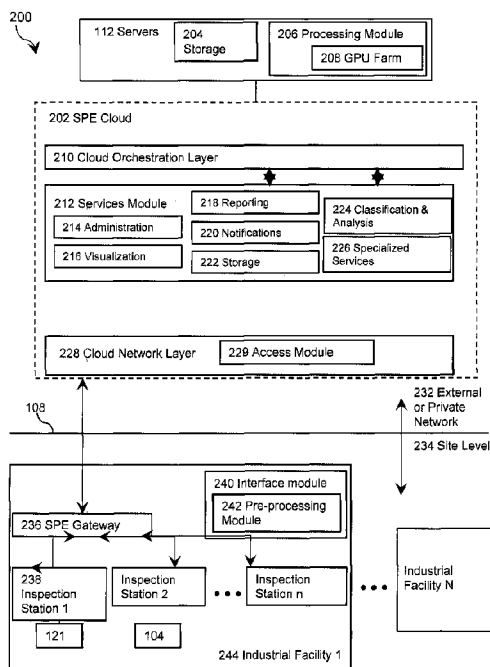




(86) Date de dépôt PCT/PCT Filing Date: 2015/06/10
(87) Date publication PCT/PCT Publication Date: 2015/12/17
(45) Date de délivrance/Issue Date: 2021/04/27
(85) Entrée phase nationale/National Entry: 2016/12/09
(86) N° demande PCT/PCT Application No.: CA 2015/050539
(87) N° publication PCT/PCT Publication No.: 2015/188275
(30) Priorités/Priorities: 2014/06/10 (US62/010,172);
2014/10/30 (US62/072,590); 2014/11/18 (US62/081,152);
2015/05/05 (US62/157,041)

(51) Cl.Int./Int.Cl. *H04L 12/24* (2006.01),
G06N 3/08 (2006.01), *H04L 12/16* (2006.01)
(72) Inventeurs/Inventors:
ALEXIUK, MARK, CA;
CASSIDY, JASON, CA;
MAVINKURVE, MAITHILI, CA;
TRENHOLM, WALLACE, CA
(73) Propriétaire/Owner:
SIGHTLINE INNOVATION INC., CA
(74) Agent: Bhole IP LAW

(54) Titre : SYSTEME ET PROCEDE POUR LE DEVELOPPEMENT ET LA MISE EN OEUVRE D'APPLICATIONS A
BASE DE RESEAU
(54) Title: SYSTEM AND METHOD FOR NETWORK BASED APPLICATION DEVELOPMENT AND IMPLEMENTATION



(57) **Abrégé/Abstract:**

An application provisioning system and method. A server provides an application provisioning service. A user of a client provides a schema defining an application. The application interacts with peripherals coupled to the client and receives input from sensors coupled to the peripherals. The sensor data is provided to the server for processing, including by neural networks. The application includes a workflow defining a finite state machine that traverses states at least partially based on the response to sensor data. The server may provide dynamic reallocation of compute resources to resolve demand for classifier training job requests; use of jurisdictional certificates to define data usage and sharing; and data fusion. Applications include manufacturing verification, medical diagnosis and treatment, genomics and viral detection.

(12) INTERNATIONAL APPLICATION PUBLISHED UNDER THE PATENT COOPERATION TREATY (PCT)

(19) World Intellectual Property
Organization
International Bureau



(10) International Publication Number
WO 2015/188275 A1

(43) International Publication Date
17 December 2015 (17.12.2015)

(51) International Patent Classification:

H04L 12/24 (2006.01) *G06N 3/08* (2006.01)
G06F 9/44 (2006.01) *H04L 12/16* (2006.01)

(21) International Application Number:

PCT/CA2015/050539

(22) International Filing Date:

10 June 2015 (10.06.2015)

(25) Filing Language:

English

(26) Publication Language:

English

(30) Priority Data:

62/010,172	10 June 2014 (10.06.2014)	US
62/072,590	30 October 2014 (30.10.2014)	US
62/081,152	18 November 2014 (18.11.2014)	US
62/157,041	5 May 2015 (05.05.2015)	US

(71) Applicant: **SIGHTLINE INNOVATION INC.** [CA/CA];
344 Westmoreland Avenue North, Suite 209, Toronto,
Ontario M6H3A7 (CA).

(72) Inventors: **TRENHOLM, Wallace**; 344 Westmoreland
Avenue North, Suite 209, Toronto, Ontario M6H3A7

(CA). **MAVINKURVE, Maithili**; 2 Fawnbrook Circle,
Markham, Ontario L3P7R9 (CA). **ALEXIUK, Mark**; 39
Lipton Street, Winnipeg, Manitoba R3G2G4 (CA). **CAS-
SIDY, Jason**; 90 Paradise Road North, Hamilton, Ontario
L8S3S8 (CA).

(74) Agent: **BHOLE, Anil**; 95 King Street East, Suite 300,
Toronto, Ontario M5C1G4 (CA).

(81) Designated States (unless otherwise indicated, for every
kind of national protection available): AE, AG, AL, AM,
AO, AT, AU, AZ, BA, BB, BG, BH, BN, BR, BW, BY,
BZ, CA, CH, CL, CN, CO, CR, CU, CZ, DE, DK, DM,
DO, DZ, EC, EE, EG, ES, FI, GB, GD, GE, GH, GM, GT,
HN, HR, HU, ID, IL, IN, IR, IS, JP, KE, KG, KN, KP, KR,
KZ, LA, LC, LK, LR, LS, LU, LY, MA, MD, ME, MG,
MK, MN, MW, MX, MY, MZ, NA, NG, NI, NO, NZ, OM,
PA, PE, PG, PH, PL, PT, QA, RO, RS, RU, RW, SA, SC,
SD, SE, SG, SK, SL, SM, ST, SV, SY, TH, TJ, TM, TN,
TR, TT, TZ, UA, UG, US, UZ, VC, VN, ZA, ZM, ZW.

(84) Designated States (unless otherwise indicated, for every
kind of regional protection available): ARIPO (BW, GH,
GM, KE, LR, LS, MW, MZ, NA, RW, SD, SL, ST, SZ,

[Continued on next page]

(54) Title: SYSTEM AND METHOD FOR NETWORK BASED APPLICATION DEVELOPMENT AND IMPLEMENTATION

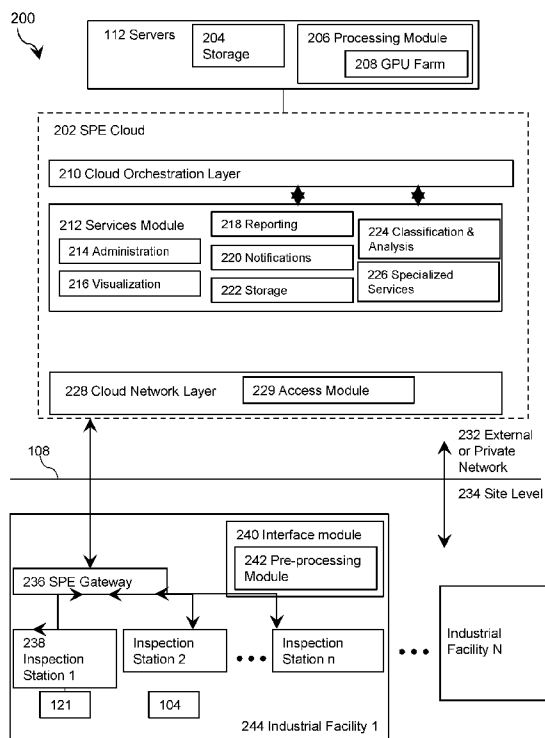


FIG. 2

(57) Abstract: An application provisioning system and method. A server provides an application provisioning service. A user of a client provides a schema defining an application. The application interacts with peripherals coupled to the client and receives input from sensors coupled to the peripherals. The sensor data is provided to the server for processing, including by neural networks. The application includes a workflow defining a finite state machine that traverses states at least partially based on the response to sensor data. The server may provide dynamic reallocation of compute resources to resolve demand for classifier training jobrequests; use of jurisdictional certificates to define data usage and sharing; and data fusion. Applications include manufacturing verification, medical diagnosis and treatment, genomics and viral detection.

WO 2015/188275 A1

TZ, UG, ZM, ZW), Eurasian (AM, AZ, BY, KG, KZ, RU, TJ, TM), European (AL, AT, BE, BG, CH, CY, CZ, DE, DK, EE, ES, FI, FR, GB, GR, HR, HU, IE, IS, IT, LT, LU, LV, MC, MK, MT, NL, NO, PL, PT, RO, RS, SE, SI, SK, SM, TR), OAPI (BF, BJ, CF, CG, CI, CM, GA, GN, GQ, GW, KM, ML, MR, NE, SN, TD, TG).

Published:

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

SYSTEM AND METHOD FOR NETWORK BASED APPLICATION DEVELOPMENT AND IMPLEMENTATION

TECHNICAL FIELD

[0001] The present invention relates generally to server systems, and more particularly to a system and method for server based application development.

BACKGROUND

[0002] Network based applications are becoming more prevalent as cloud based systems take hold. Virtualization is currently reshaping how cloud based systems are implemented. It offers new possibilities to vendors that design and supply systems and operators that license, purchase, use and manage them. Based on virtualization, system vendors no longer have to design, manufacture and maintain specific hardware. Instead, they can design their systems to run within standard virtualization systems in standard off the shelf servers and other hardware. These can be privately owned servers running in an operator's data center or the operator can contract with a cloud service provider to rent time on servers. In the cloud case, the operator does not purchase, house or maintain the physical servers. Significant cost savings to the operator can result. Significant cost savings also accrue to the operator if computations can be done close to the sensor / data acquisition. This includes derivation of features and comparison to templates e.g. CAD models or typical samples. Processing at this level also allows a layering of services around the data/feature transport such as encryption, anonymization, jurisdictional policy management.

[0003] Cloud service provided servers and others which are used for the provision of virtualization services can be difficult to develop applications for. Complex application programming interfaces and other protocols have to be learned. Moreover, typical cloud-based servers do not have access to specialized hardware such as cameras and robot control systems in a manufacturing setting. Non-optical sensing methods are considered as inputs to the system.

SUMMARY

[0004] In one aspect, a cloud based application provisioning system is provided, the system comprising a server system accessible to a client system by a network, the server system: enabling a user of the client system to define an application partially executable on the client

1 system and partially executable on the server system; and providing a plurality of application
2 provisioning services for permitting the application to access a processing module comprising
3 set of GPUs, FPGAs or ASICs for processing intensive computational tasks.

4 [0005] In another aspect, a method for cloud based application provisioning is provided, the
5 method comprising establishing a server system accessible to a client system by a network, the
6 server system: enabling a user of the client system to define an application partially executable
7 on the client system and partially executable on the server system; and providing a plurality of
8 application provisioning services for permitting the application to access a processing module
9 comprising set of GPUs, FPGAs or ASICs for processing intensive computational tasks.

10 [0006] In another aspect, a cloud based application provisioning system is provided, the
11 system comprising a network accessible server system providing a plurality of application
12 provisioning services for defining an application in accordance with a user defined schema
13 transmitted to the server system from a client system, the application being configured to
14 execute partially on the client system and partially on the server system.

15 [0007] In another aspect a method for cloud based application provisioning is provided, the
16 method comprising establishing a network accessible server system providing a plurality of
17 application provisioning services for defining an application in accordance with a user defined
18 schema transmitted to the server system from a client system, the application being configured
19 to execute partially on the client system and partially on the server system.

20 [0008] In another aspect, a system for applying a deep learning neural network to data
21 obtained from one or more sensors on a client system is provided, the system comprising: a
22 client system; a network accessible server system linked to the client system, the server system
23 configured to provision an application executable on the client system and allocating computing
24 tasks to the server system, the application causing interaction with the sensors; one or more
25 peripherals coupled to sensors for transmitting sensor data from the client system to the server
26 system the server system configure to apply a neural network based classification process to
27 the sensor data.

28 [0009] In another aspect, a method of applying a deep learning neural network to data
29 obtained from one or more sensors on a client system is provided, the method comprising:
30 linking the client system to a network accessible server system; configuring the server system to
31 provision an application executable on the client system and allocating computing tasks to the

1 server system, the application causing interaction with the sensors; transmitting sensor data
2 from the client system to the server system; applying a neural network based classification
3 process to the sensor data.

4 [0010] In another aspect, a system for network based application provisioning is provided, the
5 system comprising: a server system for maintaining schema services and classification
6 services, the schema services configured to receive a schema and create a workflow and a
7 session schema based on the schema; a hardware peripheral comprising a sensor to generate
8 sensor data; a client configured to generate a schema comprising configuration information,
9 invoke the schema services to providing a schema to the server system, and invoke the
10 classification services to provide an input to the server system from the hardware peripheral and
11 to generate an output from the classification services to effect a state change based on the
12 workflow.

13 [0011] In another aspect, a method for network based application provisioning is provided, the
14 method comprising: establishing a server system maintaining schema services and
15 classification services, the schema services configured to receive a schema and create a
16 workflow and a session schema based on the schema; coupling a hardware peripheral
17 comprising a sensor to generate sensor data to a client system; configuring the client system to
18 generate a schema comprising configuration information, invoke the schema services to
19 providing a schema to the server system, and invoke the classification services to provide an
20 input to the server system from the hardware peripheral and to generate an output from the
21 classification services to effect a state change based on the workflow.

22 [0012] In another aspect, a system for provisioning an application is provided, the system
23 comprising: an engine providing schema services; a client system accessing the schema
24 services, the client system sending a schema definition comprising metadata to the schema
25 services, the metadata comprising information on peripherals connected at a station of a
26 specialized facility; wherein: the engine generates an application and associated workflow in
27 view of the schema definition; the client system sends inputs from sensors coupled to the
28 peripherals to the engine; outputs are generated from the inputs at least in part at the engine;
29 and a state change of the workflow and application is effected in view of the inputs.

30 [0013] In another aspect, a method of provisioning an application is provided, the metho
31 comprising: a client system accessing schema services of an engine the client system sending

1 a schema definition comprising metadata to the schema services, the metadata comprising
2 information on peripherals connected at a station of a specialized facility; the engine generating
3 an application and associated workflow in view of the schema definition; the client system
4 sending inputs from sensors coupled to the peripherals to the engine; generating outputs from
5 the inputs at least in part at the engine; and effecting a state change of the workflow and
6 application in view of the inputs.

7 [0014] These and other aspects are contemplated and described herein. It will be appreciated
8 that the foregoing summary sets out representative aspects of systems, methods, apparatuses
9 for in situ electrochemical imaging to assist skilled readers in understanding the following
10 detailed description.

11 DESCRIPTION OF THE DRAWINGS

12 [0015] A greater understanding of the embodiments will be had with reference to the Figures, in
13 which:

14 [0016] FIG. 1 shows a block diagram of an example implementation of a system for network
15 based application provisioning;

16 [0017] FIG. 2 shows a block diagram of a further example implementation of a system for
17 network based application provisioning;

18 [0018] FIG. 3 shows another block diagram of the example implementation of a system for
19 network based application provisioning;

20 [0019] FIG.4 shows a block diagram of an example virtualization system;

21 [0020] FIG.5 shows a block diagram of another example virtualization system;

22 [0021] FIG.6 shows a block diagram of an example application based on the system of FIGS.
23 1 to 3;

24 [0022] FIG. 7 shows a block diagram of an example conveyor belt in an implementation for
25 monitoring manufacturing quality;

26 [0023] FIG. 8 shows an idealized diagram of an example steering wheel in the
27 implementation for monitoring manufacturing quality;

28 [0024] FIG. 9 shows example embodiment of an API configuration and functional layout for an
29 IAAS platform;

- 1 [0025] FIG. 10 shows a method for application provisioning;
- 2 [0026] FIG. 11 shows a method for application provisioning at a station of a specialized
3 facility;
- 4 [0027] FIG. 12 shows a method for training a classification model;
- 5 [0028] FIG. 13 shows example implementations of the described systems and methods;
- 6 [0029] FIG. 14 shows a plurality of flow charts illustrating exemplary implementations of the
7 described systems and methods;
- 8 [0030] FIG. 15 shows a method for generating user interfaces at a client utilizing the schema
9 services of the systems;
- 10 [0031] FIG. 16 shows a further method for generating user interfaces at a client utilizing the
11 schema services of the systems;
- 12 [0032] FIG. 17 shows data flow between a server and a client in a method for generating user
13 interfaces at a client utilizing the schema services of the systems;
- 14 [0033] FIG. 18 shows a method for monitoring manufacturing quality
- 15 [0034] FIG. 19 shows embodiment of a method to update a classification system to account
16 for changes in part numbers, versions, lighting, and other changes occurring at a station of a
17 specialized facility;
- 18 [0035] FIG. 20 shows example images of parts having rectangular foam tabs, the presence of
19 which may need to be monitored by systems described herein;
- 20 [0036] FIG. 21 shows example images of parts missing the rectangular foam tabs;
- 21 [0037] FIG. 22 shows images having a low white balance and a roughly correct white balance
22 for identifying the presence or absence of foam on a part;
- 23 [0038] FIG. 23 shows a software data flow for training a classification model and
24 implementing it at a station of specialized facility;
- 25 [0039] FIG. 24 shows software components that may be required for training a classification
26 model and implementing it at a station of specialized facility;
- 27 [0040] FIG. 25 shows an example data filtering technique;
- 28 [0041] FIG. 26 shows an example data processing technique for object classification;

1 [0042] FIG. 27 shows an example confusion matrix for monitoring classifier performance;

2 [0043] FIG. 28 shows an example data processing technique for detection of foam on a part,
3 as described with reference to Figs. 20 to 22;

4 [0044] FIG. 29 shows an exemplary block diagram of a user interface for an operator of
5 system for monitoring manufacturing quality;

6 [0045] FIG. 30 shows a block diagram of a user interface for a groundtruther to provide
7 submissions to an engine to review classification results;

8 [0046] FIG. 31 shows a block diagram of various services provided by an engine, which may
9 be accessible to different users;

10 [0047] FIG. 32 shows a block diagram of a user interface for use by an engine operator to
11 generate and manage a classification model;

12 [0048] FIG. 33 shows an example user interface landing page for logging into the engine;

13 [0049] FIG. 34 shows an example user interface page of reports available on the engine as
14 part of reporting services; and

15 [0050] FIG. 35 shows an exemplary user interface page of a parts manager page.

16 DETAILED DESCRIPTION

17 [0051] For simplicity and clarity of illustration, where considered appropriate, reference
18 numerals may be repeated among the Figures to indicate corresponding or analogous
19 elements. In addition, numerous specific details are set forth in order to provide a thorough
20 understanding of the embodiments described herein. However, it will be understood by those of
21 ordinary skill in the art that the embodiments described herein may be practised without these
22 specific details. In other instances, well-known methods, procedures and components have not
23 been described in detail so as not to obscure the embodiments described herein. Also, the
24 description is not to be considered as limiting the scope of the embodiments described herein.

25 [0052] Various terms used throughout the present description may be read and understood as
26 follows, unless the context indicates otherwise: "or" as used throughout is inclusive, as though
27 written "and/or"; singular articles and pronouns as used throughout include their plural forms,
28 and vice versa; similarly, gendered pronouns include their counterpart pronouns so that
29 pronouns should not be understood as limiting anything described herein to use,

1 implementation, performance, etc. by a single gender; “exemplary” should be understood as
2 “illustrative” or “exemplifying” and not necessarily as “preferred” over other embodiments.
3 Further definitions for terms may be set out herein; these may apply to prior and subsequent
4 instances of those terms, as will be understood from a reading of the present description.

5 [0053] Any module, unit, component, server, computer, terminal, engine or device exemplified
6 herein that executes instructions may include or otherwise have access to computer readable
7 media such as storage media, computer storage media, or data storage devices (removable
8 and/or non-removable) such as, for example, magnetic disks, optical disks, or tape. Computer
9 storage media may include volatile and non-volatile, removable and non-removable media
10 implemented in any method or technology for storage of information, such as computer
11 readable instructions, data structures, program modules, or other data. Examples of computer
12 storage media include RAM, ROM, EEPROM, flash memory or other memory technology, CD-
13 ROM, digital versatile disks (DVD) or other optical storage, magnetic cassettes, magnetic tape,
14 magnetic disk storage or other magnetic storage devices, or any other medium which can be
15 used to store the desired information and which can be accessed by an application, module, or
16 both. Any such computer storage media may be part of the device or accessible or connectable
17 thereto. Further, unless the context clearly indicates otherwise, any processor or controller set
18 out herein may be implemented as a singular processor or as a plurality of processors. The
19 plurality of processors may be arrayed or distributed, and any processing function referred to
20 herein may be carried out by one or by a plurality of processors, even though a single processor
21 may be exemplified. Any method, application or module herein described may be implemented
22 using computer readable/executable instructions that may be stored or otherwise held by such
23 computer readable media and executed by the one or more processors.

24 [0054] Referring now to Fig. 1, shown therein is a diagram of a system 100 for network based
25 application provisioning. At least one client (client 104-1, 104-2 and 104-3) is connected, via
26 network 108, to servers 112, and computing devices 114 and 116. Collectively, clients 104-1,
27 104-2 and 104-3 are referred to as clients 104, and generically as client 104.

28 [0055] Clients 104 can be based on any suitable computing environment, and the type is not
29 particularly limited so long as each client 104 is capable of receiving data from servers 112,
30 displaying data in graphical form and transmitting data to servers 112. In a present

1 embodiment, clients 104 are configured to at least execute an application that can interact with
2 the services hosted by servers 112.

3 [0056] Clients 104 can be based on any type of client computing environment, such as a
4 desktop computer, a laptop computer, a netbook, a tablet, a smart phone, a PDA, other mobile
5 client or any other platform suitable for graphical display that is known in the art. Mobile clients
6 are considered to include wearable devices which contain sensors, cameras, compute
7 resources and network connectivity such as googleGlass, GPS running watches, and health
8 monitoring devices. Each client 104 includes at least one processor connected to a non-
9 transitory computer readable storage medium such as a memory. Memory can be any suitable
10 combination of volatile (e.g. Random Access Memory ("RAM")) and non-volatile (e.g. read only
11 memory ("ROM"), Electrically Erasable Programmable Read Only Memory ("EEPROM"), flash
12 memory, magnetic computer storage device, or optical disc) memory. In one embodiment,
13 memory includes both a non-volatile memory for persistent storage computer-readable
14 instructions and other data, and a non-volatile memory for short-term storage of such computer-
15 readable instructions and other data during the execution of the computer-readable instructions.
16 Other types of computer readable storage medium external to client 104 are also contemplated,
17 such as secure digital (SD) cards and variants thereof. Other examples of external computer
18 readable storage media include resistive random access memory (RRAM), compact discs (CD-
19 ROM, CD-RW) and digital video discs (DVD).

20 [0057] Clients 104 can also include one or more input devices connected to at least one
21 processor. Such input devices are configured to receive input and provide data representative
22 of such input to the processor. Input devices can include, for example, gesture recognition,
23 voice, a keypad and a pointing device. A pointing device can be implemented as a computer
24 mouse, track ball, track wheel, touchscreen or any suitable combination thereof. In some
25 examples, clients 104 can include additional input devices in the form of one or more additional
26 buttons, light sensors, microphones and the like. More generally, any suitable combination of
27 the above-mentioned input devices can be incorporated into client 104. Clients 104 can further
28 include one or more output devices. The output devices of clients 104 can include a display.
29 When the pointing device includes a touchscreen, the touchscreen can be integrated with the
30 display. Each client 104 can also include a communications interface connected to the
31 processor. The communications interface allows client 104 to communicate with other clients,

1 for example via network 108. The communications interface is therefore selected for
2 compatibility with network 108.

3 [0058] Network 108 can comprise any network capable of linking servers 112 with clients 104
4 and can include any suitable combination of wired and/or wireless networks, including but not
5 limited to a Wide Area Network (WAN) such as the Internet, a Local Area Network (LAN), cell
6 phone networks, WiFi networks, WiMax networks and the like. Note that the networks can
7 include persistent or temporary ad hoc networks, and include sensor-compute networks where
8 sensors self-organize based on an optimal configuration to share data, compute features and
9 thereby aggregate information and use a minimal amount of network traffic for long range data
10 transfer. It is assumed that an alternate short range data communication method is also
11 provided for the sensor-compute nodes e.g. BlueTooth. It is conceived that the sensor nodes
12 communicate between themselves to determine spatial distribution of the sensors, line of sight
13 to the nearest transmission station, etc. It is conceived that nodes communicate securely, with
14 encryption, and authenticate a new sensor-node request to join the group and aggregate data.

15 [0059] In general terms, servers 112 can comprise any platform capable of processing,
16 transmitting, receiving, and storing data. In a present embodiment, servers 112 are servers
17 configured to provide network based applications and/or virtualization services. Optionally,
18 servers 112 may be hardware based servers or virtualized servers. Servers 112 can be based
19 on any desired server-type computing environment including appropriate configurations of one
20 or more central processing units (CPUs) configured to control and interact with non-transitory
21 computer readable media in the form of computer memory or a storage device. Computer
22 memory or storage device can include volatile memory such as Random Access Memory
23 (RAM), and non-volatile memory such as hard disk drives or FLASH drives, or a Redundant
24 Array of Inexpensive Disks (RAID) or cloud-based storage.

25 [0060] Servers 112 can also include one or more network interfaces, to connect to network
26 108 for example. Servers 112 can also be configured to include input devices such as a
27 keyboard or pointing device or output devices such as a monitor or a display or any of or all of
28 them, to permit local interaction. Other types of hardware configurations for servers 112 are
29 contemplated. For example, servers 112 can also be implemented as part of a cloud-based
30 computing solution, whereby the functionality of servers 112 is implemented as one or more
31 virtual machines executing at a single data center or in a mirrored form across a plurality of data
32 centers. The software aspect of the computing environment of servers 112 can also include

1 remote access capabilities in lieu of, or in addition to, any local input devices or local output
2 devices. Any desired or suitable operating system can be used in the computing environment of
3 servers 112. The computing environment can be accordingly configured with appropriate
4 operating systems and applications to effect the functionality discussed herein. Those of skill in
5 the art will now recognize that servers 112 need not necessarily be implemented as stand-alone
6 devices and can be integrated as part of a multi-purpose server or implemented entirely in
7 software, for example as virtual machines (VMs).

8 [0061] Servers 112 may be able to communicate with other computing devices 114 and/or
9 116 also connected to network 108. Computing devices 114 and/or 116 may be located at
10 specialized facilities 244 (see Fig. 2) such as manufacturing / industrial facilities or a point of
11 care (in medical implementations), and may be operably connected to peripherals 121 such as
12 a camera 120 and a robot arm 130. Specific specialized locations are described below in view
13 of particular embodiments and applications of the systems and methods described herein.
14 Clients 104 may be located at the same specialized locations and may be communicatively
15 linked to the computing devices 114 and/or 116. Alternatively, where a client 104 is located at
16 the specialized location 109, the computing devices 114 and/or 116 may be omitted and the
17 client may be operably connected to peripherals 121.

18 [0062] Network based applications are applications which are in part executed at one or more
19 servers, and partly at one or more clients 104 and/or computing devices 114, 116. More
20 particularly, network based applications are generated in view of schemas provided from the
21 client 104 and/or devices 114, 116 to the servers 112. Server portions of the applications can be
22 hosted by one or more of the servers 112. Client portions of the applications, running on one or
23 more clients and/or computing devices can access server applications as a means of interfacing
24 with the network applications for example. In some implementations, network based
25 applications may be implemented as part of a virtualization platform providing virtualization
26 services for virtualizing server functionality. In some implementations, applications may be
27 generated at the servers 112, but furnished to the client 104 and/or computing devices 114, 116
28 for execution. It will thus be appreciated that though throughout, the server based applications
29 may be described as being provided to and executed by the client 104, the applications may
30 instead be requested and executed by the computing device 116 at each facility 244 and the
31 client 104 may be located on the same network at the facility 244. Further, any exchange of
32 data may be described as being provided from the client 104 to the servers 112, but in many

1 instances data generated at a specialized facility may be provided to the servers 112 by
2 computing devices 114 and/or 116.

