

(51) International Patent Classification ⁶ :

H04L 11/04

A2

(11) International Publication Number:

WO 96/07257

(43) International Publication Date:

7 March 1996 (07.03.96)

(21) International Application Number: PCT/US95/10605

(22) International Filing Date: 18 August 1995 (18.08.95)

(30) Priority Data:

08/293,073

19 August 1994 (19.08.94)

US

(71) Applicant: PEERLOGIC, INC. [US/US]; 555 Deharo Street, San Francisco, CA 94107-2348 (US).

(72) Inventors: GREGERSON, Daniel, P.; 1500 El Granada Boulevard, Half Moon Bay, CA 94019 (US). FARRELL, David, R.; Three Bayside Village, Number 201, San Francisco, CA 94114 (US). GAITONDE, Sunil, S.; 850 Campus Drive, Number 213, Daly City, CA 94015 (US). AHUJA, Ratinder, P.; 870 Campus Drive, Daly City, CA 94015 (US). RAMAKRISHNAN, Krish; 32529 Christine Drive, Union City, CA 94587 (US). SHAFIQ, Muhammad; 430 Portola Avenue, El Granada, CA 94018 (US). WALLIS, Ian, F.; 10430 South Blaney Avenue, Cupertino, CA 95014 (US).

(74) Agent: LAURIE, Ronald, S.; McCutchen, Doyle, Brown & Enersen, Three Embarcadero Center, San Francisco, CA 94111 (US).

(81) Designated States: AM, AT, AU, BB, BG, BR, BY, CA, CH, CN, CZ, DE, DK, EE, ES, FI, GB, GE, HU, IS, JP, KE, KG, KP, KR, KZ, LK, LR, LT, LU, LV, MD, MG, MK, MN, MW, MX, NO, NZ, PL, PT, RO, RU, SD, SE, SG, SI, SK, TJ, TM, TT, UA, UG, UZ, VN, European patent (AT, BE, CH, DE, DK, ES, FR, GB, GR, IE, IT, LU, MC, NL, PT, SE), OAPI patent (BF, BJ, CF, CG, CI, CM, GA, GN, ML, MR, NE, SN, TD, TG), ARIPO patent (KE, MW, SD, SZ, UG).

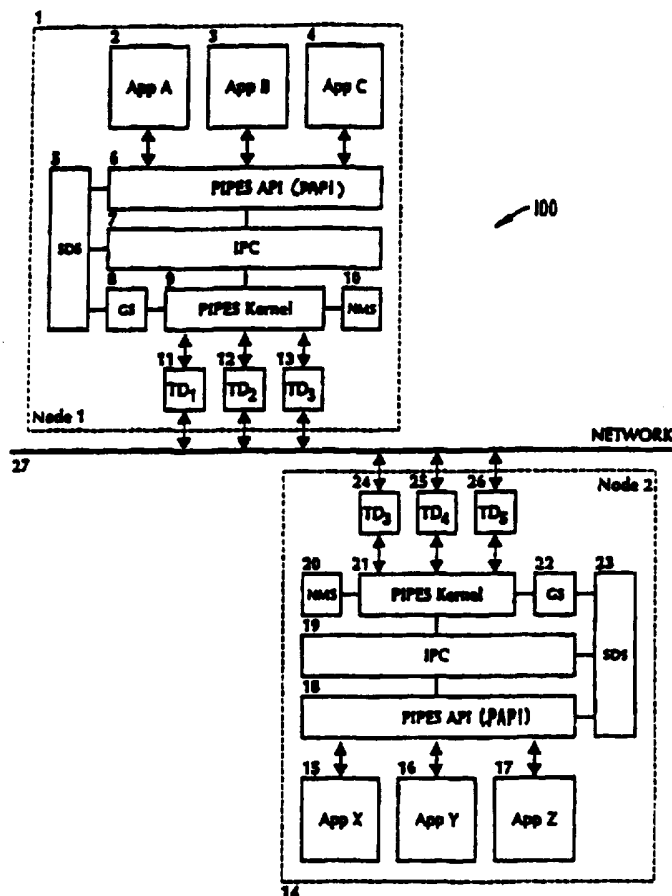
Published

Without international search report and to be republished upon receipt of that report.

(54) Title: SCALABLE DISTRIBUTED COMPUTING ENVIRONMENT

(57) Abstract

The present invention relates to distributed computing systems and is more particularly directed to an architecture and implementation of a scalable distributed computing environment which facilitates communication between independently operating nodes on a single network or on interconnected networks, which may be either homogeneous or heterogeneous. The present invention is a dynamic, symmetrical, distributed, real-time, peer-to-peer system comprised of an arbitrary number of identical (semantically equivalent) instances, i.e., kernels, that together form a logical tree. The kernels exhibit unified and consistent behavior at run time through a self-configuring and self-maintaining logical view of the network. Each kernel resides at a network node that has one or more resources associated with it. The kernels dynamically locate one another in real-time to form and maintain a hierarchical structure that supports a virtually unlimited number of independently running kernels. The system maintains its logical view of the network and user-developed programmatic resources regardless of the number and combinations of transport protocols and underlying mix of physical topologies. The system's communications services utilize a dynamic context bridge to communicate between end nodes that may not share a common transport protocol stack, thereby allowing applications residing on different stacks to communicate with one another automatically and transparently.



FOR THE PURPOSES OF INFORMATION ONLY

Codes used to identify States party to the PCT on the front pages of pamphlets publishing international applications under the PCT.

AT	Austria	GB	United Kingdom	MR	Mauritania
AU	Australia	GE	Georgia	MW	Malawi
BB	Barbados	GN	Guinea	NE	Niger
BE	Belgium	GR	Greece	NL	Netherlands
BF	Burkina Faso	HU	Hungary	NO	Norway
BG	Bulgaria	IE	Ireland	NZ	New Zealand
BJ	Benin	IT	Italy	PL	Poland
BR	Brazil	JP	Japan	PT	Portugal
BY	Belarus	KE	Kenya	RO	Romania
CA	Canada	KG	Kyrgyzstan	RU	Russian Federation
CF	Central African Republic	KP	Democratic People's Republic of Korea	SD	Sudan
CG	Congo	KR	Republic of Korea	SE	Sweden
CH	Switzerland	KZ	Kazakhstan	SI	Slovenia
CI	Côte d'Ivoire	LI	Liechtenstein	SK	Slovakia
CM	Cameroon	LK	Sri Lanka	SN	Senegal
CN	China	LU	Luxembourg	TD	Chad
CS	Czechoslovakia	LV	Latvia	TG	Togo
CZ	Czech Republic	MC	Monaco	TJ	Tajikistan
DE	Germany	MD	Republic of Moldova	TT	Trinidad and Tobago
DK	Denmark	MG	Madagascar	UA	Ukraine
ES	Spain	ML	Mali	US	United States of America
FI	Finland	MN	Mongolia	UZ	Uzbekistan
FR	France			VN	Viet Nam
GA	Gabon				

Scalable Distributed Computing Environment

Background of the Invention

The present invention relates to distributed computing systems and is more particularly directed to an architecture and implementation of a scalable distributed computing environment which facilitates communication between independently operating nodes on a single network or on interconnected networks, which may be either homogeneous or heterogeneous.

In today's business environment, corporate structures are being increasingly reshaped due to the dynamics of mergers and acquisitions, globalization and the need for real-time communication with customers, suppliers and financial institutions. In addition, immediate access to information and the need to manipulate that information quickly have become critical in establishing and maintaining competitive advantage. This requires that corporate data and the computer programs which manipulate that data be deployed in a fundamentally new way; in a distributed rather than a centralized, monolithic manner.

With distributed computing, programs and data are logically positioned so that they can be processed as near as possible to the users that interact with them. In theory, this allows the corporation to operate more reliably and efficiently by reducing communications overhead and exploiting the underutilized processing power of personal, group, and departmental computing resources. By distributing workload over many computers, information processing resources can be optimized for a given individual, work group or purpose. This approach allows data and processes to be

1 distributed and replicated so that performance and reliability can be more easily
2 maintained as the demands on the system increase. The characteristics of increased
3 granularity and scalability also provide important benefits relating to software
4 reusability, i.e., the same component may be used in several different applications, thus
5 reducing both development and maintenance time and costs.

6
7 Because of these demands, there is a movement toward enterprise-wide virtual
8 computing in which the entire resources of the network appear to the user to be locally
9 resident at his or her desktop computer or terminal. The traditional monolithic
10 centralized corporate information processing model is yielding to a distributed, fine-
11 grained approach. This transformation to virtual, dynamic enterprise computing
12 requires that mission critical core systems be implemented using a distributed
13 architecture in which localized computing resources (program elements and data) are
14 seamlessly interlinked by virtual networks.

15
16 However, in today's corporate information systems, individual applications
17 typically exist in heterogeneous environments that do not interoperate. Businesses are
18 faced with the task of connecting incompatible systems while maintaining an ever
19 increasing number of disparate operating systems and networking protocols over a wide
20 geographic area. Corporate mergers and acquisitions are again on the rise, and the
21 need to integrate installed heterogeneous networks into a single enterprise wide
22 network, not once but multiple times, is needed. Further, corporations have become
23 global entities and their information systems must now function over multiple time
24 zones, requiring those systems to be "time-independent." Moreover, as corporations
25 themselves are dynamically transformed, so are the information systems that support
26 their business operations. Thus, the corporate computing environment must be "open,"
27 i.e., it must be flexible enough to easily migrate to new standards while maintaining the
28 integrity and access to its existing "legacy" systems and data. Legacy systems typically

1 rely on the use of static tables to keep track of networked resources. Such systems do
2 not support dynamic recovery and are not easily scalable to enterprise-wide deployment
3 because of the extremely high overhead that would be required to maintain these tables
4 in a constantly changing environment.

5
6 In existing systems, in order for one resource connected to the network to
7 discover the existence of another resource, both must be "alive." As the total number
8 of resources connected to the network expands, it becomes vitally important to have a
9 mechanism for time-independent resource discovery whereby the network automatically
10 is made aware of new resources as they become available.

11
12 Existing systems are also limited by the availability of a fixed number of roles,
13 or hierarchical levels, that can be assumed by any node, e.g., machine, area, group,
14 domain, network, etc. This limitation presents significant problems when merging or
15 integrating two or more existing networks having different hierarchical structures. In
16 addition, in prior art systems, if a node assumes multiple roles, the relationship
17 between those roles is prescribed. That is, in order to function at level one (e.g.,
18 machine) and level 3 (e.g., group manager), the node must also assume the level 2
19 function (e.g., area manager). This limitation can severely degrade system performance
20 and recovery.

21
22 Prior attempts to address the problems associated with establishing robust,
23 efficient enterprise-wide computing environments, such as real time messaging,
24 message queuing, remote procedure calls, interprocess communication, and
25 broadcast/publish and subscribe represent partial solutions at best. Because true
26 distributed computing presupposes peer-to-peer communication (since master process
27 failure necessarily leads to failure of slave processes), client-server based approaches to
28 realizing the goal of enterprise computing represent suboptimal solutions. Existing

1 peer-to-peer systems utilizing static tables do not allow dynamic recovery and present
2 serious problems of scalability and maintenance.

3

4 Summary of the Invention

5

6 The present invention is a dynamic, symmetrical, distributed, real-time, peer-to-
7 peer system comprised of an arbitrary number of identical (semantically equivalent)
8 instances, i.e., kernels, that together form a logical tree. The kernels exhibit unified
9 and consistent behavior at run time through a self-configuring and self-maintaining
10 logical view of the network. Each kernel resides at a network node that has one or
11 more resources associated with it. The kernels dynamically locate one another in real-
12 time to form and maintain a hierarchical structure that supports a virtually unlimited
13 number of independently running kernels. The system maintains its logical view of the
14 network and user-developed programmatic resources regardless of the number and
15 combinations of transport protocols and underlying mix of physical topologies. The
16 system's communications services utilize a dynamic context bridge to communicate
17 between end nodes that may not share a common transport protocol stack, thereby
18 allowing applications residing on different stacks to communicate with one another
19 automatically and transparently.

20

21 The system is designed to support all forms of digitized communication,
22 including voice, sound, still and moving images, mass file transfer, traditional
23 transaction processing and any-to-any conferencing such as "groupware" applications
24 would require. The system is also designed to operate over any type of networking
25 protocol and medium, including ISDN, X.25, TCP/IP, SNA, APPC, ATM, etc. In all
26 cases, the system delivers a high percentage, typically 60-95%, of the theoretical
27 transmission capacity, i.e., bandwidth, of the underlying medium.

28

1 As new resources join (or rejoin) the network, the kernel residing at each node,
2 and thus each resource connected to that node, automatically and immediately becomes
3 accessible to all applications using the system. The role(s) assumed by any node within
4 the managerial hierarchy employed (e.g., area manager, domain manager, network
5 manager, etc.) is arbitrary, i.e., any node can assume one or multiple roles within the
6 hierarchy, and assuming one role neither requires nor precludes assumption of any
7 other role. Further, the roles dynamically change based on the requirements of the
8 network, i.e., as one or more nodes enter or leave the network. Thus, the individual
9 kernels dynamically locate one another and negotiate the roles played by the associated
10 nodes in managing the network hierarchy without regard to their physical location. In
11 addition, the number of possible roles or levels that may be assumed by any node is not
12 limited and may be selected based on the particular requirements of the networking
13 environment.

14 15 **Brief Description of the Drawings**

16
17 These and other features and advantages of the present invention will be better
18 and more completely understood by referring to the following detailed description of
19 preferred embodiments in conjunction with the appended sheets of drawings, of which:

20
21 Fig. 1 is a drawing showing a distributed computing system in accordance with
22 the present invention.

23
24 Fig. 2 is a detailed block diagram of one of the nodes in the system of Fig. 1.

25
26 Fig. 3 is a block diagram showing the structure of a kernel in accordance with
27 the present invention.

1 Fig. 4 is a flow chart of the PIPES logical network (PLN) of the present
2 invention.

3
4 Fig. 5 is a flow chart of a child login procedure in accordance with the present
5 invention.

6
7 Fig. 6 is a flow chart of a parent login procedure in accordance with the present
8 invention.

9
10 Fig. 7 is a diagram showing the login communication between different nodes in
11 accordance with the present invention.

12
13 Fig. 8 is a flow chart of a roll call procedure in accordance with the present
14 invention.

15
16 Fig. 9 is a diagram showing the roll call communication between different nodes
17 in accordance with the present invention.

18
19 Fig. 10 is a flow chart of a child monitor procedure in accordance with the
20 present invention.

21
22 Fig. 11 is a flow chart of a parent monitor procedure in accordance with the
23 present invention.

24
25 Fig. 12 is a diagram showing the "heartbeats" monitor communication between
26 different nodes in accordance with the present invention.

27

1 Fig. 13 is a flow chart of an election process in accordance with the present
2 invention.

3

4 Fig. 14 is a diagram showing the election communication between different
5 nodes in accordance with the present invention.

6

7 Fig. 15 is a flow chart of a logout process in accordance with the present
8 invention.

9

10 Fig. 16 is a diagram showing the logout communication between different nodes
in accordance with the present invention.

11

12 Fig. 17 is a diagram showing activities relating to a resource of the present
13 invention.

14

15 Fig. 18 is a flow chart of an "Add Resource" process in accordance with the
16 present invention.

17

18 Fig. 19 is a flow chart of a "Find Resource" process in accordance with the
19 present invention.

20

21 Fig. 20 is a flow chart of a "Find Resource" process at an area manager node of
22 the present invention.

23

24 Fig. 21 is a flow chart of a "Find Resource" process in accordance with the
25 present invention at a level above area manager.

26

27 Fig. 22 is a flow chart of a "Persistent Find" process at an area manager node
28 of the present invention.

1 Fig. 23 is a flow chart of a "Persistent Find" process in accordance with the
2 present invention at a level above area manager.

3
4 Fig. 24 is a flow chart of a "Clean Persistent Find" process at an area manager
5 node of the present invention.

6
7 Fig. 25 is a flow chart of a "Clean Persistent Find" process in accordance with
8 the present invention at a level above area manager.

9
10 Fig. 26 is a flow chart of a "Resource Recovery" process in accordance with the
11 present invention when an area manager goes down.

12
13 Fig. 27 is a flow chart of a "Resource Recovery" process in accordance with the
14 present invention when another managerial node goes down.

15
16 Fig. 28 is a flow chart of a "Remove Resource" process in accordance with the
17 present invention.

18
19 Fig. 29A shows the components of a context bridge of the present invention.

20
21 Fig. 29B is an example illustrating the use of context bridges for communication
22 between different protocols.

23
24 Fig. 30 is a flow chart showing a context bridge routing process in accordance
25 with the present invention.

26
27 Fig. 31 is a flow chart of a "Route Discovery" process in accordance with the
28 present invention.

1 Fig. 32 is a flow chart of a "Route Validation" process in accordance with the
2 present invention.

3
4 Fig. 33 is a flow chart of a "Route Advertisement" process in accordance with
5 the present invention.

6
7 Fig. 34 is a flow chart showing the steps performed in changing the number of
8 levels in the PIPES logical network of the present invention.

9
10 **Detailed Description of the Invention**

11
12 FIG. 1 shows a distributed computing system 100 in accordance with the present
13 invention. The implementation of system 100 by the assignee of the present application
14 is referred to as the PIPES Platform ("PIPES"). In system 100, two nodes, Node 1
15 (shown as block 1) and Node 2 (shown as block 14), communicate through a physical
16 network connection (shown as line 27). It should be obvious to a person skilled in the
17 art that the number of nodes connected to network 27 is not limited to two.

18
19 The structures of the nodes are substantially the same. Consequently, only one
20 of the nodes, such as Node 1, is described in detail. Three applications, App. A
21 (shown as block 2), App. B (shown as block 3), and App. C (shown as block 4), run
22 on Node 1. These applications are typically written by application developers to run on
23 PIPES. The PIPES software includes a PIPES Application Programmer Interface
24 ("PAPI") (shown as block 6) for communicating with Apps. A-C. PAPI 6 sends
25 messages to a single PIPES Kernel (shown as block 9) executing at Node 1 through
26 Interprocess Communication (IPC) function calls (shown as block 7). Kernel 9 sends
27 and receives messages over network 27 through transport device drivers TD₁ (shown as
28 block 11), TD₂ (shown as block 12), and TD₃ (shown as block 13).

1 Similarly, Node 2 has three applications running on it, App. X (shown as block
2 15), App. Y (shown as block 16), and App. Z (shown as block 17), and communicating
3 with a single PIPES Kernel (shown as block 21) running at Node 2 through PAPI
4 (shown as block 18) and IPC (shown as block 19). Node 2 supports three different
5 network protocols, and thus contains three transport drivers TD₃ (shown as block 24),
6 TD₄ (shown as block 25), and TD₅ (shown as block 26).

7
8 For example, if App. A at Node 1 needs to communicate with App. Z at Node
9 2, a message travels from App. A through PAPI 6, IPC 7, and kernel 9. Kernel 9 uses
10 its transport driver TD₃ to send the message over network 27 to transport driver TD₃ at
11 Node 2. The message is then passed to kernel 21 at Node 2, IPC 19, PAPI 18, and
12 finally to App. Z.

13
14 PIPES also provides generic services used by all of its component parts.
15 Network Management Services (shown as blocks 10 and 20) provides access for a
16 PIPES Network Management Agent (not shown) to monitor the kernels' network- and
17 system-level counters, attributes, and statistics. Generic Services (shown as blocks 8
18 and 22) provide a common interface for kernels 9 and 21 to operating system services,
19 including hashing, btrees, address manipulation, buffer management, queue
20 management, logging, timers, and task scheduling. System Dependent Services (shown
21 as blocks 5 and 23) provides services specific to operating system, platform,
22 environment and transports on the nodes. These services are used by Generic Services
23 (shown as blocks 8 and 22) to realize a generic service within a given operating system
24 or platform environment.

25
26 FIG. 2 shows a more detailed block diagram of the PIPES internal architecture
27 within Node 1 of system 100. The PIPES architecture is divided into three different
28 layers: the Interface Layer (shown as block 28), the Kernel Layer (shown as block

1 29), and the Transport Layer (shown as block 30). Interface Layer 28 handles queries
2 from and responses to the applications that are accessing the PIPES environment
3 through PAPI 6. Interface Layer 28 is embodied in a library which is linked to each
4 application (e.g., Apps. A-C) which accesses kernel 9. Kernel Layer 29 provides
5 programmatic resource and communication management services to applications that are
6 accessing PIPES, allowing communication between end-nodes that may not share a
7 transport protocol stack. Transport Layer 30 consist of the transport device drivers 11,
8 12, and 13 for the network protocols supported by Node 1. Each transport driver
9 provides access from kernel 9 to a network transport protocol provided by other
10 vendors, such as TCP/IP, SNA, IPX, or DLC. Transport Layer 30 handles all
11 transport-specific API issues on a given platform for a given transport discipline.
12

13 FIG. 3 illustrates the internal architecture of kernel 9. Kernel 9 contains an API
14 Interface (shown as block 31) which is the interface to PAPI 6 of Fig. 2. API Interface
15 31 handles requests from Interface Layer 28 and returns responses to those requests. It
16 recognizes an application's priority and queues an application's messages based on this
17 priority. API Interface 31 also handles responses from the Resource Layer (shown as
18 block 32) and Session Services (shown as block 35), and routes those responses to the
19 appropriate application.
20

21 Resource Layer 32 registers an application's resources within a PIPES Logical
22 Network ("PLN") layer (shown as block 33), provides the ability to find other PAPI
23 resources within PIPES, and handles the de-registration of resources within the
24 network. In addition, Resource Layer 32 implements a "Persistent Find" capability
25 which enables the locating of resources that have not yet been registered in PLN 33.
26

27 PLN 33 maintains knowledge of the logical, hierarchical relationships of the
28 nodes within PIPES to enforce a dynamic administrative framework. PLN 33 handles

1 the election of managers, the transparent reestablishment of management hierarchies as
2 a result of physical network faults. PLN 33 employs a system of "heartbeat" messages
3 which is used to monitor the status of nodes within the network and identify network
4 failures. This layer also handles requests and returns responses to Resource Layer 32
5 and an Acknowledged Datagram Service ("AKDG", shown as block 34).

6
7 AKDG 34 provides best-effort datagram service with retransmission on failures
8 for users. AKDG 34 handles the sending and receiving of messages through
9 Connectionless Messaging Service (CLMS) 36 and Session Services 35.

10
11 Session Services 35 allocates, manages, and deallocates sessions for users.
12 Session management includes sending and receiving data sent by the user in sequence,
13 ensuring secure use of the session, and maintaining the message semantics over the
14 Connection Oriented Messaging Service (COMS) stream protocol. Session Services 35
15 also multicasts PAPI application messages over sessions owned by the PAPI
16 application. Session Services 35 interacts with COMS 37 to satisfy requests from
17 AKDG 34 and API Interface 31.

