

Indice

Introduzione	1
1 Le reti di comunicazione: Internet	3
1.1 Breve storia	3
1.2 I servizi di rete	4
1.3 Architettura delle internet TCP/IP	8
1.4 Internet in Italia: INFNet e GARR	12
1.5 Le connessioni internazionali del GARR	13
1.6 Le prospettive future di Internet	16
2 Protocolli del network layer: IP e CLNP	19
2.1 I protocolli del Network Layer	19
2.1.1 IP: Internet Protocol	19
2.1.2 CLNP: ConnectionLess Network Protocol	20
2.2 Gli indirizzi	20
2.2.1 Indirizzi IP	21
2.2.2 Indirizzi CLNP	23
2.2.3 Indirizzamento piatto e gerarchico	24
2.2.4 Confronto tra gli indirizzi IP e CLNP	25
2.3 I pacchetti di dati	26
2.3.1 Datagramma IP	26
2.3.2 Pacchetto CLNP	28
2.3.3 Confronto fra le intestazioni dei pacchetti IP e CLNP .	30
2.4 Il routing	30
2.4.1 Neighbor greeting	31

2.4.2	Algoritmi di routing	33
2.4.3	Intradomain routing protocol	35
2.4.4	Interdomain routing protocol	40
3	Protocolli di routing link state	43
3.1	Caratteristiche dei protocolli di routing link state	43
3.1.1	Flooding	44
3.1.2	Algoritmo LinkState	45
3.2	OSPF	46
3.2.1	Topologia	46
3.2.2	Topological Database e Routing Table	49
3.2.3	Esecuzione	50
3.2.4	Pacchetti OSPF	51
3.3	Integrated IS-IS	58
3.3.1	Topologia	59
3.3.2	Link State Database e Forwarding Database	60
3.3.3	Routeing function	61
3.3.4	Pacchetti integrated IS-IS	63
3.4	Ship In the Night e Integrated routing	68
3.5	Confronto tra OSPF e Integrated IS-IS	69
3.5.1	Terminologia	69
3.5.2	Topologia	69
3.5.3	Designated router	75
3.5.4	Routing Table e database	76
3.5.5	Collegamenti punto a punto	76
3.5.6	Considerazioni finali	77
4	Protocolli di routing sulle reti INFNet e GARR	78
4.1	Autonomous system 137	78
4.2	I router	81
4.3	I protocolli di routing	82
4.3.1	RIP	82
4.3.2	IGRP	83
4.3.3	Integrated IS-IS	84

4.3.4	OSPF	86
4.3.5	BGP	89
4.4	Il redistribute	89
4.4.1	Percorsi asimmetrici	91
5	Prove di funzionalità dei protocolli di routing	93
5.1	Integrated IS-IS su C-LAN frame relay	93
5.1.1	Frame relay	94
5.1.2	Obiettivi dei test	95
5.1.3	Hardware e software impiegato	96
5.1.4	Funzionalità della C-LAN	97
5.1.5	Funzionalità di Frame relay switch	100
5.1.6	C-LAN in produzione	101
5.1.7	Conclusioni	108
5.2	Maglia OSPF tra Toscana e Bologna	109
5.2.1	Obiettivi del test	109
5.2.2	Topologia	109
5.2.3	Hardware e software impiegato	111
5.2.4	Misure di efficienza	111
5.2.5	Svolgimento delle prove	112
5.2.6	Conclusioni	115
5.2.7	Ultimi sviluppi	116
5.3	Confronto tra Ship In the Night e Integrated routing	117
5.3.1	Obiettivi del test	117
5.3.2	Topologia	117
5.3.3	Dimensioni dei pacchetti OSPF	118
5.3.4	Dimensioni dei pacchetti IS-IS	123
5.3.5	Dimensioni dei pacchetti integrated IS-IS	126
5.3.6	Confronto della occupazione di banda	129
6	Evoluzione dei protocolli di rete e del GARR	131
6.1	IP next generation	131
6.1.1	Breve storia	132
6.1.2	Criteri per la valutazione del nuovo protocollo	132

6.1.3	TUBA	135
6.1.4	SIPP	136
6.1.5	La nomina di SIPP a IPv6	141
6.2	Evoluzione della rete GARR	143
6.2.1	I requisiti per la nuova configurazione	143
6.2.2	I possibili approcci	144
6.2.3	Una configurazione per aree geografiche	146
6.2.4	Una configurazione per aree di affinità d'uso	147
6.2.5	Configurazione dei protocolli: possibili scenari	151
6.2.6	Conclusioni	156
7	Conclusioni	158
	Appendici:	161
A	Configurazioni router test Integrated IS-IS su C-LAN	161
B	Configurazioni dei router per il test della maglia OSPF	176
C	Reti dell'AS 137 divise per aree geografiche	179
D	Reti dell'AS 137 divise per aree di affinità d'uso	186
	Bibliografia	193
	Indice Analitico	197

Introduzione

Come tutte le tecnologie informatiche, anche le reti di comunicazione tra computer sono in continua e rapida evoluzione: in particolare le loro dimensioni sono cresciute fino a connettere calcolatori situati in continenti diversi. È il caso di Internet, una delle più importanti reti attualmente in funzione.

Per queste reti estese a livello mondiale la crescita è avvenuta in tempi rapidi e a volte in modo incontrollato, evidenziando sia l'inadeguatezza globale dei protocolli di trasmissione che la complessità gestionale delle topologie venutesi a creare.

In questa tesi è stato preso in esame il dominio italiano di Internet: esso è costituito da un “backbone” e dalle reti dei principali enti di ricerca (Università, Istituto Nazionale di Fisica Nucleare, CNR) ed è denominata GARR, dall'acronimo di *Gruppo di Armonizzazione Reti della Ricerca* che è il nome dell'autorità che la gestisce. Anche su questa rete si sta presentando la necessità di rivedere le strategie di routing ed il tipo di tecnologie di trasmissione da utilizzare.

Scopo di questa tesi è di proporre dei modelli di evoluzione della rete GARR per migliorarne il funzionamento e la gestione. A tal fine sono stati studiati i principali protocolli di routing basati sull'algoritmo “Link State”, per individuare le configurazioni più adatte alla infrastruttura disponibile che, oggi, è costituita da linee punto-a-punto e da PVC (circuiti virtuali) frame relay. Il funzionamento dei protocolli è stato anche verificato con prove effettuate sulla reti di produzione GARR e INFNet, al fine di valutarne le reali possibilità di utilizzo.

In questa tesi sono descritti i protocolli utilizzati, i test effettuati ed infine vengono presentate le proposte di evoluzione della rete GARR.

Nel primo capitolo della tesi sono descritte brevemente Internet, i servizi di rete disponibili e la rete GARR.

Nel secondo capitolo sono confrontati due protocolli responsabili della consegna dei pacchetti, l'IP ed il CLNP: IP è il protocollo usato da Internet, CLNP è invece quello della rete DECnet; questi protocolli sono eseguiti in parallelo sulla rete GARR e su INFNet.

Nel terzo capitolo sono confrontati i due protocolli di routing basati sull'algoritmo Link State: OSPF e integrated IS-IS; il primo può essere usato esclusivamente da IP, il secondo invece integra le funzioni per IP e per CLNP. Nel quarto capitolo è presentato lo stato della rete GARR: le connessioni internazionali, i collegamenti attivi, i protocolli di routing utilizzati.

Nel quinto capitolo sono descritti i test effettuati ed i risultati ottenuti.

Il sesto capitolo è strutturato in due parti: nella prima è descritto il nuovo protocollo che verrà utilizzato in un prossimo futuro in Internet; nella seconda sono formulate le proposte di configurazione dei protocolli di routing sulla rete GARR.

Buona parte del lavoro di preparazione della tesi e soprattutto i test sono stati svolti al CNAF (Centro Nazionale Analisi Fotogrammi) a Bologna, il centro dell'INFN che si occupa della gestione e dello sviluppo di INFNet. In quella sede, l'aiuto del dottore Davide Salomoni e dell'ingegnere Cristina Vistoli è stato essenziale, soprattutto per conoscere la rete e per configurare correttamente i router: ad essi va un sentito ringraziamento.

1

Le reti di comunicazione: Internet

La possibilità di comunicare è sempre stato un motivo di sviluppo e progresso per l'umanità; mai come oggi però i mezzi di comunicazione sono entrati nella globalità delle attività umane, siano esse produttive che di svago. Nel campo dell'elaborazione dell'informazione, gli sviluppi delle tecnologie hanno permesso di realizzare reti di comunicazioni per mezzo delle quali gli elaboratori si scambiano dati e condividono risorse; queste reti si sono sempre più ingrandite raggiungendo dimensioni intercontinentali.

Oggi la rete più diffusa a livello mondiale è quella che prende il nome di Internet. Essa è costituita dall'insieme di molte realtà nazionali e sovranazionali interconnesse in vari modi, modalità che nel loro insieme prendono il nome di internetworking; con questa parola infatti si intende generalmente la tecnica di connessione di reti (network) diverse. Un internet è allora una rete costituita dalla connessione di più reti locali con differenti tecnologie hardware; Internet è l'internet con la i maiuscola, cioè la principale internet nel mondo.

1.1 Breve storia

Nel 1973 la Defense Advanced Research Project Agency (DARPA) degli Stati Uniti iniziò un programma di ricerca per sviluppare una tecnica per inter-

connettere reti di vario tipo. L'obiettivo era ideare un insieme di protocolli che permettesse a computer collegati a reti distinte di comunicare (trasparentemente) attraverso altre reti connesse tra di loro. Questo progetto fu detto di *Internetting* ed il sistema di reti che risultò dalla ricerca fu chiamato *Internet*.

I protocolli che furono sviluppati divennero noti con la sigla TCP/IP dalle iniziali di Transmission Control Protocol e Internet Protocol, i due principali protocolli responsabili della trasmissione affidabile, dell'instradamento e della consegna dei pacchetti fra coppie di calcolatori.

In seguito Internet divenne parte fondamentale della infrastruttura degli istituti di ricerca degli U.S.A. per poi espandersi negli anni ottanta a livello internazionale.

Nel 1986 la National Science Foundation iniziò lo sviluppo della NFSnet realizzando così una delle maggiori backbone di Internet. Altri importanti enti di ricerca americani ed europei hanno poi creato dei backbone che, insieme, garantiscono una connessione globale ad alta velocità.

Oggi Internet raggiunge la quasi totalità delle nazioni della Terra (figura 1.1 a pagina 5) consentendo un elevato traffico di informazioni; collega più di 37.000 reti e almeno 3.800.000 host con un incremento annuo che negli ultimi tempi è stato superiore all'80%¹.

1.2 I servizi di rete

Internet è utilizzata sia da enti di ricerca che commerciali per la raccolta e lo scambio di dati, comunicazioni tra persone, esperimenti remoti. Le principali applicazioni usate sono:

Posta elettronica: consente ad un utente, che abbia accesso alla rete tramite un *account* su un calcolatore, di scrivere un messaggio e di inviarlo ad un altro utente o gruppi di utenti; i messaggi possono contenere oltre che dei testi, anche immagini (codificate opportunamente). La posta elettronica è molto più veloce della posta tradizionale, è normalmente più affidabile ed è

¹Fonte: Network Wizard, <http://www.nw.com/zone/WWW/report.html> (Dati rilevati nell'Ottobre 1994).

Sostituire con la mappa WORLD_CONNECT.PS

Figura 1.1. Nazioni connesse a Internet.

possibile assicurarsi che il destinatario abbia ricevuto il messaggio.

La posta elettronica ha avuto un tale successo che molti utenti delle reti di trasmissione dipendono da essa per la loro normale corrispondenza di lavoro; ciò è dovuto essenzialmente al fatto che il ricevente la può consultare quando vuole e con mezzi evoluti che ne permettono una comoda gestione e catalogazione.

Dalla posta elettronica sono derivate le *news*: divise per i più diversi argomenti raccolgono messaggi di posta elettronica scritti dagli utenti e che vengono diffusi a tutti coloro che hanno fatto richiesta di esserne messi a conoscenza. Le news permettono a ricercatori di tutto il mondo di scambiarsi pareri, farsi domande, porre nuovi obiettivi facilitando e velocizzando il lavoro di tutti.

Trasferimento di file: permette di trasferire file di qualsiasi tipo (documenti, archivi, immagini, programmi) anche di notevoli dimensioni, tra coppie di host. Il sistema dispone di meccanismi di sicurezza che consentono di impedire l'accesso ai file a persone non autorizzate o di consentire l'accesso a tutte le persone solo a determinati file.

Questo servizio ha permesso la nascita degli *Anonymous ftp² site*: sono “luoghi” (siti) dove sono archiviati migliaia di programmi, testi, immagini che chiunque può prelevare e usare. Un sito corrisponde ad un server che mette a disposizione alcune directory. Un utente si può collegare ad un sito utilizzando lo username *anonymous* (anonimo) ed “esplorarlo” alla ricerca dei documenti o dei programmi che gli interessano per poi eventualmente copiarli sul proprio sistema. Questo meccanismo è ora usato anche da case produttrici di software per distribuire nuove versioni o manuali limitatamente agli utenti che già possiedono le licenze d'uso.

Gli Anonymous ftp site sono oggi migliaia e molto usati; alcuni possono essere specializzati in determinati argomenti, altri mettono a disposizione la produzione locale di documenti (come le università), oppure rendono disponibili immagini particolari (ad esempio dai satelliti meteorologici). Anche in questo caso, la libertà di accesso e di uso consente una grande diffusione delle conoscenze, accelerando i progetti di ricerca.

²Ftp: File Transfer Protocol; è il protocollo più usato per il trasferimento di file.

Login remoto: permette all'utente di un computer di collegarsi ad una macchina remota e di stabilire una sessione di login interattiva. Il login remoto stabilisce una connessione inviando il codice di ogni tasto battuto dall'utente alla macchina remota e visualizzando sullo schermo del terminale ogni carattere stampato dal computer remoto.

Questo servizio è uno dei più utilizzati poiché consente l'uso di risorse di calcolo e di programma che non possono essere disponibili ovunque; ad esempio un ente di ricerca che possiede un supercomputer può rendere accessibile le sue capacità di elaborazione ad utenti di altre sedi; oppure destinando diverse macchine a diversi usi (ad esempio installandovi diversi applicativi) è possibile usarle tutte da un unico terminale.

Questo modo di lavorare ha completamente rivoluzionato l'uso dei calcolatori che inizialmente rendevano disponibili le proprie risorse solamente in loco imponendo quindi la presenza in un unico luogo di tutte le persone che ne avessero necessità. Oggi da un terminale è possibile accedere a qualunque macchina connessa: ad esempio chi dovesse gestire il routing di un autonomous system può configurare tutti i router dal computer sulla propria scrivania senza dover girare un'intera nazione o contattare decine di persone. Tutto questo si risolve in una maggiore produttività e velocità.

Multicast di pacchetti audio e video: tramite microfoni e telecamere collegate ai computer è possibile inviare suoni e immagini che diversi altri utenti possono ricevere; è possibile dunque realizzare discussioni di gruppo in tempo reale tra persone distanti anche migliaia di chilometri oppure osservare fenomeni che si verificano in qualsiasi parte della Terra. Simile alle news, consente però un contatto più immediato tra le persone: i vantaggi sono quelli di potersi trovare insieme, discutere in tempo reale, mostrarsi oggetti o foto, senza dovere fare lunghi viaggi.

Applicazioni client-server: Esiste una classe di programmi che trae profitto dall'impiego cooperativo di un internet; lo schema di iterazione fra queste applicazioni è noto come *client-server*.

Il termine *server* identifica qualsiasi programma che offre un servizio raggiungibile attraverso la rete. I server accettano le richieste che gli pervengono, eseguono il loro servizio e trasmettono i risultati al richiedente. Un

programma è un *client* quando invia una richiesta ad un server e ne attende la risposta.

Tutti questi servizi disponibili a chi utilizza un'internet hanno in comune la capacità di trasferire facilmente informazioni da un capo all'altro della Terra e questo accelera i progetti di ricerca e sviluppo.

1.3 Architettura delle internet TCP/IP

Il progetto e la realizzazione di una rete è un problema molto complesso se considerato nella sua totalità; per semplificare il lavoro si è definito un modello gerarchico costituito da diversi livelli. L'idea è che ogni livello deve fornire determinati servizi al livello superiore utilizzando i servizi che gli vengono messi a disposizione dal livello inferiore; inoltre, in ogni nodo un livello comunica con il suo pari livello di un altro nodo attraverso un protocollo di trasmissione.

4	Application layer
3	Transport layer
2	Internet layer
1	Network interface layer
	Hardware layer

Figura 1.2. TCP/IP layer.

Gli sviluppatori di TCP/IP hanno definito quattro livelli (layer) sopra alle specifiche hardware delle singole reti; essi sono riportati in figura 1.2 e sono:

0 *Hardware layer*: definisce le specifiche hardware della rete e il formato dei dati che gli devono essere passati.

1 *Network interface layer*: è l'interfaccia tra la rete fisica e l'IP cioè accetta i datagrammi IP e li trasmette su una rete specifica.

2 *Internet layer*: gestisce la comunicazione fra coppie di sistemi; accetta

le richieste di trasmissione di pacchetti dal transport layer che vengono incapsulati in un datagramma IP; a questo datagramma viene assegnata un'intestazione che contiene le informazioni necessarie ai protocolli di routing per poterlo consegnare alla destinazione. L'Internet layer gestisce anche i datagrammi in arrivo, verificandone la validità determinando se il datagramma debba essere elaborato localmente o ritrasmesso. Ai datagrammi giunti a destinazione viene tolta l'intestazione e viene determinato a quale protocollo del transport layer passarli. Infine gestisce i messaggi di errore e controllo Internet Control Message Protocol (ICMP).

3 *Transport layer*: ha il compito primario di provvedere alla comunicazione da un programma applicativo eseguito su di un nodo ad un altro, eseguito o sullo stesso nodo o su uno diverso. Può regolare il flusso di informazioni e fornire un trasporto affidabile garantendo che i dati arrivino senza errori ed in sequenza. Suddivide la sequenza di dati che devono essere trasmessi in pacchetti che i quali sono passati all'internet layer.

4 *Application layer*: realizza i programmi applicativi che interagiscono con i protocolli del Transport layer per trasmettere o ricevere i dati.

I due protocolli principali di Internet sono l'IP e il TCP.

IP (Internet Protocol)³ è il protocollo dell'Internet layer; fornisce un servizio di consegna di pacchetti senza connessione (connectionless) non affidabile e best effort.

Il servizio è definito *senza connessione* poiché ogni pacchetto contiene l'indirizzo di destinazione e viene instradato indipendentemente dagli altri. Il servizio è *non affidabile* perché la consegna dei pacchetti non è garantita: questi possono essere persi, duplicati, ritardati o consegnati fuori sequenza, ma il protocollo non rileverà tali condizioni né informerà il trasmettitore o il ricevitore che si sono verificati questo tipo di inconvenienti. *Best effort* (sforzo ottimo) significa che il software compie un serio tentativo di consegnare i pacchetti senza scartarli capricciosamente; l'inaffidabilità si verifica solo quando le risorse sono esaurite o la rete sottostante non funziona.

I pacchetti IP (o datagrammi) sono composti da una *intestazione* e un'a-

³[RFC 791]

Int. frame	Dati Frame	
	Intestazione IP	Dati IP

Figura 1.3. Pacchetto (o datagramma) IP.

rea dati; l'intestazione è utilizzata dai protocolli di routing per consegnarli a destinazione, l'area dati contiene le informazioni da trasmettere. Un datagramma IP è incapsulato in un frame della rete fisica di cui occupa l'area dati (figura 1.3).

TCP (Transmission Control Protocol)⁴ è un protocollo del transport layer; fornisce ai programmi applicativi un servizio di trasmissione affidabile dei pacchetti. Questo significa che tra due host che vogliono comunicare si stabilisce un circuito virtuale in cui entrambi verificano la disponibilità dell'altro allo scambio di informazioni, si scambiano informazioni e si confermano di averle ricevute correttamente; inoltre i dati che una applicazione passa al protocollo TCP vengono ricevuti dalla applicazione destinazione esattamente nello stesso ordine.

Intestazione IP	Dati IP	
	Int. TCP	Dati TCP

Figura 1.4. Segmento TCP.

L'unità di trasferimento è il segmento che può essere utilizzato per trasmettere informazioni di controllo o dati; esso è contenuto nell'area dati di un datagramma IP (figura 1.4).

Il modello di stratificazione delle funzioni di rete proposto per TCP/IP non è l'unico possibile o esistente e nemmeno quello più utilizzato; il modello divenuto di riferimento è piuttosto quello definito dall'ISO⁵ che lo ha de-

⁴[RFC 793]

⁵ISO: International Organization for Standardization.

nominato OSI⁶ reference model; esso è costituito dai sette livelli riportati in figura 1.5.

1 *Physical layer*: ha il compito di trasmettere sequenze di bit di informazioni

7	Application layer
6	Presentation layer
5	Session layer
4	Transport layer
3	Network layer
2	Data link layer
1	Physical layer

Figura 1.5. OSI Reference Model.

attraverso un collegamento. Definisce quindi la dimensione e il tipo di connettori, la funzione dei pin di questi, la conversione dei bit in segnali elettrici, il tipo di sincronizzazione.

2 *Data link layer*: si occupa della trasmissione di pacchetti di informazione su un collegamento fisico effettuando controlli su eventuali corruzioni dei dati ed eseguendo l'indirizzamento dei pacchetti quando più sistemi sono collegati.

3 *Network layer*: permette a due sistemi, collegati anche a reti diverse, ma connesse, di comunicare fra di loro; per fare questo deve preoccuparsi di trovare il percorso da seguire tramite un protocollo di routing. Si occupa quindi dell'instradamento dei pacchetti, della loro frammentazione-ricostruzione, del controllo della congestione della rete.

4 *Transport layer*: stabilisce una trasmissione affidabile tra coppie di sistemi; risolve i problemi dovuti a perdita, duplicazione, riordinamento e frammentazione dei pacchetti.

5 *Session layer*: offre servizi basati sulla comunicazione affidabile e full-duplex offerta dal transport layer; stabilisce, governa e termina sessioni tra applicazioni.

⁶OSI: Open System Interconnections.

6 *Presentation layer*: fornisce uno standard per permettere alle applicazioni di concordare sulla rappresentazione dei dati.

7 *Application layer*: fornisce i programmi applicativi che utilizzano la rete.

Tra il modello TCP/IP e quello OSI esistono diverse corrispondenze; le più notevoli sono quelle tra l'Internet layer TCP/IP e il Network layer OSI ed anche quelle tra il Transport layer dei due modelli. Infatti è normale definire l'IP un protocollo del network layer OSI ed il TCP un protocollo del transport layer OSI, tant'è che tali verranno considerati nel seguito della tesi.

1.4 Internet in Italia: INFNet e GARR

Il protocollo TCP/IP è diventato operativo nel 1988; esso è stato adottato sulle reti degli enti di ricerca e delle università per connettersi a Internet.

Un'importante ente di ricerca in Italia è l'Istituto Nazionale di Fisica Nucleare (INFN). È organizzato in diciannove sezioni e quattro laboratori sparsi su tutto il territorio nazionale ed opera in stretta collaborazioni con i più importanti laboratori internazionali di ricerca come il CERN di Ginevra in Svizzera e il Fermilab negli USA.

Per garantire il collegamento tra tutte le sezioni ed i ricercatori, nel 1980 è stata realizzata INFNet, una rete di trasmissione dati basata su DECnet che raggiunge tutti i siti. Gli usi a cui la rete è destinata sono molteplici:

- la trasmissione dei risultati degli esperimenti effettuati nei laboratori ai ricercatori nelle loro sedi; i risultati possono essere grandi quantità di numeri o immagini le cui notevoli dimensioni richiedono linee ad alta velocità;
- la possibilità di usare in remoto risorse di calcolo che non sono disponibili in tutte le sedi;
- la possibilità di eseguire ed osservare esperimenti in remoto;
- lo scambio di informazioni tra i vari ricercatori;
- con l'avvento delle trasmissioni multicast di pacchetti audio e video, la possibilità di effettuare discussioni di gruppo tra vari ricercatori.

INFNet oggi supporta diversi protocolli tra i quali DECnet (fase 4), TCP/IP, DECnet/OSI (fase 5), X.25. Le linee ad alta velocità sono gestite

da multiplex TDM e da pre-ATM switch. In altre linee operano router multiprotocollo Cisco e DECnis.

La rete dell'INFN è stata per alcuni anni l'unica con un'estensione nazionale; ben presto altri enti di ricerca, in particolare i laboratori di Astronomia ed il CNR, hanno chiesto e ottenuto il permesso di poter utilizzare l'infrastruttura esistente.

Nel 1988 INFN, CNR, ENEA ed altri centri di ricerca fondarono il GARR⁷ e nel 1989 avviarono la realizzazione di un'unica rete nazionale che connettesse i maggiori centri di ricerca e le università. L'unione delle forze ha permesso di incrementare la velocità di trasmissione delle principali linee. La rete GARR oggi consiste infatti di un backbone di 2 Mbps ed i suoi principali siti sono a Milano, Bologna, Pisa, Roma, Gransasso, Frascati, Bari e Napoli. La rete INFNet è completamente integrata nella rete GARR e alcuni di quelli che erano i suoi collegamenti a bassa velocità (64 kbps) costituiscono oggi le linee di *backup*⁸ per il backbone.

In figura 1.6 è riportata la mappa dei siti italiani connessi alla rete GARR.

In figura 1.7 è riportata la mappa delle principali linee nazionali e dei collegamenti internazionali.

Dal 1983 INFNet è collegato al CERN e dal 1986 al Fermilab con il protocollo DECnet. La linea col Fermilab connette INFNet con ESnet e attraverso di essa all'Internet statunitense. In Europa INFNet è connessa, attraverso il CERN, alle reti europee, a DESY, a NORDUnet, a EUROPA-net e, direttamente, con i laboratori russi di Dubna.

1.5 Le connessioni internazionali del GARR

La rete GARR non è dunque isolata ma parte integrante di Internet; in essa il GARR costituisce un singolo autonomous system identificato dal numero 137. Un autonomous system è un insieme di reti che seguono una politica di routing comune. All'interno dell'autonomous system 137 il routing è garantito dai protocolli IGRP, Integrated IS-IS, OSPF, RIP, mentre il routing

⁷GARR: Gruppo Armonizzazione Reti della Ricerca.

⁸Linee di emergenza usate in caso di guasti alle linee principali.



Figura 1.6. Mappa dei siti connessi alla rete GARR.



Figura 1.7. Mappa delle principali connessioni di INFNet e del GARR.

tra gli autonomous system è effettuato tramite il protocollo BGP.

Al CNAF di Bologna e al Cnuce di Pisa hanno sede i due router che connettono il GARR a Internet: cnaf-gw-1.infn.it è collegato con il CERN (HEPnet, AS 513), con ESnet (AS 3425) e, tramite Garr-gw-gs (presso il Laboratorio Nazionale del Gran Sasso), con il laboratorio JINR di Dubna in Russia; garr-gw.cnr.it (Cnuce) è collegato con EUROPAnet (AS 2043).

1.6 Le prospettive future di Internet

La rete Internet ha subito in questi ultimi anni una diffusione inaspettata: il numero di host collegati cresce con un ritmo esponenziale, tanto che le scelte progettuali del protocollo IP si stanno rivelando inefficienti. Sono molti i colli di bottiglia che causano crolli di prestazioni, sprecano risorse hardware, richiedono interventi manuali da parte degli operatori; questi problemi sono dovuti al fatto che i progettisti dei protocolli TCP/IP non pensavano di dovere gestire una rete di tali dimensioni.

I principali problemi che si stanno verificando sono:

- esaurimento degli indirizzi di rete (soprattutto quelli di classe B, come spiegato più avanti);
- sovraccarico di informazioni di routing, specialmente per quei router che devono conoscere l'intera topologia di Internet;
- distribuzione degli indirizzi non gerarchica;
- algoritmi di routing inadeguati;
- fragilità della rete se sottoposta a manomissioni intenzionali.

Benché il numero di host indirizzabili con i trentadue bit di un indirizzo IP sia molto elevato (più di quattro miliardi) ed effettivamente il numero di host connessi a Internet sia molto minore, gli indirizzi ancora disponibili sono pochi e con l'attuale ritmo di crescita verranno esauriti in pochi anni. Questo succede poiché gli indirizzi di rete sono suddivisi in classi (cfr. sezione 2.2.1) ed ogni classe contiene da poche centinaia (classe C) a decine di migliaia (classe A) di indirizzi IP, causando uno spreco dei numeri disponibili. Infatti,

agli enti che fanno richiesta di connettersi all'Internet, il NIC⁹ assegna un numero di rete della classe che ritiene più che sufficiente ad indirizzare tutti gli host. È dunque normale che non tutti gli indirizzi IP di un indirizzo di rete vengano usati; si aggiunga poi che nei primi tempi il NIC, quando ancora non prevedeva lo sviluppo di Internet, ha distribuito indirizzi di classe B (fino a 65.535 host) ad enti con poche centinaia di host.

Un altro problema è quello del sovraccarico di informazioni di routing che i router che conoscono la topologia mondiale devono contenere. Queste macchine infatti devono conoscere decine di migliaia di indirizzi di rete e per ogni pacchetto che devono instradare devono consultare enormi database con conseguente calo delle prestazioni.

A questo fatto si sarebbe potuto ovviare se la distribuzione degli indirizzi fosse stata gerarchica, cioè se, ad esempio, tutti gli indirizzi di ogni autonomous system avessero un prefisso comune distinto dai prefissi degli altri autonomous system; in questo modo i router che effettuano il routing globale avrebbero dovuto conoscere solo i prefissi degli autonomous system. La distribuzione gerarchica però non è avvenuta, mentre, al contrario, è stata abbastanza casuale poiché non si è stati capaci di prevederne la necessità.

Una conseguenza di questi fatti è che i protocolli di routing si sono rivelati inadeguati a trattare efficientemente una tale mole di informazioni; si è cercato allora una soluzione adottando protocolli che utilizzano gli algoritmi link state, ma anche questi presentano degli inconvenienti, richiedendo alte capacità di calcolo ai router; un'altra soluzione adottata è stata la configurazione di route statiche, le quali però presentano le loro limitazioni nel caso i collegamenti smettano di funzionare.

Altro problema è la relativa fragilità della rete a manomissioni volontarie o involontarie: ad esempio, un router che introducesse informazioni errate o inconsistenti potrebbe rendere impossibile la consegna dei pacchetti alle destinazioni; oppure, un ente che si collegasse all'Internet usando un indirizzo di rete non assegnatogli dal NIC e che fosse già in uso, potrebbe causare conflitti non risolvibili, impedendo le comunicazioni.

⁹NIC: Network Information Center; è l'ente responsabile dell'assegnamento degli indirizzi di rete e dei numeri degli autonomous system.

Per questi ed altri motivi la comunità IP sta cercando di sviluppare un nuovo protocollo del network layer denominato *IP next generation* (IPng) che risolva questi problemi, che rimanga valido per molti anni in futuro e che sia possibile interfacciare ad altri protocolli quali il CLNP (cfr. sezione 2.1.2). È necessario anche che il servizio non venga interrotto: non essendo possibile che tutti gli host adottino la nuova soluzione nello stesso giorno, questa deve poter convivere con il protocollo IP attuale.

Sono state due le principali proposte che per lungo tempo si sono contese il diritto a succedere a IP versione 4 (IPv4): *TUBA*¹⁰ e *SIPP*¹¹. TUBA proponeva sostanzialmente di adottare il CLNP, il protocollo ISO del network layer, il che avrebbe portato sostanzialmente ad un unico protocollo globale. SIPP invece ha inventato un protocollo nuovo caratterizzato da un indirizzo gerarchico di sedici ottetti e dalla possibilità di utilizzare o meno header aggiuntivi, con il vantaggio di poter creare tutte le funzionalità che si rivelassero necessarie ma precludendo la strada ad un unico protocollo globale.

Dopo avere a lungo dibattuto sulle qualità dell'uno e dell'altro, nell'incontro avvenuto a Toronto (Canada) nel Luglio 1994, i membri dello IETF¹² hanno scelto SIPP per determinare le specifiche del nuovo protocollo che è stato denominato IPv6 (IP six); attualmente molti ricercatori stanno lavorando per puntualizzare meglio le caratteristiche ed iniziare a sviluppare il software di gestione.

¹⁰[RFC 1347]

¹¹[IPv6-spec]

¹²IETF: Internet Engineering Task Force.

2

Protocolli del network layer: IP e CLNP

IP¹ (Internet Protocol) e CLNP² (ConnectionLess Network Protocol) sono due protocolli del network layer (cfr. modello OSI, figura 1.5) L'IP è forse uno dei protocolli più usati nel mondo, mentre CLNP è quello realizzato dall'ISO e utilizzato già da alcuni software di rete come ad esempio DECnet/OSI.

2.1 I protocolli del Network Layer

Il network layer è definito come il livello in cui i protocolli permettono a due sistemi collegati ad una rete di comunicare fra di loro e si occupano dell'instradamento dei pacchetti, della loro frammentazione-ricostruzione, del controllo della congestione della rete (cfr. sezione 1.3).

2.1.1 IP: Internet Protocol

L'IP è un protocollo che fornisce un servizio di consegna di pacchetti non affidabile e senza connessione tra host appartenenti a reti locali anche diverse e distanti.

¹[RFC 791]

²[ISO 8473]

Le informazioni da trasmettere sono suddivise in blocchi e inserite in pacchetti (detti datagrammi) i quali vengono spediti ognuno indipendentemente dall'altro.

Il servizio è definito *non affidabile* perché la consegna non è garantita: il pacchetto può essere perso, duplicato, ritardato, o consegnato fuori sequenza, ma il servizio non rileverà tali condizioni, né informerà il trasmettitore o il ricevitore di tali malfunzionamenti. Il servizio è *senza connessione* poiché ogni pacchetto contiene l'indirizzo destinazione e viene instradato indipendentemente dagli altri.

2.1.2 CLNP: ConnectionLess Network Protocol

CLNP è il protocollo sviluppato dall'ISO per il Network layer; anch'esso come l'IP fornisce un servizio di consegna di pacchetti non affidabile e senza connessione.

Mentre il CLNP descrive come devono essere eseguite le tipiche funzioni del network layer, il CLNS³ descrive il servizio fornito al transport layer; questo è di tipo "best effort" e cioè non garantisce che i dati non vengano persi, corrotti, consegnati fuori sequenza, duplicati.