3 [0063] Referring to FIGS. 2 to 6 virtualization services can be provided in the form of an
4 infrastructure as a service (IAAS) platform comprising an engine, such as a cloud engine 202.
5 The IAAS platform and specifically its associated engine may alternately be referred to herein
6 as the "Sightline Perception Engine", "SPE", "Stacknation", or merely "engine".

7 [0064] An IAAS platform can offer services and resources such as a virtual-machine disk
8 image library, raw block storage, file or object storage, firewalls, load balancers, IP addresses,
9 virtual local area networks (VLANs), web servers, scripting services, programming languages
10 and others to allow implementing virtualized server systems and applications. One or more
11 virtual machine monitors (VMMs), also known as hypervisors, such as OpenStack™ cloud
12 software, can facilitate the implementation of virtualized systems and IAAS platforms, as
13 illustrated in the platform 1090 illustrated in Figs. 4 to 6. VMMs will typically consist of a
14 plurality of projects, each project allowing the control of pools of processing, storage, networking
15 and other resources associated with one or more VMs. Each project may be managed or
16 provisioned through a web-based dashboard, command-line tools or Application Programming
17 Interfaces ("APIs") for example. Accordingly, VMMs can support large numbers of virtual
18 machines and the ability to scale services up and down according to customers' varying
19 requirements.

20 [0065] To deploy virtualized applications, users can install operating-system images and their
21 application software on the IAAS platform. In this model, IAAS providers typically provide IAAS
22 services on a utility computing basis where the costs reflect the amount of resources allocated
23 and consumed. Users, on the other hand maintain the operating systems and the application
24 software on the VMs allocated to them.

25 [0066] To deploy virtualized applications, instead of relying on virtual machines, a container
26 implementation of application development and deployment can be used. According to a
27 container implementation, an 'image' of each application providing an application environment,
28 including libraries required for application execution, may be deployed to clients 104. The use of
29 container implementation may reduce outward dependencies of the applications and ensure
30 reliability, control and stability over software implementation at clients 104 having a range of

1 hardware or software. Optionally, a service such as Docker™ can be used to facilitate container
2 implementation.

3 [0067] Referring now to Fig. 2, shown therein is a particular embodiment of a system 200 for
4 network based application provision. The system 200 comprises an engine 202 hosted on
5 servers 112 providing access to services of services module 212, and further comprises a
6 specialized facility 244 in communication with the engine 202. In the illustrated embodiment, for
7 the sake of clarity of illustration the client 104 for requesting services from the engine 202 and
8 executing server based applications is illustrated to be provided at the site level 234, though, as
9 described above, the client may be provided externally to the specialized facility 244 and a
10 computing device 116 at the site may at least partly execute a server based application and
11 communicate with the engine 202.

12 [0068] In use, the client 104 may request services from the engine 202. Information from
13 sensors of peripherals 121 and information input by a user at the client 104 may be sent to the
14 engine for processing in conjunction with providing services of the services module 1002.

15 [0069] As illustrated, the engine 202 may be provided on an external or private network 232
16 accessed from a network at the site level 234. The engine 202 comprises a services module
17 212, a cloud network layer 228 and a cloud orchestration layer 210. The cloud network layer
18 228 comprises an access module 229 and facilitates communication with external computing
19 devices, such as devices at facilities 244 at the site level 234. The cloud network layer may rely
20 on network communication libraries such as ZeroMQ™, CZMQ™ or D-Bus. Data transfer
21 between different facilities may be carried out over a peer to peer network. The peer to peer
22 network may be composed of different groups of peers; a peer group shares data with its peers
23 at the lowest level of data granularity. Data sharing between peer groups occurs at an
24 aggregated level of granularity or not at all; data sharing is asymmetrical. In addition data
25 sharing may be subject to jurisdictional policies for data-sharing or meta-data sharing. The cloud
26 orchestration layer 210 orchestrates the provision and deployment of services or applications.
27 DockerSwarm™ is one orchestration service that may be used. Any orchestration method used
28 will have techniques for job distribution based on methods similar to DockerSwarm™, namely:
29 spread and binpack strategies to compute rank according to a node's available CPU, GPU its
30 RAM, and the number of containers it is running. The services module 212 comprises
31 classification and analysis services 224 for analyzing and/or classifying data provided to the
32 engine; administration services 214; visualization services 216; reporting and archiving services

218 for providing reporting and archiving functions; notification services 220; storage services 222 for storing data provided to the engine and generated at the engine; and specialized services 226 for providing application provisioning in some embodiments as described below.

[0070] The engine 202 may be implemented on servers 112 comprising hardware and logic as described above in relation to Fig. 1 to execute functions of the engine 202. The servers 112 may further comprise a processing module 206 for providing the services described above. In embodiments, the processing module 206 comprises a graphics processing unit ("GPU") farm 208 which may utilize the multiple cores of GPUs to perform computationally intensive processing operations. Management of processing module 206 for providing services to various clients 104 will be described herein. The servers 112 further comprise storage 204 for storing data provided by the storage services 222 of the engine.

[0071] At the site level 234, the system 200 comprises at least one specialized facility 244, such as an industrial facility, at least one station 238 at each site, a gateway 236, and an interface module 240. The gateway 236 provides access over a network 108 to the engine 202. Each station 238 is linked to a peripheral 121 which may include at least one sensor for generating sensor readings.

[0072] The interface module 240 comprises sensor interfaces and system interfaces for transfer of data to the engine 202. Components of the interface module may be provided at each inspection station or for a particular facility 244. Sensor interfaces provide an interface between the peripherals 121 and the client 104. System interfaces provide software interfaces facilitating data transfer to the engine, optionally providing data encryption, data compression, data anonymization, auditing and logging of information, and jurisdictional certificate stamping. Jurisdictional certificates provide information (metadata) about the origin of the data, its legal ownership and the set of data-sharing agreements to which it is subject. This metadata may or may not be contained within any meta-data "header" transferred with the data itself. However, the presence of a jurisdictional certificate enables these further details about the data to be determined by a lookup table (LUT) method from a central administrative server. It is conceived that the jurisdictional certificate is related to public key encryption methods or cyclic redundancy check (CRC) in that, if the jurisdictional certificate is compromised, decryption of the data would be exceedingly difficult. Or that accidental or malicious modification of the jurisdictional certificate could be immediately and simply determined even if the data could not be recovered. Data may be anonymized prior to transfer, for example if multiple facilities have been

1 aggregated into a data sharing agreement where anonymization is required. Logging and audit
 2 data may be pushed to the cloud locally, or pulled remotely to the cloud if required, such as to
 3 analyze system performance for monitoring the effects of configuration changes. In this context
 4 configuration includes sensor configuration settings, network configuration settings, trained
 5 classification models, sample templates or computer aided design (CAD) drawings,
 6 preprocessing settings/values, classification methods and parameters, etc. Jurisdictional
 7 certificate stamping may be provided if a data or meta-data sharing agreement is in place. Meta-
 8 data in this case refers to one of: aggregate statistics that describe higher level group
 9 characteristics, behaviours, defects or tendencies; population characterizations; differences
 10 between expected population characteristics and measured population characteristics; any
 11 business or health-related decisions that are based on sample and/or aggregate values where
 12 the actual sample level data is never communicated to the external party. Certificates may be
 13 signed by trusted third parties and may have a time expiration. The certificate identifies the
 14 origin of the data and the policies to which its use is governed. It is conceived that the
 15 jurisdictional certificates may be associated with a payment mechanism between data-source
 16 and data-sink (end user) for receipt of the data or benefit from the data. It is conceived that
 17 these payments are micro-payments. An auditing method may be applied to verify the
 18 authenticity of signed certificates for sensors. An encoded jurisdiction identifier may be
 19 associated with data and metadata transferred to the engine such that jurisdictional policies can
 20 be determined and enforced. Policies may include: right to use the data at various levels of
 21 granularity; rights to publish the data; rights to share the data at various levels of granularity
 22 within a separate network or association; rights to use the data for various purposes, in various
 23 geographic locations or markets, at various times. Jurisdictions may include academic
 24 associations, research associations, government/political associations, corporate associations
 25 and non-governmental organizations NGOs. The interface module 240 may include a pre-
 26 processing module 242 coupled to peripherals 121 for pre-processing of data from any
 27 connected sensors. Pre-processing operations may include de-noising and data normalizing.
 28 Further, for sensor data processed by compute-intensive and/or iterative data-presentation or
 29 training-based methods such as neural networks, alternating optimization (clustering), or self-
 30 organizing maps (Kohonen SOM),, features may be computed locally in the factory and
 31 transferred to the cloud engine in order to train a classifier model. Pre-processing may include
 32 feature extraction at the site level (optionally, once data is normalized), carried out with feature
 33 extraction techniques by a classification model, such as by a scene parser, as described in

1 more detail below. While the term "scene parsing" implies a labelled segmentation of a visual
 2 scene, its use in this document is both literal (for optical sensing data) and metaphorical (for
 3 non-optical sensing data). In both cases it conveys the sense of isolating valuable information
 4 content, which is determinative in reaching the desired system state or behaviour, from
 5 background or noise, which can be and should be ignored in any decision making process.
 6 Thus optical and/or non-optical sensing methods are used to collect data and subsequent
 7 compute-intensive preprocessing, feature derivation and classification activities are integrated
 8 into a decision making and actuator/effector system. Feature extraction includes low and high
 9 level techniques: edge detection, thresholding, high/low pass filtering, transforms [e.g. principal
 10 component analysis (PCA), independent component analysis ICA, fast fourier transform FFT,
 11 wavelet transform], template matching, segmentation, clustering, dimensionality reduction and
 12 the like. Feature extraction at the site level may minimize data transfer to the engine.
 13 Immediate inspection of computed features at the sensor before transfer may also identify the
 14 data as duplicate, irrelevant or otherwise not needing to be transferred. This can happen if the
 15 training dataset already has sufficient number of training samples of that type. An example is in
 16 anomaly detection where the sensor and its feature-extractor are tuned to be highly sensitive
 17 and specific to a signature range of feature values. If the sample does not have feature values
 18 in this range or matching predefined attributes, the data is discarded and possibly a counter
 19 incremented for monitoring the number of normals in the sampled population. Minimization of
 20 data transfer may be conditioned upon a calculation where network traffic and compute costs
 21 are weighed against the value that the data is likely to bring to improved classification accuracy
 22 and the ultimate business need of the end user. In this instance a cost-matrix is required to
 23 determine this decision and would need to be resident in the software-accessible memory of the
 24 sensor/factory workstation. The output of the feature extraction can be logged for audit
 25 purposes and a signal may be provided to any peripheral that has been integrated with it to
 26 effect changes at the peripheral (e.g. moving a robot arm). Features may be encrypted,
 27 compressed, logged for auditing, anonymized, and associated with a jurisdictional identifier prior
 28 to transfer to and from the engine. Pre-processing can take place on a sensor or peripheral 121
 29 that has been enhanced with compute resources e.g. lab on a chip ("LoC") or system on a chip
 30 ("SoC"), or connected to an field-programmable gate array ("FPGA") fabric, application specific
 31 integrated circuit ("ASIC"), local servers at the location or cloud servers. Further, it should be
 32 noted that individual layers of a deep net can be trained or primed independently of the entire
 33 network. This follows from the showing that RBMs can be stacked and trained in a greedy

manner to form so-called Deep Belief Networks (DBN). Computations related to these layers can run in parallel as smaller GPU jobs that are allocated independently on available resources.

Classification should be understood in a larger context than simply to denote supervised learning. By classification process we convey: supervised learning, unsupervised learning, semi-supervised learning, active / ground-truth learning, reinforcement learning and anomaly detection. Classification may be multi-valued and probabilistic in that several class labels may be identified as a decision result; each of these responses may be associated with an accuracy confidence level. Such multi-valued outputs may result from the use of ensembles of same or different types of machine learning algorithms trained on different subsets of training data samples. There are various ways to aggregate the class label outputs from an ensemble of classifiers; majority voting is one method.

[0073] Management of processing module 206 comprises provision by the engine of a job queue client that dynamically assigns computational resources based on requirements of computing jobs being carried out by hardware resources available to the engine (or at the site level 234) so that compute jobs are performed at a reduced total cost. Metrics exist to quantize job affinity and use the relative speed of each instance type for a particular job, called computing rate(CR). $CR(i, j)$ is the computing rate of instance type i for job j . If $CR(i1, j)$ is twice of $CR(i2, j)$, a task of job j runs twice as fast on $i1$ than on $i2$. This and similar metrics allow estimation of job completion, job execution cost and contribute to resolving the resource allocation problem. A cost-function optimized by this re-allocation takes into account: proximity of computation resources (e.g. GPU / FPGA / ASIC resources of the pre-processing module 242) to peripherals in order to minimize data volume transfer; derivation of discriminatory features and feature reduction early on in the processing pipeline in order to reduce data transferred or to process features in a lower cost environment; availability of low cost cloud or server farm compute resources; other factors involved including estimated time for job completion based on a solving-configuration of resources and network. The pre-processing module thus enables rapid computation proximal to the sensor in terms of FPGA fabric, ASIC, GPU, etc. It is assumed that compute resources are in some sense data-agnostic and are available for processing data from a variety of sources and for a variety of purposes (training one of a variety of classifiers, analysis of performance, parameter optimization). Data-agnostic is related to multi-tenanted servers. Each hardware resource is physically located in an environment with operating costs (cooling, electricity, network data transfer charges). Each

1 compute resource has a performance envelope which describes how much of a certain type of
2 work can be executed in a given amount of time. Various resource scheduling algorithms can
3 apply the "backpack problem" to fit the maximum amount of jobs into the performance envelope
4 of the compute resources.

5 [0074] IAAS platforms can be enhanced by facilitating application development for server
6 based applications.

7 [0075] Referring now to Fig. 3, shown therein is a specific embodiment of the engine 202 of
8 an IAAS platform implemented on a server 112 in a system for application provisioning. The
9 components and functionality of the engine 202 will be listed here and described in more detail
10 below. The engine 202 comprises a services module 212, applications module 332, definitions
11 module 322, and access module 229. The services module 212 comprises specialized services
12 226, storage services 222, and classification and analysis services 224. Broadly, the specialized
13 services 226 provide services for application development such as schema services 302 for
14 creating schema definitions, workflow services 304, session services 306 and extensions 308 to
15 the services. The classification and analysis services 224 classify and/or otherwise analyze data
16 provided to the engine, optionally in conjunction with a scene parser 314 and a match engine
17 316 and according to computational modules 318, such as a neural network 320. The definitions
18 module 322 comprises schema definitions 324. Schema definitions comprise workflow
19 definitions 326 and session definitions 328 each comprising metadata for defining application
20 templates of applications 336, as well as client sessions (in the case of the session definitions).
21 The applications module 332 comprises applications 336, and application templates 334 formed
22 from schema definitions and workflow definitions. Each application has associated workflow
23 definitions, workflow data, session definitions and session data and has access to particular
24 services. The applications can be accessed through the workflow services or session services
25 from a client. The storage services 222 provide workflow data 310 and session data 312 to
26 server storage 204. The access module 229 provides access to the engine 202. The illustrated
27 access module 229 comprises APIs 330.

28 [0076] In order to facilitate application development, an IAAS can provide several specialized
29 services 226. For example, a schema service 302 can be used to facilitate the development of
30 templates for an application. Schema services can be part of rapid application development
31 services, for example, greatly increasing the speed by which server applications can be
32 developed. Accordingly, an application 336 can be defined in the form of schema definitions

1 which can form the basis of templates 334 for that application. The schema definitions can
2 relate to metadata and can include templates for sessions and for workflows. Templates can be
3 created using known schema languages such as XML.

4 [0077] Several types of schema definitions 324 can be provided to support application
5 development in a definitions module 322. For example, session definitions 328 can be provided
6 that include definitions of metadata corresponding to application sessions. Typically, session
7 metadata can be tied to a user of client 104. For example, a session can be a creation and
8 access session of the application, and the session metadata defined could include definitions for
9 protocols for accessing the application, as well as a definition of resources needed for that
10 application session. Accordingly, all of these data and metadata requirements can be defined in
11 a session definition as a template for that application. Each application can include more than
12 one session definition.

13 [0078] As a further example of schema definitions, workflow definitions 326 can be supported
14 that include definitions of data and metadata that have relevance beyond the user of an
15 application. In some implementations, workflows can define the components and the
16 interrelatedness of components that form an application.

17 [0079] An IAAS platform can provide services in addition to a schema service 302 to facilitate
18 application development. For example, a session service 306 and a workflow service 304 can
19 facilitate the conversion of schema templates such as definitions of an application into functional
20 components of an application 336. Once session data for an application is created based on
21 session definitions, session services can allow sessions with that application to be created
22 and/or used by users. For example, applications on clients 104 can access an application
23 through the session services based on the session data and metadata for that application.
24 Moreover, session services can allow access, through workflow services, to workflow data 310
25 created on the basis of workflow definitions for that application stored by storage services 222.
26 Combination of session data and the workflow data thus allows an application to be
27 implemented at the server level, which can be accessed through clients 104, and which can
28 access all resources provided by an IAAS engine 202, including access to databases, database
29 query languages, other programming and scripting languages, software services, and other
30 services, as well as any other computing devices that may be part of the system such as
31 computing devices 114 and/or 116. In some variations, at least some of the services may be
32 external to the engine 202 and hosted by computing devices 114 and/or 116.

[0080] One of the limitations of typical IAAS implemented applications is the inability to access certain hardware components such as peripherals and resources connected to interfaces and busses such as PCI, USB, EtherCAT, Modbus and others. For example, for cloud based systems, access to interfaces for robotic arms 130, cameras 120, and other hardware components are difficult to obtain, since, for example, the VMs residing on a network do not typically have drivers for such interfaces. Indeed, as shown in FIG. 1, in some implementations, the hardware components may be located on computing devices external to the servers 112.

[0081] Extensions 308 to schema services and workflow services can be used to provide access to normally inaccessible hardware components, such as peripherals 121. Referring now to FIG. 5, application workflow definitions can be extended to allow inclusion of workflow templates that define access mechanisms for interfaces and busses and such, including mechanism to access appropriate drivers. The hardware components and/or interfaces may be collocated or directly interfaced with one or more of hardware aspects of servers 112, or maybe associated with external computing devices such as computing device 116 accessible by servers 112. Extensions to workflow services can accommodate gathering hardware component metadata based on the extended workflow definitions, enabling rapid building of applications for IAAS platforms that include access to hardware interfaces and components. Applications using extended services provided by IAAS platforms can perform a wide range of functions. For example, in one implementation, applications that can perform tasks based on complex data analysis can be built. The components of these applications can be defined with workflow and session definitions created using the schema services, and accessed as fully working applications using the session and workflow services.

[0082] To illustrate an extended IAAS application based on extended services, an example of an auto parts quality verification application will be discussed below with reference to Figs. 7 and 8. It should be noted that this is only one type of application that can be implemented, and the general methods can be used to build applications with wide range of functionalities and applicability.

[0083] Referring now to FIG. 6, simplified example application architecture is shown. An API 330 of the access module 229 allows access to the workflows for the example application. The workflows include definitions and data and have access to appropriate additional services such as a scene parser 314 and a match engine 316, of the classification and analysis services 224,

1 as well as other services, such as scalable storage and databases of the storage services 222.
2 The workflows may also receive input and provide output to client applications through the
3 access module 229, and though a workflow interface and session services (not shown) as
4 appropriate. The workflows comprising the example application may also have access to
5 peripherals 121 including hardware components such as cameras, conveyor belts and robotic
6 arms.

7 [0084] Referring now to FIG. 5, a block diagram of a conveyor belt assembly is shown at 500.
8 An example application will be described with respect to Fig. 5 for illustration. The example
9 application can have access to and/or control of one or more of the peripherals and hardware
10 components of a conveyor belt assembly. Specifically, the example application has access to
11 and/or controls peripherals 121 at a specialized facility 244, the peripherals comprising a
12 conveyor belt 510, cameras 120, lights 530 and a robotic arm 130. A conveyor belt 510 conveys
13 steering wheels 520 up towards the cameras 120. A robotic arm 130 can gate the movement of
14 steering wheels for fine positioning and/or for detailed imagery by the cameras 120. Lights 530
15 can provide appropriate lighting for image capture by the cameras 120. It is assumed that all of
16 these components can be accessed and/or controlled through one or more computing devices
17 such as client 104, computing device 116 and/or servers 112 of FIG. 1. In one implementation,
18 this access and/or control can be obtained through extensions to schematic services and
19 workflow services. The extensions to schematic services can allow developing extended
20 workflow definitions that include drivers for hardware interfaces and components. Accordingly,
21 extended workflow services can allow access to hardware interface and/or component metadata
22 based on the extended workflow definitions, enabling the example application to include access
23 to hardware components and/or interfaces of the conveyor belt assembly 500. For example, the
24 example application may be able to, through workflow and session services, control direction
25 and brightness of the lights, control the motion of the conveyor belt 510, control precise
26 positioning of a steering wheel 520 being imaged through the robotic arm 130, control
27 positioning for cameras 120, and obtain images from the cameras 120, and obtain other
28 appropriate sensor information from other components not shown. Hardware access and
29 control may be facilitated, in part by the computing devices 114 and 116. Client devices 104,
30 however would only need to access the example application through the session services
31 provided through the IAAS, greatly simplifying the complexity for front end applications residing
32 on clients 104.

1 [0085] Referring to the control of positioning of cameras, steering wheel part and conveyor,
2 an alternate method is provided whereby precise control requirements are relaxed and in its
3 place a part registration software system is provided. This registration system is calibrated by
4 the passing of a known object of known dimensions on the conveyor (this object is termed the
5 pylon) under the camera view; that when the pylon is detected and recognized the system
6 determines the camera position relative to the conveyor and pylon by analysis of the image
7 dimensions of the pylon relative to its known dimensions; based on the camera position relative
8 to the conveyor and pylon, a transformation matrix is defined which can be used to compensate
9 for change in position of the camera relative to the conveyor that differs from the position that
10 was assumed and used for training the system. Also, the parts may be placed free-hand on the
11 conveyor such that a small degree of rotation does not negatively impact the ability of the
12 system to identify the steering wheel or determine the presence or absence of defects in the
13 part. This is termed: jig-free inspection in that a physical jig is not required to accurately align
14 the part with respect to the camera.

15 [0086] Referring back to FIGS. 3 and 6, the scene parser 314 is a service for parsing data
16 from a scene and providing a break-down of the scene into appropriate components, sub-
17 components and others. For example, referring to FIGS. 7 to 8 showing an implementation of
18 the systems for monitoring manufacturing quality, a steering wheel 520 may comprise the
19 scene. Components of the steering wheel may include the circumference 610 used for
20 performing the steering, and panel 620. The panel 620 may include as components an air bag,
21 behind the slit 630, and control buttons 640, which would be the sub-components of the steering
22 wheel 520. As it will be understood, there may be many levels of components. For example, a
23 scene may have sub-components, the sub-components themselves may have components,
24 which may themselves have components, and so on and so forth. The scene parser can be
25 capable of performing a breakdown of a scene to all of its components and sub-components at
26 any level. In one implementation, the analysis of a scene and the breakdown of the scene to its
27 components can be based on computational modules 318. Computational modules can be
28 implemented using any computational paradigm capable of performing data analysis based on
29 various methods such as regression, classification and others. In some variations, the
30 computational module can be learning based. One learning based computational paradigm
31 capable of performing such methods may be a neural network 320. Neural networks may
32 include Restricted Boltzmann Machines, Deep Belief Networks and Deep Boltzmann Machines.

1 Accordingly, neural network based modules 320 can be used to perform scene breakdown by
2 the scene parser. Thus, data representing a scene, such as images, video, audio and other
3 data, as well as relevant data from databases and other services, can be provided to a neural
4 network, which can perform a scene breakdown based on classification/regression or similar
5 methods.

6 [0087] In some variations, neural networks 320 can operate in at least two modes. In a first
7 mode, a training mode 402, the neural network can be trained (i.e. learn) based on known
8 scenes containing known components to perform a breakdown. The training typically involves
9 modifications to the weights and biases of the neural network, based on training algorithms
10 (backpropagation) that improve its scene parsing capabilities. In a second mode, a normal
11 mode 404, the neural network can be used to perform the breakdown of a scene. In variations,
12 some neural networks can operate in training and normal modes simultaneously, thereby both
13 performing a breakdown of a given scene, and training the network based on the breakdown
14 performed at the same time to improve its parsing capabilities. In variations, training data and
15 other data used for performing parse services may be obtained from other services such as
16 databases or other storage services. Some computational paradigms used, such as neural
17 networks, involve massively parallel computations. In some implementations, the efficiency of
18 the computational modules implementing such paradigms can be significantly increased by
19 implementing them on computing hardware involving a large number of processors, such as
20 graphical processing units.

21 [0088] Continuing with FIG. 4, a match engine 316 is a service for analyzing data from a
22 scene and/or the scene parser to identify the scene components and thus the scene. The match
23 engine may utilize data from other services such as database or other storage services to
24 perform its functions. In one implementation, the identification can be based on computational
25 modules 318. Computational modules for the match engine 316 can be implemented using any
26 computational paradigm capable of performing identification based on various methods such as
27 regression, clustering, classification and others. In some variations, the computational modules
28 can be learning based. One learning based computational paradigm capable of performing
29 such methods are known as neural networks. Accordingly, neural network 320 based modules
30 can be used to perform identification by the match engine. Considerations discussed previously
31 regarding neural network implementations in light of parser services are also applicable to the

1 implementation of a match engine using neural networks. Accordingly, the match engine may
2 also be operated in a training mode and a normal mode.

3 [0089] In embodiments, analysis and classification services may thus be implemented by
4 providing input data to a neural network, such as a feed-forward neural network, for generating
5 at least one output. The neural networks may have a plurality of processing nodes, including a
6 multi-variable input layer having a plurality of input nodes, at least one hidden layer of nodes,
7 and an output layer having at least one output node. During operation of a neural network, each
8 of the nodes in the hidden layer applies an activation/transfer function and a weight to any input
9 arriving at that node (from the input layer or from another layer of the hidden layer), and the
10 node may provide an output to other nodes (of a subsequent hidden layer or to the output
11 layer). The neural network may be configured to perform a regression analysis providing a
12 continuous output, or a classification analysis to classify data. The neural networks may be
13 trained using supervised or unsupervised learning techniques, as described below. According to
14 a supervised learning technique, a training dataset is provided at the input layer in conjunction
15 with a set of known output values at the output layer. During a training stage, the neural network
16 may process the training dataset. It is intended that the neural network learn how to provide an
17 output for new input data by generalizing the information it learns in the training stage from the
18 training data. Training may be effected by backpropagating the error to determine weights of the
19 nodes of the hidden layers to minimize the error. Once trained, or optionally during training, test
20 (verification) data can be provided to the neural network to provide an output. A neural network
21 may thus cross-correlate inputs provided to the input layer in order to provide at least one output
22 at the output layer. Preferably, the output provided by a neural network in each embodiment will
23 be close to a desired output for a given input, such that the neural network satisfactorily
24 processes the input data.

25 [0090] Workflows, comprising workflow definitions and data, can be used to define precisely
26 how to utilize scene parsing and matching in an application. For example, workflows may
27 operate using operational states, in a manner similar to a finite state machine ("FSM"). At each
28 operational state, a workflow may receive inputs from clients 104, from one or more hardware
29 components, from scalable storage, from the scene parser and from the match engine. Based
30 on all of these inputs, the workflow may generate one or more outputs and/or effect a state
31 change. The outputs may involve changes to hardware such as movement of the conveyor belt
32 or a change in lighting quality, as well as data that may serve as input to other components of

1 the application. Outputs may also involve data provided to client applications running on clients
2 104 through session services, as well as data written to scalable storage services and others
3 that will now occur to a person of skill.