18
19 CLMS 36 transfers data without a guarantee of delivery. It also interacts with
20 Context Bridge layer 38 to satisfy the requests from AKDG 34.

21
22 COMS 37 manages connections opened by Session Services 35. COMS 37
23 provides high-performance data transfer, including the fragmentation and reassembly of
24 messages for users. COMS 37 modifies message size based on maximum message
25 sizes of hops between connection endpoints.

1 Context Bridge layer 38 insulates PAPI applications from the underlying
2 networks by performing dynamic transport protocol mapping over multiple network
3 transports, thus enabling data transfer even if the end-to-end protocols are different.
4

5 The Transport Driver Interface (shown as block 39) handles communication
6 between transport-specific drivers and the CLMS 36 and COMS 37 layers. This
7 interface contains generic common code for all transport drivers.
8

9 **PLN Layer**

10

11 PLN 33 is a hierarchical structure imposed by the system administrator on a set
12 of machines executing kernels. These kernels unify at run time to form a hierarchical
13 network with dynamically elected managers that manage a given level of the hierarchy.
14 The PLN name space is divided into five different levels: normal, area, group,
15 domain, and network. All kernels at startup have normal privileges. They assume a
16 managerial role depending on their configuration in the network and such real-time
17 considerations as the number of roles already assumed. Thus, administrative functions
18 will be distributed evenly among the member kernels, leading to better performance and
19 faster recovery. It should be appreciated that the number of levels is not limited to
20 five, and any number of levels can be implemented in the system, as explained below.
21

22 In PLN 33, the primary roles played by the various managers between the
23 Network Manager and Area Manager (e.g., Domain Manager and Group Manager) are
24 essentially the same: to maintain communication with its parent and children, and to
25 route Resource Layer 32 traffic. In addition to these functions, any manager between
26 the Network Manager and Area Manager (e.g., Domain or Group) also provides
27 persistent find source caching services as described below in connection with Figs. 22
28 and 23. The Area Manager, in addition to these functions described above, provides

1 caching services for resources advertised by its children, including all of the kernels in
2 the Area Manager's name space. Therefore, the Area Manager is crucial to the orderly
3 function of PLN 33, which is built from the ground up by filling the Area Manager
4 role before any other role in the hierarchy. By default, any kernel can become an Area
5 Manager.

6
7 • As shown in FIG. 4, the PLN building and maintenance algorithm comprises
8 five main processes: Login (shown as block 100), Role Call (shown as block 200),
9 Monitor (shown as block 300), Election (shown as block 400), and Logout (shown as
10 block 500). In this description, the following terms are used in order to allow for the
11 appropriate abstraction. The number of levels in PLN 33 is defined by *MinLevel* and
12 *MaxLevel*. The kernels that have normal privileges are configured at *MinLevel* and are
13 not managers. On the other hand, a kernel that is the Network Manager is configured
14 at *MaxLevel* and has the potential to become the Network Root. The configuration
15 parameter *MaxStatus* imposes a ceiling on the highest level of which the kernel can be
16 a manager. A kernel at level n is termed to be a *child* of its *parent* kernel at level $n+1$
17 provided that the two kernels have the same name above level n .

18 19 Login

20
21 FIGS. 5 and 6 depict the Login procedure executed at the child and parent nodes
22 in PLN 33. Login is a process by which a child kernel locates and registers with a
23 parent kernel. FIG. 7 illustrates the messages passed between kernels during a
24 hypothetical execution of the Login process by which a kernel in node N7 (shown as
25 circle 37 and referred to as kernel N7) runs the Login process to enter the network.

26
27 A kernel enters the network by running the Login process to locate its parent
28 kernel. The child kernel first enters a wait period (step 101) during which the child

1 listens for other login broadcasts on the network (step 102). If a login broadcast is
2 received during the wait period (step 103), the child kernel reads the message. The
3 information in the message is sufficient for the child to ascertain the identity of its
4 parent and siblings. If the originator of the message is a sibling (step 104), the child
5 kernel modifies its Login wait period interval (step 105) in order to prevent login
6 broadcasts from inundating the network. If the originator of the message is a parent
7 (step 106), the child kernel sends a login request to the parent (step 107) and waits for
8 an acknowledgement. If a login broadcast is not received, the child kernel continues to
9 listen for a login broadcast until the end of the wait period (step 108). At the end of
10 the wait period, the child kernel sends a login broadcast on the network (step 109).

11
12 In FIG. 7, kernel N7 is attempting to login to the PIPES network by sending a
13 login broadcast message (represented by dotted line a) to a kernel in node N1
14 (represented by circle 31 and referred to as kernel N1), a kernel in node N2
15 (represented by circle 32 and referred to as kernel N2), a kernel in node N3
16 (represented by circle 33 and referred to as kernel N3), a kernel in node 4 (represented
17 by circle 34 and referred to as kernel N4), a kernel in node N5 (represented by circle
18 35 and referred to as kernel N5), and a kernel in node N6 (represented by circle 36 and
19 referred to as kernel N6). The child kernel waits for a specified time to receive a login
20 acknowledgement (step 110).

21
22 All kernels listen for login broadcast messages on the network (step 116). If a
23 login broadcast is received (step 117), the parent kernel determines whether the kernel
24 that sent the message is its child (step 118). If the originating kernel is not its child,
25 the parent continues listening for login broadcasts (step 116). However, if the
26 originating kernel is its child, the parent checks if this kernel is a duplicate child (step
27 119). If this is a duplicate child, the parent informs its duplicate children of a role

1 conflict (step 120). If not, the parent sends a login acknowledgement to its child kernel
2 (step 121).

3

4 In FIG. 7, parent kernel N4 receives kernel N7's login broadcast message a,
5 and sends a login acknowledgement message represented by line b to kernel N7.

6

7 If a login acknowledgement is received (step 110), the child kernel sends a login
8 confirmation to the first parent kernel that sends a login acknowledgement (step 114).
9 The child kernel ignores any other login acknowledgements it may receive. After
10 sending the login confirmation to its parent, the child kernel begins the Monitor process
11 with its new parent (step 115). If the parent kernel receives the login confirmation
12 (step 122), the parent kernel registers the child (step 123) and begins the Monitor
13 process with its new child (step 124). If the parent kernel does not receive the login
14 confirmation from the child (step 122), the parent kernel continues to listen for login
15 broadcasts (step 116).

16

17 In FIG. 7, after receiving parent kernel N4's login acknowledgement b, child
18 kernel N7 sends a login confirmation message represented by line c to kernel N4 and
19 begins the monitor process with its parent kernel N4.

20

21 If no parent kernel sends a login acknowledgement to the child, the child kernel
22 begins the Login process again (step 101) unless the retry threshold has been exceeded
23 (step 111). If the retry threshold has been exceeded, the child checks its *MaxStatus*
24 setting (step 112). If the child's *MaxStatus* is greater than *MinLevel*, the child begins
25 the Role Call process to assume the role of its own parent. Otherwise, the child kernel
26 will enter the Login wait period again (step 101).

27

1 *Role Call*

2

3 Role Call is a procedure by which a kernel queries the network to find out
4 vacancies in the name space hierarchy. The procedure is executed by all kernels who
5 have been configured with *MaxStatus* greater than *MinLevel*. The Role Call procedure
6 is invoked by a kernel upon startup and subsequently when there is a managerial
7 vacancy in its namespace. The Role Call algorithm is designed to minimize the number
8 of kernels simultaneously participating in the Role Call process, reducing network-wide
9 broadcasts as well as possible collisions between potential contenders for the same
10 vacancy.

11

12 The roll call procedure is shown in Fig. 8. A kernel wishing to participate in
13 Role Call goes through a forced wait period (step 201). The wait period is a function
14 of the number of roles the kernel has already assumed, whether the kernel is an active
15 context bridge, and the current state of the kernel. A random wait interval is also
16 added to the equation.

17

18 During the wait period, the kernel listens for role call broadcasts from other
19 kernels (step 202). If a role call broadcast is received for the same level of the
20 hierarchy (step 203), the kernel abandons the Role Call procedure (step 204). If a role
21 call broadcast is not received, the kernel continues to listen for role call broadcasts
22 (step 202) until the end of the wait period (step 205). At the end of the wait period,
23 the kernel sends its own role call broadcast on the network (step 206). The broadcast
24 message contains the level of the hierarchy for which the role call is being requested.
25 After sending the role call broadcast, the kernel starts a timer (step 207) and listens for
26 role call messages on the network (step 208). A kernel that is a manager of the
27 namespace for which role call is requested will respond with a point-to-point role call
28 acknowledgement message. If the kernel initiating the role call receives the

1 acknowledgement (step 209), the kernel will abandon the Role Call procedure (step
2 204). If the kernel initiating the role call instead receives another role call broadcast
3 for the same level of the hierarchy (step 210), the kernel reads the message. If the
4 originator of the message has higher credentials (step 211), the kernel will abandon the
5 Role Call procedure (step 204). The credentials of a particular kernel are a function of
6 the number of roles the kernel has already assumed, whether the kernel is an active
7 context bridge, and the current state of the kernel. At the end of the timeout period
8 (step 212), the kernel assumes the vacant managerial role for which it requested role
9 call (step 213).

10
11 FIG. 9 depicts an example of the Role Call procedure. Kernel N4, represented
12 by circle 44, becomes isolated from the network due to physical connection problems.
13 Kernel N7, represented by circle 47, detects the absence of kernel N4 as a result of its
14 Monitor process (described in detail below) with its parent kernel N4. Kernel N7 goes
15 into the forced wait period and listens for role call broadcast traffic on the network. If
16 kernel N5, represented by circle 45, had started its Role Call process before kernel N7,
17 kernel N7 would abort its Role Call after receiving kernel N5's role call broadcast
18 message, represented by dotted line i. However, assuming that kernel N7 started its
19 Role Call first, kernel N7 sends out its broadcast message, represented by dotted line h,
20 at the end of the role call wait period.

21
22 If kernel N5 sends its own role call broadcast message after kernel N7 has
23 already done so, kernel N7 compares its credentials with those of kernel N5. If kernel
24 N5's credentials are higher, kernel N7 abandons Role Call and kernel N5 assumes the
25 managerial role left vacant by the disappearance of kernel N4. If kernel N7's
26 credentials are higher, kernel N5 abandons Role Call and kernel N7 assumes kernel
27 N4's vacant managerial role at the end of the timeout period.

28

1 If kernel N4 has reappeared on the network and has received kernel N5's
2 broadcast message i or kernel N7's broadcast message h, kernel N4 responds by
3 sending an acknowledgement message to kernel N5, represented by line j, or to kernel
4 N7, represented by line k. If kernel N4 has not reappeared on the network, kernel N5
5 and kernel N7 continue their Role Call processes.

6
7 *Monitor*

8
9 FIGS. 10 and 11 depicts the child and parent Monitor processes, which is used
10 to keep track of one another.

11
12 The parent has its own "heartbeat" timer set to the slowest heartbeat interval of
13 all of its children. The parent initially resets its heartbeat timer at the beginning of the
14 Monitor process (step 312) and listens for heartbeat messages from its children (step
15 313). A child participating in the Monitor process with its parent first sends a
16 heartbeat message to its parent (step 301) and waits for an acknowledgement. If a
17 heartbeat message is received by the parent (step 314), the parent will send a heartbeat
18 acknowledgement to the child (step 315) and check off the child in its list of children
19 (step 316). The acknowledgement message contains a heartbeat offset value to scatter
20 the heartbeat intervals among its children. If the child receives the heartbeat
21 acknowledgement (step 302), the child modifies its heartbeat interval (step 306) and
22 enters a wait period (step 307). If the child does not receive a heartbeat
23 acknowledgement, it sends another heartbeat message to its parent (step 303). If a
24 heartbeat acknowledgement is received (step 304) at this time, the child then modifies
25 its heartbeat interval (step 306) and enters the wait period (step 307). If the child still
26 does not receive a heartbeat acknowledgement, the child assumes that it has become
27 orphaned and begins the Login process (step 305).

1 When the parent's heartbeat timer expires (step 317), the parent checks its list of
2 children for missing heartbeat messages (step 318). If the parent detects a missing
3 heartbeat, the parent sends a heartbeat message to the missing child (step 319). If the
4 parent does not receive a heartbeat acknowledgement from the missing child (step 320),
5 the parent de-registers the child (step 321).

6
7 During its wait period (step 307), the child listens for a heartbeat message from
8 its parent (step 308). If a heartbeat message is received by the child (step 309), the
9 child sends a heartbeat acknowledgement to its parent (step 310), modifies its heartbeat
10 interval (step 306), and enters the wait period again (step 307). At the end of the wait
11 period (step 311), the child begins the Monitor process once again (step 301).

12
13 FIG. 12 shows the periodic check-in messages, or "heartbeats," passed between
14 the parent and child during the Monitor process. In FIG. 12, kernels N3 and N4
15 (represented by circles 53 and 54, respectively) are the child of kernel N2 (represented
16 by circle 52). Kernel N2 is in turn the child of kernel N1 (represented by circle 51).
17 Messages d_1 through d_3 represent heartbeat messages from child to parent, while
18 messages e_1 through e_3 represent heartbeat acknowledgements from parent to child.
19 Messages f_1 through f_3 represent heartbeat messages from parent to child, while
20 messages g_1 through g_3 represent heartbeat acknowledgements from child to parent.

21 *Election*

22
23
24 PIPES kernels engage in a distributed Election (FIG. 13) to determine the
25 winner when role conflicts arise. Two or more managers may claim managerial
26 responsibility over the same namespace when there are problems in the underlying
27 physical connections that cause fragmentation of the network. Collisions in the
28 namespace are primarily detected through either role call or login broadcasts, described

1 above. When a kernel detects a namespace collision, it will inform the principals that
2 in turn execute the Election process. New participants may join an Election that is
3 already in progress. Because the Election is fully distributed, each kernel separately
4 conducts the Election and arrives at the result.

5
6 When a kernel detects a role conflict or is informed of one, the kernel begins
7 the Election process by starting an election timer and opening an election database (step
8 401). The kernel stores the election participants known so far, and sends an election
9 request to each one (step 402). This message consists of all known kernels that are
10 participating in the election. The kernel then listens for any election traffic on the
11 network (step 403). If the kernel receives an election response (step 404), which
12 contains a list of known participants, the kernel stores any new election participants in
13 the database and sends each one an election request (step 402). If another election
14 request is received (step 405), the kernel sends an election response to the originator
15 (step 406), updates the election database, and sends election requests to the new
16 participants (step 402). When the election timer expires (step 407), the kernel queries
17 its election database to determine the winner (step 408). The winner of an election
18 depends on the number of roles each participating kernel has already assumed, whether
19 the participating kernels are active context bridges, and the current state of each kernel.
20 If the kernel is the winner of the election (step 409), the kernel sends an election result
21 message to all election participants (step 410). If the kernel loses the election, the
22 kernel will resign its post as manager (step 411), informing all of its children of their
23 new parent. All participants in the election verify the election result and finally close
24 their election databases (step 412).

25
26 FIG. 14 illustrates an example of the Election process. Suppose that kernels A
27 and B (represented by circles 61 and 62, respectively) have detected role conflicts
28 independently. Kernel A will send an election request message (arrow 1) to kernel B.

1 This message will consist of participants known to kernel A, at this point being just
2 kernels A and B. When kernel B receives this message, kernel B will send kernel A an
3 election response message (arrow m). Later, kernel C detects a role conflict with
4 kernel B. Kernel C will then send an election request message (arrow n) to kernel B.
5 Kernel B will update its election database with the new entrant kernel C and will send
6 an election response message (arrow o) back to kernel C. This message will contain
7 the election participants known to kernel B at this point, namely, kernels A, B, and C.
8 When kernel C receives this message, it will detect the new contestant kernel A, update
9 its election database, and send an election request message (arrow p) to kernel A. At
10 this point, kernel A will become aware of the new contestant (from its perspective),
11 update its database with kernel C's credentials, and respond to kernel C's request
12 (arrow q). In the same fashion, when kernel D enters the election only aware of kernel
13 A, it will soon be aware of kernels B and C.

14

15 *Logout*

16

17 Logout (FIGS. 15 & 16) is a procedure by which a kernel de-registers from its
18 parent. Logout may be initiated as part of the kernel shutdown logic, or as a result of
19 resigning as a manager of a particular level of the hierarchy. A child kernel (shown as
20 kernel N2 in Fig. 16) sends a logout request (represented by arrow x) to its parent,
21 shown as kernel N1 in Fig. 16 (step 501). When the parent receives the logout request
22 from its child (step 506), it sends a logout acknowledgement (shown as arrow y in Fig.
23 16) to the child (step 507) and de-registers the child (step 508). If the child is a
24 manager (step 503), the child will send messages (represented by messages z_1 through
25 z_3 in Fig. 16) inform all of its children (i.e., kernels N3, N4, and N5 in Fig. 16) that it
26 is no longer their parent (step 504). In addition, the parent kernel will nominate a
27 successor from among its children by nominating the winner of an election process
28 which it performs on its children (step 505).

1 **Resource Layer**

2
3 The Resource Layer (block 32 in FIG. 3) is responsible for managing all of the
4 resources distributed throughout the PIPES network hierarchy. A resource is a
5 functional subset of a PIPES application that is made available to other PIPES
6 applications executing at other nodes on the network. A PIPES resource can be thought
7 of as a well-defined service element, where one or more elements, when considered as
8 a whole, combine to form a complete service.

9
10 FIG. 17 describes the life cycle of a resource in PIPES. A resource enters the
11 network through the Add Resource process (block 600). In order to utilize the services
12 provided by a resource, an application must execute the Find Resource Process (block
13 700) to determine its location within the PIPES address space. For example, after
14 executing a Find Query and obtaining the address of an available resource, an
15 application might attempt to establish a session with the resource through Session
16 Services 35.

17
18 If a resource is not available at the time an application executes a Find Query,
19 the application might alternatively execute a Persistent Find Query, which will notify
20 the application of a resource's availability as soon as a resource meeting the search
21 criteria enters the network through the Add Resource Process. In this case, Area
22 Managers in PIPES maintain caches of pending Persistent Find Queries to facilitate an
23 immediate response to such a query. If an Area Manager were to become disconnected
24 from the rest of the PIPES hierarchy through a physical network failure, a recovery
25 mechanism (block 800) is employed to recreate the persistent find cache at the new
26 Area Manager that takes over the disconnected manager's responsibilities.

27

1 During its lifetime on the network, a resource is available to provide services to
2 applications on the network. If the application that owns the resource removes the
3 resource from the network, the Resource Layer executes the Remove Resource process
4 (block 900).

5
6 *Add Resource Process*

7
8 FIG. 18 illustrates the Add Resource process which is used to introduce an
9 application's resource into PLN 33. The node at which the resource originates first
10 checks its local resource database to determine whether a resource with the same name
11 already exists (step 601). If such a resource does exist, the originating node returns an
12 ERROR to the user's application (step 602). If the resource does not exist, the
13 originating node adds an entry for the resource in its local database (step 603). The
14 resource then checks its persistent find query cache to determine whether an application
15 executing at the node is waiting for a resource (step 604). If the new resource matches
16 any of the search criteria in the persistent find cache, then the originating node sends
17 the new resource's attributes to the originating user's application that initiated the
18 Persistent Find Query (step 605). The originating node then removes from the cache
19 the Persistent Find Query for which the new resource matched the search criteria (step
20 606). If the scope of the newly removed persistent find query is greater than machine
21 level (step 607), then the originating node sends a Clean Persistent Find Query to its
22 parent node (step 608). At the end of the Persistent Find processing, or if no Persistent
23 Find Query was matched by the new resource, the originating node sends an add
24 resource request to its parent Area Manager (step 609).

25
26 If an Area Manager receives an add resource request from one of its children
27 (step 610), the Area Manager adds the resource to its own separate resource cache (step
28 611). The Area Manager then checks its own persistent find cache to determine

1 whether the new resource matches any of the criteria of a query in the cache (step 612).
2 If so, the Area Manager sends the resource's attributes to the node that originated the
3 Persistent Find Query (step 613) and removes the Query from its persistent find cache
4 (step 614). If the scope of that Query is greater than area level (step 615), then the
5 Area Manager sends a Clean Persistent Find Query to its parent Group Manager (step
6 616).

7 8 *Find Resource Process*

9
10 An application searching for a resource within the PLN 33 may specify one of
11 three different options for the Find Query which it sends to the PIPES Kernel: Find,
12 Find Next, or Persistent Find. A Find Query will begin searching for resources at the
13 local machine, moving to the area level if no resources are found at the machine level.
14 If no resources are found at the area level, the search continues at the group level, and
15 so on up the PIPES network hierarchy. If a resource is found at a particular level, that
16 resource's attributes are sent to the application requesting the resource. If the
17 application later issues a Find Next Query, the search will continue where the previous
18 search had left off within the PIPES hierarchy.

19
20 If the user issues a Persistent Find Query, the originating node first converts it
21 into a regular Find Query, which travels the network just like any other Find Query. If
22 any resource is returned to the user, the Find Query will not persist within the network;
23 however, if no resource is found within the PIPES hierarchy, the Persistent Find Query
24 is stored within the PIPES hierarchy in the Area Managers' persistent find caches.

25
26 FIG. 19 depicts the Find Resource process as it executes at the originating node.
27 If a Find or Persistent Find Query is initiated, the originating node clears a resource
28 cache which is used as a buffer to store the resource attributes satisfying the query's

1 search criteria (step 701). Because a Find Query is completely coordinated by the
2 originator of the query, and no state is maintained at any of the intermediate nodes,
3 each query data packet must carry sufficient information to enable the intermediate
4 nodes to conduct their searches. Some of the most important pieces of information is
5 the originating node's location within the network, the maximum number of matches
6 that is desired by the originating node (*MaxMatches*), the current number of matches
7 that have been returned to the originating node (*CurrMatches*), the scope of the search
8 (*Scope*), the level at which the search was last conducted (*Level*), and the status of the
9 last search at that level (*Level Status*). When the search begins with a Find Query or a
10 Persistent Find Query, the originating node initializes some of these variables to begin
11 the search at the machine level (step 702). Because a Find Next Query is designed to
12 begin the next search where the previous search left off, a Find Next Query causes the
13 originating node to skip these initialization steps.