2.2 Gli indirizzi

L'operazione che avviene più comunemente sulle reti, è la connessione tra coppie di host che devono scambiarsi informazioni. Per stabilire la comunicazione deve essere possibile identificare i due host tra tutti quelli disponibili; per permettere questo, ad ognuno è assegnato un *indirizzo* (address): normalmente si tratta di un numero che codifica le informazioni necessarie a localizzare l'host nel dedalo delle reti connesse.

³CLNS: ConnectionLess Network Service.

2.2.1 Indirizzi IP⁴

L'IP permette la comunicazione tra due qualsiasi host connessi ad un internet. Per individuare un host è quindi necessario assegnargli un indirizzo che lo identifichi univocamente (per cui diverso da ogni altro indirizzo di host dell'internet). Tale indirizzo è costituito da un numero di 32 bit suddiviso in due campi: identificatore della rete (net-id) e identificatore dell'host (host-id) (figura 2.1).



Figura 2.1. Indirizzo IP.

Il *net-id* è il numero che identifica una rete locale mentre l'*host-id* è il numero che identifica un host all'interno della rete specificata dal net-id. Questa suddivisione permette di semplificare le operazioni di instradamento dei pacchetti e costituiscono l'unico livello di gerarchia di questi indirizzi.

Negli indirizzi IP un campo composto o da tutti zero o da tutti uno assume un particolare significato; in particolare tutti zero significano “questo”, mentre tutti uno significano “tutti”. Ad esempio, un indirizzo con campo host-id composto da tutti zero identifica la rete indicata nel net-id; se l'host-id è costituito da tutti uno allora l'indirizzo identifica tutti gli host della rete indicata nel net-id; un indirizzo con net-id di tutti zero identifica l'host indicato nel host-id che appartiene alla stessa rete dell'host che ha generato quell'indirizzo. Questo meccanismo è utile per indirizzare un pacchetto a più host oppure per comunicazioni rapide tra host di una stessa rete.

Più che un host, un indirizzo IP identifica un interfaccia di rete, per cui host connessi a più reti (cioè con più interfacce di rete) avranno più indirizzi che lo identificano.

La notazione degli indirizzi IP è composta da quattro numeri decimali separati da punti; ogni numero corrisponde a otto bit. Ad esempio l'indirizzo 10000011 10011010 00000001 00000101 è scritto normalmente 131.154.1.5

⁴[RFC 1060]

Classi

Nei diversi indirizzi, il campo net-id non ha sempre la stessa dimensione; in questo modo sono definibili reti che possono contenere un diverso numero massimo di host.

Gli indirizzi sono suddivisi in cinque classi a seconda della lunghezza dei campi net-id e host-id (figura 2.2); ognuna è identificata tramite i bit iniziali.

	0		7		15		23		31			
Classe A	0		Net-id			Host-id						
Classe B	1		0		Net-id			Host-id				
Classe C	1		1		0		Net-id			Host-id		
Classe D	1		1		1		0		Indirizzi multicast			
Classe E	1		1		1		1		0		Riservato	

Figura 2.2. Classi degli indirizzi IP.

Gli indirizzi multicast della classe D identificano gruppi di host; questo significa che un pacchetto che ha per destinazione un indirizzo multicast, giungerà a tutti gli host che appartengono al gruppo identificato da quell'indirizzo.

Subnetwork⁵

È possibile fare in modo che un singolo indirizzo di rete identifichi più di una rete fisica; questo avviene con una tecnica detta di *subnetting*. In pratica una parte del campo host-id viene utilizzata per identificare le singole reti fisiche. Per determinare quale parte dell'indirizzo è l'identificatore di sottorete è necessario accompagnarlo con una maschera di 32 bit, in cui gli uno identificano i bit dell'indirizzo di rete e sottorete, mentre gli zero l'host-id. Non è necessario che gli uno della maschera siano contigui.

Ad esempio, assegnato un indirizzo di rete di classe B 131.154.0.0, è possibile utilizzarlo per identificare varie sottoreti usando il terzo numero decimale come identificatore di sottorete; la maschera che dovrà accompagnarlo sarà

⁵[RFC 950]

255.255.255.0 (255 corrisponde al numero binario 11111111). In questo modo si ricavano 255 sottoreti di dimensioni uguali a quelle di classe C.

Supernetwork⁶

Ormai gli indirizzi di classe B sono esauriti; purtroppo però quelli di classe C per molte realtà hanno un numero insufficiente di indirizzi di host. Si è trovata una soluzione temporanea a questo problema con la tecnica di *supernetting*: ad una rete vengono assegnati diversi indirizzi di classe C contigui, cioè con un prefisso comune, e tali che non esistano altre reti di classe C con quel prefisso. Accompagnando l'indirizzo con una maschera che evidenzia questo prefisso comune, è possibile considerare quell'insieme di reti come una rete unica.

2.2.2 Indirizzi CLNP

Il servizio del network layer è fornito al transport layer attraverso un punto di accesso (concettuale) localizzato al confine tra il transport layer ed il network layer; tale punto è chiamato NSAP⁷. Esiste un NSAP per ogni protocollo del transport layer.

Ogni NSAP può essere raggiunto nella internet ISO tramite il suo NSAP address. Se in un end system⁸ sono attivi più protocolli del transport layer, esso avrà tanti NSAP address quanti sono i protocolli; essi differiranno però solamente nell'ultimo byte (il Selector).

A volte è utile poter utilizzare l'indirizzo di un host senza far riferimento al protocollo del transport layer, come, ad esempio, nel caso di un host che esegue un protocollo di routing; tale indirizzamento è eseguito tramite il NET⁹. Il NET coincide con l'NSAP address in tutto tranne nel campo selector, che è uguale a zero.

⁶[RFC 1519]

⁷NSAP: Network Service Access Point.

⁸End System (ES) è il termine ISO per indicare un host che non esegue funzioni di routing.

⁹NET: Network Entity Title.

La maggior parte degli end system e degli intermediate system¹⁰ hanno un unico NET, diversamente dai router IP che hanno un indirizzo per ogni interfaccia.

IDP		DSP	
AFI	IDI	DSP	
Area		ID	Sel

Figura 2.3. NSAP address.

Gli NSAP address sono assegnati gerarchicamente per semplificare il routing. Un NSAP address ha una lunghezza variabile (ma generalmente è lungo venti ottetti) ed è principalmente diviso in due parti: l'*Initial Domain Part* (IDP) e il *Domain Specific Part* (DSP). L'IDP è ulteriormente diviso nell'*Authority and Format Identifier* (AFI) e l'*Initial Domain Identifier* (IDI).

L'AFI fornisce informazioni a riguardo della struttura e del contenuto dei campi IDI e DSP, in particolare se l'IDI ha lunghezza variabile o se il DSP usa la notazione decimale o quella binaria. Il campo IDI specifica l'autorità responsabile di stabilire il formato del DSP. Il DSP è infatti anch'esso suddiviso in sottocampi: una prima parte contribuisce, insieme all'IDP, a determinare l'area di appartenenza; la rimanente è divisa nell>ID, che identifica un end system all'interno di un'area, e nel selector.

In figura 2.3 è riportata lo schema di un NSAP address.

2.2.3 Indirizzamento piatto e gerarchico

Si parla di *indirizzamento piatto* quando il numero che costituisce l'indirizzo di un host è significativo per raggiungerlo solo se considerato nella sua interezza.

Si parla invece di *indirizzamento gerarchico* quando è possibile utilizzare sottoparti sempre più grandi di un indirizzo per avvicinarsi sempre più alla

¹⁰Intermediate System (IS) è il termine ISO per indicare un host che esegue funzioni di routing.

destinazione, ovvero quando tutti gli host di una zona hanno un prefisso comune sempre più piccolo man mano che si considera una zona più grande.

Un indirizzo è piatto o gerarchico non per come è fatto ma per come viene assegnato. Gli indirizzi IP infatti sono piatti poiché, quando sono stati distribuiti, non ci si è preoccupati di assegnare indirizzi con prefissi uguali a reti che risiedevano nello stesso autonomous system, in modo da poter identificare con un unico numero tutti gli host interni. Gli indirizzi NSAP invece sono gerarchici perché fin dalla loro ideazione si è imposto che prefissi sempre più lunghi avrebbero individuato insieme sempre più piccoli di host contigui.

L'indirizzamento gerarchico è sicuramente migliore di quello piatto perché permette di riassumere le informazioni di routing diminuendo la mole di dati che i router devono memorizzare e gestire.

2.2.4 Confronto tra gli indirizzi IP e CLNP

La principale differenza tra i due indirizzi è la lunghezza: quattro ottetti il primo, normalmente venti il secondo; calcolando il numero di indirizzi possibili, sembrerebbe che quattro ottetti siano più che sufficienti per indirizzare tutti i computer della Terra, ma in realtà, a causa della cattiva politica di allocazione esposta nel paragrafo 1.6, questi indirizzi stanno per essere esauriti. I venti ottetti dell'NSAP address possono apparire una esagerazione, ma permettono di gestire un routing gerarchico soddisfacendo le esigenze del routing globale e contemporaneamente quelle di ogni singolo ente.

Altra notevole differenza è quella già esposta per cui gli indirizzi IP sono piatti, mentre quelli NSAP sono gerarchici.

Sia gli indirizzi IP che quelli ISO sono suddivisi in due campi: mentre in IP il primo campo identifica tutti gli host di una singola rete fisica ed il secondo un singolo host di quella rete fisica, negli NSAP la prima parte è gerarchica ed identifica un'area (insieme di reti fisiche) mentre la seconda è piatta e identifica un host che appartiene ad un gruppo di reti fisiche.

2.3 I pacchetti di dati

Nelle reti fisiche la trasmissione dei dati avviene tramite frame di bit, cioè sequenze di bit ben delimitate e accompagnate da alcune informazioni di controllo. I protocolli del network layer sfruttano questi pacchetti messi a disposizione dal data link layer inserendovi, nella parte riservata ai dati, blocchi di informazioni e valori di controllo.

2.3.1 Datagramma IP

Il *datagramma* è l'unità fondamentale di trasferimento del protocollo IP. I datagrammi sono costituiti da due parti fondamentali: una *intestazione* (header) e un'area *dati*; l'intestazione contiene informazioni necessarie per far giungere il pacchetto integro a destinazione, quali la dimensione e gli indirizzi sorgente e destinazione; l'area dati contiene invece i dati da trasmettere. Un datagramma IP è incapsulato in un frame della rete fisica di cui occupa l'area dati (figura 1.3).

0	4	8	16	19	24	31
Version	H. len.	Type of Service	Total Length			
Identification			Flag	Fragment offset		
Time To Live		Protocol	Header checksum			
Source address						
Destination address						
Options					Padding	

Figura 2.4. Intestazione di un datagramma IP.

I campi dell'intestazione sono riportati in figura 2.4 e sono:

Version: contiene la versione del protocollo IP usata per creare il datagramma; la versione attuale è la 4.

Header length: indica la lunghezza dell'intestazione misurata in multipli di 32 bit.

Type of Service (TOS): specifica il modo in cui deve essere gestito il datagramma ed è suddiviso in alcuni sottocampi: tre bit per la *precedence* che indica la precedenza nella trasmissione e va da 0 (precedenza normale) a 7 (controllo di rete); quattro bit *D*, *T*, *R*, *C* che specificano la caratteristica da privilegiare per la trasmissione: *D* indica di minimizzare il ritardo (Delay), *T* di massimizzare il throughput, *R* di massimizzare l'affidabilità (reliability), *C* di minimizzare il costo (cost); l'ultimo bit non è definito.

Total length: indica la lunghezza totale del datagramma misurata in ottetti.

Identification: è un numero sequenziale imposto dall'host che ha generato il pacchetto e serve per riassemblare eventuali pacchetti frammentati.

Fragment offset: quando, durante il cammino verso l'host destinazione, la lunghezza di un datagramma supera quella massima consentita da una rete intermedia, il datagramma può essere frammentato in datagrammi più piccoli. Se il datagramma è un segmento di un datagramma più grande, questo campo specifica la posizione nel datagramma iniziale.

Flag: sono flag usati per la frammentazione-ricostruzione dei datagrammi.

Time To Live (TTL): indica il numero di secondi per cui il datagramma può restare nell'Internet prima di venire eliminato; esso è decrementato da ogni router che tratta il datagramma. Questo contatore è utile per impedire che un pacchetto malformato o destinato a reti irraggiungibili vaghi nell'Internet per un tempo indeterminato.

Protocol: indica il protocollo del transport layer che sta usando l'IP; ad esempio, il numero 6 identifica il protocollo TCP.

Checksum: è un numero calcolato eseguendo un'operazione nota sui byte dell'intestazione; serve per assicurarsi che l'intestazione non è stata corrotta durante la trasmissione.

Source address: contiene l'indirizzo IP dell'host che ha generato il pacchetto.

Destination address: contiene l'indirizzo IP dell'host a cui deve essere recapitato il pacchetto.

Options: è un campo a lunghezza variabile e contiene richieste di particolari trattamenti nell'instradamento del pacchetto; serve soprattutto per il test e il debug del funzionamento di una rete.

Padding: sono bit settati a zero necessari per compensare la lunghezza va-

riabile del campo Options affinché l'header sia di lunghezza multipla di 32 bit.

2.3.2 Pacchetto CLNP

Come in IP, il pacchetto è l'unità fondamentale di trasferimento. I pacchetti CLNP sono abbastanza simili a quelli IP; sono costituiti da due parti fondamentali: una *intestazione* (header) e un' *area dati*; l'intestazione contiene informazioni necessarie al routing, mentre l'area dati contiene le informazioni da trasmettere.

Network Layer protocol ID				1
Length Indicator				1
Version				1
Lifetime				1
SP	MS	E/R	Type	1
Segment length				2
Header checksum				2
Dest. address length indicator				1
Destination address				1-20
Source address length indicator				1
Source address				1-20
Data unit identifier				2 o 0
Segment offset				2 o 0
Total length				2 o 0
Options				variabile

Figura 2.5. Intestazione di un pacchetto CLNP; a fianco di ogni campo è indicata la dimensione in ottetti.

I campi dell'intestazione sono riportati in figura 2.5 e sono:

Network Layer protocol ID: identifica il protocollo del Network layer che ha generato il pacchetto.

Length indicator: indica la lunghezza dell'intestazione in ottetti.

Version: indica la versione del protocollo che ha generato il pacchetto.

Lifetime: indica il numero di mezzi secondi per cui il pacchetto può restare nella rete prima di venire eliminato; esso è decrementato da ogni intermediate system che tratta il pacchetto.

SP: è un flag che indica la presenza o meno del campo Data unit identifier.

MS: è un flag che indica se il pacchetto è stato frammentato e se non è l'ultimo frammento.

E/R: è un flag che, se settato, indica che la sorgente richiede di sapere se un pacchetto non è stato consegnato.

Type: è un numero che indica il tipo di pacchetto; 28 indica un normale pacchetto di dati; 1 indica un pacchetto con indicazioni di errori.

Segment length: indica la lunghezza del pacchetto in ottetti; se il pacchetto è stato frammentato indica la lunghezza del frammento.

Header checksum: è un numero calcolato eseguendo un'operazione nota sui byte dell'intestazione; serve per assicurarsi che l'intestazione non sia stata corrotta durante la trasmissione.

Destination address length indicator: poiché gli NSAP address non hanno una lunghezza fissa, questo campo indica il numero di ottetti che compongono l'indirizzo dell'host destinazione.

Destination address: contiene l'NSAP address dell'host destinazione.

Source address length indicator: poiché gli NSAP address non hanno una lunghezza fissa, questo campo indica il numero di ottetti che compongono l'indirizzo dell'host sorgente.

Source address: contiene l'NSAP address dell'host sorgente.

Data unit identifier: è un numero assegnato ai pacchetti dall'host sorgente, diverso per ogni pacchetto; serve per permettere la frammentazione/ricostruzione di un pacchetto.

Segment offset: indica la posizione del frammento nel pacchetto originale.

Total length: indica la lunghezza totale iniziale di un pacchetto frammentato.

Questi ultimi tre campi sono presenti solo se il pacchetto originale è stato frammentato.

Options: contiene alcune richieste per particolari funzioni di routing: tipo di servizio, source routing, security. Il tipo di servizio, o *Quality of Service*

(QoS), è codificato in maniera differente dall'IP; la presenza dell'ottetto del QoS è indicata dall'Option code posto uguale a 201; se l'ottetto inizia con 11 (binario) è selezionato il *globally defined quality of service option* per cui il pacchetto è trasmesso il più velocemente possibile; altrimenti diventano rilevanti tre bit così definiti: *D/C transit delay vs. cost* indica di favorire percorsi a minor ritardo a discapito del costo; *E/D residual error probability vs. delay* indica di favorire i percorsi a minor probabilità di errore a discapito del ritardo; *E/C residual error probability vs. cost* indica di favorire i percorsi a minor probabilità di errore a discapito del costo. Le combinazioni dei tre danno le diverse possibili precedenze.

2.3.3 Confronto fra le intestazioni dei pacchetti IP e CLNP

Le intestazioni dei pacchetti IP e CLNP sono abbastanza simili, sia come funzioni che come campi. Una delle differenze forse fondamentale è che alcuni campi dell'header CLNP possono essere presenti o meno; questo dover verificare la presenza o meno di un campo diminuisce le prestazioni dei router.

Per il resto i campi con la stessa funzione si differenziano per minimi particolari: il calcolo della checksum, la determinazione del tipo di servizio da privilegiare, l'unità di misura della lunghezza della intestazione e del tempo di vita. Il campo protocol dell'intestazione IP non è presente in quella CLNS poiché coincide con il campo Sel degli NSAP address.

L'intestazione dei pacchetti CLNP ha dei campi ridondanti o inutili.

2.4 Il routing

Una delle funzioni principali dei protocolli del Network layer è quella di trasportare i pacchetti da un host sorgente all'host destinazione. Per fare questo è necessario che particolari host siano in grado di effettuare il *routing* e cioè di conoscere la topologia della rete. Questi host speciali sono i *router* (nella terminologia IP) e gli *intermediate system* (IS) (nella terminologia CLNP) e la loro particolarità è quella di essere collegati a più di una rete

fisica.

Un pacchetto che deve essere trasmesso passa dall'host sorgente al router (o IS) della stessa rete fisica; poi viene passato da router a router (da IS a IS) finché non giunge a quello che risiede sulla stessa rete dell'host destinazione al quale viene consegnato.

Per compiere queste operazioni, ogni router deve possedere informazioni sulla locazione di ogni host dell'internet, informazioni che però sono parziali: un router conosce solamente il router successivo da attraversare per giungere ad una certa destinazione e non tutta la catena dei router; è a questo router che viene passato il pacchetto, il quale dovrà fare altrettanto. Queste informazioni sono mantenute in speciali tabelle (database) dette *routing table*.

Il routing avviene a diversi livelli per ridurre la mole di informazioni topologiche; i livelli sono denominati *neighbor greeting*, *intradomain* e *interdomain*.

2.4.1 Neighbor greeting

Al primo livello del routing sta il *neighbor greeting* e cioè il processo per cui gli host riconoscono gli indirizzi direttamente raggiungibili, scoprono i router adiacenti e associano gli indirizzi del network layer degli host ai rispettivi indirizzi del data link layer.

ARP¹¹

L'ARP (Address Resolution Protocol) è un protocollo per l'IP che permette di stabilire una corrispondenza tra un indirizzo IP di un host e il suo indirizzo fisico in una LAN (ad esempio, il numero della sua scheda ethernet).

Un host che vuole comunicare con un altro e di cui conosce l'indirizzo IP, necessita dell'indirizzo fisico della destinazione; per ottenerlo trasmette in broadcast un *ARP query*, cioè un messaggio che riporta l'indirizzo IP dell'host destinazione e che tutti gli host della LAN ricevono. L'host che individua il suo indirizzo IP nella query risponde con un *ARP response* che contiene il suo indirizzo fisico.

¹¹[RFC 826]

Ogni host mantiene una cache di dimensione finita di coppie <indirizzo IP, indirizzo fisico> per non dover inviare una query per ogni pacchetto che deve spedire.

L'ARP è utilizzato da un host quando, osservando il net-id dell'indirizzo destinazione, questi si accorge che risiede sulla stessa LAN dell'host con cui vuole comunicare. Altrimenti il pacchetto viene passato al router della network che utilizzerà un protocollo di routing di livello superiore. Il router a cui passare i pacchetti per altre LAN è definito staticamente.

ES-IS¹²

L' ES-IS (End System to Intermediate System) è un protocollo per il CLNP che permette ad un end system di raggiungere un altro end system sulla stessa LAN o, tramite un intermediate system, un end system della LAN con lo stesso intermediate system in comune.

Fornisce tre tipi di messaggi: *End System Hello* (ESH), trasmesso periodicamente da ogni end system per annunciare la loro presenza sulla LAN; *Intermediate System Hello* (ISH), trasmesso periodicamente dagli intermediate system per annunciare la loro presenza sulla LAN; *Redirect*, utilizzato da un intermediate system che vuole informare un end system di utilizzare un altro IS per instradare un pacchetto per una determinata destinazione.

Ogni end system pone gli indirizzi degli intermediate system che hanno trasmesso un ISH nella propria *router cache*; ogni end system mantiene anche una *destination cache* in cui sono memorizzate le coppie <indirizzo network layer, indirizzo data link layer> degli altri end system sulla LAN che hanno spedito un ESH.

Quando un end system deve spedire un pacchetto controlla se la destinazione è nella destination cache e se la trova trasmette il pacchetto direttamente, altrimenti lo trasmette ad un intermediate system il cui indirizzo è nella router cache. Se questo intermediate system conosce un intermediate system migliore per instradare tale pacchetto, lo comunica all'end system tramite un messaggio di Redirect.

¹²[ISO 9542]

Nel caso che l'indirizzo dell'end system destinazione non sia nella destination cache e non vi siano indirizzi di intermediate system nella router cache (questo può essere il caso di una LAN isolata, senza router), allora il pacchetto da trasmettere viene inviato all'indirizzo di gruppo che comprende tutti gli end system, con indicato l'indirizzo di network layer dell'end system sconosciuto. Se un end system che riceve un pacchetto indirizzato a tutti gli end system scopre il suo indirizzo all'interno di tale pacchetto allora invia immediatamente un ESH all'end system che ha generato il pacchetto, che può in tal modo aggiornare la sua destination cache.

L'ES-IS fornisce un meccanismo di autoconfigurazione per cui un end system è in grado di determinare il suo indirizzo di area dall'IS a cui fa riferimento.

Confronto tra ARP e ES-IS

Esiste una sostanziale differenza tra la soluzione IP e quella CLNP: in IP il riconoscimento dei vicini avviene su una singola LAN, cioè un host è in grado di indirizzare pacchetti a host che hanno il medesimo indirizzo di rete, altrimenti si rivolge al router che risiede sulla LAN, il quale poi usa un altro protocollo di routing; in CLNP un end system è invece in grado di indirizzare end system che risiedono anche su LAN diverse sempre che queste abbiano un intermediate system in comune.

Tutti e due i protocolli visti introducono una buona dose di messaggi di controllo sulle LAN: nell'ARP è dovuto alle ARP query e alle risposte; nell'ES-IS è dovuto agli ESH. Nella soluzione IP inoltre tutti gli host devono esaminare le query e buttare quelle che non li riguardano sprestando tempo della CPU, cosa che invece non accade per CLNP.

2.4.2 Algoritmi di routing

I protocolli di routing usano essenzialmente due algoritmi per calcolare le tabelle di routing: il *distance vector* e il *link state* (anche detto *shortest path first*).

Distance Vector

Questo algoritmo è derivato da quello noto come “Bellman-Ford”. Ogni router che esegue il protocollo determina quali router gli sono vicini, cioè quali router sono direttamente connessi a lui, ed invia a tutti questi un vettore (il vettore delle distanze) con l’elenco delle reti a cui lui stesso è connesso.

Il *vettore delle distanze* è costituito da coppie contenenti l’indirizzo di una rete destinazione e la distanza da essa, distanza che ha per unità di misura il numero di router da attraversare (*hop*) per raggiungerla.

Dai vettori che ha ricevuto e che ha generato, ogni router costruisce la propria tabella di routing costituita da triple: la rete destinazione, la distanza da essa, il primo router del percorso ad essa. Dalla tabella di routing viene costruito il nuovo vettore da inviare ai vicini e contenente le indicazioni di tutte le reti conosciute.

Quando un router riceve un nuovo vettore controlla che tutte le reti riportate siano già nella sua tabella di routing: se una network non c’è, viene aggiunta con la distanza riportata incrementata di uno e, come primo router del percorso, con il router che ha inviato il vettore; se c’è già, vengono confrontate le distanze e se quella del vettore è più breve allora la tabella di routing viene modificata, mettendo come distanza quella del vettore incrementata di uno e come primo router del percorso quello che ha inviato il vettore; altrimenti la tabella rimane invariata.

Dopo un certo periodo di tempo (detto di convergenza) le tabelle di routing diventano stabili e consistenti. I router continuano però a scambiarsi i vettori ad intervalli regolari, nel caso che qualche collegamento si interrompa e quindi sia necessario aggiornare le tabelle di routing.

I principali difetti di questo algoritmo sono che il tempo di convergenza può essere relativamente lungo, per cui si possono avere istanti in cui le tabelle dei vari router sono inconsistenti e causano dei percorsi ad anello (*loop*); inoltre, se le reti sono molte i vettori delle distanze risultano essere molto grandi, causando un notevole incremento del traffico sull’internet e quindi un rallentamento generale.

I principali vantaggi sono la relativa semplicità di implementazione ed il fatto che, essendo il primo algoritmo ad essere stato utilizzato, è ben col-

laudato.

Link State

L'algoritmo link state è derivato da quello ideato da Dijkstra e determina in un grafo pesato i cammini più brevi da un nodo a tutti gli altri. Nelle internet, i nodi sono i router, i collegamenti tra i nodi sono i collegamenti fisici, i pesi sono le metriche assegnate ad ogni collegamento. Da questi percorsi viene costruita la tabella di routing in cui sono riportate le destinazioni, le distanze ad esse e il primo router del cammino a quella destinazione. In IP le destinazioni sono gli indirizzi di rete, in CLNP sono i numeri di area o sottoparti di questi.

È fondamentale che i grafi ricostruiti da tutti i router siano uguali per garantire la consistenza delle tabelle di routing. Per fare questo inizialmente ogni router si occupa di scoprire i suoi vicini e di conoscerne il nome; dopodiché costruisce un pacchetto detto Link State Packet (LSP) in cui sono elencati i vicini con il costo del collegamento. I LSP sono inviati periodicamente a tutti gli altri router ed ogni router colleziona l'ultimo LSP generato da ogni altro router. Sulla base dei LSP collezionati, ogni router ricostruisce la topologia della rete (cioè il grafo pesato) e calcola i cammini di costo minimo per ogni destinazione.

I principali vantaggi di questi algoritmo sono che converge molto rapidamente e che il traffico introdotto sulla rete è molto ridotto. Per contro, l'algoritmo che calcola i cammini di costo minimo impegna notevolmente la CPU. Attualmente però, i protocolli basati su questo algoritmo sono preferiti perché la convergenza veloce è utile a garantire i percorsi di backup; inoltre, la capacità di calcolo dei router ha raggiunto livelli tali da eseguire velocemente l'algoritmo anche per grafi con molti nodi.

2.4.3 Intradomain routing protocol

Al secondo livello dell'instradamento stanno i protocolli detti *intradomain*, cioè interni ad un dominio.

Le internet possono avere dimensioni enormi e causare diversi problemi ai protocolli di routing; ad esempio Internet è estesa a livello mondiale e

comprende migliaia di reti, per cui per raggiungere tutte le destinazioni ogni router dovrebbe mantenere tabelle di routing gigantesche. Si è visto però che la maggior parte del traffico è tra reti relativamente vicine (come ad esempio quelle di una nazione) e dunque è inutile che tutti i router siano a conoscenza della intera topologia.

Un internet può allora essere suddivisa in zone (dette *autonomous system* in IP e *routing domain* in CLNP) dove i router che vi appartengono conoscono solo la topologia interna. Solo pochi router per ogni autonomous system conoscono la topologia dell'intera internet e ad essi i router interni si rivolgono per instradare pacchetti diretti a host appartenenti ad altri autonomous system. Gli intradomain routing protocol sono allora quei protocolli utilizzati dai router interni per conoscere la topologia dell'autonomous system e costruirsi le tabelle di routing.

RIP¹³

Il RIP (Routing Information Protocol) è un protocollo per IP che ebbe una grande diffusione nei primi anni di vita di Internet, ma che oggi viene progressivamente abbandonato perché poco adatto ad autonomous system di grandi dimensioni.

RIP utilizza l'algoritmo distance vector. Un router che lo esegue trasmette un vettore delle distanze ogni trenta secondi; la distanza è misurata in salti (hop) cioè con il numero di router da attraversare per giungere a destinazione. Non tutte le network però lavorano alla stessa velocità per cui un percorso di due salti che attraversa due linee seriali non è detto che sia più veloce di uno che attraversa tre ethernet. Per risolvere questo problema, alcune implementazioni effettuano un conteggio dei salti artefatto, per cui connessioni lente misurano più di un salto.

Per impedire che un percorso che non esiste più rimanga in una tabella di routing, ad ognuno di essi è associato un timer che viene riavviato ogni volta che un vettore conferma quel percorso. Se non dovesse essere confermato entro centottanta secondi, esso viene scartato.

L'instabilità si verifica quando la caduta di un collegamento provoca un

¹³[RFC 1058]

aumento a catena delle distanze dei percorsi che attraversano quello stesso collegamento, sino a raggiungere la massima distanza consentita. Per prevenire l'instabilità il numero massimo di salti in un percorso è di sedici. È un numero molto basso che impedisce l'utilizzo di RIP in autonomous system di grandi dimensioni.

Durante i periodi di instabilità le tabelle di routing sono inconsistenti e si possono verificare dei loop, cioè dei percorsi ad anello che non portano alla destinazione. Per evitare questi loop sono state realizzate alcune tecniche. La *split horizon* prevede di differenziare i vettori che un router invia ai vicini, in particolare di non inviare ad un vicino quei percorsi che sono stati appresi da vettori inviati dallo stesso vicino. La *poison reverse*, quando una connessione si interrompe fa in modo che il router connesso l'annunci per alcune volte con un costo infinito. La *holddowns*, quando riceve un messaggio di irraggiungibilità di una rete, non accetta, per un certo periodo, informazioni riguardanti quella rete, in modo da consentire all'informazione di irraggiungibilità di raggiungere tutti i router.

È in definitiva un protocollo robusto e collaudato ma che converge lentamente, richiede un'eccessiva larghezza di banda nell'invio dei vettori, è inadatto ad autonomous system di grandi dimensioni. Tutti questi difetti ne stanno causando un progressivo abbandono.

IGRP¹⁴

L'IGRP (Interior Gateway Routing Protocol) è un protocollo per IP basato sull'algoritmo distance vector. È un miglioramento del RIP di proprietà della Cisco, abbastanza diffuso e soggetto tuttora a miglioramenti.

Innanzitutto può essere usato in grandi autonomous system con topologia complessa senza il limite dei sedici salti; può scegliere come default route quella con metrica più bassa tra alcuni autonomous system boundary router contrassegnati.

Per determinare il costo di un collegamento usa una combinazione pesata di quattro costi: ritardo, larghezza di banda, affidabilità e carico; permette di avere più strade con metriche che rientrano in un certo range per la stessa

¹⁴[Hedrick]

destinazione.

Per prevenire i percorsi errati usano i triggered updates e gli holddowns: un *triggered update* è un nuovo vettore spedito immediatamente da un router che riscontra un cambiamento. *Holddowns* è come in RIP e permette alle informazioni dei triggered update di diffondersi.

Per prevenire i loop usa split horizon e route poisoning: split horizon è come per RIP. *Route poisoning* controlla nei vettori se le metriche subiscono aumenti consistenti rispetto ai valori precedenti; in tal caso suppone che su un tale percorso si sia verificato un loop: lo rimuove dalla tabella di routing e lo pone in holddown.

Recentemente è stato sviluppato l'EIGRP (Enhanced Interior Gateway Routing Protocol), un estensione dell'IGRP di cui aumenta l'efficienza e la velocità di convergenza. I miglioramenti tendono a minimizzare le informazioni scambiate tra i router e il calcolo delle tabelle di routing: quando lo stato di un collegamento cambia viene inviato un vettore contenente solo le variazioni e non l'intero elenco delle destinazioni; in più effettua la trasmissione affidabile di pacchetti solo se necessario. Inoltre, classifica alcuni vicini come *successor* e cioè sicuramente non appartenenti ad un loop: se avviene un cambiamento sui link i percorsi vengono ricalcolati solo se queste modifiche coinvolgono tutti i successor.

Un difetto che però entrambi non sono in grado di risolvere è la bassa velocità di convergenza dovuta al tipo di algoritmo che usano. Un ulteriore limite è il fatto di essere protocolli proprietari, per cui disponibili solo su macchine Cisco; questo ne impedisce un utilizzo su reti con macchine di produttori diversi, fatto molto frequente negli autonomous system estesi.

OSPF¹⁵

OSPF (Open Shortest Path First) è un protocollo per IP che utilizza l'algoritmo link state.

Suddivide il routing di un autonomous system in due livelli: al primo

¹⁵[RFC 1583]

livello ci sono le aree, insiemi di reti locali il cui routing è affidato a router interni che non conoscono la topologia del resto dell'autonomous system; al secondo livello sta il backbone, un insieme di router che garantisce il routing tra le aree. Backbone e aree si intersecano negli area border router. In ogni area e sul backbone sono eseguite copie indipendenti dell'algoritmo. Questa suddivisione riduce il traffico introdotto sulla rete e velocizza il calcolo delle tabelle di routing.

Ulteriori dettagli su OSPF e gli altri protocolli link state sono nel capitolo 3.

IS-IS¹⁶

IS-IS (Intermediate System to Intermediate System) è l'unico protocollo intradomain per CLNP ed usa l'algoritmo link state.

Suddivide il routing di un routing domain in due livelli: al livello più basso stanno tutti gli end-system e gli intermediate-system con lo stesso Area-id; al suo interno il routing è effettuato dagli IS interni (che vengono detti di Livello 1) sulla base del campo ID del NSAP address. Al secondo livello stanno gli IS di livello 2 che effettuano il routing tra le varie aree sulla base del campo Area del NSAP address.

