pieces; and a rendering module to:
9 determine a first rendering time of a first clip of an animation of the avatar as
10 approximately equivalent to a predetermined acceptable rendering latency, the first clip
11 comprising one or more sequential pieces of the response, a first playing time of the first
12 clip determined as approximately the first rendering time multiplied by a multiplicative
13 factor; render the first clip of the animation of the avatar; determine a subsequent
14 rendering time for each of one or more subsequent clips of the animation of the avatar,
15 each of the subsequent clips comprising one or more sequential pieces of the response
16 that succeed the preceding clip of the animation of the avatar, each subsequent
17 rendering time is determined to be approximately equivalent to the predetermined
18 acceptable rendering latency plus the total playing time of the preceding clips, each
19 subsequent playing time is determined to be approximately the rendering time of the
20 respective subsequent clip multiplied by the multiplicative factor; and render the one or
21 more subsequent clips of the animation of the avatar.

22 [0013] In a particular case, the multiplicative factor being an approximation of the
23 ratio between a playing time of a representative clip and a rendering time of that
24 respective clip.

25 [0014] In another case, the first clip and at least one of the one or more subsequent
26 clips are rendered approximately contemporaneously on separate processors.

27 [0015] In yet another case, the playing time of the first clip is reduced such that the
28 first clip ends at a natural break in speech in the response, and wherein the playing time

1 of each of the subsequent clips are reduced such that and each of the subsequent clips
2 end at a natural break in speech in the response.

3 [0016] In yet another case, the natural breaks in speech comprise a break between
4 words or at a punctuation.

5 [0017] In yet another case, the one or more processors of the system are on a
6 remote computing device that is remote to a local computing device connected to the
7 user interface, the remote computing device in communication with the local computing
8 device over a computer network.

9 [0018] In yet another case, the parsing module deconstructs the response into each
10 of the pieces, wherein each of the pieces are phonemes of the response.

11 [0019] In another aspect, there is provided a method for rendering of an animated
12 avatar on one or more computing devices using one or more animated delay clips
13 between responses of the animated avatar, the method comprising: generating an
14 avatar delay graph (ADG) by associating each of the animated delay clips with a
15 directed edge in the ADG, associating a playing length of the animated delay clip with
16 the respective edge, each edge connected to at least one other edge via a node, each
17 node associated with a point at which the animated delay clips associated with the
18 edges terminating and emanating at the node can be stitched together; selecting an
19 initial node of the ADG to be a current node; determining whether a response is being
20 processed, and while there is no response being processed: rendering one or more
21 animated delay clips using the ADP, the rendering comprising: stochastically selecting
22 one of the edges emanating from the current node; updating the current node to be the
23 node at which the selected edge is terminated; and rendering the animated delay clip
24 associated with the selected edge; and communicating the rendered one or more
25 animation delay clips to be displayed.

26 [0020] In a particular case, the rendering further comprising repeatedly:
27 stochastically selecting one of the edges emanating from the current node; updating the

1 current node to be the node at which the selected edge is terminated; and rendering the
2 animated delay clip associated with the selected edge.

3 [0021] In another case, an expressive state vector is an encoding of an expressive
4 state of the animated avatar as perceived by the user, a current expressive state being
5 a current value of the expressive state vector, each of the edges are associated with a
6 value for the expressive state vector, the method further comprising selecting an initial
7 expressive state vector as the current expressive state vector and the rendering further
8 comprising updating the current expressive state vector based on the expressive state
9 vector associated with the selected edge when such edge is selected.

10 [0022] In yet another case, the edges are selected using a probability inversely
11 proportional to a distance between the current expressive state and expressive state
12 values associated with each of the respective selectable edges.

13 [0023] In another case, a system for rendering of an animated avatar using one or
14 more animated delay clips between responses of the animated avatar, the animated
15 avatar displayed on a user interface, the system comprising one or more processors
16 and a data storage device, the one or more processors configured to execute a delay
17 module to: generate an avatar delay graph (ADP) by associating each of the animated
18 delay clips with a directed edge in the ADG, associating a playing length of the
19 animated delay clip with the respective edge, each edge connected to at least one other
20 edge via a node, each node associated with a point at which the animated delay clips
21 associated with the edges terminating and emanating at the node can be stitched
22 together; select an initial node of the ADG to be a current node; determine whether a
23 response is being processed, while there is no response being processed: render one
24 or more animated delay clips using the ADP, the rendering comprising: stochastically
25 selecting one of the edges emanating from the current node with a probability inversely
26 proportional to a distance between an expressive state vector associated with the
27 respective edge and a vector of the same rank associated with the animated delay clip;
28 updating the current node to be the node at which the selected edge is terminated; and

1 rendering the animated delay clip associated with the selected edge; and communicate
2 the rendered one or more animation delay clips to the user interface.

3 [0024] In a particular case, the one or more processors of the system are on a
4 remote computing device that is remote to a local computing device connected to the
5 user interface, the remote computing device in communication with the local computing
6 device over a computer network, and wherein at least one of the animated delay clips is
7 locally cached on the local computing device.

8 [0025] These and other aspects are contemplated and described herein. It will be
9 appreciated that the foregoing summary sets out representative aspects of systems and
10 methods to assist skilled readers in understanding the following detailed description.

11 DESCRIPTION OF THE DRAWINGS

12 [0026] A greater understanding of the embodiments will be had with reference to the
13 Figures, in which:

14 [0027] FIG. 1 is a schematic diagram of a system for rendering of an animated
15 avatar, in accordance with an embodiment;

16 [0028] FIG. 2 is a schematic diagram showing an exemplary operating environment
17 for the system of FIG. 1;

18 [0029] FIG. 3 is a flow chart of a method for rendering of an animated avatar, in
19 accordance with an embodiment;

20 [0030] FIG. 4 is a diagram of an example of an avatar delay graph (ADG); and

21 [0031] FIG. 5 is a flow chart of a method for rendering of an animated avatar, in
22 accordance with another embodiment.