14
15 The originating node compares *CurrMatches* to *MaxMatches* to determine
16 whether the user has already received the maximum number of matches for which it
17 asked (step 703). If *CurrMatches* is not equal to *MaxMatches* (*CurrMatches* can never
18 exceed *MaxMatches*), then the originating node checks its resource to see if any more
19 resources are available to return to the user (step 704). Resources may be left over in
20 the local cache because although a distributed Find Query may return more than one
21 resource to the originating node, the originating node returns resources to the user one
22 at a time. If there are resources left in the local cache, the originating node returns the
23 first resource to the user (step 705). If the resource cache is empty, the originating
24 node checks the *Level Status* to determine where the last search left off (step 707).
25 *Level Status* is set to EOF (i.e., end of find) if there are no resources available at that
26 level. If the *Level Status* is EOF, the originating node increments *CurrLevel* to
27 continue the search at the next level of the hierarchy (step 710). If the *Level Status* is
28 not EOF, the originating node checks *CurrLevel* to determine whether to begin the

1 search at the local machine before beginning a distributed search (step 708). If
2 *CurrLevel* is set to Machine, the originating node searches its local resource database to
3 see if local resource may match the search criteria (step 709). If a local resource is
4 available, the originating node copies up to *MaxMatches* resources' attributes to the
5 query's resource cache, and sets *CurrMatches* to the number of matches found and
6 copied to the cache (step 706). The originating node then returns the first resource
7 from the cache to the user that requested the resource (step 705). If no local resources
8 are found, the originating node sets the *Level Status* to EOF (step 711), and then
9 increments *CurrLevel* to continue the search at the next level (step 707).

10

11 If *CurrLevel* exceeds *MaxLevel* (step 712) or *Scope* (step 716), then search has
12 either worked its way through the complete PIPES hierarchy or exceeded the scope of
13 the original query. Thus, if either of these conditions have been met, the search is
14 complete. If not, the originating node sends the Find Query to its parent, the Area
15 Manager to begin the distributed search (step 713). If resources' attributes are returned
16 in response (step 714), the originating node copies the resources' attributes to the
17 query's resource cache (step 718) and returns the first to the user (step 717). If the
18 search completes unsuccessfully, the originating node checks *CurrMatches* to see if any
19 resources have been returned to the user (step 715). If *CurrMatches* is greater than
20 zero, then the user has received all of its resources, and the originating node returns an
21 EOF to the user (step 723). If *CurrMatches* is zero, and no resources were found on
22 the network, the originating node distributes a Persistent Find Query if the user has so
23 specified (step 719). This entails adding the query to a listing of Persistent Find
24 Queries pending at the node in order to keep track of the sources of the Persistent Find
25 Queries (step 720). If a resource existing at the local machine could possibly match the
26 search criteria of the Query (step 721), the originating node adds the query to its
27 persistent find cache (step 722), which is used to keep track of the search criteria so
28 that resources that meet those criteria may be returned as soon as they are added to

1 PIPES. If the scope of the query is greater than machine level (step 724), then the
2 Persistent Find Query is send to the Area Manager (step 725).

3
4 FIGS. 20 and 21 illustrate how the Resource Layer routes a Find Query
5 throughout PLN 33. FIG. 20 shows the process which is executed at the Area Manager
6 level. When the Area Manager receives a Find Query (step 726), the Area Manager
7 checks *CurrLevel* to determine the level at which a search is requested (step 727). If
8 *CurrLevel* is less than Area (step 728), then the Area Manager returns an error to the
9 node that sent the Find Query because the Area Manager received the query by mistake
10 (step 729). If *CurrLevel* is greater than Area (step 728), the Area Manager will
11 forward the Find Query to its parent (step 732) if the Area Manager received the Find
12 Query from one of its children (step 731). Thus, the Area Manager is just passing on
13 the Find Query because the search should continue at a higher level of the hierarchy.
14 If the search should continue at this level, the Area Manager analyzes the search
15 criteria to determine whether a resource in this area could satisfy the criteria (step 730).
16 If not, the Area Manager returns the Find Query to the sender (step 738). In addition,
17 if *CurrMatches* is already equal to *MaxMatches* (step 733), the Area Manager also
18 returns the Find Query to the sender (step 738). Otherwise, the Area Manager searches
19 its resource database looking for a match that is visible to the originating node (step
20 734). The user that adds a resource to PIPES can specify which applications can utilize
21 its services, or its "visibility" within PIPES. If visible matches are found, a maximum
22 of *MaxMatches* resources' attributes are copied to the Find Query (step 735). If more
23 than *MaxMatches* resources are found (step 737), the Area Manager sets the *Level*
24 *Status* to OK (step 739) so that the search will continue at this level the next time a
25 Find Next Query is issued. Otherwise, the Area Manager sets the *Level Status* to EOF
26 to notify the originating node that no more resources are available at this level (step
27 736). Finally, the Area Manager returns the Find Query to the sender (step 738).

28

1 The Find Query Process at managerial levels higher than Area Manager in the
2 PLN hierarchy (FIG. 21) is similar to that at the Area Manager level, except that no
3 searching occurs because only machines and Area Managers possess resources
4 databases. Steps 740 through 747 in FIG. 21 are the same as steps 726 through 733 in
5 FIG. 20. In each case, the node determines whether the search should continue at this
6 level or at a higher level. In this case, a search at this level consists of forwarding the
7 Find Query to each of the manager's children in turn. If any more children have not
8 yet seen the Find Query (step 748), the manager sends the Find Query to the next child
9 (step 749). When no more children are left, the manager sets the *Level Status* to EOF
10 (step 751) and returns the Find Query to the sender (step 750).

11
12 FIGS. 22 and 23 illustrate the process of adding a Persistent Find Query
13 throughout the network, and FIGS. 24 and 25 depict a similar "clean-up" process used
14 to remove a Persistent Find Query from the network. In FIG. 22, an Area Manager
15 node processes a Persistent Find Query received over PLN 33 (step 752). First, if the
16 Area Manager received the Query from one of its children (step 753), the Area
17 Manager adds the query to its source list of pending persistent finds (step 754). If a
18 resource in this area could satisfy the Persistent Find Query's search criteria (step 755),
19 then the Area Manager adds the query to its persistent find cache. If the *Scope* of the
20 Query is greater than Area level (step 757), the Area Manager sends the Persistent Find
21 Query to its parent (step 758). Similarly, in FIG. 23, a manager at a level higher than
22 Area receives a Persistent Find Query (step 759). If the sender is one of the manager's
23 children (step 760), the manager adds the Query to its source list of pending persistent
24 finds (step 761). If this level is within the search criteria specified in the Query (step
25 762), the manager forwards the Query to its children (except possibly the child that sent
26 the Query) (step 763). If the *Scope* of the Query is greater than this level (step 764),
27 then the manager sends the Persistent Find Query to its parent (step 765).

1 Similar processes are illustrated in FIGS. 24 and 25 that "clean-up" Persistent
2 Find Queries by removing them from nodes' source lists of pending persistent finds
3 (steps 768 and 775) and removing them from Area Managers' persistent find caches
4 (step 770).

5
6 *Persistent Find Recovery Process*

7
8 Because important information about distributed Persistent Find Queries is kept
9 at the Area Manager nodes, and to a lesser extent at the other managerial nodes, a
10 recovery process must be used when one of these nodes crashes or becomes
11 disconnected from the rest of the PLN hierarchy. FIGS. 26 and 27 represent the
12 processes used to provide recovery when the Area Manager (FIG. 26) or another
13 managerial node (FIG. 27) goes down.

14
15 When a machine logs in to its new parent Area Manager, selected by the
16 Election Process, the child machine sends its source list of pending persistent finds to
17 its new parent (step 800). The new Area Manager receives this list (step 801) and
18 updates its own source list of pending persistent finds using the information received
19 from its children (step 802). The new Area Manager then sends a replenish cache
20 request to its parent (step 803). The other managers receive the request (step 805) and
21 send it to all of its children in the manager's source list of pending persistent finds (step
22 806). If the sender is the manager's child (step 807), the manager sends the request up
23 the PLN hierarchy to its parent (step 808). Eventually, the other Area Managers in
24 PLN 33 receive the replenish cache request (step 809), and if the new Area Manager
25 has a Query in its persistent find cache (step 810), the receiving Area Manager replies
26 to the new Area Manager with matching queries from its persistent find cache (step
27 811). The new Area Manager then updates its own Persistent Find Cache with the
28 replies from other Area Managers in PLN 33 (step 804).

1 FIG. 27 describes the situation that exists when a manager other than an Area
2 Manager goes down. The new manager's children send their source lists of pending
3 persistent finds to the new manager (step 812). The new manager receives these lists
4 (step 813) and update its list of pending persistent finds with the information sent from
5 its children (step 814). If any of the queries are scoped higher than this level (step
6 815), then the queries are sent up the PLN hierarchy to the new manager's parent (step
7 816). The new manager's parent verifies its source list of pending persistent finds with
8 the information obtained from its new child (step 817).

9 10 *Remove Resource Process*

11
12 When an application withdraws its resources from the PLN hierarchy, Resource
13 Layer 33 executes the Remove Resource Process illustrated in FIG. 28. The node at
14 which the resource originated first check to see if the resource exists in its resource
15 database (step 901). If the resource exists, the originating node removes the resource
16 from the database (step 903) and sends the remove resource request to its parent Area
17 Manager (step 904). If not, the originating node returns an error to the user (step 902).
18 The Area Manager receives the remove resource request (step 905) and removes the
19 resource from its area manager resource cache (step 906).

20 21 Context Bridge Layer

22
23 FIG. 29A illustrates the components of Context Bridge Layer 38. The main
24 function of Context Bridge Layer is the Routing Process (block 1000), which routes a
25 Protocol Data Unit ("PDU") from a source node to a destination node. The source
26 node and the destination node may share a routable protocol. A routable protocol is
27 defined as a protocol that allows a decision about where a PDU must be sent in order
28 to reach its destination to be made solely from the destination address. The source

1 node merely transfers the PDU to the routable protocol, and the routable protocol itself
2 determines how to get the PDU to its destination by parsing the destination address.
3 Thus, no knowledge of the intermediate nodes used to forward a PDU from the source
4 to the destination is necessary. Within PIPES, TCP/IP and SNA are routable
5 protocols, whereas IPX, NetBios and DLC are non-routable protocols.
6

7 If the source node and the destination node share a non-routable protocol, or if
8 the source and destination do not share any protocol at all, intermediate nodes must be
9 used to "bridge" the source and destination nodes. In this case, the Routing Process
10 uses the Routing Information Database ("RIDB", shown as block 1400) to determine
11 how to route a PDU from source to destination. The RIDB contains the information
12 necessary to route a PDU to a non-routable protocol or to a protocol that the source
13 node does not support. The RIDB contains two caches: a source routing cache (block
14 1401) is used for non-routable protocols, and a next-hop routing cache (block 1402) is
15 used for dissimilar protocol bridging. The source routing cache is populated through
16 the Route Discovery Process (block 1100) and is validated through the Route Validation
17 Process (block 1200). The next-hop routing cache is populated through the Route
18 Advertisement Process (block 1300).
19

20 FIG. 29B illustrates a system 1600 in which the context bridge of the present
21 invention can be advantageously used. The context bridges can be used to route
22 packets generated by nodes using protocols of different levels, as defined in the
23 International Organization of Standardization ("ISO") Reference Model. For example,
24 system 1600 contains two nodes 1610 and 1630 which use the SNA (APPC) and DLC
25 protocols, respectively. These two protocols are at different ISO levels: the SNA is at
26 the presentation level while the DLC is at the data link level. In order to route packets
27 from node 1610 to node 1630 through a network 1640, it is necessary to use a node
28 1620 containing a context bridge which can bridge the SNA (APPC) and DLC

1 protocols. Thus, the packet generated by node 1610 is first routed to node 1620 via
2 path 1642, which then routes the packet to node 1630 via path 1643.

3
4 Similarly, if it is desirable to route a message generated by node 1610 to a node
5 1650 which uses the UDP protocol (at ISO transport level), it is necessary to use a
6 node 1660 containing a context bridge which can bridge the SNA and UDP protocols.
7 Thus, the packet generated by node 1610 is first routed to node 1660 via path 1645,
8 which then routes the packet to node 1650 via path 1646.

9
10 *Routing Process*

11
12 FIG. 30 depicts a flowchart of the Context Bridge Routing Process. When the
13 source node's Context Bridge Layer receives a PDU to be sent to a given destination
14 node, the source node looks at the destination address to determine whether the
15 destination has a routable protocol (step 1001).

16
17 If the destination has a routable protocol, the source node determines whether or
18 not it supports the same routable protocol as the destination (step 1002). If the source
19 and destination share the same routable protocol, the source sends the PDU to the
20 destination using the transport driver for the shared routable protocol (step 1003). If
21 the source and destination do not share the same routable protocol, the source searches
22 its RIDB next-hop routing cache for a route to the destination (step 1004). The source
23 node then checks to see whether a route exists in the RIDB (step 1006). If a route is
24 found, the source sends the PDU to the intermediate node specified by the route found
25 in the RIDB (step 1007). If a route is not found, the source returns an error stating
26 that the destination is not reachable (step 1009).

1 If the destination has a non-routable protocol, the source searches its RIDB
2 source routing cache for a route to the destination (step 1005). The source node then
3 checks to see whether a route exists in the RIDB (step 1008). If a route is found, the
4 source sends the PDU to the intermediate node specified by the route found in the
5 RIDB (step 1007). If a route is not found, the source executes the Route Discovery
6 Process to find a route to the destination (step 1011). The source node then ascertains
7 whether a route was found by the Route Discovery Process (step 1012). If a route was
8 found by Route Discovery, the source node updates its RIDB source routing cache (step
9 1010), and sends the PDU to the intermediate node specified by the route (step 1007).
10 If a route was not found, the source node returns an error that the destination is not
11 reachable (step 1009).

12 13 *Route Discovery Process*

14
15 FIG. 31 describes the Route Discovery Process, which is used to update the
16 RIDB source routing cache with source routes to individual destinations. A source
17 node initiates the Route Discovery Process when a route to a destination with a non-
18 routable protocol needs to be found. First, a source node sends a Route Discovery
19 Packet to all of the active context bridges about which it has information (step 1101).
20 A node is an active context bridge if it supports more than one protocol; the node acts
21 as a bridge between the protocols found at that node. All of the nodes in the network
22 find out about active context bridges through the Route Advertisement Process.

23
24 A context bridge that receives the source node's Route Discovery Packet first
25 determines whether it is a reply packet (step 1107). If it is a reply packet, the
26 intermediate node forwards the packet back to the source node using the route specified
27 in the reply packet (step 1112). If it is not a reply packet, the node receiving the Route
28 Discovery Packet inserts its own address into the packet (step 1108). The node then

1 checks to see if it is the intended destination of the packet (step 1109). If the node is
2 the intended destination of the packet, the end node changes the type of the packet to
3 **REPLY** (step 1111), and forwards the packet back to the source using the route
4 specified in the Route Discovery Packet (step 1112). If the receiving node is not the
5 destination, the intermediate node forwards the packet to all context bridges to which it
6 is connected except the context bridge from which it originally received the packet (step
7 1110).

8
9 The source node is waiting to see if a reply is received (step 1102). If no reply
10 is received within a specified time period, the source returns an error that the
11 destination is unreachable (step 1103). If a reply is received, the source node checks if
12 there is already a valid route to the destination (step 1104). If there is already a valid
13 route, the source discards the reply packet (step 1105). Otherwise, the source node
14 updates its RIDB source routing cache with the route specified in the reply packet (step
15 1106).

16 17 *Route Validation Process*

18
19 FIG. 32 illustrates the Route Validation Process, which is used to check the
20 validity of the routes contained in the RIDB source routing cache. The source node
21 sends a Route Validation Packet to all of the destination nodes in its RIDB source
22 routing cache that have not been marked as valid (step 1201). The source then sets a
23 timer (step 1202) and listens for validation replies (step 1203).

24
25 The end nodes also listen for Route Validation Packets (step 1209) and checks to
26 see if a Validation Packet is received (step 1210). If a Validation Packet is not
27 received within a specified time period, the end nodes continue listening for Route
28 Validation Packets (step 1209). If a Validation Packet is received, the end nodes

1 validate the route specified in the Route Validation Packet (step 1211) and return the
2 Packet to the sender (step 1212).

3
4 The source node checks to see whether a validation reply has been received
5 (step 1204). If a validation reply is received, the source node marks the source route to
6 the destination as valid in the RIDB source routing cache (step 1205). If a validation
7 reply is not received, the source node checks the timer (step 1206). If the timer has
8 not expired, the source node continues to listen for validation replies (step 1203). If
9 the timer has expired, the source node will reset the timer (step 1202) if the retry
10 threshold has not been exceeded (step 1207). If the retry threshold has been exceeded,
11 the source node removes the invalid source route from the RIDB source routing cache
12 (step 1208).

13

14 *Route Advertisement Process*

15

16 FIG. 33 represents the Route Advertisement Process, a process which is
17 executed intermittently at every active context bridge and end node. Each context
18 bridge periodically sends a broadcast message known as a Routing Advertisement
19 Packet ("RAP") (step 1301), and each end node listens for RAP broadcasts (step 1305).
20 The RAP preferably contains the following information: the protocols that can be
21 handled by the context bridge and the number of hops required. All context bridges
22 and end nodes then wait until a RAP broadcast is received (steps 1302 and 1306). If a
23 RAP broadcast is received, the node receiving the broadcast determines if there is any
24 change in routing information by comparing the RAP broadcast with its RIDB next-hop
25 routing cache (steps 1303 and 1307). If changes are necessary, the receiving node
26 updates its RIDB next-hop routing cache (steps 1304 and 1308).

27

Unlimited Levels

In the preferred embodiment of the present invention, the number of levels in the PLN hierarchy is not limited. FIG. 34 illustrates the steps that is preferred taken by developer of system 100 (the system developer), the application developer, and the end user to implement a larger number of levels than the default number of levels (e.g., five). The maximum number of levels of a certain implementation is set when the PIPES kernel and PAPI library code is compiled. If it is desirable to have greater flexibility in their PIPES and greater number of levels in the hierarchy, the PIPES kernel and PAPI library need to be customized.

The system developer changes the *MinLevel* and *MaxLevel* parameters that are hard-coded in a header file of the software (step 1501). The PAPI library (step 1502) and PIPES kernel (step 1503) will be recompiled, and the new PAPI library and PIPES kernel are distributed to the application developer (step 1504).

The application developer receives these components from the system developer (step 1505) and makes any necessary modifications to their own PIPES application (step 1506). The application developer then recompiles its own PIPES application with the new PAPI library (step 1507) and distributes the new PIPES application and PIPES kernel to the end user (step 1508).

The end user receives these components from the application developer (step 1509) and installs them on all of the nodes in the PLN (step 1510). After making any necessary modifications to its PIPES configuration (step 1511), the end user finally restarts the system by loading the PIPES kernel (step 1512) and the PIPES application (step 1513). At this point, the end user can realize the number of levels desired in the PLN hierarchy.

1 While the present invention has been described with what is presently considered
2 to be the preferred embodiments, it is to be understood that the appended claims are not
3 to be limited to the disclosed embodiments, but on the contrary, are intended to cover
4 modifications, variations, and equivalent arrangements which retain any of the novel
5 features and advantages of the invention.

WHAT IS CLAIMED IS

1 1. A method for independently executing software components in a node of
2 a network containing a plurality of nodes, the method comprising the steps of:
3 generating a logical hierarchy of the roles of the nodes in the network wherein
4 any node can assume one or multiple roles, the assumption of which neither requires or
5 precludes the assumption of any other role; and
6 negotiating the role of the nodes when there is a change in the configuration of
7 the network.

1 2. The method of Claim 1 wherein a node having a managerial role leaves
2 the network and at least one of the remaining nodes participates in a negotiation process
3 of determining which node assumes the managerial role, the negotiating step further
4 comprising the steps, performed by the participating node, of:
5 broadcasting a message indicating the participating node's interest in assuming
6 the managerial role;
7 listening, subsequent to the broadcasting step, for messages on the network; and
8 assuming the managerial role if there is no message on the network which
9 indicates that another node is better qualified to assume the managerial role.

1 3. The method of Claim 1 further comprising the steps performed by the
2 participating node:
3 listening, prior to the said broadcasting step, for a specified period of time for
4 messages sent by other participating nodes; and
5 withdrawing from the process when the messages indicates that there is at least
6 one participating node which is more qualified to assume the managerial role.

1 4. The method of Claim 1 wherein at least two conflicting nodes claim the
2 same managerial role and at least one of the conflicting nodes participates in a process
3 of determining a node which assumes the managerial role, the negotiating step further
4 comprising the steps, performed by each participating node, of:

5
6 setting up a database containing the names of all known nodes participating in
7 the process;

8 transmitting election messages to nodes included in the data base, the election
9 messages containing information relating to the participating nodes known to the
10 sending node;

11 receiving election messages from other participating nodes;

12 updating the database using the information contained in the received election
13 messages; and

14 determining, based on the information contained in the updated database, which
15 one of the participating node assumes the managerial role.

1 5. The method of Claim 1 wherein one of the nodes is a parent node, the
2 method further comprising the step of searching for a parent node when a node enters
3 the network.

1 6. The method of Claim 5 wherein the searching step further comprises the
2 steps, performed by the entering node, of:

3
4 listening to messages for a specified period of time;

5 determining, if a message is received, the entering node's parent based on the
6 received message;

7 broadcasting, if no parent is found upon expiration of the specified period of
8 time, a message for searching its parent;

9 listening for responses to the broadcasted message and determining whether if
10 any one of the responses originates from its parent;
11 assuming the role as its own parent when no response is received.

1 7. The method of Claim 1 wherein one of the nodes is a parent node, the
2 method further comprising the step of registering a child upon its entering the network.

1 8. The method of Claim 7 wherein said registering step further comprises:
2 listening to messages send by entering nodes;
3 determining whether one of the messages is sent by a child node or a duplicate
4 child node;
5 if a duplicate child node is detected, informing the duplicate child role of a role
6 conflict;
7 if a child node is detected, sending an acknowledge message to the child node.

1 9. The method of Claim 1 wherein one of the nodes is a parent and one of
2 the remaining nodes is a child, said method further comprising the step of monitoring
3 the status of the parent and the child.

1 10. The method of Claim 9 wherein the monitoring step further comprises
2 the step of:
3
4 exchanging status messages between the parent and the child at specified time
5 intervals;
6 searching for a new parent node when the child does not receive status messages
7 from the parent within a predetermined period of time.

1 11. The method of Claim 9 further comprising the step of de-registering the
2 child when the parent does not receive status message from the child within a
3 predetermined period of time.

1 12. The method of Claim 1 wherein the logical hierarchy consists of an
2 arbitrary number of levels.

1 13. The method of Claim 12 wherein the number of levels is changeable.

1 14. The method of Claim 1 wherein the roles of the nodes in the network is
2 changeable contingent on the requirements of the network.