Quando un router di livello 1 riceve un pacchetto, esamina il campo "indirizzo di area" dell'indirizzo destinazione e, se questo coincide con il suo indirizzo di area, lo instrada basandosi sul campo ID; altrimenti lo invia al router di livello 2 più vicino.

Integrated IS-IS¹⁷

Integrated IS-IS è la versione del protocollo IS-IS in grado di effettuare il routing contemporaneamente sia per CLNP che per IP.

L'algoritmo di routing è applicato una sola volta per entrambi i protocolli al grafo costituito dagli IS, i quali sono identificati dai loro NSAP address; in più, per la parte IP, ogni IS aggiunge ai propri LSP l'elenco degli indirizzi

¹⁶[ISO 10589]

¹⁷[RFC 1195]

IP di rete a cui è connesso. Vengono poi costruite due distinte tabelle di routing, una per IP e una per CLNP.

2.4.4 Interdomain routing protocol

Al livello più alto del routing stanno i protocolli detti di *interdomain*, cioè che collegano i vari autonomous system (routing domain).

In ogni autonomous system, un piccolo numero di router esegue un ulteriore protocollo di routing in cui impara la topologia globale dell'internet. È a questi router che quelli interni si rivolgono quando devono consegnare pacchetti a host che non risiedono nel loro stesso autonomous system.

EGP¹⁸

L'EGP (Exterior Gateway Protocol) fu il primo protocollo per IP ad essere utilizzato per connettere i diversi autonomous system; esso è impiegato dai router vicini esterni, cioè da router connessi insieme ed appartenenti ad autonomous system differenti.

È stato progettato per una topologia in cui un insieme dei router esterni costituiscono un nucleo e si scambiano informazioni su tutte le destinazioni che conoscono; al protocollo partecipano poi anche tutti gli altri router esterni i quali però annunciano solo le destinazioni interne all'autonomous system a cui appartengono.

EGP usa un algoritmo di routing simile al distance vector con la differenza che la misura della distanza riportata nei messaggi non è considerata per determinare il cammino più breve, ma solamente per specificare l'esistenza di un percorso ad una certa rete (distanza 255 significa che una rete è irraggiungibile).

Non è quindi possibile effettuare un instradamento vero e proprio; la struttura dei cammini che viene ricostruita è un albero con i router del nucleo alla radice e ogni destinazione alle foglie. Quindi per ogni destinazione esiste un solo percorso.

¹⁸[RFC 904]

Questa struttura ha vari problemi; ad esempio, un guasto nel nucleo può pregiudicare l'intero sistema; inoltre, essendo unico il cammino, tutto il traffico per una certa destinazione si riverserà su un unico collegamento anche se ne esistono altri. Un altro problema dell'EGP sono le grosse dimensioni dei messaggi di aggiornamento che i router del nucleo si scambiano. Attualmente questo protocollo è sempre meno utilizzato e viene sostituito da BGP.

BGP¹⁹

Il BGP (Border Gateway Protocol) è il protocollo per IP che fu sviluppato per sopperire ai limiti dell'EGP. Usa un algoritmo distance vector in cui il costo di un cammino corrisponde al numero degli autonomous system da attraversare.

Ogni router calcola il suo percorso migliore per un particolare indirizzo IP e lo comunica ai suoi vicini. Insieme a questo percorso sono inclusi alcuni attributi, tra cui il più importante è l'*autonomous system path* che riporta il numero di sistemi autonomi nel percorso e che viene utilizzato per sopprimere i loop.

Per diminuire la banda occupata dalle informazioni di routing BGP usa gli *aggiornamenti incrementali*: dopo un iniziale scambio completo di informazioni di routing, le coppie di router si informano solo dei cambiamenti che avvengono. Questa tecnica però richiede che le comunicazioni avvengano in maniera affidabile per cui il BGP utilizza il transport layer per trasmettere le informazioni.

IDRP²⁰

L'IDRP (InterDomain Routing Protocol) fu sviluppato dal BGP per essere utilizzato con gli indirizzi ISO ed utilizza un algoritmo distance vector.

I router che partecipano all'IDRP sono detti *Boundary Intermediate System* (BIS) e ce ne deve essere almeno uno per ogni routing domain; due BIS sono vicini esterni se appartengono a routing domain differenti ma sono di-

¹⁹[RFC 1654]

²⁰[ISO 10747]

rettamente connessi. La metrica utilizzata è quella del numero di routing domain da attraversare.

IDRP realizza un protocollo di consegna affidabile simile ad un protocollo del transport layer.

3

Protocolli di routing link state

I protocolli di routing usano essenzialmente due algoritmi per calcolare i percorsi migliori: il *distance vector* e il *link state*.

Il distance vector fu il primo algoritmo ad essere utilizzato per la sua relativa semplicità: fu adottato da RIP, HELLO, IGRP. Esso ha però una grossa limitazione che è la sua bassa velocità di convergenza.

I più recenti protocolli, come IS-IS e OSPF, usano invece l'algoritmo link state che converge molto rapidamente e introduce poco traffico nella rete. Ha il difetto che richiede una grossa capacità elaborativa da parte dei router, problema che però è sempre meno influente grazie ai progressi che si sono ottenuti in tale direzione.

3.1 Caratteristiche dei protocolli di routing link state

I passi principali del protocollo sono:

- Ogni router si occupa di scoprire i suoi vicini e di conoscerne l'identificatore.
- Ogni router costruisce un pacchetto, detto *Link State Packet* (LSP), in cui sono elencati i vicini ed il costo del collegamento con essi.
- Il LSP è inviato a tutti i router raggiungibili; ogni router colleziona l'ultimo LSP generato da ogni altro router.
- Sulla base dei LSP ricevuti, ogni router ricostruisce la topologia della rete

e calcola un cammino di costo minimo per ogni router.

3.1.1 Flooding

Il flooding è la funzione che permette ai LSP di raggiungere tutti i router in modo che tutti ricostruiscano il medesimo grafo. La complessità di questa funzione sta nel minimizzare il numero di ritrasmissioni di un pacchetto e nel fare in modo che queste abbiano termine quando tutti i router hanno ricevuto il pacchetto.

Per fare questo, quando un router riceve un LSP lo ritrasmette a tutti i vicini (tranne a quello da cui lo ha ricevuto) solo se non era già stato ricevuto e ritrasmesso prima; questo confronto è facile da verificare poiché ogni router mantiene un database dei LSP più recenti.

Un problema che si verifica è che un router non può sapere se il LSP appena ricevuto è anche l'ultimo ad essere stato generato. Una soluzione può essere quello di associare ad ogni LSP un *timestamp* dell'istante in cui è stato generato: un router che riceve un LSP, controlla che il timestamp non sia troppo lontano nel futuro, poiché potrebbe essere errato ed in seguito nessun LSP con timestamp minore ma successivo verrebbe accettato; oppure lontano nel passato, in tal caso si può dedurre che il router che lo ha generato non è attivo. La difficoltà di questa soluzione è di mantenere gli orologi di tutti i router sincronizzati. Un'altra soluzione prevede di associare ad ogni LSP un *Sequence Number*: il LSP con il sequence number più alto è l'ultimo ad essere stato generato. Essendo i sequence number finiti, una volta giunti al valore massimo riprendono da zero; per poter affermare che il sequence number b è più grande del sequence number a si deve verificare che:

$$|a - b| < n/2 \text{ e } a < b \text{ oppure } |a - b| > n/2 \text{ e } a > b$$

con n il valore più grande raggiungibile da un sequence number.

Per prevenire errori dovuti ad un reinizio del conteggio o altri per cui fallirebbe il confronto sul sequence number, viene associato ad ogni LSP un campo *Age* che viene decrementato con il passare del tempo durante il quale permane in un database di un router. Giunto a zero viene comunque considerato più vecchio di qualsiasi LSP anche con sequence number minore.

3.1.2 Algoritmo LinkState

L'algoritmo link state determina in un grafo pesato i cammini più brevi da un nodo a tutti gli altri. In una rete i *nodi* corrispondono ai router, i *collegamenti* tra i nodi ai link fisici, i *pesi* alle varie metriche assegnabili, come la larghezza di banda, il ritardo introdotto, l'affidabilità.

L'algoritmo usato è quello inventato da Dijkstra e conosciuto come *Shortest Path First* (SPF). Ha una complessità computazionale proporzionale al quadrato del numero dei nodi del grafo; in reti non altamente connesse può essere ridotta al numero di collegamenti per il logaritmo del numero dei nodi.

Per effettuare il calcolo l'algoritmo necessita di alcuni database:

- *PATHS*: rappresenta un grafo aciclico orientato di cammini minimi dal nodo *S* che ha effettuato il calcolo. Esso è memorizzato sotto forma di triple nella forma $\langle N, d(N), \{Adj(N)\} \rangle$ dove: *N* è l'identificatore di un nodo, *d(N)* è la distanza di *N* da *S*, $\{Adj(N)\}$ è l'insieme degli identificatori dei nodi direttamente connessi a *S* che *S* può usare per raggiungere *N* (è un insieme poiché possono esistere diversi cammini ad uguale costo).

Quando un nodo compare come *N* in una tripla di *PATHS* allora il cammino descritto per quel nodo è il migliore possibile.

- *TENT*: è un database di triple della forma $\langle N, d(N), \{Adj(N)\} \rangle$ dove *N*, *d(N)*, $\{Adj(N)\}$ hanno lo stesso significato di quelle definite per *PATHS*.

Contiene i nodi per i quali il cammino indicato *potrebbe* essere il migliore; se viene accertato che il cammino è davvero migliore, allora la tripla viene spostata in *PATH*.

Esecuzione

All'inizio dell'esecuzione dell'algoritmo, il nodo che lo esegue pone sé stesso alla radice dell'albero dei cammini minimi, in pratica mette la tripla $\langle RT_ID, 0, \emptyset \rangle$ in *PATH* (*RT_ID* è l'identificatore del router che sta effettuando il calcolo). *TENT* viene precaricato con la lista dei nodi adiacenti.

Bisogna notare che un nodo non è posto in *PATHS* a meno che non esista un cammino più breve ad esso. Quando un nodo *N* è posto in *PATHS*, viene esaminato il cammino ad ogni altro nodo *M* vicino di *N* a cui viene assegnato un costo dato da quello fino a *N* con aggiunto quello da *N* a *M*. Se in *PATH*

c'è già M , il nuovo percorso viene cancellato perché sicuramente più lungo. Se invece il nodo M è in $TENT$ e il nuovo cammino a M è più breve allora il vecchio percorso viene eliminato da $TENT$ e sostituito con il nuovo. Se il nuovo cammino ha invece la stessa lunghezza di quello che è in $TENT$ allora il campo $\{Adj(M)\}$ viene posto uguale all'unione dell' $\{Adj(M)\}$ già in $TENT$ e il nuovo $\{Adj(N)\}$. Se M non è in $TENT$ vi viene aggiunto.

In seguito l'algoritmo ricerca in $TENT$ la tripla $\langle N, x, \{Adj(N)\} \rangle$ con x minimo che viene posta in $PATHS$. Questo può essere fatto poiché tutti i cammini attraverso i nodi in $PATHS$ sono già stati considerati e quelli attraverso nodi in $TENT$ sono sicuramente più lunghi in quanto x è minimo.

Quando $TENT$ è vuoto, $PATHS$ è completo e l'algoritmo è terminato.

3.2 OSPF¹

OSPF (Open Shortest Path First) è un protocollo routing di intradomain basato sull'algoritmo link state ed utilizzato da IP.

3.2.1 Topologia

La topologia gestita da OSPF è costituita da un autonomuos system suddiviso in aree collegate da un backbone.

Autonomous system

Un *autonomous system* è costituito da tutti i router che eseguono lo stesso protocollo OSPF e da tutte le reti connesse a questi router.

Aree

OSPF permette di raggruppare insieme network contigue, ovvero LAN che tra coppie di esse hanno almeno un router in comune; tali gruppi sono detti *aree*. Per determinare a quale area appartiene un router, ad esso deve essere dato un particolare comando di configurazione che gli specifica l'identificatore dell'area di appartenenza.

¹[RFC 1583]

I router appartenenti ad un'area eseguono una copia dell'algoritmo indipendente da quella dei router appartenenti ad altre aree; la topologia di un'area è sconosciuta all'esterno dell'area stessa e allo stesso modo i router di un'area non conoscono i dettagli del resto dell'autonomous system. Questo isolamento permette una riduzione del traffico delle informazioni di routing e di limitare la portata di eventuali errori nel calcolo del routing all'interno dell'area stessa.

Un router che appartiene ad una sola area è detto *Internal Router* ed è in grado di determinare il percorso ad ogni altro host dell'area; le comunicazioni tra host di aree differenti avvengono invece attraverso il backbone. Un router connesso ad una o più aree e al backbone è detto *Area Border Router*: esso conosce la topologia di tutte le aree a cui è connesso e del backbone ed è in grado di effettuare il routing interarea.

Una area può anche essere configurata come *Stub Area*: essa è definita dal fatto che non vi entrano informazioni di routing esterne e per raggiungere host esterni viene utilizzato sempre un'area border router di default.

È possibile riassumere gli indirizzi di sottorete che escono da un'area con un unico indirizzo comune; se, ad esempio, un'area contiene tutte le sottoreti di una rete di classe B, allora è sufficiente che all'esterno venga conosciuto solo l'indirizzo di classe B.

Backbone

Il *backbone* dell'autonomous system è costituito dagli area border router e dai router connessi alle network che non appartengono a nessuna area; tutti questi router sono detti *backbone router*.

I router del backbone devono essere connessi; se esistono aree che ne spezzano la connettività, allora è possibile configurare un *virtual link* tra due backbone router che abbiano un'interfaccia su tale area. Un virtual link è considerato come una rete punto-a-punto unnumbered².

Il backbone ha le stesse proprietà di un'area (è detto anche backbone area) ed è responsabile di effettuare il routing interarea: i backbone router

²Le interfacce su reti unnumbered prendono per indirizzo lo stesso indirizzo di un'altra interfaccia dello stesso router.

ricostruiscono la topologia del backbone e calcolano i cammini minimi ad ogni router. Gli area border router includono nei loro pacchetti l'elenco degli indirizzi IP delle reti che appartengono alle loro area.

I router sono così classificati:

Internal Router: sono quelli che appartengono o solo ad un'area o solo al backbone ed eseguono una sola copia dell'algoritmo link state.

Area Border Router: sono i router connessi ad almeno un'area ed al backbone; eseguono un algoritmo di calcolo per ogni area a cui appartengono e uno per il backbone. Sul backbone si scambiano informazioni circa la raggiungibilità di tutti gli indirizzi IP di rete dell'autonomous system.

Backbone Router: sono i router che hanno almeno un'interfaccia sul backbone.

Autonomous System Boundary Router: sono i backbone router che eseguono anche protocolli di routing interdomain e attraverso i quali si raggiungono gli altri autonomous system.

Nella figura 3.1 è riportato un esempio di autonomous system e della denominazione dei router che vi appartengono.

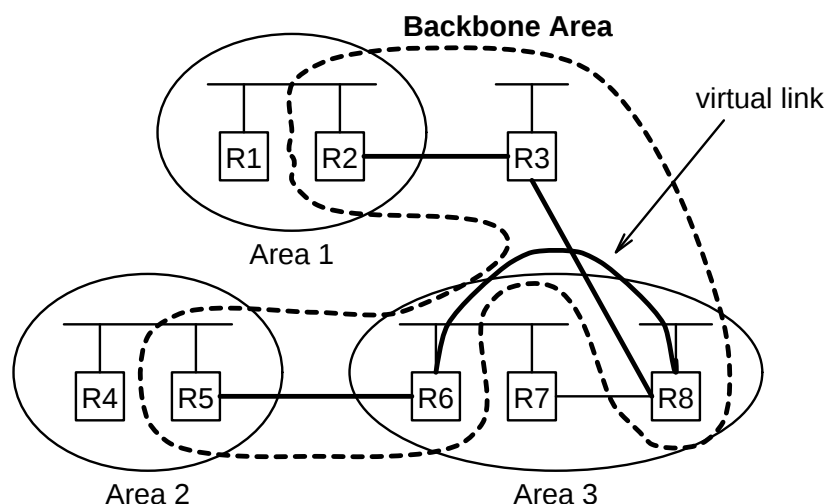


Figura 3.1. Esempio di topologia di un autonomous system: R2, R3, R5, R6 e R8 formano il backbone; R2, R5, R6 e R8 sono anche area border router.

3.2.2 Topological Database e Routing Table

Ogni router distribuisce le informazioni sullo stato dei suoi collegamenti inviando all'interno della propria area dei *Link State Advertisement* (LSA); inoltre mantiene un database che raccoglie tutti i LSA ricevuti e che descrive la topologia dell'area di appartenenza; questo database è detto *Topological Database*. Esso descrive un grafo orientato; i nodi di tale grafo sono costituiti dai router e dalle LAN con più router, le quali costituiscono uno pseudonodo (tale funzione è svolta da uno speciale router detto *Designated Router*); le connessioni del grafo sono i collegamenti punto a punto tra due router o i collegamenti tra un router di una LAN e il relativo pseudonodo (fig. 3.2). Ad ogni connessione è associato un costo configurabile.

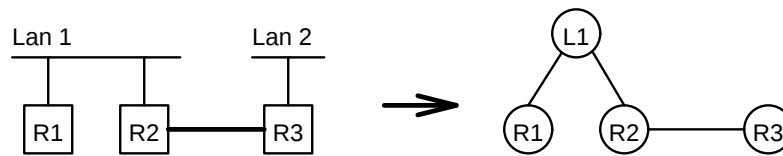


Figura 3.2. Esempio di una topologia e il grafo corrispondente.

Da questo grafo ogni router costruisce un albero di cammini minimi con sé stesso alla radice e come foglie tutti gli altri nodi; il grafo e l'algoritmo applicato devono essere identici in tutti i router di un'area per mantenere la consistenza dei cammini. Se esistono cammini con costo uguale, il traffico viene equamente distribuito su ognuno di essi. I cammini minimi vengono calcolati per ogni TOS.

L'albero di cammino minimo è riassunto nella *Routing Table* in cui per ogni destinazione è riportato solamente il router successivo da raggiungere. In figura 3.3 sono riportati i campi della routing table.

Type	Dest	Area	Path type	Cost	Next hop	Advertis. router
------	------	------	-----------	------	----------	------------------

Figura 3.3. Routing table: type è il tipo di destinazione (Network, Area border router, AS boundary router), path type indica se il percorso è interno o esterno all'area o all'autonomous system.

Ogni area border router genera un *Summary Link Advertisement* per ogni indirizzo di rete esterno all'area ma appartenente al resto dell'autonomous system e uno per ogni AS boundary router esterni all'area; questi sono ricevuti da tutti i router dell'area e descrivono il percorso per raggiungerli (questo non avviene se l'area è stub). Tali informazioni vengono anch'esse riassunte nella routing table.

Allo stesso modo, gli AS boundary router generano un AS external link advertisement per ogni destinazione esterna all'autonomous system che raggiunge ogni router dell'autonomous system (tranne quelli interni alle stub area); ad esse possono essere associate due metriche: il tipo 1 è equivalente alla metrica delle connessioni interne che quindi viene sommata alle metriche interne per calcolare i percorsi; il tipo 2 invece prevale sulle metriche interne, per cui se due router annunciano la connessione con una stessa rete esterna, tutti i router interni sceglieranno quello che annuncia il costo di tipo 2 minore, indipendentemente dalla distanza interna.

Ogni router ha quindi nella propria tabella di routing ogni indirizzo di rete e sottorete dell'autonomous system, oltre a tutte quelle conosciute dagli AS boundary router.

3.2.3 Esecuzione

La prima azione che compie ogni router è di usare il protocollo di Hello per scoprire quali altri router sono a lui vicini (cioè direttamente connessi); nelle reti broadcast (es. ethernet) questo protocollo è usato anche per eleggere il Designated Router (il pseudonodo). Poi il router stabilisce le *adiacenze* con i suoi vicini; queste sono necessarie poiché i pacchetti del protocollo di routing sono distribuiti solo sulle adiacenze e su di esse avviene il controllo della distribuzione dei pacchetti di routing. Ogni router informa gli altri dello stato delle sue adiacenze inviando periodicamente dei Link State Update (LSU).

Nelle LAN, il designated router è adiacente ad ogni router della LAN e genera un LSU con la lista di tutti questi router. Quando il designated router riceve un LSU lo ritrasmette in broadcast sulla LAN. Tutti gli altri router della LAN nei loro LSU riportano solo l'adiacenza con il designated router.

Un algoritmo di flooding provvede a far conoscere a tutti i router di un area tutti i LSU in modo che tutti i router ricostruiscano lo stesso topological database. Sulla base di questo database viene calcolato il cammino minimo per ogni destinazione.

3.2.4 Pacchetti OSPF

Tutti i pacchetti OSPF hanno un'intestazione comune di 24 ottetti; essa è riportata in figura 3.4. *Version* è la versione di OSPF (attualmente 2); *packet type* è un numero che corrisponde ai tipi descritti qui di seguito; *packet length* indica il numero di ottetti che compongono l'intero pacchetto; *router ID* è l'indirizzo IP del router che genera il pacchetto; *area ID* è il numero che identifica l'area OSPF a cui appartiene il router; *authentication type* indica il tipo di dato contenuto nel campo *authentication data*.

1	version
1	packet type
2	packet length
4	router ID
4	area ID
2	checksum
2	authentication type
8	authentication data

Figura 3.4. Intestazione comune a tutti i pacchetti OSPF; a fianco dei campi è indicata la dimensione in ottetti.

I tipi di pacchetti del protocollo di routing sono cinque:

1 Hello Packets: usati per stabilire e mantenere relazioni di vicinato e anche per eleggere il designated router. Sono spediti periodicamente ogni *HelloInterval* secondi su ogni interfaccia e sui virtual link ai soli vicini. Hanno una dimensione di 44 ottetti più 4 per il numero dei router dai quali si è ricevuto un valido hello.

In figura 3.5 sono riportati i campi del pacchetto: *Network Mask* è la maschera dell'indirizzo della rete a cui appartiene il router; *hello interval* è il numero di secondi tra un hello e l'altro; *option* comprende un bit *T* che indica se sono supportate le metriche multiple e un bit *E* che indica se l'area è stub; *DR priority* è un numero proporzionale alla possibilità di essere eletto designated router; *dead intervall* è il numero di secondi entro i quali se un hello non viene ricevuto da un vicino, questo viene dichiarato non attivo; *designated router* è l'indirizzo IP del designated router (se c'è); *backup DR* è l'indirizzo IP del backup designated router (se c'è); *neighbor* sono gli indirizzi IP dei vicini da cui si sono ricevuti hello.

24	common header
4	network mask
2	hello intervall
1	option: E, T
1	router priority
4	dead intervall
4	designated router
4	backup DR
4x	neighbor

Figura 3.5. Hello packet; il campo *neighbor* è ripetuto per ogni vicino.

2 Database Description: riassumono il contenuto del topological database e servono per inizializzare una adiacenza. La spedizione di questi pacchetti dipende dallo stato del vicino: se è in stato *exstart* i pacchetti sono vuoti e sono spediti ogni RxmtInterval secondi; se è in stato *exchange* allora i pacchetti sono completi. Quando due router hanno terminato di scambiarsi i database questi pacchetti non vengono più ritrasmessi. Hanno una dimensione di 32 ottetti più 12 per ogni record del topological database.

In figura 3.6 sono riportati i campi del pacchetto: *Flag* e *DD sequence number* contengono dati sulla sequenza di trasmissione dei pacchetti; *LSP header* sono le intestazioni di tutti i LSP presenti nel database.

24	common header
2	reserved
1	option: E, T
1	flag
4	DD sequence number
20x	LSP header

Figura 3.6. Database description packet; il campo *LSP header* è ripetuto per ogni LSP contenuto nel database.

3 Link State Request: dopo lo scambio dei Database Description packet uno dei router può accorgersi che parte del suo database non è più aggiornato; questi pacchetti servono per richiedere al vicino un aggiornamento di specifiche parti del database. Se un router non ottiene risposta rispedisce questi pacchetti ogni RxmtInterval secondi. Hanno una dimensione di 24 ottetti più 12 per ogni record di cui si richiede l'aggiornamento.

In figura 3.7 sono riportati i campi del pacchetto; i campi sono descritti nella sezione dei link state advertisement.

24	common header
4	link state type
4	link state ID
4	advertising router

Figura 3.7. Link state request; i campi *link state type*, *link state ID* e *advertising router* sono ripetuti per ogni LSP richiesto.

4 Link State Update: hanno il compito di distribuire i Link State Advertisement i quali sono contenuti al loro interno. Vengono trasmessi ogni *LSRefreshTimer* secondi. Hanno una dimensione di 28 ottetti più la dimensione di ogni LSA contenuto e vengono ritrasmessi da ogni router finché tutti quelli di un area non lo hanno ricevuto.

In figura 3.8 sono riportati i campi del pacchetto: *Numero LS advertisement contenuti* indica quanti LSA sono contenuti di seguito nel pacchetto.

24	common header
4	number of LS advertisement
Nx	link state advertisement

Figura 3.8. Link state update; il campo *link state advertisement* è ripetuto per ogni LSA, ognuno con una propria dimensione differente.

5 Link State Acknowledgments: usati per informare della ricezione dei Link State Update packet. Un LS Acknowledgements può contenere il riscontro di più LS Advertisement. Riportano al loro interno gli header dei LS Advertisement che riscontrano. Hanno una dimensione di 24 ottetti più la dimensione dell'header di ogni LS Advertisement contenuto e vengono ritrasmessi da ogni router finché tutti quelli di un'area non li hanno ricevuti.

In figura 3.9 sono riportati i campi del pacchetto.

24	common header
20x	link state adv. header

Figura 3.9. Link state acknowledgment; il campo *link state header* è ripetuto per ogni link state da riscontrare.

I Link State Advertisement sono di cinque tipi, ma hanno tutti un intestazione comune riportata in figura 3.10: *LS age* è una stima dei secondi trascorsi da quando il LSA è stato generato; *link state type* è un numero da uno a cinque che identifica il tipo di LSA: il numero corrisponde a quelli riportati a fianco della descrizione qui di seguito; *link state ID* è definito a seconda del tipo di LSA: se di tipo 1, contiene l'identificatore del router che ha generato il LSA; se di tipo 2, contiene l'indirizzo IP del designated router; se di tipo 3, contiene l'indirizzo IP della destinazione; se di tipo 4, l'identificatore dell'AS boundary router; se di tipo 5, l'indirizzo IP della de-

stinazione. *Advertising router* è l'identificatore del router che genera il LSA; *length* è il numero di ottetti di cui è composto il LSA.

2	LS age
1	option: E, T
1	link state type
4	link state ID
4	advertising router
4	LS sequence number
2	LS checksum
2	length

Figura 3.10. Intestazione link state advertisement.

1 Router Links Advertisement: originati da ogni router di un'area, descrivono lo stato e il costo di un link ad un altro router dell'area. Rimangono all'interno dell'area in cui sono stati generati. Hanno una dimensione di 24 ottetti, più, per ogni interfaccia, 12 più 4 per ogni TOS.

In figura 3.11 sono riportati i campi del pacchetto: il *flag E* indica se il router che ha generato il LSP è un AS boundary router, il *flag B* indica se il router è un area border router; *number of link* indica il numero di link che il router ha nell'area; *link type* è 1 se link è punto a punto, è 2 se il link è una LAN con più di un router (*transit LAN*), è 3 se è LAN con un solo router (*stub LAN*), è 4 se è un virtual link; *link ID* è l'identificatore del router vicino se il link è di tipo 1 o di tipo 4, è l'indirizzo IP del designated router se il link è di tipo 2, è l'indirizzo di sottorete se il link è di tipo 3; *link data* contiene la maschera di sottorete se il link è di tipo 3, altrimenti contiene l'indirizzo IP sul collegamento del router che ha generato il LSP; *number of TOS* indica il numero di TOS attivi.

2 Network Links Advertisement: originati dai designated router, contengono la lista dei router connessi ad una LAN. Rimangono all'interno dell'area dove sono stati generati. Hanno una dimensione di 24 ottetti più 4 per ogni router sulla LAN.

20	link state header
1	flag: E, B
1	reserved
2	number of link
4	link ID
4	link data
1	type
1	number of TOS
2	TOS 0 metric
1	TOS
1	reserved
2	metric

Figura 3.11. Router link advertisement; gli ultimi tre campi sono ripetuti per ogni TOS attivo; i campi da *link ID* fino all'ultimo sono ripetuti per ogni link riportato.

In figura 3.12 sono riportati i campi del pacchetto: *Attached router* è l'identificatore di ogni router connesso alla LAN.

20	LS header
4	network mask
4x	attached router

Figura 3.12. Network link advertisement; il campo *attached router* è ripetuto per ogni router connesso alla LAN.

3,4 Summary Link Advertisement: originati dagli area border router e inviati all'interno della loro area, descrivono i percorsi per destinazioni fuori dall'area ma dentro all'autonomous system; ne vengono generati uno per ogni destinazione. Sono di tipo 3 se la destinazione è una network IP, sono di tipo 4 se la destinazione è un AS boundary router. Hanno una dimensione di 24 ottetti più 4 per ogni TOS.

In figura 3.13 sono riportati i campi del pacchetto.

20	LS header
4	network mask
1	TOS
3	metric

Figura 3.13. Summary link advertisement; i campi *TOS* e *metric* sono ripetuti per ogni TOS attivo.

5 Autonomous System External Link Advertisement: originati dagli AS boundary router e inviati all'interno dell'autonomous system, ne viene generato uno per ogni destinazione esterna all'autonomous system. Hanno una dimensione di 24 ottetti più 12 per ogni TOS.

In figura 3.14 sono riportati i campi del pacchetto: *Forwarding address* contiene l'indirizzo di un altro router se questo ha un percorso migliore per la stessa destinazione; *external route tag* non è utilizzato ma è a disposizione dei protocolli interdomain.

20	LS header
4	network mask
1	TOS
3	metric
4	forwarding address
4	external route tag

Figura 3.14. Autonomous system external link advertisement; gli ultimi quattro campi sono ripetuto per ogni TOS attivo.

3.3 Integrated IS-IS³

Integrated IS-IS (Intermediate System to Intermediate System) è un protocollo di routing interdomain basato sull'algoritmo link state ed è utilizzato sia da IP che da CLNP; il prefisso integrated deriva proprio dal fatto che può essere usato come unico protocollo di routing in un internet gestita tramite questi due protocolli. È comunque possibile impiegarlo anche come protocollo solo per IP; questo significa che, in internet IP che non usino il CLNP, è comunque possibile utilizzare integrated IS-IS assegnando ad ogni router un NSAP di identificazione; nei pacchetti mancheranno solo le informazioni proprie del CLNP.

Le funzioni di IS-IS sono divise in due gruppi:

Subnetwork Independent Functions: realizzano la trasmissione full-duplex dei NPDU⁴ tra ogni coppia di sistemi vicini. In particolare la *Routeing function* determina i percorsi che devono seguire i NPDU e la *Congestion Control function* controlla l'utilizzo delle risorse in ogni intermediate system (IS).

Subnetwork Dependent Functions: determinano gli indirizzi dei vicini, inizializzano i collegamenti, frammentano i pacchetti.

³[ISO 10589] [RFC 1195]

⁴NPDU: Network Packet Data Unit.

3.3.1 Topologia

La topologia di Integrated IS-IS consiste di un routing domain suddiviso in aree a cui appartengono IS di livello 1 e 2.

Routing domain

Un *routing domain* è costituito da tutti gli IS che eseguono lo stesso protocollo e da tutte le reti connesse a questi IS.

Livello 1: Aree

Un routing domain è un insieme di aree; un'area è insieme di network in cui un host comunica con un altro della stessa area attraverso un IS di livello 1; lo stesso host comunica con un host di un'altra area attraverso il più vicino IS di livello 2. All'interno di un'area è sconosciuta la topologia del resto del routing domain.

Ogni host è identificato da un NSAP e un'area è costituita dagli host in cui il campo area address dell'NSAP è uguale; gli indirizzi IP delle LAN che appartengono ad un'area non è necessario che abbiano una particolare correlazione l'un con l'altro.

In ogni area gli IS di livello 1 sono in grado di determinare il percorso ad ogni altro end-system dell'area; questo percorso viene determinato dal Level 1 Forwarding Database ricostruito per mezzo dei Level 1 Link State PDU⁵ che gli IS di livello 1 di un'area si scambiano. È necessario che all'interno di un'area vi sia almeno un IS di livello 2 per garantire la connessione con le altre aree.

Ad ogni area possono essere associati più indirizzi che insieme costituiscono il *manual area addresses*. Esso viene distribuito nei Link State PDU (LSP) degli IS di livello 1. Due IS di livello 1 diventano adiacenti se i loro manual area addresses hanno almeno un indirizzo di area in comune. Il manual area address è utile nel caso si prevedano modifiche alla topologia: ad esempio per unire due aree in una unica oppure per partizionarne una in due.

⁵PDU: Packet Data Unit.

I LSP degli IS di livello 1 contengono triple del tipo <indirizzo_IP, subnet_mask, metrica>, una per ogni indirizzo IP direttamente raggiungibile dall'IS che li genera.

Anche per integrated IS-IS come per OSPF è possibile riassumere gli indirizzi di sottorete che escono da un'area con un unico indirizzo comune.

Livello 2

La connessione tra le aree è realizzata dagli IS di livello 2; essi determinano i percorsi dal Level 2 Forwarding Database che viene ricostruito sulla base dei Level 2 Link State PDU che essi stessi si scambiano. Per effettuare questo scambio è necessario che ogni IS di livello 2 sia connesso fisicamente ad un altro IS di livello 2. Tutti insieme essi costituiscono un sottodominio detto *Level 2 Subdomain*.

Attraverso questo sottodominio è possibile effettuare la riparazione delle aree partizionate: se un collegamento fra IS di livello 1 si interrompe provocando la partizione dell'area a cui appartengono (cioè viene a mancare l'unico percorso tra host di una stessa area), allora questo fatto viene rilevato dal (dagli) IS di livello 2 dell'area, il quale provvede a ristabilire la connessione tra le partizioni creando un collegamento virtuale tra IS del Level 2 Subdomain. Gli IS di livello 2 elencano nei loro LSP tutti gli indirizzi raggiungibili all'interno della loro area.

