

(19)



SUOMI - FINLAND  
(FI)

PATENTTI- JA REKISTERIHALLITUS  
PATENT- OCH REGISTERSTYRELSEN  
FINNISH PATENT AND REGISTRATION OFFICE

(10) **FI 130697 B1**  
(12) **PATENTTIJULKAISU**  
**PATENTSKRIFT**  
**PATENT SPECIFICATION**

|      |                                                                                                                                                                                                                                                                                                                                    |            |
|------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------|
| (45) | Patentti myönnetty - Patent beviljats - Patent granted                                                                                                                                                                                                                                                                             | 24.01.2024 |
| (51) | Kansainvälinen patenttiluokitus - Internationell patentklassificering - International patent classification<br><b>H04L 41/0663</b> ( 2022 . 01 )<br><b>H04L 41/0823</b> ( 2022 . 01 )<br><b>H04L 41/149</b> ( 2022 . 01 )<br><b>H04W 16/20</b> ( 2009 . 01 )<br><b>H04W 24/02</b> ( 2009 . 01 )<br><b>G06N 3/092</b> ( 2023 . 01 ) |            |
| (21) | Patenttihakemus - Patentansökan - Patent application                                                                                                                                                                                                                                                                               | 20225692   |
| (22) | Tekemispäivä - Ingivningsdag - Filing date                                                                                                                                                                                                                                                                                         | 28.07.2022 |
| (23) | Saapumispäivä - Ankomstdag - Reception date                                                                                                                                                                                                                                                                                        | 28.07.2022 |
| (41) | Tullut julkiseksi - Blivit offentlig - Available to the public                                                                                                                                                                                                                                                                     | 24.01.2024 |

(73) Haltija - Innehavare - Holder  
**1• Elisa Oyj**, Ratavartijankatu 5, 00520 Helsinki, (FI)

(72) Keksijä - Uppfinnare - Inventor  
**1• Nikitin, Alexander**, HELSINKI, (FI)

(74) Asiamies - Ombud - Agent  
**Moosedog Oy**, Vähäheikkiläntie 56 C, 20810 Turku

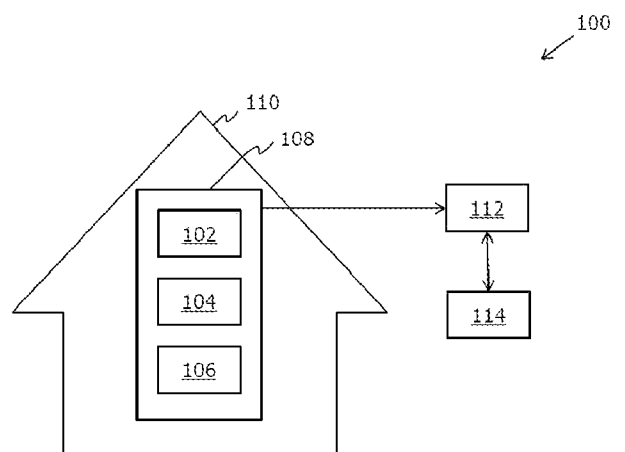
(54) Keksinnön nimitys - Uppfinningens benämning - Title of the invention  
**JÄRJESTELMÄ JA MENETELMÄ TILAAJ JÄRJESTETYN TIETOLIIKENNEVERKON LAITTEIDEN HALLINNOIMISEKSI**  
**SYSTEM OCH METOD FÖR HANTERING ENHETER I KOMMUNIKATIONSNÄTVERK DET IMPLEMENTERAS I LOKALEN**  
**SYSTEM AND METHOD FOR MANAGING DEVICES IN COMMUNICATION NETWORK THAT IS IMPLEMENTED AT PREMISES**

(56) Viitejulkaisut - Anförda publikationer - References cited  
US 2022066456 A1, US 2021409976 A1, US 2021192768 A1, US 2018317095 A1, TALEB, S. M. et al. Nodes placement in wireless mesh networks using optimization approaches: a survey. In: Neural Computing and Applications. Springer [online], 2022-02-17, Vol. 34, No. 7, pages 5283–5319, KARUNARATNE, S. et al. An Overview of Machine Learning Approaches in Wireless Mesh Networks. In: IEEE Communications Magazine. IEEE [online], 2019-04-30, Vol. 57, No. 4, pages 102-108.

(57) Tiivistelmä - Sammandrag - Abstract

Keksinnön kohteena on järjestelmä (100), laitteiden (102, 104, 106) viestintäverkossa (108) tapahtuvaa hallinnointia varten, joka viestintäverkko on toteutettu kiinteistössä (110). Järjestelmä sisältää vähintään yhden palvelimen (112), joka on määritetty hankkimaan kiinteistösuunnitteludataa ja verkkodataa palvelusta (228), joka liittyy kiinteistöön, jossa kiinteistösuunnitteludata liittyy vähintään yhteen kiinteistön pohjapiirrokseen ja verkkodata liittyy laitteisiin viestintäverkossa ja viestintäverkon suorituskykyyn; käsittelemään kiinteistösuunnitteludataa ensimmäisen prosessointiputken (202) kautta tuottaakseen ensimmäisen tuloksen; käsittelemään verkkodataa toisen prosessointiputken (204) kautta tuottaakseen toisen tuloksen; ja käsittelemään ensimmäisen ja toisen tuloksen vahvistusoppimiseen perustuvalla prosessointisilmukalla (224) määrittämään vähintään yhden toiminnan, joka tulee toteuttaa viestintäverkon laitteiden hallitsemiseksi.

Disclosed is a system (100) for managing devices (102, 104, 106) in a communication network (108) that is implemented at a premises (110). The system comprises at least one server (112) configured to obtain premises planning data and network data from a service (228) associated with the premises, wherein the premises planning data pertains to at least one floor plan of the premises and the network data pertains to the devices in the communication network and a performance of the communication network; process the premises planning data via a first processing pipeline (202) to generate a first output; process the network data via a second processing pipeline (204) to generate a second output; and process the first output and the second output using a reinforcement learning-based processing loop (224) to determine at least one action that is to be implemented for managing the devices in the communication network.



## SYSTEM AND METHOD FOR MANAGING DEVICES IN COMMUNICATION NETWORK THAT IS IMPLEMENTED AT PREMISES

### TECHNICAL FIELD

5 The present disclosure relates to systems for managing devices in a communication network that is implemented at a premises. The present disclosure relates to methods for managing devices in a communication network that is implemented at a premises. The present disclosure relates to computer program products for managing devices in a communication  
10 network that is implemented at a premises.

### BACKGROUND

Recently, mesh networks are being extensively used to connect devices to each other and to cloud-based services. Mesh networks are becoming increasingly popular to set up a new network that spans more than single  
15 router. The mesh networks use mesh devices with wireless processors that contain software, thereby enabling the wireless processors to configure themselves automatically, with negligible or null central coordination. Applications of mesh network may be in, but are not limited to, military organisations, smart cities, industrial automation, emergency  
20 services, agriculture, and the like. In a first example, the military organisations may use the mesh network to avoid breakdowns in communication. In a second example, smart cities may use mesh network to monitor traffic flow, sewage treatment, and control street lighting.

Currently, the mesh networks have easy scalability due to the presence  
25 of the mesh devices, that act as nodes. Hence, a size of the mesh network can be easily and quickly changed based on the availability and number of the mesh devices in a property (for example, such as a building). However, this increases a workload on each mesh device, as the mesh

devices are configured to work like a router, along with sending messages to the other mesh devices. This leads to unwanted complications due to the increase in the workload, thereby leading to a mismanagement of the mesh devices. Beneficially, the mesh network reduces redundancy by configuring a given mesh device to both receive and translate information received from other mesh devices. Hence, even if another mesh device stops working due to an unforeseen circumstance, such as, obstruction by walls or object present within the property, the mesh network will use the given mesh device, to continue functioning in the mesh network. However, the given mesh device may experience latency due to a decrease in a number of mesh devices in the mesh network. Advantageously, the mesh networks can easily increase range by connecting the nodes (i.e., the given mesh device) to gateways, wherein the gateways transfer the messages to the other mesh devices in the mesh network. However, this increases power consumption for each node, as the node is configured to function both as an endpoint a route. Hence, the given mesh device will need to draw more power than required to operate efficiently, which impacts an electricity bill of the property.

In patent application US 2022066456 A1 (EBRAHIMI AFROUZI ALI [US] et al. "*Obstacle recognition method for autonomous robots*") there is described an obstacle recognition method for autonomous robotic devices such as a drone, a robotic vacuum cleaner, a robotic lawn mower, a robotic mop. To operate autonomously or with minimal (or less than fully manual) input and/or external control within an environment, methods such as mapping, localization, object recognition, and path planning methods, among others, are used so that robotic devices may autonomously create a map of the environment, subsequently use the map for navigation, and devise intelligent path and task plans for efficient navigation and task completion. Additionally, neural networks may be trained to recognize different types of objects within an environment and the Wi-Fi router the robot connects to may be used by the processor as

an indication of where the robot is. In consumer homes and commercial establishments, wireless routers may be replaced by a mesh of wireless/Wi-Fi repeaters/routers.

Therefore, in light of aforementioned discussion, there exists a need to  
 5 overcome the aforementioned drawbacks associated with management of mesh devices in a property.

## SUMMARY

The present disclosure seeks to provide a system for managing devices in a communication network that is implemented at a premises. The  
 10 present disclosure also seeks to provide a method for managing devices in a communication network that is implemented at a premises. The present disclosure also seeks to provide a computer program product for managing devices in a communication network that is implemented at a premises. An aim of the present disclosure is to provide a solution that  
 15 overcomes at least partially the problems encountered in prior art.