1 15. A method for finding a resource requested by a node in a scalable system
2 interconnecting a plurality of nodes on a digital network, each node being associated
3 with one or more resource, each resource having an active state in which it is available
4 to other nodes and an inactive state in which it is not available, the method comprising
5 the steps of:

6 automatically identifying resources that join the network by switching from the
7 inactive to the active state; and

8 automatically informing the requester that the requested resource has become
9 available.

1 16. The method of Claim 15 wherein at least one of said plurality of nodes
2 contain a cache for storing a list of requests for resources, said method further
3 comprising the step of storing a request for resources in the cache.

1 17. The method of Claim 16 further comprising the step of removing a
2 request from the cache when a resource becomes active and satisfies the request.

1 18. The method of Claim 15 wherein said plurality of nodes are arranged in
2 at least two levels, and wherein nodes in a first level contain information relating to the
3 resources present in nodes in a second level.

1 19. The method of Claim 18 wherein a plurality of nodes in the first level
2 contain a cache for storing a list of requests for resources, said method further
3 comprising the step of storing a request for resources in the cache.

1 20. The method of Claim 19 further comprising the step of removing a
2 request from the cache when a resource becomes active and satisfies the request.

1 21. The method of Claim 18 wherein the first level is a parent level and the
2 second level is a child level, each node in the parent level being associated with one or
3 more nodes in the child level and each node in the child level being associated with one
4 node in the parent level.

1 22. A method for determining routing paths in a context bridge which is able
2 to route packets between nodes having different communication protocols at different
3 levels, the context bridge being one of many context bridges in a heterogeneous
4 network containing a plurality of nodes, the method comprising the steps of:

5 Setting up a list of context bridges and the communication protocols handled by
6 each context bridge in the list;

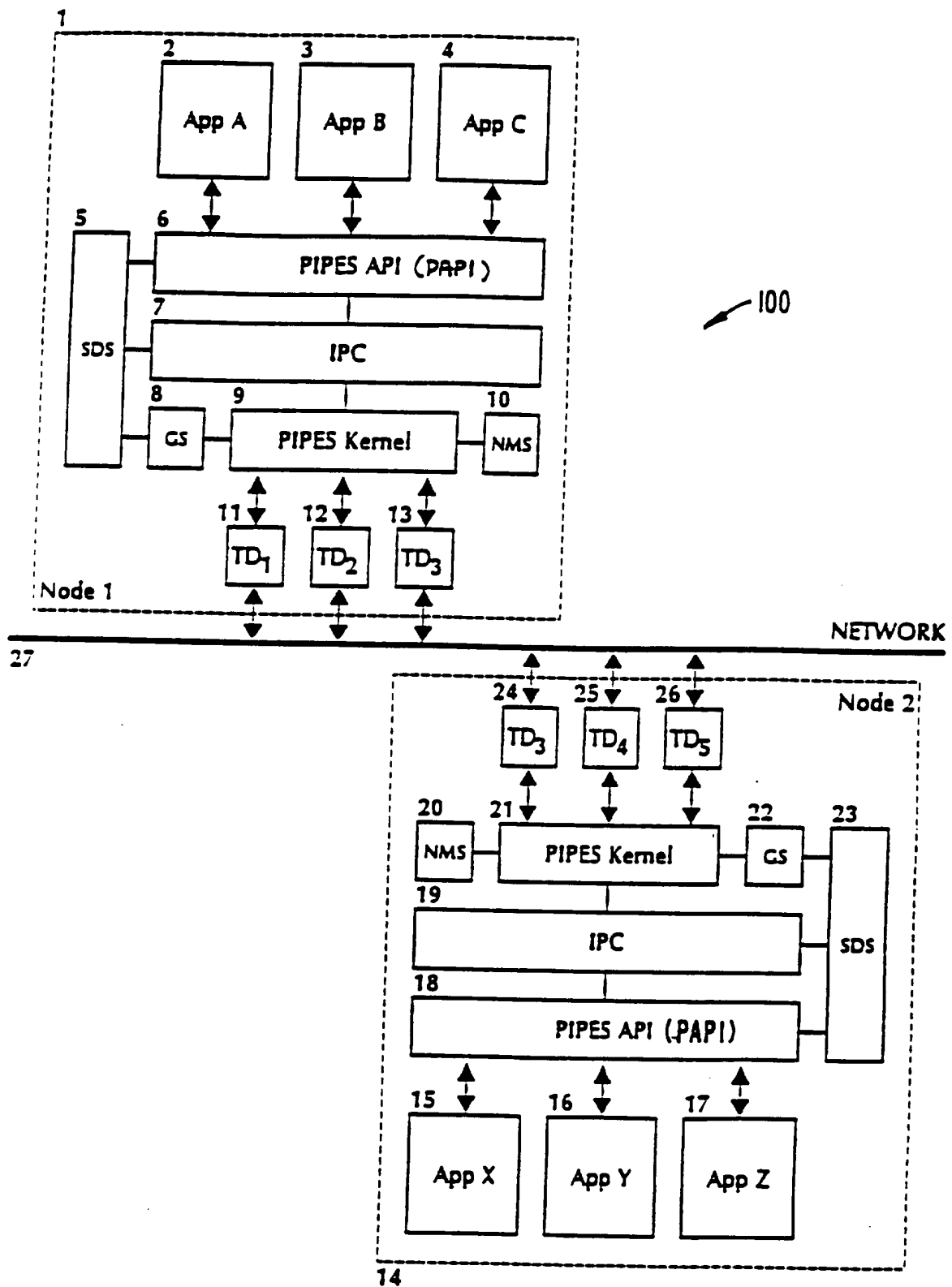
7 listening for routing information packets which are periodically broadcast by
8 other context bridges informing recipients the communication protocols handled by the
9 broadcasting context bridge; and

10 updating the database using the information contained in the received routing
11 information packets.

1 23. A method for routing packets from a source node to a destination node
2 using a context bridge, the routing protocol of the source node being at a different level
3 as the routing protocol of the destination node, the context bridge being one of many
4 context bridges in a heterogeneous network containing a plurality of nodes, the method
5 comprising the steps of:
6 determining whether the destination node has a routable protocol;
7 routing the packet using a source route containing the addresses of the context
8 bridges from the source to the destination, if the destination has no routable protocol,
9 the addresses being discovered by the source sending a route discovery packet and the
10 destination replying to the discovery packet; and
11 routing the packet through at least one context bridge containing routable
12 protocols if the destination has routable protocol.