Un IS può essere o solo di livello 1, o solo di livello 2, oppure appartenere ad entrambi i livelli. Nella figura 3.15 è riportato un esempio di routing domain.

3.3.2 Link State Database e Forwarding Database

Ogni IS di livello 1 distribuisce le informazioni sullo stato dei propri collegamenti tramite i Level 1 Link State PDU: in essi sono riportati il campo ID del suo NSAP e di quelli degli IS a lui connessi, oltre a tutti gli indirizzi IP con relative maschere delle LAN a cui esso appartiene. Ogni IS raccoglie tutti i LSP ricevuti nei *Link State Database*, uno per ogni metrica utilizzata.

Questi database descrivono un grafo in cui i nodi sono gli IS di livello 1 e le connessioni i link tra coppie di essi; a questo grafo ogni IS applica

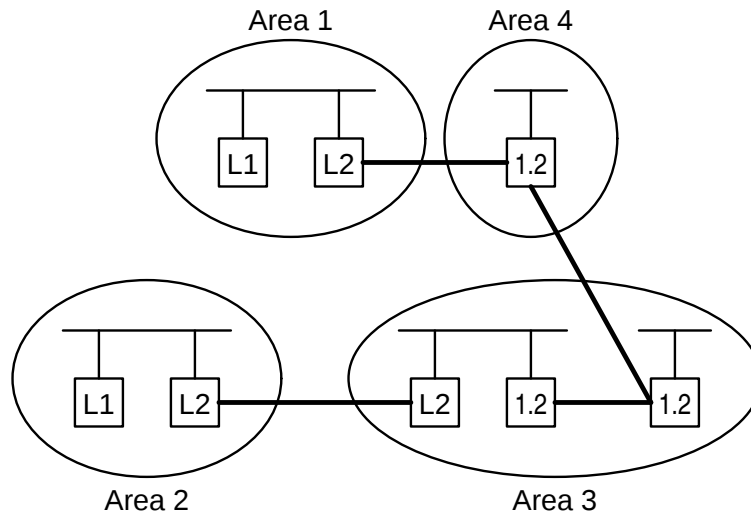


Figura 3.15. Schema di un routing domain; all'interno degli IS è indicato il loro livello.

un algoritmo di cammino minimo per determinare i percorsi da lui ad ogni altra destinazione dell'area. Ogni IS applica una copia indipendente dello stesso algoritmo al proprio link state database; per garantire la consistenza dei percorsi (per evitare loop) e il determinismo del loro calcolo, è necessario che i link state database coincidano negli IS di una stessa area.

I cammini sono riassunti nel *Level 1 Forwarding Database* che per ogni destinazione indica l'IS di livello 1 da raggiungere per primo; le destinazioni sono gli indirizzi IP delle LAN connesse ad ogni nodo (IS) del grafo. Nel forwarding database è contenuta anche l'informazione per raggiungere l'IS di livello 2 più vicino. I Level 1 Forwarding Database in ogni IS sono uno per ogni metrica.

Identico discorso per gli IS del level 2 subdomain: ognuno possiede i Link State Database che raccolgono le informazioni portate dai Level 2 Link State PDU e da essi vengono determinati i *Level 2 Forwarding Database*. Una differenza fondamentale è che negli LSP sono indicati gli area address dell'NSAP degli IS e non i campi id.

3.3.3 Routeing function

La routeing function è costituita da quattro processi:

Decision: calcola i percorsi per ogni destinazione del dominio. È eseguita separatamente a livello 1 e 2 e per ogni metrica supportata. Usa il Link State Database per calcolare il cammino più breve per ogni IS nell'area; esso consiste di informazioni a riguardo dei più recenti LSP ricevuti dagli IS dell'area ed è costruito dal processo di Update. I risultati del calcolo formano il Forwarding Database.

I Link State Database e i Forwarding Database sono sia di livello 1 che di livello 2; inoltre esiste un Forwarding Database per ogni metrica.

Nelle broadcast network (LAN) viene eletto un *Designated IS* in modo da costituire uno pseudonodo: il designated IS è adiacente ad ogni IS della LAN, mentre gli IS della LAN sono adiacenti solo al designated IS. Il designated IS genera allora alcuni LSP con elencati i collegamenti a tutti gli IS della LAN ai quali è assegnato costo zero, mentre gli IS generano LSP dove indicano solo l'adiacenza con il designated IS.

Quando un area viene spezzata in più parti a causa dell'interruzione di un collegamento, gli IS di livello 2 di quell'area sono in grado di rimediare al guasto creando un collegamento virtuale di livello 1 nel level 2 subdomain. Gli IS di livello 2 che intervengono nel virtual link si dichiarano *partition designated level 2 IS* e sfruttano il fatto che nella loro area possono funzionare come degli IS di livello 1. Essi annunciano nei loro level 1 LSP un'adiacenza con ogni partition designated level 2 IS dell'area, assegnandole un indirizzo virtuale.

Update: costruisce, riceve e propaga i LSP; questi contengono informazioni sull'identità e sulla metrica dei sistemi adiacenti all'IS che ha generato il LSP. I LSP sono inviati ad intervalli regolari di *LSPGenerationInterval* ed ogni volta che avviene un cambiamento nella topologia. I LSP generati da un IS di livello 1 sono propagati a tutti gli IS all'interno di un'area, mentre i LSP generati da un IS di livello 2 sono trasmessi solo agli IS di livello 2 del routing domain.

Forwarding: si occupa di far arrivare i pacchetti IP e CLNP a destinazione basandosi sul forwarding database.

Receive: si occupa della ricezione dei pacchetti e decide come devono essere trattati; inoltre gestisce le informazioni derivate dal protocollo ES-IS .

3.3.4 Pacchetti integrated IS-IS

I pacchetti usati da Integrated IS-IS hanno tutti la stessa intestazione riportata in figura 3.16 e possono avere dei campi variabili la cui struttura è riportata in figura 3.17.

1	protocoll identifier
1	header length
1	version
1	ID length
1	packet type
1	version
1	riserved
1	maximum area addresses

Figura 3.16. Intestazione pacchetti integrated IS-IS.

Protocol identifier è una costante uguale a 132; i due campi *version* indicano entrambi la versione del protocollo; *ID length* indica il numero di ottetti del campo ID del NSAP del router.

Nei campi variabili, *Code* è un numero che ne specifica il contenuto, *value* contiene dei particolari valori e può essere ripetuto più volte.

1	code
1	length
var	value

Figura 3.17. Parte variabile dei pacchetti integrated IS-IS.

Di seguito sono descritti i pacchetti usati da Integrated IS-IS; per determinarne la dimensione è stato considerato un NSAP address di venti ottetti di cui tredici per l'area, sei per l'ID, uno per il selector.

Level 1 LAN IS to IS hello PDU: serve agli IS di livello 1 per scoprire l'identità degli altri IS di livello 1 connessi alla stessa rete broadcast. Sono

trasmessi ogni *HelloTimer* secondi su ogni interfaccia. Hanno una dimensione di 48 ottetti più 1 per ogni protocollo supportato, 4 per ogni interfaccia IP dove sono inviati e 6 per ogni interfaccia CLNP (in IS-IS la dimensione fissa è di 46 ottetti e non vi è l'indicazione del protocollo usato).

In figura 3.18 sono riportati i campi del pacchetto: *Circuit type* indica il tipo di collegamento su cui è inviato il pacchetto: 1 identifica un circuito solo di livello 1, 2 solo di livello 2, 3 sia di livello 1 che 2; *source ID* è il campo ID del NSAP; *holding time* è il tempo di vita degli hello; una volta trascorso, il router viene dichiarato non attivo. *Priority* è un numero che indica la priorità a diventare designated router; *LAN ID* è l'identificatore assegnato alla LAN dal designated router. Nelle parti variabili, con codice 1 sono indicati tutti gli indirizzi di area, con codice 129 i protocolli supportati, con codice 132 l'indirizzo IP dell'interfaccia sul quale il pacchetto è trasmesso, con codice 133 i valori di autenticazione.

8	common fixed header
1	circuit type
length ID	source ID
2	holding time
2	packet length
1	priority
lengthID + 1	LAN ID

Figura 3.18. Level 1 e 2 LAN IS to IS hello PDU, parte fissa.

Level 2 LAN IS to IS hello PDU: serve agli IS di livello 2 per scoprire l'identità degli altri IS di livello 2 di un circuito broadcast. Trasmessi ogni *HelloTimer* secondi su ogni interfaccia. Hanno una dimensione di 48 ottetti più 1 per ogni protocollo supportato, 4 per ogni interfaccia IP dove sono inviati e 6 per ogni interfaccia CLNP (in IS-IS la dimensione fissa è di 46 ottetti e non vi è l'indicazione del protocollo usato). I campi dei pacchetti sono riportati in figura 3.18 (sono gli stessi di quelli dei Level 1 IS to IS Hello PDU).

Point-to-point IS to IS hello PDU: trasmesso dagli IS sulle reti punto-a-punto; serve per determinare se il vicino è un IS di livello 1 o 2. Sono ritrasmessi ogni HelloTimer secondi. Hanno una dimensione di 41 ottetti più 1 per ogni protocollo supportato, 4 per ogni interfaccia IP dove sono inviati e 6 per ogni interfaccia CLNP.

I campi del pacchetto sono riportati in figura 3.19: *Circuit type* indica il tipo di collegamento su cui è inviato il pacchetto: 1 indica un circuito solo di livello 1, 2 solo di livello 2, 3 sia di livello 1 che 2; *source ID* è il campo ID del NSAP; *holding time* è il tempo di vita degli hello; una volta trascorso, il router viene dichiarato non attivo. *Local circuit ID* è un identificatore per l'interfaccia dove è inviato l'hello. Nella parte variabile i campi sono gli stessi del Level 1 LAN IS to IS hello.

8	common fixed header
1	circuit type
length ID	source ID
2	holding time
2	packet length
1	local circuit ID

Figura 3.19. Point-to-point IS to IS hello PDU, parte fissa.

Level 1 link state PDU: sono generati dagli IS di livello 1 e 2 e trasmessi all'interno dell'area di appartenenza agli altri IS. Ogni IS indica nei propri pacchetti lo stato delle sue adiacenze con gli altri IS direttamente connessi. Sono trasmessi periodicamente con un intervallo di tempo compreso fra *MinimumLSPGenerationInterval* e *MaximumLSPGenerationInterval*. Hanno una dimensione di 52 ottetti più 1 per ogni protocollo supportato, 4 per ogni indirizzo IP abilitato sull'IS, 12 per ogni subnet raggiungibile direttamente dal IS e 11 per ogni vicino di livello 1 (in IS-IS la dimensione fissa è di 50 ottetti e non vi è l'indicazione del protocollo usato).

In figura 3.20 sono riportati i campi del pacchetto: *LSP ID* è un numero che identifica univocamente il pacchetto; insieme a *remaining lifetime* e a

sequence number servono per il flooding affidabile; *P* è un flag che indica se l'IS è un partition designated level 2 IS; *Att* indica le metriche supportate; *OL* indica se l'IS ha esaurito lo spazio del suo database di LSP; *IS type* indica se il router è di livello 1 o 2. Nella parte variabile con codice 1 sono indicati tutti gli indirizzi di area, con codice 2 sono elencati i campi ID degli NSAP degli IS vicini, con codice 3 sono elencati i campi ID degli NSAP degli ES vicini (solo nei LSP di livello 1), con codice 5 sono elencati i prefissi degli indirizzi degli IS di livello 2 vicini (solo nei LSP di livello 2), con codice 129 i protocolli supportati, con codice 132 l'indirizzo IP dell'interfaccia sulla quale il pacchetto è trasmesso, con codice 128 gli indirizzi IP (con relative maschere di sottorete) delle destinazioni raggiungibili all'interno del routing domain e connesse all'IS.

8	common fixed headers
2	packet length
2	remaining lifetime
length ID + 2	LSP ID
4	sequence number
2	checksum
1	P, Att, OL, IS type

Figura 3.20. Level 1 e 2 link state PDU, parte fissa.

Level 2 link state PDU: generati dagli IS di livello 2 e propagati agli IS di livello 2 del routing domain. Ogni IS indica nei suoi pacchetti lo stato delle sue adicenze con altri IS di livello 2. Trasmessi periodicamente con un intervallo di tempo compreso fra `MinimumLSPGenerationInterval` e `MaximumLSPGenerationInterval`. Hanno una dimensione di 54 ottetti più 1 per ogni protocollo supportato, 4 per ogni interfaccia IP abilitata sull'IS, 12 per ogni indirizzo IP interno al routing domain raggiungibile direttamente dall'IS, 12 per gli indirizzi IP esterni al routing domain raggiungibili direttamente e 18 per ogni vicino di livello 2 (in IS-IS la dimensione fissa è di 50 ottetti e non vi è l'indicazione del protocollo usato).

In figura 3.20 sono riportati i campi della parte fissa del pacchetto; essi sono uguali a quelli dei level 1 link state PDU.

Level 1 complete sequence number PDU: sono generati dai designated IS di livello 1 delle LAN di un'area e sono indirizzati agli IS della LAN stessa. Contengono una lista dei LSP nel database dell'IS; in questo modo un'altro IS può confrontare questa lista con il suo database e scoprire eventuali differenze. Trasmessi periodicamente ogni *CompleteSNPIInterval* secondi. Hanno una dimensione di 33 ottetti più 16 per ogni LSP ricevuti.

In figura 3.21 sono riportati i campi del pacchetto: *Start LSP ID* e *End LSP ID* servono per la frammentazione dell'elenco dei LSP che può essere molto lungo e sono presenti solo nei complete sequence number PDU. Nella parte variabile, con codice 9 sono elencati i LSP ID dei pacchetti nel link state database, con codice 133 le informazioni di autenticazione.

8	common fixed header
2	packet length
length ID + 1	source ID
length ID + 2	start LSP ID
length ID + 2	end LSP ID

Figura 3.21. Level 1 e 2 complete e partial sequence number PDU, parte fissa.

Level 2 complete sequence number PDU: generati dal designated IS di livello 2 di una LAN e indirizzati agli IS di livello 2 della LAN stessa, contengono una lista dei LSP nel database del designated IS. Sono inviati periodicamente ogni *CompleteSNPIInterval* secondi. Hanno una dimensione di 33 ottetti più 16 per ogni LSP ricevuto.

In figura 3.21 sono riportati i campi della parte fissa del pacchetto.

Level 1 partial sequence number PDU: sono generati dagli IS di livello 1 e hanno due scopi principali: su collegamenti punto a punto sono usati come esplicito acknowledge per la ricezione di un LSP; su reti broadcast servono per richiedere al designated IS la ritrasmissione di LSP che sono stati trovati nel CSNP appena ricevuto, ma non nel database locale. Hanno una dimensione

di 17 ottetti più 16 per ogni LSP ricevuto.

In figura 3.21 sono riportati i campi del pacchetto.

Level 2 partial sequence number PDU: sono generati da IS di livello 2 e hanno gli stessi scopi del level 1 PSN PDU ma a livello 2. Hanno una dimensione di 17 ottetti più 16 per ogni LSP ricevuto.

In figura 3.21 sono riportati i campi del pacchetto.

3.4 Ship In the Night e Integrated routing

Ship In the Night (abbreviato in S.I.N.) è un metodo che implementa il routing multiprotocollo (IP, CLNP) in cui i protocolli di routing sono eseguiti in parallelo, come se ci fosse un router dedicato per ognuno di essi.

Integrated routing invece prevede che venga eseguito un singolo protocollo di routing; in pratica l'algoritmo che calcola i percorsi ad ogni altro router viene applicato una sola volta, mentre ogni protocollo aggiunge nei LSP le informazioni necessarie a far conoscere gli indirizzi raggiungibili.

Il vantaggio del routing Integrated è proprio quello che, eseguendo un solo algoritmo di calcolo, si diminuisce il carico sulla CPU dei router; inoltre è necessaria meno memoria fisica: con topologie molto complesse infatti potrebbe risultare impossibile per macchine non particolarmente potenti eseguire due protocolli link state contemporaneamente. In più l'occupazione della banda da parte dei pacchetti di controllo è molto minore dato che sono eliminate le informazioni ridondanti. Il risparmio di risorse fisiche e di banda trasmissiva si risolvono in un risparmio economico.

Un'altro vantaggio del routing integrated è che in realtà in S.I.N. i due protocolli sono sì indipendenti ma condividono le stesse risorse, per cui l'instabilità dell'uno si può ripercuotere gravemente sull'altro. Per contro non è possibile avere topologie differenti per i due protocolli; ad esempio, non si può definire un routing domain IP diverso da quello CLNP e neppure escludere determinati nodi in un protocollo e lasciarli attivi in un altro.

3.5 Confronto tra OSPF e Integrated IS-IS

In questa sezione è stato eseguito un confronto tra le implementazioni dei due protocolli.

3.5.1 Terminologia

I due protocolli usano una differente terminologia che però in certi termini coincide; nella tabella 3.1 sono riportati in ogni riga le corrispondenze fra alcuni termini dei due protocolli.

OSPF	integrated IS-IS
Autonomous System	Routing Domain
Topological Database	Link State Database
Routing Table	Forwarding Database
Internal router non di backbone	IS di livello 1
Internal router di backbone	IS di livello 2

Tabella 3.1. Termini corrispondenti tra OSPF e IS-IS.

3.5.2 Topologia

I due protocolli hanno due modi differenti di configurare una topologia fisica esistente, anche se entrambi la gestiscono a due livelli.

Livello 1

Al livello più basso, i router delle aree OSPF conoscono tutti gli indirizzi dell'autonomous system e, se nell'area ci sono più di un area border router, sanno a quale rivolgersi per raggiungere più velocemente una rete di un'altra area; in integrated IS-IS invece un router di livello 1 non conosce nessuno degli indirizzi esterni alla sua area e se deve trasmettere un pacchetto ad un host esterno, invia il pacchetto al router di livello 2 più vicino a sé stesso.

Ognuno dei due funzionamenti ha aspetti positivi e negativi: OSPF è in grado di scegliere i percorsi interarea ottimali mentre in integrated IS-IS un host passa sempre per lo stesso router per qualsiasi destinazione esterna, rischiando di percorrere strade più lunghe di quelle ottimali; per contro questo implica un maggior traffico di pacchetti di informazioni introdotto da OSPF all'interno di un'area ed una maggiore dimensione delle tabelle di routing dei router interni. Da notare che in OSPF è possibile configurare un'area come *stub* (cfr. sezione 3.2.1) il cui funzionamento, per quanto riguarda la conoscenza di destinazioni esterne all'area, è simile a quello delle aree integrated IS-IS.

Livello 2

Al livello più alto, i backbone router OSPF e gli IS di livello 2 integrated IS-IS conoscono entrambi tutti gli indirizzi dell'autonomous system, per cui il traffico introdotto e la dimensione delle tabelle di routing è più o meno lo stesso. Una differenza sostanziale sta nel modo in cui avvengono le comunicazioni tra il livello più basso e quello più alto: in OSPF deve esistere un router che appartenga sia ad un'area che al backbone e che quindi deve eseguire due processi di calcolo dei cammini minimi e mantenere due tabelle di routing. In integrated IS-IS invece i due livelli sono nettamente separati ed è sufficiente che ad ogni area appartenga un router anche solo di livello 2; un router di livello 1 e 2 esegue anch'esso due algoritmi per il calcolo dei cammini minimi e mantiene due tabelle di routing, ma non c'è nessuno scambio di informazione diretto tra i due processi: è come se i due processi fossero eseguiti da due router differenti.

In OSPF, un'area border router può anche appartenere a più di un'area e per ognuna di esse deve eseguire un processo di calcolo differente; in integrated IS-IS invece un IS sia di livello 1 che 2 può appartenere ad una sola area a causa del suo NSAP. In OSPF un backbone router può non appartenere a nessuna area mentre invece in integrated IS-IS un IS di livello 2 appartiene sempre ad un'area.

Queste differenze hanno importanti risvolti: in OSPF devono esistere router che eseguono comunque due processi di routing, il che implica una

maggior occupazione di memoria ed un maggior impiego della CPU e quindi la necessità di router comunque più potenti di quelli di IS-IS, in cui non sono richieste macchine di interfacciamento.

Inoltre, il fatto che in OSPF un router può appartenere a più di un'area oltre che al backbone, può avere come conseguenza che un router deve eseguire troppi processi di calcolo con il rischio di esaurire le sue risorse fisiche (CPU e memoria) causandone il crollo delle prestazioni. Bisogna dunque prestare attenzione quando in un autonomous system già definito vengono connesse nuove aree: se si utilizza come area border router uno che lo era già lo si costringe ad eseguire un altro algoritmo. Meglio prevedere di connettere nuove aree solamente su backbone router interni tramite dei nuovi area border router (vedere esempio di figura 3.22).

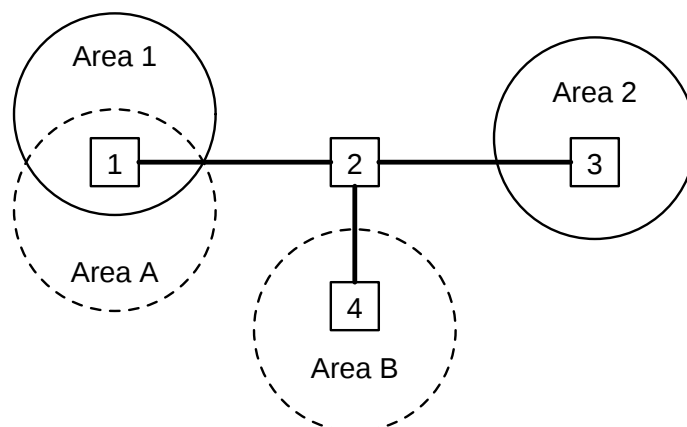


Figura 3.22. OSPF. Le aree A e B sono state connesse in seguito all'autonomous system originale formato dalle aree 1 e 2 e dai router 1,2,3 che formano il backbone. L'area A appesantisce il compito del router 1 costringendolo ad eseguire un algoritmo in più mentre l'area B non affatica il router 2 avendo un proprio area border router (4).

In integrated IS-IS non esiste il rischio che un IS esegua troppi algoritmi poiché esso può appartenere ad una sola area: al massimo ne eseguirà due se è sia di livello 1 che 2. Quando però una nuova area viene connessa ad un routing domain esistente bisogna prevedere di configurare all'interno di essa almeno un IS di livello 2 (vedere esempio di figura 3.23).

Un'altra conseguenza, più di termini che di fatti, è che in OSPF i confini delle aree sono sui router, nel senso che un link appartiene sempre interamente

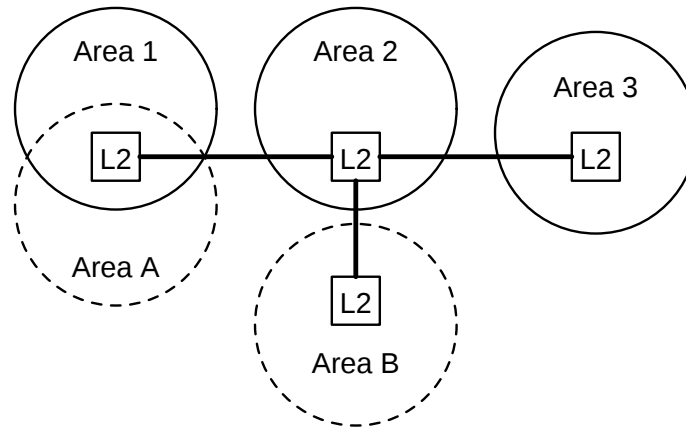


Figura 3.23. Integrated IS-IS. Le aree A e B sono state connesse in seguito al routing domain originale formato dalle aree 1,2,3. L'area A condivide l'IS di livello 2 dell'area 1 e l'area B ha un proprio IS di livello 2 connesso direttamente ad uno qualsiasi degli IS di livello 2 preesistenti.

ad un'area, mentre è il router che appartiene a diverse aree a seconda di come sono configurate le sue interfacce (figura 3.24).

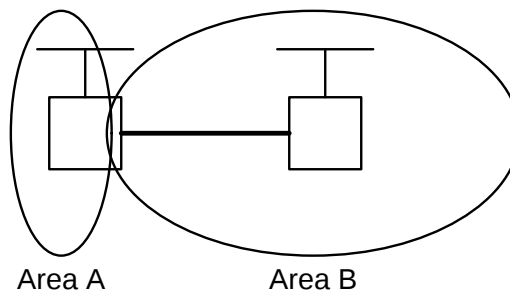


Figura 3.24. OSPF: i confini delle aree sono sui router.

In Integrated IS-IS i confini delle aree sono sui collegamenti, nel senso che un IS può appartenere solo ad un'area essendo unico il suo NSAP (figura 3.25).

Virtual Link

In OSPF non è necessario che i backbone router siano fisicamente connessi l'uno all'altro ma è possibile configurare dei virtual link tra coppie di backbone router che appartengono alla stessa area. Per creare un virtual

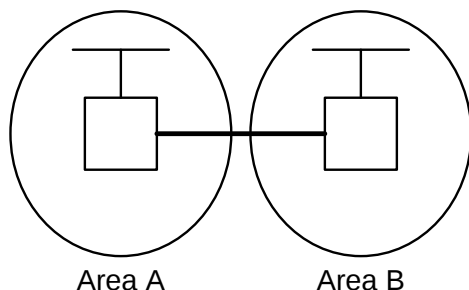


Figura 3.25. Integrated IS-IS: i confini delle aree sono sui link.

link è necessario configurarli entrambi informando l'uno dell'identificatore dell'altro. Questo provoca la dipendenza della configurazione di un router da quella dell'altro e viceversa, per cui bisogna prestare attenzione ad eventuali modifiche poiché i cambiamenti effettuati su un router si ripercuotono sull'altro.

In integrated IS-IS, invece, gli IS di livello 2 devono essere fisicamente connessi l'uno all'altro e non è possibile definire dei link virtuali tra coppie di essi. Questo implica una maggiore semplicità nella configurazione di integrated IS-IS, ma nello stesso tempo richiede un numero maggiore di IS di livello 2 all'interno di un'area rispetto ad OSPF; nell'esempio di figura 3.26 sono presenti due router connessi con altre aree ma non tra di loro. In OSPF è sufficiente configurare un virtual link tra i due; in integrated IS-IS è necessario creare un percorso di IS di livello 2. Questo significa che in integrated IS-IS devono essere presenti più IS che effettuano anche il routing di livello 2 con conseguente spreco di risorse.

Configurazione di router su LAN multicast

In OSPF la configurazione di un router prevede di assegnare ogni interfaccia ad un'area, per cui se a tutte le interfacce è assegnato lo stesso numero il router è interno, altrimenti è un area border router. Quando si assegna ad una LAN un numero di area, questo numero deve essere lo stesso su tutti i router della LAN multicast, altrimenti la LAN verrà annunciata appartenere a due aree diverse e non verrà utilizzata nelle comunicazioni interarea (figura 3.27). Il fatto che una LAN appartenga a due aree non è di per sé negativo,

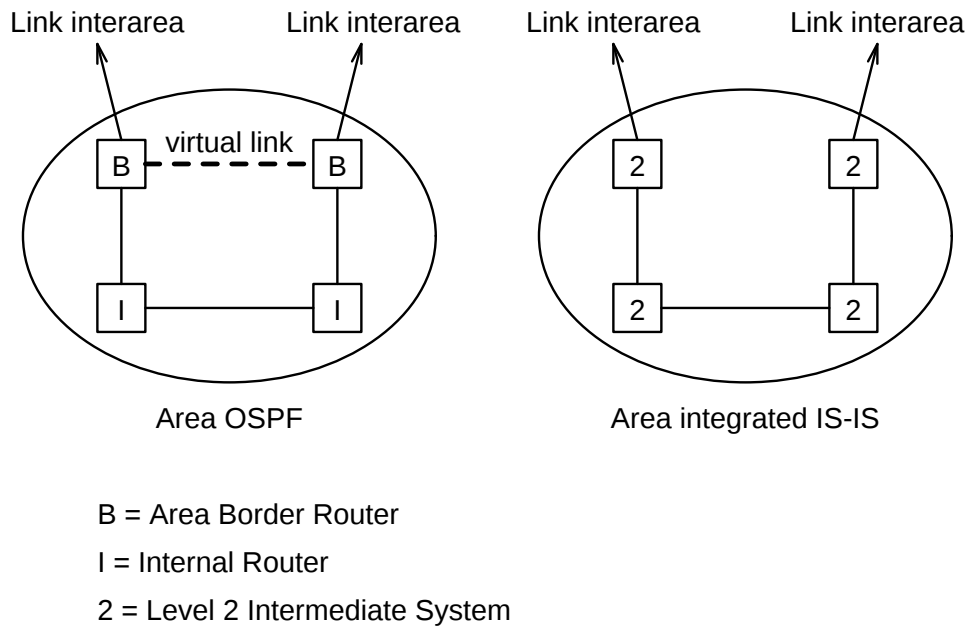


Figura 3.26. Esempio delle connessioni intrarea tra border router OSPF e integrated IS-IS.

implica però un maggiore traffico di routing e l'impossibilità di riassumere più sottoreti con un indirizzo unico.

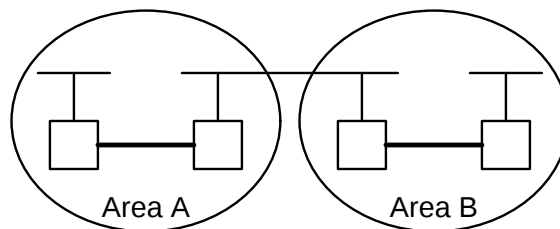


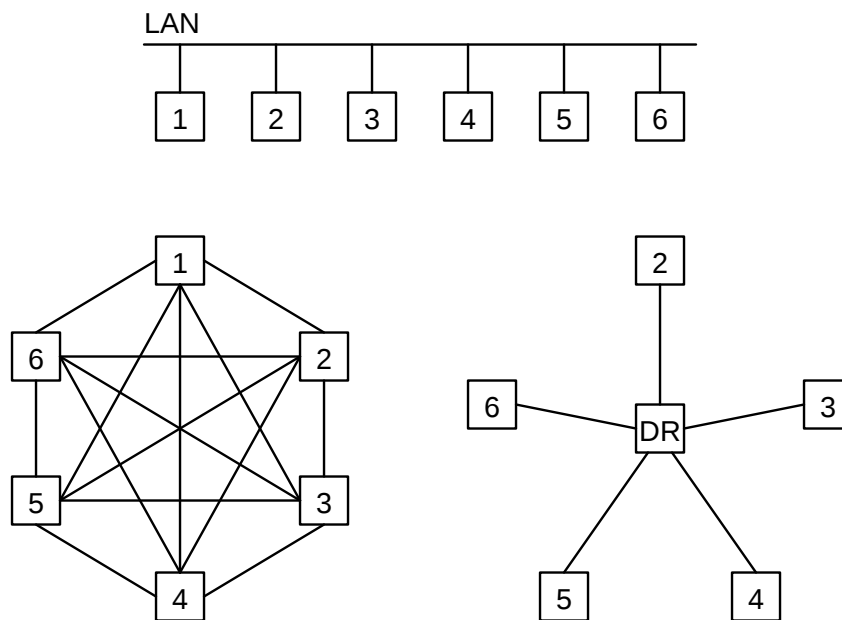
Figura 3.27. OSPF. Due aree che benché connesse direttamente da una LAN multicast non la utilizzano per le comunicazioni interarea poiché su di essa non vi è un area border router.

In integrated IS-IS un IS appartiene ad un area a seconda del suo NSAP e del suo manual area addresses, per cui non può appartenere a più di un area. Una LAN può allora appartenere a più di un area se gli IS ad essa connessi hanno area address differenti; in tal caso la LAN è utilizzata per colloquiare solo se gli IS sono di livello 2. Lo svantaggio è lo stesso di OSPF in quanto l'indirizzo IP della LAN viene annunciata in aree differenti.

In definitiva una LAN è utilizzata da due IS per colloquiare solo se sono entrambi di livello 1 e appartengono alla stessa area oppure se sono entrambi di livello 2.

3.5.3 Designated router

Entrambi i protocolli posseggono un meccanismo di *designated router* per semplificare la topologia di aree dove esistano diversi router connessi alla stessa LAN multicast. Uno di questi router viene designato pseudonodo, cioè che annuncia sé stesso come nodo a cui sono connessi direttamente tutti gli altri; si evita così che ogni router annuncii tutti i collegamenti ad ogni altro router della LAN complicando la topologia del grafo che rappresenta l'area (vedere esempio in figura 3.28).



Grafo senza Designated Router

Grafo con Designated Router

Figura 3.28. Una LAN multicast ed il grafo corrispondente senza e con il designated router.

Il designated router OSPF differisce da quello integrated IS-IS per come effettua il riscontro degli LSP dei router a lui sottoposti: in OSPF il desig-

nated router è responsabile di collezionare espliciti acknowledgement inviati dagli altri router della LAN per confermare la ricezione degli LSP inviati dal designated router. Quando un router su una LAN riceve un LSP lo invia al designated router il quale poi provvede ad inviarlo in multicast a tutti gli altri router della LAN; questi devono poi rispondere ognuno con un proprio ack.

In integrated IS-IS, il designated IS genera i CSNP per il riscontro dei LSP ricevuti su una LAN: gli altri router osservano il CSNP e i loro database di LSP e se vi sono differenze richiedono aggiornamenti, altrimenti non viene generato nessun pacchetto.

La soluzione di integrated IS-IS permette di ridurre la banda occupata per la trasmissione affidabile dei LSP.

3.5.4 Routing Table e database

Nella routing table (forwarding database) sono elencati gli indirizzi di rete noti e i router direttamente connessi da cui passare per arrivarci; i percorsi però possono differire per il tipo di servizio (TOS) da privilegiare. In OSPF la routing table è unica e se i percorsi per una certa destinazione differiscono a seconda del TOS vi sono elencati tutti. Integrated IS-IS invece ha un forwarding database per ogni metrica.

Anche il topological database è unico per OSPF mentre Integrated IS-IS ne mantiene uno per ogni metrica.

Le due soluzioni si equivalgono anche se quella di OSPF occupa meno memoria nei router.

3.5.5 Collegamenti punto a punto

In OSPF un collegamento punto a punto può essere unnumbered o avere un indirizzo di sottorete; in entrambi i casi è utilizzato se sui due router alle estremità del link le interfacce sono configurate per appartenere alla stessa area.