4 [0091] As a simplified example, the workflow components may define moving the next
5 steering wheel into position to be imaged by the cameras as a first application state (position
6 object). The positioning can be monitored through inputs from the cameras and controlled
7 through outputs to control the motion of the conveyor belt using the extended workflow services.
8 Once the steering wheel is appropriately positioned, the application may enter, based on the
9 workflow definitions and data, the next operational state, the recognition state. If the object
10 cannot be positioned an error may be indicated to perform manual inspection of the conveyor
11 belt. In the recognition state, the application may cause, as outputs, lighting to be chosen such
12 that the scene parser module can properly break down the steering wheel 520 into its
13 components, such as the steering radius 610 and the panel 620. The scene parser may provide
14 input to the operational state indicating any changes that might be needed in lighting and
15 position until the breakdown can be successfully performed. Upon indication of success, the
16 matching engine may be invoked to match the identified components of the steering wheel so
17 that the steering wheel can be recognized as a known type of steering wheel, i.e. indicating a
18 particular part number. The matching engine may use data services to obtain data (or the
19 appropriate matching engine in the form of a neural network for example) to match the identified
20 components to one of the known types of steering wheels. If the operation fails, the wheel may
21 be rejected, indicated as such (for example through the use of the robotic arm), and the next
22 wheel may be brought up for inspection. If the matching operation succeeds, the application
23 may enter the next state of operation where the wheel is inspected for manufacturing defects.
24 The workflows can continue in this manner until the wheel is rejected or accepted as of
25 satisfying quality.

26 [0092] In some high speed applications it may not be possible to subject the part to rapid
27 acceleration and deceleration in order to obtain a still picture for inspection. Consider low friction
28 surfaces and the extra equipment required to control part movement. In some embodiments a
29 high speed camera or video capture device is used to obtain a stream of images (video stream).
30 In this case the system analyzes the series of frames in order to select an optimal frame for
31 further analysis. Frames can be selected such that the background does not confound
32 subsequent image analysis, e.g. when a seam is present in the conveyor belt under the part this

1 can confound foreground/background separation or segmentation. Frames can be selected
2 when the part is in an optimal position, e.g. centered under the camera. Determination of when
3 the part is centered in the frame can be computed by thresholding the image and counting a
4 distribution of binary values. Also, a video stream can image-stitch a long part into a single
5 image when at any one time only a section of the part is visible and in the field of view of the
6 camera.

7 [0093] In some embodiments it may not be desirable to use a lighting source which is
8 diffused. Thus glare can occur on glossy surfaces. Image processing techniques can be
9 confounded by glare since it introduces high intensity values and high pixel intensity gradients.
10 This artifact of poor lighting can introduce errors into feature detection methods based on
11 edges, contours etc. In such embodiments, the frame from a video stream may be selected
12 such that the frame is consistently in the same position relative to the lighting which reduces the
13 effect of the glare to be that of a non-physical feature that may or may not be useful in
14 determining the identity of the part or the presence of defects. Further, the glare can be
15 detected and the area effectively masked from contributing features for classification purposes.

16 [0094] In some embodiments, multiple support vector machine or other classifiers can run
17 simultaneously, i.e. an ensemble of classifiers may be used. A classification ensemble may be
18 homogeneous or heterogeneous. A homogeneous ensemble consists of classifiers of the same
19 machine learning type e.g. a number of support vector machines (SVM). Each classifier in a
20 homogeneous ensemble will have different parameter values and may have been trained in a
21 distinct subset of the samples in the training set. A heterogeneous ensemble consists of
22 classifiers that belong to a variety of machine learning algorithms e.g. a deep neural network as
23 well as K-means clustering and an SVM. Classifiers in heterogeneous ensembles may be
24 trained on the same training data or may have been trained on distinct subsets of the training
25 data. If a multiplicity of a machine learning algorithms exists in a heterogeneous ensemble, each
26 instance of the multiplicity may have been trained on some samples unique only to that
27 instance. All classifiers can inherit an abstract base class which defines an API. A configuration
28 file may be read at execution time that loads a saved classifier model and assigns it to a mode,
29 such as an active or passive mode. The list of classifiers can be iterated through at runtime due
30 to the common inherited API. Active mode classifiers can be trusted with respect to their
31 performance capabilities and may modify some aspect of the manufacturing execution system.
32 For example, active mode classifiers may activate a peripheral, such as a robot-arm to remove

1 a defective part from the factory conveyor belt. Majority vote or similar resolution strategies may
 2 be applied when multiple active classifiers are enabled. Further, a confidence value may be
 3 assigned to each prediction of a classifier, such as part number predicted by a matching engine
 4 classifier. A message consisting of all predicted part numbers may be communicated along with
 5 confidence score. Logic based on a production cost function can be implemented to effect
 6 desired system behavior. Passive mode classifiers can compute features and run their decision
 7 making process but simply log the result to file. Passive classifiers may provide a means to
 8 assess alternate parameter values or classifier training methods or features. To minimize
 9 processing time, any features that used by more than one classifier in the ensemble may not be
 10 recomputed but a feature value can be cached. Caching can be accomplished, for example,
 11 with a dictionary data structure of key:value pairs. If decision making depends upon features
 12 derived from comparing a sample part to a template, features associated with the template may
 13 be static and may be stored as a constant value for runtime access. This may include FFT or
 14 wavelet coefficients of the template.

15 [0095] In manufacturing automation implementations of quality inspection systems, it is often
 16 the case in quality verification systems that a sample part must first be physically placed into a
 17 rigid jig in order to achieve registration and comparison with a quality reference. According to
 18 the systems and methods provided herein, in view of the embodiment described with reference
 19 to Figs. 7 to 8 and described in more detail below, it will be appreciated that the components of
 20 a scene may not need to be provided in a fixed presentation to sensing peripherals (e.g. a
 21 camera system) for quality verification. The image processing software components may have
 22 an orientation normalization module which computes an offset angle. The orientation
 23 normalization may occur as a pre-processing step before feature vectors are computed. This
 24 adaptability on the part of the system reduces the placement tolerances in a part feed system
 25 (e.g. conveyor belt).

26 [0096] Although the example application discussed above only uses session, workflow,
 27 schema, hardware access and data storage/access services, other applications may use other
 28 services provided by the IAAS. For example, VMs can be used to manage application
 29 requirements, and the application may manage its VM requirements dynamically during
 30 operation based on its workload. For example, when scene parser or matching engine require
 31 large amounts of memory to perform complex scene analysis, the memory use of the VMs may
 32 be increased dynamically to accommodate the altered requirements. In variations, the

1 application can be implemented on hardware based servers 112, as opposed to virtualized
 2 systems. In yet other variations, only portions of the application may reside in virtualized
 3 systems, the rest operating on hardware based servers. Accordingly, performance penalties
 4 that may occur due to virtualization may be minimized for portions of the application that are
 5 sensitive to such performance penalties. Other variations will now occur to a person of skill.

6 [0097] The above described embodiments provide an overview of the systems, hardware and
 7 methods. In the following, with respect to Figs. 9 to 35, specific implementations and
 8 components of the systems will be described in additional detail. Further, particular example
 9 methods for use and implementations of the systems will be described.

10 [0098] Referring now to Fig. 9, shown therein is an example embodiment of an API
 11 configuration and functional layout for an IAAS platform, and more particularly for providing
 12 access to services of an engine for application provisioning such as engine 202. As described
 13 above, the engine 202 may comprise an access module 229 having APIs 230. Particularly, an
 14 API 230 may facilitate assessing engine services, including uploading of schema definitions to
 15 provide the basis for an application and its associated workflow and session services, as well as
 16 providing access to various libraries to implement the described functionality. The illustrated
 17 configuration 800 comprises a pipeline API 802 and a Runtime API 818. Pipeline API broadly
 18 facilitates uploading media to the engine and accessing classification and analysis services 224.
 19 The Runtime API facilitates uploading schema definitions comprising metadata to provision
 20 applications, as well as implementing applications.

21 [0099] The runtime API 818 comprises a runtime.core 820, runtime.compute 830,
 22 runtime.store 836, runtime.cloud 842 and runtime.controller 848.

23 [00100] The Runtime.Core 820 API receives metadata from client 104 to provide application
 24 development for a client and session. Runtime.Core comprises Meta.FSM 822, Meta.UX 826,
 25 Meta.ML 824 and Meta.cloud 828. Meta.FSM provides a nested Finite State Machine ("FSM")
 26 implementation extended to represent planner or control logic that can drive metadata based
 27 instantiation during bootstrap procedures. Meta.UX 826 comprises metadata constructs used to
 28 describe a user experience ("UX") layout and operations. The UX layout and operations are
 29 interpreted and may be overwritten to produce a desired UX. Meta.ML 824 comprises additional
 30 metadata constructs to describe machine learning ("ML") operations and to tie Pipeline.Analysis
 31 abstractions together, including storage and distribution requirements, to resources that may be

1 dynamically discovered and removed from the system as made available by interaction with the
 2 Runtime.Compute API. Meta.Cloud 828 comprises additional metadata constructs to describe
 3 the interaction, dynamic allocation and deallocation, failover, high availability, and other
 4 “complete system” state changes. Meta.Cloud is the root level description of a deployment,
 5 which may or may not include multi-tenanting and multiple interrelated applications across
 6 tenants.

7 [00101] The Runtime.Compute API 830 comprises a root API 832 that provides resource
 8 directory service and a job queue client that manages jobs sent to the engine, and .driver that
 9 provides access to drivers including GStreamer™, Hardware IO (via e.g. PCI, DMA, Fieldbus),
 10 acceleration via FPGA, ASIC and/or GPUs, and provides system library access for machine
 11 learning libraries, such as Theano™, Caffe™, etc.

12 [00102] The Runtime.Cloud API 842 comprises .Master 844 and .Driver 846. .Master provides
 13 an extension of the Runtime.Controller class with a base implementation that loads a
 14 Meta.Cloud definition, linking the cloud orchestration UX, and each individual link to
 15 Runtime.Controller in the system. .Driver provides drivers for the UX such as via JsNation (a
 16 javascript framework designed to facilitate Big Data operations) / XAG (XML Accessibility
 17 Guidelines), as well as cloud orchestration drivers, such as for Ansible™, Docker™,
 18 OpenStack™, AWS or Ixc.

19 [00103] The Runtime.Controller 848 class comprises .Bootstrap and .Trace. .Bootstrap 850 is
 20 a root class containing runtime nodes. A base implementation loads root FSM which instantiates
 21 all runtime objects from metadata driven references. .Trace 852 allows capture, record, storage
 22 and playback of all inter object communication for simulation, testing, debug, monitoring, etc.

23 [00104] Runtime.Store 838 comprises a Root API providing access to create, read, update and
 24 delete (“CRUD”) operations, as well as Object Serialization. Runtime.Store also comprises
 25 .Driver 840 which provides access to a framework, such as JsNation’s Component Object
 26 System COS (“COS”) providing the capability to work with divergent data sources.

27 [00105] The Pipeline API 802 comprises Pipeline.Media 804 and Pipeline.Analysis 812.

28 [00106] Pipeline.Media comprises Device I/O 806, Signal/Filter 808 and Content 810.
 29 Device/IO streams input from sensors, and may be implemented via drivers to a multimedia
 30 handling framework, such as GStreamer™ as indicated in the description of the Runtime API.
 31 The Signal/Filter 808 provides image adjustment, including exposure, compensation, stitching,

1 panorama, etc. and may similarly be implemented via drivers to a multimedia handling
 2 framework. Content 810 provides access to a data stores driven by Multi-Purpose Internet Mail
 3 Extensions ("MIME")-type like content identification and may be implemented by drivers to
 4 Jsnation per the runtime API.

5 [00107] Pipeline.Analysis 812 comprises a feature extraction abstraction which may provide
 6 feature extraction techniques such as Oriented FAST and Rotated BRIEF ("ORB"), Scale-
 7 invariant feature transform ("SIFT") or Extract histogram of oriented gradients ("HoG"), which
 8 may be implemented via drivers to computer vision / image processing libraries such as
 9 Skimage, OpenCV, SciPy, scikit-image, Tesseract etc., as per the runtime API.

10 Pipeline.Analysis also comprises a machine learning abstraction 816 comprising a capture
 11 mode, training mode, and normal mode of machine learning for classification, segmentation
 12 and/or regression. The machine learning abstraction may implement Deep Neural Networks
 13 ("DNN"), Support Vector Machines ("SVM") and/or Convolutional Neural Networks ("CNN")
 14 which can be implemented via drivers to Theano™, Caffe™, etc. as per the Runtime API.

15 [00108] In addition to the libraries described above, the functionality of the systems and
 16 methods may be implemented in conjunction with enterprise integration libraries such as
 17 Structured Query Language ("SQL") and Enterprise Service Bus ("ESB"); control integration
 18 libraries, such as the Robot Operating System, and EC-master stack; video/camera libraries
 19 such as Video4Linux2, Ximea xiAPI and GigE Vision. Databases may be implemented by
 20 MongoDB to facilitate data integration.

21 [00109] Referring now to Figs. 10 to 14, shown therein are example methods and
 22 implementations of the embodiments described above.

23 [00110] Referring now to Fig.10, the systems and methods described above provide a method
 24 1000 for application provisioning according to which, at block 1002 a client accesses schema
 25 services of an engine and optionally an extension to the schema services; at block 1004 a
 26 schema definition comprising metadata is sent from a client 104 to an engine, optionally
 27 comprising information on peripherals connected at a station of a specialized facility 244; at
 28 block 1006 an application and associated workflow are generated in view of the schema
 29 definition; at block 1008 inputs are sent from the client to the engine from sensors coupled to
 30 any connected peripherals, optionally after being pre-processed by a pre-processing module; at
 31 block 1010 outputs are generated from the inputs at least in part at the engine; and, at block

1 1012 the outputs may effect a state change of the workflow and application in view of the inputs.
2 Processing of the inputs at the engine, and pre-processing at the facility 244, may be carried out
3 by computational modules such as neural networks, which may be trained in a training mode
4 before normal operation in a normal mode to process runtime data from peripherals to generate
5 outputs. Processing of inputs at block 1006 may result in dynamic changes to the application
6 (such as state changes), workflows and management of server resources.

7 [00111] Referring now to Fig. 11, shown therein is an example embodiment of a method for
8 application provisioning at a station 238 of a specialized facility 244 and training a classification
9 model for use at the station in a supervised training mode 402. At blocks 1102 to 1106 a
10 connection is established to an engine by a client; metadata and other configuration information
11 is sent using schema services from the station 238 and from the client, which may not be
12 located on the same network as the station. An application, session and associated workflows
13 are generated on the basis of the uploaded information and may be sent to a computing device
14 communicatively linked to the station for execution. At block 1108 peripherals connected to the
15 station may be configured for use with the application and workflows. At block 1110 sensors
16 connected to peripherals at the station 238 collect training data. At block 1116 the training data
17 is sent to the engine. At block 1118 the engine uses the training data to train a computational
18 module, which may comprise a neural network. The training data may be labeled and used as
19 reference data to train the computational module, such as a classification model, in a
20 supervised learning method. At blocks 1120 to 1122 a trained model of the computational
21 module is sent back to the station 238 for use in a normal mode 404. As illustrated, at blocks
22 1112 to 1114 the raw data may be pre-processed at the site and data may be labeled and
23 feature extraction techniques may be performed locally before sending the features to the
24 engine, minimizing required uploading of data. Alternatively, training may be accomplished at
25 the site and a trained model may be uploaded to the engine. In an alternate embodiment
26 illustrated by block 1124, unsupervised learning techniques may be implemented to generate a
27 trained computational module.

28 [00112] Ground truthing may be implemented to ensure classification result accuracy
29 according to an active learning technique. Specifically, results from classification models may be
30 rated with a confidence score, and high uncertainty classification results can be pushed to a
31 ground truther to verify classification accuracy. Optionally, classification outputs can periodically
32 be provided to ground truthers to ensure accuracy.

[00113] Features may thus in some instances be calculated at each of the stations 238 and in some instances utilizing classification and analysis services 224 of the engine. Local computational resources communicatively linked to the sensors may thus be provided to complement computational resources provided in a cloud engine. Feature vectors may thus be computed locally or raw data may be transferred to the cloud prior for computing of feature vectors. Training data may in some instances be collected from local inspection stations or may be collected from cloud storage. Once a model is trained, the classification model may be encrypted, compressed, logged for auditing, anonymized and/or associated with a jurisdictional identifier before transfer to/from the cloud. In unsupervised learning, raw data will be used to train models with little to no human intervention. In supervised learning, users may manually select feature sets to be used in the training process. Locally trained models may be activated when classification accuracy reaches acceptable levels. Locally trained models may be employed on the machines they are trained on, or deployed on other local machines. Locally trained models may also be pushed to the cloud for reconfiguration, further training, or analysis. Alternatively, classification models can be trained in the cloud. Cloud training offers the benefits of dedicated hardware (such as a GPU farm 208) for larger models and faster training times. Once cloud trained models are ready, they will be deployed by pushing to a site remotely, or pulling from the cloud from the site. Once deployed, cloud trained models may be pushed back to the cloud for reconfiguration, further training, or analysis. Further, even if a classification model is trained locally, an enhanced classification model can be trained in the cloud to enhance performance, or to account for changes to the configuration of the station (such as manufacture of different parts). An enhanced model can thus be trained in the cloud and then sent to the site 244.

[00114] Referring now to Fig. 12, shown therein is a method for training a classification model. At block 1202 raw data is collected from peripherals and any connected sensors. At blocks 1204 and 1205, raw data may be processed for feature extraction locally or sent for feature extraction to a cloud engine. At block 1206, interface module 240 may anonymize, encrypt, compress data and/or apply jurisdictional identifiers before transferring data to the engine. At block 1208 the data is transferred to the engine. At block 1210 the data is used as training data for training a model in a training mode 402. At block 1212 a trained model is sent back to the specialized facility 244 for use, and particularly for use at the station 238.

1 [00115] Referring now to Fig.13, embodiments of the systems and methods described above,
 2 and particularly the engine 202, may be utilized for myriad implementations, including medical
 3 implementations 1302, such as for pathogen detection, and manufacturing implementations
 4 1304.

5 [00116] Referring now to Figs. 14(a) to 14(c), shown therein are flow charts illustrating
 6 exemplary implementations of the systems and methods. Fig. 14(a) provides an implementation
 7 for monitoring manufacturing quality, such as for detecting contamination in food manufacturing
 8 or verifying parts quality as in Figs. 7 to 8. Fig. 14(b) provides an implementation for sample
 9 analysis for medical diagnostics. Fig. 14(c) provides an implementation for security and
 10 surveillance of contraband at an airport and border points.

11 [00117] In each of Figs. 14(a) to 14(c) inputs are provided from peripherals at block 1401. The
 12 inputs may then be pre-processed such as to provide image stitching of optical inputs at block
 13 1408, or other pre-processing at block 1424, which may be carried out by dedicated hardware
 14 (e.g. FPGA). Supplemental data may be received from a database, such as an enterprise
 15 information system at block 1412 of Fig. 14(a), a patient information system in Fig. 14(b) or a
 16 traveler information system in Fig. 14(c). The data inputs and supplemental data may be
 17 aggregated at a content aggregation block 1426. A feature extraction technique is then applied
 18 to the data at block 1416. The extracted data may then be classified in a normal mode at block
 19 1422, or used for training of a computational module (e.g. a classification model) at block 1418.
 20 In a normal mode outputs from the module can be provided as an output, optionally to effect a
 21 state change of an application.

22 [00118] In distributed systems, various blocks of the system may be executed at different
 23 locations. For example, data inputs at blocks 1402 and blocks 1404 may be received from
 24 different locations. Further, feature extraction at block 1416 may be effected at a different
 25 location than the classification blocks 1418 and 1422. It may be desirable that certain blocks,
 26 such as blocks 1424, 1418 and 1422 be carried out with hardware acceleration performed by,
 27 for example, FPGA, ASIC or GPU resources. It will be appreciated that different types of inter-
 28 component communications may provided, including TCP/IP, fieldbus, or other types of data
 29 communications. Each of blocks 1410 and 1411 may relate to a specific API providing
 30 dynamically configurable aspects of a coordinated cloud system, as described above.
 31 Particularly, Pipeline.Media objects 1410 relates to the Pipeline.Media API 804 described
 32 above. Pipeline.Analysis objects relates to the Pipeline.Analysis API 812 described above.

[00119] Different types of sensors / peripherals may be provided depending on the particular implementation to generate inputs for analysis and classification. As illustrated, for the detection of contamination in food manufacturing in Fig. 14(a), various optical inputs may be provided, such as an X-ray imaging input 1414. For medical diagnostics, sensors may include interferometric, spectroscopic and other optical inputs, or other sensors as described below. For security and surveillance of contraband, similarly, interferometric and spectroscopic, as well as X-ray, and other optical input can be provided to the systems.

[00120] Various specific implementations of systems and methods will now be described. Figs. 15 to 17 provide systems and methods for user interface generation. Figs. 18 to 28 broadly provide embodiments of the systems in implementations for monitoring manufacturing quality. Illustrative user interface screens and wireframe block diagrams will then be described with reference to Figs. 29 to 35.

[00121] Systems and methods for application provisioning are described above, particularly with reference to Figs. 1 to 6. Schema services are utilized by a client device for uploading data to a platform comprising an engine for application provisioning. Referring now to Figs. 15 to 17, shown therein are particular embodiments of systems and methods for generating user interfaces at a client utilizing the schema services described above. Particularly, the illustrated method generates a user interface utilizing a user interface generation technique referred to herein as "XmlAppGen".

[00122] Referring specifically to Fig. 15, at block 1502 a user interface definition is generated by an engine from metadata provided with schema services to the engine from a client. At block 1504 the user interface definition is provided to the client. At block 1506 XmlAppGen processes the user interface definition in association with a pre-defined library of widget configurations to define customized components of a user interface. At block 1508 user interface widgets are provided as a user interface for a client device.

[00123] XmlAppGen includes a set of code-generators that produce application components in various target platforms, including: Sencha Touch™ and Sencha ExtJS™ via JsNation, Node.js server components via JsNation, Python server components via StackNation™, and iOS™ and OSX™ native components using JsNation or StackNation™ server connectivity. Sencha Touch™ is a Model–view–controller ("MVC")-based JavaScript framework for building cross-platform mobile web applications. Sencha Touch™ leverages hardware acceleration techniques

1 to provide high-performance UI components for mobile devices. Sencha Ext JS™ is an
 2 MVC/MVVM JavaScript framework for building cross-platform web applications. Ext JS
 3 leverages HTML5 features on modern browsers while maintaining compatibility and functionality
 4 for legacy browsers.

5 [00124] Referring now to Fig. 16, XmlAppGen 1608 is an XML schema that can be used to
 6 declare the high-level structure and functionality of a network distributed application. JSNation
 7 1606 encompasses XmlAppGen which contains an internal library 1610 and, in an embodiment,
 8 a Sencha driver 1610. Within the Sencha driver there exists a pre-defined library of widget
 9 configurations 1604 along with a method that allows users to define their own custom
 10 components 1602. These configurations allow the use of tags in the Xml definition 1612.
 11 Finally, the Xml can be translated into a Sencha ExtJs object 1614 matching the previously
 12 defined configuration.

13 [00125] Referring now to Fig. 17, shown therein is an overview of how data is transferred to
 14 generate user interfaces according to the methods described in relation to Figs. 15 and 16. The
 15 web to server tier is facilitated by JsNation's Remote Procedure Call mechanism. JsNation's
 16 RPC framework for javascript is called the Remote Component Framework (RCF) 1706. The
 17 RCF can be used to distribute functionality across a network in the form of session objects.
 18 Encapsulating functionality in session objects conforms to Object Oriented principles and
 19 produces a centralized and in turn a more secure application. The RCF runs on top of socket.io,
 20 which comprises multi-platform library-based web sockets. Websockets are an HTTP extension
 21 which facilitate bi-directional communication streams. JsNation's RCF also facilitates peer-to-
 22 peer communication, as opposed to a traditional client/server model, where efforts are
 23 distributed between peers creating less traffic over the network. The underlying communication
 24 of RCF occurs over HTTP or HTTPS, thus having access to secure communications through
 25 TLS/SSL. This functionality may be driven by the XMLAppGen code generators.

26 [00126] Various embodiments thus rely on features of JsNation. JsNation is a JavaScript
 27 library, developed by the applicant of the present application, facilitating a rapid development
 28 framework for Node.js and HTML5 interfaces, by providing several different aspects that are
 29 brought together to speed up the development process. First, a bi-directional RPC system
 30 based on Socket.io and HTTP/HTTPS WebSockets. Also, a schema ORM that integrates SQL
 31 and no-SQL datastores in order to make converting data types a simpler procedure. Finally, UI
 32 generation tools for Sencha Touch and Sencha ExtJS which minimize development time.

1 JsNation allows software systems to seamlessly connect to components written in various
2 languages on different platforms. By utilizing an enterprise computing framework, integrating
3 with big data stores and cloud deployment systems, JsNation maximizes developer productivity,
4 fostering the flexibility and security. In essence, JsNation is a javascript framework designed to
5 facilitate Big Data operations in online multi-tiered network environments. With the emergence
6 of technologies such as nodejs, HTML5, and websockets, JsNation makes it possible to
7 implement a solution that can ease access to Big Data solutions like MongoDB™, HBase™,
8 and Cassandra™. JsNation is made of up several components that can be used independently
9 or in concert, consisting of the Remote Invocation Framework (RIF), the Object Distribution
10 Layer (ODL), Schema Framework, Component Object System (COS) and the Big Data
11 Abstraction Layer (BDAL). The RIF provides peer to peer communication in the style of remote
12 method invocation, allowing objects to live and communicate with each other on any tier. Due to
13 its flexible infrastructure, objects can also be written in other languages. That is to say, for
14 example, one peer could be written in javascript and another in Python. In order to focus data
15 and data processing, JsNation has a Schema Framework. High level information about data
16 models is defined in schemas, which allow complex operations to be performed on them simply
17 through their schema definition. This makes integration of new data models fast and robust
18 since all the information needed to handle a new data model is defined in its schema. New
19 features can be added since they can use the schema information to carry out their task and
20 then be immediately available to all data models. The ODL meanwhile fosters security, access
21 control, monitoring, logging and tracing, overall enhancing the power and flexibility of the RIF.
22 For example, real time monitoring of communication occurring in a JsNation multi-tier
23 application is possible, allowing fine-tuning of performance features to meet user demands.
24 The ability to work with divergent data sources is fulfilled in JsNation with the COS. It is often
25 the case that such data needs to be corralled and molded into a different format to fit the needs
26 of a specific application. Typically this has been done by building data warehouses that migrate
27 and reformat data. The problem that arises is keeping the data in sync, since every change in
28 the source databases must be migrated into the data warehouse. The COS may resolve this
29 issue by creating a virtual data image directly from the original databases, so any change in the
30 source may be reflected in the virtual image designed for the special application usage.
31 Applications relying on this infrastructure remain current and can allow write backs to persist on
32 the original datastores in real-time. JsNation eases complexity of working with big data clusters
33 with the BDAL, offering clean interface for talking with a backend big data implementation.

Whether using HBase™, MongoDB™ or Cassandra™, applications can be written to one standard, and implementations can be changed later. Closely linked to the BDAL is the Cross Datastore Transaction Manager, allowing the user to define transactions to ensure business operations are carried out correctly. This is a two-phase transaction management system that checks all relevant sources before committing data changes to a cluster.