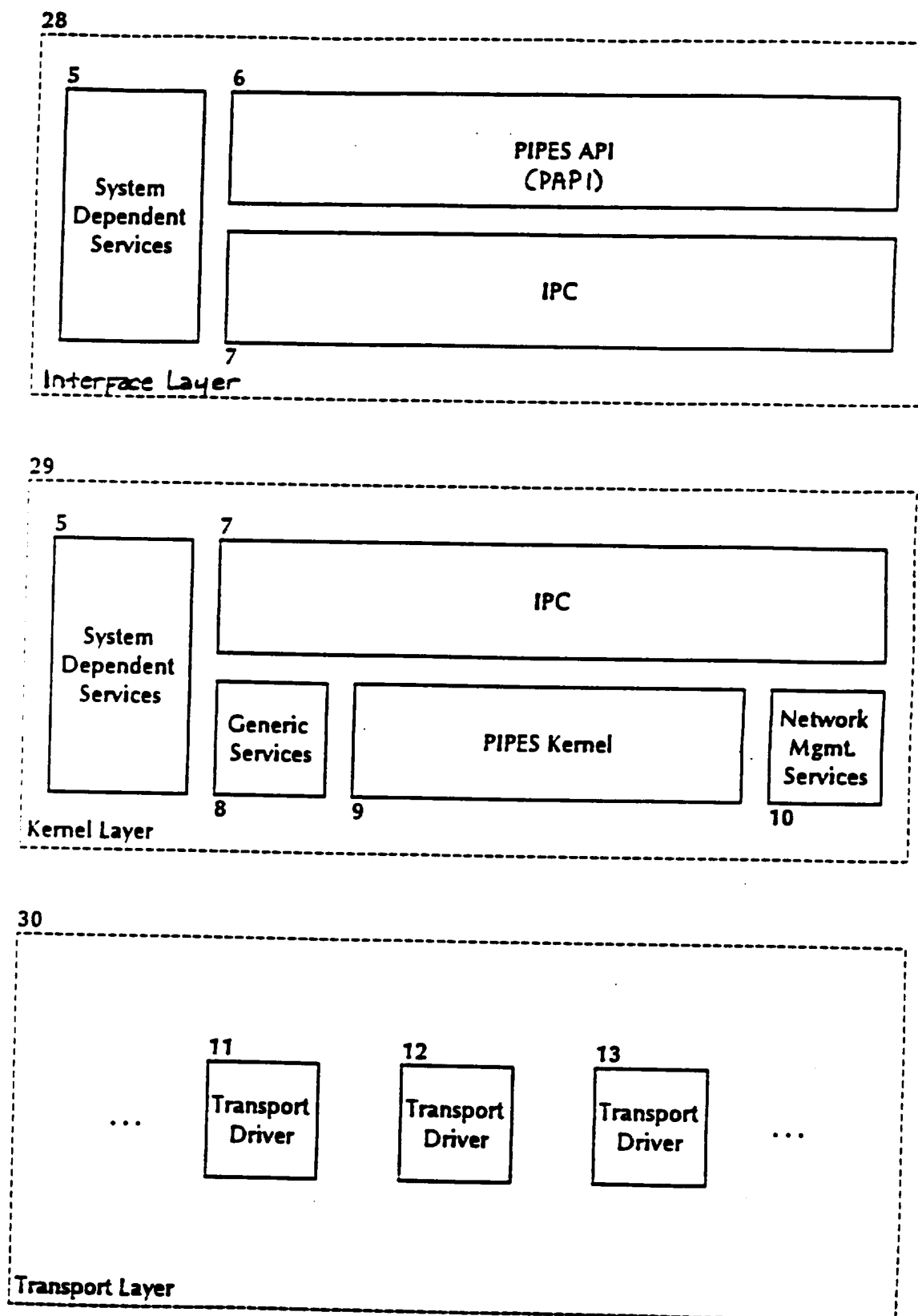
1/34

FIG. 1



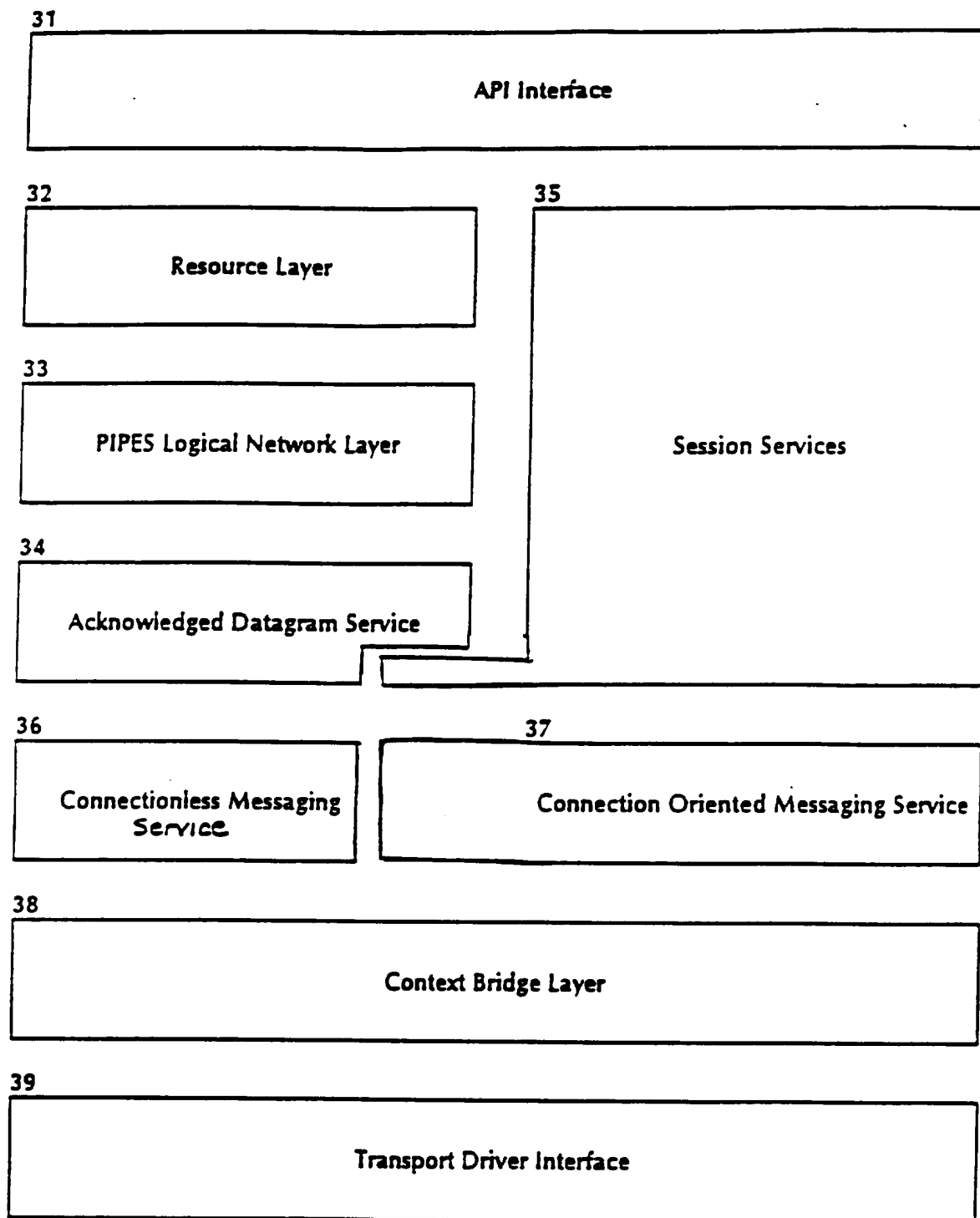
2/34

FIG. 2



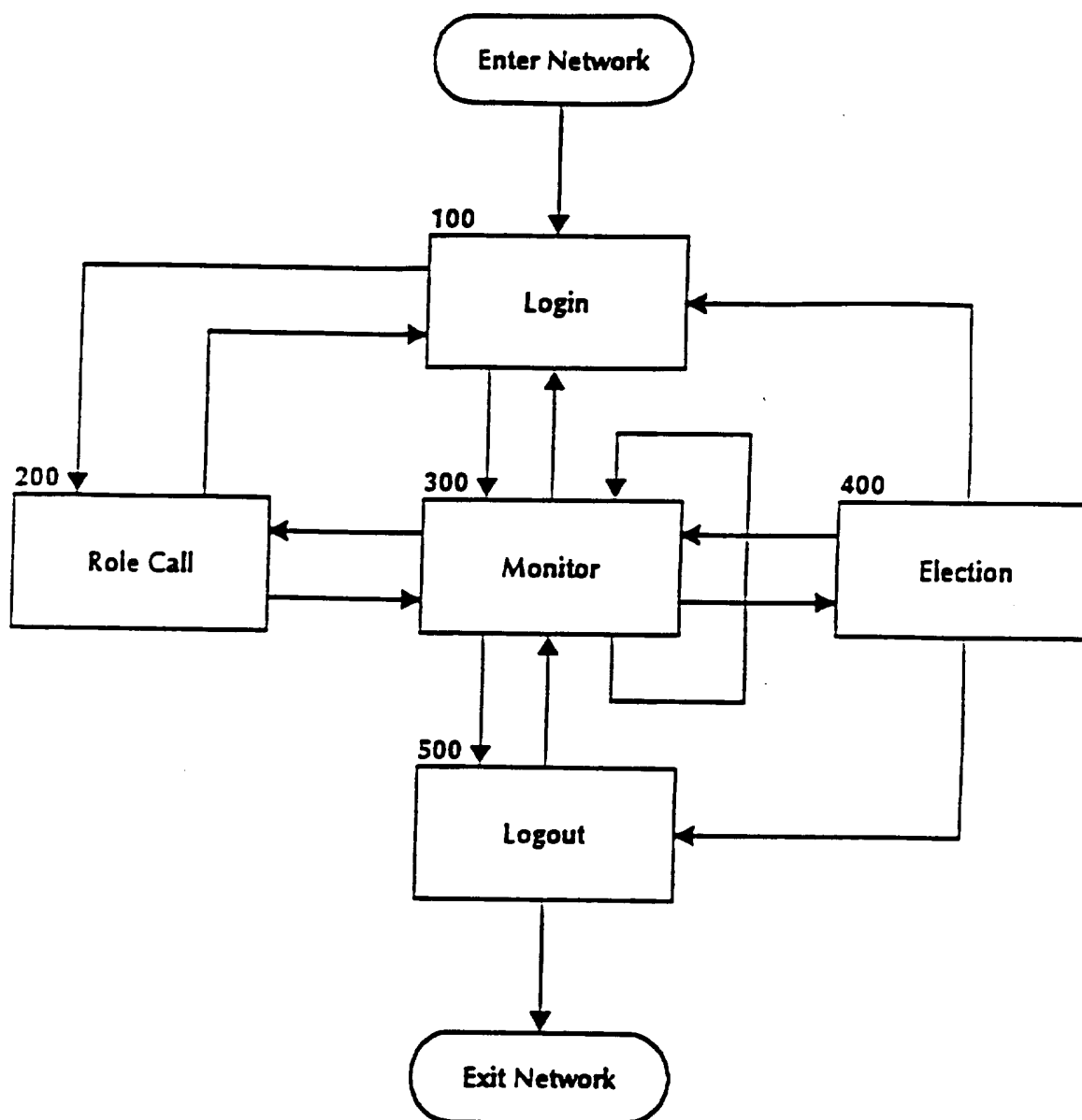
3/34

FIG. 3



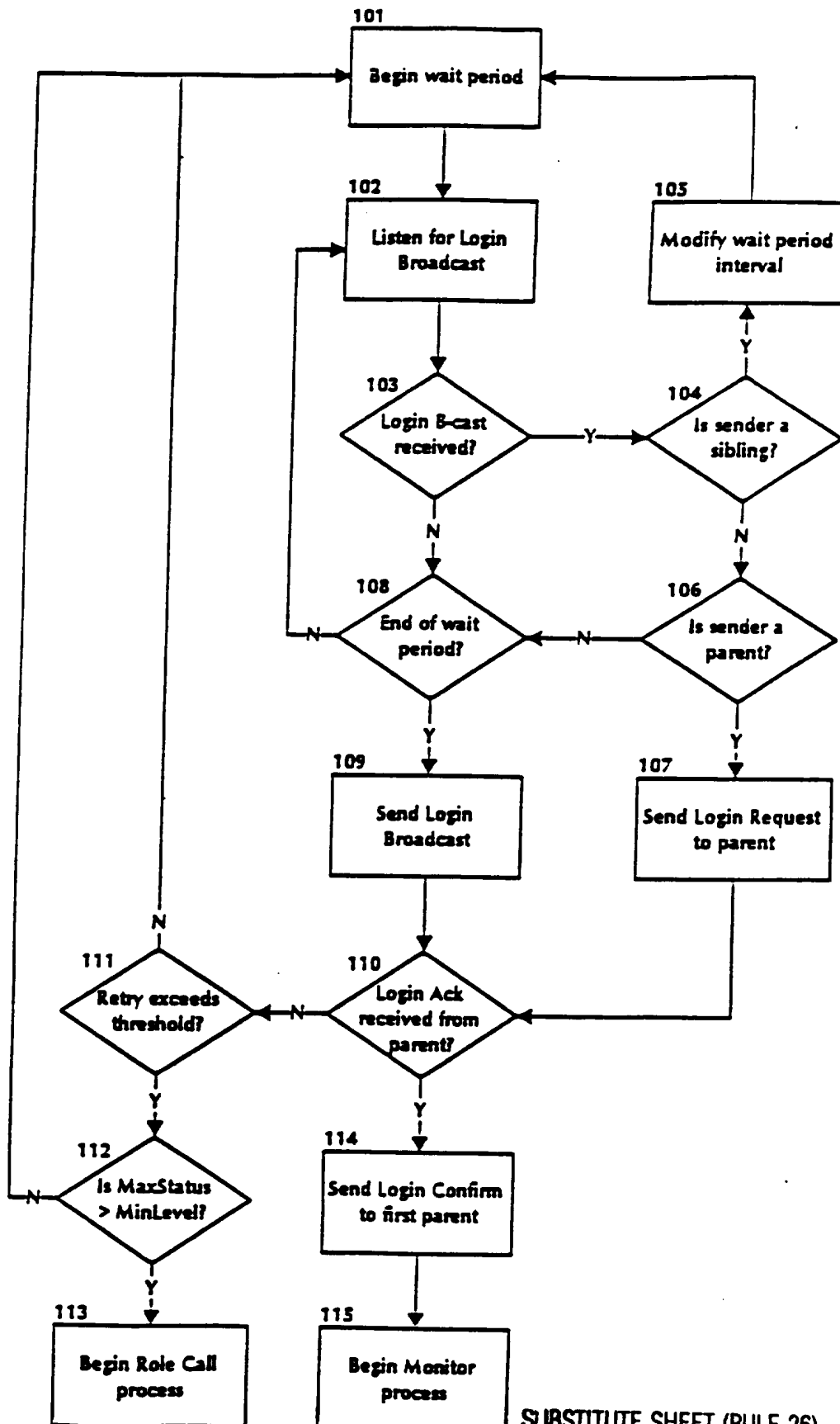
4/34

FIG. 4



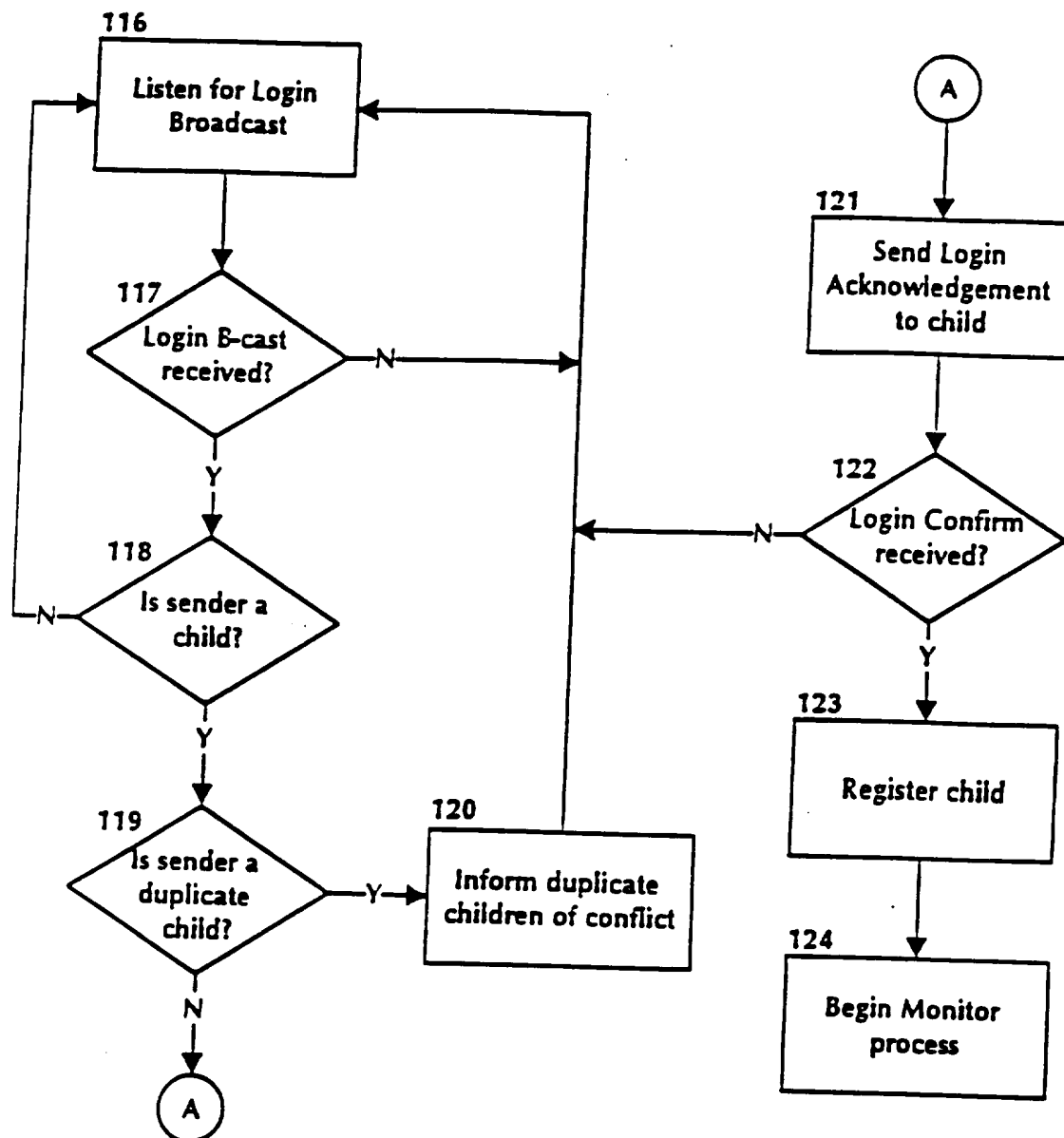
5/34

FIG. 5



6/34

FIG. 6



7/34

FIG. 7

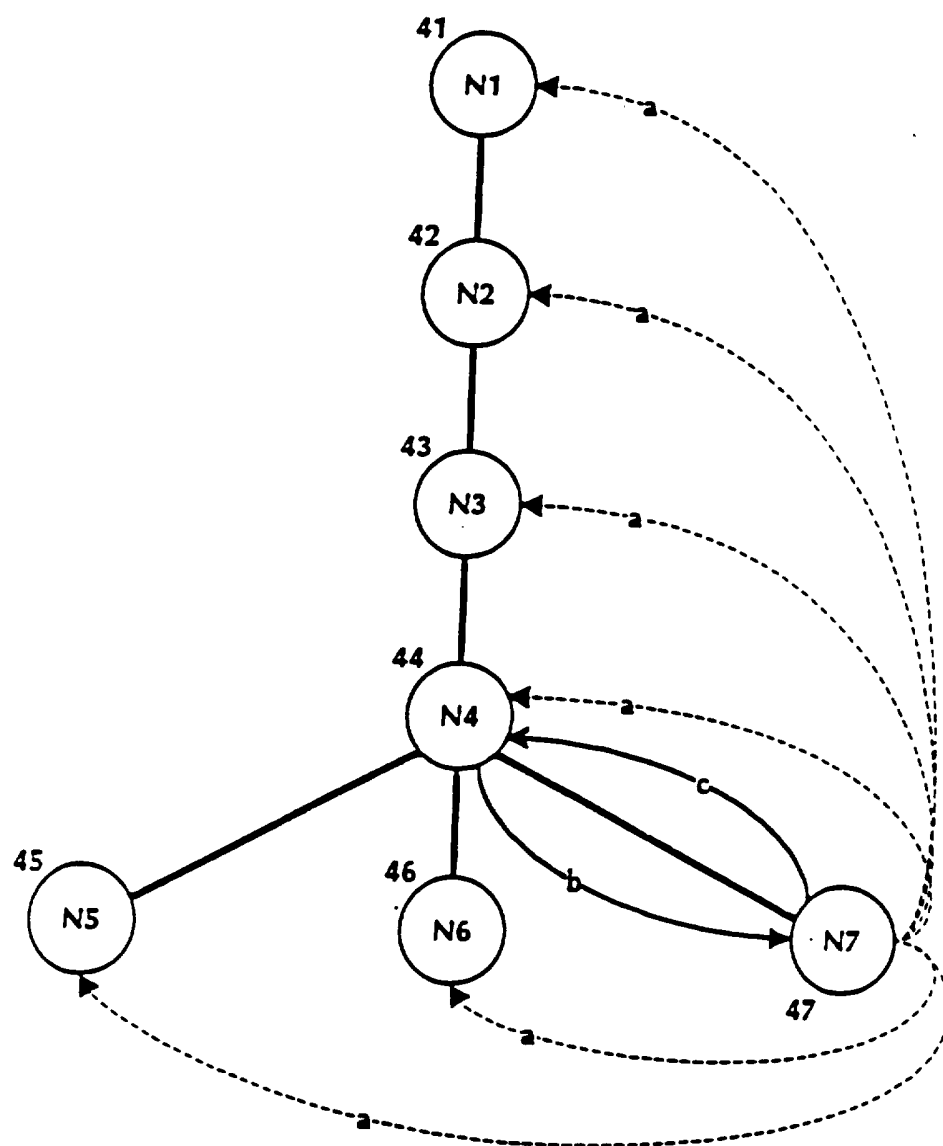
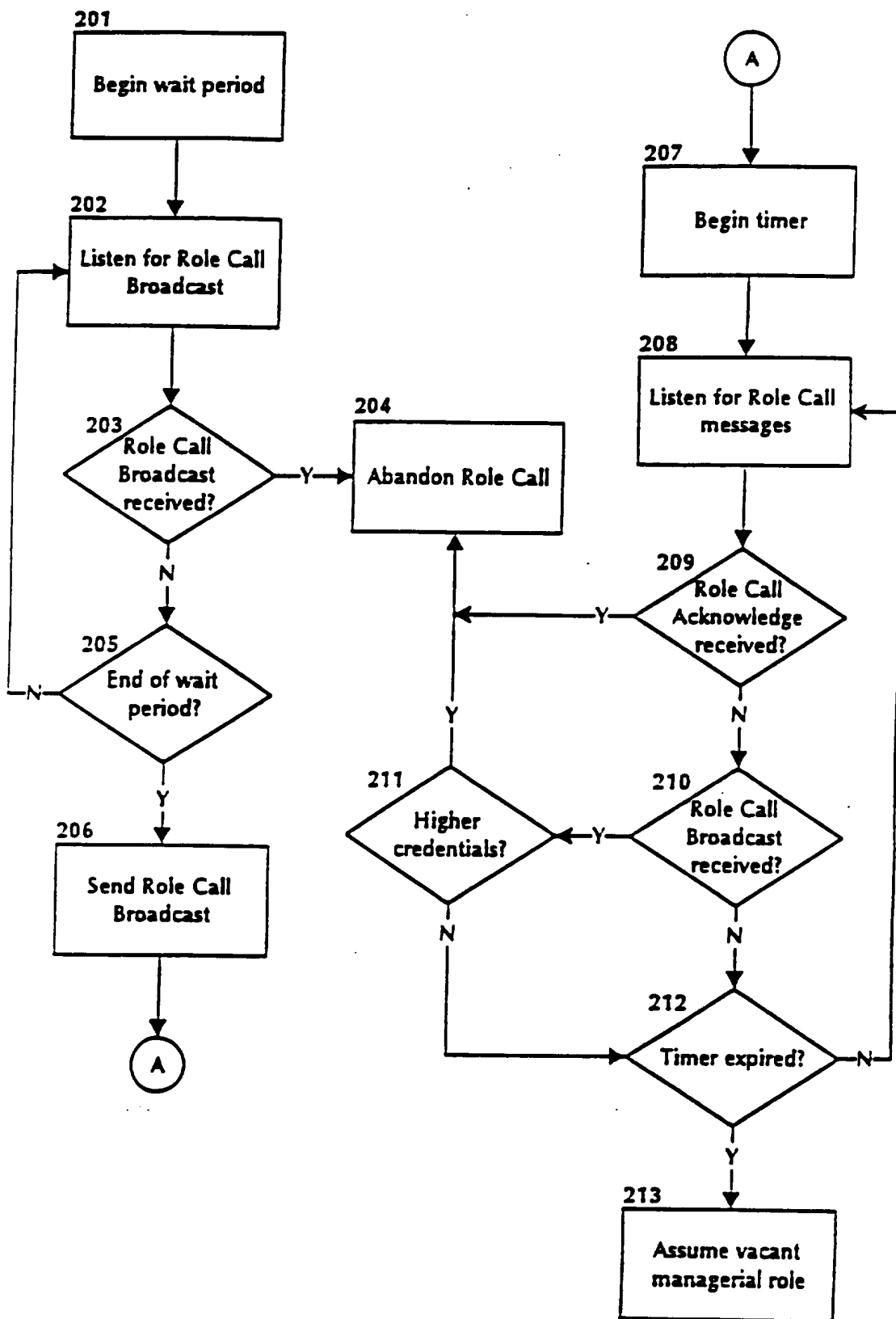
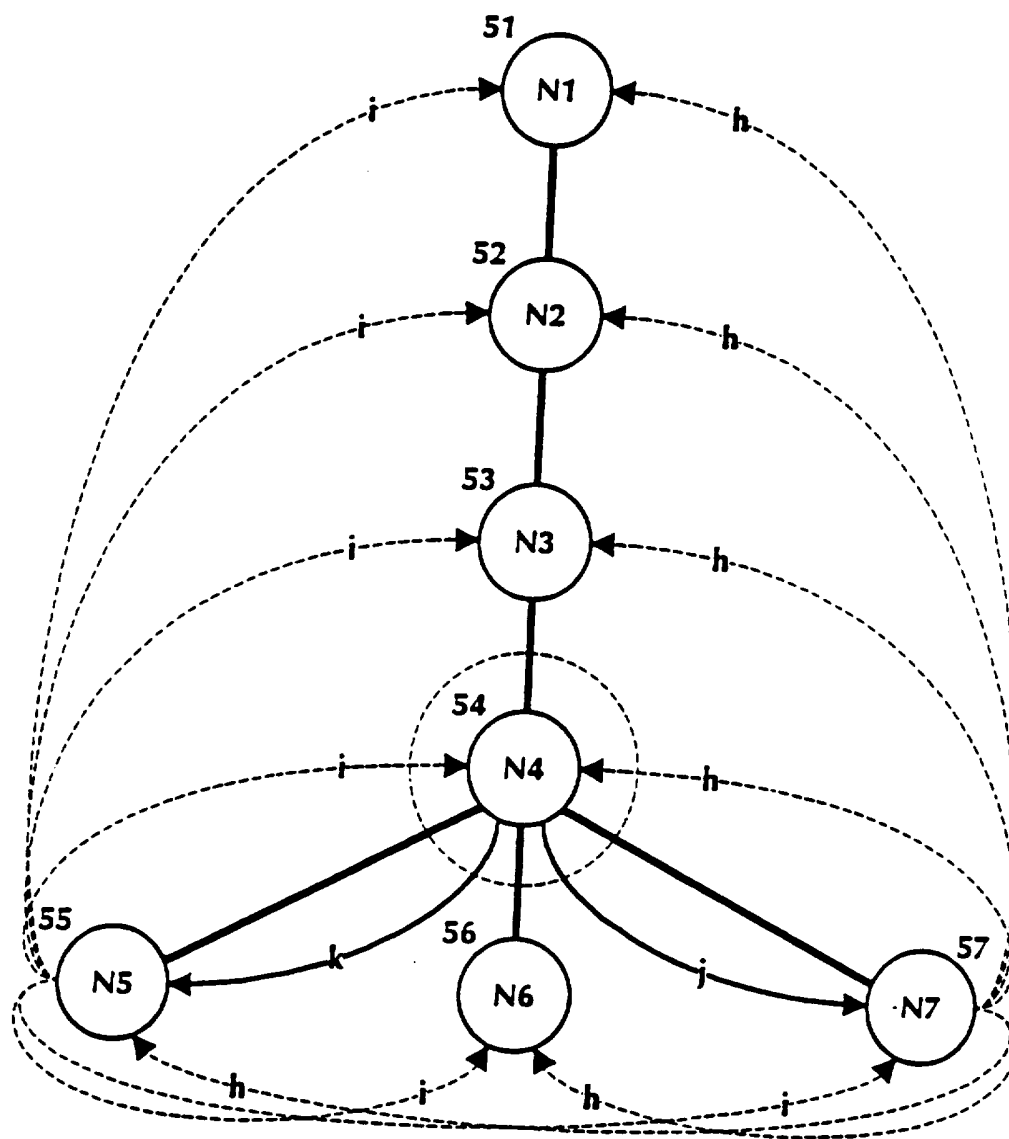


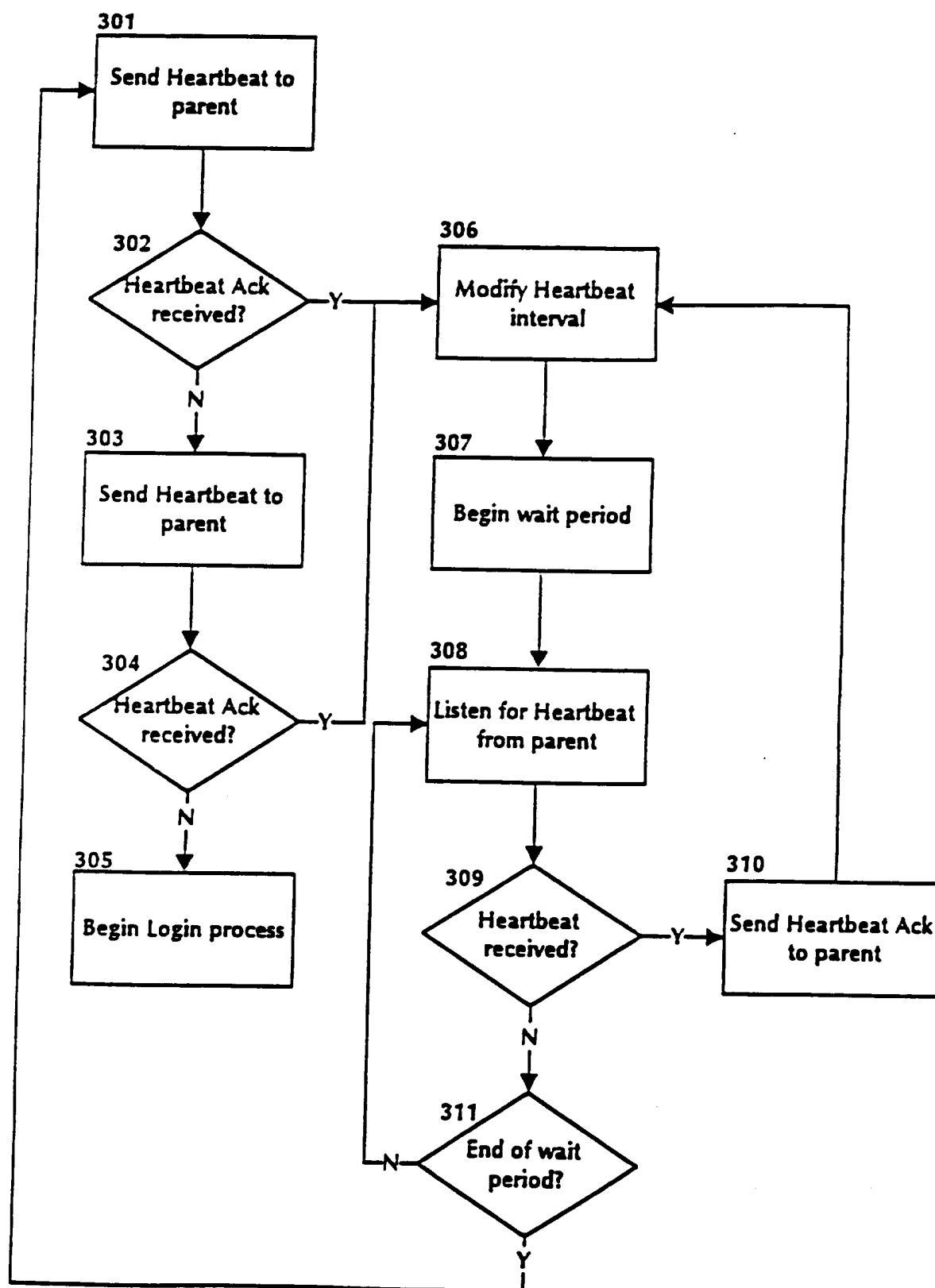
FIG. 8
8/34

9/34
(FIG. 9)



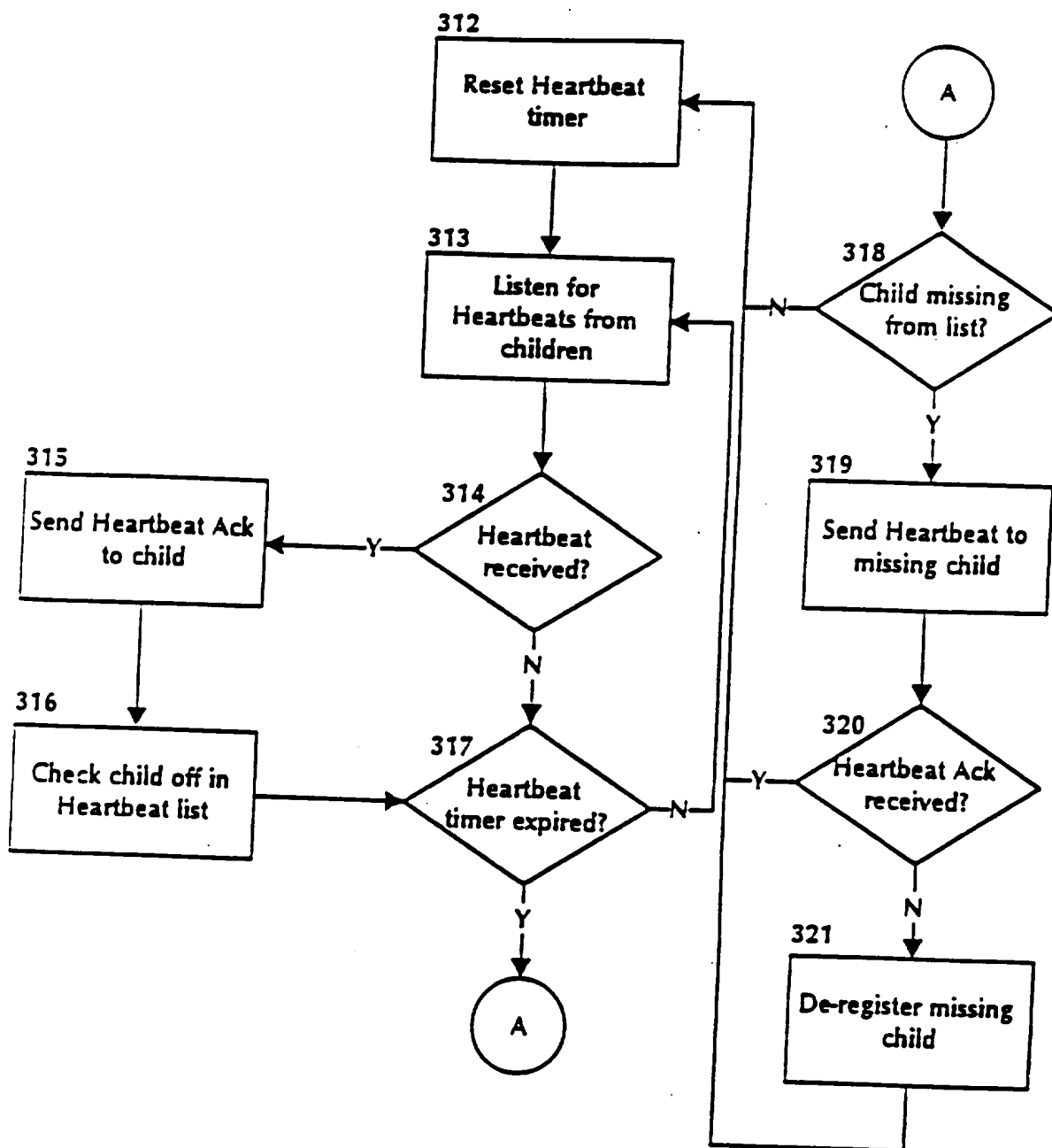
10/34

FIG. 10



11/34

FIG. 11



12/34

FIG. 12

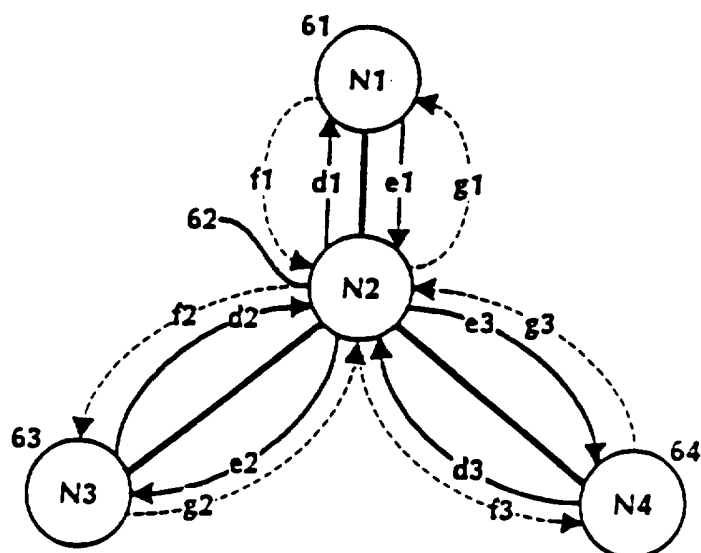
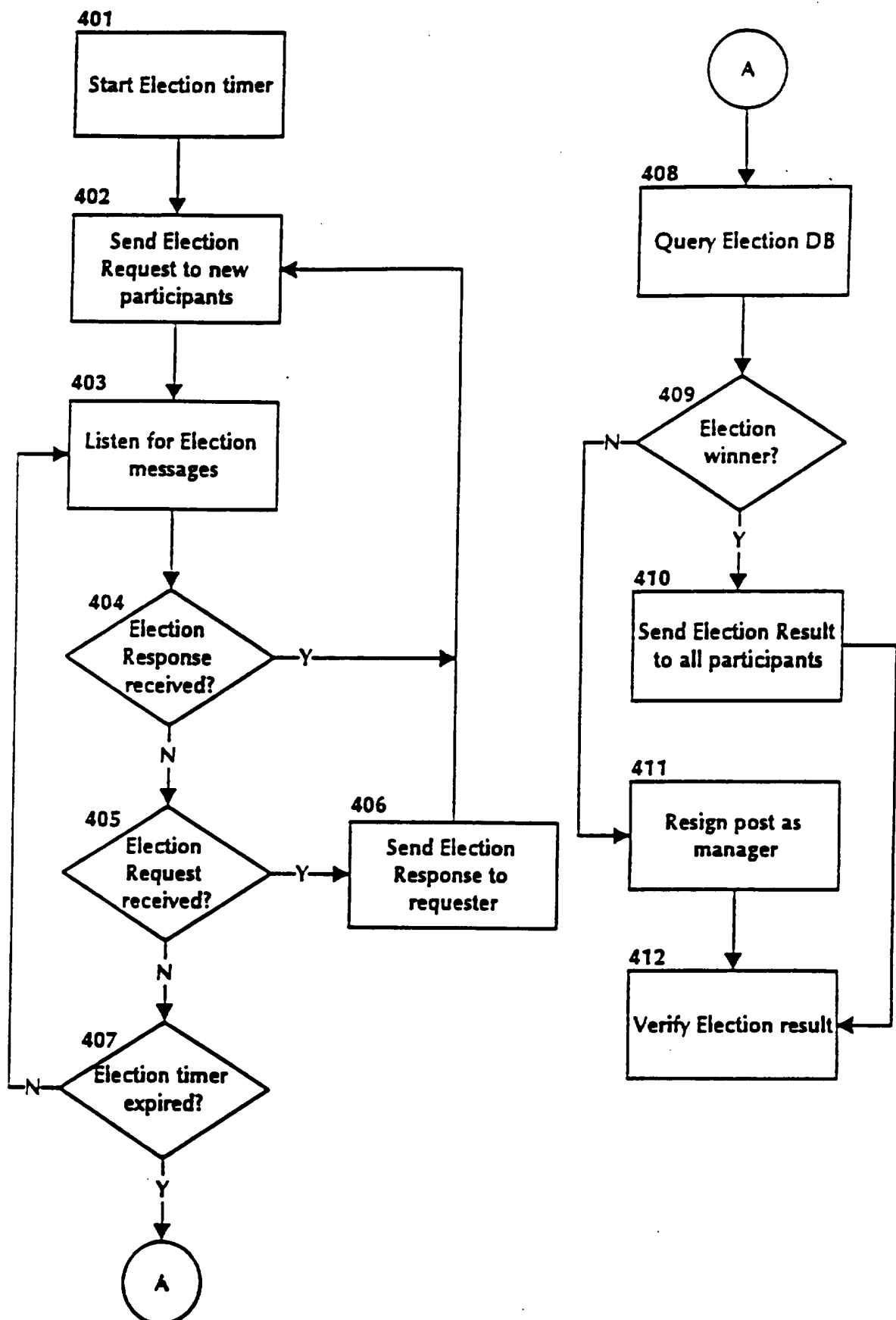