In integrated IS-IS un collegamento punto a punto consente ai due IS alle estremità di essere vicini sempre che siano entrambi di livello 1 e ap-

partengano alla stessa area, oppure se sono entrambi di livello 2.

In integrated IS-IS la configurazione è più semplice non essendo necessario definire indirizzi delle interfacce o aree di appartenenza.

3.5.6 Considerazioni finali

Quale dei due protocolli sia migliore in modo assoluto rimane difficile da stabilire; ciascuno dei due possiede alcune qualità che non possiede l'altro, ma, in definitiva, nessuno dei due supera l'altro in tutto. Tra le qualità, integrated IS-IS ha quelle di essere più lineare da configurare e di occupare meno banda trasmissiva con i suoi pacchetti; OSPF invece è più efficiente nel routing interarea. Una eventuale scelta deve basarsi sulla infrastruttura esistente in cui il protocollo verrà usato e sugli aspetti che si vogliono prediligere.

Un fatto che sembra però essere accertato è che, al momento, OSPF è quello tra i due protocolli a cui vengono prestate maggiori attenzioni e che rimane soggetto di continui studi e sviluppi. Prova ne è il fatto che la RFC⁶ che ne contiene le specifiche è stata aggiornata nel marzo 1994; inoltre, è stata scritta una proposta di OSPF per IPng e si sta studiando il suo interfacciamento con i protocolli interdomain. In definitiva, l'ambiente degli sviluppatori dell'IETF preferisce lavorare senza i vincoli dati a integrated IS-IS dall'ISO che ne è "proprietaria".

⁶RFC: Request For Comment; sono documenti che contengono standard, misure, proposte, specifiche dei protocolli.

4

Protocolli di routing sulle reti INFNet e GARR

Le reti INFNet e GARR sono pienamente integrate fra di loro e insieme formano l'autonomous system 137 che collega i principali centri di ricerca in Italia. Come però verrà evidenziato in seguito, non è utilizzato un unico protocollo di routing, ma bensì alcuni diversi protocolli per soddisfare le differenti necessità che la presenza di IP e CLNP richiede.

4.1 Autonomous system 137

Per semplificare le operazioni di routing l'Internet è suddivisa in *autonomous system*; un autonomous system è un insieme di reti che seguono una politica di routing comune e dove ogni host sa raggiungere ogni altro nel modo migliore. All'interno di un autonomous system sono eseguiti i protocolli di routing *intradomain*. Le comunicazioni tra i vari autonomous system sono effettuate tramite gli *Autonomous System Boundary Router* e gestite dai protocolli *interdomain*.

Il GARR e INFNet costituiscono un unico autonomous system identificato dal numero 137 (AS 137); al suo interno il routing è garantito dai protocolli IGRP, Integrated IS-IS, OSPF, RIP. Questi protocolli sono anche in grado di instradare pacchetti destinati ad altri autonomous system informando ogni router interno su come raggiungere l'autonomous system boundary router

più vicino.

Nell'AS 137 gli AS boundary router sono due: garr-gw.cnr.it al Cnuc di Pisa e Cnaf-gw-1.infn.it al CNAF¹ di Bologna; al primo è connessa la linea con EUROPAnet (AS 2043); al secondo la linea col CERN (HEPnet AS 513), quella con ESnet (AS 3425) e, tramite Garr-gw-gs, quella con il laboratorio JINR di Dubna in Russia. In figura 4.1 sono riportate queste connessioni e quelle tra gli autonomous system connessi al GARR.

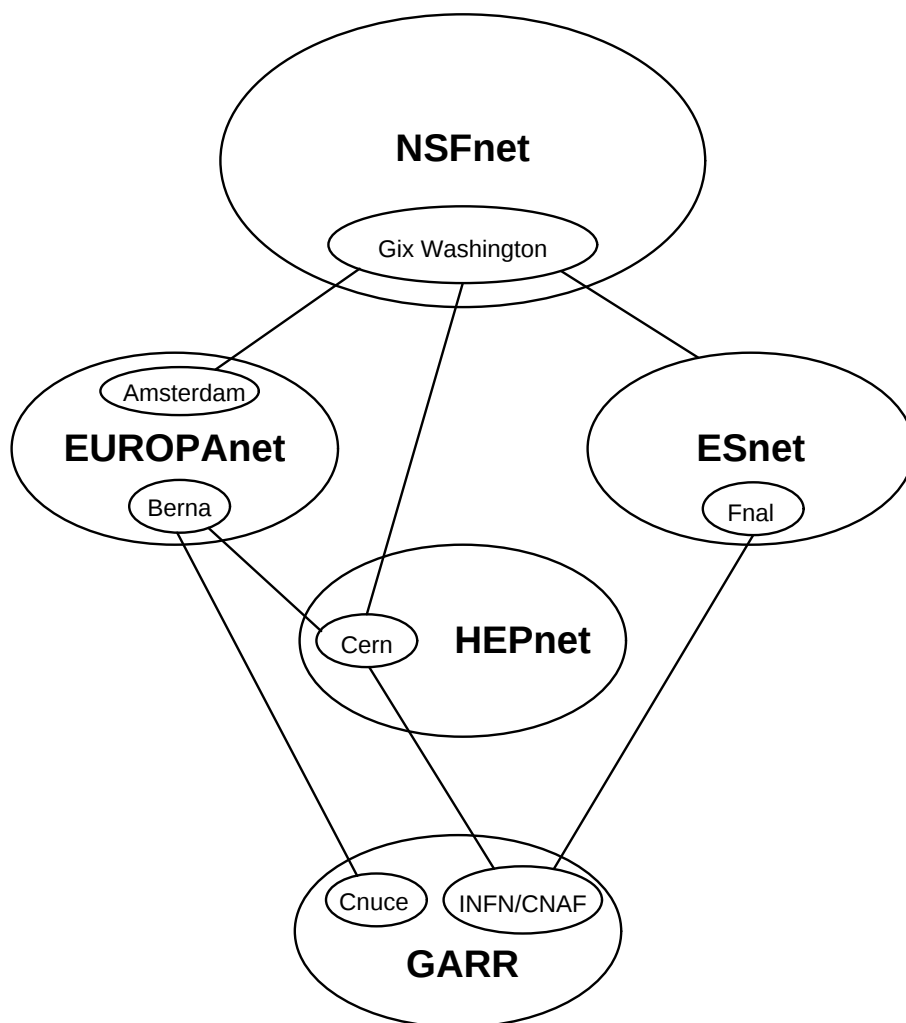


Figura 4.1. Collegamenti internazionali dell'AS 137.

¹CNAF: Centro Nazionale Analisi Fotogrammi; è una sezione dell'INFN con sede a Bologna che si occupa essenzialmente della gestione e dello sviluppo dell'INFNet.

Gli autonomous system confinanti con l'AS 137 e con cui vengono scambiate informazioni di routing sono l'AS 2043 (EUROPAnet) e l'AS 2038; quest'ultima ha il compito di interfacciare il GARR con HEPnet (CERN) ed ESnet (vedere figura 4.2). Gli autonomous system confinanti con l'AS 2038 sono l'AS 3425 (ESnet autonomous system), l'AS 513 (CERN autonomous system) e l'AS 2875 (Dubna autonomous system). Il protocollo di routing interdomain utilizzato è il BGP-4.

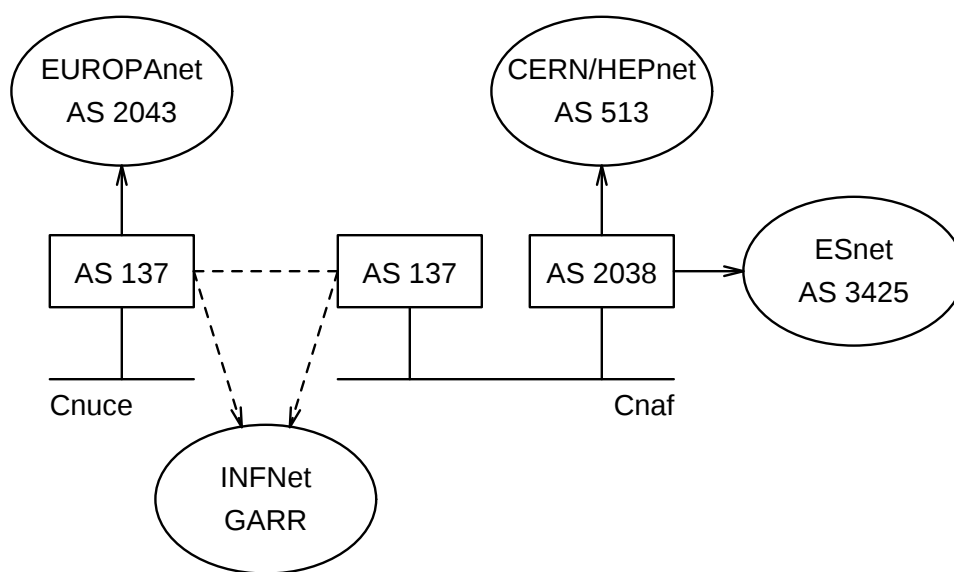


Figura 4.2. Autonomous system confinanti con l'AS 137.

Le reti normalmente nascono per soddisfare le esigenze di enti che le finanziano, per cui su tali reti viaggia principalmente il traffico di informazioni che interessa appunto gli enti proprietari; ad esempio, HEPnet garantisce il collegamento con tutti i laboratori HEP, mentre EUROPAnet trasporta traffico interdisciplinare ed ha una interconnettività molto estesa. L'INFN ha link con il CERN e con il Fermilab di Chicago perché è suo interesse avere collegamenti diretti con questi importanti centri di ricerca di fisica nucleare. Il CNR invece ha un link con EUROPAnet per attestare la presenza dell'Italia nella rete che connette le nazioni europee.

Gli indirizzi dell'AS 137 vengono fatti conoscere dal protocollo BGP in tutti gli autonomous system connessi (EUROPAnet, HEPnet, ESnet) e ad altri confinanti con questi. Gli autonomous system che non conoscono gli

indirizzi dell'AS 137 si rivolgono al GIX² di Washington per raggiungere gli host del GARR. Questo GIX è anche usato dal GARR per raggiungere reti i cui indirizzi non gli sono noti; infatti all'interno del GARR, quando si devono contattare reti sconosciute, si raggiunge sempre ESnet tramite cnaf-gw-1.infn.it (al CNAF); se tali reti non sono note nemmeno ai router di ESnet questi passano al GIX di Washington il quale sicuramente le conosce (se esistono).

All'interno del GARR tutti gli host usano gli stessi link per raggiungere una medesima destinazione; questo significa che, ad esempio, per raggiungere un host di HEPnet tutti gli host del GARR passano per i router del CNAF, compresi quelli del CNUCE; allo stesso modo per raggiungere un host di Funet (Finlandia) tutti gli host del GARR passano per il CNUCE, compresi quelli del CNAF che avrebbero una strada più corta passando per il CERN.

Il link con il CERN viene normalmente usato per raggiungere HEPnet e cioè i centri di ricerca di fisica nucleare, mentre il link con EUROPA-net viene utilizzato per connettersi alle altre reti europee. Il link con il Fermilab serve invece per raggiungere tutte le altre reti nel mondo.

Osservando la figura 4.1 si può notare come i collegamenti tra i diversi autonomous system siano magliati; questo permette di avere percorsi di backup anche per le destinazioni internazionali. Si nota, ad esempio, come il GIX di Washington sia raggiungibile direttamente sia da EUROPA-net che da HEPnet e da ESnet e che tutti questi tre autonomous system sono direttamente connessi con il GARR; questo implica l'esistenza di tre percorsi differenti.

4.2 I router

All'interno del GARR sono in funzione essenzialmente due tipi di router: quelli costruiti dalla Digital e quelli costruiti dalla Cisco. I Digital sono stati i primi ad essere acquistati poiché storicamente le workstation usate nell'INFN erano di questa casa costruttrice e INFNet era sorta come rete DECnet. I router Cisco vengono usati da quando sulla rete GARR è in

²GIX: Global Interconnections eXchange; router che conosce tutti gli indirizzi di rete di tutta l'Internet.

funzione anche il protocollo IP.

Purtroppo, benché i protocolli di routing siano definiti da standard di pubblico dominio, spesso si verificano incompatibilità nella connessione di macchine di produttori diversi; inoltre il software che gestisce i router non è ancora ottimizzato per cui durante l'utilizzo si rileva abbastanza frequentemente la presenza di “buchi”³. Per questo, le configurazioni studiate sulla carta a volte non possono essere realizzate nella realtà costringendo ad adottare soluzioni non sempre ottimali.

4.3 I protocolli di routing

I protocolli di routing utilizzati sono diversi: RIP, IGRP, Integrated IS-IS, OSPF, BGP. La scelta dell'uno o dell'altro ha innanzitutto ragioni tecniche: facilità di configurazione, affidabilità, protocolli del network layer che li necessitano, tipo di routing che devono effettuare. Purtroppo la convivenza di tanti protocolli genera diversi problemi di interfacciamento risolvibili con una attenta configurazione, ma sempre a rischio di malfunzionamenti ogni volta che si effettuano modifiche alle configurazioni o si aggiungono nuovi siti.

4.3.1 RIP

RIP è un protocollo molto limitato in quanto non adatto a reti di grandi dimensioni, poco reattivo ai cambiamenti di topologia, non sempre affidabile per i lunghi tempi di convergenza ma comunque facile da configurare; per questo è sempre meno usato e rimane solo dove questi limiti sono meno influenti. Essenzialmente infatti viene utilizzato localmente in siti con diversi router che usano diversi protocolli, per effettuare il redistribute di informazioni tra di questi .

Ad esempio al CNAF (in figura 4.3 è riportata la mappa dei router del CNAF) tutti i router sono connessi alla LAN 131.154.1.0 e tale LAN è usata da tutti i protocolli di routing; RIP è usato su niscn3, cnaf-gw-2 e garr-gw-bo; in niscn3, integrated IS-IS è ridistribuito in RIP; negli altri due, RIP è

³Buchi: difetti nel software.

ridistribuito in IGRP; in questo modo avviene il passaggio di informazione da integrated IS-IS a IGRP senza che vi sia un diretto redistribute tra i due.

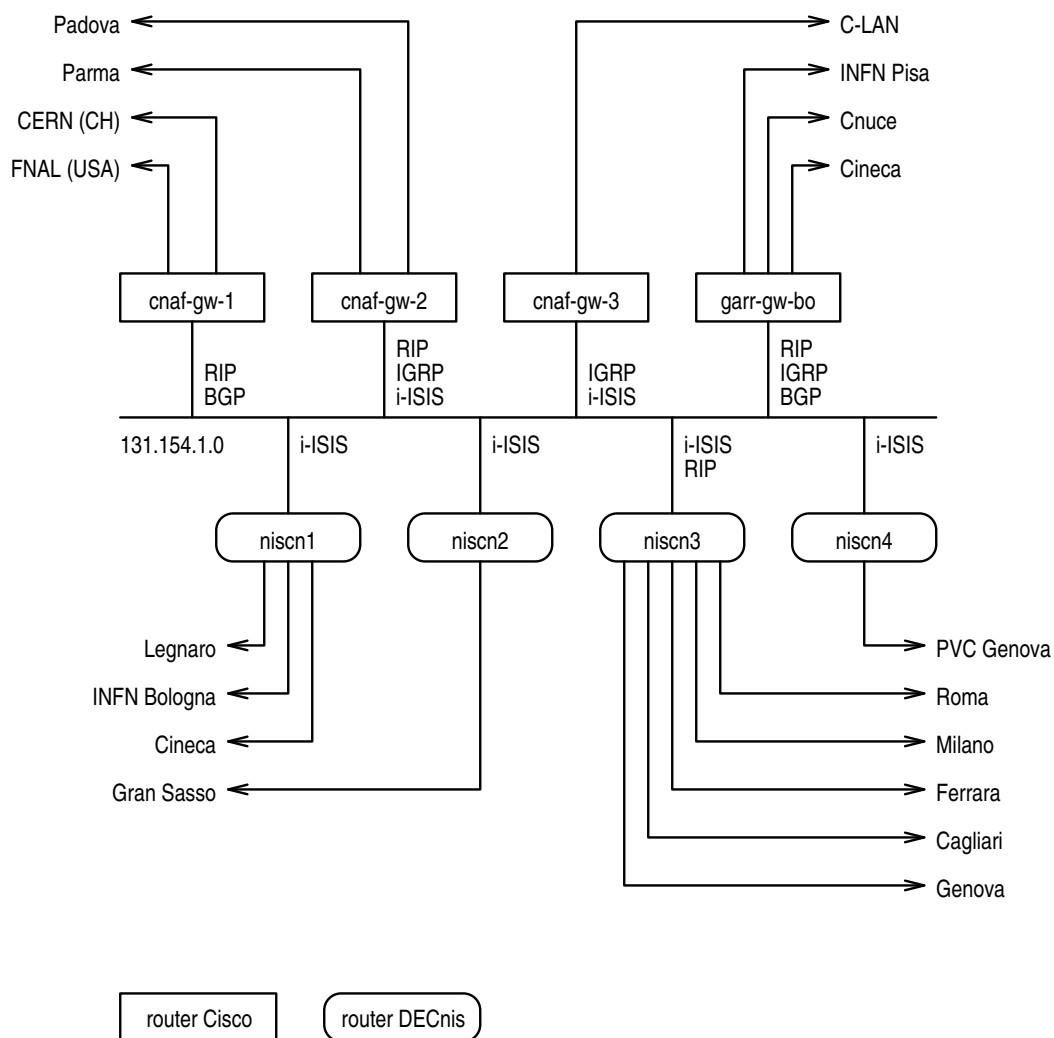


Figura 4.3. Mappa dei router del CNAF

4.3.2 IGRP

IGRP è il protocollo più usato all'interno dell'autonomous system 137 benché sia di proprietà della Cisco; infatti la maggior parte dei router del GARR è di questa casa produttrice.

IGRP è un protocollo che usa l'algoritmo distance vector e che risente di tutti i limiti che ne derivano: bassa velocità di convergenza, rischi di loop; nonostante questo è usato per la sua semplicità di configurazione, per il suo funzionamento collaudato e soprattutto perché è stato il primo protocollo che si potesse usare con affidabilità su ampi autonomous system. La tendenza oggi è comunque quella di sostituirlo con protocolli che usano l'algoritmo link state come OSPF e Integrated IS-IS che assicurano una maggiore velocità di convergenza e che sono ormai pienamente affidabili.

IGRP 137 è il protocollo usato da tutti i router del GARR denominati *garr-gw* (abbreviazione di “garr gateway”⁴). Questo protocollo è responsabile del routing sulla dorsale IP, costituita dalle linee a 2 Mbps tra Roma (*garr-gw-rm.infn.it*), Gransasso (*garr-gw-gs.infn.it*), Bologna-CNAF (*garr-gw-bo.infn.it*), Bologna-Cineca (*garr-gw.cineca.it*), Pisa (*garr-gw.cnr.it*), Milano-Cilea (*garr-gw.cilea.it*), Torino (*garr-gw.polito.it*), Napoli (*garr-gw.unina.it*), Csata (*garr-gw.csata.it*). È inoltre usato in altri siti raggiunti da linee a velocità più bassa come Milano-INFN (*garr-gw-mi.infn.it*), Bari (*garr-gw.ba.infn.it*), L'Aquila (*garr-gw-sci.aquila.infn.it*), Perugia (*garr-gw-pg.unipg.it*), Lecce (*garr-gw.unile.it*) oltre a Ferrara, Pavia, Firenze, Sassari, Cagliari.

Nella figura 4.4 è riportata la mappa dei siti dove oggi è utilizzato IGRP 137. Come si può osservare, il routing del backbone IP del GARR dipende da IGRP. IGRP permette di assegnare diversi costi ai link; ad esempio a velocità più alte si può fare corrispondere un costo più basso in modo da forzare l'utilizzo delle linee ad alta velocità come quelle tra CNAF, Cineca, Roma e Gransasso per comunicazioni tra siti non direttamente connessi.

4.3.3 Integrated IS-IS

Integrated IS-IS è un'altro dei protocolli più usati all'interno dell'autonomous system 137 e sta lentamente sostituendo IGRP. Il suo utilizzo è dovuto innanzitutto alla presenza del CLNP di cui IS-IS è l'unico protocollo di routing: per avere un protocollo link state anche per IP e per evitare di averne

⁴Gateway è il termine IP per indicare i router.



Figura 4.4. Mappa dei siti del GARR che impiegano IGRP 137.

due completamente separati l'adozione di Integrated IS-IS è stata obbligata. Si deve anche aggiungere che nella rete sono utilizzati molti router Digital che non possono eseguire IGRP e che inizialmente hanno avuto delle difficoltà ad eseguire correttamente OSPF. Dunque la scelta di Integrated IS-IS è risultata quasi forzata.

Anche integrated IS-IS ha avuto però inizialmente dei problemi di funzionamento; questi, insieme a quelli di OSPF, sono dovuti al fatto che i protocolli link state sono tutt'ora poco usati in Internet e quindi il software non è ancora stato provato su topologie complicate.

All'interno dell'autonomous system 137, integrated IS-IS è usato in quasi tutti i siti dell'INFN poiché sono connessi anche tramite DECnet/OSI che ha come protocollo di rete il CLNP; la rete IP infatti ha affiancato una infrastruttura DECnet già presente e funzionante che ha determinato le scelte di sviluppo. I percorsi attraverso il dominio integrated IS-IS è dunque sempre preferito per le comunicazioni tra host di questi siti.

Nella figura 4.5 è riportata la mappa dei siti dove oggi è utilizzato integrated IS-IS.

4.3.4 OSPF

OSPF è un protocollo link state che all'interno del GARR è usato poco. I motivi per cui tale protocollo è stato ignorato sono da ricercare tra quelli che hanno spinto alla diffusione di integrated IS-IS; una convivenza di IS-IS e OSPF di tipo Ship In the Night su router non particolarmente potenti risulta infatti problematica esaurendo le risorse fisiche e quindi provocando crolli di prestazione. Questo problema esisteva al tempo in cui è stato preferito integrated IS-IS; oggi i nuovi router sono in grado di supportarli entrambi per cui un suo impiego è stato preso di nuovo in esame dovendo essere definita una nuova politica di routing dell'autonomous system 137.

All'interno della rete GARR, OSPF è usato tra Pisa, Firenze, San Piero a Grado (INFN) e Sassari (CNR).

Nella figura 4.6 è riportata la mappa dei siti dove è utilizzato OSPF.

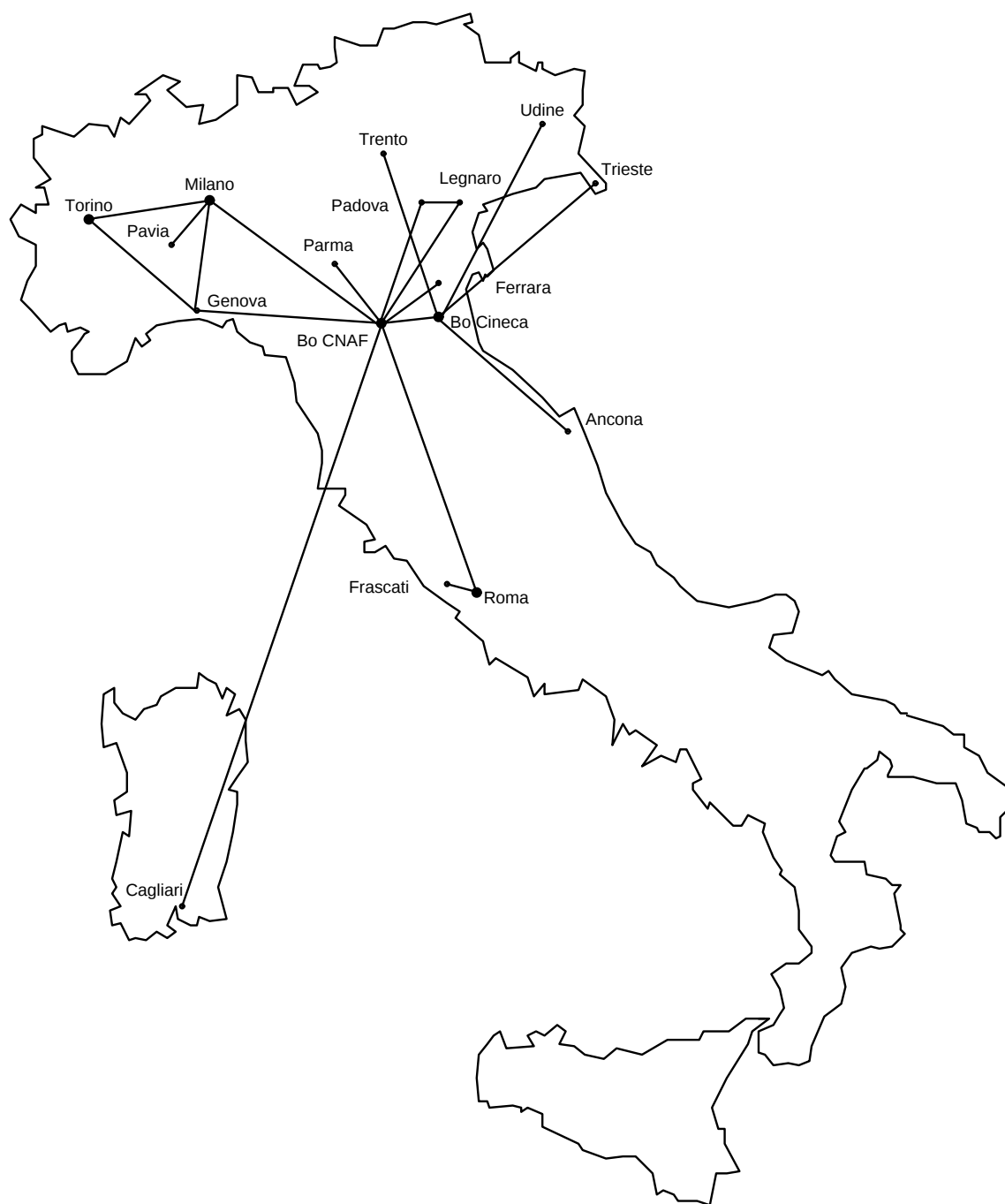


Figura 4.5. Mappa dei siti GARR che impiegano integrated IS-IS.



Figura 4.6. Mappa dei router del GARR che impiegano OSPF.

4.3.5 BGP

BGP è il protocollo interdomain utilizzato dai router del GARR collegati con altri autonomous system; la sua scelta è stata obbligata essendo il protocollo già in uso nei router a cui sono collegati quelli di confine dell'autonomous system 137.

BGP è utilizzato a Bologna (cnaf-gw-1.infn.it e garr-gw-bo.infn.it) e a Pisa (garr-gw.cnr.it); effettua il routing internazionale come riportato nella sezione 4.1.

4.4 Il redistribute

Per consentire lo scambio di informazioni tra i diversi protocolli di routing, è necessario effettuare dei *redistribute* tra coppie di essi su alcuni dei router dove sono attivi entrambi.

L'operazione di redistribute consiste nel fare in modo che un protocollo annunci a quello in cui viene redistribuito tutte le destinazioni che ha nella routing table, con i relativi costi. Il protocollo che riceve queste nuove informazioni le tratta come destinazioni esterne al proprio routing domain; questo significa che se un indirizzo è già presente viene comunque preferito il percorso interno al routing domain anche se più lungo.

Nella terminologia di integrated IS-IS e OSPF, routing domain e autonomous system hanno lo stesso significato di “insieme di router che eseguono lo stesso protocollo di routing”; nella terminologia strettamente IP, *routing domain* ha ancora lo stesso significato, mentre con *autonomous system* si intende un insieme di router che eseguono un routing coerente e che esternamente (da altri autonomous system) è visto come un unico blocco, benché al suo interno possano essere eseguiti diversi protocolli di routing. In sostanza un routing domain forma un singolo autonomous system, ma un autonomous system può essere composto da più di un routing domain.⁵

Nell'AS 137 sono eseguiti tre principali algoritmi di routing intradomain: IGRP, integrated IS-IS e OSPF; ognuno determina un routing domain. Per

⁵[Bgp4ospf]

estendere i routing domain a tutto l'autonomous system è quindi necessario fare dei redistribute tra i vari protocolli. Queste operazioni però richiedono una attenta analisi sul dove effettuarli, in quanto si possono verificare problemi di inconsistenza nelle tabelle di routing; infatti, se una rete viene annunciata in due redistribute diversi, oppure viene prima ridistribuita da un protocollo in un altro e successivamente dal secondo al primo, si possono verificare dei loop, cioè dei percorsi ad anello in cui non si giunge mai alla destinazione.

Per integrated IS-IS e IGRP la politica seguita è stata quella che, in ogni sito dove sono entrambi presenti questi protocolli, integrated IS-IS annuncia ad IGRP solo le reti di competenza del sito. Ad esempio, a Milano integrated IS-IS annuncia a IGRP solo le reti di Milano il cui routing dipende da integrated IS-IS; a Bologna integrated IS-IS annuncia a IGRP le reti di Genova, Bologna e Legnaro poiché in tutti questi siti non è eseguito IGRP.

La redistribuzione avviene o direttamente su quei router che eseguono entrambi i protocolli oppure attraverso RIP, come indicato nell'esempio del paragrafo 4.3.1. Nel GARR questi redistribute avvengono a Torino, Roma, Bologna (CNAF e Cineca) e Milano.

Lo scambio di informazioni da IGRP a integrated IS-IS avviene in modo differente: non viene effettuato nessun redistribute da IGRP a integrated IS-IS ma in ogni sito i router che eseguono integrated IS-IS hanno una route statica al router (Cisco) che esegue IGRP. In questo modo se un router del routing domain integrated IS-IS non trova la destinazione nella propria routing table, passa al routing domain IGRP attraverso il router di default e da lì l'instradamento avviene sulla base delle tabelle di routing IGRP.

Questo modo di procedere serve ad evitare che una stessa rete del dominio integrated IS-IS venga annunciata a due router IGRP differenti

OSPF esegue il routing per le reti della Toscana dove è l'unico protocollo eseguito. A Pisa, precisamente su garr-gw.cnr.it, è eseguito il redistribute di OSPF in IGRP; in questo modo IGRP viene a conoscere tutte le reti della Toscana. Di conseguenza un host del dominio IGRP passerà sempre per garr-gw.cnr.it per raggiungere un host del dominio OSPF.

Su questo router avviene anche l'operazione inversa, cioè IGRP viene ridistribuito in OSPF. Questo non crea problemi poiché il punto di contatto tra i due protocolli è unico e non vi sono indirizzi in comune.

Tra OSPF e integrated IS-IS non vi è alcun contatto e quindi nemmeno redistribute.

4.4.1 Percorsi asimmetrici

L'utilizzo di diversi protocolli in un unico autonomous system richiede un'attenta configurazione degli stessi e l'uso di tecniche come i redistribute e le route statiche. Queste però possono avere delle conseguenze indesiderate come ad esempio i percorsi asimmetrici.

Si ha un *percorso asimmetrico* quando, durante la comunicazioni tra due host i pacchetti in una direzione seguono un percorso diverso da quelli nella direzione opposta. La caduta di un link può interrompere il flusso in una direzione, ma non in quella opposta e senza che venga ricostruito il percorso di backup simmetrico.

Nella figura 4.7 è riportata una topologia reale di confine tra integrated IS-IS e IGRP; durante la migrazione dell'INFN di Bologna da IGRP a integrated IS-IS era stata realizzata la configurazione seguente che causava dei percorsi asimmetrici. Il redistribute di integrated IS-IS in IGRP avveniva al CNAF tramite RIP; sulla ethernet del dipartimento di fisica non vi era redistribute tra i due router ed i suoi host avevano come default gateway il cisco.

Quando vaxbo e gpxbof volevano comunicare, i pacchetti da vaxbo a gpxbof attraversavano solo bofisir poiché a questo era connessa la rete del dipartimento di fisica ed inoltre vaxbo puntava a bofisir. I pacchetti da gpxbof a vaxbo invece passavano per il Cineca e per il CNAF poiché gpxbof puntava al cisco e tra questi e bofisir non vi era nessuno scambio di informazioni. In questa situazione, se uno dei link tra Fisica, Cineca e CNAF cessava di funzionare, per gpxbof diventava impossibile raggiungere vaxbo ma non viceversa.

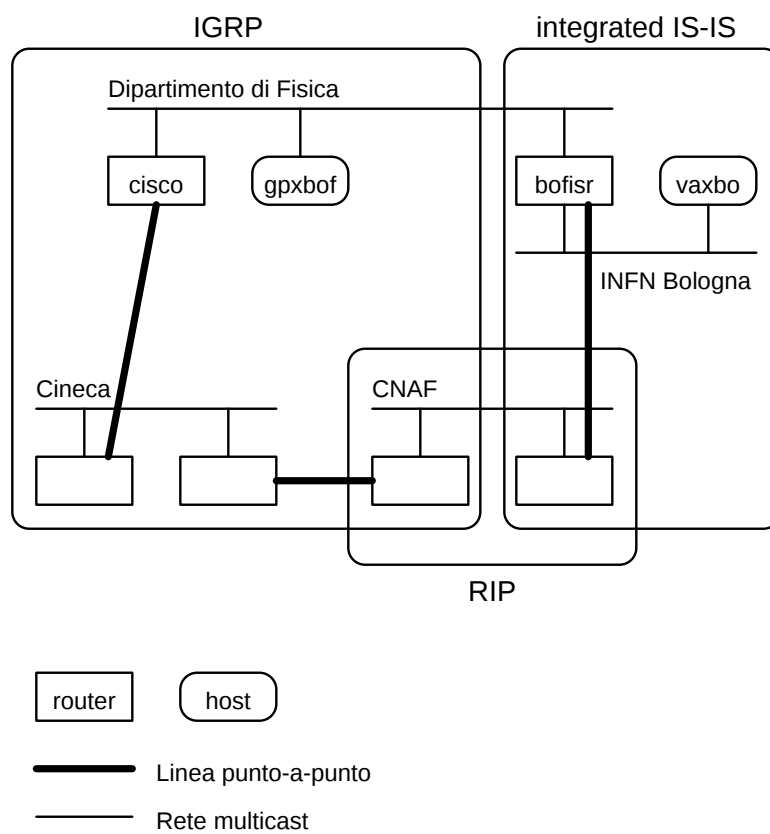


Figura 4.7. Topologia reale di una zona di confine tra IGRP e integrated IS-IS.