[00127] Referring now to Figs. 18 to 28, implementations of the systems for monitoring manufacturing quality will now be described.

[00128] Referring now Fig. 18, shown therein is a method for monitoring manufacturing quality of parts in a manufacturing system, which may not require the use of a rigid jig for holding the parts. The method comprises: At block 1802, capturing input data from a peripheral, such as camera video stream of the manufacturing system. At block 1804, evaluating sequential frames from the video stream; and rejecting frames for which the part is not fully contained, which may be achieved by monitoring pixel values in a line pattern, accounting for glare upon the part which introduce higher pixel values than the part has under other lighting conditions, estimating the beginning and end of the part based on difference in pixel values compared to the conveyor background. Further, any frame that indicates an operator error is rejected e.g. the placement of a sample part on a section of the conveyor belt that has a seam which can affect segmentation of the foreground/background and subsequent feature calculations. At block 1806, selecting a frame for which the part is centered in the frame and as such is a consistent positioning of the part for comparison with a template image. At block 1808, optionally converting the image to a different image file format (eg. RGB, BGR, JPEG, PNG, etc.) such that the video frame undergoes a pixel value modification or compression and such that the sample image frame is provided in the same format as the images that the classifier model has been trained on. At block 1810, computing the outer contour of the part against the background e.g. conveyor. At block 1812, comparing the outer contour of the sample to a template image (reference image) of the part. This reference image could be a computer aided design (CAD) reference file. At block 1814, estimating the orientation of the part to left-handed and right-handed versions of the part by computing pixel-moments or geometrical features which distinguish right- and left-handedness. At block 1816, estimating the degree of rotation of the sample part with respect to a reference image by a method of moments; determining if the sample is the same or opposite handedness based on a threshold applied to the rotation estimate; rotating the sample image such that it is aligned to the template image and calculating further features on the sample

1 image at that point. At block 1818, allowing for a degree of rotation of the sample part with
2 respect to the reference image; rotating the sample image; running a template matching
3 algorithm on the rotated sample image; storing similarity scores for a range of rotations;
4 selecting the maximum similarity score and using the associated rotation value to rotate the
5 image before further features are calculated.

6 [00129] Referring now to Fig. 19, shown therein is an embodiment of a method to update a
7 classification system to account for changes in part numbers, versions, lighting, and other
8 changes occurring at a station 238. At block 1902, number and type of sample parts are
9 recorded. At block 1904, an independent factory system records the number and type of sample
10 parts. At block 1906, the two records are compared and a confusion matrix is generated. Fig. 27
11 illustrates an example confusion matrix. At block 1908, the performance accuracy as
12 documented in the confusion matrix is compared to client-vendor specified acceptable
13 performance levels. At block 1910, log files are reviewed for internal state details that have
14 bearing on the feature values derived and classification decisions. At block 1912, when
15 deficiencies in performance are noted, a system operator and a classification designer in charge
16 of designing the classification model are notified. Deficiencies may result from behavioral
17 deviance of an operator, changes in lighting and introduction of a new / modified part. At block
18 1914, behavioral deviance of the system operator from defined quality procedures can be
19 determined, such as by measuring placement of the part on the conveyor, e.g. whether the part
20 is on a seam or is too rotated for measurement purposes. These deviances can be logged and
21 manufacturing quality officers notified. At block 1916, lighting changes can be determined by
22 comparing pixel histograms of acquired images with reference images from a calibration
23 session. The operator can be advised to restore an acceptable lighting environment or image
24 transformation methods can be applied such that the trained classification system is robust to
25 such changes. At block 1918, when a new part is introduced, a whitelist file can be used to
26 enable capture of images of certain part types and these images can be transferred to a training
27 server where a new classifier can be developed and re-deployed to the factory.

28 [00130] When training a classification model a predetermined number of images may be
29 required to ensure classification model accuracy. A prerequisite when acquiring, training and
30 deploying systems may be to determine the best lighting conditions, camera position and focus.
31 Once determined, changes to lighting, position or focus may result in misclassification. For
32 example, referring now to Fig. 20, shown therein are example images of parts having

1 rectangular foam tabs, the presence of which may need to be monitored. Fig. 20 shows parts
2 with single rectangular foam 2002. Fig. 21 shows images of parts missing the presence of the
3 foam. Once trained, a classification model could validate the presence of the foam 2002. Fig.
4 22(a) and Fig. 22(b) show images having a low white balance, and a roughly correct white
5 balance, respectively, for identifying the presence or absence of foam on a part. In particular
6 implementations the number of images required to train a new part may be approximately one
7 hundred. Acquiring images can be done through a user interface, by selecting the part number
8 which comes from a database and recording these images. For an exemplary implementation
9 of parts which have foam configurations, foam and no foam configurations can be recorded.
10 After recording images of all parts that the client determines will be needed for the system, the
11 system can be trained on the recorded images. Once trained, the system can classify any of
12 the parts the system has been trained on.

13 [00131] Updating a part can be accomplished by retraining the classifier model. If a new part is
14 identical to a previous part, a part number translator can be used for a one-to-one replacement.
15 Specifically, a part label 2004 may be changed through a translator module which may map the
16 new part label to the old part label and visa-versa. If the replacement part has visible changes,
17 in addition to updating the part label, the classifier model should be tuned. In the embodiments
18 illustrated with respect to Figs. 20 to 22, for example, where the presence of foam is monitored,
19 tuning the model may comprise updating foam position (in case of change in foam position),
20 updating foam classification threshold (in case of change in foam material) and retraining
21 Principal component analysis ("PCA") and SVM for tab foam detection. If the new part is a part
22 for which the system has never before been trained, then a full training data image set may
23 need to be obtained for the new part. The training data image set can be sent to a cloud server
24 in order to train a model, and the trained model can be returned to the station for use in parts
25 analysis.

26 [00132] Referring now to Figs. 23 to 27, shown therein are block diagrams of software data
27 flows and components of an illustrative embodiment for use in various applications described
28 herein for providing prediction from a computational module, such as a machine learning model.
29 The embodiments will be described, for clarity of illustration, with reference to monitoring
30 manufacturing quality but can be used for other implementations of the systems and methods
31 described herein.

[00133] Referring now to Fig. 23, a training mode 2302 of a model for use at a station of a specialized location at the site level 234 comprises: receiving sensor data at block 2304 from connected peripherals and any associated sensors; applying a customized data filtering technique at block 2306; uploading filtered data to the cloud at block 2308 and downloading a trained model from the cloud at block 2310. A cloud engine 2322 of servers 112 may provide training of a model by, acquiring data from the specialized location at block 2324, applying customized data processing techniques to the data at block 2326, and training the model at block 2328, before sending it back to the station. Once the model is trained and provided to the station, the station can operate in a normal mode for manufacturing production 2312, wherein the model is applied to classify data by: receiving data from sensors at block 2314, applying the customized data filtering techniques at block 2316, applying the customized data processing techniques at block 2318, and generating a prediction from the model at block 2320.

[00134] Referring now to Fig. 24, software components at the site level 234 and at the cloud servers 112 are illustrated. Particularly, software components provide the functionality of each of the blocks described in relation to Fig. 23. Additionally, a software component 2402 may be provided at the server in order to manage server resources.

[00135] Referring now to Fig. 25, shown therein is an example data filtering technique for use particularly in an application where an image is provided from a peripheral for analyzing a part shown in the image. Data filtering of each frame may select only image frames wherein the part is shown, the left and right edges of the image frame are not overlapped by the part, and where multiple frames are provided for a part. Further, the first frame where the center of the object is on the right side of the frame may be selected.

[00136] Referring now to Fig. 26, shown therein is an example data processing technique for object classification and particularly for distinguishing different parts (or other objects) using image frames which may not be distinguishable by the human eye. At blocks 2602 to 2608 a template pool of images can be prepared. At block 2602 objects in received images can be masked from the background. At block 2604 edge detection can be applied. At block 2606 key points detection may be carried out. At block 2608 key points from a number of image frames are collected as a template pool. Once a template pool is established, the blocks illustrated at blocks 2610 to 2620 may then be carried out to process filtered image frames. At block 2610 the object in each frame can be masked from the background. At block 2612 edge detection can be applied. At block 2614 key points detection may be carried out. At block 2616 the key points can

1 be matched with key points in the template pool (e.g. using Hamming distance), and the
 2 presence of the key points in the template pool can be used as input to a machine learning
 3 technique. At block 2618 a geometry characteristic, such as length, and an object surface
 4 characteristic can be calculated and used as input to a computational module, such as a
 5 machine learning classification model. At block 2620 a machine learning model may be applied.

6 [00137] Referring now to Fig. 27, classifier performance can be visualized with a confusion
 7 matrix. In the illustrated example, a matrix of parts is shown with each cell representing how the
 8 classifier assigns a part label to the sample-under-test, such as samples under test on a
 9 conveyor belt in a manufacturing quality implementation.

10 [00138] Referring now to Fig. 28, shown therein is an exemplary data processing technique for
 11 use in an example implementation for detecting a subcomponent of a part, such as for detecting
 12 foam as described with reference to Figs. 20 to 22. At blocks 2802 to 2804, template images of
 13 a part may be prepared. At block 2802 a template image of a part may be masked from the
 14 background. At block 2804 the masked part may be moved to the middle of a square image with
 15 a black background. At blocks 2806 to 2822 the technique can be used to prepare images of a
 16 part for use in image classification. At block 2806 an image of a part can be masked from the
 17 background. At block 2808 the masked part can be moved to the middle of a square image with
 18 a black background. At block 2810 the resultant image can be registered with a prepared
 19 template. At block 2812 a histogram of the registered image can be examined at a possible
 20 foam location. At block 2814, for parts with foam, affine transformations can be applied that
 21 were used for registration to the original image. At block 2816 the square image can be cropped
 22 to a possible foam region. At block 2818, histogram oriented gradients techniques can be
 23 applied to the cropped region. At block 2820, the dimensions of the resultant histogram oriented
 24 gradients vector can be reduced using PCA. At block 2822, a trained classifier (such as an SVM
 25 classifier) can be applied.

26 [00139] Referring now to Figs. 29 to 35, shown therein are illustrative user interface screens
 27 and wireframe block diagrams illustrating how operators of the systems described herein can
 28 carry out the functionality of those systems.

29 [00140] Fig. 29 provides an exemplary wireframe block diagram of a user interface for an
 30 operator of system for monitoring manufacturing quality, comprising blocks 2902 to login and

1 modify parts (e.g. add / delete parts) that can be analyzed by the engine, as well as blocks 2904
2 for operators to access to monitor parts quality during production.

3 [00141] Fig. 30 provides a wireframe block diagram of a user interface of a mobile application
4 for ground truthing. According to a contemplated embodiment, a mobile application may be
5 provided to facilitate ground truthing by providing submissions to an engine to review
6 classification results. The mobile application may be communicatively linked to an IAAS platform
7 comprising an engine and associated databases and classification model results. As described
8 above, ground truthing may be implemented to ensure classification result accuracy according
9 to an active learning technique. Specifically results from classification models may be rated with
10 a confidence score, and high uncertainty classification results can be pushed to a ground truther
11 to verify classification accuracy. A threshold on confidence level can automatically trigger a job
12 to ground-truth. Optionally, classification outputs can periodically be provided to ground truthers
13 to ensure accuracy. If a classification output is to be sent to a ground truther, a given input and
14 current classification labeling can be entered into a job queue and provided to a ground truther
15 for review. Further, results of submissions from ground truthers may be sent to the engine to
16 dynamically modify the configuration of a particular classification model. The mobile application
17 may provide a sign up 3002 feature so that individuals can provide ground truthing for
18 classification models. A possible ground truther may sign up to the application and provide
19 some description about their expertise, availability, interests, geographic location. The ground
20 truther may subscribe to different types of tasks. Upon receipt of a submission from a
21 groundtruther, an operator of the classification model could be alerted automatically that a
22 review should begin. The operator could review the submission and update any relevant client
23 databases based on the new information.

24 [00142] Fig. 31 provides a wireframe block diagram of various services provided by an engine,
25 which may be accessible to different users. For example, an administrator may have access to
26 configuration services.

27 [00143] Fig. 32 provides a wireframe block diagram of a user interface for use by an engine
28 operator to generate and manage a classification model.

29 [00144] Fig. 33 provides an exemplary user interface landing page for logging into the engine.

30 [00145] Fig. 34 provides an exemplary user interface page of reports available on the engine
31 as part of reporting services.

1 [00146] Fig. 35 provides an exemplary user interface page of a parts manager page.

2 [00147] Various implementations are contemplated for embodiments of the systems and
3 methods. Some example implementations will be described below. The listed implementations
4 are merely illustrative and are not limiting. Further implementations would be contemplated by
5 those of skill in the art.

6 [00148] Various implementations for monitoring manufacturing quality are contemplated, as
7 described with reference to Fig. 14(a). Surface Inspection implementations of two and three-
8 dimensional objects are contemplated, optionally over time, such as for automotive glass
9 inspection. For example, exterior / interior surface inspection of a vehicle for variable and
10 random defects in paint could be carried out. Coating and painting inspection implementations
11 are further contemplated. With regards to medical coatings, inspection may be required at the
12 two microns level. Implementations of optical interference patterns to detect surface defects are
13 contemplated, such as to provide inspection of contact lenses to detect surface defects, such as
14 holes. Braid inspection implementations are contemplated, such as for inspecting braided hoses
15 in lengths of up to 1 mile which may require inspection for missing braid strands or braid strand
16 sub components at up to 1 ft/sec. In addition a braid may need to be inspected for variance to
17 the angle of the braid relative to a hose. Microscopic inspection of defects to identify a type of
18 contamination for root cause analysis is contemplated, such as for manufacturing forensics.
19 Motor stator wire winding pattern inspection implementations are contemplated to detect
20 random bunching or contamination that could lead to premature failure. Image processing for
21 some mobile sensor networks, such as for unmanned aerial vehicles (UAV) is contemplated,
22 such as by using a deep convolutional network providing real time classification and using on
23 board processing via FPGA for computation. Compressed sensing may be provided to make
24 downlink of data possible, and to use distributed computation.

25 [00149] With respect to implementations relating to the detection of contamination in food
26 manufacturing, as illustrated in Fig 14(a), in order to provide a uniform product to consumers,
27 industrial food production requires a method to ensure ingredients are consistently mixed. By
28 combining inputs of multiple sensors (force, elasticity, density, opacity/translucency, pressure,
29 temperature, flow, proximity, displacement, moisture, hyperspectral imaging, 3D scanning, X-
30 ray, spectroscopy, Timers, Ultrasound, Camera, and Hyperspectral), embodiments of the
31 engine could inform an operator that the mixture contains the correct amount of each ingredient
32 and that the ingredients been completely mixed. Further, processed food suppliers have

1 regulatory requirements to inspect products for various contaminations before and during
2 production. This is typically done as batch testing by humans of raw materials and finished
3 goods, however this methodology allows for the potential of some contamination reaching the
4 consumer. By scanning ingredients at the raw material and production stages with hyperspectral
5 sensors, the engine could be implemented to detect contaminants such as bone, bacteria, and
6 hair that would otherwise go undetected. Additionally data can be stored and used for audit
7 purposes to isolate sources of contamination.

8 [00150] Referring again to Fig. 14(c), implementations for inspection for security and
9 surveillance are also contemplated. Inspection for contamination, pathogens, or quality
10 deviation in food manufacturing are contemplated, using embodiments of the engine, and
11 optical imaging, X-ray imaging, Terahertz imaging, millimetre wave imaging, spectrographic
12 imagers or sensors, interferometric imagers or sensors, electronic nose, and/or other sensors.
13 Further, inspection for contamination, pathogens, or quality deviation in pharmaceutical
14 manufacturing are contemplated, using embodiments of the engine, and optical imaging, X-ray
15 imaging, Terahertz imaging, millimeter wave imaging, spectrographic imagers or sensors,
16 interferometric imagers or sensors, electronic nose, and/or other sensors. Further, inspection for
17 contamination, pathogens, or contraband at the border or for airport luggage inspection are
18 contemplated, including for food, plants, fungal contamination, etc., and other items prohibited,
19 using embodiments of the engine, and optical imaging, X-ray imaging, terahertz imaging,
20 millimetre wave imaging, spectrographic imagers or sensors, interferometric imagers or
21 sensors, electronic nose, and/or other sensors. Further, augmentation of current airport and
22 border screening processes, leveraging collaborative machine learning, active machine
23 learning, and realtime training data set creation by processing of human operator actions are
24 contemplated, using embodiments of the engine, and optical imaging, X-ray imaging, terahertz
25 imaging, millimetre wave imaging, spectrographic imagers or sensors, interferometric imagers
26 or sensors, electronic nose, and/or other sensors are contemplated. It is contemplated that an
27 alert module could activated as a result of the output from an engine applied for security and
28 surveillance. For example, if a pathogen is detected, a local security or medical professional
29 could be notified.

30 [00151] Referring again to Fig. 14B, various medical diagnostics / analysis implementations
31 are contemplated. Detection of pathogens (but not merely of pathogens) may be carried out by
32 receiving detection signals from peripherals / associated sensors ("detectors") which may take

1 any form that is suitable for the detection of a desired substance from a sample, or for detection
2 of other relevant features. Typically, a sample may contain concentrations of a substance to be
3 detected. The desired substance to be detected can be pathogens to be detected or indicators
4 of the presence of the pathogens to be detected. Imaging modalities and appropriate sensors
5 will be selected for detecting the presence of the pathogen or indicators of the presence of the
6 pathogen. In particular embodiments, computation modules, such as neural networks are
7 provided with input signals from sensors. The computational modules may be trained to classify
8 the input signals, such as by indicating whether input signals relate to reference signals
9 indicating the presence or absence of a substance.

10 [00152] A detector could be a molecular sensing array designed to be sensitive to one or more
11 desired substances. Further, the detector could perform Surface-Enhanced Raman
12 Spectroscopy (SERS), Surface Plasmon Resonance (SPR), Surface Plasmon Resonance
13 Imaging (SPRi), Localised Surface Plasmon Resonance (LSPR), Optofluidic Nanoplasmonic,
14 Optical waveguide-based sensing, Optical ring resonator-based sensing, Photonic crystal-based
15 sensing, Nanosensitive OCT sensing, Lensless digital holographic imaging, Superresolution
16 microscopy techniques, piezoelectric sensing, nano-cantilever sensing, Raman spectroscopy
17 (RS), Resonance Raman spectroscopy (RRS), and infrared spectroscopy (IRS). In variations,
18 the detector could perform interferometer-based detection, such as by using a Mach-Zehnder
19 Interferometer, Young's interferometer, Hartman interferometer, Interferometric scattering
20 microscopy (iSCAT), Single Particle Interferometric Reflectance Imaging (SPIRIS) and
21 backscattering interferometry. Other detectors based on other detection techniques will now
22 occur to a person of skill and are contemplated.

23 [00153] Aside from imaging sensors, other inputs relating to a patient may be provided for
24 medical diagnostics implementations including temperature sensing, blood pressure, heart rate,
25 patient medical history (as illustrated in block 1412), height, weight, demographic information
26 etc.

27 [00154] According to a first embodiment, a neural network interprets received detection signals
28 from a detector. The selected neural network may be configured as a convolutional feed-forward
29 neural network. Optionally, the neural network may receive at least one detection signal as an
30 input and output an indication of whether a particular detection signal relates to a reference
31 signal indicating the presence of a bound desired substance, or a reference signal indicating
32 that a desired substance has not bound the sensing surface, as described above. Accordingly,

1 during use at least a measured detection signal, or some scaled or otherwise modified value
2 thereof, will be provided to the neural network as an input. Optionally, additional data may be
3 provided to the input layer of the neural network to assist in interpreting received detection
4 signals from a detector. Combinations of data could be provided at the input layer, including:
5 protein interaction data (e.g. of the pathogen), and genomic / nucleic acid data of the pathogen,
6 subject and/or desired substance (i.e. biomarker). Accordingly, high-throughput genomic
7 sequencing of the subject / pathogen may be required, but could be performed by remote
8 computers 160 and need not be carried out at the local device 140. Further input data could
9 include mass spectrometry data (e.g. from pathogen protein sequencing). This embodiment
10 may thus cross-correlate various inputs to provide an output to aid in interpreting a detection
11 signal to determine whether a pathogen has been detected.

12 [00155] An output indicative of detection of a pathogen may result in notification being
13 generated to alert a local medical professional; alert a medical professional already associated
14 with the patient; alert an expert in the healthcare field with special knowledge of the specimen;
15 and, alerting local health or police authorities as required by law for diagnosis of health
16 conditions. Further, an output indicative of detection of a pathogen may result in generating a
17 request for human ground-truthing of the detection signal / sample. For example, a microscopic
18 image of a sample can be electronically transmitted to ground truther for assessment. Further,
19 the patient may be advised of any immediate actions that they should take for their own
20 immediate health and safety and for the public in the vicinity.

21 [00156] In another embodiment, a neural network is applied to compensate for nanoscale and
22 quantum realm detection limitations. This is super-resolution or sub-pixel imaging or
23 compressed sensing photography or sub-wavelength imaging. Particularly, a detection signal is
24 provided to a neural network, with a desired output compensating for defects in the detection
25 signal that may be caused by limitations of imaging in the nano-realm. The input layer may
26 receive data relating to the detection modality and an input detection signal for detection of
27 viruses, bacteria, fungi, parasites, human host (i.e. subject) cells, disease biomarkers, etc. The
28 neural network may be trained such that the output layer provides a clean detection signal
29 compensating for signal defects. Particularly, the neural network may be trained with a training
30 dataset comprising, at the input layer, detection signals comprising nano-realm defects, and
31 with associated clean detection signals at the output layer for viruses, bacteria, fungi, parasites,
32 human host cells, disease biomarkers for detection modalities. The output of the trained neural

1 network may provide a processed detection signal similar to known reference signals for
 2 particular detection modalities such that processing by the neural network remedies some
 3 defects and limitations of received detection signals.

4 [00157] In another embodiment, a neural network is applied to drive evolution of the choice of
 5 reaction substances provided for pathogen detection with a sample. Particularly, selection of a
 6 reaction substance may compensate for mutation, pleomorphism and polymorphism of
 7 pathogens to ensure that appropriate reaction substances are selected to maximize likelihood of
 8 detecting a pathogen. Accordingly, inputs including a combination of time series genomic data
 9 of the pathogen and/or human host cells, and data relating to a plurality of desired substances
 10 (e.g. disease biomarkers) may be provided to a neural network trained to provide an output
 11 indicating which reaction substance(s) should be selected. A sensing surface comprising, for
 12 example, a selected biosensor immunoassay antigen capture mechanism could then be
 13 provided to each reaction chamber of a collection unit. The embodiment may similarly require
 14 high-throughput genomic sequencing of the subject / pathogen, as well as mass spectroscopy
 15 data.

16 [00158] In another embodiment, a neural network based predictive output machine is provided.
 17 Particularly, the machine learning predictive output machine may receive inputs comprising time
 18 series genomic data of a subject in order to provide an output indicative of a clinical outcome.
 19 To provide time series inputs, samples may be taken and sequenced from a subject and/or
 20 pathogen over a period of time to maintain or improve the accuracy of the neural network over
 21 time. To train the neural network a training dataset may comprise known inputs of the specified
 22 data types as well known associated clinical outcomes. Further data inputs may include, time
 23 series genomic data of various pathogens and protein interaction in the pathogen and host,
 24 subject history (e.g. flight history or medical history), subject condition (e.g. resistivity to aids).

25 [00159] In view of the above implementations for medical diagnostics / analysis, a method is
 26 contemplated wherein: foreground / background scene parsing , foreground and background
 27 separation in image and non-image processing techniques for a classification model that is
 28 extended to genomic analysis, proteomic analysis, protein-protein interaction, through the use of
 29 machine learning techniques including deep neural networks; the genomics ("ACGT") string
 30 mapping problem is translated to a format which is suitable for solution by deep nets, which may
 31 include a mapping such as closed polygon to partial string techniques; a generative model for

1 genetic coding anomalies is constructed, and outputs of generative models are compared in an
2 optimized manner to sample data using hashes at the sensor location (Bloom filters).

3 [00160] Implementations for precision medicine are contemplated. Precision medicine is an
4 emerging approach for disease treatment and prevention that takes into account individual
5 variability in genes, environment, and lifestyle for each person. With inputs from sensors,
6 embodiments of the engine may be implemented to analyze variation in human genomic data,
7 genomic expression, proteomic data, protein expression, phenotypic data, symptoms, and
8 external data inputs including environmental factors, diet, symptoms, demographic information,
9 exercise, etc. Based on the data collected, tracked and provided to the engine, analyzed data
10 may help provide information to tailor a treatment approach for a patient.

11 [00161] Implementations for analyzing pathogen variation are contemplated. In embodiments,
12 the engine can analyze variation in pathogen (including virus, bacteria, parasite, fungus, viroid,
13 prion) genomic data, genomic expression, proteomic data, protein expression, phenotypic data,
14 symptoms, and external human and / or animal data inputs including environmental factors, diet,
15 symptoms, demographic information, exercise, etc. to identify, track and or help treat
16 pathogenic infections. Based on the data collected, tracked and provided to the engine,
17 analyzed data will help provide information to identify a pathogen, classify a new pathogen and
18 or recommend treatment approaches. A particular implementation of a method for analyzing
19 pathogen variation may comprise polling scientific literature to assess pathogen variation.
20 Particularly, the method may comprise selecting a relevant set of scientific journals that will be
21 sufficient to update a database with respect to a specific pathogen, its evolution, its relation to
22 drug efficacy, and population cohorts that are especially susceptible. The engine may download
23 new content and run a text mining technique on the data to note references to the pathogen.
24 Text association rules may be applied. A summary of inferences may be provided for review by
25 an appropriate operator of the engine. Upon approval by the human operator, the content may
26 be added to a curated database. An audit trail may be maintained about the content used to
27 make the decision, who made the decision to include the new content.

28 [00162] Implementations for drug discovery are contemplated. Drug discovery is the process
29 by which new candidate medications are discovered, often through the screening of a library of
30 molecules or compounds in the thousands. Coupling and measuring drug candidate profiles
31 against known and unknown disease markers helps identify either compounds with similar
32 efficacy profiles or new drug targets. Based on the data collected, tracked and provided to the

1 engine, analyzed data will help provide information to identify suitable drug candidates in a drug
2 library or help identify a target(s) that are associated with a particular disease state.

3 [00163] Implementations to discover potential uses for drugs are contemplated. Often drugs,
4 once approved for clinical use, are limited in terms of the indications (uses) given to them by
5 geographic regulatory bodies. To help determine whether a candidate drug has other potential
6 uses, thereby broadening its indication uses, a drug's usefulness and commercial lifespan can
7 be expanded thereby increasing its commercial and therapeutic value. Based on the drug data
8 collected, tracked and provided to the engine, analyzed data will help provide information to
9 identify other suitable drug targets and therapeutic uses by identifying comparative mechanisms
10 of actions against other existing drugs, and or identifying new target genes and mechanisms of
11 action.

12 [00164] Implementations to assess drug efficacy are contemplated. Drug efficacy must be
13 proven to regulatory bodies in order to achieve market approval. In addition, the stronger the
14 results, the higher the drug's value and potential price point. Assessing drug efficacy is a time-
15 consuming process of drug and clinical study analysis. Based on the drug data collected,
16 tracked and provided to the engine, analyzed data will help provide information to identify other
17 suitable drug targets and therapeutic uses by identifying comparative mechanisms of actions
18 against other existing drugs, and or identifying new target genes and mechanisms of action.