In a first aspect, the present disclosure provides a system for managing devices in a communication network that is implemented at a premises, the system comprising at least one server configured to:

- obtain premises planning data and network data from a service  
 20 associated with the premises, wherein the premises planning data pertains to at least one floor plan of the premises and the network data pertains to the devices in the communication network and a performance of the communication network, wherein the devices in the communication network are communicably coupled in a mesh network topology;
- 25 - process the premises planning data via a first processing pipeline to generate a first output;
- process the network data via a second processing pipeline to generate a second output; and

- process the first output and the second output using a reinforcement learning-based processing loop to determine at least one action that is to be implemented for managing the devices in the communication network;
- 5 wherein when processing the network data via the second processing pipeline (204), the at least one server (112) is configured to:
  - extract, from the network data, metadata pertaining to the devices (102, 104, 106) in the communication network (108) and the performance of the communication network;
  - 10 - convert the network data, using a mesh converter (216) into an internal representation of the devices (102, 104, 106) in the communication network (108) and the performance of the communication network, wherein a format of the internal representation is different from at least one format of the network data; and
  - 15 - transform the internal representation of the devices (102, 104, 106) in the communication network (108) and the performance of the communication network, using a mesh embedder (218), into a second vector representation, wherein the second vector representation and the metadata pertaining to the devices in the communication network
  - 20 collectively constitute the second output of the second processing pipeline (204).

In a second aspect, the present disclosure provides a method for managing devices in a communication network that is implemented at a premises, the method comprising:

- 25 - obtaining premises planning data and network data from a service associated with the premises, wherein the premises planning data pertains to at least one floor plan of the premises and the network data pertains to the devices in the communication network and a performance of the communication network, wherein the devices in the communication
- 30 network are communicably coupled in a mesh network topology;

- processing the premises planning data via a first processing pipeline to generate a first output;
  - processing the network data via a second processing pipeline to generate a second output; and
  - 5 - processing the first output and the second output using reinforcement learning-based processing loop to determine at least one action that is to be implemented for managing the devices in the communication network;
- wherein when the processing the network data via the second processing
- 10 pipeline (204) comprises:
- extracting, from the network data, metadata pertaining to the devices (102, 104, 106) in the communication network (108) and the performance of the communication network;
  - converting the network data, using a mesh converter (216) into an
  - 15 internal representation of the devices (102, 104, 106) in the communication network (108) and the performance of the communication network, wherein a format of the internal representation is different from at least one format of the network data; and
  - transforming the internal representation of the devices (102, 104,
  - 20 106) in the communication network (108) and the performance of the communication network, using a mesh embedder (218), into a second vector representation, wherein the second vector representation and the metadata pertaining to the devices in the communication network collectively constitute the second output of the second processing pipeline
  - 25 (204).

In a third aspect, the present disclosure provides a computer-program product for managing devices in a communication network that is implemented at a premises, wherein the computer program product comprises a non-transitory machine-readable data storage medium

30 having stored thereon program instructions that, when accessed by a

processing device, cause the processing device to execute steps of a method of the second aspect.

Embodiments of the present disclosure substantially eliminate or at least partially address the aforementioned problems in the prior art, and  
5 enable efficient management of devices in the communication network in the premises.

Additional aspects, advantages, features and objects of the present disclosure would be made apparent from the drawings and the detailed description of the illustrative embodiments construed in conjunction with  
10 the appended claims that follow.

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.

#### 15 BRIEF DESCRIPTION OF THE DRAWINGS

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  
20 drawings. However, the present disclosure is not limited to specific methods and instrumentalities disclosed herein. Moreover, those skilled in the art will understand that the drawings are not to scale. Wherever possible, like elements have been indicated by identical numbers.

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

FIG. 1 is a block diagram of a system for managing devices in a communication network, in accordance with an embodiment of the present disclosure;

FIG. 2 is an exemplary block diagram for managing devices in a communication network that is implemented at a premises, in accordance with an embodiment of the present disclosure; and

FIG. 3 is an illustration of a flowchart depicting steps of a method for managing devices in a communication network, in accordance with an embodiment of the present disclosure.

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.

#### DETAILED DESCRIPTION OF EMBODIMENTS

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 practising the present disclosure are also possible.

In a first aspect, the present disclosure provides a system for managing devices in a communication network that is implemented at a premises, the system comprising at least one server configured to:

- obtain premises planning data and network data from a service associated with the premises, wherein the premises planning data pertains to at least one floor plan of the premises and the network data pertains to the devices in the communication network and a performance

of the communication network, wherein the devices in the communication network are communicably coupled in a mesh network topology;

- process the premises planning data via a first processing pipeline to generate a first output;
  - 5 - process the network data via a second processing pipeline to generate a second output; and
  - process the first output and the second output using a reinforcement learning-based processing loop to determine at least one action that is to be implemented for managing the devices in the
  - 10 communication network;
- wherein when processing the network data via the second processing pipeline (204), the at least one server (112) is configured to:
- extract, from the network data, metadata pertaining to the devices (102, 104, 106) in the communication network (108) and the
  - 15 performance of the communication network;
  - convert the network data, using a mesh converter (216) into an internal representation of the devices (102, 104, 106) in the communication network (108) and the performance of the communication network, wherein a format of the internal representation
  - 20 is different from at least one format of the network data; and
  - transform the internal representation of the devices (102, 104, 106) in the communication network (108) and the performance of the communication network, using a mesh embedder (218), into a second vector representation, wherein the second vector representation and the
  - 25 metadata pertaining to the devices in the communication network collectively constitute the second output of the second processing pipeline (204).

In a second aspect, the present disclosure provides a method for managing devices in a communication network that is implemented at a

30 premises, the method comprising:

- obtaining premises planning data and network data from a service associated with the premises, wherein the premises planning data pertains to at least one floor plan of the premises and the network data pertains to the devices in the communication network and a performance of the communication network, wherein the devices in the communication network are communicably coupled in a mesh network topology;
  - processing the premises planning data via a first processing pipeline to generate a first output;
  - processing the network data via a second processing pipeline to generate a second output; and
  - processing the first output and the second output using reinforcement learning-based processing loop to determine at least one action that is to be implemented for managing the devices in the communication network;
- wherein when the processing the network data via the second processing pipeline (204) comprises:
- extracting, from the network data, metadata pertaining to the devices (102, 104, 106) in the communication network (108) and the performance of the communication network;
  - converting the network data, using a mesh converter (216) into an internal representation of the devices (102, 104, 106) in the communication network (108) and the performance of the communication network, wherein a format of the internal representation is different from at least one format of the network data; and
  - transforming the internal representation of the devices (102, 104, 106) in the communication network (108) and the performance of the communication network, using a mesh embedder (218), into a second vector representation, wherein the second vector representation and the metadata pertaining to the devices in the communication network collectively constitute the second output of the second processing pipeline (204).

In a third aspect, the present disclosure provides a computer-program product for managing devices in a communication network that is implemented at a premises, wherein the computer program product comprises a non-transitory machine-readable data storage medium  
5 having stored thereon program instructions that, when accessed by a processing device, cause the processing device to execute steps of a method of the second aspect.

The present disclosure provides the aforementioned system, method and computer program product to manage the devices in the premises. The  
10 communication network ensures continuous connectivity of the devices within the premises, by constantly updating itself regarding floorplan of the premises. The system constantly updates itself to optimally tune the devices in a required manner, based on the floor plan of the premises and a quality of signal from different devices connected in the  
15 communication network. As a result of which, existing devices in the communication network is adjusted in a better manner to ensure continuous connectivity and an effective coverage area. Furthermore, to maintain the continuous connectivity and an effective coverage area, adding new devices in the communication network in the premises, will  
20 not disrupt existing connectivity in the communication network. Additionally, alternatively, the devices may be tuned automatically according to the update to function in a required manner in the premises.

Throughout the present disclosure, the term "*communication network*" refers to an arrangement of interconnected programmable and/or non-  
25 programmable components that are configured to facilitate data communication between a plurality of devices of the system. Furthermore, the communication network may include, but is not limited to, one or more peer-to-peer network, a hybrid peer-to-peer network, local area networks (LANs), radio access networks (RANs), metropolitan  
30 area networks (MANs), wide area networks (WANs), all or a portion of a

public network such as the global computer network known as the Internet, a private network, a cellular network and any other communication system or systems at one or more locations. The devices in the communication network are communicably/operably coupled in a mesh network topology, wherein a given device may be connected directly or indirectly to other devices within the premises. Examples of the devices may include, but are not limited to, routers, switches, modems, bridges. As a first example, there may be three devices A, B, C and D in the communication network that is implemented at the premises. The device A may be directly connected to the device B, but indirectly connected to the device D through the device C. Beneficially, the communication network can create multiple routes of communication for information to travel among the devices. Hence, resilience of the communication network increases in case there is a device or a communication failure, as an alternate route may be available for communication. Continuing with the first example, there may further be a device E in the communication network. The device E may be directly connected to device B, but may also be indirectly connected through the device C to the device D. A communication failure may occur in the direct connection between the devices B and E. Consequently, communication between the devices B and E is carried out using an alternate route. Herein, the alternate route is the indirect connection of the devices. Additionally, the communication network may include wired or wireless communication that can be carried out via any number of known protocols, including, but not limited to, Web Real-Time Communication (WebRTC) protocols, Internet Protocol (IP), Wireless Access Protocol (WAP), Frame Relay, or Asynchronous Transfer Mode (ATM).

Throughout the present disclosure, the term "*premises*" refers to a physical location encompassing a given area. Examples of the premises may include, but are not limited to, a house (for example, such as a

bungalow, villa, an apartment, and the like), a housing complex (for example, a multi-storey structure having multiple apartments), a hospital, a commercial establishment (for example, such as a shop, a shopping mall, an office premises, and the like), an educational institution  
5 (for example, such as a school, an university), an industrial establishment (for example, such as a factory, a warehouse, and the like).