5

Prove di funzionalità dei protocolli di routing

In questo capitolo sono descritte alcune prove sui dei protocolli link state effettuate nella rete GARR nei mesi di settembre e ottobre 1994. Lo scopo di questi test è stato la verifica della funzionalità ed efficienza dei protocolli in previsione della possibile riconfigurazione del routing nel GARR.

5.1 Integrated IS-IS su C-LAN frame relay

Frame relay è un protocollo di rete che si colloca nel Data link layer del modello OSI. Esso è utilizzato dalla Telecom Italia (ex SIP) sulla rete pubblica a commutazione di pacchetto denominata C-LAN. Su questa rete pubblica è possibile definire delle reti private virtuali con accessi a velocità compresa tra 64 kbps a 2 Mbps e con PVC che sostituiscono le normali linee punto a punto.

L'INFN ha i siti di Bologna, Milano, Torino, Roma, Genova e Cagliari connessi alla C-LAN; i collegamenti (PVC) tra le varie sedi sono riportati in figura 5.1.

Questa prova ha voluto verificare la funzionalità della C-LAN e l'utilizzo su di essa del protocollo di routing link state Integrated IS-IS.

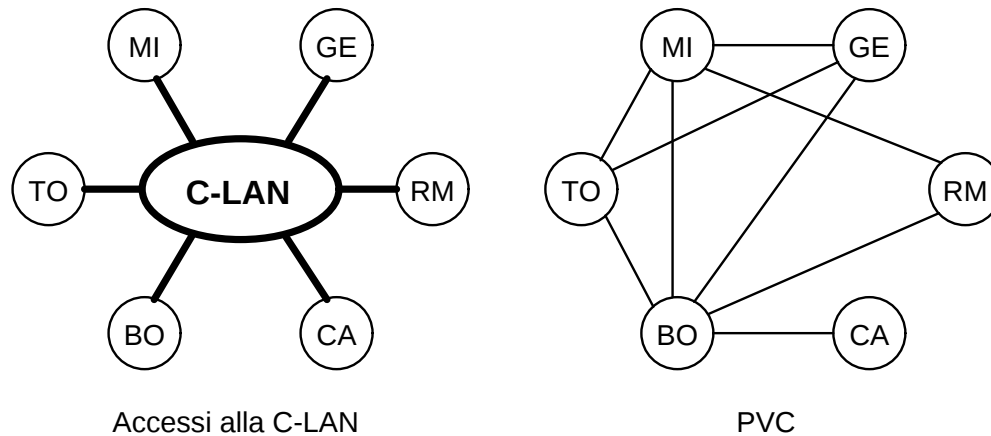


Figura 5.1. Accessi alla C-LAN e PVC tra alcuni siti INFN.

5.1.1 Frame relay¹

La tecnologia *frame relay* fornisce una comunicazione a commutazione di pacchetto ad alta affidabilità ed efficienza. È una soluzione economica e veloce al collo di bottiglia costituito dal collegamento delle LAN attraverso le WAN: infatti in ogni sito connesso viene pagato solamente l'accesso alla rete frame relay e non il numero dei PVC realizzati; inoltre, utilizzando linee ad alta affidabilità come le fibre ottiche, non sono sprecate risorse in pesanti controlli sulla correttezza delle trasmissioni, fornendo così una maggiore velocità di comunicazione.

Frame relay è un protocollo che necessita di un protocollo di linea come l'HDLC² o il PPP³ per realizzare tutte le funzioni del data link layer; in figura 5.2 è riportato lo schema di come frame relay si inserisce nel modello OSI.

Frame relay realizza una comunicazione a commutazione di pacchetto che viene usata come interfaccia tra l'utente (DTE⁴) e la rete fisica (DCE⁵); gli switch frame relay sono in grado di eseguire un multiplexing statistico di trasmissioni per diverse destinazioni su di un singolo collegamento fisico;

¹[Black] capitolo 13 e [Cisco] capitolo 14.

²HDLC: High-level Data Link Protocol.

³PPP: Point-to-Point Protocol.

⁴DTE: Data Terminal Equipment.

⁵DCE: Data Circuit-terminating Equipment.

Transport layer	TCP	
Network layer	IP	
Data Link layer	HDLC	PPP
	Frame Relay	

Figura 5.2. Il protocollo frame relay nel modello OSI.

con *multiplexing statistico* si intende la capacità di allocare la banda per ogni connessione in maniera flessibile, in base a quante trasmissioni sono effettivamente in corso in modo da ottimizzare in ogni istante l'uso della banda disponibile.

I collegamenti fra i router connessi alla rete frame relay sono costituiti da circuiti virtuali, configurati e gestiti dall'amministratore del servizio: essi sono chiamati *Permanent Virtual Circuit* (PVC). I PVC sono identificati tramite i *Data Link Connection Identifier* (DLCI) codificati in dieci bit.

Insieme alle funzioni di base del protocollo frame relay, sono state specificate alcune estensioni LMI che facilitano la gestione di reti estese e complesse; le principali sono:

- *virtual circuit status messages*: informano sullo stato dei PVC;
- *multicasting*: permettono di spedire pacchetti a diverse destinazioni;
- *Simple flow control*: realizzano un meccanismo di controllo del flusso.

5.1.2 Obiettivi dei test

L'utilizzo di una rete formata da PVC frame relay invece che dai normali collegamenti punto-a-punto basati su linee dedicate o affittate ha diversi vantaggi:

- dal punto di vista economico, a parità di velocità di trasmissione il costo è minore sulle lunghe distanze; inoltre viene pagato un unico accesso alla rete qualunque sia il numero di PVC definiti;
- la funzionalità dei PVC è garantita dal fornitore del servizio per cui il back-up di un collegamento fisico che non funziona è trasparente all'utente;

- se l'infrastruttura fisica è in grado di garantire un'alta affidabilità di trasmissione, frame relay consente di raggiungere più alti valori di throughput poiché vengono persi, e quindi ritrasmessi, meno pacchetti.
- i PVC sono facilmente riconfigurabili (ad opera del fornitore del servizio) per cui si può valutare nel tempo e col mutare delle esigenze quali collegamenti virtuali siano necessari e cambiarli rapidamente di conseguenza.

Gli obiettivi del test sono stati diversi: innanzitutto di provare la funzionalità, la stabilità e l'efficienza della C-LAN; queste sono state verificate prima con il semplice routing statico per non introdurre troppi fattori di complessità; poi con l'utilizzo del protocollo di routing link state integrated IS-IS configurato per il routing domain costituito dai siti connessi alla C-LAN, senza e con il traffico di produzione; con *traffico di produzione* si intendono tutte le trasmissioni effettuate da tutti gli host dei vari siti, in pratica il traffico degli utilizzatori finali dell'internet. Durante queste prove è anche stato verificato il funzionamento dei router Cisco come frame relay switch.

5.1.3 Hardware e software impiegato

L'hardware ed il software dei router coinvolti nei test è stato il seguente:

- Cnaf-gw-3 è un router Cisco AGS+ con 16 MB di memoria e software GS versione 10.0(2.7), situato presso il CNAF/INFN di Bologna.
- Garr-gw-mi è un router Cisco AGS+ con 16 MB di memoria e software GS versione 10.0(2.7), situato presso la sezione INFN di Milano.
- Cisco.clan.to è un router Cisco AGS+ con software GS versione 10.0 (2.7), situato presso la sezione INFN di Torino.
- Gw1.crs4.it è un router Cisco AGS+ con software GS versione 10.0(2.7), situato presso il CRS4⁶ di Cagliari.
- Garr-gw-rm è un router Cisco AGS+ con 16 MB di memoria e software GS versione 10.0(2.7), situato presso l'INFN di Roma.
- Rt2ge5 è un router Digital DECnis600 con software Decnis versione 2.3.2, situato presso la sezione INFN di Genova.

⁶CRS4: Centro di Ricerca, Sviluppo e Studi Superiori in Sardegna.

5.1.4 Funzionalità della C-LAN

Scopo di questo test è stato di verificare la funzionalità, l'affidabilità e la stabilità della rete C-LAN e delle interfacce frame relay dei router connessi con il semplice routing statico effettuando misure di throughput sui vari PVC. Il routing statico è stato utilizzato al fine di non introdurre ulteriori complessità nella valutazione delle prestazioni.

I siti interessati in questa prova sono stati Milano, Genova, Bologna, Torino e Roma; in figura 5.3 è riportato il particolare della topologia di questo test.

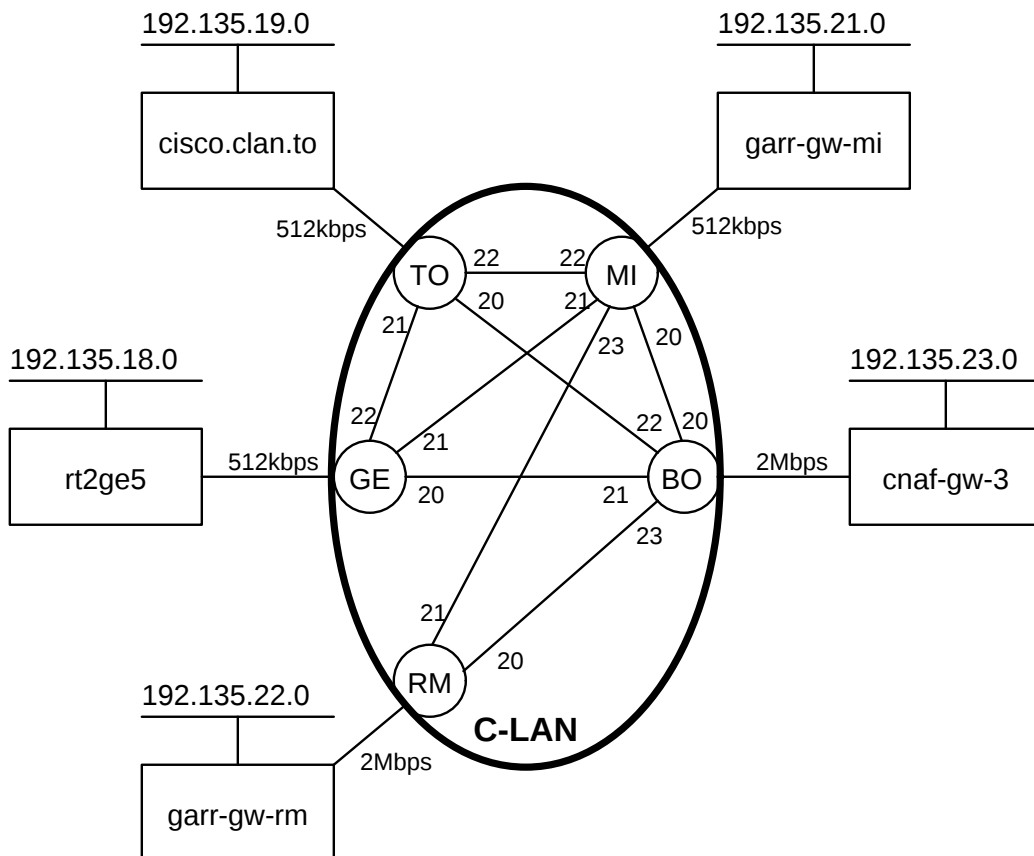


Figura 5.3. Topologia test funzionalità C-LAN con routing statico; i numeri a fianco dei PVC indicano i DLCI.

La configurazione dei router per questo test è stata abbastanza semplice e non ha dato problemi: è stato sufficiente abilitare le interfacce frame relay

e impostare il routing statico per ogni PVC. Per le prove, in ogni sito è stata connessa una macchina sulla stessa ethernet del router; da notare che i numeri di rete utilizzati erano visibili esclusivamente all'interno della C-LAN e che le macchine di prova non potevano raggiungere nessun altro indirizzo esterno.

La tabella 5.1 riporta i valori massimi di throughput realizzati fra i siti direttamente connessi.

da \ a	Milano	Genova	Torino	Bologna	Roma
Milano	-	463	456	457	465
Genova	460	-	470	460	-
Torino	470	470	-	430	-
Bologna	471	461	471	-	1100
Roma	460	-	-	1100	-

Tabella 5.1. Throughput tra i vari siti espresso in kbps.

Tenendo conto dell'overhead del protocollo (stimato circa 8,4%), negli accessi a 512 kbps ci si è avvicinati al massimo throughput consentito. Sul PVC fra Roma e Bologna invece il throughput massimo rilevato è stato di 1,1 Mbps, molto inferiore al massimo consentito (2 Mbps); questo fatto non è da imputare alle macchine coinvolte che permettono prestazioni migliori.

Eseguendo trasferimenti contemporanei da diversi siti che però confluivano tutti sull'accesso Bologna, si sono raggiunti valori di throughput maggiori, ma comunque mai superiori a 1,67 Mbps.

Nella tabella 5.2 sono riportati i valori di throughput ottenuti effettuando trasferimenti da Bologna a tutti gli altri siti; i valori sono stati rilevati negli istanti in cui i trasferimenti erano in corso contemporaneamente.

Per quanto riguarda il dato di Roma che risulta addirittura più basso degli altri nonostante l'accesso sia il più veloce, resta un problema aperto.

Sono state effettuate anche prove di ripartizione della banda di un singolo PVC effettuando la trasmissione di un file di dimensioni fisse (10240 pacchetti da 512 byte) da ogni sito a Milano; i risultati di queste prove sono nella tabella

	minimo	massimo	medio
Milano	173.912	470.592	287.609
Genova	166.664	444.448	268.920
Torino	285.712	470.592	380.792
Roma	93.024	333.336	200.976

Tabella 5.2. Throughput ottenuti da Bologna ai siti elencati espresso in bps.

5.3.

Origine	Throughput
Bologna	119,8
Genova	113,3
Torino	119,4
Roma	113,3
Aggregato	465,8

Tabella 5.3. Ripartizione della banda sull'accesso di Milano. Il throughput è espresso in kbps.

Si può osservare come la distribuzione della banda sia abbastanza uniforme e la somma dei throughput si avvicini al valore nominale dell'accesso di Milano (512 kbps).

Nelle misure di ripartizioni sugli accessi a 2 Mbps si sono avuti risultati peggiori. Nella tabella 5.4 sono riportati i risultati di un'analogo test con destinazione Bologna.

In questo caso la distribuzione della banda non è propriamente uniforme, con differenze anche di 131 kbps (circa il 36 % rispetto al throughput medio). Inoltre la somma dei throughput rimane molto al di sotto del valore nominale dell'accesso (2 Mbps).

In definitiva i router e la C-LAN hanno dimostrato la loro funzionalità con

Origine	Throughput
Milano	309
Genova	436
Torino	333
Roma	440
Aggregato	1518

Tabella 5.4. Ripartizione della banda sull'accesso di Bologna. Il throughput è espresso in kbps.

un buon funzionamento generale ma con performance migliori sugli accessi a 512 kbps che su quelli a 2 Mbps, evidenziando un probabile sottodimensionamento del servizio da parte del fornitore (Telecom).

5.1.5 Funzionalità di Frame relay switch

Durante questo test si è voluta provare la funzionalità dei router Cisco come frame relay switch, cioè la capacità di questi router di commutare i pacchetti provenienti da un PVC su di un'altra sua interfaccia, svolgendo così la funzione di un DCE frame relay. In figura 5.4 è riportato il particolare della topologia di questo test.

Il router configurato come frame relay switch è Cnaf-gw-3: ad esso è collegato il PVC da Genova che viene commutato su una sua interfaccia seriale che connette il router Niscn4.

Il test ha presentato alcune difficoltà nella impostazione della giusta configurazione delle macchine essendo di due case produttrici diverse: Cisco (cnaf-gw-3) e Digital (rt2ge5 e niscn4). Come LMI è stato settato il tipo "Cisco" su cnaf-gw-3 e il tipo "Joint" sui DECnis poiché perfettamente compatibili; come HDLC è stato configurato il CHDLC⁷ su entrambi.

La funzionalità è stata provata in quanto le macchine sono state perfettamente in grado di colloquiare insieme; niscn4 risultava direttamente connesso

⁷CHDLC: è una estensione del protocollo HDLC ideata dalla Cisco.

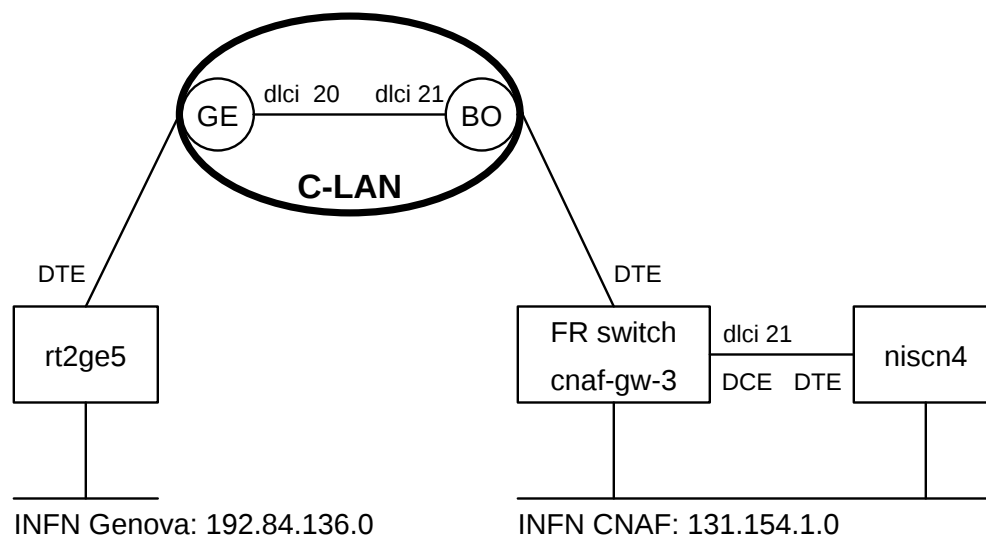


Figura 5.4. Topologia test funzionalità frame relay switch.

con rt2ge5. Non sono state notate variazioni di performance rispetto al collegamento diretto tra niscn4 e rt2ge5 che era attivo prima della prova. Anche il router Cisco non è sembrato essere stato influenzato negativamente (con cali di performance) dal suo utilizzo come frame-relay switch.

Il vantaggio di questa configurazione è che, benché in un sito l'accesso alla C-LAN sia unico (cioè quello di cnaf-gw-3), è comunque possibile configurare la macchina connessa come un frame realy switch e ad essa collegare direttamente altre macchine, le quali risultano, in pratica, direttamente affacciate alla C-LAN. Esiste quindi un vantaggio economico in quanto viene pagato un unico accesso alla C-LAN; è anche tecnicamente conveniente dato che si semplifica il routing evitando un salto (hop); inoltre, non facendo passare il traffico sulla ethernet che collega lo switch al DTE (cnaf-gw-3 a niscn4), si può aumentare l'efficienza delle comunicazioni.

5.1.6 C-LAN in produzione

Questa prova ha voluto sperimentare il routing dinamico sui router di una rete fortemente magliata, con collegamenti costituiti dai PVC della C-LAN.

I siti interessati in questa prova sono stati Milano, Bologna, Torino e Roma; in figura 5.5 è riportato il particolare della topologia di questo test.

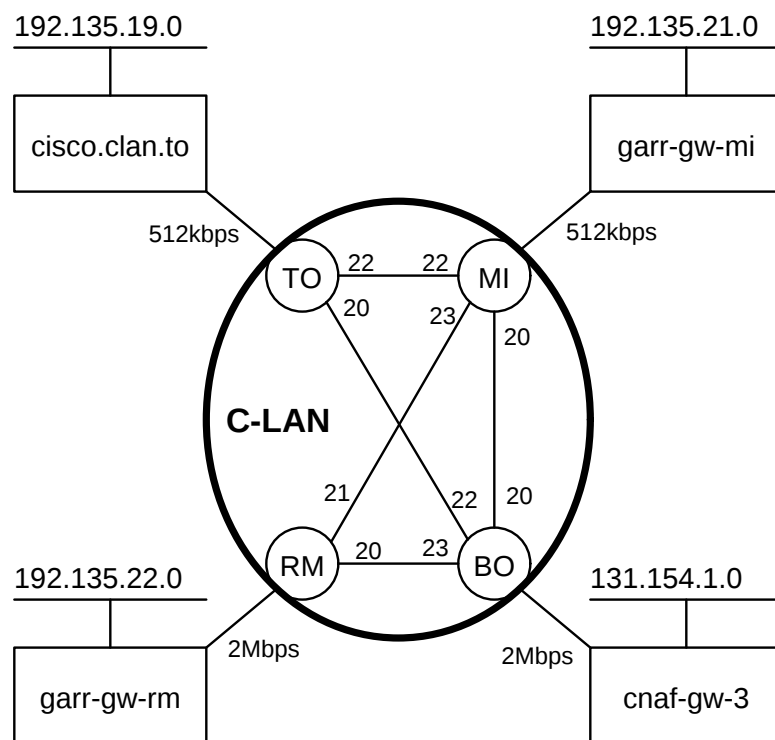


Figura 5.5. Topologia test C-LAN in produzione; i numeri a fianco dei PVC indicano i DLCI.

Configurazione del routing domain

I router coinvolti nel test sono stati: *cisco.clan.to*, *garr-gw-mi*, *cnaif-gw-3*, *garr-gw-rm*. In essi è stato configurato Integrated IS-IS per effettuare solo il routing IP (cfr. 3.3). I router erano configurati di livello 1 e 2 e costituivano un routing domain indipendente. Oltre che gli indirizzi IP, ad ogni router è stato assegnato un extended NSAP poiché necessario a integrated IS-IS per determinare le gerarchie di routing; infatti ogni NSAP aveva un diverso numero di area per cui i router effettuavano solo il routing di livello 2.

Garr-gw-mi: Milano ha dei PVC con Roma, Torino, Bologna; prima del test eseguiva i protocolli IGRP 137 per IP e ISIS area_31 per CLNS. Durante il test è stato disattivato IS-IS per CLNS e configurato integrated IS-IS solo per IP per poterne esaminare meglio il funzionamento. Tra integrated IS-IS e IGRP non è stato fatto nessun redistribute delle informazioni di routing e gli annunci integrated IS-IS sono stati fatti preferire a quelli IGRP; in tal modo la C-LAN è stata usata solo per raggiungere siti affacciati ad essa. Questa configurazione era già sufficiente a far passare il traffico di produzione sulla C-LAN poiché *garr-gw-mi* era già utilizzato come default gateway degli host del sito.

Cnaif-gw-3: Bologna ha PVC con Roma, Torino, Milano; prima del test eseguiva solo il protocollo IGRP 137 per IP. Anche qui tra integrated IS-IS e IGRP non è stato fatto nessun redistribute e gli annunci integrated IS-IS sono stati fatti preferire a quelli IGRP; in tal modo la C-LAN è stata usata solo per raggiungere siti affacciati ad essa. Per avere il traffico di produzione sulla C-LAN è stato necessario modificare il default gateway degli host del sito da *Niscn3* a *Cnaif-gw-3*.

Garr-gw-rm: Roma ha PVC con Bologna e Milano; prima del test eseguiva solo il protocollo IGRP 137 per IP. Tra integrated IS-IS e IGRP non è stato fatto nessun redistribute e gli annunci integrated IS-IS sono stati fatti preferire a quelli IGRP; in tal modo è stata usata la C-LAN solo per raggiungere siti affacciati ad essa. Questa configurazione era già sufficiente a far passare il traffico di produzione sulla C-LAN poiché *garr-gw-rm* era già utilizzato come default gateway dagli host del sito.

Cisco.clan.to: Torino ha PVC con Bologna, Milano, Genova; prima del test

era utilizzato solo con il routing statico. L'uso di integrated IS-IS ha fatto in modo che fosse utilizzata la C-LAN. Per far passare il traffico di produzione sulla C-LAN è stato necessario fare arrivare le informazioni di routing integrated IS-IS al Cisco 7000 di Torino, default gateway per le macchine del sito; per fare questo è stato avviato sui due router un processo IGRP 999, distinto da IGRP 137. Su Cisco.clan.to, integrated IS-IS è stato ridistribuito in IGRP 999; sul Cisco 7000 non è stato fatto nessun redistribute tra i due IGRP, ma gli annunci del 999 sono stati fatti preferire a quelli del 137 (figura 5.6).

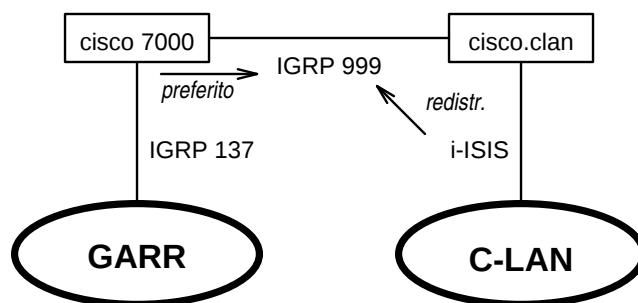


Figura 5.6. Protocolli sui router di Torino.

Routing sulla C-LAN

Con queste configurazioni il traffico di produzione da e per siti interni al routing domain attraversava sempre i PVC, evitando di passare per il resto del GARR. Invece il traffico per siti non affacciati alla C-LAN seguiva gli annunci di IGRP 137 continuando a preferire il resto della rete GARR.

Tutti gli indirizzi di rete dei siti connessi alla C-LAN sono stati annunciati anche nella rete del GARR a causa della presenza su ognuno di essi del protocollo IGRP 137. Siti esterni al routing domain, ma che facevano riferimento a router affacciati alla C-LAN, conoscevano i percorsi attraverso IGRP, ma giunti ai router del test usavano i PVC. Ad esempio, Napoli conosceva le reti del CNAF attraverso IGRP che gli diceva di passare per Roma; giunto a Roma però seguiva la strada indicata da integrated IS-IS.

Back-up sul GARR

Al fine di non isolare siti in seguito a guasti della C-LAN, è stato previsto un meccanismo di back-up delle linee sulla rete del GARR.

In seguito a mal funzionamenti della C-LAN o delle interfacce dei router che vi accedono, possono diventare irraggiungibili (attraverso la C-LAN) alcune o tutte le reti connesse ai router; in tal caso vengono a mancare anche gli annunci di integrated IS-IS di questi siti.

Queste reti però sono ancora annunciate tramite IGRP nel resto della rete GARR e tali annunci arrivano anche ai router connessi alla C-LAN; mancando integrated IS-IS, questi annunci sono ascoltati e gli indirizzi posti nella tabella di routing.

Questo meccanismo garantisce il back-up della C-LAN sulla rete GARR.

Il back-up del GARR sulla C-LAN non avviene poiché non vi è redistribute di integrated IS-IS in IGRP.

Connessioni tra C-LAN e GARR

Nella figura 5.7 è riportata una schematizzazione di come alcuni siti sono connessi alla C-LAN e al GARR. È possibile che si verifichino asimmetrie nei percorsi dei siti che raggiungono la rete GARR solo attraverso router connessi anche alla C-LAN (ad esempio Napoli).

Risultati prove

Una prima prova empirica è stata quella di disattivare dei PVC e osservare quanto tempo trascorrevano prima che le tabelle di routing venissero corrette. Questa correzione è risultata immediata: giusto il tempo di dare il comando di shutdown di una interfaccia e subito dopo il comando per visionare le tabelle di routing che risultavano già aggiornate. Ciò evidenzia come i protocolli link state rilevino velocemente variazioni nella topologia della rete ed altrettanto velocemente siano in grado di ricalcolare i percorsi ottimali; questo pregio è tanto più utile in una rete fortemente magliata come quella che può essere realizzata sulla C-LAN.

Altre prove hanno voluto misurare la funzionalità della C-LAN; a tale

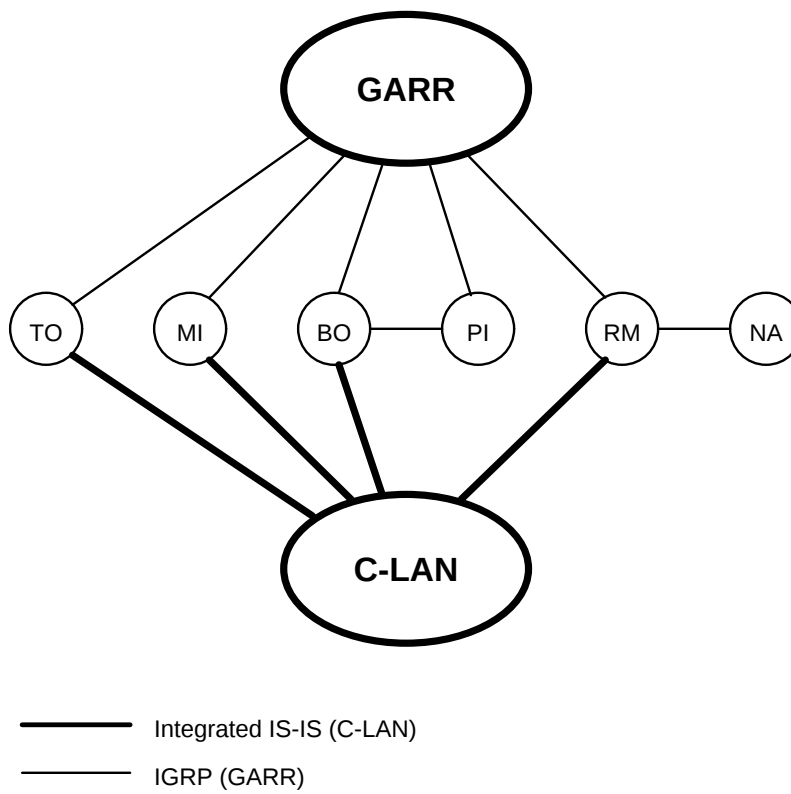


Figura 5.7. Schema dei collegamenti di alcuni siti alla C-LAN e al GARR.

proposito sono state effettuate misure del round-trip dei ping⁸ e misure di throughput.

Nella tabella 5.5 sono riportati i valori medi di ping eseguiti da ogni router ad ogni altro.

da \ a	Milano	Torino	Bologna	Roma
Milano	-	22	21	24
Torino	24	-	28	52
Bologna	21	28	-	28
Roma	24	44	24	-

Tabella 5.5. Round trip medio dei ping tra i vari router espresso in ms.

Il valore quasi doppio del round trip dei ping tra Torino e Roma è dovuto al fatto che vengono attraversati un PVC ed un router in più rispetto agli altri.

Per le misure del throughput è stato effettuato un trasferimento di file tra un host a Bologna e uno a Roma; sull'interfaccia del router cnaf-gw-3 è stato rilevato un throughput massimo di 1250 kbps.

Usando anche un host di Torino ed effettuando trasferimenti tra Bologna e Roma e tra Bologna e Torino contemporaneamente, è stato rilevato un throughput massimo di 1400 kbps sull'interfaccia di cnaf-gw-3.

Usando due macchine a Bologna e due a Roma, sull'interfaccia di cnaf-gw-3 è stato rilevato un throughput di 1400 kbps.

Effettuando, in contemporanea, due trasferimenti tra Bologna e Roma, uno tra Bologna e Torino e uno tra Bologna e Milano, sull'interfaccia di cnaf-gw-3 è stato rilevato un throughput di 1630 kbps.

Anche in questo caso si è rilevato l'impossibilità di raggiungere il throughput nominale dell'accesso di Bologna (2Mbps) confermando quanto misurato nella prova di funzionalità.

⁸Ping: speciali pacchetti inviati da un host ad un altro, il quale deve rispedirli immediatamente al mittente; tramite essi è possibile misurare il tempo di andata e ritorno e la percentuale di pacchetti persi nel tragitto.

Configurazioni

Le configurazioni dei router sono riportate nell'appendice A.

5.1.7 Conclusioni

Dalle prove effettuate si possono trarre diverse conclusioni.

Il *management* è risultato abbastanza carente sia dal punto di vista degli switch della Telecom che dell'interrogazione dei router Cisco, in quanto i primi sono inaccessibili ed i secondi hanno fornito risposte non attendibili. Sono comunque prevedibili ampi miglioramenti in funzione dell'esperienza acquisita e dell'evoluzione dei prodotti.

Le *performance* sono accettabili per gli accessi a 512 kbps; sono invece difficilmente interpretabili i valori riscontrati sugli accessi a 2 Mbps dove è complicato raggiungere i valori massimi del throughput che rimane comunque molto inferiore ai 2 Mbps. Se l'infrastruttura interna della C-LAN non verrà potenziata è facile prevedere che qualora venisse usata anche da altre realtà si avrebbe un ulteriore calo delle performance.

La *stabilità* dei PVC è risultata abbastanza soddisfacente, anche se alcune volte sono state riscontrate interruzioni di alcuni circuiti anche per diverse ore.

Per quanto riguarda il *routing*, integrated IS-IS ha dimostrato rapidità di convergenza e semplicità di configurazione; in topologie così fortemente magliate è necessario adottare un protocollo di routing link state.

La funzionalità di *frame relay switching* del software dei router Cisco è risultato perfettamente funzionante ed in grado di fornire un utilizzo locale molto flessibile dei PVC disponibili.

5.2 Maglia OSPF tra Toscana e Bologna

In questo test si è voluto provare la funzionalità di un backbone OSPF, realizzato effettuando collegamenti diretti e virtuali fra router di produttori diversi, in particolare Cisco e Digital.

5.2.1 Obiettivi del test

Attualmente, sulla rete GARR, OSPF è utilizzato solamente in Toscana per effettuare il routing tra il CNR di Pisa, l'università di Pisa e di Firenze, l'INFN di Pisa a San Piero a Grado e il CNR di Sassari; i primi tre appartengono ad un'area, Sassari ad un'altra. Per effettuare questa prova è stata realizzata una terza area OSPF che comprende alcuni router del CNAF di Bologna e sono stati poi creati i collegamenti di backbone tra questa e l'area toscana.

Obiettivo era la verifica delle effettive possibilità di realizzazione di questa topologia essendo presenti router di marche diverse che in passato avevano già avuto difficoltà a connettersi insieme; queste difficoltà erano dovute in genere al non perfetto sviluppo del software dei router; si è rilevato che questi difetti permangono tutt'ora e proprio per questo motivo il test non è stato portato a termine. Una volta configurato il routing domain OPSF, si sarebbero dovute effettuare misure di funzionalità del protocollo.