23 DETAILED DESCRIPTION

24 [0032] It will be appreciated that for simplicity and clarity of illustration, where
25 considered appropriate, reference numerals may be repeated among the Figures to

1 indicate corresponding or analogous elements. In addition, numerous specific details
2 are set forth in order to provide a thorough understanding of the embodiments
3 described herein. However, it will be understood by those of ordinary skill in the art that
4 the embodiments described herein may be practised without these specific details. In
5 other instances, well-known methods, procedures and components have not been
6 described in detail so as not to obscure the embodiments described herein. Also, the
7 description is not to be considered as limiting the scope of the embodiments described
8 herein.

9 [0033] It will be appreciated that various terms used throughout the present
10 description may be read and understood as follows, unless the context indicates
11 otherwise: "or" as used throughout is inclusive, as though written "and/or"; singular
12 articles and pronouns as used throughout include their plural forms, and vice versa;
13 similarly, gendered pronouns include their counterpart pronouns so that pronouns
14 should not be understood as limiting anything described herein to use, implementation,
15 performance, etc. by a single gender. Further definitions for terms may be set out
16 herein; these may apply to prior and subsequent instances of those terms, as will be
17 understood from a reading of the present description.

18 [0034] It will be appreciated that any module, unit, component, server, computer,
19 terminal or device exemplified herein that executes instructions may include or
20 otherwise have access to computer readable media such as storage media, computer
21 storage media, or data storage devices (removable and/or non-removable) such as, for
22 example, magnetic disks, optical disks, or tape. Computer storage media may include
23 volatile and non-volatile, removable and non-removable media implemented in any
24 method or technology for storage of information, such as computer readable
25 instructions, data structures, program modules, or other data. Examples of computer
26 storage media include RAM, ROM, EEPROM, flash memory or other memory
27 technology, CD-ROM, digital versatile disks (DVD) or other optical storage, magnetic
28 cassettes, magnetic tape, magnetic disk storage or other magnetic storage devices, or
29 any other medium which can be used to store the desired information and which can be
30 accessed by an application, module, or both. Any such computer storage media may be

1 part of the device or accessible or connectable thereto. Further, unless the context
2 clearly indicates otherwise, any processor or controller set out herein may be
3 implemented as a singular processor or as a plurality of processors. The plurality of
4 processors may be arrayed or distributed, and any processing function referred to
5 herein may be carried out by one or by a plurality of processors, even though a single
6 processor may be exemplified. Any method, application or module herein described
7 may be implemented using computer readable/executable instructions that may be
8 stored or otherwise held by such computer readable media and executed by the one or
9 more processors.

10 [0035] In accordance with the foregoing, in one aspect, a system and method for
11 rendering of an animated avatar is provided.

12 [0036] While some artificially intelligent devices, such as smart speakers, interact
13 with a user via audio-only responses, this may not be ideal. Generally, humans interact
14 best when the other party is represented both auditorily and visually. In this way, visual
15 cues can be exchanged to provide a more meaningful and realistic interaction.

16 [0037] However, animating an audio signal, such as those generated as a response
17 to an artificially intelligent device, to correspond with an anthropomorphic avatar is an
18 especially challenging technical problem.

19 [0038] An exemplary approach for animating an avatar using an audio signal
20 involves decomposing the audio signal into basic components; for example, phonemes.
21 The audio signal can be decomposed using, for example, natural language processing
22 on the audio signal to generate the corresponding text, which can be parsed into
23 sequences of phonemes. For each phoneme, there is a database of one or more
24 corresponding avatar animations to execute. If these animations are sufficiently
25 synchronized with the audio signal, the avatar can appear to generally realistically talk.

26 [0039] The above approach can be augmented by encoding into the audio signal
27 being generated a collection of hints as to what the avatar should be doing; for example,

1 should it simulate being happy or sad at a certain point in the sequence. This can be
2 used to fine tune the animations that are being generated.

3 [0040] A limitation of the above approach can be that it requires substantive
4 computational resources in the computing pipeline in order to graphically render the
5 animation. Further, where the generated response is somewhat long, a user is typically
6 going to be annoyed having to wait for the full animation to be generated and rendered
7 before being able to view it. Accordingly, this can significantly affect uptake of animated
8 response technology. Even if a system starts playing part-way through rendering of the
9 full response, the user will nonetheless generally have to wait until a sufficiently long
10 sequence has been generated.

11 [0041] FIG. 2 shows an exemplary computing environment 10 of the embodiments
12 described herein. In this example, a local computing device 26 communicates with, and
13 accesses content located on, a remote computing device 32 over a network, such as
14 the internet 24. The remote computing device 32 can be a centralized server or a
15 distributed computing architecture, such as a cloud computing resource. In further
16 embodiments, embodiments of methods and systems described herein can be run on
17 the remote computing device 32 or run partially on the remote computing device 32 and
18 partially on the local computing device 26. It is understood that the remote computing
19 device 32 may be in communication with multiple local computing devices 26, and *vice*
20 *versa*.

21 [0042] FIG. 1 shows various physical and logical components of an embodiment of a
22 system 100 for rendering of an animated avatar. As shown, the system 100 has a
23 number of physical and logical components, including at least one central processing
24 unit ("CPU") 102 (comprising one or more processors), random access memory ("RAM")
25 104, a network interface 110, non-volatile storage 112, and a communication link 114
26 enabling CPU 102 to communicate with the other components. The communication link
27 114 can be, for example, a local bus, a network communication link, or the like. CPU
28 102 executes an operating system, and various modules, as described below in greater
29 detail. RAM 104 provides relatively responsive volatile storage to CPU 102. The

1 network interface 110 permits communication with other systems, such as other
2 computing devices and servers remotely located from the system 100. In some cases,
3 the network interface 110 communicates with a user interface 130 located on the local
4 computing device 32. Non-volatile storage 112 stores the operating system and
5 programs, including computer-executable instructions for implementing the operating
6 system and modules, as well as any data used by these services. Additional stored
7 data, as described below, can be stored in a database 116. During operation of the
8 system 100, the operating system, the modules, and the related data may be retrieved
9 from the non-volatile storage 112 and placed in RAM 104 to facilitate execution.