Throughout the present disclosure, the term "*server*" refers to a structure and/or module that includes programmable and/or non-programmable components configured to store, process and/or share information.  
10 Specifically, the at least one server includes any arrangement of physical or virtual computational entities capable of enhancing information to perform various computational tasks. Furthermore, it should be appreciated that the at least one server may be a single hardware server and/or plurality of hardware servers operating in a parallel or distributed  
15 architecture. In an example, the at least one server may include components such as memory, at least one processor, a network adapter and the like, to store, process and/or share information with other computing components, such as a user device and/or a user equipment. Optionally, the server is implemented as a computer program that  
20 provides various services (such as database service) to other devices, modules or apparatus.

The service associated with the premises acts as a storage of a condition of the communication network and the premises. For example, installation of a new device in the communication network may trigger a  
25 change in the service, and a condition of the communication network corresponding to the change is stored in the service. Similarly, a change in positioning of furniture within the premises may trigger a change in the service, and a condition of the premises corresponding to the change is duly stored in the service. The premises planning data and the network  
30 data are information provided to the system for processing. The premises

planning data comprises information pertaining to the at least one floor plan of the premises. The at least one floor plan is a technical drawing drawn to scale, showing relationships between rooms, spaces, and other physical features of one level of the premises. Examples of the physical features may include, but are not limited to, walls, floors, doors, objects, furniture. It will be appreciated, that the at least one floor plan includes locations of the devices present within the premises (for example, location of a router).

Optionally, the information pertaining to the devices may be indicative of at least one of : manufacturer details of the device, an operation mode of the device, a range of the device, a network technology of the device, a network type of the device, a network protocol of the device, antenna features of the device, physical attributes of the device, security details of the device, physical dimensions of the device, values of at least one parameter associated with the devices, a functionality of the device, a count of a number of devices in the communication network, voltage rating of the device, maximum operating temperature of the device. The at least one parameter associated with the devices may be at least one of: a bandwidth, a transmission rate, a frequency band, a memory space.

Optionally, the performance of the communication network comprises information indicative of at least one of: a throughput, a latency, a backhaul performance, received signal strength indication (RSSI), a roaming time, a coverage area, a scalability, reliability.

The term "*processing pipeline*" refers to a pipeline having a set of interconnected data elements and at least one instruction that are connected in series, which simultaneously process data within the pipeline. It will be appreciated that in most processing pipelines, the data elements are executed in a parallel or time-sliced fashion. The processing pipeline may be at least one of: an instruction pipeline, a graphic pipeline, a software pipeline. In the instruction pipeline the at least one instruction

may be executed by overlapping phases of the processing pipeline. In the graphic pipeline, a three-dimensional (3D) image is transformed into a two-dimensional (2D) image. Herein, the phases of the processing pipeline are at least one of: fetching, decoding and executing, a given instruction. A technical effect is that the processing pipeline increases throughput of the system. For example, the processing pipeline may fetch a given instruction from a memory while previous instructions are being executed simultaneously. The processing pipeline may be divided into segments of a certain duration. The duration may be equal or unequal.

10 In an instance, the processing pipeline is divided into the first processing pipeline and the second processing pipeline.

Optionally, the first processing pipeline processes the premises planning data. The premises planning data is received as input to the first processing pipeline. The first processing pipeline is divided into stages that are responsible for processing information from the at least one floor plan to generate the first output. The first output is a vector representation of the premises planning data. The vector representation of the premises planning data comprises points, lines, and curves that are based on mathematical formulae. Examples of common file format for vector representation may include, but are not limited to, Adobe Illustrator (.ai), Encapsulated PostScript (.eps), Portable Document Format (.pdf), Scalable Vector Graphics (.svg). Moreover, the vector representation is easily understood by the reinforcement learning-based, which assists in accurate results.

25 Optionally, the premises planning data is in form of at least one image of the at least one floor plan of the premises, wherein when processing the premises planning data via the first processing pipeline, the at least one server is configured to:

- extract, from the at least one image, metadata pertaining to the premises;

- split the at least one image, using a premises plan splitter, into a plurality of images representing floor plans of a plurality of units of interest in the premises;
- convert the plurality of images, using a premises plan converter,  
5 into an internal representation of the premises, wherein a format of the internal representation is different from an image file format of the plurality of images; and
- transform the internal representation of the premises, using a premises plan embedder, into a first vector representation, wherein the  
10 first vector representation and the metadata pertaining to the premises collectively constitute the first output of the first processing pipeline.

In this regard, the metadata comprises at least one of: a year of a building in the premises, a number of floors of the building, a total area of the building, a total area of the premises. Such metadata may be  
15 represented as at least one: a textual representation, a graphical representation, a tabular representation. The premises plan splitter splits the at least one image into the plurality of images representing floor plans of the plurality of units of interest in the premises. Optionally, the units of interest comprise at least one of: rooms, floors, position of floors, in  
20 the building. For example, the premises planning data may be an image representing a floor plan of a building in the premises. The premises plan splitter may split the at least one image into a plurality of images, wherein each image may represent a room of the floor. A technical advantage of using the premises plan splitter is that splitting the plurality of images  
25 into the plurality of units of interest makes it easier for the at least one server to decipher information accurately.

Optionally, the term "*premises plan converter*" refers to a digital converter that receives the plurality of images from the premises plan splitter and converts the plurality of images into an internal representation, such that  
30 a format of the internal representation is a computer-understandable

format. Optionally, the format of the internal representation is a JavaScript Object Notation (JSON) format.

The term "*premises plan embedder*" refers to an embedder that digitally embeds the internal representation of the premises as vectors into the first vector representation. Optionally, the premises plan embedder is implemented using at least one of: an autoencoder architecture, a self-supervised learning architecture. For example, some elements in the internal representation may be masked, and the premises plan embedder may be trained to predict the masked internal representation, and such predicted masked internal representation is thereon sent to the reinforcement learning-based processing loop. Optionally, the premises plan embedder is a floor plan embedder. A technical advantage of processing the premises planning data via the first processing pipeline in the above-mentioned manner is that the first output obtained in this manner is easily understood by the reinforcement learning-based processing loop to determine the at least one action.

Optionally,

- the premises plan splitter is implemented using a first algorithm, the first algorithm being at least one of: a deep neural network algorithm, a computer vision algorithm;
- the premises plan converter is implemented using a plurality of second algorithms, the plurality of second algorithms being implemented as a semantic segmentation algorithm and a classification algorithm; and
- the premises plan embedder is implemented using a third algorithm, the third algorithm being at least one of: convolutional neural network algorithm, self-supervised learning algorithm.

In this regard, the at least one server employs the first algorithm, the plurality of second algorithms and the third algorithm to collectively determine the first output. Additionally, the at least one server employs

a given algorithm to compute the first processing pipeline. The given algorithms enable the at least one server to become more accurate in predicting outcomes and/or performing tasks, without being explicitly programmed. Specifically, the given algorithms are employed to  
 5 artificially train the at least one server so as to enable them to automatically learn and improving performance from experience, without being explicitly programmed.

Optionally, the first algorithm is a deep neural network algorithm, wherein the deep neural network algorithm is used for end-to-end (E2E)  
 10 learning. The E2E learning refers to training a complex learning system (i.e., the premises plan splitter) represented by a deep neural network model that represents a target system, bypassing intermediate layers usually present in a given processing pipeline. The deep neural network algorithm may be used when a given floor plan of the premises  
 15 represented in the at least one image is simple. Examples of the deep neural network algorithm may include, but is not limited to, U-Net algorithm, SegFormer<sup>®</sup> algorithm. The first algorithm can be a computer vision algorithm, wherein the computer vision algorithm mimics working of a human brain. The computer vision algorithm can be used for at least  
 20 one of: classification, identification, tracking, of the plurality of units of interest in the premises. The computer vision algorithm may be used instead of the deep neural network algorithm when a given floor plan of the premises represented in the at least one image is simple. Examples of the computer vision algorithm may include, but are not limited to, Line  
 25 segment detector, Ramer-Douglas-Peucker algorithm.

Optionally, the plurality of second algorithms is a semantic segmentation algorithm. The semantic segmentation algorithm assigns a label to every pixel in the plurality of images. It will be appreciated that each pixel represents a given unit of interest that can be classified into a class and  
 30 the given unit of interest is identifiable unambiguously. Suitably, the

semantic segmentation algorithm treats multiple units of interest of the same class as a single entity. For example, a given image may represent a floor plan of room A, room B, floor 1 and floor 2. The given image may be composed of 1000 pixels. Subsequently, upon implementing the  
 5 semantic segmentation algorithm, 200 pixels of the given image may be attributed to room A, 200 pixels may be attributed to room B, 300 pixels may be attributed to floor 1, and 300 pixels may be attributed to floor 2, thereby labelling the plurality of units of interest in the given image.

Optionally, the plurality of second algorithms is the classification  
 10 algorithm. Herein, the plurality of images is sent to the classification algorithm, wherein the classification algorithm predicts a class or a category of the plurality of units of interest represented in the plurality of images. The class of the plurality of units of interest refers to a category of a given unit of interest sharing similar characteristics. The class may  
 15 be pre-known or may be defined by the classification algorithm.

Optionally, the third algorithm is the convolutional neural network algorithm. The convolutional neural network algorithm is a type of artificial neural network used in image recognition. The convolutional neural network algorithm comprises at least one convolutional layer. The  
 20 convolutional neural network algorithm is used for at least one of: processing, classifying, segmenting of the internal representation of the premises. The third algorithm can be the self-supervised learning algorithm. The self-supervised learning algorithm is a machine learning algorithm, wherein the premises plan embedder trains itself to learn part  
 25 of an input (i.e., the internal representation) is programmed to refine its own performance. A technical advantage of using the given algorithm is that it makes the system agile and robust.