19 [00165] Further, it is contemplated that the engine can be used to analyze variation in a
20 patient's endogenous biomarker and cellular profile and as such detect and monitor various
21 disease states over time, before, during and after treatment for various medical conditions,
22 including Infectious Diseases, Rheumatological Disease, Cardiovascular Disease, Respiratory
23 Disease, Endocrine Disease, Neurological Disease, Nephrological Disease, Urological Disease,
24 Gynecological Disease, Obstetrical Disease, Hematological Disease, Gastroenterological
25 Disease and for Oncology. The engine can be used to analyze a patient's genomic data,
26 genomic expression, proteomic data, protein expression, phenotypic data, metabolomic profile,
27 tissue biopsies, signs & symptoms, and external human and / or animal data inputs including
28 environmental factors, diet, symptoms, demographic information, exercise, etc. Based on the
29 data collected, tracked and provided to the engine, analysis will help provide information to
30 identify risk factors, pathophysiology, appropriate management and prognosis of disease via the
31 monitoring of disease states before, during and following therapeutic intervention.

1 [00166] Further, the engine can be implemented to analyze and interpret normal vs. abnormal
2 patient images collected and/or seen through various imaging and microscopy modalities
3 including but not limited to: Radiography, CT, MRI, PET, Ultrasound, Dermatoscopy,
4 Colonoscopy, Endoscopy, Cystoscopy, Colposcopy, Mammography, Confocal, electron, digital
5 holographic microscopy.

6 [00167] By way of brief review of some of the embodiments described above, the systems and
7 methods described herein thus provide a nested FSM structure or other data structure (i.e. a
8 workflow), including a structure expressed as an interconnection of directed graph structures,
9 where such directed graph structures may be recursive and therefore nested, and where state
10 transitions can be driven by mathematical or logic evaluation, which may or may not result in
11 static or dynamic structural adjustment based on any of: an internal state evaluation, including a
12 nested FSM evaluation that involves a reflective or child calculation, IO connectivity to any
13 hardware device, and/or realtime feedback of system or resource state from components in the
14 system. The structures may be used to represent internal object state attributes or fields, or
15 initialize values of said targets. The structures may be represented via a markup language, an
16 in-memory data structure, or a database or other serialization of said structure. The structures
17 may be used to implement a metadata configuration system which may perform dynamic
18 runtime provisioning, bootstrapping, instantiation, and configuration of system components,
19 including to perform automatic UI generation potentially with a statically or dynamically defined
20 overlay on an automatically generated UI, and to perform element debugging, storing, tracing,
21 and testing.

22 [00168] A system is provided which provides metadata driven bootstrapping of computation
23 and/or storage nodes, a metadata driven cloud orchestration controller which may dynamically
24 adjust a configuration, and a metadata driven dynamic UI generation which may be adjusted via
25 calculations of the system.

26 [00169] A system is also provided that coordinates, enables, or is enabled by the definition of a
27 stream or message processing pipeline where, said pipeline which may be distributed across
28 compute nodes connected by TCP/IP, DMA, PCI, USB, Bluetooth or similar connectivity.

29 [00170] Message based communications of components of the systems may be streamed
30 between the configured nodes for networked distribution to accelerate operations that may or
31 may not leverage systems based on FPGA, ASIC, GPU, Embedded/SoC.

1 [00171] Further message based communications may be streamed between nodes for static,
 2 or dynamic configuration or reconfiguration of a networked distribution which may: improve
 3 system performance or accuracy via active learning (e.g. ground truthing), and dynamic
 4 reconfiguration of capture modes, training modes, and normal modes of classification,
 5 segmentation, and regression; optimize a system configuration by statically or dynamically
 6 adjusting node bootup, shutdown, or interconnection based on node availability, computing
 7 resource availability on nodes (including FPGA, GPU, and CPU state), system fault or network
 8 conditions; and enables coordinated static or dynamic deployment of runtime components by
 9 leveraging server, desktop, mobile, and or SoC/embedded platforms, and direct physical or
 10 networked access to sensors.

11 [00172] Access to sensors may provide detection of or a derivative computation from
 12 electromagnetic sensing spanning the spectrum (e.g. sub-optical (high energy), IR, microwave,
 13 RF) which may involve imagery analysis, interferometric analysis, chemical composition
 14 analysis based, detectable by input from EM detection (such as sensors having sensing modes
 15 including optical, spectroscopy, NEMS, MEMS and electronic noise). Sensors include but are
 16 not limited to any combination of the following: optical, biometrics, X-ray, MEMS, NEMS,
 17 interferometry that is based on NEMS or MEMS sensor systems including microfluidics (e.g.,
 18 saliva analysis), LoC, SoC, mass spectrometry, next-generation sequencing methods [genome
 19 sequencing, genome resequencing, transcriptome profiling (RNA-Seq), DNA-protein
 20 interactions (ChIP-sequencing), and epigenome characterization, massively parallel signature
 21 sequencing (or MPSS), polony sequencing, pyrosequencing, Illumina dye sequencing, SOLiD
 22 sequencing (sequencing by ligation), Ion semiconductor sequencing, DNA nanoball sequencing,
 23 Heliscope single molecule sequencing, single molecule real time (SMRT) sequencing, nanopore
 24 DNA sequencing, tunneling currents DNA sequencing, sequencing by hybridization, sequencing
 25 with mass spectrometry (Matrix-assisted laser desorption ionization time-of-flight mass
 26 spectrometry, or MALDI-TOF MSa), microfluidic Sanger sequencing, RNAP sequencing, In vitro
 27 virus high-throughput sequencing]; Radiological: CT, MRI, X-ray, PET, endoscopy,
 28 bronchoscopy; tissue biopsy, dermatoscopy, microscopy, fundoscopy, colposcopy, cystoscopy,
 29 optical coherence tomography.

30 [00173] Imaging and image capture may comprise video recorders and cameras, lighting
 31 equipment, thermal imagers, X-ray.

1 [00174] A tracing system is provided that intercepts stream or messaging connections between
2 components that may be integrated in a fashion that obfuscates or abstracts network
3 distribution; that may record and playback of internode communication for the purposes of
4 simulation, development activity or operations testing and/or deployment, and debugging and/or
5 monitoring; that may have integration with developer operations, testing, and monitoring
6 processes or operations; and may integrate with any of the above systems to create a feedback
7 loop that affects any aspect of system performance via links that are statically coded, or driven
8 by further metadata definition and/or any system calculation.

9 [00175] The system may further be implemented for real time and dynamic output to augment
10 or replace human inspection and quality control operations applicable to manufacturing
11 verification in auto, aviation, food, and pharmaceutical inspections, and for travel and border
12 transit auditing and verification including border or airport screening for contraband or
13 agricultural restricted material, including but not limited to fungal contamination. Systems can be
14 used to analyze a stationary or moving target, optionally providing analysis while a stationary or
15 moving target is exposed to materials in motion, and may normalize imagery input in realtime.

16 [00176] The systems may further be implemented to attempt to accomplish objectives
17 including: compensation for differences between processing pipeline requirements for imagery
18 spanning a micro or nano resolution cell size; realtime control of sample collection and
19 preparation of a mechatronic or robotics based system; and evolution of material specification or
20 production of any functionalized materials on a micro or nano surface.

21 [00177] The systems may further be used for pathogen detection, prediction of medical
22 diagnostic outcomes where the system may integrate with paper or electronic collections of
23 patient information, driving background and foreground separation of genetic material for a
24 required analysis. The systems may be used for applying imagery segmentation techniques to
25 genetic data, and for using feature definition and extraction based on computer vision derived
26 methods. The systems may drive static or dynamically improving correlations/predictions
27 between protein markers, genetic information, and biological outcomes, and may be used in
28 applications that augments or replaces protein interaction analysis with genetic analysis.

29 [00178] The systems may further be used for antimicrobial resistance ("AMR"), drug testing,
30 animal testing, food safety testing, predicting biological outcomes for drug resistance, drug or
31 vaccine effectiveness and pathogen evolution.

[00179] When the system is used for AMR, an exemplary system is one: where sensors are used with a patient data collection system with a classification system; where the classification system is augmented with a database; where the database contains information on class (classification) labels associated with data samples; where the database is updated periodically with new information content and/or corrections to previous content; where the database is derived from a collection of information that includes journal publications that have been text mined; where a "doctor in the loop" workflow model is implemented and all automated text mined results are presented to a medical healthcare professional for approval to update the database with the mined content; where all content updates trigger a retrospective review to determine which patient treatment programs are affected; where the database has a list of patients and their primary care providers (PCP); where a notification service is provided such that when a patient's healthcare program is impacted by new data updates; where all "doctor in the loop" activities are logged in a manner suitable for auditing; where the database contains information about antimicrobial resistance; where the database additionally contains information about patient susceptibility to different diseases/pathogens; where the database additionally contains information on which treatments are likely to be most effective for the diagnosed pathogen given information relative to treatment efficacy on a patient cohort that is similar to the patient from which the data sample was obtained; where the patient similarity to the patient cohort includes genetic similarity; where the data samples contain information about pathogen exposure and/or presence; where the data samples contain other health-related information; where the data samples collected can be processed to determine a patient's healthcare profile and likelihood that a certain treatment program is optimal with respect to health outcomes; where a peer to peer network exists that includes health research organizations and point of care (POC) locations; where the POC stations are field stations in pathogen outbreak regions; where data is anonymized and aggregated, data processed and health alerts are issued to human mass transit control points; where real-time alerts are issued for limiting the spread of persons with highly contagious diseases; where the real-time aspect of the alert is facilitated by a resource orchestration agent such that clinical decision support is enabled with ad hoc, heterogeneous networks of compute resources; where the alert is for an individual or a disease susceptible group of persons; where the diagnosis is considered individualized due to the analysis of genomics information derived from patient samples; where the data transfer and alert notification conform to health regulations such as privacy; where the analysis system may consist of a deep net; where use of the described system with data which has associated

1 metadata in the form of a jurisdictional certificate; where the jurisdictional certificate is used to
2 determine access to and use of the data element by downstream compute nodes or
3 associations of end users, whether use or access it at the sample level or a higher aggregate
4 granularity level; where the jurisdictional certificate is associated with a micro-payment
5 mechanism; where the jurisdictional certificate can be authenticated by independent third
6 parties in point of care or point of service (POS) distributed locations e.g. airports; where the
7 jurisdictional certificate is combined with the encryption technique to allow only authorized use
8 of the associated data; where the jurisdictional certificate is associated with a notification list for
9 various organizations. e.g. airport alert about a local pathogen outbreak; where the jurisdictional
10 certificate is associated with various combinations of research, academic, governmental, non-
11 governmental and/or corporate organizations; where the jurisdictional certificate is associated
12 with specified geographic areas or political boundaries; where the jurisdictional certificate is
13 associated with different disease / pathogens and health care policies; where the jurisdictional
14 certificate is authenticated at mass transit locations such as an airport, shipyard or train station;
15 where the jurisdictional certificate allows for negotiated access and use of the data element (or
16 aggregated data values) based on a pre-defined negotiation strategy/cost per data formula.

17 [00180] The above systems and implementations may be provided with hardware acceleration
18 via embedded systems or SoC, FPGA, ASIC, GPU acceleration, other hardware systems.

19 [00181] The systems may provide improved accuracy of the described functions via active
20 learning, ground truthing, and dynamic reconfiguration of capture, train, and normal mode of
21 classification, segmentation, and regression dynamic reconfiguration of capture, train, and
22 normal mode of classification, segmentation, and regression

23 The above-described implementations are intended to be examples and alterations and
24 modifications may be effected thereto, by those of skill in the art, without departing from the
25 scope which is defined solely by the claims appended hereto. For example, methods, systems
26 and embodiments discussed can be varied and combined, in full or in part.

CLAIMS

We claim:

1. A system for applying a deep learning neural network to data obtained from one or more sensors on a client computing system, the one or more sensors sensing components of an industrial scene, the system comprising:

a network accessible server system linked to the client computing system, the server system configured to provision and send an application to be executable on the client computing system, the application allocating computing tasks to the server system, the server system provisioning the application based on a schema definition comprising metadata received from the client computing system such that the application is customized for the client computing system and causes interaction with the sensors particular to the client computing system; and

one or more peripheral devices, at least one of the peripheral devices comprising one or more of the sensors, each peripheral device associated with the scene, each peripheral device connected to the client computing system, the client computing system transmitting the sensor data to the server system, the server system building a one or more trained classification models to extract one or more features using the sensor data sensed by the sensors connected to the client computing system, the one or more classification models applying a neural network based scene parsing and scene matching to a portion of the sensor data transmitted from the client computing system to parse and match the one or more extracted features as feature vectors, training of the one or more classification models considering only sensor data from the one or more sensors connected to the client computing system, the client computing system parsing and matching one or more features from other portions of the sensor data transmitted from the client computing system based on feature vectors computed using the one or more trained classification models received from the server system, the matching using previously known components, the client computing system controlling positioning, motion, or obtainment of sensor information of at least one of the peripheral

devices based on the output of the matching from the one or more trained classification models.

2. The system of claim 1, wherein the sensors comprise image capture devices.
3. The system of claim 1, wherein the sensor data is transmitted by the client computing system to the server system through interfaces that provide anonymization, encryption, compression and jurisdictional stamping.
4. The system of claim 1, wherein the sensor data comprises a plurality of data points aggregated from a plurality of sensors.
5. The system of claim 1, wherein results of the classification process are transmitted to a ground truther to evaluate model accuracy to implement an active learning process.
6. The system of claim 5, wherein the system further provides a confidence score to the ground truther.
7. The system of claim 1, wherein the server system applies a plurality of neural networks simultaneously and provides a list of classifiers.
8. The system of claim 7, wherein the classifiers in the list are provided with confidence values.
9. A method of applying a deep learning neural network to data obtained from one or more peripheral devices connected to a client computing system, at least one of the peripherals devices comprising one or more sensors, the one or more sensors sensing components of an industrial scene, each peripheral device associated with the scene, the method comprising:

linking a network accessible server system to the client computing system;
configuring the server system to provision and send an application to be executable on the client computing system, the application allocating computing tasks to the server system, the server system provisioning the application based on a schema definition comprising metadata received from the client computing system such that the application is customized for the client computing system and causes interaction with the sensors particular to the client computing system;

receiving sensor data from the one or more peripheral devices connected to the client computing system at the server system;

building one or more trained classification models at the server system to extract one or more features using the sensor data sensed by the sensors, the one or more classification models applying a neural network based scene parsing and scene matching to a portion of the sensor data transmitted from the client computing system to parse and match the one or more extracted features as feature vectors, training of the one or more classification models considering only the sensor data from the one or more sensors connected to the client computing system;

transmitting the one or more trained classification models to the client computing system to be usable at the client computing system to parse and match one or more features from other portions of the sensor data based on feature vectors computed using the one or more trained classification models, the matching using previously known components; and

outputting the matching from the one or more trained classification models to be used by the client computing system to control positioning, motion, or obtainment of sensor information of at least one of the peripheral devices based on the output of the matching.

10. The method of claim 9, wherein the sensors comprise image capture devices.
11. The method of claim 9, wherein the sensor data is transmitted by the client computing system to the server system through interfaces that provide anonymization, encryption, compression and jurisdictional stamping.
12. The method of claim 9, wherein the sensor data comprises a plurality of data points aggregated from a plurality of sensors.
13. The method of claim 9, wherein results of the classification process are transmitted to a ground truther to evaluate model accuracy to implement an active learning process.
14. The method of claim 13, further comprising providing a confidence score to the ground truther.
15. The method of claim 13, wherein the server system applies a plurality of neural networks

simultaneously and provides a list of classifiers.

16. The method of claim 15, wherein the classifiers in the list are provided with confidence values.

1/35

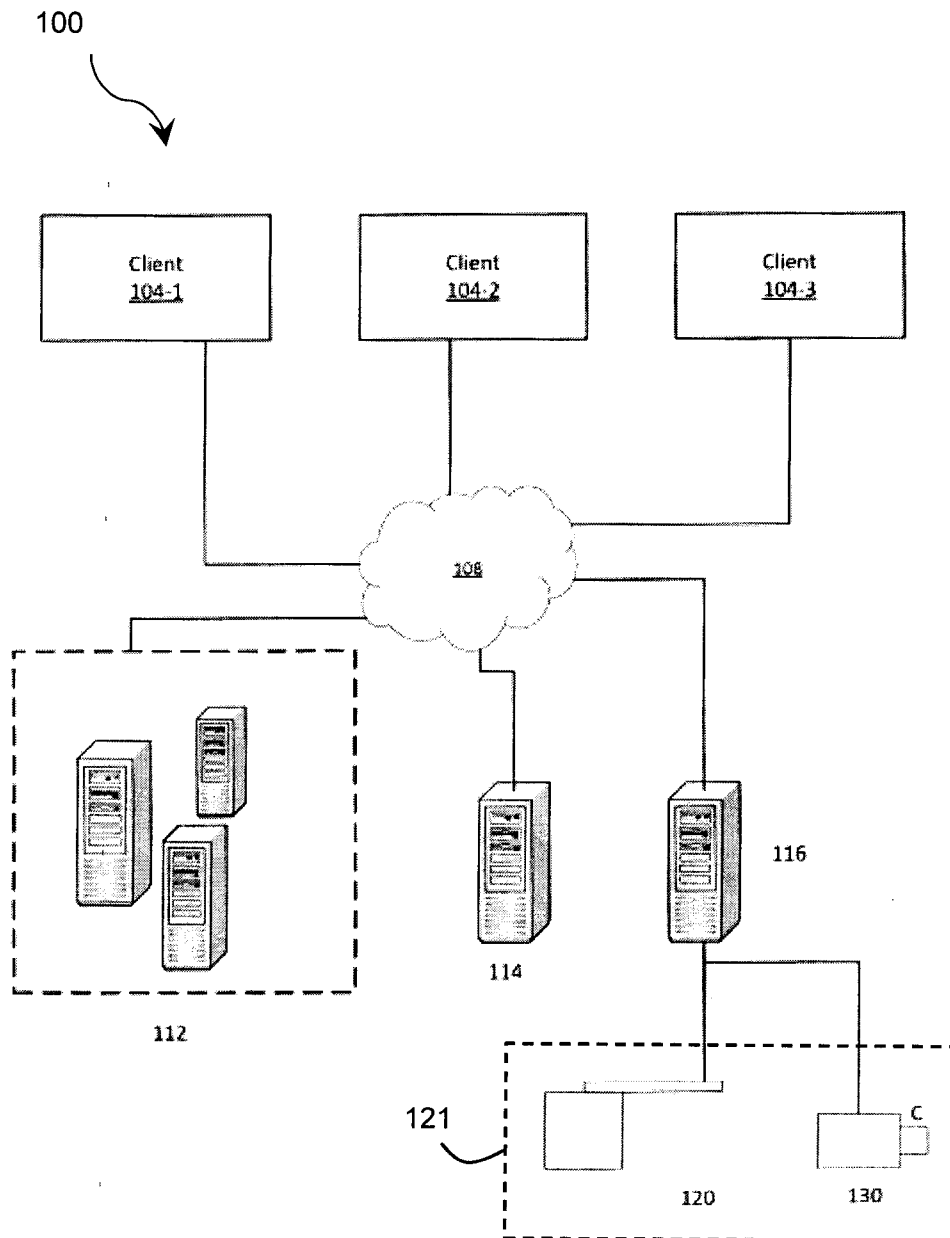
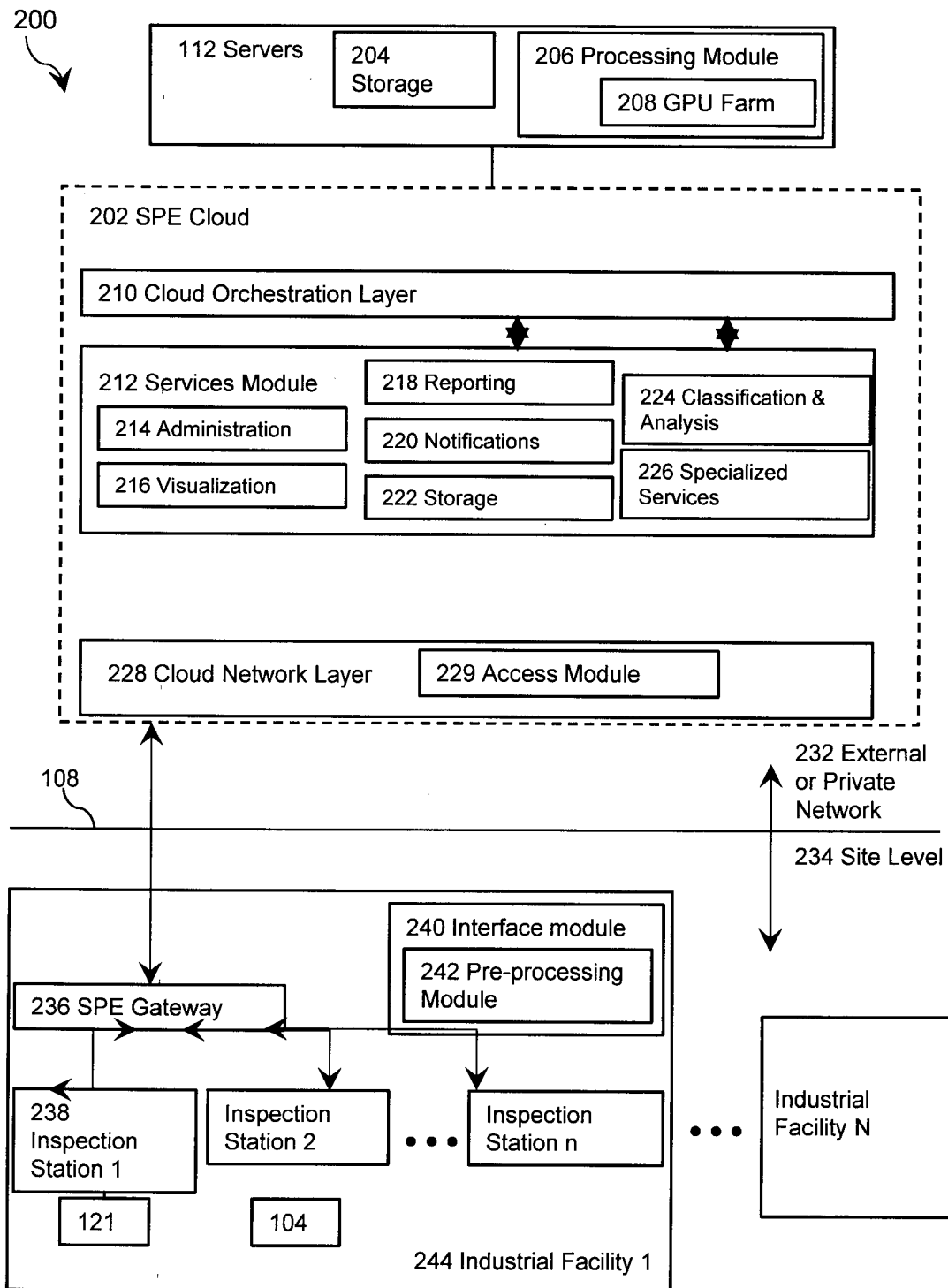
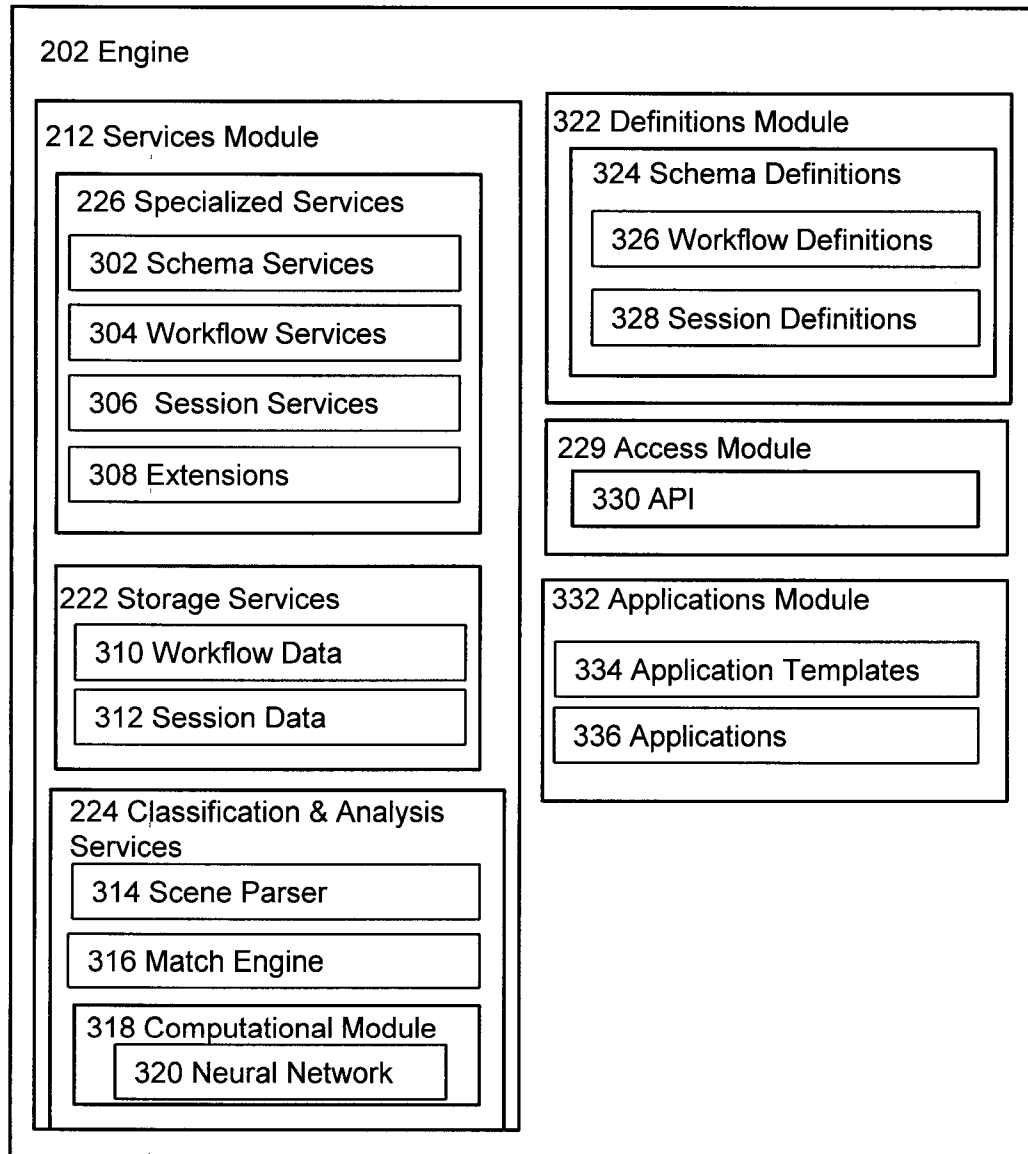