10 [0043] In an embodiment, the system 100 further includes a parsing module 120, a
11 rendering module 122, and a delay module 124. In some cases, some or all of the
12 operations and/or functions of the various modules 120, 122, 124 may be executed
13 either all on the remote computing device 32, all on the local computing device 26, or
14 partly on the remote computing device 32 and partly on the local computing device 26.

15 [0044] Advantageously, the system 100 can parallelize rendering of the avatar. The
16 parsing module 120 can deconstruct a determined response into smaller pieces. The
17 rendering module 122 can render those pieces in parallel. These rendered clips can
18 then be communicated to the user interface 130, via the network interface 110, where it
19 can be presented sequentially to the user. "Clip," as referred to herein, refers to a
20 sequence of animation frames animating the avatar.

21 [0045] If the relationship between playing time, T_p , and rendering and network
22 latency time, T_r , is approximated as a multiplicative factor (κ), so $T_p = \kappa T_r$. If there is
23 also a predetermined acceptable rendering latency (T), then a first rendering stream
24 generally has T seconds to render a first clip; resulting in a length of κT of animated
25 video.

26 [0046] In some cases, the multiplicative factor (κ) can be determined experimentally
27 and can model an efficiency for the rendering module 122. For example, if $\kappa = 1$ then the
28 rendering module 122 is able to render in real time (playing time of the animation), if

$\kappa > 1$ then it can render in greater than real time, and if $\kappa < 1$ then it is less efficient than real time. In many cases, κ also includes communication latency between the rendering module 122 and the user interface 130. The acceptable latency value T models generally a length of time a hypothetical user is willing to wait for a response. In an example, T values between 500 milliseconds and 1.5 seconds would be acceptable latency values.

[0047] In some cases, a second rendering stream can also begin rendering a second clip right away, the second clip being for a portion of the animation starting after the first clip. This second rendering stream generally has an initial latency period, plus the first clip's playing time, within which to render. Thus, the second rendering stream has $T + \kappa T$ seconds of rendering time and produces $\kappa (T + \kappa T)$ seconds of rendered animated video. In a particular case, the second rendering stream is rendered on a separate processor or computing device than the first rendering stream such that they can be rendered in parallel.

[0048] More generally for n rendering streams, and in some cases, n processors or computing devices rendering the n rendering streams:

$$T_r^n = T + \sum_{i=0}^{n-1} T_p^i.$$

Where T_r^n is the rendering time of the n 'th rendering stream and T_p^n is the playing time of the n 'th clip. Thus, the above equation indicates that the n 'th rendering component has rendering time T (the latency to start) plus the playing time of all the clips preceding the start of clip n . Under the assumption that $T_p = \kappa T_r$, then:

$$T_r^n = T + \kappa \sum_{i=0}^{n-1} T_r^i.$$

[0049] The above second equation illustrates that the above can be represented in terms of rendering time. Thus, a rendering time for a first rendering stream is T , the second rendering stream is $T + \kappa T_r$, and so on. Advantageously, this provides break

1 points in the video to be played such that each rendering task can distributed over a
2 number of processors. Further, the above equation can provide resource allocation by
3 providing a maximum number of processors that need to be allocated to the task of
4 rendering a given avatar response.

5 [0050] In some cases, it is desirable to stitch sequential clips together when playing
6 them so that arbitrary clip points can be avoided. In these cases, instead of using the
7 break points identified as above, being the playing time of each clip, the system 100 can
8 treat the theoretical break points above as maximum values and seek the next earliest
9 point in the response that corresponds to a word break, punctuation, or other natural
10 break in speech. Advantageously, the use of natural speech break points can provide
11 more natural break points in rendering of the animation. In an example, suppose there
12 is a break point T_p identified as described above. Rather than splitting the response at
13 this point, the parsing module 120 can scan backwards (towards the beginning of the
14 respective clip) searching and selecting a first break in the response; for example, either
15 a punctuation or a space between words. In this example, the time moving backwards
16 until the first word break is referred to as T_B and the time until the first punctuation is
17 referred to as T_P . Each of the times are weighted by K_B and K_P respectively. The
18 rendering module 122 selects which of $T_B K_B$, $T_P K_P$, and V_{max} has the smallest value as
19 the break point. In this case, V_{max} is a maximum weighted distance to backup. In some
20 cases, larger backup values can reduce the effectiveness of parallelism provided by the
21 system 100. Thus, a value of V_{max} may be a small number of seconds in some cases.
22 While, generally, this is not a large issue for English text as word break occurs quite
23 frequently, it may be more prevalent where there are very long words. In the case of
24 long words, it can be desirable to break the utterance in the middle of the word. Note
25 that in some cases, especially for very short duration clips, one or more of T_B and T_P
26 may not exist.

27 [0051] FIG. 3 shows an embodiment of a method 300 for rendering of an animated
28 avatar. At block 302, a determined response (also referred to as an utterance) is
29 received from a conventional response generation engine 130. The response
30 generation engine 130 can be executed on the remote computing device 32 or on

1 another computing device in communication with the remote computing device. The
2 response generation engine 130 can receive an input, such as an auditory query, from a
3 user. Utilizing a machine learning technique, the response generation engine 130 can
4 digitally parse and identify words from the input and use a machine learning technique
5 to determine a response to the input.

6 [0052] At block 304, the parsing module 120 deconstructs the determined response
7 into smaller response pieces. In most cases, the smaller cases can be phonemes. In
8 further cases, the smaller pieces can be other demarcations of language, such as each
9 piece being a particular word. In further cases, the determined response can be
10 received from the response generation engine already in the smaller pieces.

11 [0053] At block 306, the rendering module 122 renders a first clip of the avatar's
12 animation. The first clip comprises one or more sequential response pieces. The overall
13 length of playing time of the first clip is determined by the rendering module 122 as a
14 multiplicative factor multiplied by an acceptable rendering latency time. The
15 multiplicative factor being an approximation of the ratio between a playing time of a
16 representative clip and a rendering time of that respective clip. In some cases, the
17 representative clip can be an experimental clip used to determine the multiplicative
18 factor. In other cases, the representative clip can be the first clip. In some cases, the
19 multiplicative factor can be an approximation of the ratio between a playing time of a
20 representative clip and a rendering time, plus a network latency time, of that respective
21 clip. The network latency time being approximately the latency between the remote
22 computing device 32 and the local computing device 26.