The second processing pipeline processes the network data. The network data is received as input to the second processing pipeline. The second  
 30 processing pipeline is divided into stages that are responsible for

processing information pertaining to the devices in the communication network and the performance of the communication network to generate the second output. The network sub-pipeline is a segmented portion of a network pipeline, which allows the network data to travel a shorter distance to achieve the second output. It will be appreciated that the network sub-pipeline makes the network more efficient and the system agile. Similar to the first output, the second output is a vector representation of the network data. The vector representation of the network data comprises individual points that may be stored as coordinate pairs to indicate a physical location of the devices in the premises. Examples of common file format for vector representation of the second output is similar to the examples of the common file format for vector representation of the first output.

When processing the network data via the second processing pipeline, the at least one server is configured to:

- extract, from the network data, metadata pertaining to the devices in the communication network and the performance of the communication network;
- convert the network data, using a mesh converter into an internal representation of the devices in the communication network and the performance of the communication network, wherein a format of the internal representation is different from at least one format of the network data; and
- transform the internal representation of the devices in the communication network and the performance of the communication network, using a mesh embedder, into a second vector representation, wherein the second vector representation and the metadata pertaining to the devices in the communication network collectively constitute the second output of the second processing pipeline.

Optionally, the metadata comprises at least one of: a signal strength, the physical location, network quality, Internet Protocol (IP) address, Media Access Control (MAC) address, Universal Unique Identifier (UUID), of the device. The mesh converter converts the network data into the internal representation of the devices using a simple data processing pipeline (i.e., the second processing pipeline). The mesh converter changes the format of the internal representation into a format understandable by the mesh embedder. Optionally, the format of the internal representation is a JSON format. The mesh embedder transforms the internal representation into the second vector representation. The second vector representation is a high-dimensional embedding from the internal representation of the devices in the communication network and the performance of the communication network. Herein, the high-dimensional embedding is a high-dimensional translation of a high-dimensional vector representation of the internal representation. A technical advantage of processing the network data via the second processing pipeline in the above-mentioned manner is that the second output obtained in this manner is easily understood by the reinforcement learning-based processing loop to determine the at least one action.

Optionally, the mesh embedder is implemented using at least one of: a transformer model, an autoencoder model. The transformer model is a neural network model that learns context (and thereby, meaning) by tracking relationships between the devices in the network data. The autoencoder model is a neural network model that seeks to learn a compressed representation of the internal representation. A technical advantage of using such algorithms for the mesh embedder is that it makes the system agile and robust. Indeed a technical advance is that presented methods can produce embeddings and allow scaling along with a growth in the collected data.

Optionally, the at least one server is further configured to train a given model using an optimization algorithm, the optimization algorithm being at least one of: a gradient descent optimization algorithm, an Adam optimization algorithm, a stochastic gradient descent algorithm. The term

5 "*optimisation algorithm*" refers to a training algorithm which trains the given model to optimise its results. The gradient descent algorithm is an algorithm for training the given model by minimizing errors between an actual result and an expected result. The Adam optimization algorithm is a replacement optimization algorithm for stochastic gradient descent

10 which can handle sparse gradients on noisy problems. The stochastic gradient descent algorithm is an optimization algorithm used to train the given model to find the parameters that correspond to a best fit between a predicted output and an actual output. A technical advantage of using the above-mentioned algorithms for training the given model is that these

15 algorithms are computationally fast and require less memory space, hence keeping the system agile. Another technical benefit of training the given model using such algorithms is that they train the given model to optimise its results, making them more accurate. Optionally, the given model is trained optimization algorithms, known in the art.

20 Optionally, the at least one server is further configured to train the premises plan splitter, the premises plan converter, the premises plan embedder, the mesh converter, and the mesh embedder with linear warmup and cosine annealing. Herein, linear warmup is a learning rate schedule comprising a learning rate of the given model. The learning rate

25 is increased linearly from a low rate to a constant rate, so as to reduce volatility in early stages of training the given model. Moreover, cosine annealing is an aggressive learning rate schedule. Herein, the learning rate of the given model starts high, and drops relatively rapidly to a minimum value near zero before being rapidly increased again to a

30 maximum. Optionally, the at least one server is further configured to train the given model using a backpropagation algorithm. Herein, the given

model back-propagates errors from output of the premises plan embedder and the mesh embedder to input of the premises plan splitter and the mesh converter, respectively. Optionally, the given model may be trained on another server, and thereafter be transferred to the server  
 5 for inference.

The term "*reinforcement learning-based processing loop*" refers to an infinite loop which recommends actions based on a given situation. The reinforcement learning-based processing loop may either recommend such actions on a request, or periodically. Notably, the reinforcement  
 10 learning-based processing loop evaluates a quality of an ongoing situation, chooses an action from learned state-transition dynamics and thereafter recommends an action to an entity. Herein, the entity may be a human or a robot. Moreover, the entity thereon performs the action, triggering a change in the communication network at the premises, which  
 15 thereafter triggers a further iteration. Beneficially, such performing of the action by the entity assist in managing the devices in the communication network. It will be appreciated that the first output and the second output are processed analogously.

Optionally, the at least one server is further configured to:

- 20 - obtain at least one other data from at least one other data source;
  - process the at least one other data is via at least one other processing pipeline to generate at least one other output; and
  - process the at least one other output using reinforcement learning based processing loop to determine the at least one action.
- 25 Optionally, the reinforcement learning-based processing loop is implemented using a reinforcement learning algorithm, the reinforcement learning algorithm being one of: a Q-learning algorithm, a Deep Q-learning algorithm, a Deep Deterministic Policy Gradient algorithm, a Q-learning algorithm with Normalized Advantage Functions. The term

"*reinforcement learning algorithm*" refers to a machine learning-based training algorithm which supervises the training by rewarding desired results and/or punishing undesirable results. The Q-learning algorithm is a reinforcement learning algorithm which assists the given model to learn a value of an action in a given state. The Deep Q-learning algorithm is a deep-learning based reinforcement learning algorithm which assists the given model to map given input states into pairs. The Deep Deterministic Policy Gradient (DDPG) algorithm is a model-free off-policy reinforcement learning algorithm which assists the given model to learn continuous actions. A technical benefit of using the reinforcement learning algorithms as mentioned above is that such reinforcement learning algorithms maximize performance of the system.

Optionally, the at least one action comprises at least one of: tune a given device using a given setting, re-position a given device, add a given device, remove a given device. Herein, tuning the given device using the given setting refers to adjusting a frequency of the given device. A technical advantage of the at least one action comprising the above-mentioned is that accurate identification of issues and performing of such action allows the device to function appropriately, such that the communication network exhibits superior performance.

Optionally, when processing the first output and the second output using the reinforcement learning-based processing loop, the at least one server is configured to:

- evaluate a current quality of the communication network, based on the premises planning data and the network data;
- determine whether the current quality of the communication network is less than a predefined quality; and
- when the current quality of the communication network is less than the predefined quality, use a function which maps a set of actions that are feasible corresponding to a current state of the communication

network with an expected quality of the communication network which would be achieved upon implementation of the set of actions, wherein the actions in the set are identified as the at least one action when the expected quality of the communication network is equal to or greater  
 5 than the predefined quality.

In this regard, the at least one action is recommended accurately by the reinforcement learning-based processing loop. Optionally, the machine learning algorithm is chosen based on the quality of the communication network. The current quality refers to a quality of the communication  
 10 network, in real time. Optionally, a quality of the communication network is measured as at least one of: a connection speed, an upload speed, a download speed, a network latency. The predefined quality refers to a quality of the communication network that is required for a required performance of the devices in the premises. As a first example, if the  
 15 quality is measured as a network latency, the current quality may be 150 milliseconds, and the predefined quality may be 100 milliseconds. Herein, the current quality of the communication network is less than the predefined quality, since network latency is a negative value. Herein, negative value of the network latency means a reduction in response time  
 20 in prediction of the at least one action by the reinforcement learning-based processing loop. The function estimates the quality of the communication network, by mapping the current quality with the predefined quality. For example, a function may be " $Q = S * A * R$ ", wherein S is a set of states having positions of the devices and a structure of the  
 25 premises, A is a set of actions, and R is a measurement of the quality of the communication network. The quality of the communication network may be expressed in terms of a corresponding action. In an instance, the measurement of the quality of the communication network may be speed of the communication network. A corresponding set of actions may  
 30 comprise at least one of: change a position of the device, add a new device, adjust frequency of the device. The expected quality refers to a

quality of the communication network that is required for the required performance. Optionally, the expected quality is calculated using at least one of: a Bellman equation, a deep Q-learning approach. In furtherance to the first example, the expected quality may be 30 milliseconds. A technical advantage of processing the first output and the second output in this manner is that it maximises performance of the communication network (and the devices therein), since the set of actions result in an improved quality of the communication network. Optionally, a machine learning algorithm is chosen based on the quality of the communication network.

Optionally, the at least one server is further configured to provide a communication indicating the at least one action to be implemented in a form of at least one: a notification, a web interface, a speech interface, to a user device communicably coupled to the at least one server. The term "*user device*" refers to a communication device which is accessible by a user. Optionally, the user is the entity who performs the action for managing devices in the communication network. Examples of the user device include, but are not limited to, a laptop, a computer, a tablet, a phablet, a pager, a smartphone, a smartwatch, a smart device. For example, the communication may be an email notification or an alert by a virtual assistant. A technical advantage of providing the communication is that it allows an entity to perform the at least one action, resulting in accurate and useful performance of the devices. The communication may be provided to the user device in real time. The communication may include at least one of: location, network quality, requirement of the device, in the premises. Upon receiving such communication, the at least one action may be initiated.