FIG. 13 13/34



14/34

FIG. 14

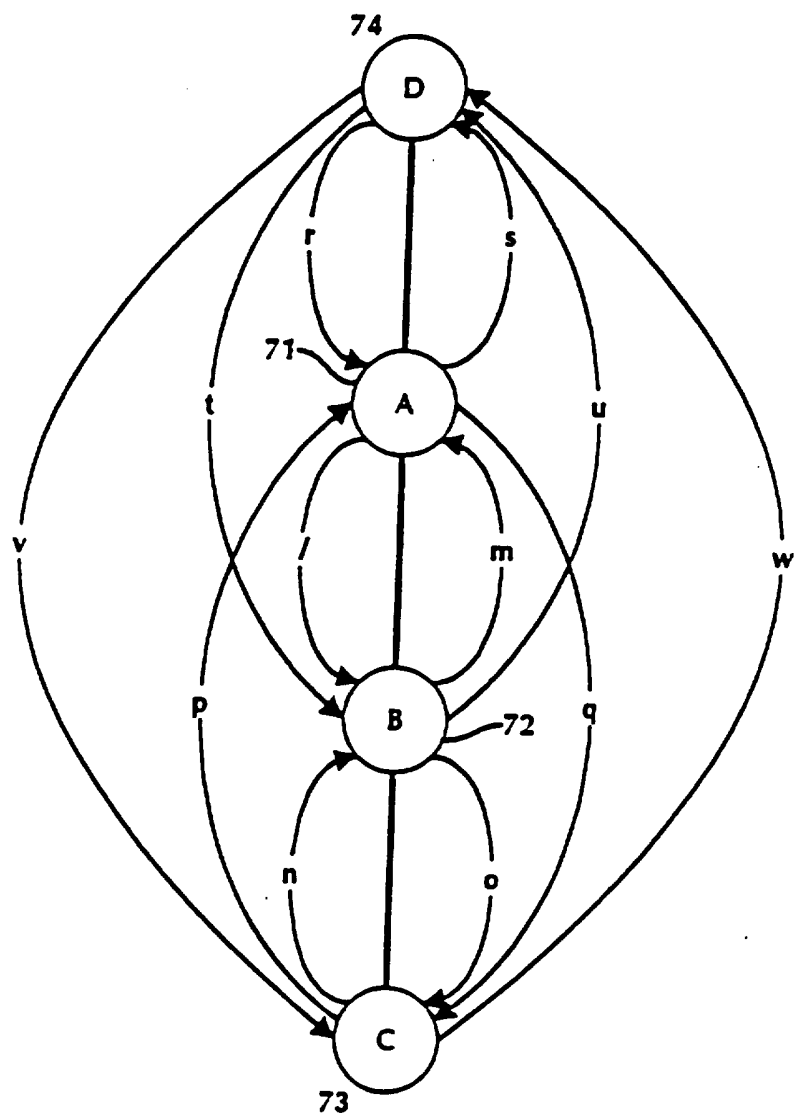


FIG. 15

15/34

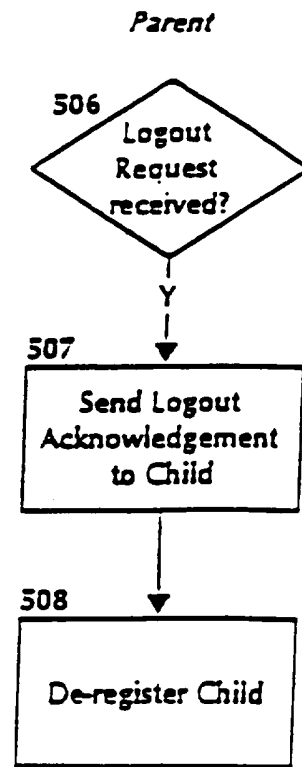
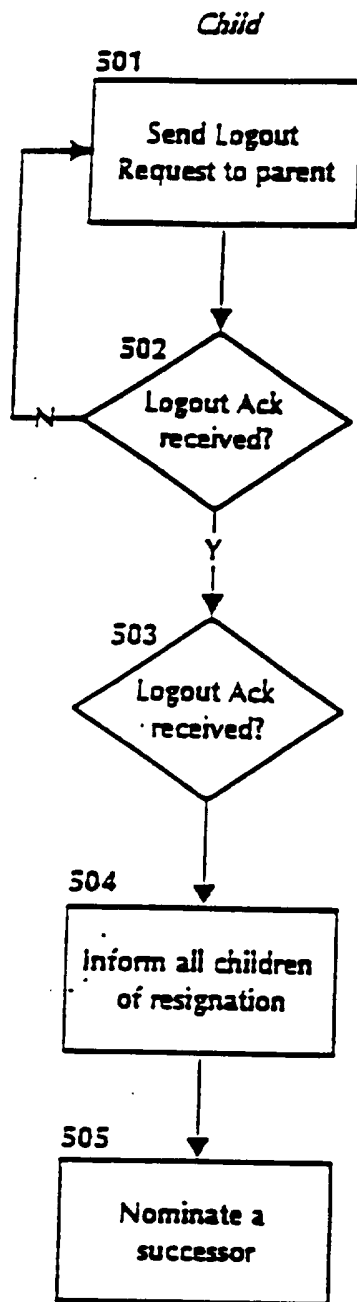
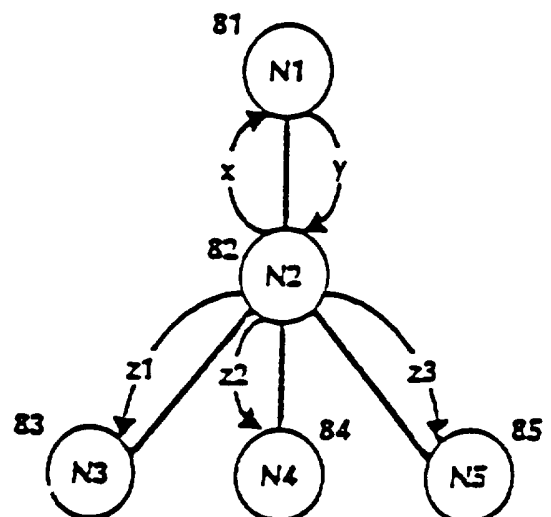
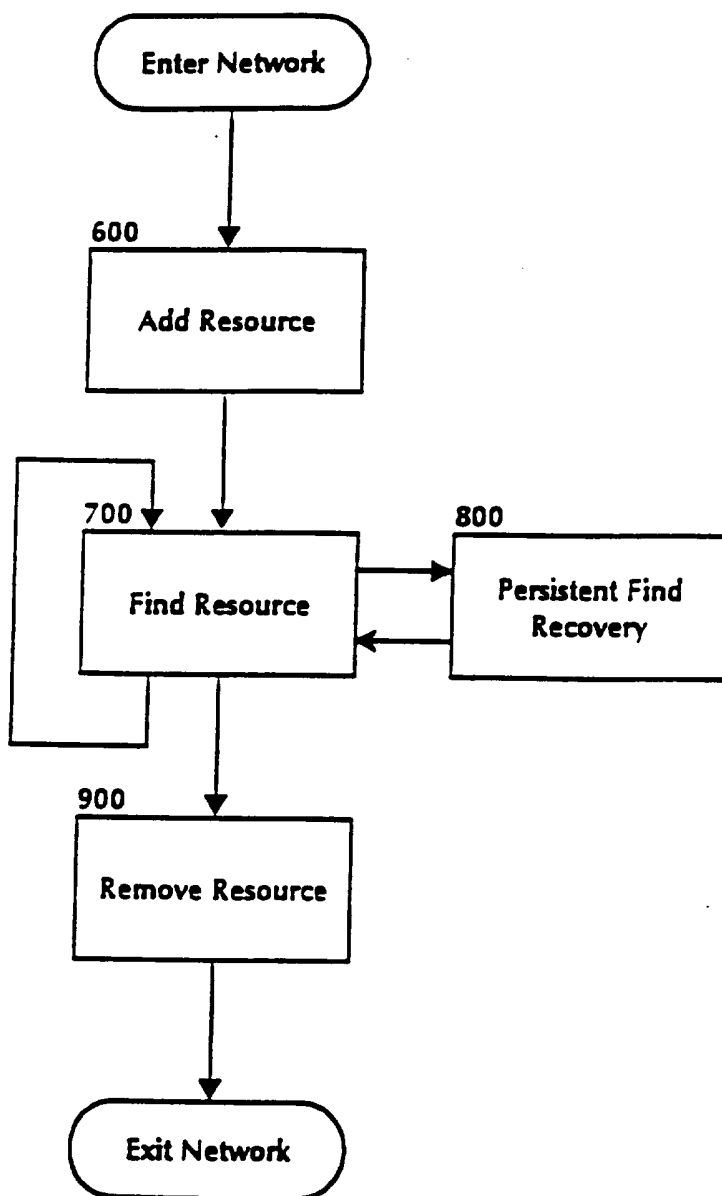


FIG. 16



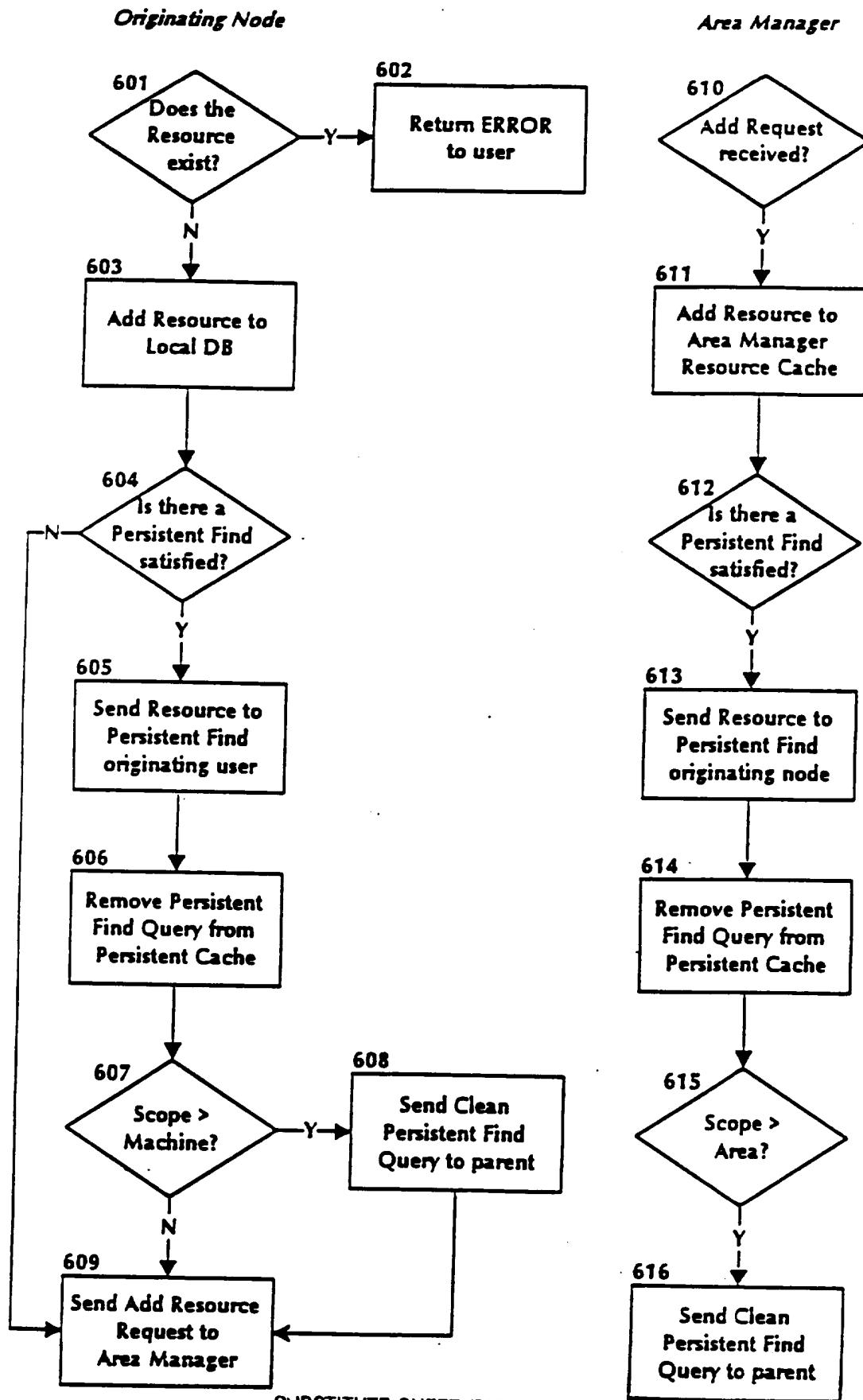
16/34

FIG. 17



17/34

FIG. 18



SUBSTITUTE SHEET (RULE 26)

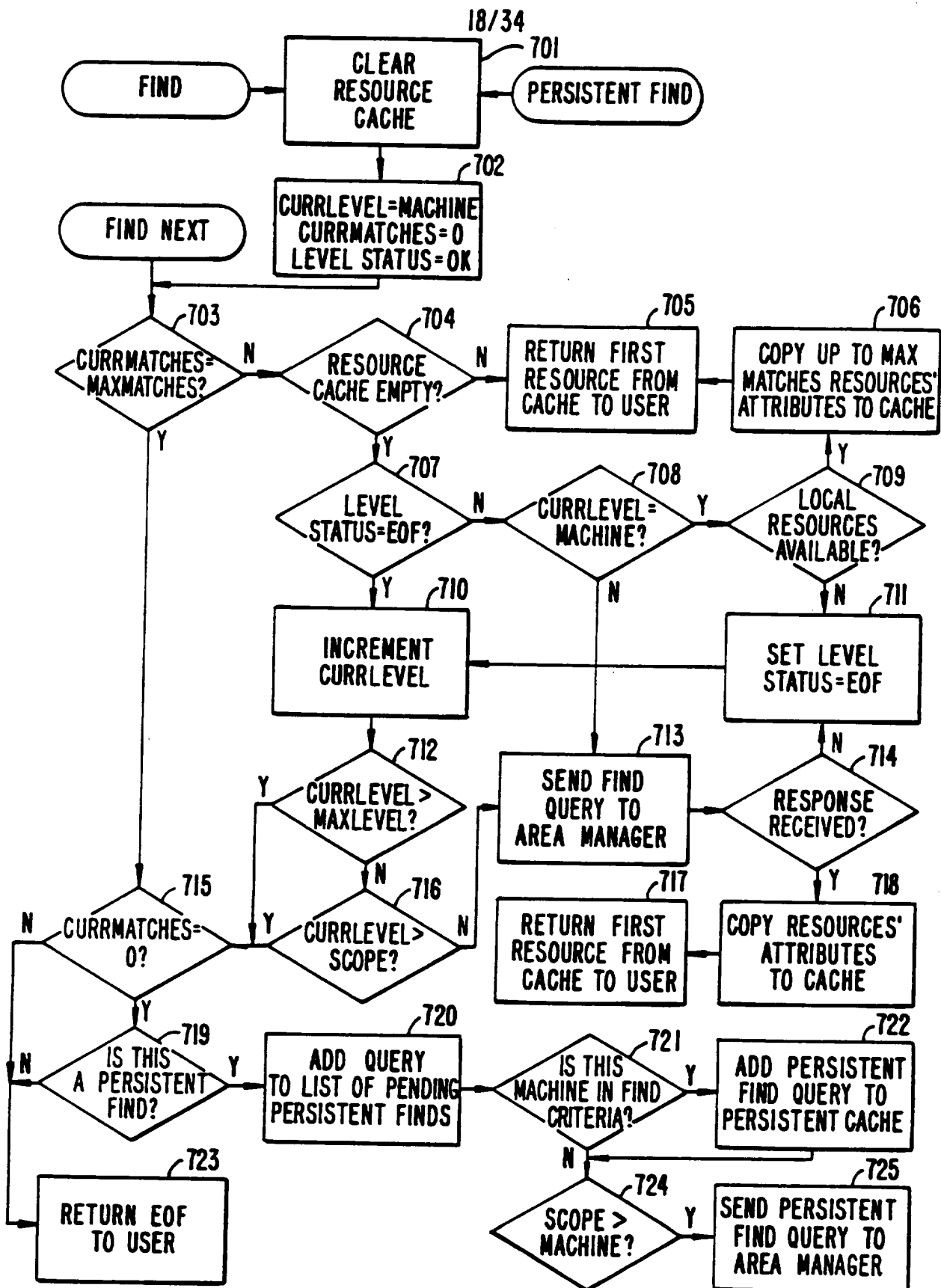
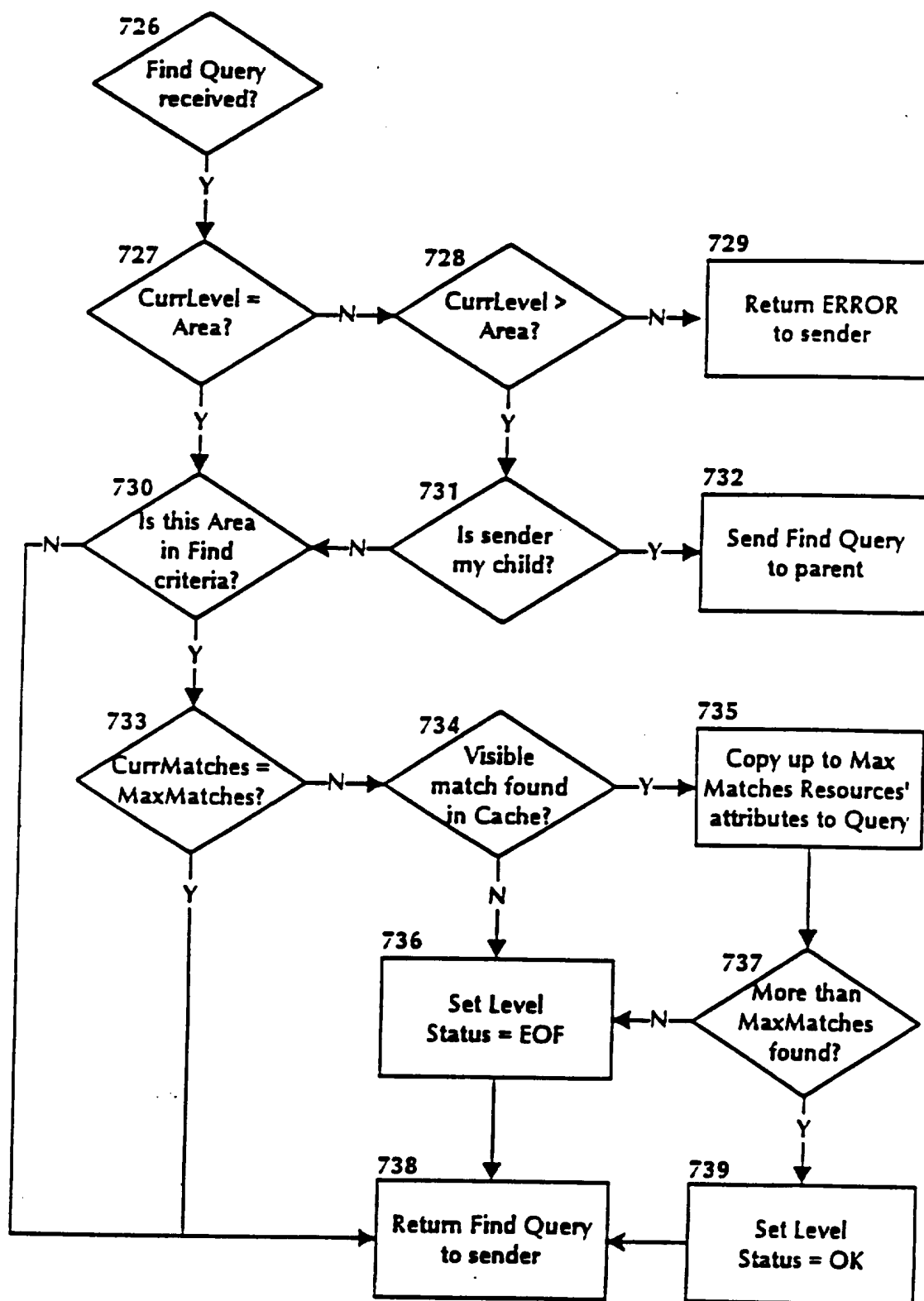


FIG. 19.