23 [0054] At block 308, the rendering module 122 renders one or more subsequent
24 clips of the avatar's animation. Each of the subsequent clips being a portion of the
25 animation starting after the clip that precedes it; for example, a second clip being the
26 portion of the animation that follows the first clip, a third clip being the portion of the
27 animation that follows the second clip, and so on until, in some cases, the end of the
28 determined response is reached. Each of the subsequent clips has a rendering time that
29 is equal to or less than the totality of the playing times of the preceding clips plus a

1 predetermined acceptable rendering latency. The total playing time of each clip is equal
2 to the respective rendering time multiplied by the multiplicative factor.

3 [0055] At block 310, when each of the animation clips are rendered, each respective
4 clip is communicated to the user interface 130 via the network interface 110 to be
5 displayed by the user interface 130 to the user in sequential order received, producing a
6 full animation of the determined response.

7 [0056] In some cases, the delay module 124 can stall, or add unintended latency, to
8 the animated video being generated where desirable. In a particular case, this delay can
9 be obscured by cyclically playing the animated video back and forth a small amount in
10 order to avoid the appearance of the animated avatar being stuck or stuttering to the
11 user. Such cyclically playing (also referred to as “rolling”) of the animated video
12 backwards and forwards can be used to hide unexpected latency.

13 [0057] In some cases, between utterances, the avatar should not be still. Rather, the
14 system 100 should render animations for the avatar to engage in apparently normal
15 motion when not providing a response or engaged with the user. In some cases, the
16 system 100 should render the avatar to transit from this delay behavior to utterance
17 behaviour approximately seamlessly. The delay module 124 can accomplish this
18 behaviour by pre-rendering, and in some cases, sending to the user interface 130 and
19 caching, a plurality of idle renderings that can be played when the avatar is idle. These
20 idle renderings can be combined together by the delay module 124 to make arbitrarily
21 long sequences of idle behaviour.

22 [0058] In an embodiment, an avatar delay graph (ADG) can be used by the delay
23 module 124 to provide a formal structure to encode short idle animation sequences.
24 These idle animation sequences can be played at the user interface 130 to provide an
25 animation of the avatar between utterances. In some cases, the short idle animation
26 sequences can be locally cached on the local computing device 26. The ADG can also
27 be used to provide a mechanism within which to obscure rendering and transmission
28 latencies, which are generally unavoidable given the distributed rendering of the avatar.

[0059] The ADG is modelled as a labelled directed graph: $G = (V, E)$, where $V = \{x_1, x_2, \dots, x_n\}$ and $E = \{e_1, e_2, \dots, e_n\}$. Nodes, labelled $x_1, x_2, \dots, x_n$, correspond to points at which specific animation sequences can be stitched together smoothly. Edges, labelled $e_1, e_2, \dots, e_n$, model individual animation sequences. Each edge, for example $e = (x_a, x_b)$, is labelled with tau $\tau(e)$, where the length of time required to play or present the animation sequence, tau $\tau(e)$, corresponds to edge e . When the avatar is animated with the animation sequence corresponding to edge e , the avatar's representation within the ADG transits from one edge to another, for example x_a to x_b . In most cases, also associated with edge e is an "expressive state" $es = (s_1, s_2, \dots, s_p)$, which is an encoding of the nature of the avatar as it is perceived by a user. The expressive state for each graph can have a predetermined dimensionality to allow the graph to represent more or less complex expressive state transitions; the dimensionality of es can be avatar dependent.

[0060] Initially, animation of the avatar is in some node x and has some avatar state S . When the avatar is not animated providing a response or uttering an expression, the animation of the avatar notionally traverses the ADG in a stochastic manner, as described below. When in node x , one of the edges departing from x is selected. For each candidate edge e_i , the delay module 124 determines a distance from S to $es(e_i)$, represented as $d_i = |S - es(e_i)|$. The delay module 124 then selects randomly from each of the incident edges with a probability inversely proportional to this distance. Specifically, with a probability proportional to $1 / (d_i + \epsilon)$. Once an edge e_{best} is selected, the avatar's state S is updated using $S' = \lambda S + (1 - \lambda)es(e_{best})$, where e_{best} is the outgoing edge chosen. Generally, ϵ is selected to be a relatively small number in order to avoid the computation $1 / d_i$ becoming infinite when d_i is zero. In an example, ϵ can be approximately 0.001. Generally, λ is a number between 0 and 1 that represents how much the avatar's expressive state is changed when traversing an edge. In most cases, λ is a predetermined value. For example, if $\lambda=0$, then the avatar's expressive state becomes that of the edge that is traversed, $es(e_{best})$. If $\lambda=1$, then the avatar's expressive state is unchanged even though the selected edge, e_{best} , is traversed. In an example, λ can be approximately 0.9.

[0061] An example of an ADG and its operation are illustrated in FIG. 4. In this example, the graph has two nodes $V = \{x_1, x_2\}$, with multiple edges connecting x_1 and x_2 to themselves and transitions between x_1 and x_2 . In this example, the dimensionality of es is 1, so the values are (1), (0.5), (-1), (-0.5); with the one dimension of es representing 'happiness' running from -1 (sad) to +1 (happy). In further examples, each dimension of es can represent a different expressive state; for example, es might have a dimensionality of 2, so $es=(a,b)$, where the a dimension can be happiness and the b dimension can be engagement.