Optionally, the at least one server is further configured to:

- perform a simulation and/or A/B testing of the communication network to obtain a reference output; and

- select the reinforcement learning algorithm based on the reference output.

As an example, A/B testing can be used as one selection criteria to find best/suitable approach. For example, autoencoders or transformers for making the embeddings can be used. All the users can be split into two sets; for the first set of users, autoencoders could be used, and for the second set of users– transformers could be used. Then, metrics over these sets can be calculated (e.g., number of complaints, quality of service, etc.) and then statistically compare which one is better. A/B testing in this example would help to evaluate all the components of the system. Results of the A/B system can be used when selecting which reinforcement learning algorithm is selected.

In this regard, a technical advantage of performing such a simulation is testing the communication network, for selecting an appropriate reinforcement learning algorithm. Herein, the simulation test is used to predict performance of the communication network, and to generate the reference output. Similarly, the A/B testing is a method of comparing two versions of the communication network against each other to determine a version out of the two versions that performs better with respect to the other version, thereby generating the reference output. Beneficially, the simulation and/or the A/B testing are used to perform evaluation of the communication network. Alternatively, optionally, the at least one server selects the reinforcement learning algorithm based on a survey.

Optionally, the at least one server is further configured to detect at least one of: a change in the premises, a change in the at least one floor plan of the premises, a change in the devices in the communication network, a change in performance of the devices in the communication network, wherein the premises planning data and the network data are obtained upon such detection. When at least one of such occurrences is detected, there occurs a change in the service associated with the premises either

automatically or manually. The premises planning data and the network data are updated accordingly. In such case, such updated premises planning data and network data are processed again to identify the at least one action to be performed. Subsequently, updating the premises  
5 planning data and the network data is repeated over time to ensure accurate performance of the devices when the at least one action is performed. In this manner, the communication network is tuned such that each device performs accurately. A technical advantage of such detection is that it allows errors to be identified timely, such that the  
10 communication network is tuned accurately. Moreover, an architecture is chosen using cross-validation or assessing performance of the at least one action.

Optionally, the at least one server comprises at least one graphical processing unit (GPU), wherein the at least one GPU is configured to  
15 implement at least one processing task. The term "*graphical processing unit*" refers to a processing unit which assists in creating and rendering computer graphics. Herein, the at least one GPU is configured to implement at least one of: processing the premises planning data, processing the network data, processing the output data. A technical  
20 advantage of utilizing the at least one GPU is that it allows for faster processing, improving performance of the system. Alternatively, optionally, the at least one server does not comprise at least one graphical processing unit (GPU). Herein, the at least one processing task is implemented by at least one central processing unit (CPU). Optionally,  
25 the at least one server comprises a tensor processing unit (TPU), wherein modelling processing is implemented on the TPU.

Optionally, the at least one server comprises a plurality of servers communicably coupled to each other, the plurality of servers being located in a same geographical region and/or in proximal geographical  
30 regions. It will be appreciated that the given model requires the at least

one GPU for quick processing, so the plurality of servers are implemented, wherein each server comprises the at least one GPU. A technical advantage of having the plurality of servers is robustness of the system along with quick processing, such that none or a negligible latency is  
5 observed. It will be appreciated that it makes the system agile. Moreover, communication between the plurality of servers is resource-intensive and therefore requiring the plurality of servers to be located in proximity.

Optionally, the system further comprises at least one balancer communicably coupled to the plurality of servers to enable load balancing  
10 across the plurality of servers. The term "*balancer*" refers to a device that distributes processing traffic across a number of servers. Examples of the balancer include, but are not limited to, an Application Load Balancer, a Network Load Balancer, a Classic Load Balancer, a Gateway Load Balancer. A technical advantage of using the at least one balancer is that  
15 it distributes the traffic equally amongst the plurality of servers; and prevents any single server from getting overloaded and possibly breaking down.

Optionally, the system further comprises a data repository communicably coupled to the at least one server, wherein the at least one server is  
20 configured to store, at the data repository, at least one of: the premises planning data, the network data, the first output, the second output, the reference output, information pertaining to the at least one action. A technical advantage of the data repository is that it allows the at least one server to store lesser data, making it lighter and thereby facilitating  
25 faster processing of the at least one server.

The present disclosure also relates to the method as described above. Various embodiments and variants disclosed above, with respect to the aforementioned method, apply *mutatis mutandis* to the method.

The present disclosure also relates to the computer program product as described above. Various embodiments and variants disclosed above, with respect to the aforementioned method, apply *mutatis mutandis* to the computer program product.

- 5 The computer program product for managing devices in a communication network that is implemented at a premises, the computer program product comprising a non-transitory machine-readable data storage medium having stored thereon program instructions that, when accessed by a processing device, cause the processing device to execute steps of
- 10 the aforementioned method. The term "*computer program product*" refers to a software product comprising program instructions that are recorded on the non-transitory machine-readable data storage medium, wherein the software product is executable upon a computing hardware for implementing the aforementioned steps of the method for managing
- 15 devices in a communication network that is implemented at a premises.

In an embodiment, the non-transitory machine-readable data storage medium can direct a machine (such as computer, other programmable data processing apparatus, or other devices) to function in a particular manner, such that the program instructions stored in the non-transitory

20 machine-readable data storage medium cause a series of steps to implement the function specified in a flowchart corresponding to the instructions. Examples of the non-transitory machine-readable data storage medium includes, but is not limited to, an electronic storage device, a magnetic storage device, an optical storage device, an

25 electromagnetic storage device, a semiconductor storage device, a portable computer diskette, a hard disk, a random access memory (RAM), a read-only memory (ROM), an erasable programmable read-only memory (EPROM or Flash memory), a static random access memory (SRAM), a portable compact disc read-only memory (CD-ROM), a digital

30 versatile disk (DVD), a memory stick, a floppy disk, a mechanically

encoded device such as punch-cards or raised structures in a groove having instructions recorded thereon, or any suitable combination thereof.

#### EXPERIMENTAL PART

- 5 As mentioned above, the format of the internal representation is a JSON format. Exemplary embodiments of the JSON format are as follows:

Example 1:

```

{
  "plan": {
10     "walls": [
        {
            "dimensions": [10, 2, 0.5], #l., h., w.
            "x_start": 0.2,
            "y_start": 0.4,
15     "x_end": 0.9,
            "y_end": 1.0,
        },
        ....
    ]
20     "furniture": [
        {
            "type": "bed",
            "position": [0.1, 0.2],
            "dimensions": [2, 0.5, 1.5],
25     },
        ....
    ]
  },
  "metadata": {

```

```

        "n_floors": 5,
        "house_build_date": 2014,
        ....
    },
5   }

```

Example 2:

```

{
    "devices": [
10     {
        "type": "mesh",
        "model": "TP link",
        "version": "Deco M4",
        "frequency": "3.2, #GHz
15     "speed": 120, #Mbps
        "position_x ": 0.1,
        "position_y ": 0.2,
        ....
    },
20     ....
    ]
    "metadata": {
        "n_devices": 5,
        "house_build_date": 2014,
25     ....
    },
}

```

## DETAILED DESCRIPTION OF THE DRAWINGS

Referring to FIG. 1, there is shown a block diagram of a system **100** for managing devices (depicted as devices **102**, **104**, **106**) in a communication network **108**, in accordance with an embodiment of the present disclosure. The system **100** is implemented at a premises **110**.  
 5 The system **100** comprises at least one server (depicted as a server **112**). The system **100** further comprises a data repository **114** communicably coupled to the at least one server **112**.

Referring to FIG. 2, there is shown an exemplary block diagram **200** for managing devices in a communication network that is implemented at a premises, in accordance with an embodiment of the present disclosure. The exemplary block diagram **200** may provide premises planning data and network data for processing to a first processing pipeline **202** and a second processing pipeline **204**. In the first processing pipeline **202**,  
 10 firstly, a premises plan splitter **206** may be used to split at least one image into a plurality of images. Secondly, a premises plan converter **208** may be used to convert the plurality of images into an internal representation of the premises. Thirdly, a premises plan embedder **210** may be used to transform the internal representation of the premises into  
 15 a first vector representation (as shown by an arrow **212**). Finally, metadata is extracted (as shown by an arrow **214**) from the at least one image. The first vector representation (as shown by the arrow **212**) and the metadata constitute a first output of the first processing pipeline **202**.

In the second processing pipeline **204**, a mesh converter **216** may be used to convert the network data into an internal representation of the devices in the communication network and the performance of the communication network. Subsequently, a mesh embedder **218** may be used to transform the internal representation of the devices in the communication network and the performance of the communication  
 25 network into a second vector representation (as shown by an arrow **220**).

Finally, metadata pertaining to the devices in the communication network and the performance of the communication network is extracted (as shown by an arrow **222**). The second vector representation (as shown by the arrow **220**) and the metadata constitute a second output of the second processing pipeline **204**. The first output and the second output are processed by a reinforcement learning-based processing loop **224**. The reinforcement learning-based processing loop **224** may be an exemplary cyclic process flow. At step **A**, a current quality of the communication network may be evaluated, based on a given input. At step **B**, an action is chosen based on the current quality of the communication network. At step **C**, the action is recommended to a user device **226**. Simultaneously, output of step **C** may further be an input to step **A**. The user device **226** performs the action as a service **228** associated with the premises.

FIG. 2 is merely an example, which should not unduly limit the scope of the claims herein. A person skilled in the art will recognize many variations, alternatives, and modifications of embodiments of the present disclosure.