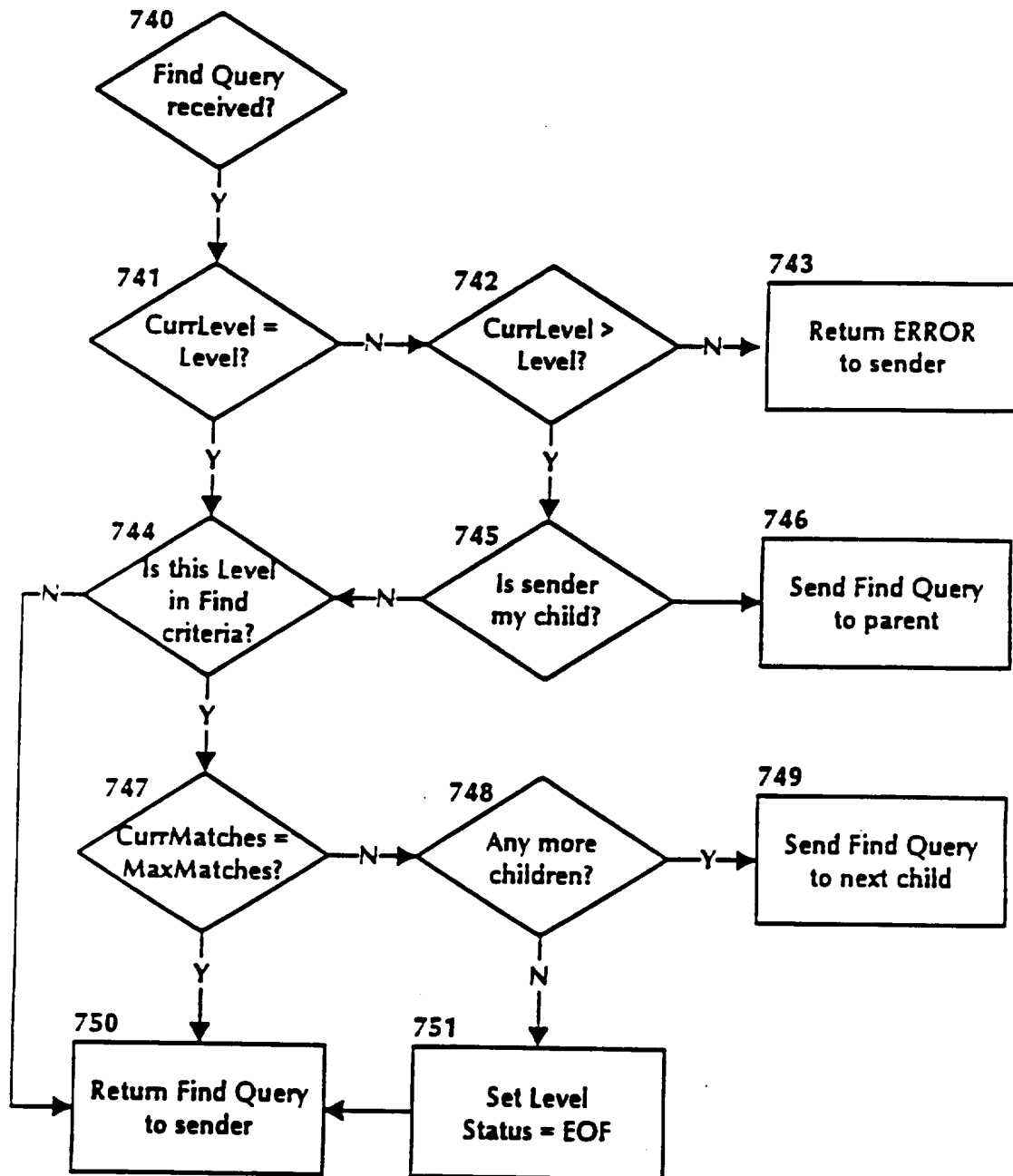
19/34

FIG. 20

Find Query Process at Area Manager Node

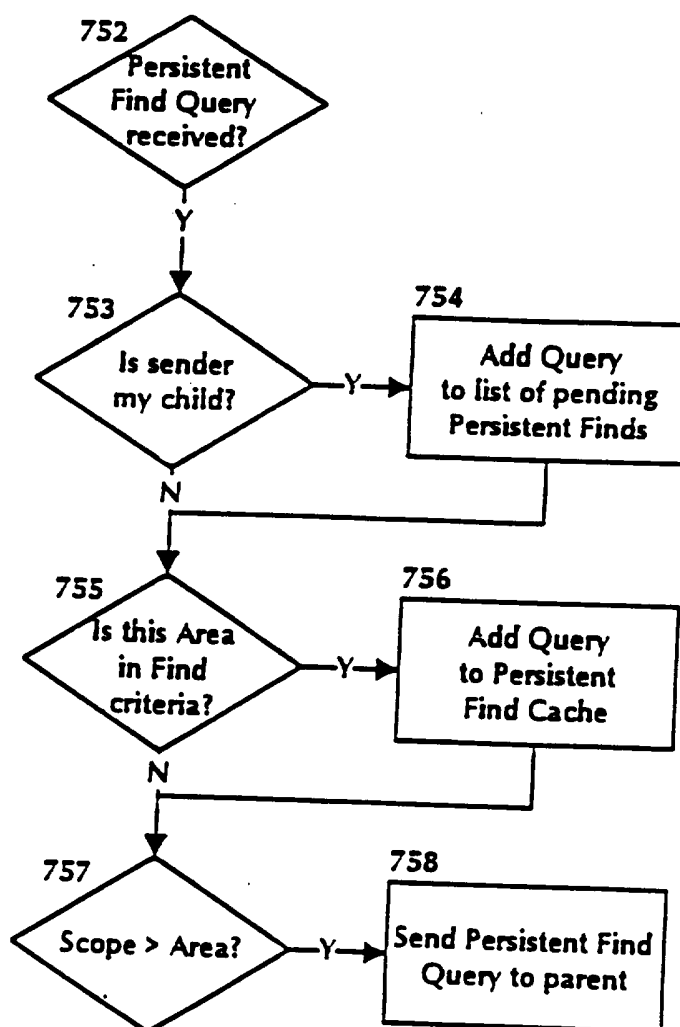
20/34

FIG. 21

Find Query Process at (Manager > Area Manager) Node

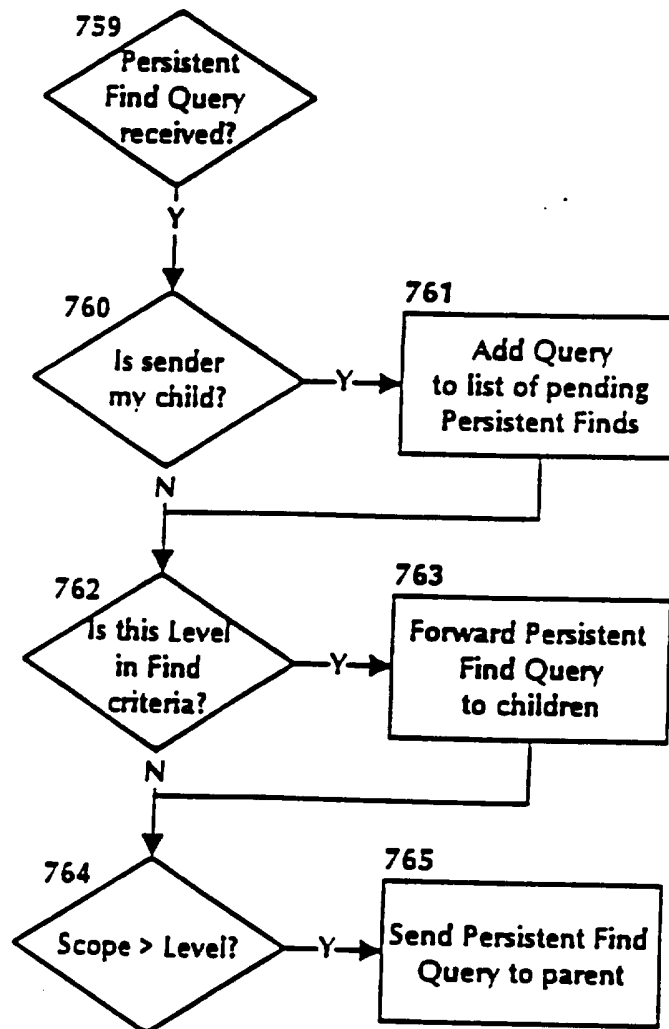
21/34

FIG. 22

Persistent Find Query Process at Area Manager Node

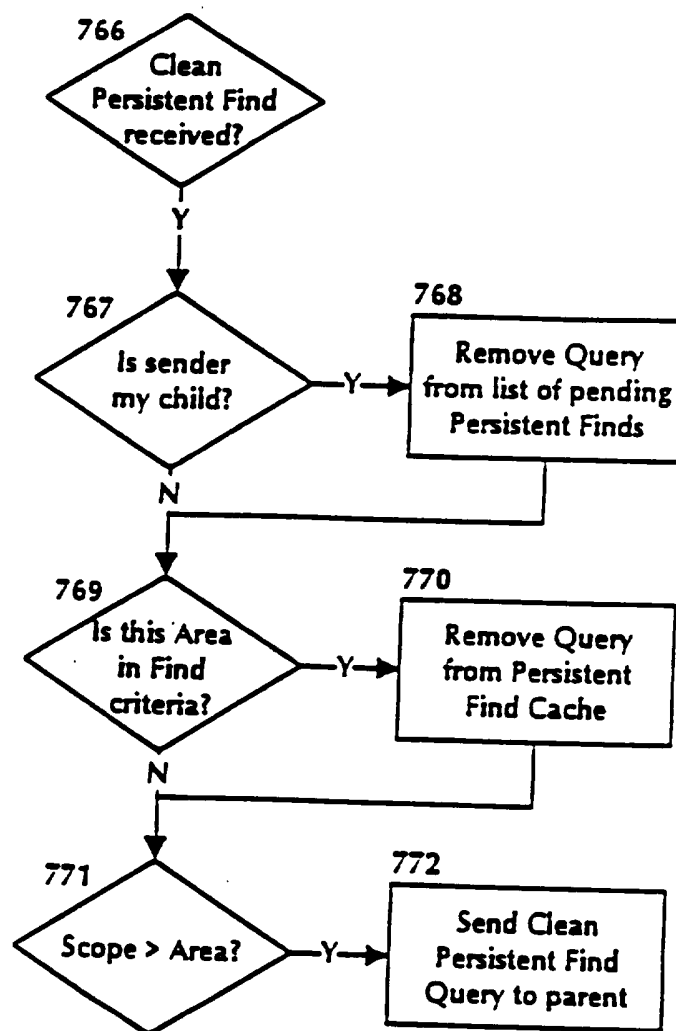
22/34

FIG. 23

Persistent Find Query Process at (Manager > Area Manager) Node

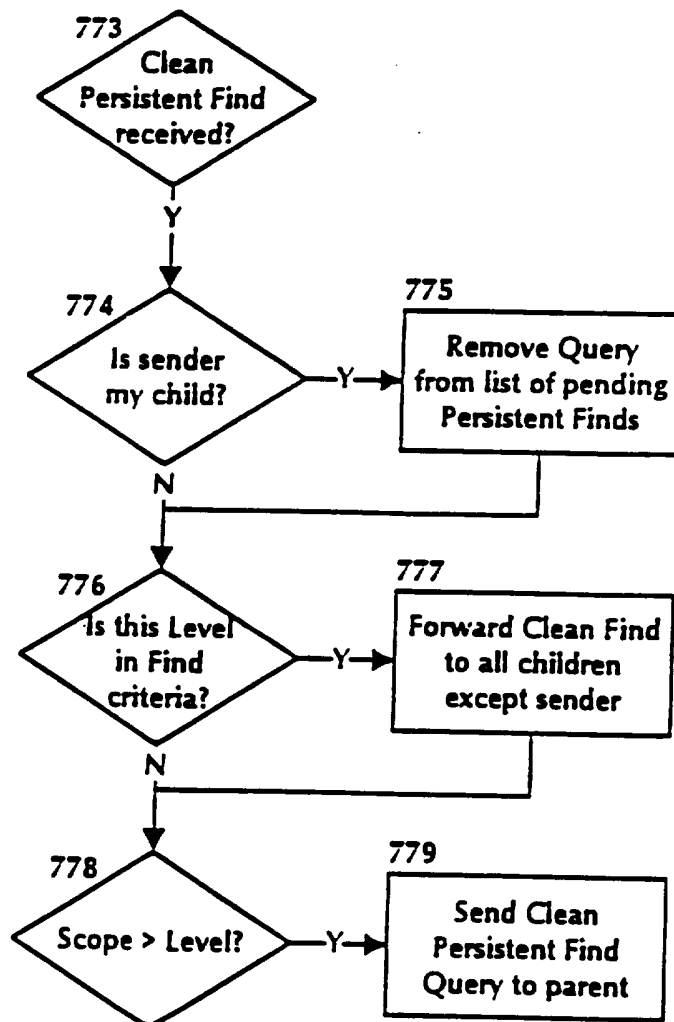
23/34

FIG. 24

Clean Persistent Find Query Process at Area Manager Node

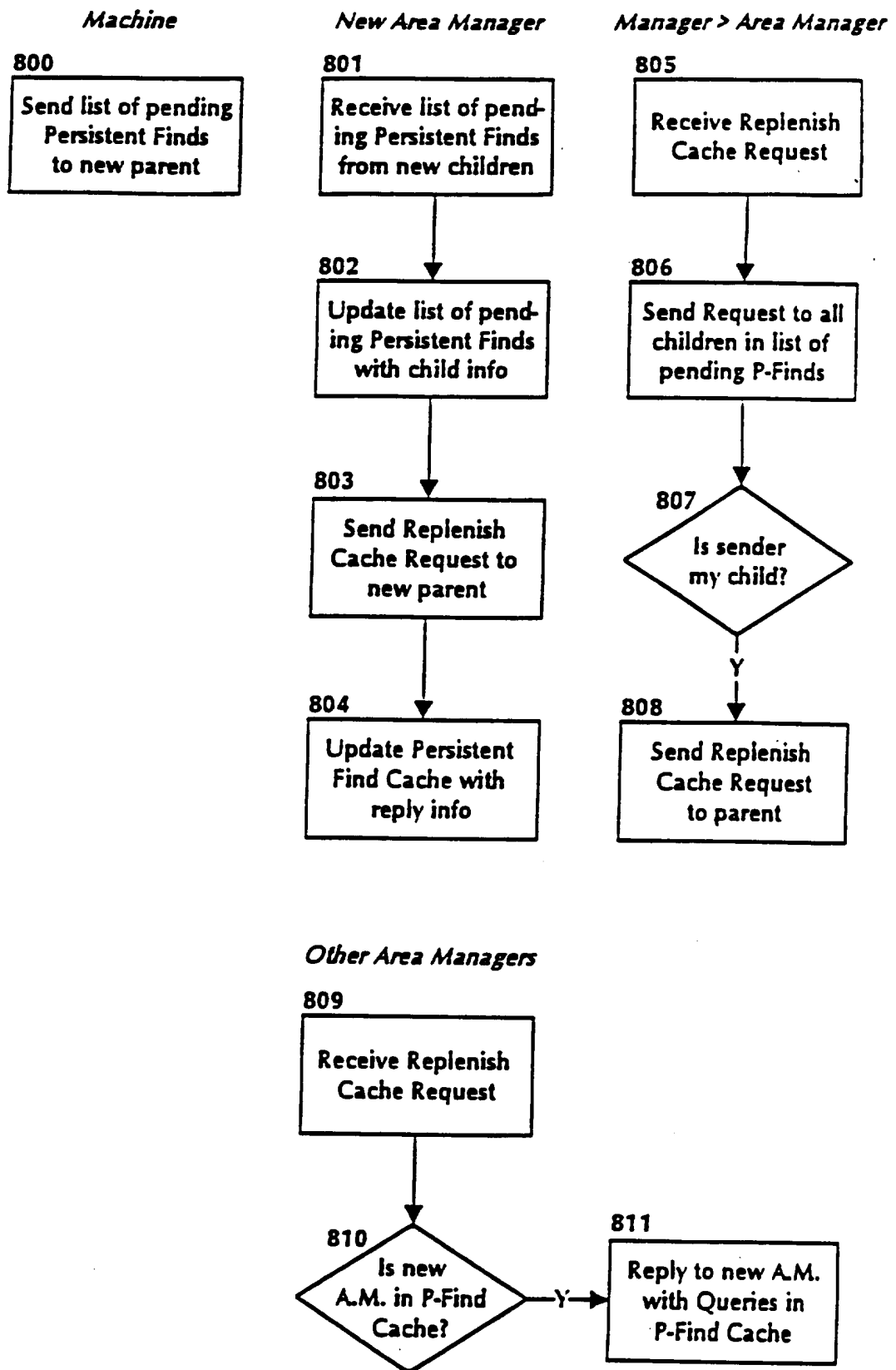