[0062] In the example of FIG. 4, suppose the avatar animation is at x_1 with an expressive state $S = 1$. There are three possible transitions that can follow from x_1 : edge A which leads back to x_1 , edge B that leads back to x_1 and edge D that leads to x_2 . Thus, the next animation sequence to be played will be one of A, B, and D. The delay module 124 determines a distance from its current state S to each of these three edges, A, B, and D, $d_A = 0$, $d_B = 0$, and $d_D = 0.5$ respectively. The delay module 124 stochastically selects either of A, B or D based on relative probabilities using the above distances, $P_A = 1/\epsilon$, $P_B = 1/\epsilon$, and $P_D = 1/(0.5+\epsilon)$ respectively. In an example, suppose ϵ is 0.5, then the probability proportionality values are 2, 2, and 1; which normalize to $P_A = 2/5$, $P_B = 2/5$, $P_D = 1/5$. Suppose that B is chosen. Then the B animation sequence is displayed (in this case for a duration of 3 seconds), S is updated as $S' = \lambda S = (1 - \lambda)es(B)$ and the above steps can be repeated.

[0063] In some cases, vertices in the ADG can be labelled as being a starting or a terminating node to aid in merging ADG transitions and renderings with renderings associated with responses. A node can be both an initial and terminating node. When response is to be generated, an appropriate starting and terminating node is also identified from the nodes labelled as being initial or terminating respectively.

[0064] In the present embodiments, advantageously, the system 100 renders the avatar always doing something; which it does by traversing the ADG stochastically. When the user interacts with the avatar soliciting a response, the system 100 must transition from its stochastic background appearance to one that represents interaction

1 with the user. In most cases, the response should be presented as 'fitting in' with what
2 the avatar is currently doing. In some cases, the system 100 can do this by having the
3 delay module 124 identify a node in the ADG that can be used to branch out of the ADG
4 into the utterance and then another node in the ADG to where it will return after the
5 utterance is complete. Nodes that might be used as start points for this are generally
6 called 'initial' nodes. Similarly, nodes that can be used to re-enter the ADG once the
7 utterance is complete are called 'terminating' nodes. In some cases, all nodes can be
8 predetermined to be initial and terminating nodes, or some subset of the nodes can be
9 predetermined to be a initial node, a terminating node, or both.

10 [0065] In some cases, the delay module 124 can be executed on the local computing
11 device 26, or some functions of the delay module 124 can be executed on the local
12 computing device 26 and some on the remote computing device 32. In some cases, the
13 avatar delay graph (ADG) approach described herein can be made more sophisticated
14 by caching only portions of the graph on the local computing device 26 and the updating
15 them as the state of the avatar changes. When the avatar is to render some response, a
16 new temporary edge $E = (start, end)$ can be constructed. Here the *start* and *end* nodes
17 can be selected from the set of initial and terminating nodes in the ADG. The end node
18 is chosen such that it has a terminating label and a mean of $|es(end, x_k) - S|$ is minimized.
19 Thus, when the response is generated, it can terminate in a state where there is a good
20 exiting edge in the ADG.

21 [0066] The choice of start node is similar; however, it is also necessary to identify a
22 node that can be accessed quickly in terms of transitions in the ADT in order to avoid
23 the introduction of abrupt changes in the avatar's appearance. The start node is chosen
24 such that it has an initial label and the cost of $\sum \alpha \tau(e) + (1-\alpha) |es(e) - S|$ is minimized.
25 Where α is a parameter than can be used to tune between the desirability of quickly
26 moving from the ADG to begin uttering the response ($\alpha = 1$) and making the transition as
27 smooth as possible ($\alpha = 0$). Where the sum is over a path in the ADG from the avatar's
28 current state to the *start* node. In essence, this selects a nearby start node such that the
29 *es* values are similar to the current state of the avatar *S*. Note that selecting the start

1 node also enables the determination of the expected delay before it is necessary to start
2 rendering the response.

3 [0067] Once the start and end nodes have been identified, the delay module 124
4 begins to move deterministically through the ADT to the start node following the
5 sequence identified in the process of identifying this node. When the delay module 124
6 reaches the start node it then signifies to the rest of the system 100 to execute the
7 rendered utterance. The delay module 124 can then re-enter the ADG at the end node.
8 Generally, the value of S can remain unchanged, although it would be possible to
9 associate a change in S with each utterance. Once at the end node, the delay module
10 124 continues its stochastic traverse through the ADG until the next response is
11 available and the above is repeated.

12 [0068] FIG. 5 shows another embodiment of a method 500 for rendering of an
13 animated avatar using one or more delay clips between utterances of the animated
14 avatar. At block 501, the delay module 124 generates the avatar delay graph (ADG) by
15 associating each of the animated delay clips with an edge in the ADG and determining a
16 playing length of the animated delay clip with the respective edge. Each edge is
17 connected to at least one other edge via a node, each node being at a point at which
18 the animated delay clips associated with the edges terminating and emanating at the
19 node can be stitched together. In some cases, each node is connected to each other
20 node via an edge. In some cases, each node also has an edge that emanates from it
21 and terminates at itself.

22 [0069] At block 502, the delay module 124 selects an initial node as a current node
23 of the ADG and communicates the associated clip to the user interface 130. In some
24 cases, the initial node can be predetermined or selected stochastically among the
25 available nodes.

26 [0070] At block 503, the delay module 124 determines whether a response is being
27 processed, where a response is being processed if a response has been received from
28 the response generation engine 130 or a response is currently being rendered by the
29 rendering module 122. At block 504, while the above is negative, the delay module 124

1 renders one or more delay animation clips using an avatar delay graph (ADG). At block
2 506, when each of the delay animation clips are rendered, each respective clip is
3 communicated to the user interface 130 via the network interface 110 to be displayed by
4 the user interface 130 to the user in sequential order received.

5 [0071] As part of block 504, at block 514, the delay module 124 stochastically
6 selects one of the edges emanating from the current node. At block 516, the delay
7 module 124 updates the current node to be the node at which the selected edge is
8 terminated. The delay module 124 communicates the clip associated with the selected
9 edge to the user interface 130 to be played after the previous clip communicated to the
10 user interface 130.

11 [0072] The delay module repeats blocks 514 and 516 while the condition at block
12 502 is negative.