Referring to FIG. 3, there is shown an illustration of a flowchart depicting steps of a method for managing devices in a communication network, in accordance with an embodiment of the present disclosure. At step **302**, premises planning data and network data are obtained from a service associated with the premises. The premises planning data pertains to at least one floor plan of the premises and the network data pertains to the devices in the communication network and a performance of the communication network. At step **304**, the premises planning data is processed via a first processing pipeline to generate a first output. At step **306**, the network data is processed via a second processing pipeline to generate a second output. At step **308**, the first output and the second output are processed using reinforcement learning-based processing loop

to determine at least one action that is to be implemented for managing the devices in the communication network.

The aforementioned steps are only illustrative and other alternatives can also be provided where one or more steps are added, one or more steps  
5 are removed, or one or more steps are provided in a different sequence without departing from the scope of the claims herein.

Modifications to embodiments of the present disclosure described in the foregoing are possible without departing from the scope of the present disclosure as defined by the accompanying claims. Expressions such as  
10 "including", "comprising", "incorporating", "have", "is" used to describe and claim the present disclosure 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.

## CLAIMS

1. A system (100) for managing devices (102, 104, 106) in a communication network (108) that is implemented at a premises (110), the system comprising at least one server (112) configured to:

- 5 - obtain premises planning data and network data from a service (228) associated with the premises, wherein the premises planning data pertains to at least one floor plan of the premises and the network data pertains to the devices in the communication network and a performance of the communication network, wherein the devices in the communication network
  - 10 are communicably/operably coupled in a mesh network topology;
    - process the premises planning data via a first processing pipeline (202) to generate a first output;
    - process the network data via a second processing pipeline (204) to generate a second output; and
  - 15 - process the first output and the second output using a reinforcement learning-based processing loop (224) to determine at least one action that is to be implemented for managing the devices in the communication network;
- wherein when processing the network data via the second processing
- 20 pipeline (204), the at least one server (112) is configured to:
- extract, from the network data, metadata pertaining to the devices (102, 104, 106) in the communication network (108) and the performance of the communication network;
  - convert the network data, using a mesh converter (216) into an
  - 25 internal representation of the devices (102, 104, 106) in the communication network (108) and the performance of the communication network, wherein a format of the internal representation is different from at least one format of the network data; and

- transform the internal representation of the devices (102, 104, 106) in the communication network (108) and the performance of the communication network, using a mesh embedder (218), into a second vector representation, wherein the second vector representation and the metadata pertaining to the devices in the communication network collectively constitute the second output of the second processing pipeline (204).

2. The system (100) according to claim 1, wherein the at least one server (112) is further configured to provide a communication indicating the at least one action to be implemented in a form of at least one: a notification, a web interface, a speech interface, to a user device (226) communicably coupled to the at least one server.

3. The system (100) according to claim 1 or 2, wherein the at least one action comprises at least one of: tune a given device using a given setting, re-position a given device, add a given device, remove a given device.

4. The system (100) according to any of the preceding claims, wherein the premises planning data is in form of at least one image of the at least one floor plan of the premises (110), wherein when processing the premises planning data via the first processing pipeline (202), the at least one server (112) is configured to:

- extract, from the at least one image, metadata pertaining to the premises;
- split the at least one image, using a premises plan splitter (206), into a plurality of images representing floor plans of a plurality of units of interest in the premises;
- convert the plurality of images, using a premises plan converter (208), into an internal representation of the premises, wherein a format of the

internal representation is different from an image file format of the plurality of images; and

- transform the internal representation of the premises, using premises plan embedder (210), into a first vector representation, wherein the first vector representation and the metadata pertaining to the premises collectively constitute the first output of the first processing pipeline (202).

5. The system (100) according to any of the preceding claims, wherein:

- the premises plan splitter (206) is implemented using a first algorithm, the first algorithm being at least one of: a deep neural network algorithm, a computer vision algorithm;
- the premises plan converter (208) is implemented using a plurality of second algorithms, the plurality of second algorithms being implemented as a semantic segmentation algorithm and a classification algorithm; and
- the premises plan embedder (210) is implemented using a third algorithm, the third algorithm being at least one of: convolutional neural network algorithm, self-supervised learning algorithm.

6. The system (100) according to any of the preceding claims, wherein the mesh embedder (218) is implemented using at least one of: a transformer model, an autoencoder model.

7. The system (100) according to any of claims 3-6, wherein the at least one server (112) is further configured to train a given model using an optimization algorithm, the optimization algorithm being at least one of: a gradient descent optimization algorithm, an Adam optimization algorithm, a stochastic gradient descent algorithm.

8. The system (100) according to any of the preceding claims, wherein when processing the first output and the second output using the

reinforcement learning-based processing loop (224), the at least one server (112) is configured to:

- evaluate a current quality of the communication network (108), based on the premises planning data and the network data;
- 5 - determine whether the current quality of the communication network (108) is less than a predefined quality; and
- when the current quality of the communication network (108) is less than the predefined quality, use a function which maps a set of actions that are feasible corresponding to a current state of the communication network
- 10 with an expected quality of the communication network which would be achieved upon implementation of the set of actions, wherein the actions in the set are identified as the at least one action when the expected quality of the communication network is equal to or greater than the predefined quality.

- 15 9. The system (100) according to any of the preceding claims, wherein the reinforcement learning-based processing loop (224) is implemented using a reinforcement learning algorithm, the reinforcement learning algorithm being one of: a Q-learning algorithm, a Deep Q-learning algorithm, a Deep Deterministic Policy Gradient algorithm, a Q-learning
- 20 algorithm with Normalized Advantage Functions.

10. The system (100) according to claim 9, wherein the at least one server (112) is further configured to:

- perform a simulation and/or A/B testing of the communication network (108) to obtain a reference output; and
- 25 - select the reinforcement learning algorithm based on the reference output.

11. The system (100) according to any of the preceding claims, wherein the at least one server (112) is further configured to detect at least one of: a change in the premises (110), a change in the at least one floor plan of the premises, a change in the devices (102, 104, 106) in the communication  
 5 network (108), a change in performance of the devices in the communication network, wherein the premises planning data and the network data are obtained upon such detection.

12. The system (100) according to any of the preceding claims, further comprising a data repository communicably coupled to the at least one  
 10 server (112), wherein the at least one server is configured to store, at the data repository, at least one of: the premises planning data, the network data, the first output, the second output, information pertaining to the at least one action.

13. A method for managing devices (102, 104, 106) in a communication  
 15 network (108) that is implemented at a premises (110), the method comprising:

- obtaining premises planning data and network data from a service (228) associated with the premises, wherein the premises planning data pertains to at least one floor plan of the premises and the network data  
 20 pertains to the devices in the communication network and a performance of the communication network, wherein the devices in the communication network are communicably/operably coupled in a mesh network topology;
- processing the premises planning data via a first processing pipeline (202) to generate a first output;
- 25 - processing the network data via a second processing pipeline (204) to generate a second output; and
- processing the first output and the second output using reinforcement learning-based processing loop (224) to determine at least one action that

is to be implemented for managing the devices in the communication network;

wherein when the processing the network data via the second processing pipeline (204) comprises:

- 5     -     extracting, from the network data, metadata pertaining to the devices (102, 104, 106) in the communication network (108) and the performance of the communication network;
  - converting the network data, using a mesh converter (216) into an internal representation of the devices (102, 104, 106) in the communication  
10    network (108) and the performance of the communication network, wherein a format of the internal representation is different from at least one format of the network data; and
  - transforming the internal representation of the devices (102, 104, 106) in the communication network (108) and the performance of the  
15    communication network, using a mesh embedder (218), into a second vector representation, wherein the second vector representation and the metadata pertaining to the devices in the communication network collectively constitute the second output of the second processing pipeline (204).
- 20    14.    A computer-program product for managing devices (102, 104, 106) in a communication network (108) that is implemented at a premises (110), wherein the computer program product comprises a non-transitory machine-readable data storage medium having stored thereon program instructions that, when accessed by a processing device, cause the  
25    processing device to execute steps of a method of claim 13.

## PATENTTIVAATIMUKSET

1. Järjestelmä (100) tilaan (110) toteutetussa viestintäverkossa (108) olevien laitteiden (102, 104, 106) hallitsemiseksi, joka järjestelmä käsittää vähintään yhden palvelimen (112), joka on
- 5 konfiguroitu:
- hankkimaan tilan suunnitteludataa ja verkkodataa tilaan liittyvästä palvelusta (228), jolloin tilan suunnitteludata koskee vähintään yhtä tilan pohjapiirrosta ja verkkodata koskee viestintäverkon laitteita ja viestintäverkon suorituskykyä, jolloin viestintäverkon laitteet
  - 10 on kytketty kommunikatiivisesti/toiminnallisesti mesh-verkkotopologiaan;
  - prosessoimaan tilan suunnitteludata ensimmäisen prosessointikäytävän (202) kautta ensimmäisen lähdön muodostamiseksi;
  - 15 - prosessoimaan verkkodata toisen prosessointikäytävän (204) kautta toisen lähdön muodostamiseksi; ja
  - prosessoimaan ensimmäinen lähtö ja toinen lähtö vahvistusoppimiseen perustuvaa prosessointisilmukkaa (224) käyttämällä vähintään yhden sellaisen toimenpiteen määrittämiseksi,
  - 20 joka on määrä toteuttaa viestintäverkon laitteiden hallitsemiseksi; jolloin kun prosessoidaan verkkodataa toisen prosessointikäytävän (204) kautta, kyseinen vähintään yksi palvelin (112) on konfiguroitu:
  - poimimaan verkkodatasta metadataa, joka koskee viestintäverkon
  - 25 (108) laitteita (102, 104, 106) ja viestintäverkon suorituskykyä;
  - muuntamaan verkkodata mesh-muunninta (216) käyttämällä viestintäverkon (108) laitteiden (102, 104, 106) ja viestintäverkon suorituskyvyn sisäiseksi esitykseksi, jolloin sisäisen esityksen formaatti poikkeaa vähintään yhdestä verkkodatan formaatista; ja
  - 30 - muuntamaan viestintäverkon (108) laitteiden (102, 104, 106) ja viestintäverkon suorituskyvyn sisäinen esitys mesh-upotinta (218)

käyttämällä toiseksi vektoriesitykseksi, jolloin toinen vektoriesitys ja viestintäverkon laitteita koskeva metadata muodostavat yhdessä toisen prosessointikäytävän (204) toisen lähdön.