5.2.2 Topologia

Il test è stato provato sulla topologia riportata in figura 5.8: garr-gw-bo.infn.it e niscn4.infn.it appartengono all'area bolognese denominata 131.154.0.0 mentre garr-gw.cnr.it, cis1pi.pi.infn.it e faeta.unipi.it all'area toscana 131.114.0.0.

Garr-gw-bo.infn.it, niscn4.infn.it, garr-gw.cnr.it e cis1pi.pi.infn.it sono anche backbone router con un virtual link (OSPF) tra garr-gw-bo.infn.it e niscn4.infn.it e uno tra garr-gw.cnr.it e cis1pi.pi.infn.it

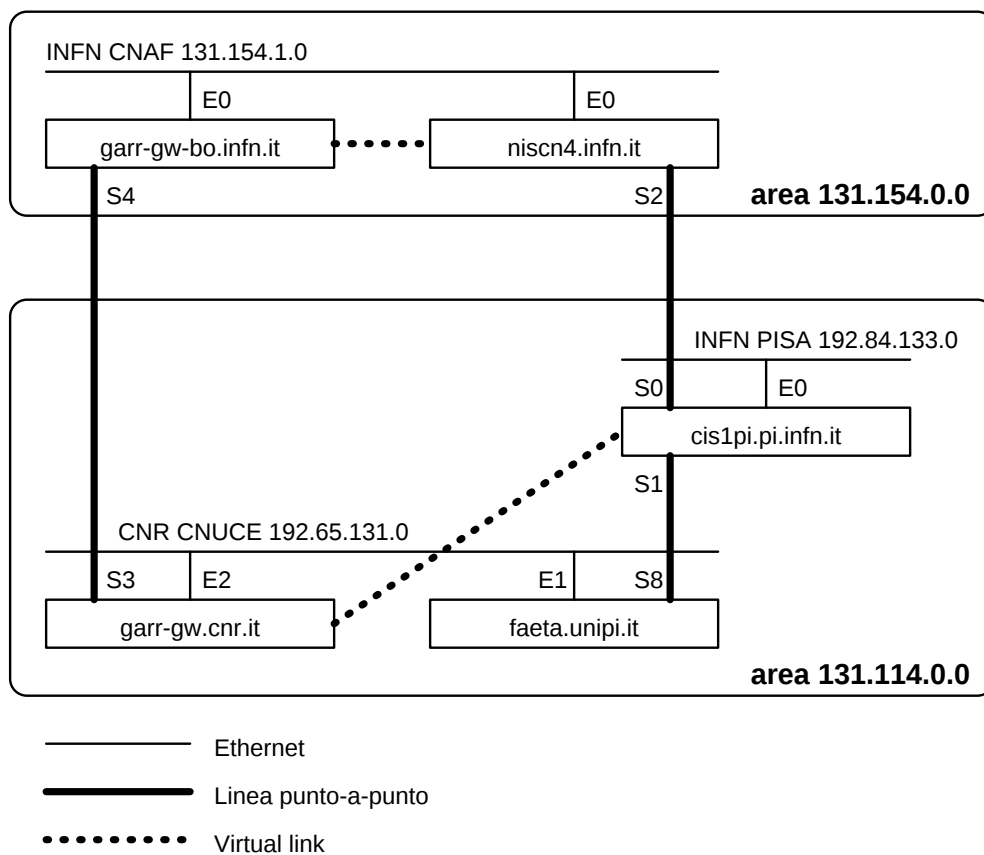


Figura 5.8. Topologia test della maglia OSPF.

5.2.3 Hardware e software impiegato

L'hardware e il software dei router coinvolti nei test è stato il seguente:

- Garr-gw-bo.infn.it è un router Cisco AGS+ con 16 MB di memoria e software GS versione 9.21(5430), situato presso il CNAF/INFN di Bologna.
- Niscn4.infn.it è un router Digital Decnis600 con software Decnis versione 2.3.2, situato presso il CNAF/INFN di Bologna.
- Garr-gw.cnr.it è un router Cisco AGS+ con 16 MB di memoria e software GS versione 9.21(5430), situato presso il Cnuce (CNR) di Pisa.
- Faeta.unipi.it è un router Cisco AGS+ con 16 MB di memoria e software GS versione 9.1(11.9), situato presso l'università di Pisa.
- Cislpi.pi.infn.it è un router Cisco 4000 con 4 MB di memoria più 2 MB di flash eprom e software 4000 versione 9.1(11.9), situato presso la sezione INFN di Pisa a San Piero a Grado.

5.2.4 Misure di efficienza

Oltre a verificare la funzionalità della configurazione proposta si volevano effettuare alcune misure per valutarne anche l'efficacia; per fare questo avrebbero dovuto essere utilizzate alcune procedure scritte a proposito.

- Tramite DECMcc è possibile chiedere informazioni ai router Cisco sulla loro memoria libera e sul carico della loro CPU (istantaneo e mediato su un certo intervallo di tempo); è stata preparata una procedura per interrogare garr-gw-bo.infn.it, garr-gw.cnr.it, cislpi.pi.infn.it e faeta.unipi.it ogni dieci minuti sulla memoria libera, sull'occupazione istantanea della CPU e su quella mediata su uno e cinque minuti.

- È possibile chiedere ai router Cisco il contenuto delle loro tabelle di routing, delle dimensioni dei database, dei contatori; è stata scritta una procedura che interroga automaticamente i quattro Cisco sul loro stato ogni quindici minuti.

Queste due misure sono utili per valutare le reazioni di quei router che, come garr-gw.cnr.it, eseguono diversi protocolli di routing contemporaneamente.

- Tramite il trasferimento di file con FTP, è possibile avere una misura del throughput; è stata scritta una procedura per effettuare un trasferimento

di un file da cnaf43.infn.it (macchina presso il CNAF) a vspcid.pi.infn.it (macchina di S. Piero a Grado) e viceversa ogni venti minuti.

- Tramite Ping è possibile osservare la percentuale di pacchetti persi ed il tempo minimo, medio e massimo necessario ad un pacchetto per andare fino ad un host prefissato e tornare indietro; è stata scritta una procedura che invia cento ping da cnaf43.infn.it a vspcid.pi.infn.it ogni venti minuti.

- Tramite la funzione di traceroute è possibile osservare i router attraversati dai pacchetti per raggiungere la destinazione; è stata scritta una procedura che effettua dei trace da cnaf43.infn.it e da vspcid.pi.infn.it a tutti e cinque i router del test.

Conoscendo la velocità di trasmissione nominale delle linee interessate e queste misure è possibile stabilire se sono usate al meglio.

5.2.5 Svolgimento delle prove

La prova non ha avuto esito positivo e non è stato possibile portarla a termine a causa di una incompatibilità tra il software Cisco e quello Digital unitamente ad un problema del router Digital. Questo ha impedito di completare la configurazione e di effettuare le misure di efficienza e funzionalità. Qui di seguito è comunque descritto dettagliatamente quanto accaduto.

Martedì 11 ottobre 1994.

La configurazione di OSPF sui router toscani era già attiva e collaudata. Garr-gw.cnr.it era già configurato per il collegamento con garr-gw-bo; l'unica modifica da effettuare sarebbe stata la configurazione del collegamento di backbone tra cislpi.pi.infn.it e niscn4.infn.it.

Il primo passo è stato dunque quello di attivare OSPF sui router del CNAF: su garr-gw-bo e niscn4 è stata definita un'area denominata 131.154.-0.0 contenente la sottorete 131.154.1.0. Sui due router è stata poi avviata l'area di backbone attivando un virtual link tra di loro. Tramite il debug dei router Cisco e le rilevazioni dei messaggi inviati dai router Digital è stato possibile controllare eventuali malfunzionamenti. Fino a questo punto non vi sono stati problemi ed il funzionamento dei due router è stato quello aspettato. È da rilevare che in luglio era già stata effettuata questa prova ma si era

verificato un malfunzionamento del virtual link in quanto niscn4 comunicava di ricevere pacchetti malformati (bad packet) su di esso. Nel frattempo il software del DECnis era stato aggiornato dalla versione 2.2.2 alla versione 2.3.2. Si è dedotto che la nuova versione doveva avere risolto il problema.

Attestato il buon funzionamento, si è avviata la connessione di backbone tra garr-gw-bo.infn.it e garr-gw.cnr.it. Da quest'ultimo sono fluite tutte le numerose informazioni sulle network della Toscana e di tutto il dominio IGRP (poiché al CNR viene effettuato un redistribute di IGRP in OSPF). Questo ha fatto disabilitare OSPF da niscn4 poiché i suoi database erano sottodimensionati (prevedeva al massimo 5 destinazioni esterne al dominio OSPF mentre ne sono giunte circa 70).

Questo fatto ha avuto però una grave e misteriosa conseguenza: niscn1.infn.it (altro router connesso all'ethernet del CNAF con indirizzo di rete 131.154.1.0) ha iniziato a inviare messaggi in cui affermava ripetutamente di non volere più fare il designated router per integrated IS-IS per poi subito smentirsi. Oltre a questo, niscn1 non era più in grado di instradare pacchetti per il CERN, il Fermilab e Trieste. L'INFN di Bologna era invece raggiungibile discontinuamente.

Per risolvere la situazione, su garr-gw-bo e niscn4 è stato disabilitato OSPF ed è stato necessario riavviare completamente niscn1.

Mercoledì 12 ottobre 1994.

La prima operazione è stata quella di ridimensionare come segue i database di niscn4:

IP local adjacencies:	50
Number of IP reachable destinations:	50
Number of IP Level 1 destinations:	100
IP area connectivity:	50
Number of IP Level 2 destinations:	1000
IP domain connectivity:	50
Number of external destinations:	2000
OSPF maximum connected areas:	3
OSPF average connected routers:	10
OSPF maximum area interfaces:	5
OSPF average area networks:	100

OSPF maximum network routers:	10
OSPF maximum system networks:	500
OSPF maximum boundary routers:	5
OSPF maximum external routes:	1000
OSPF average external connectivity:	5
OSPF maximum destinations:	2000
OSPF maximum adjacencies:	25

È stato poi riavviato OSPF tra garr-gw-bo.infn.it e niscn4 limitatamente al CNAF senza problemi.

Appena è stato collegato garr-gw-bo con garr-gw.cnr.it, niscn4 ha disabilitato OSPF indicando un overflow di ASBR⁹ link state packet, esattamente come il giorno precedente; questa volta però senza conseguenze disastrose per niscn1.

OSPF è stato disabilitato anche da garr-gw-bo e ed è stato ridimensionato il numero di *OSPF maximum boundary routers* di niscn4 da 5 a 70.

È stato riabilitato OSPF al CNAF senza problemi e poi avviato il collegamento con la Toscana. Questa volta niscn4 è riuscito a gestire tutti gli annunci ed è rimasto attivo. Dopo pochi secondi però, è sorto un'altro problema: garr-gw-bo e niscn4 non si sono più intesi sul virtual link e mentre garr-gw-bo ritrasmetteva un LSP, niscn4 lo scartava come malformato richiedendone la ritrasmissione, esattamente come avveniva in luglio.

Imputando il malfunzionamento ad un baco nel software di uno dei due router di Bologna, si è deciso di effettuare una prova senza il virtual link; per fare questo è stato necessario disabilitare l'area 131.154.0.0 e configurare la subnetwork 131.154.1.0 in area di backbone. Effettuata la modifica e ricollegato garr-gw-bo con garr-gw.cnr.it il funzionamento è stato perfetto. È stato notato che su niscn4 il processo OSPF aveva cambiato in falso il settaggio di area border router; la cosa era corretta poiché, con la configurazione presente, niscn4 risultava essere solo backbone router.

È stato fatto un'altro tentativo e si è riconfigurata l'area 131.154.0.0 senza i collegamenti con la Toscana. Appena è stato configurato il virtual link tra niscn4 e garr-gw-bo, niscn4 ha avvertito che stava rilevando alcuni errori nella configurazione dell'area di backbone; nel frattempo garr-gw-bo aveva

⁹ASBR: Autonomous System Boundary Router.

stabilito l'adiacenza con niscn4 ma il virtual link rimaneva in fase di inizializzazione. In questa situazione è stato notato che su niscn4 il settaggio area border router permaneva falso: in questo caso il settaggio era errato poiché niscn4 era configurato appartenere all'area 0 e all'area di backbone; il malfunzionamento è stato imputato a questo fatto. Purtroppo non esiste un comando di linea per imporre al router di essere area border router, ma è necessario riavviare il router per ripristinare la situazione esatta (normalmente il passaggio da area border router a internal router e viceversa dovrebbe essere automatico).

Dopo un ulteriore tentativo con il medesimo esito si è deciso di interrompere la prova.

Configurazioni

Le configurazioni dei router sono riportate nell'appendice B.

5.2.6 Conclusioni

Non è stato possibile insistere oltre nell'effettuare tentativi di configurazioni poiché le prove hanno interessato linee e router di fondamentale importanza per il routing nazionale ed internazionale e quanto avvenuto ha avuto gravi conseguenze, essendo state più volte interrotte le comunicazioni.

Benché non portata a termine, questa prova è comunque utile per stabilire quali configurazioni siano possibili nella realtà. Allo stato attuale dello sviluppo del software non è possibile effettuare un backbone OSPF con virtual link tra coppie di router Cisco-Digital. Il problema è stato comunicato agli sviluppatori di software di queste aziende e forse in un futuro più o meno prossimo verrà risolto. Resta comunque il fatto che la realizzazione di reti con apparecchiature "multivendor" comporta sempre maggiori problemi di funzionamento rispetto a quelle con apparecchiature "monovendor", oltre che per il loro interfacciamento, anche per il raddoppio delle normali operazioni di gestione come l'aggiornamento delle versioni software.

Questa situazione è abbastanza disdicevole in quanto i produttori di software dicono di attenersi agli standard, ma in realtà non lo fanno: un po'

per limitare i costi, un po' perché molte funzionalità non sono utilizzate da nessuno. Quando però si scoprono queste inadempienze (come nel caso di questo test) gli effetti negativi si ripercuotono sugli utilizzatori finali, rendendo difficile realizzare in pratica quanto progettato sulla base di documenti universalmente approvati.

5.2.7 Ultimi sviluppi

Con la versione 2.3.6 del software per i router DECnis è stato risolto il problema del settaggio *area border router*: in novembre è stata provata la configurazione limitata all'area di Bologna e niscn4 è risultato in grado di adattare automaticamente questo parametro a seconda della presenza o meno dell'area di backbone.

La Digital ha anche individuato il problema relativo al bad packet sul virtual link: sembra che il pacchetto in questione sia un AS external link LSA inviato dal Cisco sul virtual link, la qual cosa è espressamente proibita in [RFC 1583] (specifiche di OSPF). Ora è stato chiesto alla Cisco di risolvere il problema.

5.3 Confronto tra Ship In the Night e Integrated routing

I protocolli del network layer maggiormente usati sulla rete INFNet sono l'IP e il CLNP: ognuno di questi necessita di propri protocolli di routing. Come già scritto nella sezione 3.4, questi possono convivere indipendenti l'uno dall'altro nella modalità nota come Ship In the Night (S.I.N.) oppure si può utilizzarne uno unico con doppia funzionalità (integrated).

In un test previsto si volevano provare questi due tipi di convivenze usando OSPF, IS-IS e integrated IS-IS nella maglia della prova riportata nella sezione 5.2, per osservare come la banda occupata dai pacchetti generati da ogni protocollo influenzasse le performance della rete. Purtroppo non è stato possibile configurare OSPF e quindi questa prova non è risultata fattibile.

In alternativa si è eseguito un calcolo sulla base delle specifiche dei tre protocolli per determinare l'occupazione di banda teorica nella configurazione del test.

5.3.1 Obiettivi del test

Obiettivo di questa prova è stato il confronto tra l'occupazione di banda dei pacchetti generati dai protocolli OSPF e IS-IS in convivenza S.I.N., con quella dei pacchetti generati da integrated IS-IS, per determinare quale delle due configurazioni utilizza meno risorse della rete.

5.3.2 Topologia

Per calcolare la dimensione dei pacchetti è stato necessario definire una topologia della rete ed alcune grandezze.

La topologia presa in esame ha ricalcato quella reale che è riportata in figura 5.8; nei calcoli la si è considerata come un routing domain isolato (cioè che non confina con nessun altro routing domain) in cui sono definite due aree (area BO e area PI) che contengono solo le sottoreti connesse ai router e che sono riportate nella figura. Per la configurazione IP si è assunto che su ogni interfaccia fosse definito un indirizzo, cioè che sia le reti broadcast

che quelle punto a punto fossero definite come subnetwork, ognuna distinta dalle altre. Per la configurazione CLNP si è assunto che su ogni router fosse definito un NSAP address di venti ottetti, tredici per l'area, sei per l'ID, uno per il Selector (cfr. figura 2.3).

In figura 5.9 è riportata la configurazione per OSPF: garr-gw-bo e niscn4 appartengono sia all'area BO che al backbone; tra i due è quindi definito un virtual link; si assume che garr-gw-bo sia eletto designated router per la subnetwork 131.154.1.0. Garr-gw.cnr.it, faeta.unipi.it e cislpi.pi.infn.it appartengono invece all'area PI e garr-gw.cnr e cislpi anche al backbone; anche tra loro è definito un virtual link. Si assume che garr-gw.cnr sia eletto designated router per la subnetwork 192.65.131.0.

In figura 5.10 è riportata la configurazione per IS-IS e integrated IS-IS: garr-gw-bo e niscn4 appartengono all'area BO e sono di livello 1 e 2. Si assume che garr-gw-bo sia eletto designated router per la rete broadcast del CNAF. Garr-gw.cnr.it, faeta.unipi.it e cislpi.pi.infn.it appartengono invece all'area PI e sono tutti di livello 1 e 2. Si assume che garr-gw.cnr sia eletto designated router per la rete broadcast del CNUCE.

I pacchetti di Hello sono ritrasmessi ad intervalli regolari lunghi un numero di secondi definito dalla variabile *HelloInterval* per OSPF e *HelloTimer* per IS-IS e integrated IS-IS. Gli LSP invece, in condizioni di stabilità dei collegamenti, sono generati ad intervalli regolari lunghi un numero di secondi indicati dalla variabile *LSPRefreshTimer* per OSPF e *LSPGenerationInterval* per IS-IS e integrated IS-IS.

5.3.3 Dimensioni dei pacchetti OSPF

Facendo riferimento al paragrafo 3.2.4, si è calcolato la dimensione dei pacchetti. È stata considerata la configurazione in una situazione stabile, per cui non sono stati presi in esame i pacchetti che i router si scambiano inizialmente per stabilire le connessioni.

Hello Packet

Gli hello packet che un router genera sono utilizzati per confermare la propria esistenza; non esistono differenze tra quelli usati all'interno delle aree e quelli

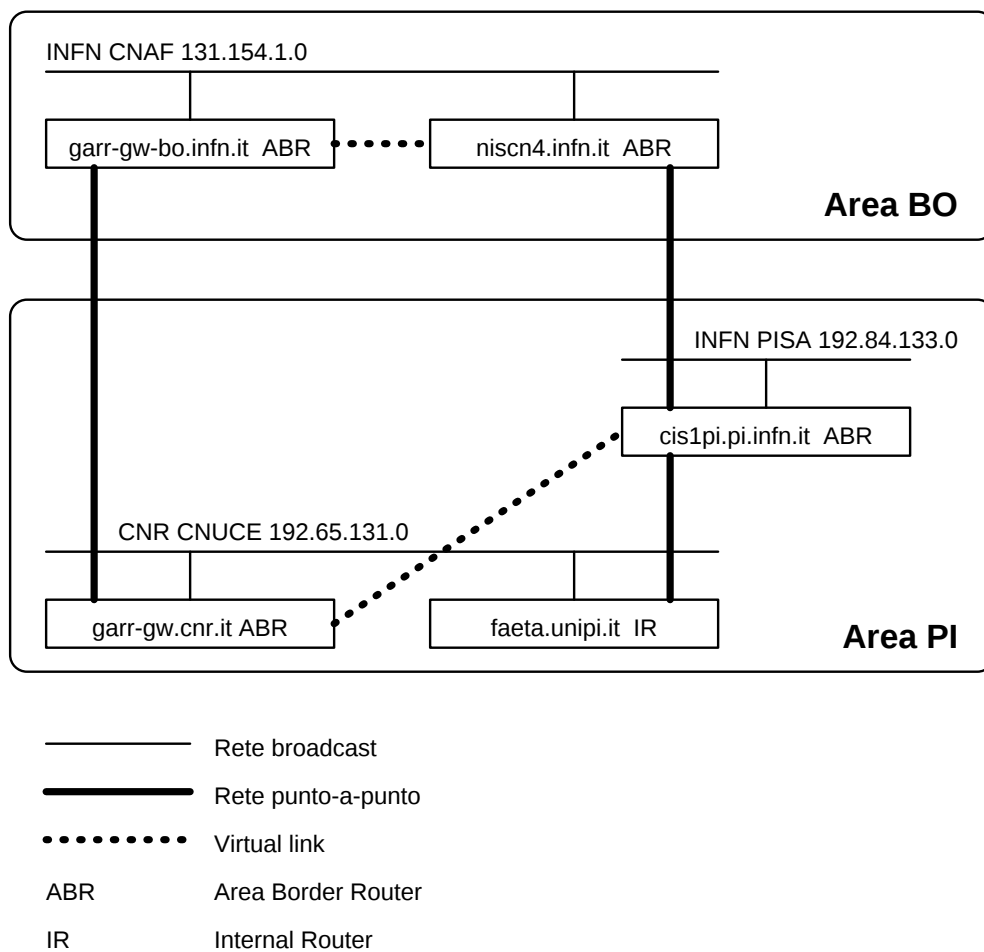


Figura 5.9. Configurazione OSPF.

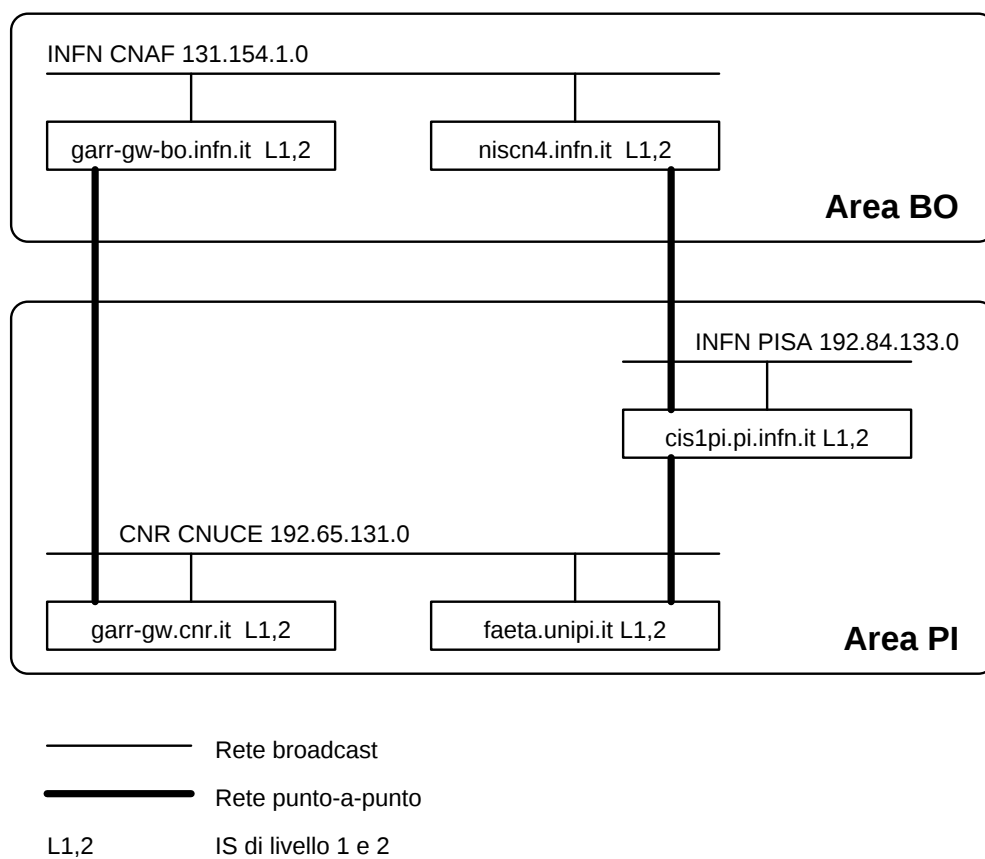


Figura 5.10. Configurazione IS-IS e integrated IS-IS.

del backbone.

- Garr-gw-bo.infn.it genera tre Hello packet (uno per il collegamento punto a punto, uno per quello broadcast e uno per il virtual link) di 56 ottetti (44 per l'header più 4 per ogni vicino), per un totale di 166 ottetti.
- Niscn4.infn.it genera gli stessi pacchetti e della stessa dimensione di quelli di garr-gw-bo, per un totale di 166 ottetti.
- Garr-gw.cnr.it genera gli stessi pacchetti e della stessa dimensione di quelli di garr-gw-bo, per un totale di 166 ottetti.
- Cis1pi.pi.infn.it genera gli stessi pacchetti e della stessa dimensione di quelli di garr-gw-bo, per un totale di 166 ottetti.
- Faeta.unipi.it genera due Hello packet (uno per il collegamento punto a punto e uno per quello broadcast) di 52 ottetti (44 per l'header più 4 per ogni vicino) per un totale di 104 ottetti.

In totale i cinque router immettono nel routing domain 760 ottetti ogni HelloInterval secondi per la gestione degli Hello.

Area BO

All'interno dell'area BO, i router si scambiano i Link State Update packet (LSU) per informarsi dello stato dei collegamenti; la ricezione di questi è confermata tramite l'invio di un Link State Acknowledge packet (LSA).

- Garr-gw-bo.infn.it genera un LSU di 144 ottetti (28 per l'header, 36 per il RLA¹⁰, 32 per il NLA¹¹ in quanto designated router, 48 per il SLA¹² in quanto area border router che conosce le due subnetwork dell'area PI). Genera anche un LSA di 64 ottetti per confermare la ricezione dei due advertisement di niscn4.

- Niscn4.infn.it genera un LSU di 112 ottetti (28 per l'header, 36 per il RLA, 48 per il SLA in quanto area border router che conosce le due subnetwork dell'area PI). Genera anche un LSA di 84 ottetti per confermare la ricezione dei tre advertisement di garr-gw-bo.

¹⁰RLA = Router Link Advertisement.

¹¹NLA = Network Link Advertisement.

¹²SLA = Summary Link Advertisement.

In totale, i due router immettono nell'area BO 404 ottetti ogni LSRefreshTimer secondi per la gestione dei Link State packet.

Area PI

All'interno dell'area PI, i router si scambiano i Link State Update packet e Link State Acknowledge packet.

- Garr-gw.cnr.it genera un LSU di 120 ottetti (28 per l'header, 36 per il RLA, 32 per il NLA in quanto designated router, 24 per il SLA in quanto area border router che conosce la subnetwork dell'area BO); questo LSU raggiunge Faeta.unipi che lo rispedisce a cis1pi. Genera anche un LSA di 64 ottetti per confermare la ricezione dei due advertisement di cis1pi e un LSA di 44 ottetti per confermare la ricezione dell'advertisement di faeta.unipi.
- Faeta.unipi.it genera due LSU di 76 ottetti (28 per l'header, 48 per il RLA), uno per garr-gw.cnr e uno per cis1pi. Genera anche un LSA di 64 ottetti per confermare la ricezione dei due advertisement di cis1pi e un LSA di 84 ottetti per confermare la ricezione dell'advertisement di garr-gw.cnr.
- Cis1pi.pi.infn.it genera un LSU di 88 ottetti (28 per l'header, 36 per il RLA, 24 per il SLA in quanto area border router che conosce la subnetwork dell'area BO); questo LSU raggiunge Faeta.unipi che lo rispedisce a garr-gw.cnr. Genera anche un LSA di 84 ottetti per confermare la ricezione dei tre advertisement di garr-gw.cnr e un LSA di 44 ottetti per confermare la ricezione dell'advertisement di faeta.unipi.

In totale i tre router immettono nell'area PI 952 ottetti ogni LSRefreshTimer secondi per la gestione dei Link State packet.

Backbone

All'interno del backbone, i router si scambiano i Link State Update packet e Link State Acknowledge packet come in una normale area.

- Garr-gw-bo.infn.it genera un LSU di 108 ottetti (28 per l'header, 48 per il RLA, 32 per il NLA in quanto designated router); questo LSU raggiunge tutti gli altri tre backbone router. Genera anche tre LSA di 44 ottetti per confermare la ricezione degli advertisement degli altri tre router (ne invia uno per ciascuno).

- Niscn4.infn.it genera un LSU di 76 ottetti (28 per l'header, 48 per il RLA); questo LSU raggiunge tutti gli altri tre backbone router. Esso genera anche due LSA di 44 ottetti per confermare la ricezione degli advertisement di cis1pi e garr-gw.cnr e uno di 64 ottetti per confermare la ricezione dell'advertisement di garr-gw-bo.
- Garr-gw.cnr.it genera un LSU di 76 ottetti (28 per l'header, 48 per il RLA); questo LSU raggiunge tutti gli altri tre backbone router. Esso genera anche due LSA di 44 ottetti per confermare la ricezione degli advertisement di cis1pi e niscn4 e uno di 64 ottetti per confermare la ricezione dell'advertisement di garr-gw-bo.
- Cis1pi.pi.infn.it genera un LSU di 76 ottetti (28 per l'header, 48 per il RLA); questo LSU raggiunge tutti gli altri tre backbone router. Esso genera anche due LSA di 44 ottetti per confermare la ricezione degli advertisement di garr-gw.cnr e niscn4 e uno di 64 ottetti per confermare la ricezione dell'advertisement di garr-gw-bo.

In totale i quattro router immettono nel backbone 1596 ottetti ogni LSRefreshTimer secondi per la gestione dei Link State packet.

5.3.4 Dimensioni dei pacchetti IS-IS

Facendo riferimento al paragrafo 3.3.4 si è calcolato la dimensione dei pacchetti. È stata considerata la configurazione in una situazione stabile.

Hello Packet

Gli hello packet che un router genera sono utilizzati per confermare la propria esistenza; in IS-IS, gli hello di livello 1 sono distinti da quelli di livello 2.

- Garr-gw-bo.infn.it genera un L1-LAN-Hello¹³ sulla rete broadcast per il livello 1 di 52 ottetti; un L2-LAN-Hello¹⁴ sulla rete broadcast per il livello 2 di 52 ottetti; un PtP-Hello¹⁵ sulla rete punto-a-punto di 45 ottetti. In totale genera 149 ottetti.

¹³L1-LAN-Hello = Level 1 LAN IS to IS Hello Packet Data Unit.

¹⁴L2-LAN-Hello = Level 2 LAN IS to IS Hello Packet Data Unit.

¹⁵PtP-Hello = Point-to-Point IS to IS Hello Packet data Unit.

- Niscn4.infn.it genera gli stessi pacchetti e della stessa dimensione di quelli di garr-gw-bo, per un totale di 149 ottetti.
- Garr-gw.cnr.it genera gli stessi pacchetti e della stessa dimensione di quelli di garr-gw-bo, per un totale di 149 ottetti.
- Faeta.unipi.it genera gli stessi pacchetti e della stessa dimensione di quelli di garr-gw-bo, per un totale di 149 ottetti.
- Cislpi.pi.infn.it genera un L1-LAN-Hello sulla rete broadcast per il livello 1 di 52 ottetti; un L2-LAN-Hello sulla rete broadcast per il livello 2 di 52 ottetti; un PtP-Hello di 51 ottetti per ciascuna rete punto-a-punto. In totale genera 206 ottetti.

In totale, i cinque router immettono nel routing domain 802 ottetti ogni HelloTimer secondi per la gestione degli Hello.

Area BO

All'interno dell'area BO (di livello 1), i router si scambiano i Level 1 link state Packet Data Unit (L1-LSP) per informarsi dello stato dei collegamenti; la ricezione di questi è confermata tramite l'invio dei Level 1 partial sequence number Packet Data Unit (L1-PSN) sulle linee punto-a-punto o con dei Level 1 complete sequence number Packet Data Unit (L1-CSN) da parte dei designated router sulle reti broadcast.

- Garr-gw-bo.infn.it genera un L1-LSP di 61 ottetti (50 per l'header e 11 per indicare il vicino) e un L1-CSN di 49 ottetti (33 per l'header e 16 per l'unico L1-LSP ricevuto) in quanto designated router.
- Niscn4.infn.it genera solo un L1-LSP di 61 ottetti (50 per l'header e 11 per indicare il vicino).

In totale i due router immettono nell'area BO 171 ottetti ogni LSPGenerationInterval secondi per la gestione dei link state packet di livello 1.

Area PI

All'interno dell'area PI, i router si scambiano i Level 1 link state Packet Data Unit, i Level 1 partial sequence number Packet Data Unit ed i Level 1 complete sequence number Packet Data Unit.

- Garr-gw.cnr.it genera un L1-LSP di 61 ottetti (50 per l'header e 11 per indicare il vicino) che raggiunge Faeta.unipi il quale lo rispedisce a cis1pi; genera anche un L1-CSN di 65 ottetti (33 per l'header e 32 per i due L1-LSP ricevuti) in quanto designated router.
- Faeta.unipi.it genera un L1-LSP di 72 ottetti (50 per l'header e 22 per indicare il vicino) che raggiunge garr-gw.cnr e cis1pi; genera anche un L1-PSN di 33 ottetti (17 per l'header e 16 per il L1-LSP ricevuto da cis1p1).
- Cis1pi.pi.infn.it genera un L1-LSP di 61 ottetti (50 per l'header e 11 per indicare il vicino) che raggiunge faeta il quale lo rispedisce a garr-gw.cnr; genera anche un L1-PSN di 49 ottetti (17 per l'header e 32 per i L1-LSP ricevuti).

In totale, i tre router immettono nell'area PI 535 ottetti ogni LSPGenerationInterval secondi per la gestione dei link state packet di livello 1.

Level 2 subdomain

Per effettuare il routing di livello 2, i router si scambiano i Level 2 link state Packet Data Unit (L2-LSP) per informarsi dello stato dei collegamenti; la ricezione di questi è confermata tramite l'invio dei Level 2 partial sequence number Packet Data Unit (L2-PSN) sulle linee punto-a-punto o con dei Level 2 complete sequence number Packet Data Unit (L2-CSN) da parte dei designated router.