13 [0073] The embodiments described herein advantageously provide a more realistic
14 and interactive mechanism for human-robot interaction. The embodiments can thus be
15 deployed in a range of different applications; for example, service roles where humans
16 seek information from a greeter, help desk or receptionist. In one exemplary application,
17 a greeter in a service-oriented company can be provide 24-7 by the animated avatar of
18 the embodiments described herein. the animated avatar of the embodiments described
19 herein can advantageously provide visually accurate, realistic, and consistent
20 interaction with users. In some cases, the embodiments described herein can be
21 deployed in either a fixed installation (for example, an information kiosk) or as part of an
22 autonomous robot.

23 [0074] Although the foregoing has been described with reference to certain specific
24 embodiments, various modifications thereto will be apparent to those skilled in the art
25 without departing from the spirit and scope of the invention as outlined in the appended
26 claims.

CLAIMS

1. A method for rendering of an animated avatar with a response on one or more computing devices, the method comprising:
 - receiving the response, the response comprising a plurality of pieces;
 - determining a first rendering time of a first clip of an animation of the avatar as approximately equivalent to a predetermined acceptable rendering latency, the first clip comprising one or more sequential pieces of the response, a first playing time of the first clip determined as approximately the first rendering time multiplied by a multiplicative factor;
 - rendering the first clip of the animation of the avatar;
 - determining a subsequent rendering time for each of one or more subsequent clips of the animation of the avatar, each of the subsequent clips comprising one or more sequential pieces of the response that succeed the pieces of a clip of the animation of the avatar that directly precedes the corresponding subsequent clip or the pieces of the first clip of the animation of the avatar where the first clip directly precedes the corresponding subsequent clip, each subsequent rendering time is determined to be approximately equivalent to the predetermined acceptable rendering latency plus the total playing time of all the clips that precede the corresponding subsequent clip, each subsequent playing time is determined to be approximately the subsequent rendering time of the respective subsequent clip multiplied by the multiplicative factor; and
 - rendering the one or more subsequent clips of the animation of the avatar.
2. The method of claim 1, wherein the multiplicative factor is an approximation of a ratio between a playing time of a representative clip and a rendering time of the representative clip.
3. The method of claim 1, wherein the first clip and at least one of the one or more subsequent clips are rendered approximately contemporaneously.

Replacement Sheet

4. The method of claim 1, wherein the playing time of the first clip is reduced such that the first clip ends at a natural break in speech in the response, and wherein the playing time of each of the subsequent clips are reduced such that each of the subsequent clips each end at other natural breaks in speech in the response.
5. The method of claim 4, wherein the natural breaks in speech comprise a break between words or at a punctuation.
6. The method of claim 5, wherein the natural breaks in speech comprise a break that is closest in time to a total respective rendering time of the respective first clip or subsequent clip.
7. The method of claim 1, wherein each of the pieces are phonemes of the response.
8. A system for rendering of an animated avatar displayable on a user interface with a response received from a response generation engine, the system comprising one or more processors and a data storage device, the one or more processors configured to execute:

a parsing module to receive the response, the response comprising a plurality of pieces; and

a rendering module to:

determine a first rendering time of a first clip of an animation of the avatar as approximately equivalent to a predetermined acceptable rendering latency, the first clip comprising one or more sequential pieces of the response, a first playing time of the first clip determined as approximately the first rendering time multiplied by a multiplicative factor;

render the first clip of the animation of the avatar;

determine a subsequent rendering time for each of one or more subsequent clips of the animation of the avatar, each of the subsequent clips comprising one or more sequential pieces of the response that succeed the pieces of a clip of the animation of the avatar that directly precedes the corresponding subsequent clip or the pieces of the first clip of the animation of the avatar where the first clip directly precedes the

Replacement Sheet

corresponding subsequent clip, each subsequent rendering time is determined to be approximately equivalent to the predetermined acceptable rendering latency plus the total playing time of all the clips that precede the corresponding subsequent clip, each subsequent playing time is determined to be approximately the subsequent rendering time of the respective subsequent clip multiplied by the multiplicative factor; and render the one or more subsequent clips of the animation of the avatar.

9. The system of claim 8, wherein the multiplicative factor is an approximation of the ratio between a playing time of a representative clip and a rendering time of the representative clip.
10. The system of claim 9, wherein the one or more processors comprises a plurality of processors and wherein the first clip and at least one of the one or more subsequent clips are rendered approximately contemporaneously on separate processors of the plurality of processors.
11. The system of claim 9, wherein the playing time of the first clip is reduced such that the first clip ends at a natural break in speech in the response, and wherein the playing time of each of the subsequent clips are reduced such that each of the subsequent clips each end at other natural breaks in speech in the response.
12. The system of claim 11, wherein the natural breaks in speech comprise a break between words or at a punctuation.
13. The system of claim 9, wherein the one or more processors of the system are on a remote computing device that is remote to a local computing device connected to the user interface, the remote computing device in communication with the local computing device over a computer network.
14. The system of claim 9, wherein the parsing module deconstructs the response into each of the pieces, wherein each of the pieces are phonemes of the response.