FIG. 1

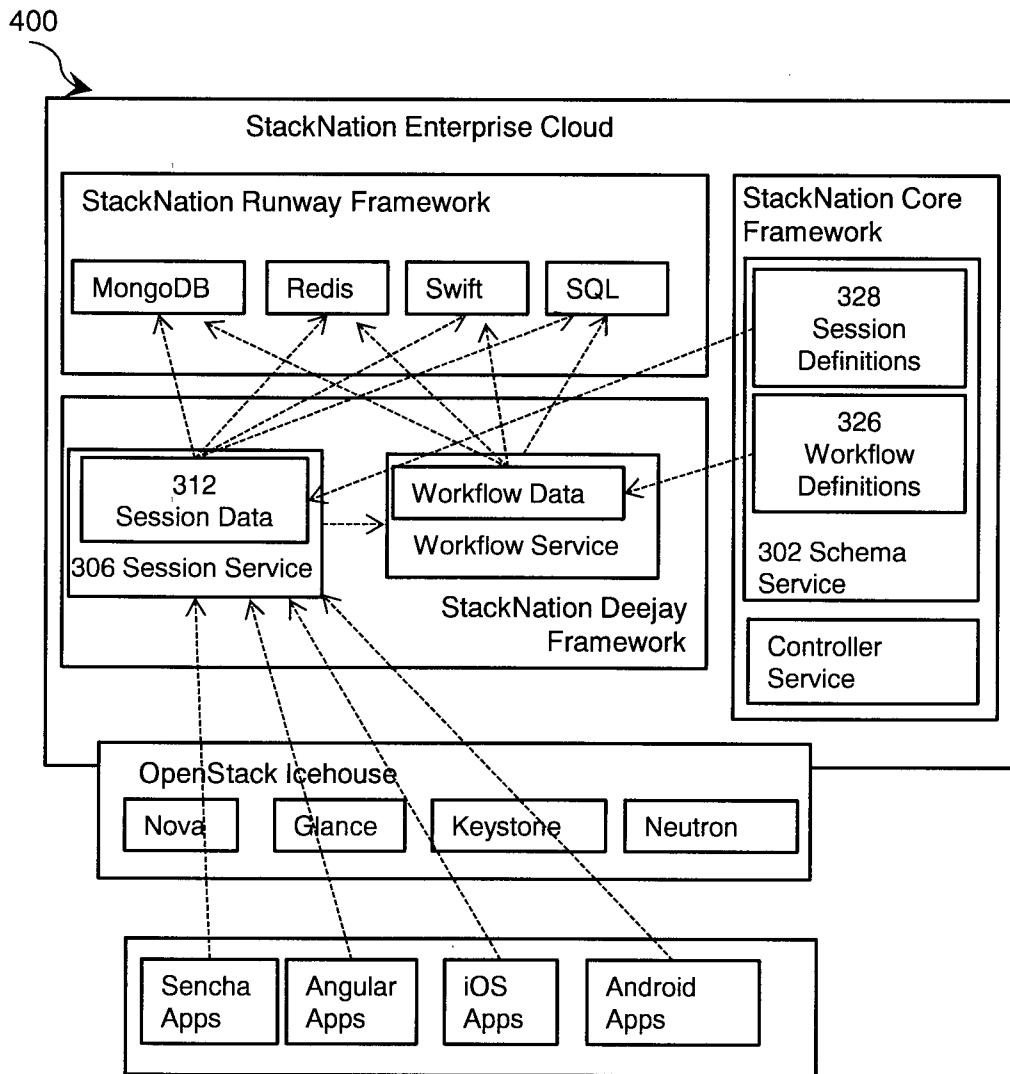
2 / 35

**FIG. 2**

3/35

**FIG. 3**

4/35

**FIG. 4**

5/35

400

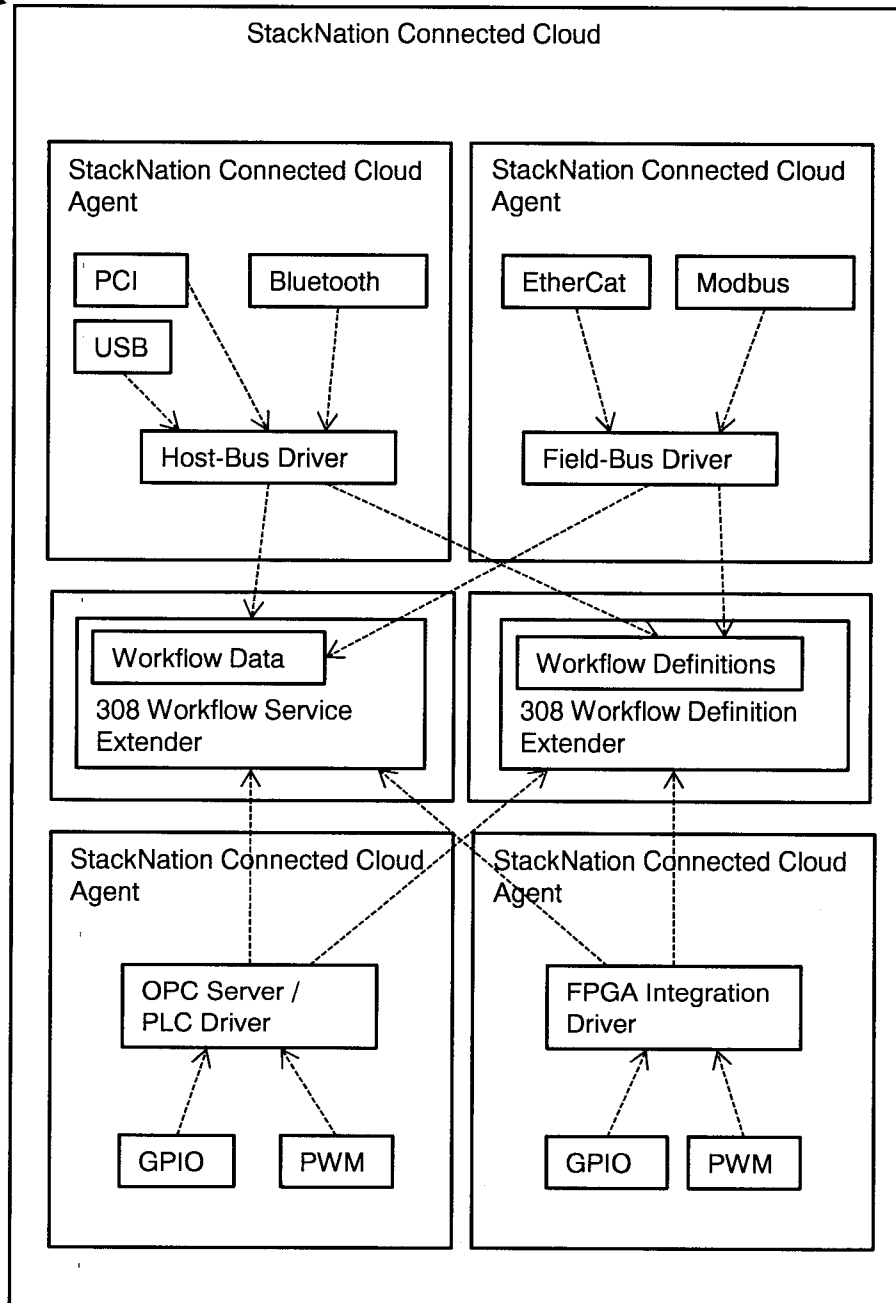


FIG. 5

6/35

400

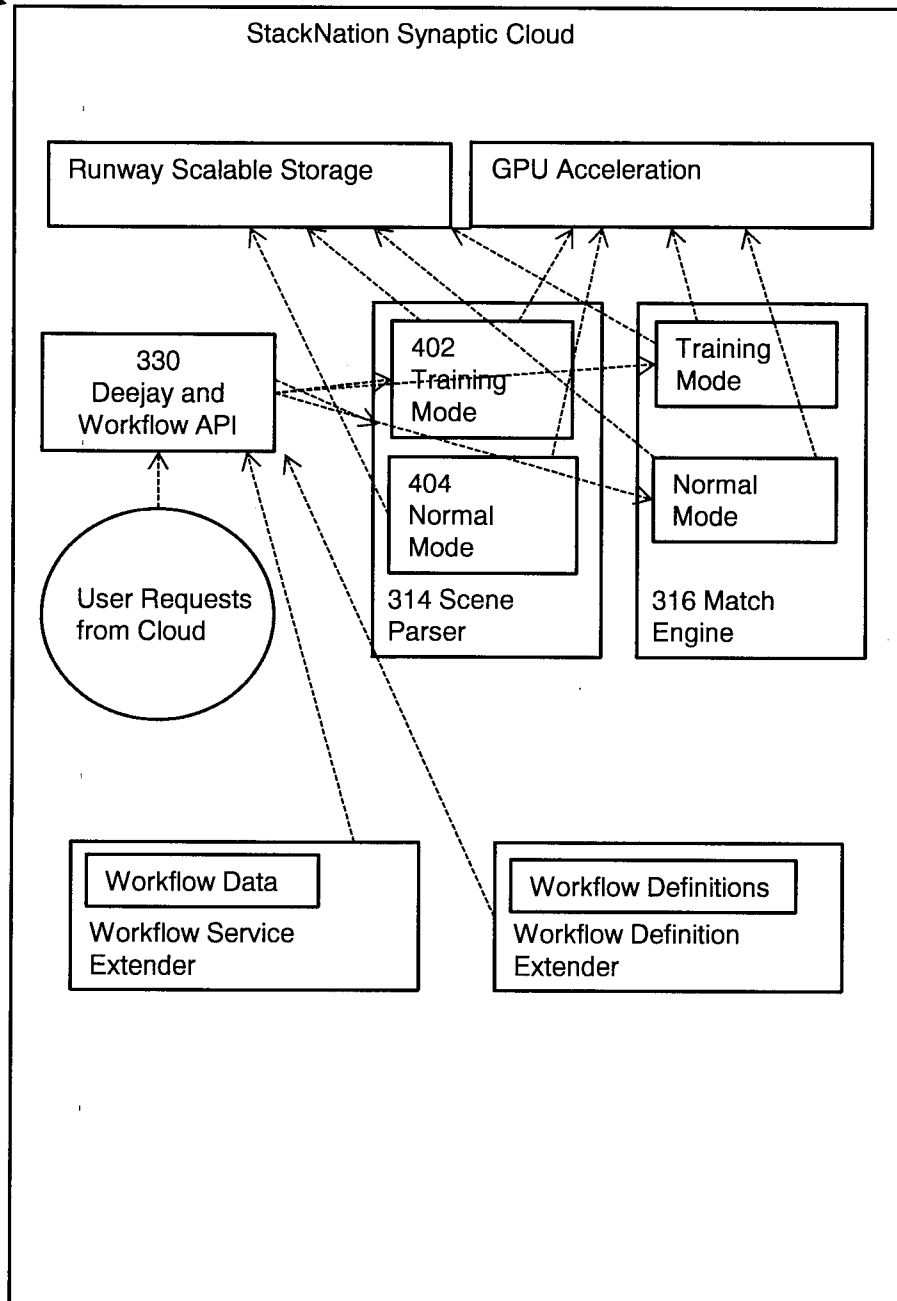


FIG. 6

7/35

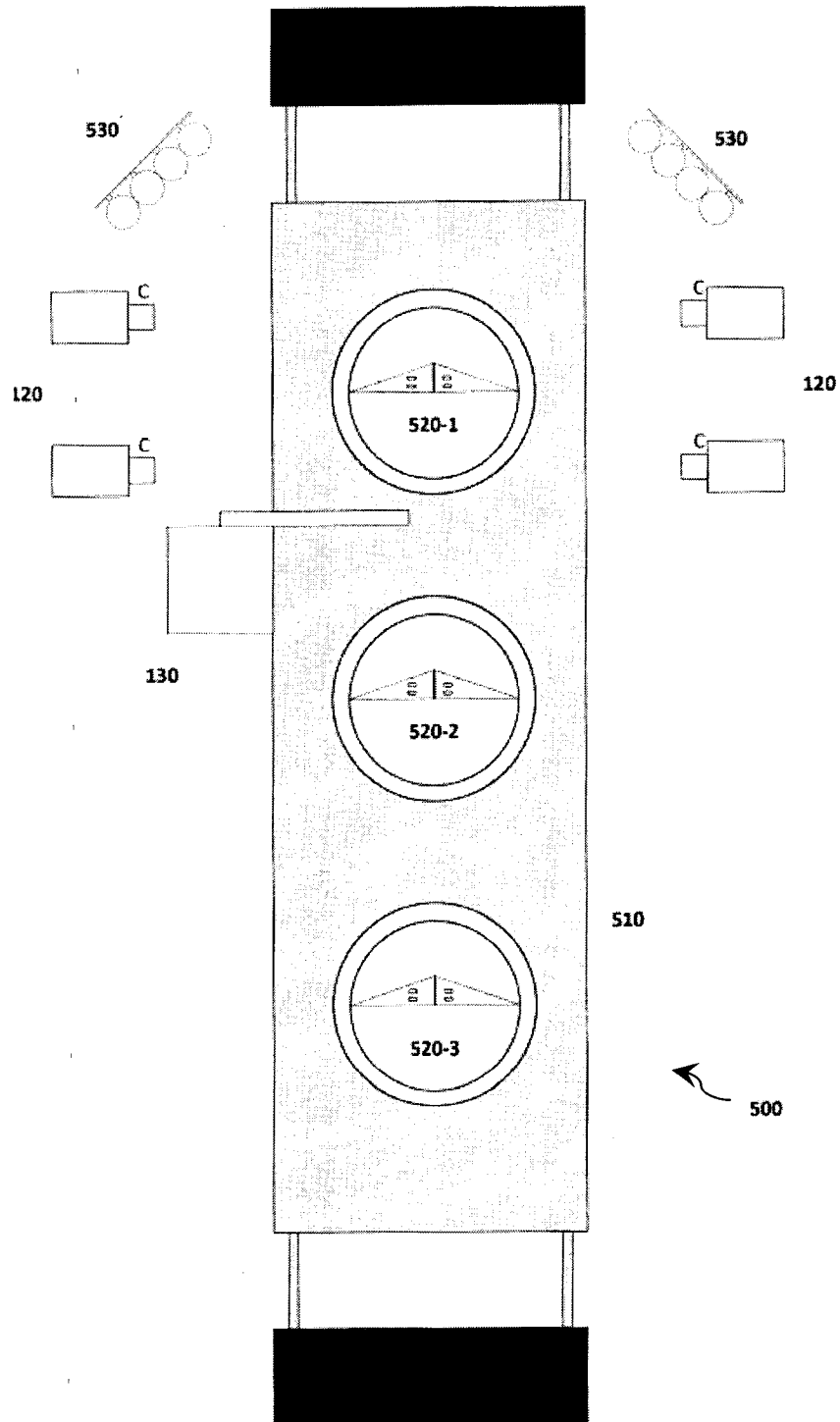


FIG. 7

8/35

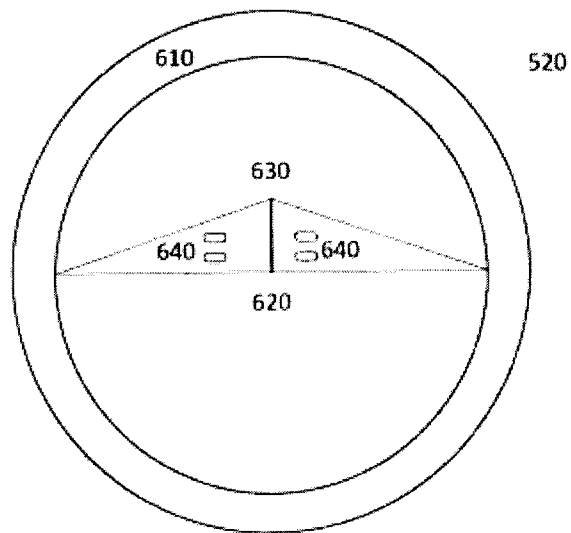


FIG. 8

9/35

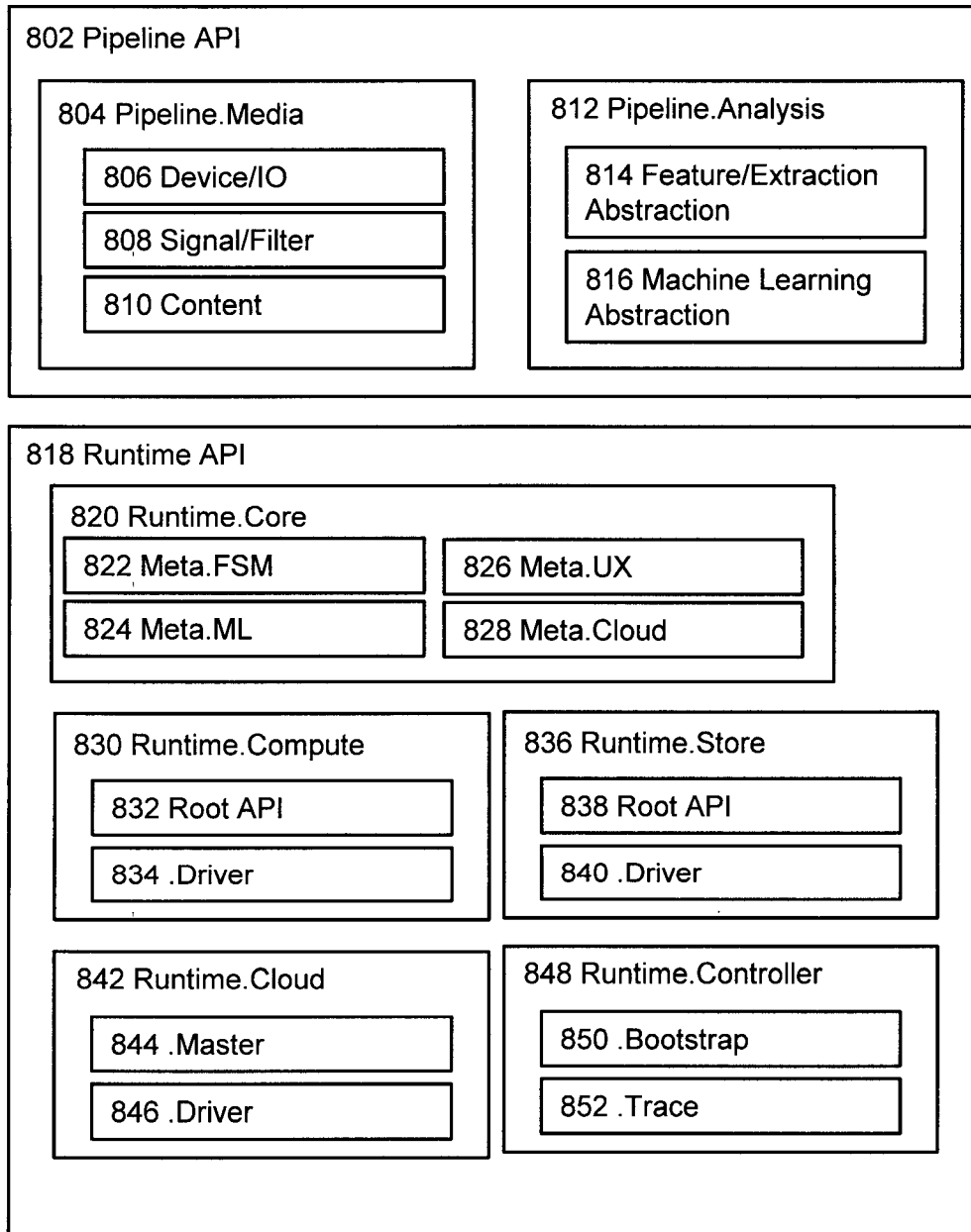
800 

FIG. 9

10/35

1000 

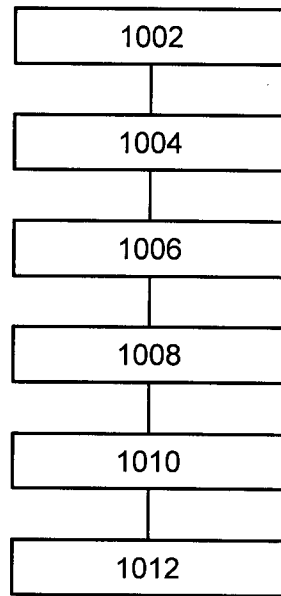
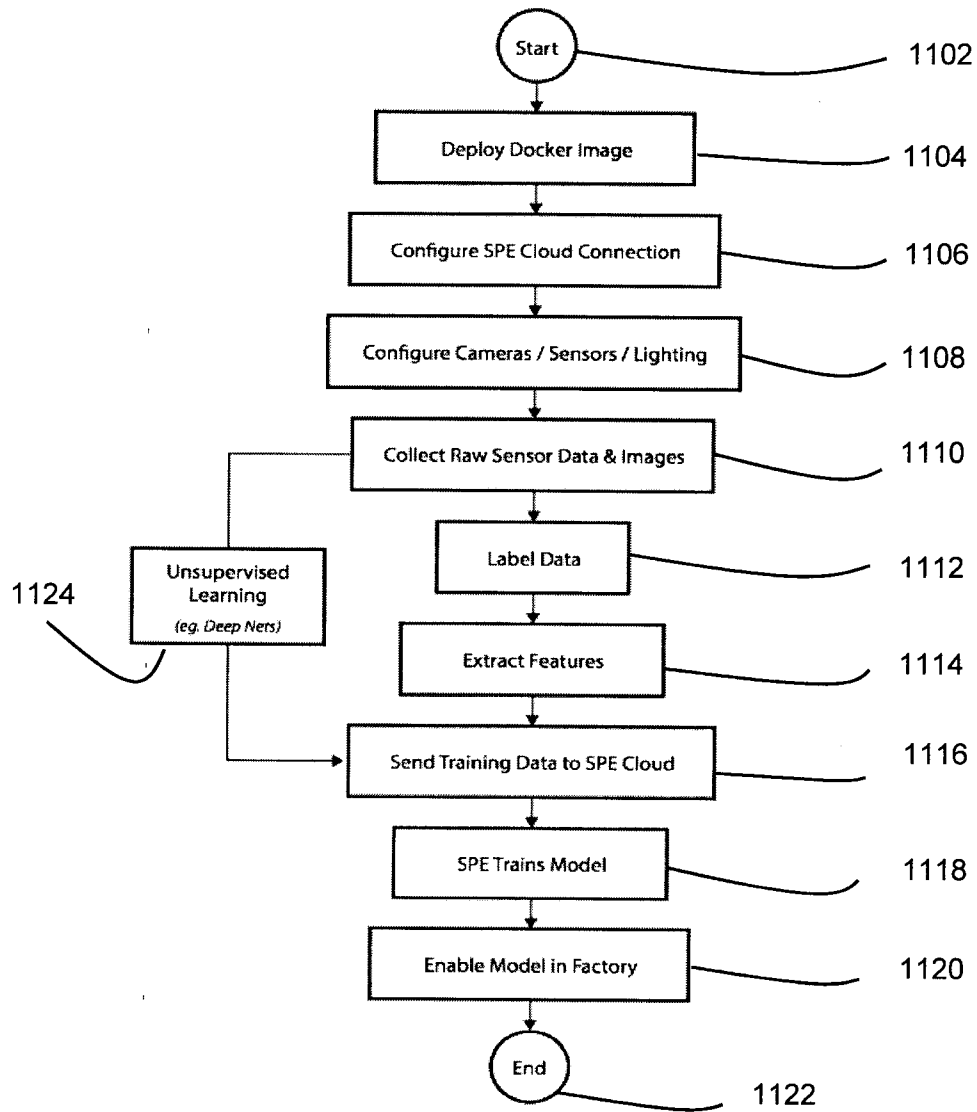