2. Patenttivaatimuksen 1 mukainen järjestelmä (100), jossa
  - 5 kyseinen vähintään yksi palvelin (112) on lisäksi konfiguroitu aikaansaamaan viestintä, joka ilmaisee kyseisen vähintään yhden toimenpiteen, joka on määrä toteuttaa vähintään yhden seuraavista muodossa: ilmoitus, web-käyttöliittymä, puhekäyttöliittymä, käyttäjälaitteeseen (226), joka on kytketty kommunikatiivisesti
    - 10 kyseiseen vähintään yhteen palvelimeen.
3. Patenttivaatimuksen 1 tai 2 mukainen järjestelmä (100), jossa kyseinen vähintään yksi toimenpide käsittää vähintään yhden seuraavista: säätää tietty laite tiettyä asetusta käyttämällä, uudelleensijoittaa tietty laite, lisätä tietty laite, poistaa tietty laite.
- 15 4. Jonkin edellisen patenttivaatimuksen mukainen järjestelmä (100), jossa tilan suunnitteludata on vähintään yhden kuvan muodossa, joka esittää vähintään yhtä tilan (110) pohjapiirrosta, jolloin kun tilan suunnitteludataa prosessoidaan ensimmäisen prosessointikäytävän (202) kautta, kyseinen vähintään yksi palvelin (112) on konfiguroitu:
  - 20 - poimimaan kyseisestä vähintään yhdestä kuvasta tilaa koskevaa metadataa;
  - jakamaan kyseinen vähintään yksi kuva tilan pohjan jakajaa (206) käyttämällä useiksi kuviksi, jotka edustavat useiden tilassa olevien kiinnostavien yksiköiden pohjapiirroksia;
  - 25 - muuntamaan useat kuvat tilan pohjan muunninta (208) käyttämällä tilan sisäiseksi esitykseksi, jolloin sisäisen esityksen formaatti poikkeaa kyseisten useiden kuvien kuvatiedostoformaattista; ja
  - muuntamaan tilan sisäinen esitys tilan pohjan upotinta (210)
    - 30 käyttämällä ensimmäiseksi vektoriesitykseksi, jolloin ensimmäinen

vektoriesitys ja tilaa koskeva metadata muodostavat yhdessä ensimmäisen prosessointikäytävän (202) ensimmäisen lähdön.

5. Jonkin edellisen patenttivaatimuksen mukainen järjestelmä (100), jossa:

- 5 - tilan pohjan jakaja (206) on toteutettu käyttämällä ensimmäistä algoritmia, jolloin ensimmäinen algoritmi on vähintään yksi seuraavista: syväneuroverkkoalgoritmi, tietokonenäköalgoritmi;
- tilan pohjan muunnin (208) on toteutettu käyttämällä useita toisia algoritmeja, jolloin useat toiset algoritmit on toteutettu semanttisena
- 10 segmentointialgoritmina ja luokittelualgoritmina; ja
- tilan pohjan upotin (210) on toteutettu käyttämällä kolmatta algoritmia, jolloin kolmas algoritmi on vähintään yksi seuraavista: konvoluutioneuroverkkoalgoritmi, itsevalvottu oppimisalgoritmi.

6. Jonkin edellisen patenttivaatimuksen mukainen järjestelmä (100),  
15 jossa mesh-upotin (218) on toteutettu käyttämällä vähintään yhtä seuraavista: muuntajamalli, autoenkooderimalli.

7. Jonkin patenttivaatimuksista 3–6 mukainen järjestelmä (100), jossa kyseinen vähintään yksi palvelin (112) on lisäksi konfiguroitu kouluttamaan tiettyä mallia optimointialgoritmia käyttämällä, jolloin  
20 optimointialgoritmi on vähintään yksi seuraavista: gradientin laskeutumisen optimointialgoritmi, Adamin optimointialgoritmi, stokastinen gradientti laskeutumisalgoritmi.

8. Jonkin edellisen patenttivaatimuksen mukainen järjestelmä (100), jolloin kun prosessoidaan ensimmäistä lähtöä ja toista lähtöä  
25 vahvistusoppimiseen perustuvaa prosessointisilmukkaa (224) käyttämällä, kyseinen vähintään yksi palvelin (112) on konfiguroitu:

- arvioimaan viestintäverkon (108) nykyinen laatu tilan suunnitteludatan ja verkkodatan perusteella;

- määrittämään, onko viestintäverkon (108) nykyinen laatu alhaisempi kuin ennalta määritetty laatu; ja
- kun viestintäverkon (108) nykyinen laatu on alhaisempi kuin ennalta määritetty laatu, käyttämään toimintoa, joka kartoittaa sarjan toimenpiteitä, jotka ovat sopivia vastaamaan viestintäverkon nykyistä tilaa viestintäverkon odotetulla laadulla, joka saavutettaisiin toimenpiteiden sarjan toteutuksen yhteydessä, jolloin sarjan toimenpiteet tunnistetaan kyseiseksi vähintään yhdeksi toimenpiteeksi, kun viestintäverkon odotettu laatu on yhtä suuri tai suurempi kuin ennalta määritetty laatu.

9. Jonkin edellisen patenttivaatimuksen mukainen järjestelmä (100), jossa vahvistusoppimiseen perustuva prosessointisilmukka (224) on toteutettu käyttämällä vahvistusoppimisalgoritmia, jolloin vahvistusoppimisalgoritmi on jokin seuraavista: Q-oppimisalgoritmi, syvä Q-oppimisalgoritmi, Deep Deterministic Policy Gradient -algoritmi, Q-oppimisalgoritmi, jossa on normalisoituja etufunktioita.

10. Patenttivaatimuksen 9 mukainen järjestelmä (100), jossa kyseinen vähintään yksi palvelin (112) on lisäksi konfiguroitu:

- suorittamaan viestintäverkon (108) simulointi ja/tai A/B-testaus referenssilähdön saamiseksi; ja
- valitsemaan vahvistuksen oppimisalgoritmi referenssilähdön perusteella.

11. Jonkin edellisen patenttivaatimuksen mukainen järjestelmä (100), jossa kyseinen vähintään yksi palvelin (112) on lisäksi konfiguroitu havaitsemaan vähintään yksi seuraavista: muutos tilassa (110), muutos vähintään yhdessä tilan pohjapiirroksessa, muutos viestintäverkon (108) laitteissa (102, 104, 106), muutos viestintäverkon laitteiden suorituskyvyssä, jolloin tilan suunnitteludata ja verkkodata saadaan tällaisen havaitsemisen yhteydessä.

12. Jonkin edellisen patenttivaatimuksen mukainen järjestelmä (100), joka käsittää lisäksi tietovaraston, joka on kytketty kommunikatiivisesti kyseiseen vähintään yhteen palvelimeen (112), jolloin kyseinen vähintään yksi palvelin on konfiguroitu tallentamaan tietovarastoon  
 5 vähintään yksi seuraavista: tilan suunnitteludata, verkkodata, ensimmäinen lähtö, toinen lähtö, kyseistä vähintään yhtä toimenpidettä koskeva informaatio.

13. Menetelmä tilaan (110) toteutetussa viestintäverkossa (108) olevien laitteiden (102, 104, 106) hallitsemiseksi, joka menetelmä  
 10 käsittää seuraavaa:

- hankitaan tilan suunnitteludataa ja verkkodataa tilaan liittyvästä palvelusta (228), jolloin tilan suunnitteludata koskee vähintään yhtä tilan pohjapiirrosta ja verkkodata koskee viestintäverkon laitteita ja viestintäverkon suorituskykyä, jolloin viestintäverkon laitteet on  
 15 kytketty kommunikatiivisesti/toiminnallisesti mesh-verkkotopologiaan;
- prosessoidaan tilan suunnitteludata ensimmäisen prosessointikäytävän (202) kautta ensimmäisen lähdön muodostamiseksi;
- prosessoidaan verkkodata toisen prosessointikäytävän (204)  
 20 kautta toisen lähdön muodostamiseksi; ja
- prosessoidaan ensimmäinen lähtö ja toinen lähtö vahvistusoppimiseen perustuvaa prosessointisilmukkaa (224) käyttämällä vähintään yhden sellaisen toimenpiteen määrittämiseksi, joka on määrä toteuttaa viestintäverkon laitteiden hallitsemiseksi;  
 25 jolloin kun verkkodatan prosessointi toisen prosessointikäytävän (204) kautta käsittää seuraavaa:
- poimitaan verkkodatasta metadataa, joka koskee viestintäverkon (108) laitteita (102, 104, 106) ja viestintäverkon suorituskykyä;
- muunnetaan verkkodata mesh-muunninta (216) käyttämällä  
 30 viestintäverkon (108) laitteiden (102, 104, 106) ja viestintäverkon

suorituskyvyn sisäiseksi esitykseksi, jolloin sisäisen esityksen formaatti poikkeaa vähintään yhdestä verkkodatan formaatista; ja

- muunnetaan viestintäverkon (108) laitteiden (102, 104, 106) ja viestintäverkon suorituskyvyn sisäinen esitys mesh-upotinta (218) käyttämällä toiseksi vektoriesitykseksi, jolloin toinen vektoriesitys ja viestintäverkon laitteita koskeva metadata muodostavat yhdessä toisen prosessointikäytävän (204) toisen lähdön.