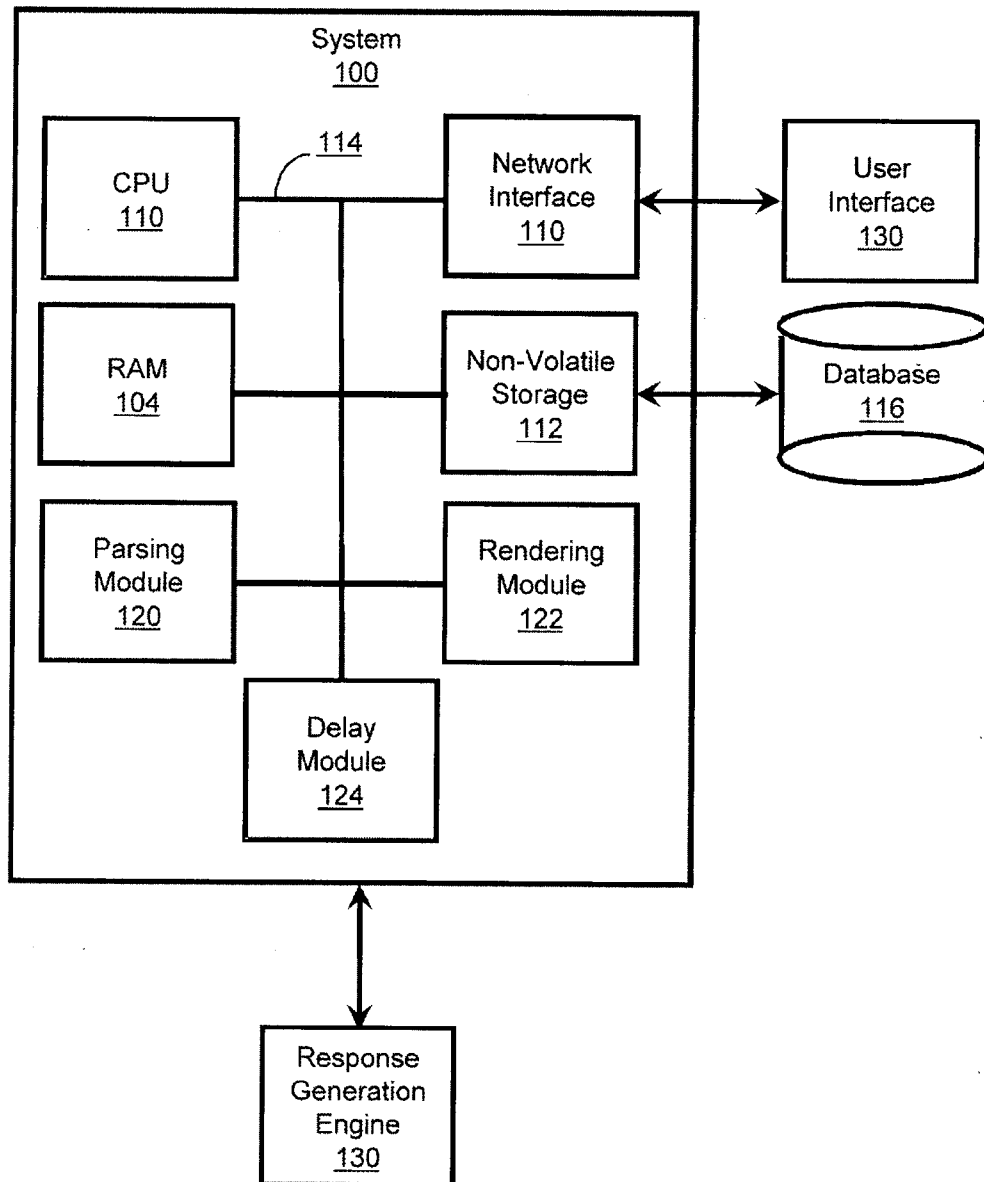


FIG. 1

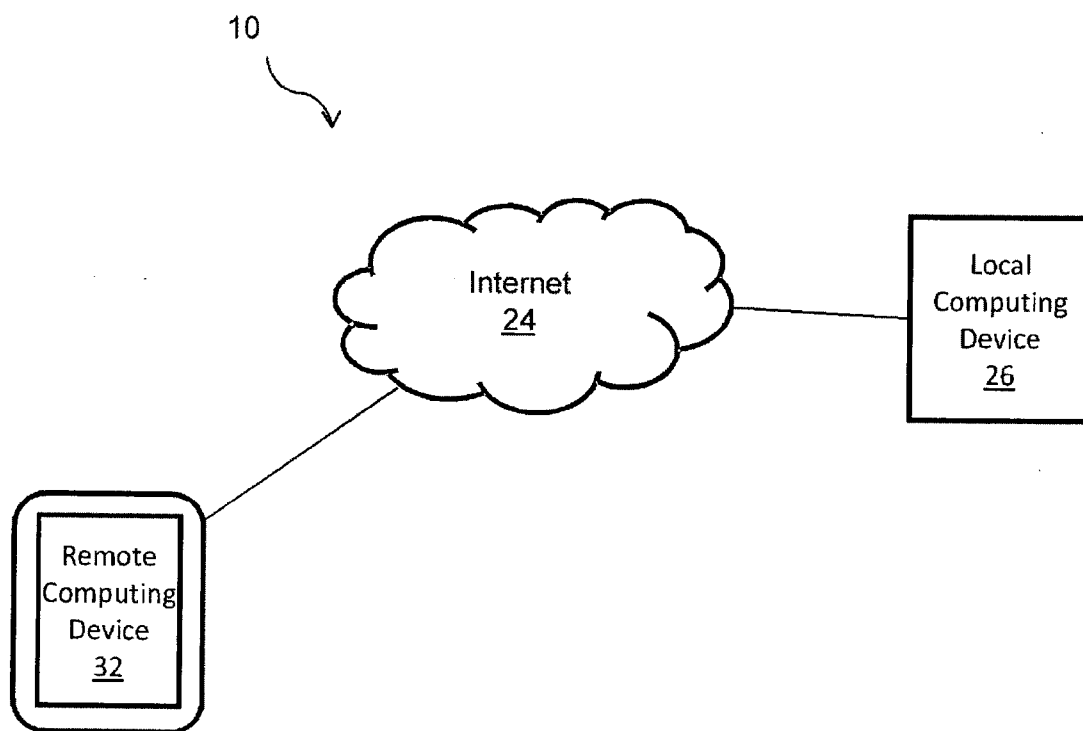


FIG. 2

300
↘

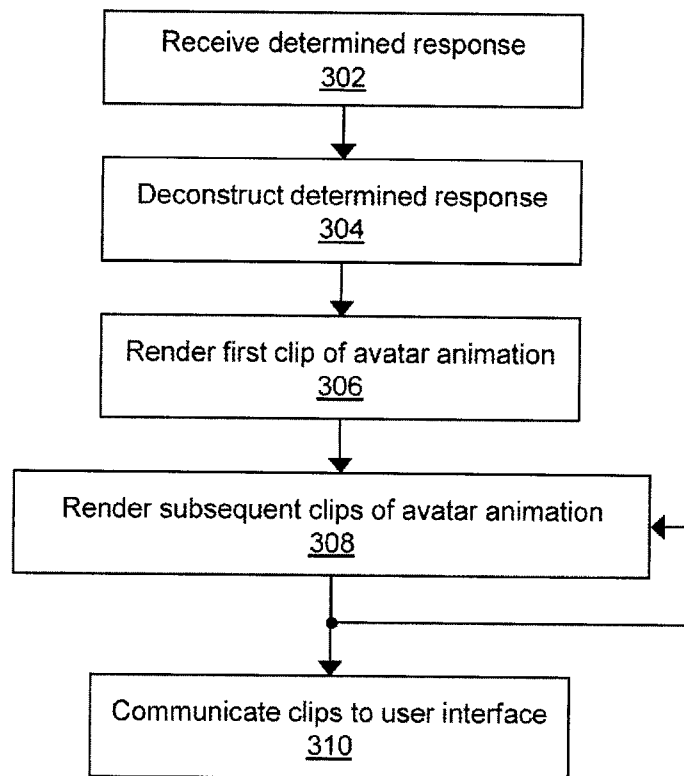