FIG. 10

11/35

Inspection Station Setup

**FIG. 11**

12/35

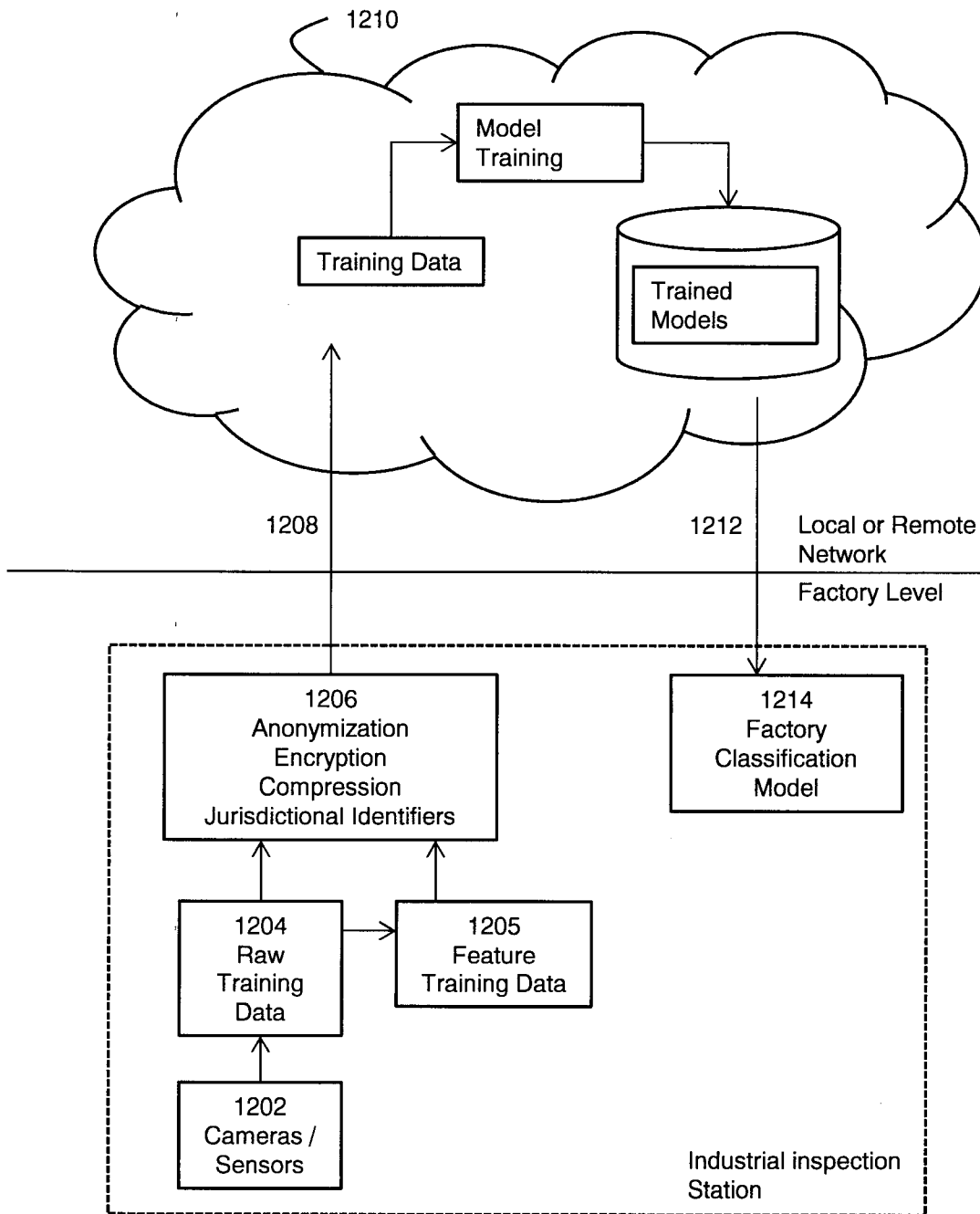
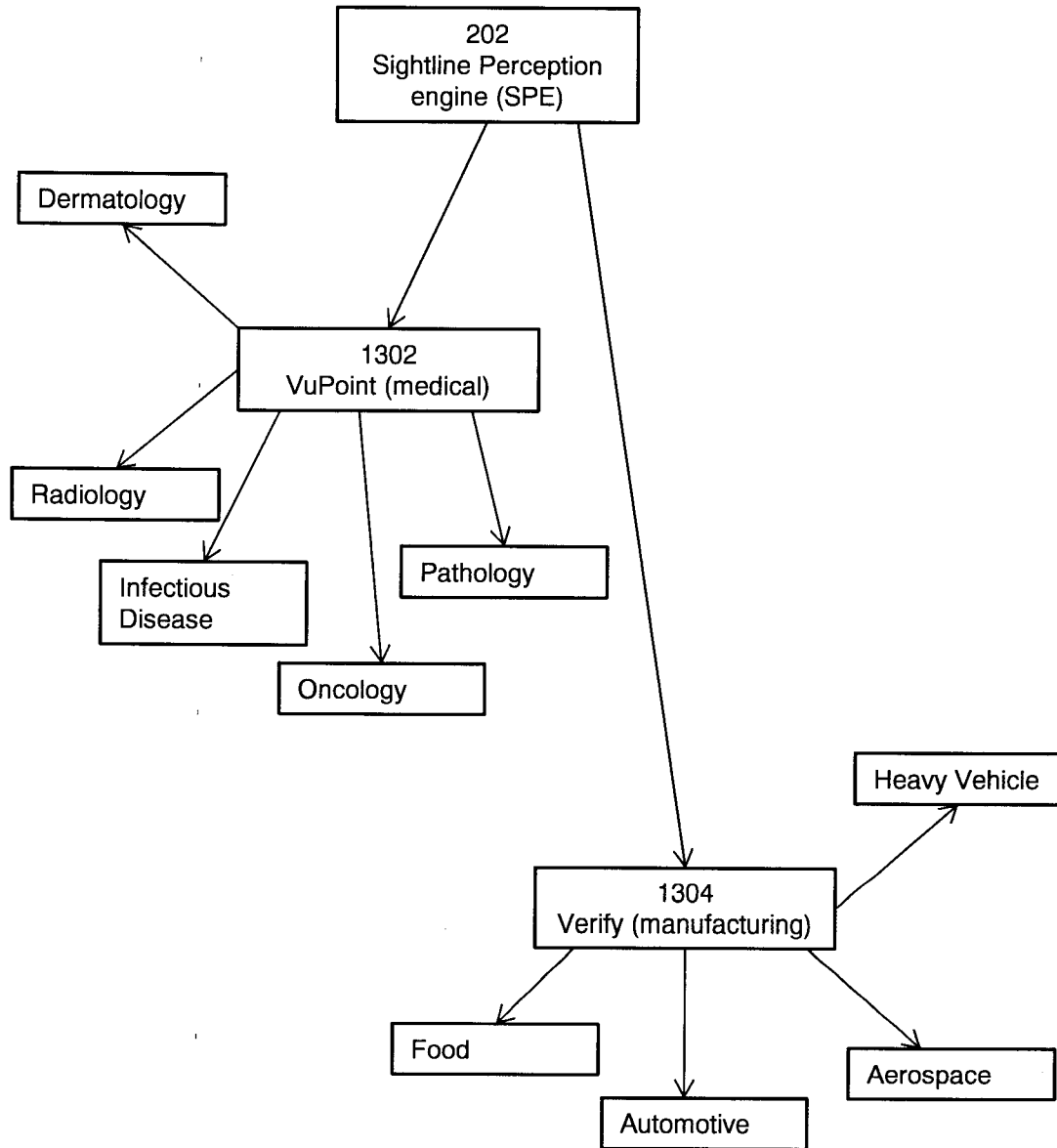
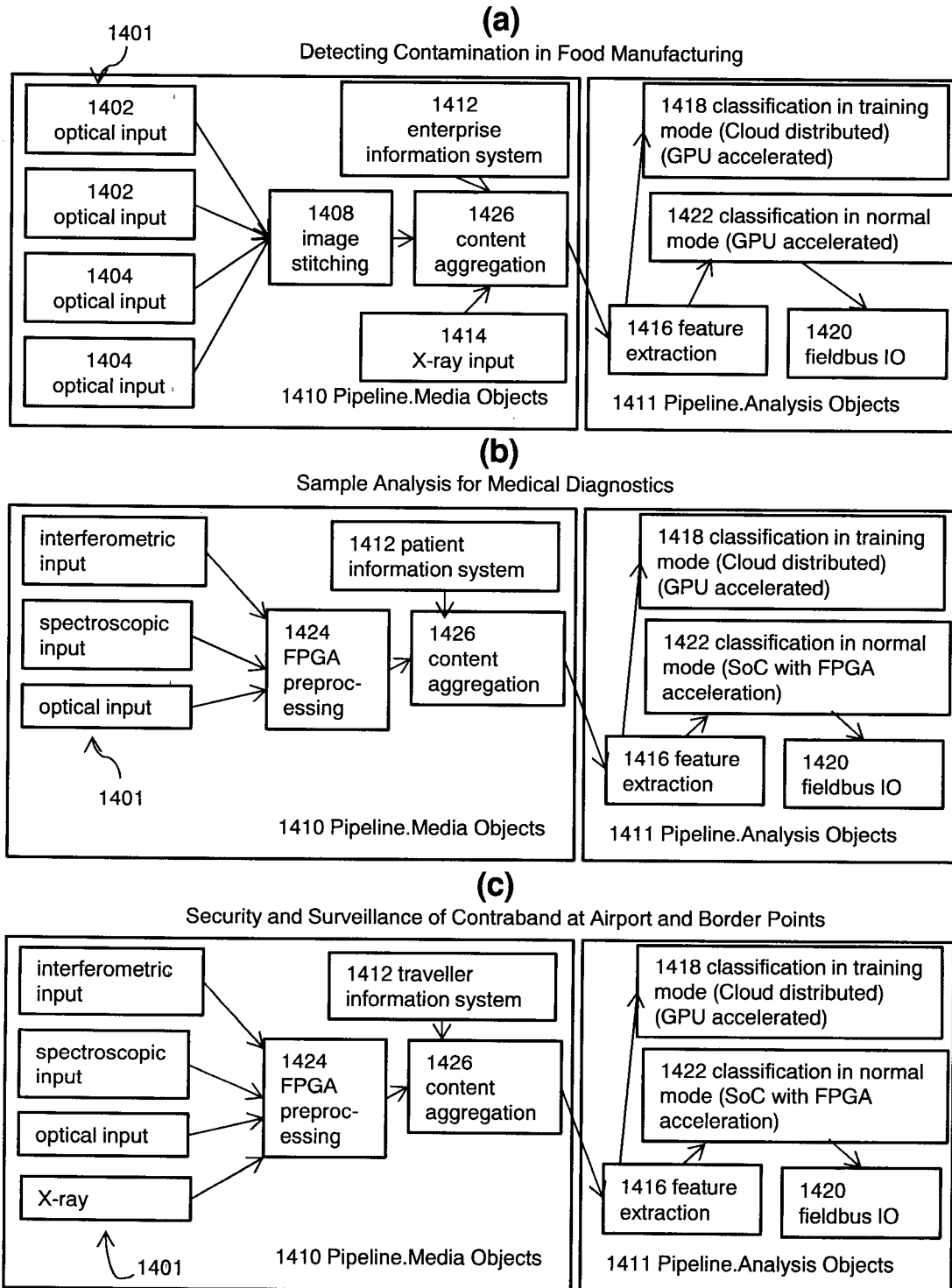


FIG. 12

13/35

**FIG. 13**

14/35

**FIG. 14**

15/35

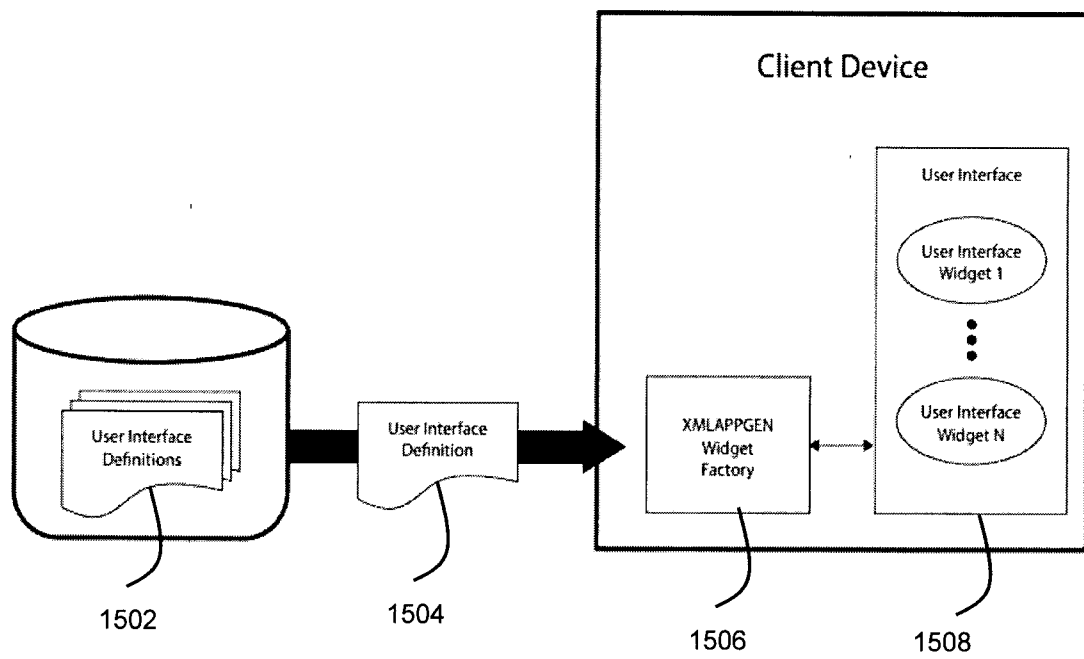
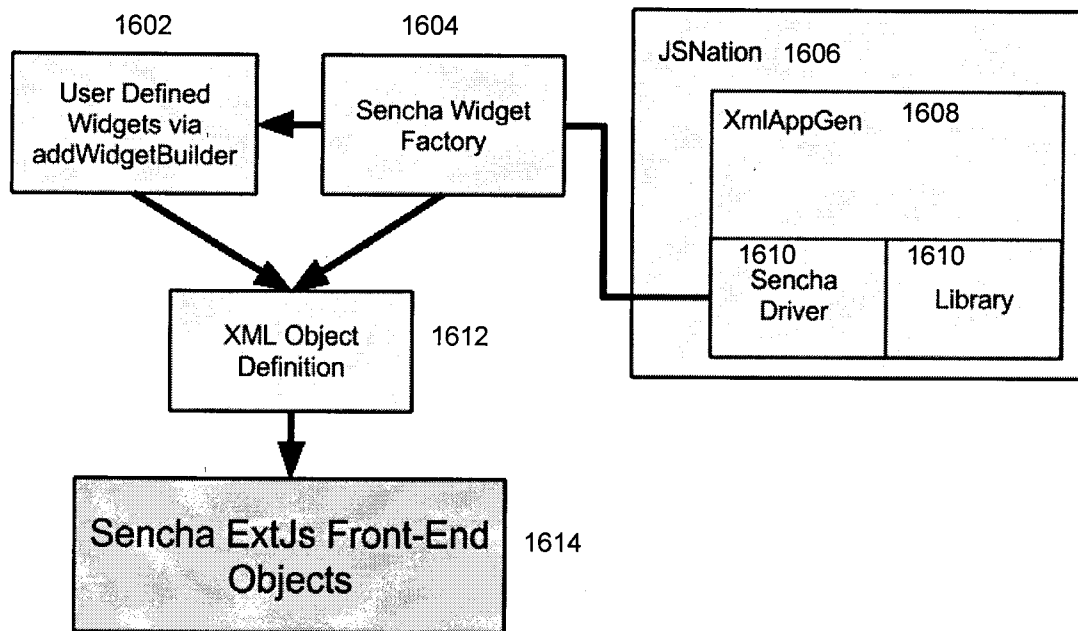


FIG. 15

16/35

**FIG. 16**

17/35

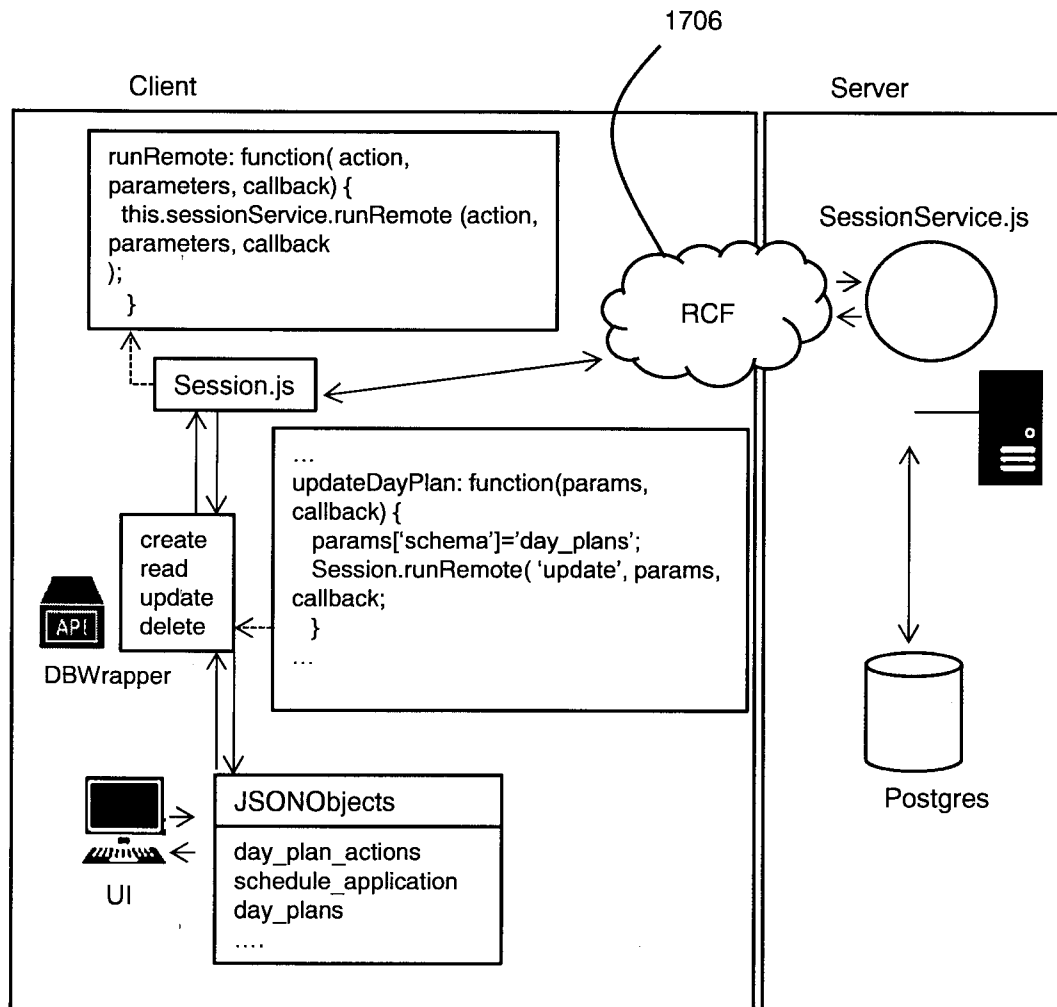


FIG. 17

18/35

1800 

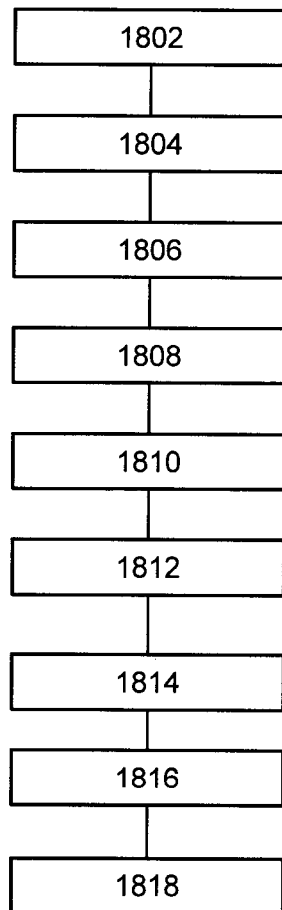


FIG. 18

19/35

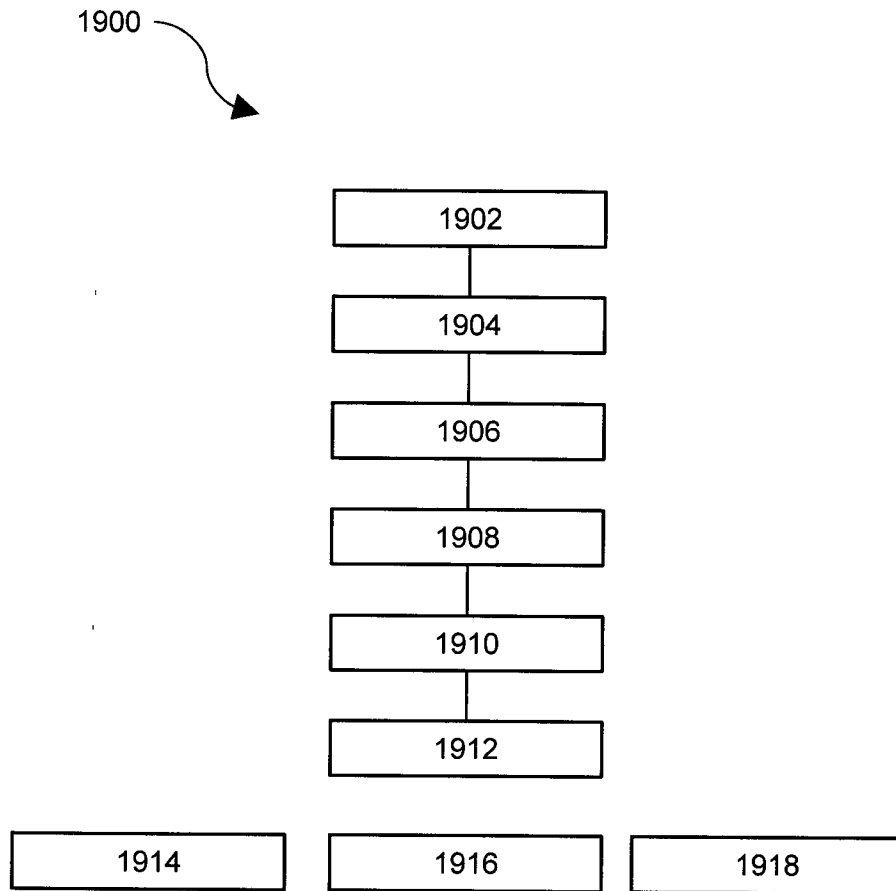


FIG. 19

20/35

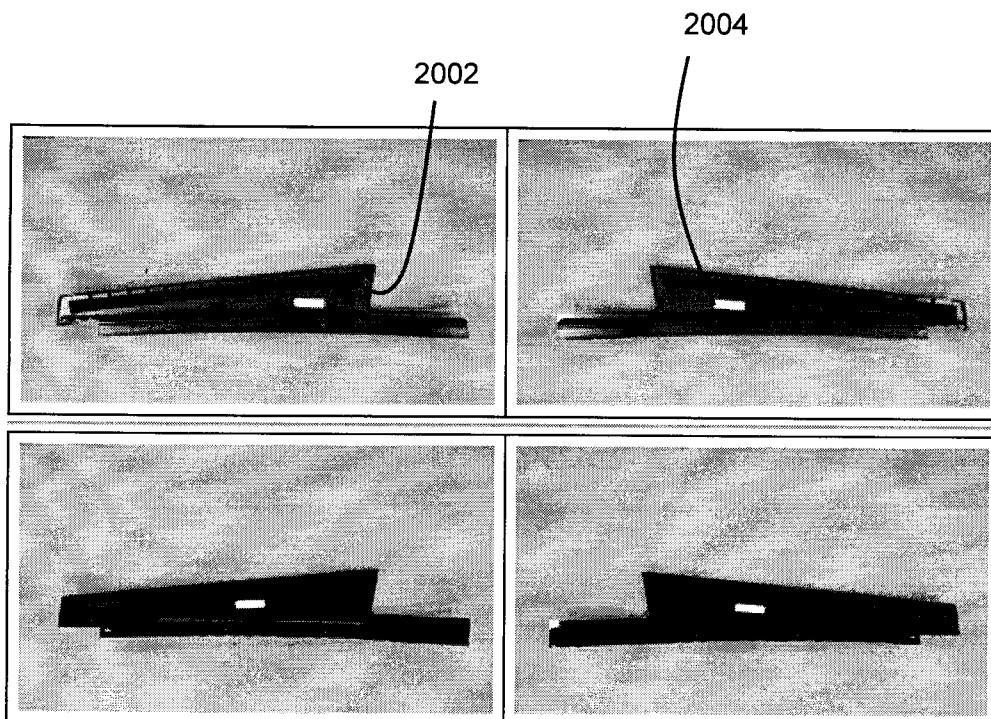


FIG. 20

21 /35

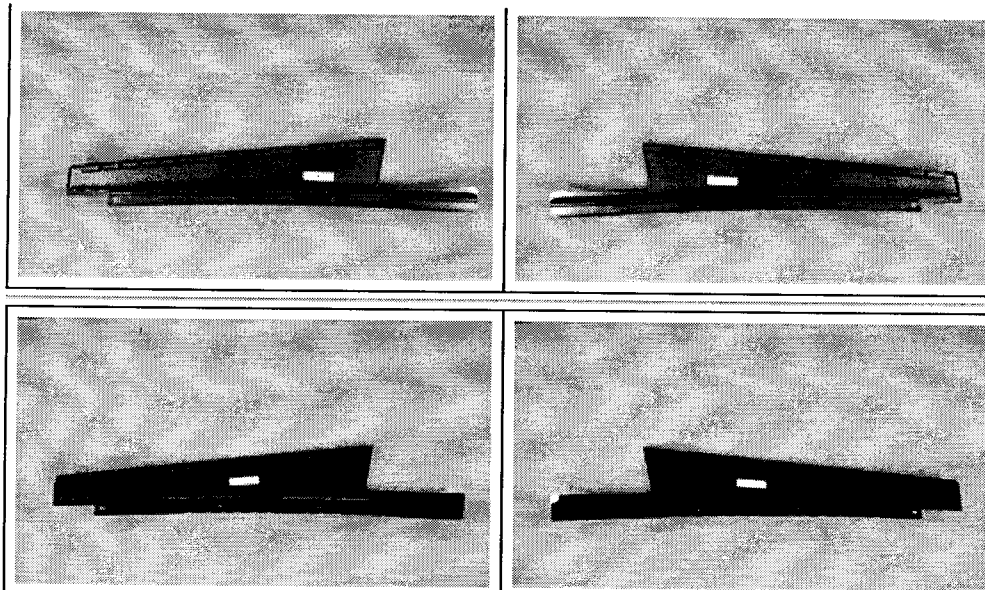
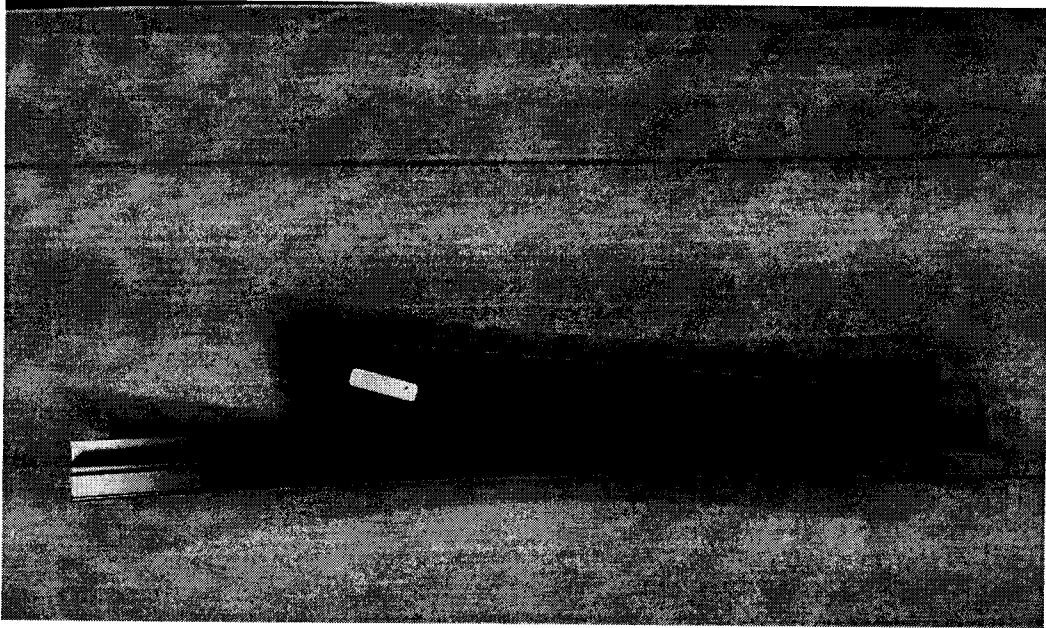
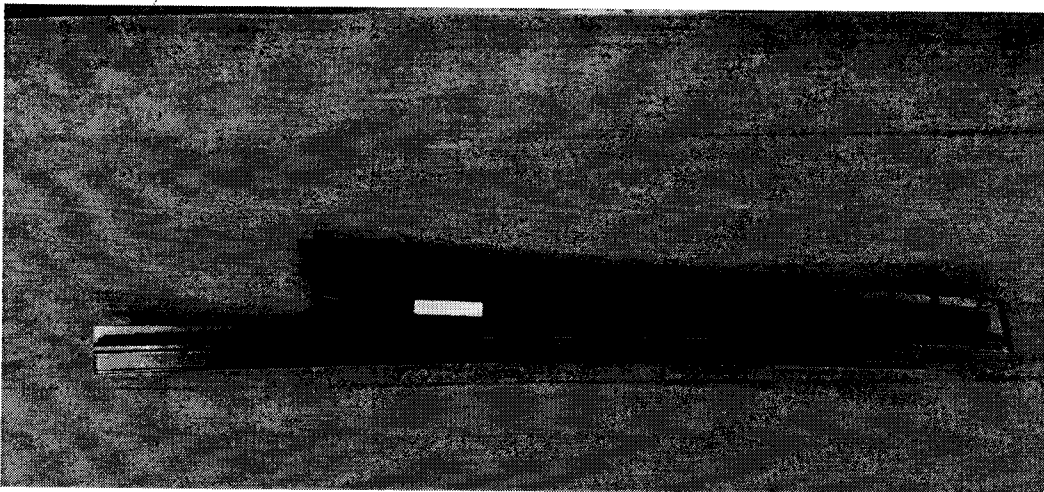


FIG. 21

22/35



(a)



(b)