24/34

FIG. 25

Clean Persistent Find Query Process at (Manager > Area Manager) Node

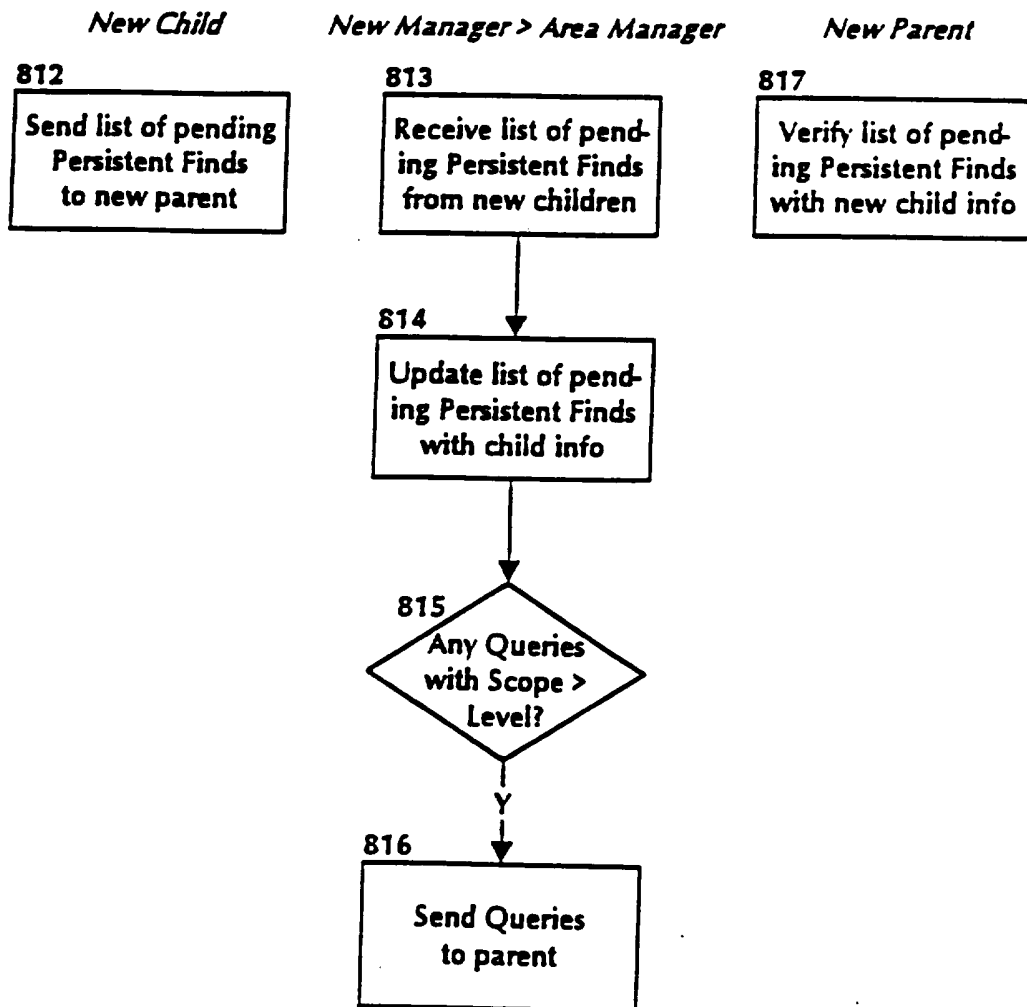
25/34

FIG. 26



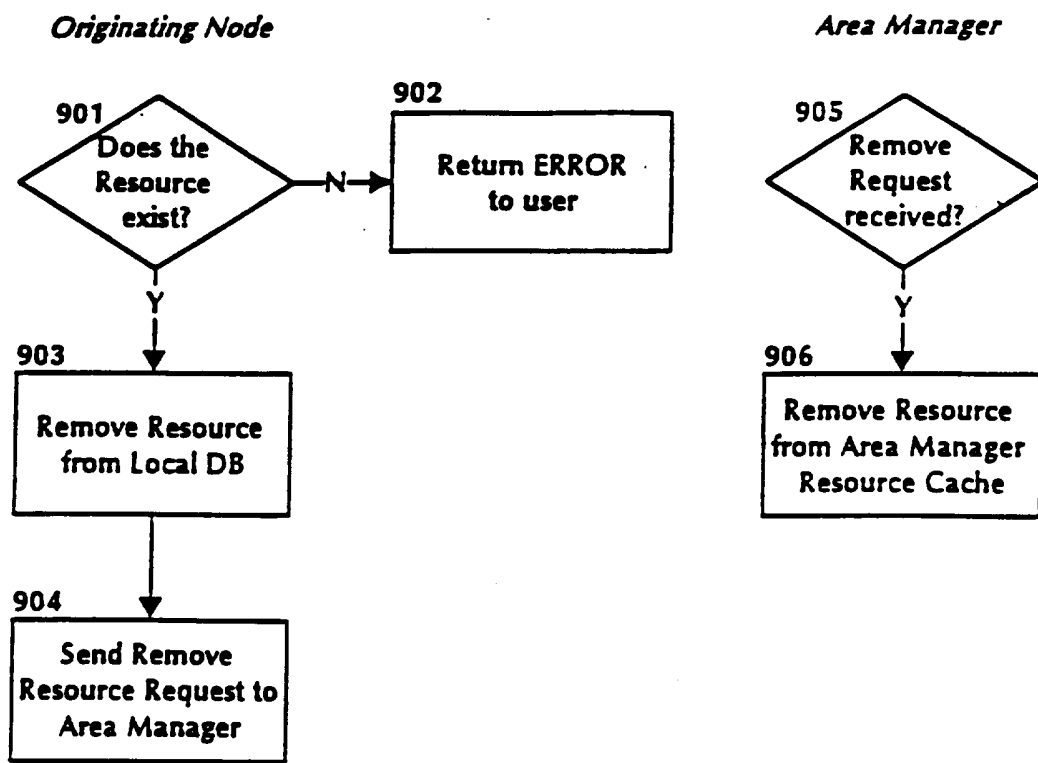
26/34

FIG. 27



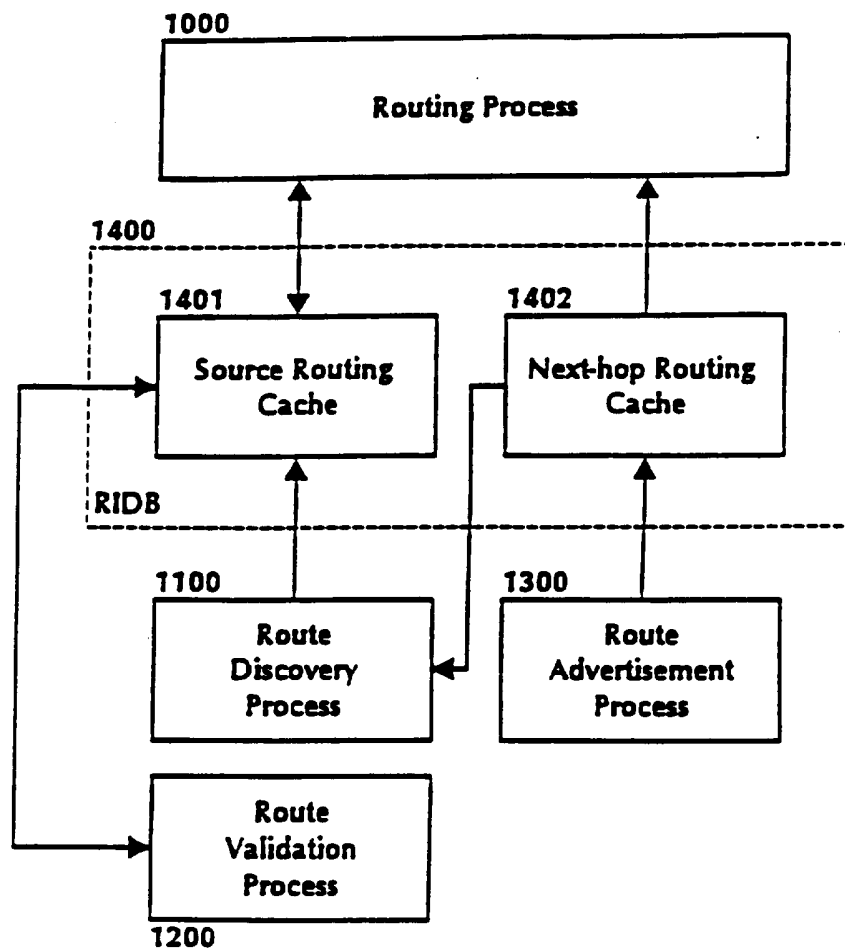
27/34

FIG. 28



28/34

FIG. 29A



29/34

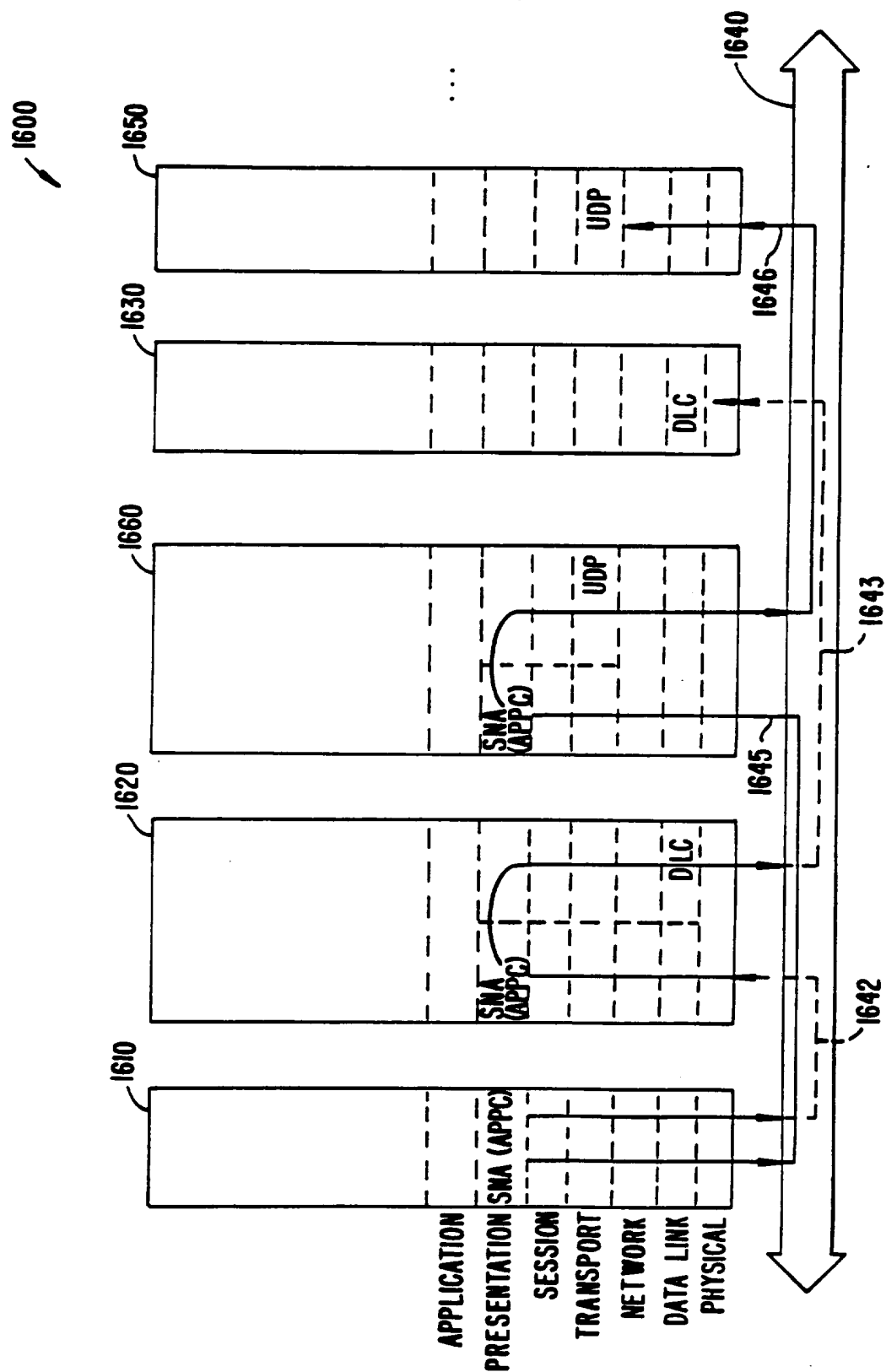
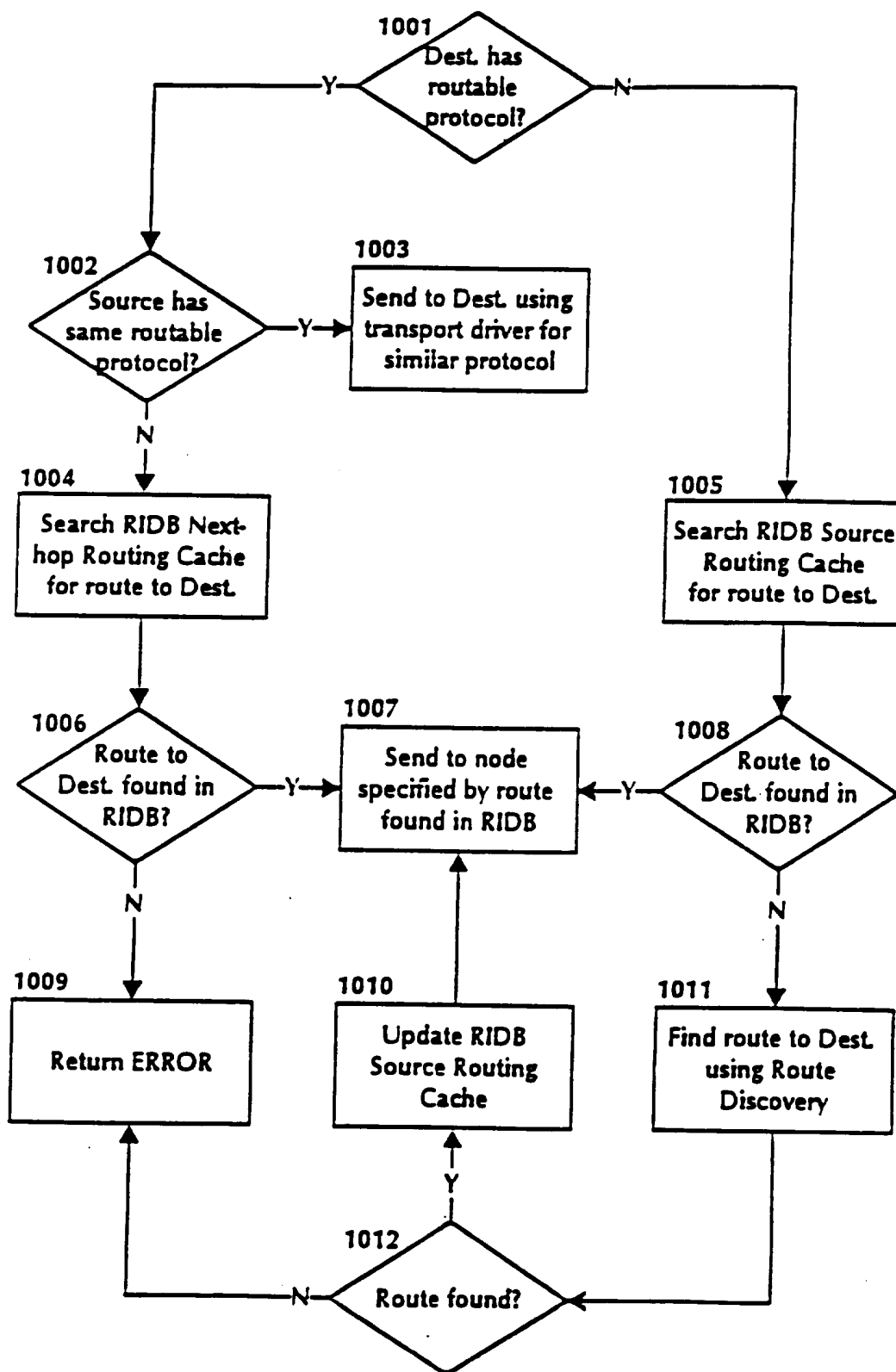


FIG. 29B.

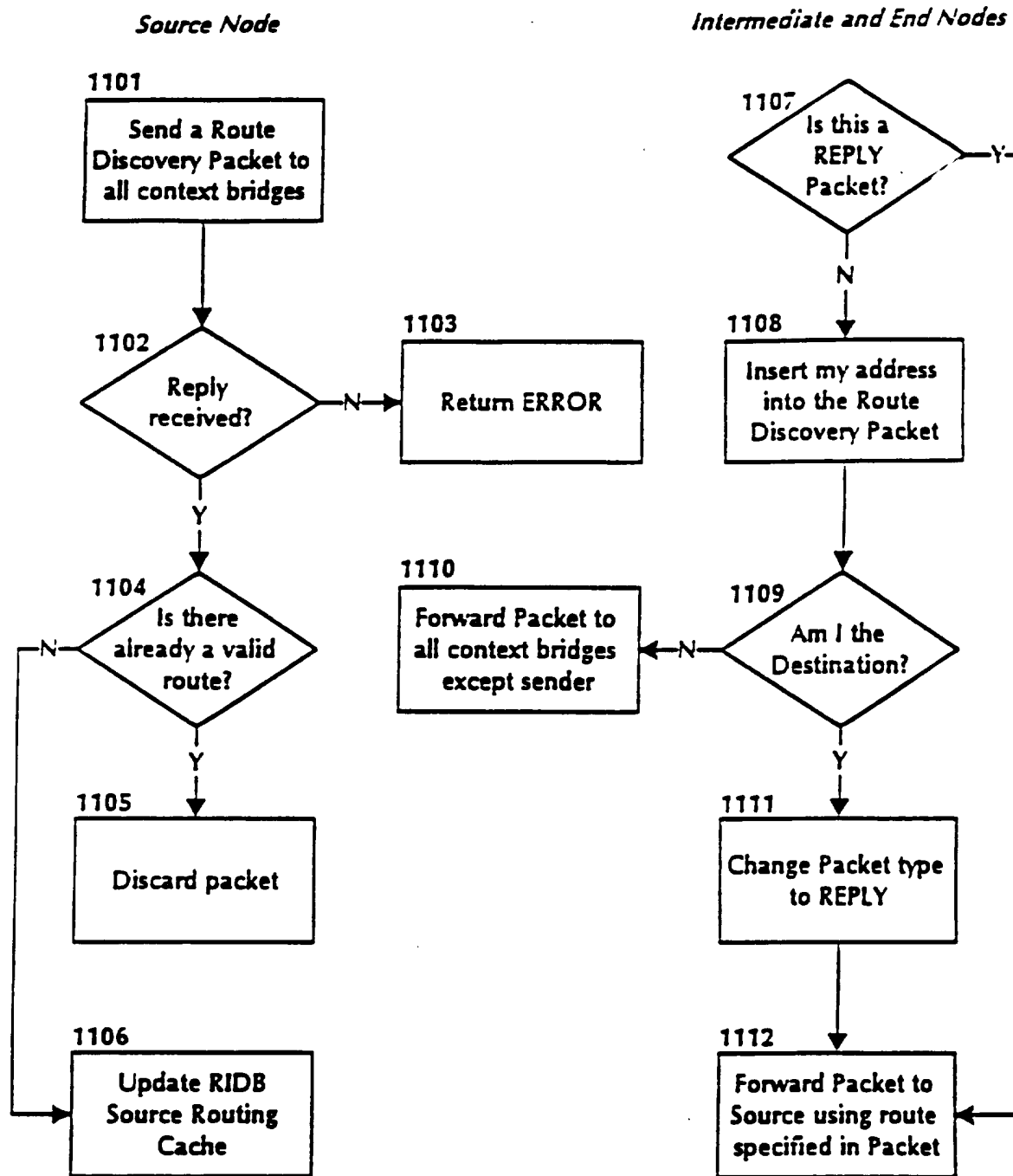
30/34

FIG. 30



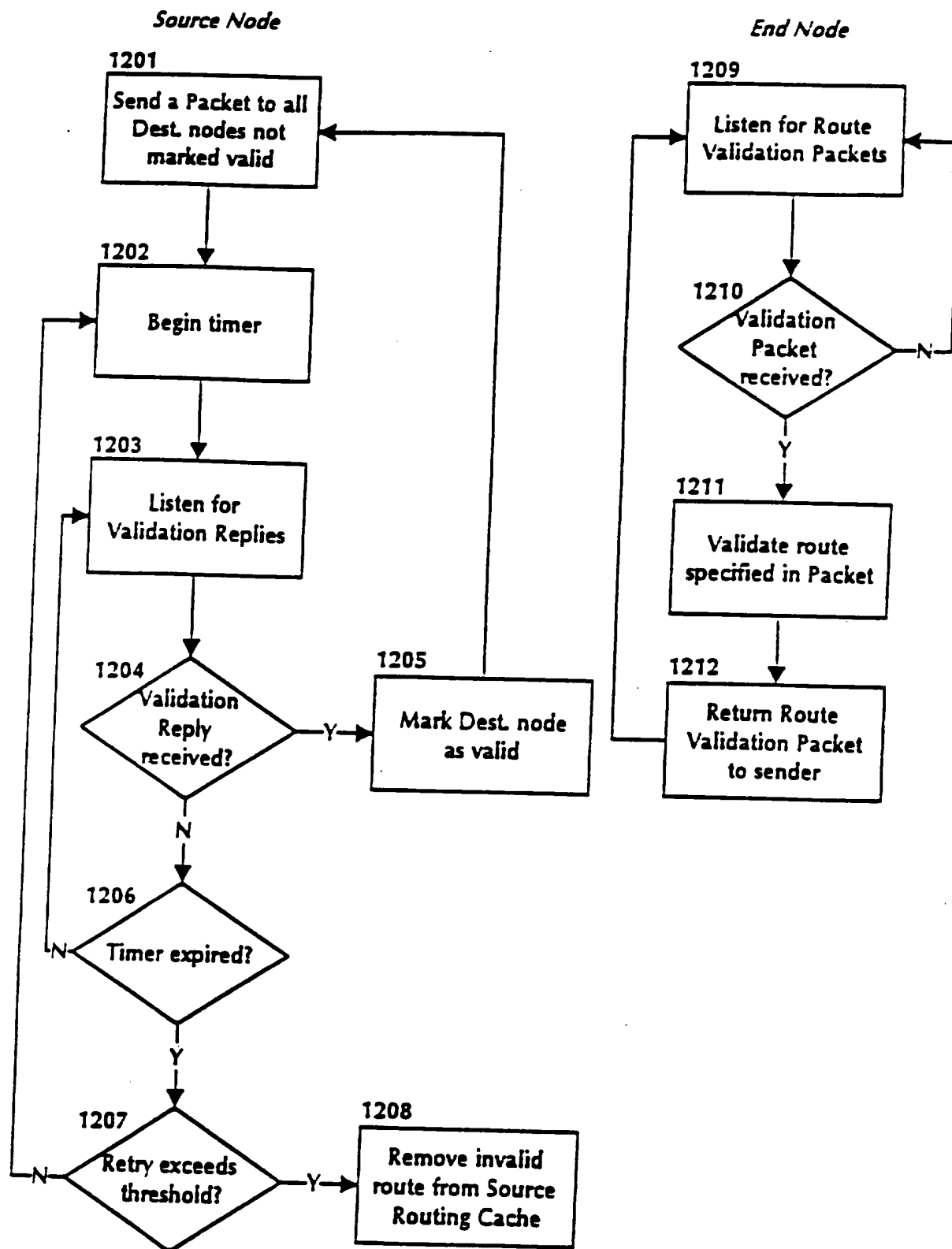
31/34

FIG. 31



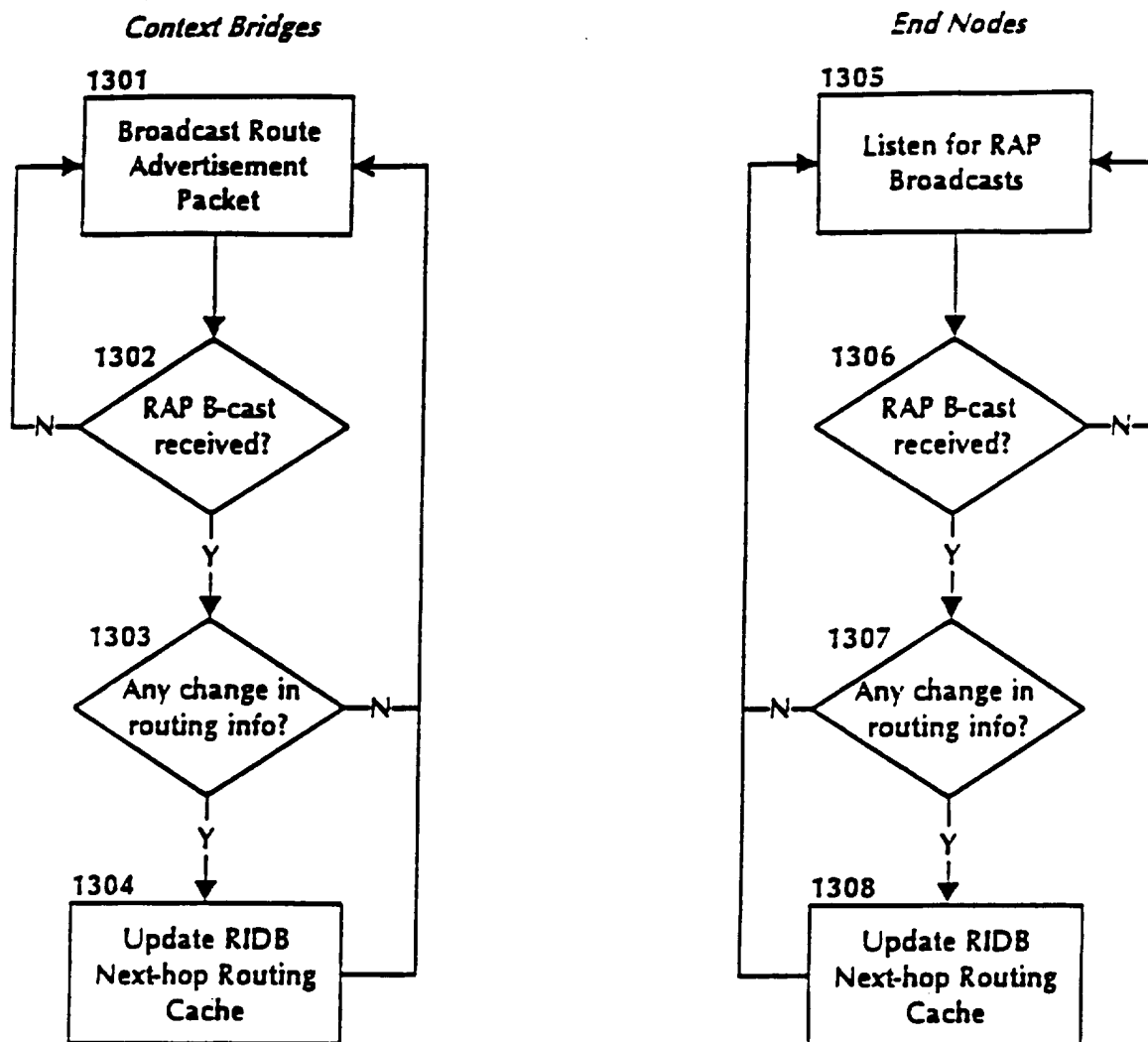
32/34

FIG. 32



33/34

FIG. 33



34/34

FIG. 34