FIG. 3

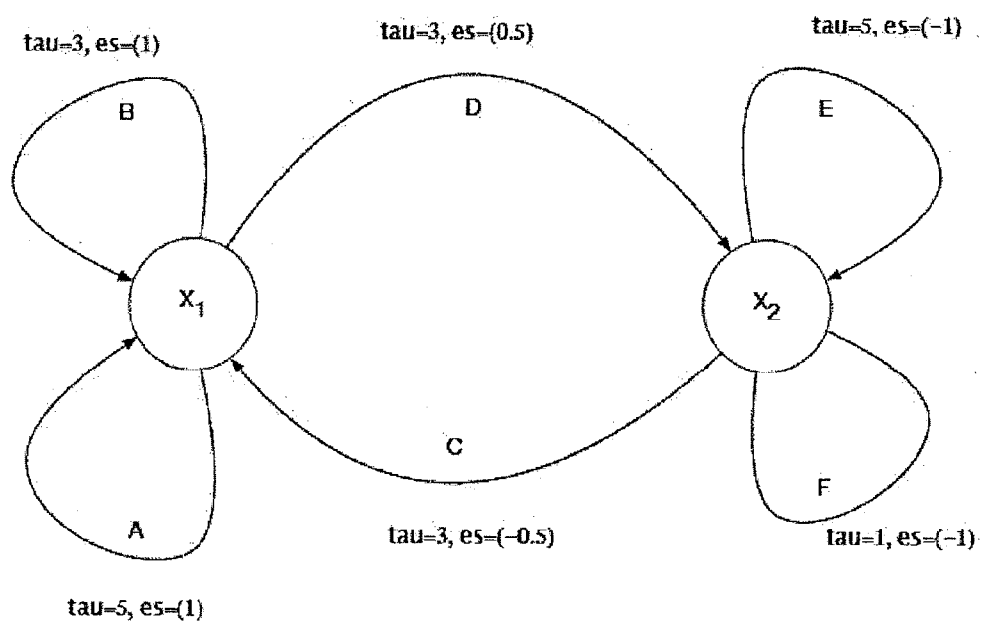


FIG. 4

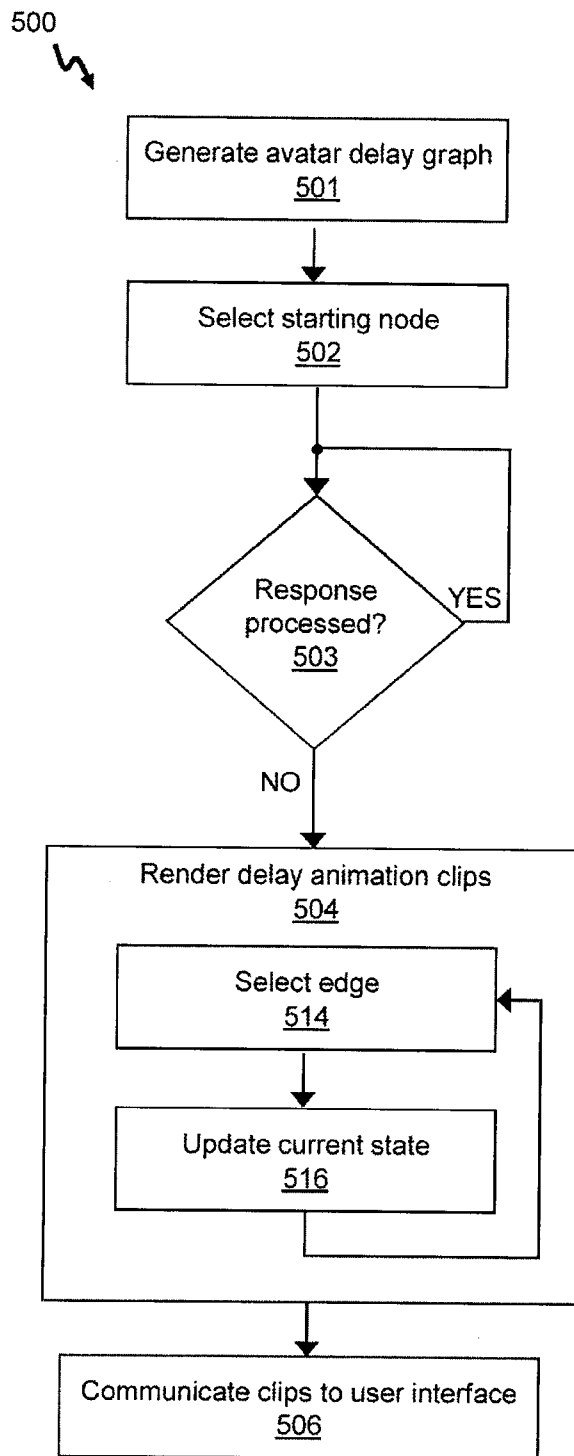


FIG. 5