FIG. 22

23/35

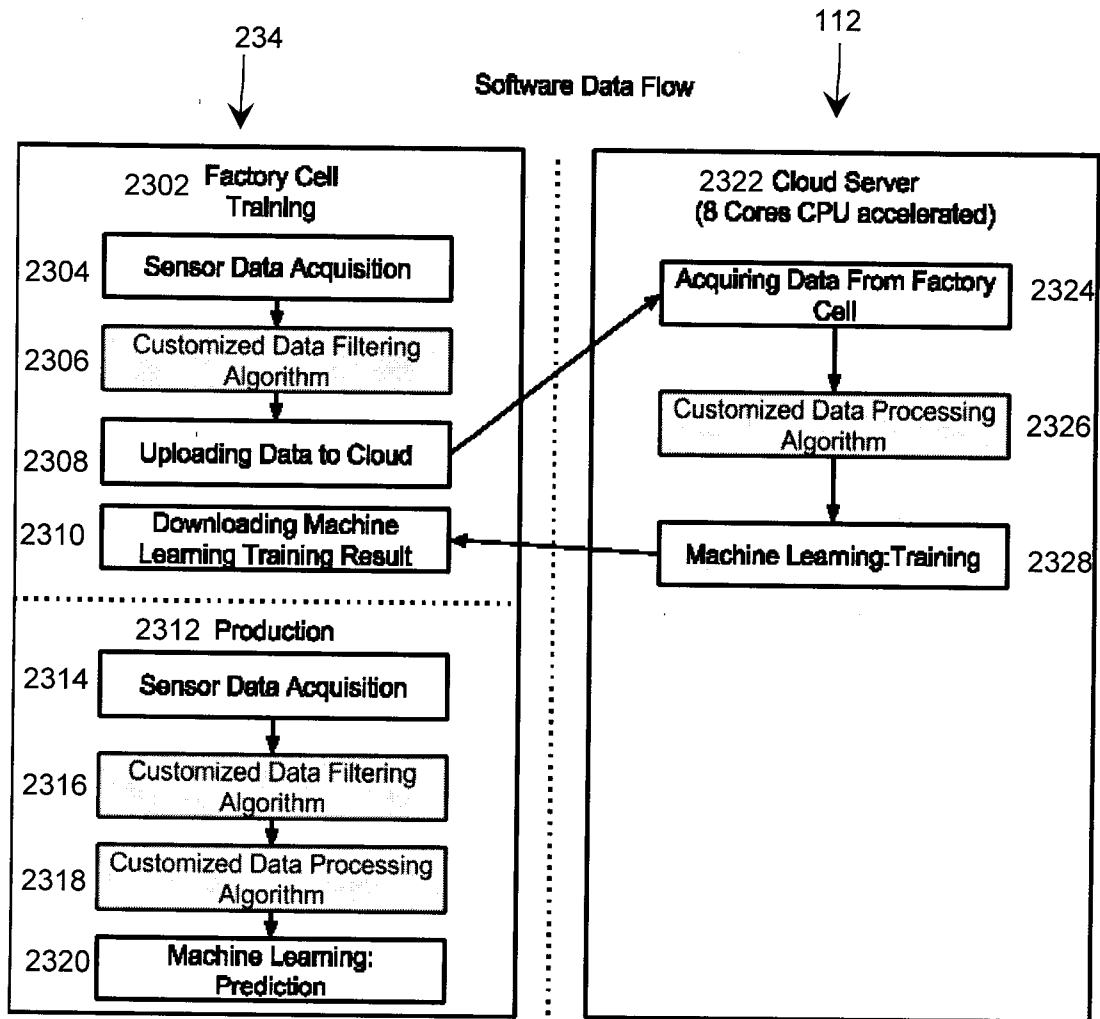


FIG. 23

24/35

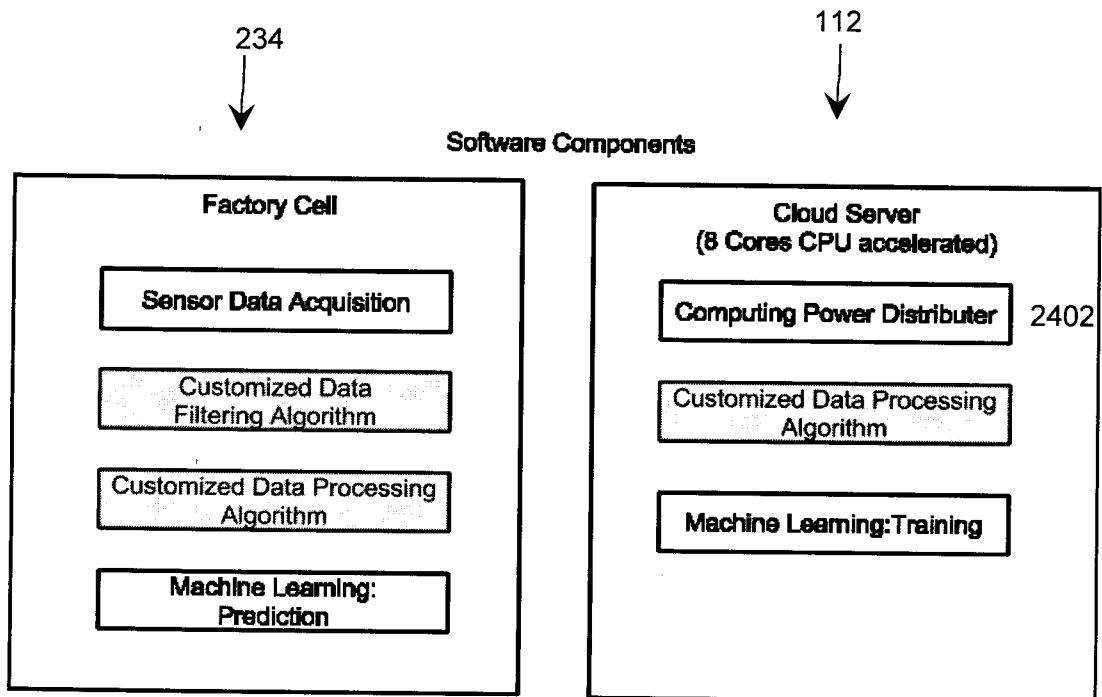
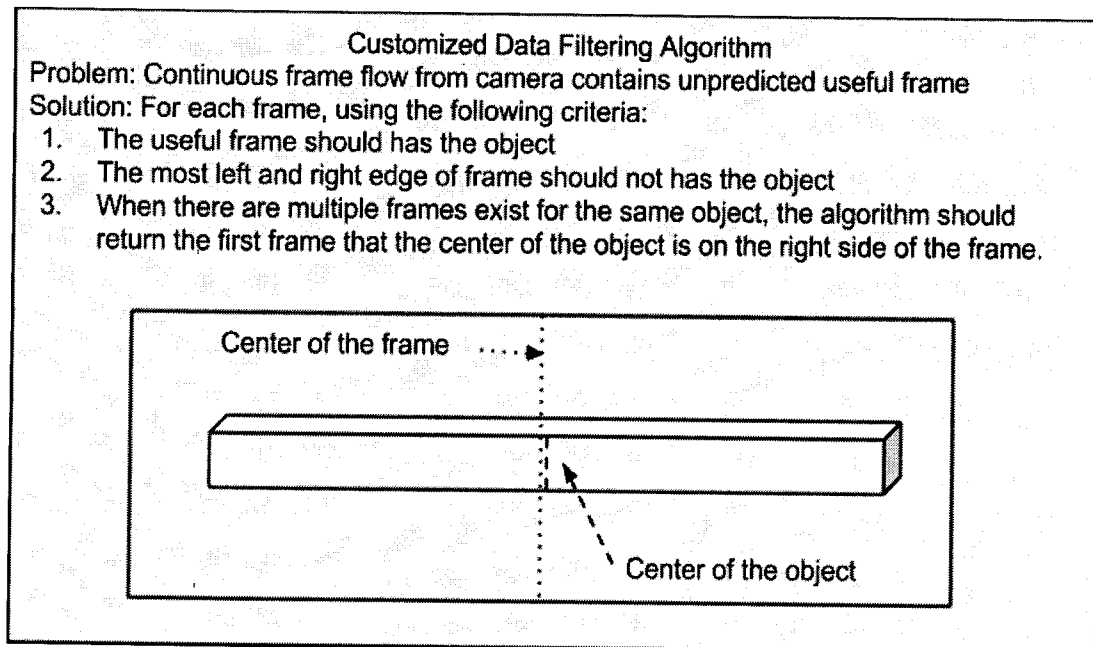


FIG. 24

25/35

**FIG. 25**

26 /35

	Customized Data Processing Algorithm (for Object Classification)	
	Problem:	
	Distinguish different objects, using the frames which are not distinguishable by human eye	
	Solution:	
	Preparing the algorithm:	
2602	1.	Mask the object from background (Otsu Thresholding)
2604	2.	Apply edge detection to the object (Canny Edge Detection)
2606	3.	Apply key points detection (Orb)
2608	4.	Collect many key points from a number of frames as a template pool
	For each filtered frame, using the following process	
2610	1.	Mask the object from background (Otsu Thresholding)
2612	2.	Apply edge detection to the object (Canny)
2614	3.	Apply key points detection (Orb)
2616	4.	Use these key points to match with the key points in the template pool (Hamming Distance), and use the presence of each key points in the template pool as part of the input of machine learning algorithm.
2618	5.	Calculate the geometry characteristic (Length) and object surface characteristic (HSV LAB Channel), and use as the input of machine learning algorithm.
2620	6.	Apply machine learning algorithm.

FIG. 26

27/35

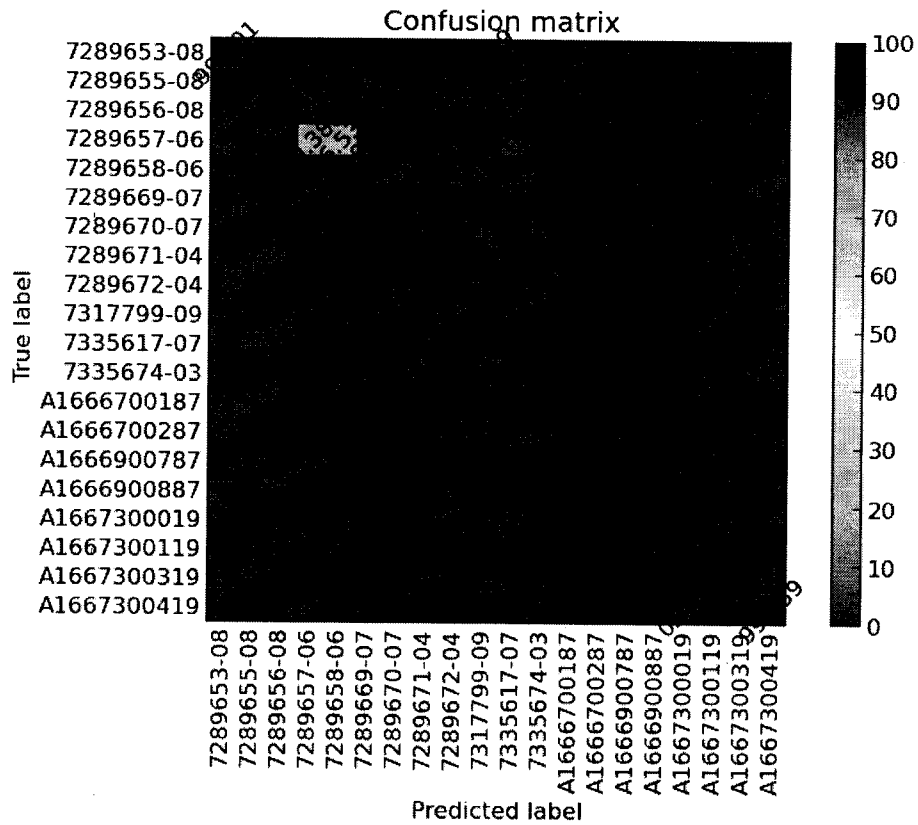
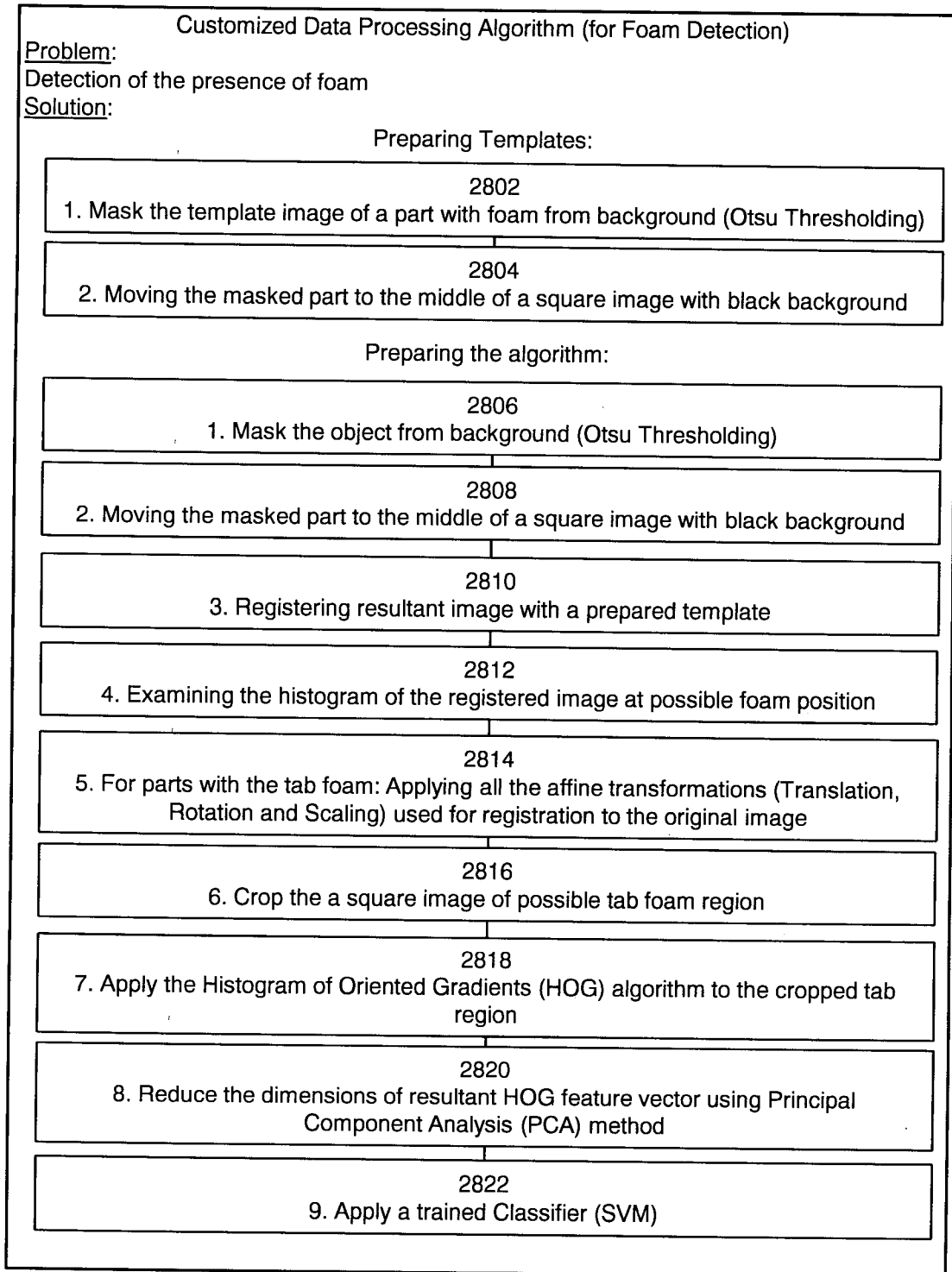


FIG. 27

28/35

**FIG. 28**

29/35

VERIFY QUALITY INSPECTION wireframe

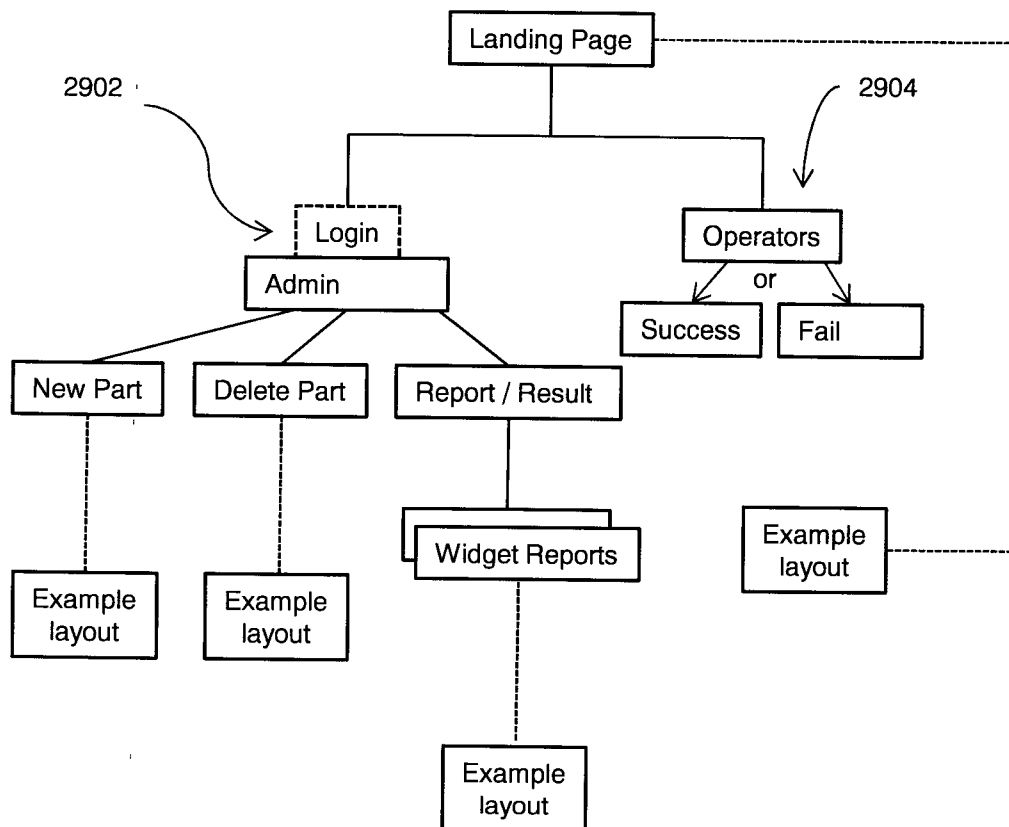


FIG. 29

30/35

GROUNDTRUTHER wireframe

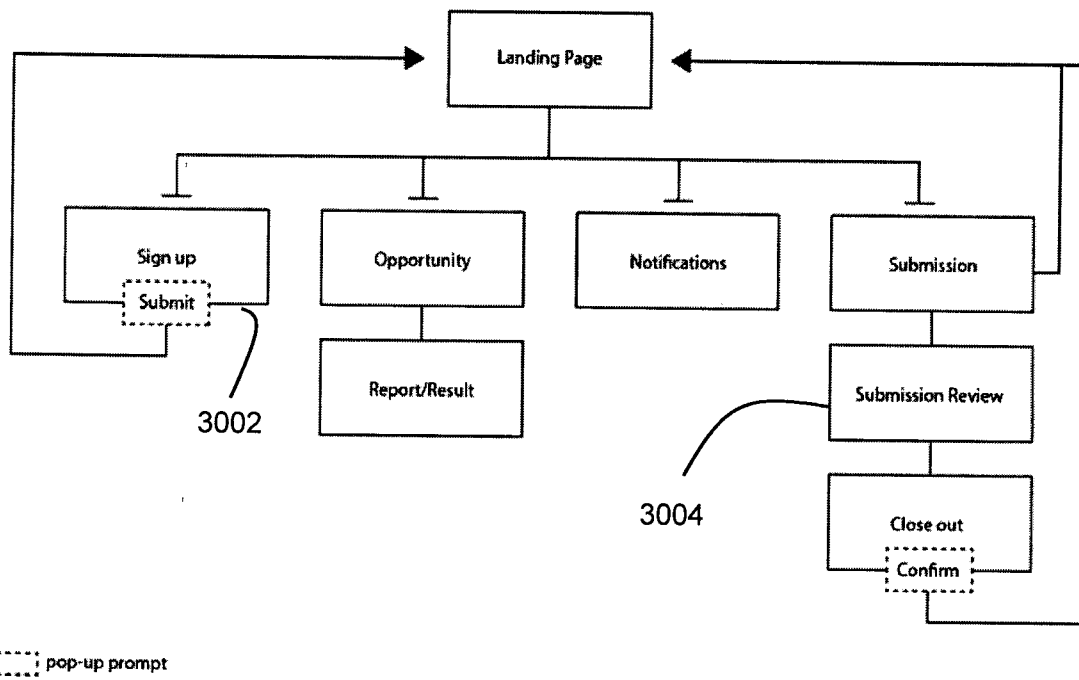


FIG. 30

31/35

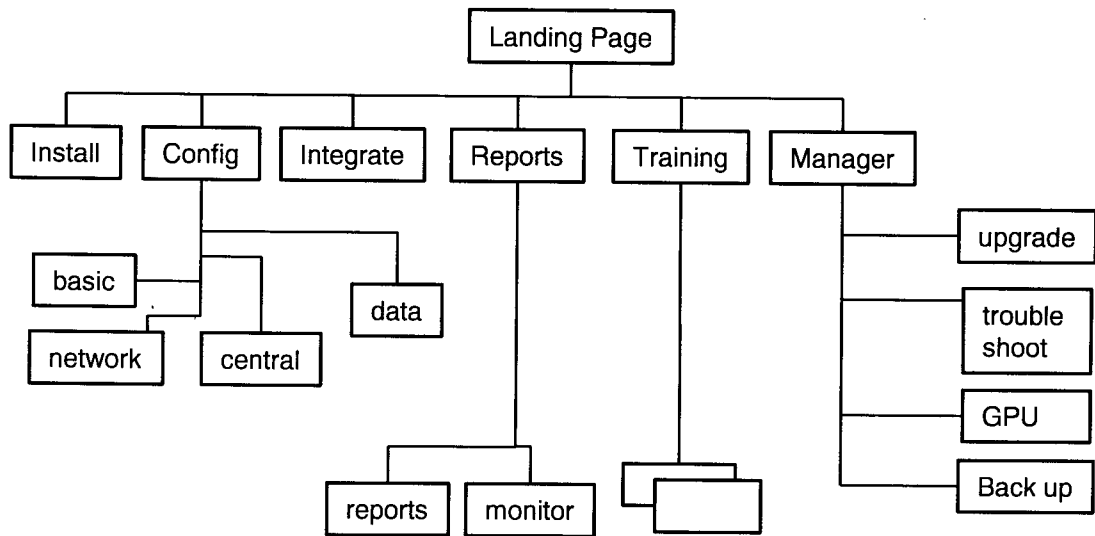
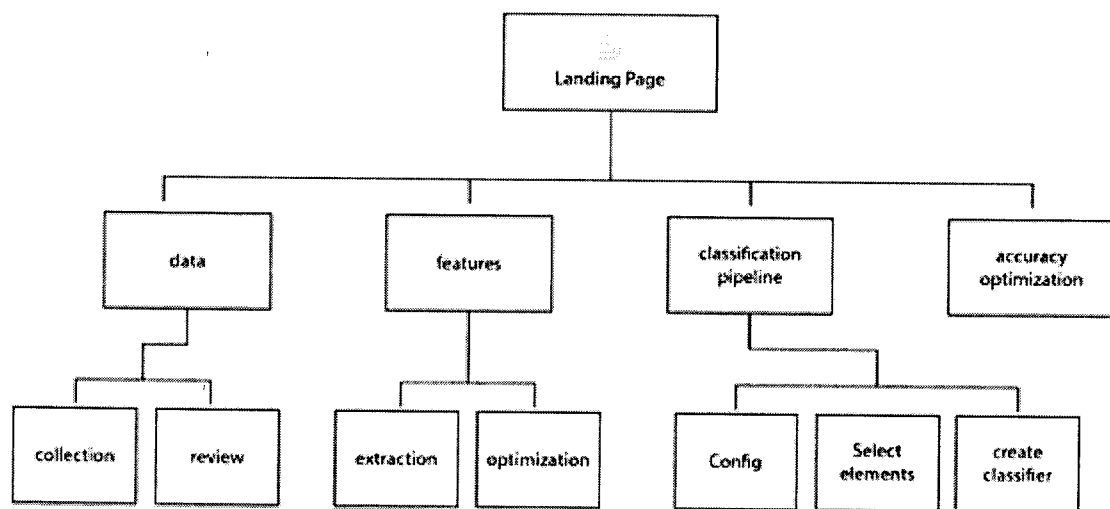


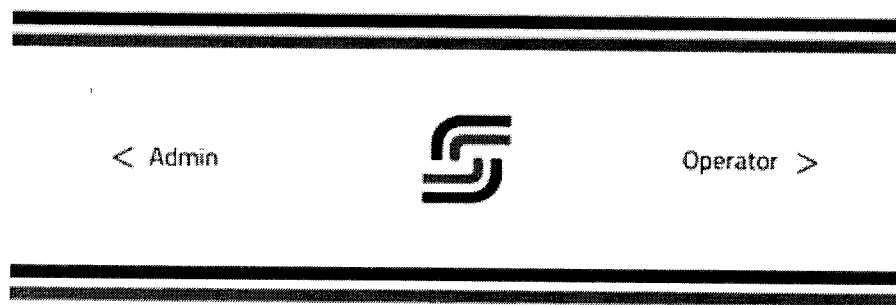
FIG. 31

32/35

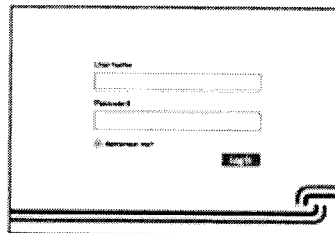
2.4 Use cases - SPE - Sightline Data Scientist wireframe

**FIG. 32**

33/35



Verify Landing Page



Verify Admin Login Page

FIG. 33

34 /35

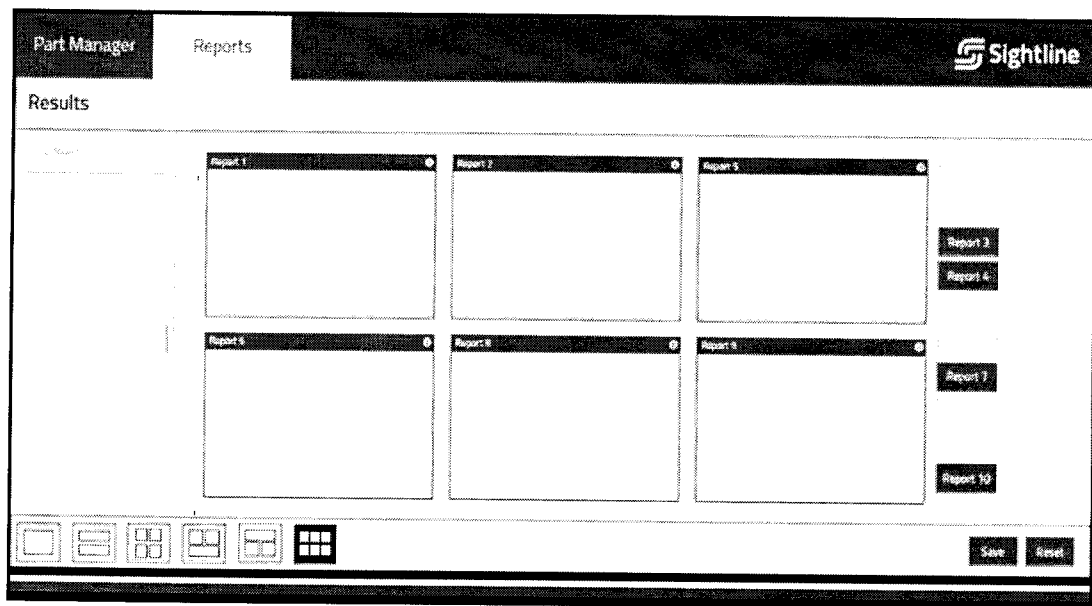
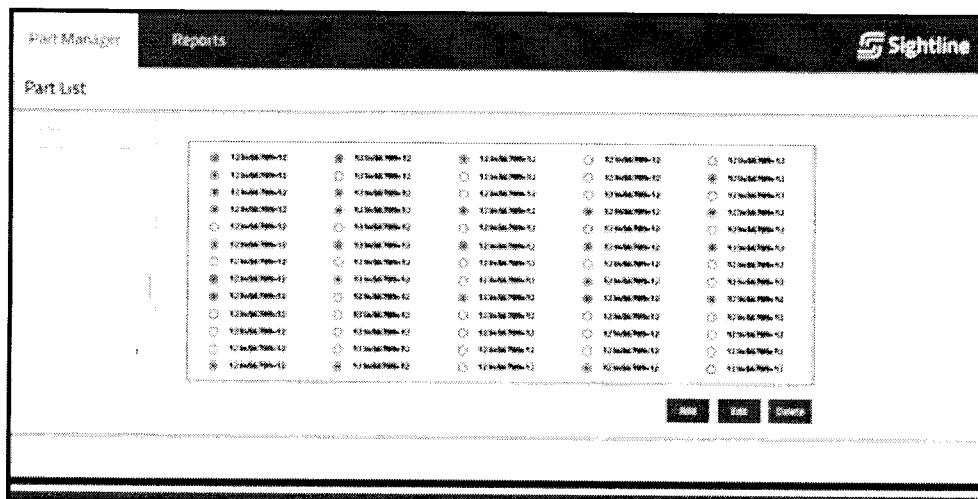


FIG. 34

35 / 35



Verify Part Manager Page

FIG. 35

200