14. Tietokoneohjelmatuote tilaan (110) toteutetussa viestintäverkossa (108) olevien laitteiden (102, 104, 106) hallitsemiseksi, jolloin
- 10 tietokoneohjelmatuote käsittää pysyvän koneellisesti luettavan tiedontallennusvälineen, johon on tallennettu ohjelmakäskyjä, jotka prosessointilaitteen päästessä niihin saavat prosessointilaitteen suorittamaan patenttivaatimuksen 13 mukaisen menetelmän vaiheita.

1/2

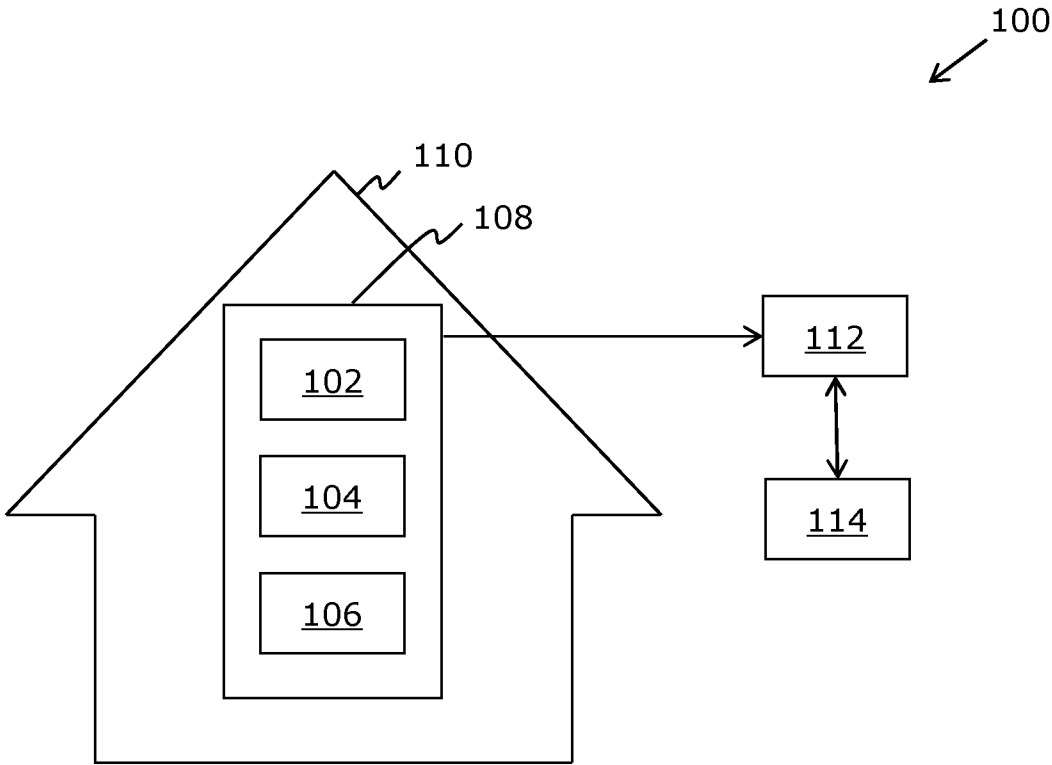


FIG. 1

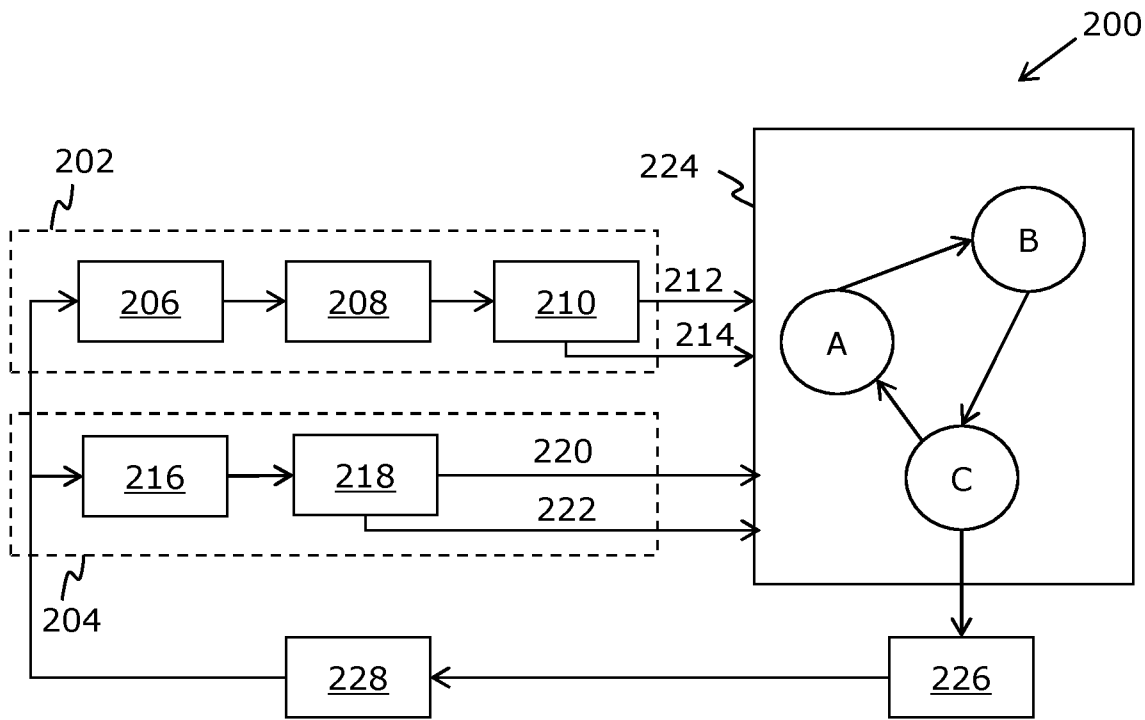


FIG. 2

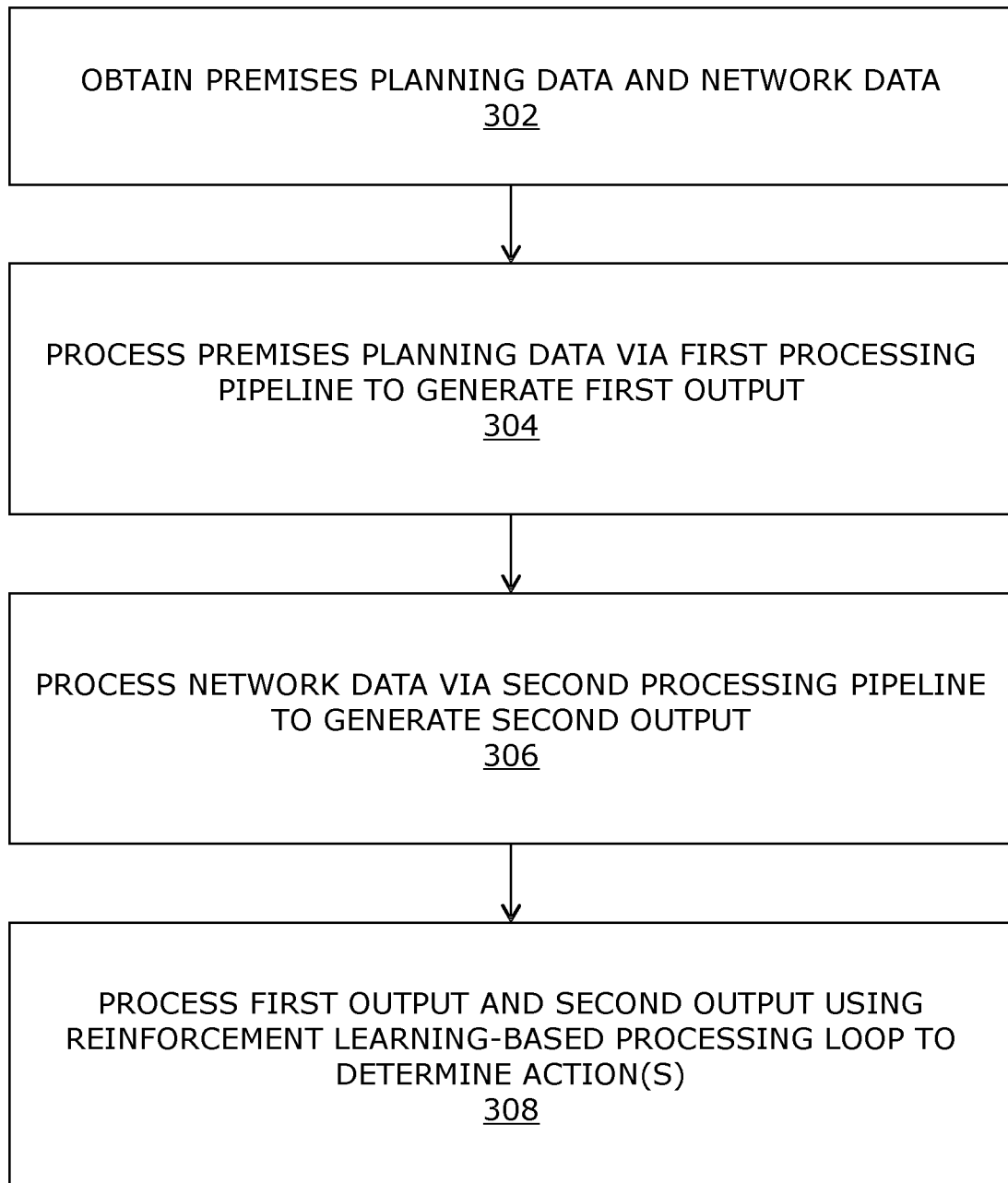


FIG. 3