- Garr-gw.bo.infn.it genera un L2-LSP di 94 ottetti (58 per l'header e 36 per indicare i due vicini) che raggiunge tutti gli altri quattro router; genera un L2-CSN di 97 ottetti (33 per l'header e 64 per i quattro L2-LSP ricevuti) in quanto designated router; genera anche un L2-PSN di al più 81 ottetti (17 per l'header e 64 per i quattro L2-LSP ricevuti).
- Niscn4.infn.it genera un L2-LSP di 94 ottetti (58 per l'header e 36 per indicare i due vicini) che raggiunge tutti gli altri quattro router; genera anche un L2-PSN di 81 ottetti (17 per l'header e 64 per i quattro L2-LSP ricevuti).
- Garr-gw.cnr.it genera un L2-LSP di 94 ottetti (58 per l'header e 36 per indicare i due vicini) che raggiunge tutti gli altri quattro router; genera un L2-CSN di 97 ottetti (33 per l'header e 64 per i quattro L2-LSP ricevuti) in quanto designated router; genera anche un L2-PSN di al più 81 ottetti (17

per l'header e 64 per i quattro L2-LSP ricevuti).

- Faeta.unipi.it genera un L2-LSP di 94 ottetti (58 per l'header e 36 per indicare i due vicini) che raggiunge tutti gli altri quattro router; genera anche un L2-PSN di 81 ottetti (17 per l'header e 64 per i quattro L2-LSP ricevuti).
- Cis1pi.infn.it genera un L2-LSP di 94 ottetti (58 per l'header e 36 per indicare i due vicini) che raggiunge tutti gli altri quattro router; genera anche due L2-PSN di 81 ottetti (17 per l'header e 64 per i quattro L2-LSP ricevuti), uno per ogni linea punto-a-punto.

In totale, i cinque router immettono nel level 2 subdomain 2560 ottetti ogni LSPGenerationInterval secondi per la gestione dei link state packet di livello 2.

5.3.5 Dimensioni dei pacchetti integrated IS-IS

Facendo riferimento al paragrafo 3.3.4 si è calcolato la dimensione dei pacchetti. È stata considerata la configurazione in una situazione stabile.

Hello Packet

Gli hello packet che un router genera sono utilizzati per confermare la propria esistenza; in IS-IS, gli hello di livello 1 sono distinti da quelli di livello 2.

- Garr-gw-bo.infn.it genera un L1-LAN-Hello sulla rete broadcast per il livello 1 di 60 ottetti; un L2-LAN-Hello sulla rete broadcast per il livello 2 di 60 ottetti; un PtP-Hello sulla rete punto-a-punto di 53 ottetti. In totale genera 173 ottetti.
- Niscn4.infn.it genera gli stessi pacchetti e della stessa dimensione di quelli di garr-gw-bo, per un totale di 173 ottetti.
- Garr-gw.cnr.it genera gli stessi pacchetti e della stessa dimensione di quelli di garr-gw-bo, per un totale di 173 ottetti.
- Faeta.unipi.it genera gli stessi pacchetti e della stessa dimensione di quelli di garr-gw-bo, per un totale di 173 ottetti.
- Cis1pi.pi.infn.it genera un L1-LAN-Hello sulla rete broadcast per il livello 1 di 60 ottetti; un L2-LAN-Hello sulla rete broadcast per il livello 2 di 60 ottetti; due PtP-Hello di 63 ottetti, uno per ciascuna rete punto-a-punto. In totale genera 246 ottetti.

In totale, i cinque router immettono nel routing domain 938 ottetti ogni HelloTimer secondi per la gestione degli Hello.

Area BO

All'interno dell'area BO (di livello 1), i router si scambiano i Level 1 link state Packet Data Unit (L1-LSP) per informarsi dello stato dei collegamenti; la ricezione di questi è confermata tramite l'invio dei Level 1 partial sequence number Packet Data Unit (L1-PSN) sulle linee punto-a-punto o con dei Level 1 complete sequence number Packet Data Unit (L1-CSN) da parte dei designated router sulle reti broadcast.

- Garr-gw-bo.infn.it genera un L1-LSP di 85 ottetti (52 per l'header, 2 per indicare i due protocolli supportati, 11 per indicare il vicino, 8 per gli indirizzi IP configurati, 12 per l'indirizzo della subnet del CNAF) e un L1-CSN di 49 ottetti (33 per l'header e 16 per l'unico L1-LSP ricevuto) in quanto designated router.

- Niscn4.infn.it genera solo un L1-LSP di 85 ottetti (52 per l'header, 2 per indicare i due protocolli supportati, 11 per indicare il vicino, 8 per gli indirizzi IP configurati, 12 per l'indirizzo della subnet del CNAF).

In totale i due router immettono nell'area BO 219 ottetti ogni LSPGenerationInterval secondi per la gestione dei link state packet di livello 1.

Area PI

All'interno dell'area PI i router si scambiano i Level 1 link state Packet Data Unit, i Level 1 partial sequence number Packet Data Unit ed i Level 1 complete sequence number Packet Data Unit.

- Garr-gw.cnr.it genera un L1-LSP di 85 ottetti (52 per l'header, 2 per indicare i due protocolli supportati, 11 per indicare il vicino, 8 per gli indirizzi IP configurati, 12 per l'indirizzo della subnet del CNUCE) che raggiunge Faeta.unipi il quale lo rispedisce a cis1pi; genera anche un L1-CSN di 65 ottetti (33 per l'header e 32 per i due L1-LSP ricevuti) in quanto designated router.

- Faeta.unipi.it genera un L1-LSP di 96 ottetti (52 per l'header, 2 per indicare i due protocolli supportati, 22 per indicare i due vicini, 8 per gli indirizzi IP

configurati, 12 per l'indirizzo della subnet del CNUCE) che raggiunge garr-gw.cnr e cis1pi; genera anche un L1-PSN di 33 ottetti (17 per l'header e 16 per il L1-LSP ricevuto da cis1p1).

- Cis1pi.pi.infn.it genera un L1-LSP di 89 ottetti (52 per l'header, 2 per indicare i due protocolli supportati, 11 per indicare il vicino, 12 per gli indirizzi IP configurati, 12 per l'indirizzo della subnet dell'INFN) che raggiunge faeta il quale lo rispedisce a garr-gw.cnr; genera anche un L1-PSN di 49 ottetti (17 per l'header e 32 per i L1-LSP ricevuti).

In totale, i tre router immettono nell'area PI 687 ottetti ogni LSPGenerationInterval secondi per la gestione dei link state packet di livello 1.

Level 2 subdomain

Per effettuare il routing di livello 2, i router si scambiano i Level 2 link state Packet Data Unit (L2-LSP) per informarsi dello stato dei collegamenti; la ricezione di questi è confermata tramite l'invio dei Level 2 partial sequence number Packet Data Unit (L2-PSN) sulle linee punto-a-punto o con dei Level 2 complete sequence number Packet Data Unit (L2-CSN) da parte dei designated router.

- Garr-gw-bo.infn.it genera un L2-LSP di 118 ottetti (60 per l'header, 2 per indicare i due protocolli supportati, 36 per indicare i due vicini, 8 per gli indirizzi IP configurati, 12 per l'indirizzo della subnet del CNAF) che raggiunge tutti gli altri quattro router; genera un L2-CSN di 97 ottetti (33 per l'header e 64 per i quattro L2-LSP ricevuti) in quanto designated router; genera anche un L2-PSN di al più 81 ottetti (17 per l'header e 64 per i quattro L2-LSP ricevuti).

- Niscn4.infn.it genera un L2-LSP di 118 ottetti (60 per l'header, 2 per indicare i due protocolli supportati, 36 per indicare i due vicini, 8 per gli indirizzi IP configurati, 12 per l'indirizzo della subnet del CNAF) che raggiunge tutti gli altri quattro router; genera anche un L2-PSN di 81 ottetti (17 per l'header e 64 per i quattro L2-LSP ricevuti).

- Garr-gw.cnr.it genera un L2-LSP di 118 ottetti (60 per l'header, 2 per indicare i due protocolli supportati, 36 per indicare i due vicini, 8 per gli indirizzi IP configurati, 12 per l'indirizzo della subnet del CNUCE) che raggiunge tutti

gli altri quattro router; genera un L2-CSN di 97 ottetti (33 per l'header e 64 per i quattro L2-LSP ricevuti) in quanto designated router; genera anche un L2-PSN di al più 81 ottetti (17 per l'header e 64 per i quattro L2-LSP ricevuti).

- Faeta.unipi.it genera un L2-LSP di 118 ottetti (60 per l'header, 2 per indicare i due protocolli supportati, 36 per indicare i due vicini, 8 per gli indirizzi IP configurati, 12 per l'indirizzo della subnet del CNUCE) che raggiunge tutti gli altri quattro router; genera anche un L2-PSN di 81 ottetti (17 per l'header e 64 per i quattro L2-LSP ricevuti).

- Cis1pi.infn.it genera un L2-LSP di 122 ottetti (60 per l'header, 2 per indicare i due protocolli supportati, 36 per indicare i due vicini, 12 per gli indirizzi IP configurati, 12 per l'indirizzo della subnet del CNUCE) che raggiunge tutti gli altri quattro router; genera anche due L2-PSN di 81 ottetti (17 per l'header e 64 per i quattro L2-LSP ricevuti), uno per ogni linea punto-a-punto.

In totale i cinque router immettono nel level 2 subdomain 2560 ottetti ogni LSPGenerationInterval secondi per la gestione dei link state packet di livello 2.

5.3.6 Confronto della occupazione di banda

Nella tabella 5.6 sono riportati i numeri totali di ottetti dei pacchetti di hello generati dai tre protocolli ad ogni intervallo di tempo; per effettuare il confronto si sono considerati HelloInterval e HelloTimer uguali.

	OSPF	IS-IS	integrated IS-IS	OSPF + IS-IS
Hello	760	802	938	1562

Tabella 5.6. Banda occupata dagli Hello espressa in ottetti per intervallo di ritrasmissione.

Nella tabella 5.7 sono riportati il numero di ottetti totali dei LSP e dei loro acknowledge generati dai tre protocolli ad ogni intervallo di tempo; per effettuare il confronto si sono considerati LSRefreshTimer e LSPGenerationInterval uguali.

	OSPF	IS-IS	integrated IS-IS	OSPF + IS-IS
Area BO	404	171	219	575
Area PI	952	535	687	1487
Backbone	1596	2560	3056	4156
Totale	2952	3266	3962	6218

Tabella 5.7. Banda occupata dai LSP espressa in ottetti per intervallo di ritrasmissione.

Riguardo a integrated IS-IS, bisogna notare che la differenza della banda occupata tra un suo uso solo IP e l'uso IP + CLNP è minima, poiché gli NSAP address sono usati in entrambi i casi.

Da queste tabelle si può allora osservare che tra i protocolli IP, OSPF occupa in totale meno banda rispetto a integrated IS-IS, anche se è vero l'inverso se si considera solo il livello 1. La differenza maggiore è nel routing di livello 2 (backbone), dove integrated IS-IS risente della presenza di un router in più (faeta.unipi.it non è un backbone router OSPF ma è di livello 2 per integrated IS-IS) e dell'uso degli NSAP che sono lunghi ben 20 ottetti.

Se invece si confronta integrated IS-IS con Ship In the Night (OSPF in parallelo a IS-IS), si nota subito la sproporzione tra i due: l'occupazione di banda di S.I.N. è quasi doppia rispetto a integrated IS-IS. Questo accade poiché molte informazioni sono ridondanti nei pacchetti OSPF e IS-IS, soprattutto quelle di intestazione per l'identificazione dei pacchetti e dei router che li hanno inviati.

Naturalmente questa è una stima dell'occupazione di banda che non confronta la funzionalità né l'affidabilità delle possibili scelte. Scelte che dovranno tenere conto, oltre che dell'eventuale carico di CPU al crescere delle dimensioni della rete e dei protocolli usati, anche delle "politiche" di routing nazionali ed internazionali.

6

Evoluzione dei protocolli di rete e del GARR

L'Internet è in rapida crescita e di conseguenza in rapida evoluzione: i protocolli di comunicazione e di routing in uso sono stati sviluppati per una rete completamente diversa da quella attuale e quindi non sempre in grado di risolvere efficientemente i problemi che si verificano oggi a causa delle elevate velocità di trasmissione, dell'elevato numero di host connessi e della complessità del sistema.

È questo un momento di transizione: a livello mondiale decine di progettisti stanno ideando le specifiche del nuovo protocollo che sostituirà IP. In Italia, nel GARR, si sta pensando a quali protocolli di routing adottare nella topologia dell'infrastruttura trasmissiva, tenendo anche conto del nuovo protocollo IP che nel giro di poco tempo sarà disponibile.

6.1 IP next generation

Come già riferito nella sezione 1.6, si sta cercando una soluzione ai problemi che la crescita esponenziale degli host collegati a Internet sta causando. Piuttosto che delle modifiche al protocollo IP, si è deciso di metterne a punto uno nuovo che risponda alle attuali e prossime esigenze, in particolare che permetta di indirizzare molti più host e che realizzi un efficiente routing gerarchico.

6.1.1 Breve storia

Nell'autunno 1992 la comunità di Internet presentò quattro diverse proposte per IPng; esse erano *CNAT*, *IP encaps*, *Nimrod* e *Simple CLNP*. Nel dicembre 1992 seguirono altre tre proposte: *The P Internet Protocol* (PIP), *The Simple Internet Protocol* (SIP) e *TP/IX*. Nella primavera del 1993 Simple CLNP divenne *TCP and UDP with Bigger Address* (TUBA) e IP encaps divenne *IP Address Encapsulation* (IPAE). Nell'autunno 1993 IPEA e SIP si unirono mantenendo sempre il nome di SIP; in seguito SIP e PIP si unirono anch'esse nel gruppo che si chiamò *Simple Internet Protocol Plus* (SIPP); più o meno nello stesso tempo TP/IX cambiò nome in *Common Architecture for the Internet* (CATNIP).

Il 25 luglio 1994 l'IETF ha deciso che il nuovo protocollo, denominato IPv6, dovrà essere sviluppato a partire dalle proposte di SIPP¹.

6.1.2 Criteri per la valutazione del nuovo protocollo²

L'IETF³, per orientare i progettisti e poi effettuare una scelta, ha stabilito alcuni principi generali da rispettare e diversi criteri di valutazione per il nuovo protocollo.

Principi generali

Semplicità architetturale: in modo che possa essere usato da diversi protocolli di livello più alto ed utilizzare più protocolli di livello inferiore.

Un protocollo globale: l'Internet deve fornire un servizio di connettività globale per più protocolli in modo di permettere a tutti di connettersi.

Lunga vita: è molto complesso cambiare un protocollo diffuso come l'IP; è necessario che IPng sia progettato pensando che dovrà essere usato per una

¹Dal documento ufficiale: " 11. IPng RECOMMENDATION: After a great deal of discussion in many forums and with the consensus of the IPng Directorate, we recommend that the protocol described in Simple Internet Protocol Plus (SIPP) Spec. (128 bit ver) be adopted as the basis for IPng, the next generation of the Internet Protocol." [IPng-rec]

²[IPng-crit]

³IETF: Internet Engineering Task Force.

ventina d'anni e che permetta all'Internet di crescere ulteriormente senza problemi.

Cooperazione decentralizzata: il successo dell'Internet è dovuto in parte al fatto che non c'è un punto di controllo centralizzato dell'intera rete, ma ognuno sviluppa quanto gli interessa cercando solo di mantenere la connettività; questo deve essere anche per IPng.

Criteri

Espandibilità: IPng deve permettere di indirizzare più di mille miliardi di host e consentire una crescita rapida del numero di rete connesse. È proprio questo uno dei motivi che ha portato alla necessità di cambiare l'IP.

Topologie flessibili: IPng deve permettere molte e differenti topologie di rete e consentire che l'intera Internet abbia un grosso diametro.

Efficienza: I router devono essere in grado di gestire il traffico IPng ad una velocità tale da utilizzare pienamente i dispositivi di comunicazione ad alta velocità oggi disponibili e prevedere di fare altrettanto con quelli che verranno sviluppati.

Robustezza: Il network service e gli associati protocolli di routing e di controllo devono essere robusti, cioè devono comportarsi correttamente anche in presenza di pacchetti malformati, perdita di informazioni, interruzione di collegamenti, non funzionamento di host e di router. In caso di configurazioni errate, IPng deve accorgersi di questi errori e al limite correggerli.

Transizione: Deve essere previsto un piano di transizione da IP a IPng che permetta ai due protocolli di convivere per un certo periodo; questo poiché è impossibile stabilire un giorno in cui tutte le macchine IP passino a IPng contemporaneamente. Durante questo periodo la rete deve mantenere le caratteristiche di robustezza.

Indipendenza dai mezzi di collegamento: Il protocollo deve poter lavorare su reti costituite da qualsiasi tipo di LAN, MAN o WAN, ognuna con la propria velocità di trasmissione. Deve supportare collegamenti multiaccesso e punto a punto, oltre che a circuiti permanenti e commutati.

Consegna dei datagrammi di tipo connectionless: Il protocollo deve fornire un servizio di consegna dei datagrammi non affidabile e senza connessione

come IP poiché si è constatato essere la soluzione migliore.

Configurazione: Il protocollo dovrebbe potersi configurare in modo automatico, almeno su zone della rete non troppo complesse. Inoltre deve permettere una facile allocazione degli indirizzi e consentire di cambiare facilmente la topologia di un routing domain.

Sicurezza: Non deve essere preclusa la sicurezza delle operazioni. In particolare si deve impedire la possibilità di inserire nella rete pacchetti che distruggano il routing; deve essere inoltre possibile consentire o meno l'accesso a determinati percorsi ai pacchetti a seconda di chi li ha generati.

Indirizzi unici: IPng deve assegnare indirizzi che siano unici in tutta l'Internet e che non siano ambigui.

Documentazione: I protocolli che definiscono IPng e tutti gli altri protocolli ad essi associati, come ad esempio quelli di routing, devono essere pubblicati su documenti standard quali sono gli RFC e come tali essere di libera distribuzione.

Indirizzi multicast: Il protocollo deve permettere sia l'indirizzamento unicast che quello multicast. Il routing del multicast deve essere automatico e dinamico. Deve essere possibile il multicast su aree molto estese ma deve anche essere possibile limitare il multicast a specifiche network.

Estensibilità: Il protocollo deve essere estensibile, cioè permettere che in futuro vi siano aggiunti servizi che si rivelino necessari. Deve però essere possibile fare queste estensioni senza che tutti le realizzino contemporaneamente. Esse devono poter riguardare: gli algoritmi usati dal protocollo (i quali devono essere disaccoppiati dal protocollo stesso), le intestazioni dei pacchetti, le strutture dati, i tipi di pacchetti.

Servizi di rete: Il protocollo deve permettere alla rete di associare pacchetti a particolari classi di servizi e questi pacchetti devono poter essere trattati diversamente a seconda del loro servizio.

Mobilità: Il protocollo deve consentire la mobilità degli host e delle reti. Questa mobilità è di due tipi: host che vengono staccati da una rete e collegati ad un'altra, oppure host che pur muovendosi rimangono sempre collegati alla stessa rete.

Protocolli di controllo: Come ICMP per IP, deve essere definito un protocollo

che realizzi funzioni di debug, test e controllo della rete.

Tunneling: IPng deve supportare il tunneling, cioè essere in grado di trasportare protocolli non IP attraverso zone in cui questi non sono abilitati.

6.1.3 TUBA⁴

TUBA (TCP and UDP with Bigger Addresses) è, insieme a SIPP, una delle due principali proposte per IPng che a lungo sono state dibattute. Essa prevede di sostituire il protocollo IP con quello CLNP, eventualmente modificato in accordo con l'ISO; punto di forza è la prospettiva di avere un unico protocollo globale e un unico sistema di indirizzamento visto che il CLNP è già attivo su diverse reti e che molti studi sugli assegnamenti degli indirizzi NSAP sono già stati fatti. L'INFN ed il GARR hanno appoggiato questa soluzione⁵ essendo già in funzione sulle loro reti il CLNP come network protocol di DECnet/OSI (phase V). Essa però è stata scartata.

Caratteristiche⁶

TUBA prevede di adattare il protocollo TCP all'uso di CLNP come network layer protocol; il CLNP è descritto nella sezione 2.1.2 alla quale si rimanda.

CLNP fornisce sostanzialmente gli stessi servizi di IP oltre ad avere molte caratteristiche in comune: la massima dimensione dei datagrammi è la stessa e come IP prevede un meccanismo di frammentazione/ricostruzione; in entrambi una checksum controlla la correttezza delle trasmissioni e un "tempo di vita" (time to live) evita che dei pacchetti sopravvivano troppo a lungo all'interno dell'internet senza essere consegnati; entrambi prevedono delle opzioni per controllare e modificare il flusso dei datagrammi.

CLNP non realizza tutto quello che al nuovo protocollo è richiesto (cfr. sezione 6.1.2); in particolare gli indirizzi multicast e un meccanismo per l'indirizzamento di host mobili non sono presenti, ma sono stati a lungo studiati per poter essere inseriti nello standard ISO.

⁴[RFC 1347]

⁵[RFC 1676]

⁶[RFC 1561]

Transizione da IP a TUBA⁷

Dovendo CLNP subentrare ad IP senza interrompere il servizio, ma non potendo avvenire la sostituzione su tutti gli host contemporaneamente, è stato necessario sviluppare un meccanismo di transizione in cui IP e CLNP convivono finché tutti gli host non saranno stati aggiornati. TUBA ha previsto di realizzare un meccanismo denominato *dual stack* il quale prevede che, su ogni macchina in transizione, IP e TUBA siano eseguiti parallelamente (figura 6.1): a seconda del protocollo supportato dall'host destinazione viene scelto quello dei due in grado di colloquiare. Questa scelta viene fatta in base al tipo di indirizzo restituito dal *DNS server*⁸, NSAP o IP. La comunicazione avverrà utilizzando CLNP solo se sia la sorgente che la destinazione lo supportano già.

Applicazioni Internet	
TCP / UDP	
IP	CLNP

Figura 6.1. TUBA dual stack.

Nel periodo iniziale, durante il quale la rete CLNP non sarà completamente connessa, sarà necessario effettuare un tunneling dei pacchetti CLNP nelle aree solo IP; in seguito, quando sarà la rete IP a non essere più connessa, verrà effettuato il tunneling dei pacchetti IP.

6.1.4 SIPP⁹

SIPP è il protocollo che è stato designato a succedere all'IP attuale (versione 4, IPv4); esso viene ora denominato IPng o IP versione 6 (IPv6). Esso è stato progettato come evoluzione di IPv4 e non come radicale cambiamento, adattandolo alle nuove esigenze ma alleggerendolo delle cose meno necessarie; il

⁷[Tuba-trans]

⁸DNS: Domain Name Server; server che restituisce agli host che gli trasmettono un nome di dominio (es. `vaxbo.bo.infn.it`), il corrispondente indirizzo IP (es. `131.154.10.60`).

⁹[IPv6-spec] e [IPng-view]

fatto di essere un protocollo “su misura” e di non dovere rispettare specifiche già presenti (come TUBA nei riguardi di CLNP) è stato uno dei motivi per cui è stato preferito.

Le principali differenze di SIPP da IPv4 sono:

Espansione delle capacità di indirizzamento e di routing: la dimensione degli indirizzi passa da 32 a 128 bit per permettere il supporto di diversi livelli di indirizzamento gerarchico e di indirizzare un numero maggiore di host. È stato definito un nuovo tipo di indirizzo, denominato *cluster address*, che identifica un gruppo di nodi con un prefisso di indirizzo comune.

Semplificazione del formato dell'intestazione: alcuni campi dell'intestazione di IPv4 sono stati tolti o resi opzionali per ridurre il tempo di processo di ogni pacchetto e anche per non aumentare la occupazione di banda nonostante la crescita delle dimensioni degli indirizzi. Infatti, nonostante gli indirizzi siano quattro volte più grandi dei precedenti, l'intestazione SIPP è solamente il doppio di quella IPv4.

Migliore utilizzo delle opzioni: il formato dell'intestazione SIPP permette di utilizzare al meglio le opzioni senza limiti di dimensioni, inoltre consente di aggiungerne delle nuove in futuro.

Maggiori capacità di controllo dei servizi: è stata aggiunta la capacità di etichettare i pacchetti per richiedere ai router uno speciale trattamento durante il loro cammino.

Indirizzi

Gli indirizzi SIPP, lunghi 128 bit, identificano singoli nodi o insiemi di nodi; un numero così alto di bit permette di indirizzare una quasi infinita quantità di host (più di seicentomila miliardi di miliardi di indirizzi per metro quadrato sulla Terra); in realtà essi servono per consentire un efficiente routing gerarchico.

Ci sono tre tipi di indirizzi: unicast, cluster, multicast; la loro distinzione avviene sui primi bit (da tre a otto) dell'indirizzo.

Unicast address: identificano un singolo nodo. Per realizzare il routing gerarchico è necessario che siano assegnati con un criterio per cui host di una stessa zona abbiano prefissi comuni; nella figura 6.2 è riportato uno schema

di assegnazione basato su chi lo fornisce (provider based).

3	n bit	m bit	p bit	125-n-m-p
010	Provider ID	Subscriber ID	Subnet ID	Node ID

Figura 6.2. Provider based unicast address.

Cluster address: identificano gruppi di nodi con un prefisso comune nei loro indirizzi unicast; sono tali che un pacchetto inviato ad un indirizzo di cluster viene consegnato solo al “più vicino” dei nodi che svolgono la funzione di boundary router nel cluster. Possono essere usati solamente come indirizzi destinazione. In figura 6.3 è riportato il formato di un cluster address.

n bit	128-n bit
Cluster prefix	00000000000000

Figura 6.3. Cluster address.

Multicast address: identificano gruppi di nodi e sono tali che un pacchetto inviato ad un indirizzo multicast è ricevuto da tutti i nodi del gruppo; un nodo può appartenere a più gruppi. In figura 6.4 è riportato il formato di un multicast address; i flag indicano se l'indirizzo è assegnato permanentemente o solo temporaneamente; lo Scope delimita la portata dell'indirizzo ovvero se deve essere consegnato solo all'interno di un'area oppure di un sito oppure di aree sempre più vaste, via via fino alla globalità dell'Internet.

8 bit	4 bit	4 bit	112 bit
11111111	Flag	Scope	Group Identifier

Figura 6.4. Multicast address.

Intestazione dei pacchetti SIPP

L'intestazione principale dei pacchetti SIPP è stata ridotta al minimo; è però possibile inserire in cascata altre intestazioni per eventuali opzioni. In figura 6.5 è riportato lo schema dell'intestazione SIPP.

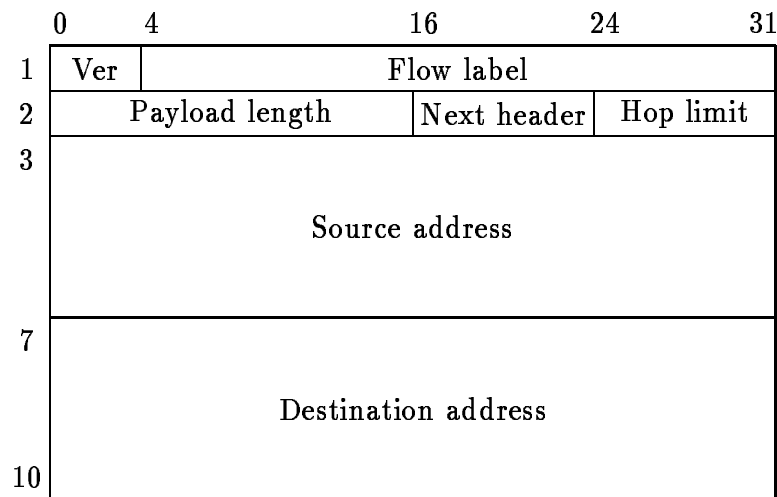


Figura 6.5. Intestazione di un pacchetto SIPP.

Version: indica la versione del protocollo; per SIPP è uguale a 6.

Flow label: è usato dalla sorgente per etichettare quei pacchetti per i quali si vuole richiedere uno speciale trattamento da parte dei router.

Payload length: indica la lunghezza del pacchetto in ottetti esclusa l'intestazione principale.

Next header: identifica il tipo di intestazione che segue quella principale.

Hop limit: indica il numero massimo di router che un pacchetto può attraversare nel suo tragitto; è decrementato di uno da ogni router.

Source address: contiene l'indirizzo del nodo che ha generato il pacchetto.

Destination address: contiene l'indirizzo del nodo destinazione.

Le opzioni sono inserite in intestazioni indipendenti poste tra quella principale e l'area dati. Quasi tutte queste intestazioni sono esaminate solamente dai nodi sorgente e destinazione in modo da accrescere la velocità di trasmissione. Le intestazioni-opzioni finora previste riguardano il source routing, la

frammentazione/ricostruzione, l'autenticazione; le opzioni esaminate da tutti i router di un percorso risiedono in un'intestazione denominata Hop-by-Hop.

Routing

Il routing di SIPP è abbastanza simile a quello di IP con la differenza della dimensione dell'indirizzo. Con diverse modifiche, tutti i protocolli di routing di IP (OSPF, IDRP, IS-IS, RIP) possono essere usati anche da SIPP.

È prevista l'opzione di *source routing* per cui è possibile indicare tutti i router o i cluster che un pacchetto deve attraversare per giungere a destinazione; inoltre è in progetto la gestione degli host mobili.

Transizione da IP a SIPP¹⁰

La transizione da IPv4 a IPv6 deve avvenire nel modo più trasparente possibile senza che si verifichino interruzioni nel servizio; il *Simple Ip version six Transition* (SIT) è un insieme di protocolli che realizzano questo passaggio; esso è derivato da IPAE.

SIT prevede che inizialmente gli host supportino sia IPv4 che IPv6; questa fase può avvenire a diverse velocità nei vari siti e può rivelarsi anche piuttosto lunga. La fase successiva in cui le aree diventeranno *solo-IPv6* potrà avvenire solamente quando la maggior parte degli host sarà sia IPv4 che IPv6.

Le macchine in transizione supportano IPv4 e IPv6 in parallelo; questo significa che il TCP ha a disposizione due protocolli del network layer per comunicare con altri host e che le comunicazioni tra coppie di questi avvengono usando il protocollo che hanno in comune. Ad esempio una macchina IPv4/IPv6 colloquia con una solo-IPv4 tramite IPv4 e con una solo-IPv6 tramite IPv6. Per sapere che protocollo usare, un host osserva l'indirizzo di destinazione che gli viene restituito dal DNS che può essere IPv4 o IPv6.

SIT prevede alcuni meccanismi che consentano anche ad host solo-IPv4 di colloquiare con quelli solo-IPv6; questo avviene tramite dei router in grado di tradurre i pacchetti IPv4 in pacchetti IPv6 e viceversa. Per fare questo gli indirizzi IPv4 vengono mappati in un indirizzo IPv6 come riportato in figura

¹⁰[IPv6-trans]

6.6.

96 bit	32 bit
0:0:0:0:0:0	IPv4 address

Figura 6.6. Indirizzo IPv4 mappato in uno IPv6.

Gli indirizzi IPv6 invece devono essere in un formato *compatibile IPv4* con la struttura riportata in figura 6.7. Essi vengono assegnati solamente a quei nodi solo IPv6 che vogliono colloquiare con quelli solo-IPv4.

96 bit	32 bit
0:0:0:0:0:ffff	IPv4 address

Figura 6.7. Indirizzo IPv6 compatibile IPv4.

Se le comunicazioni tra coppie di host IPv6 avvengono attraverso zone solo-IPv4 allora i router IPv4/IPv6 agli estremi di queste zone possono effettuare il *tunneling* dei pacchetti, ovvero il primo router di bordo della zona solo-IPv4 incapsula il pacchetto IPv6 in uno IPv4 e lo trasmette, utilizzando IPv4, al router all'estremità opposta della zona solo-IPv4; questo estrae il pacchetto IPv6 e torna ad instradarlo utilizzando IPv6.

6.1.5 La nomina di SIPP a IPv6

La scelta del protocollo che doveva sostituire IP è stata lunga e travagliata: due partiti, i sostenitori di SIPP e quelli di TUBA, si sono a lungo confrontati ognuno decantando i pregi del proprio favorito e sottolineando i difetti del rivale. La decisione, che era doveroso prendere e in tempi brevi, ha quindi lasciato degli scontenti. Nell'incontro dell'IETF tenutosi a Toronto (Canada) nel luglio 1994, si è deciso che sarà SIPP a ispirare le specifiche per IPv6 e che su di esso dovranno concentrarsi gli sforzi dei progettisti.

I motivi di questa scelta sono solo in parte tecnici; in effetti SIPP si è liberamente ispirato ad IPv4, ne ha mantenuto gli aspetti funzionali, ha

eliminato quelli inutili o poco efficienti e ha aggiunto quanto si è ritenuto necessario nell'immediato e prossimo futuro. Ma la sostanza non è particolarmente diversa dal CLNP di TUBA e comunque non garantisce l'unicità del protocollo.

Si è anche detto che il CLNP non rispettava tutte le caratteristiche richieste al nuovo protocollo (cfr. sezione 6.1.2); questo è vero, ma nemmeno SIPP, allo stato attuale, le rispetta tutte. Entrambi i protocolli hanno ancora bisogno di essere affinati, ma le modifiche al CLNP avrebbero dovuto essere mediate con l'ISO, il che risultava scomodo.

I veri motivi sono probabilmente più politici o commerciali ma comunque non ufficialmente dichiarati. Probabilmente l'IETF ha voluto assicurarsi il pieno potere sulle nuove specifiche, riservandosi la facoltà di apportarvi modifiche in piena libertà senza dovere mediare la gestione con l'ISO, l'ente di standardizzazione che agli occhi degli americani (e quindi della maggior parte dei componenti dell'IETF) ha il difetto di essere europeo.

6.2 Evoluzione della rete GARR

Anche la rete GARR deve far fronte ai problemi derivati dal sempre maggior numero di nuove reti e di host che diventano parte dell'autonomous system 137; queste aumentano la complessità dell'internet italiana e di conseguenza la sua connessione con le reti internazionali.

In questa sezione, sulla base dei test effettuati e delle caratteristiche del protocollo IPv6, si vuole formulare una proposta di evoluzione dell'utilizzo dei protocolli di routing nella rete del GARR.

6.2.1 I requisiti per la nuova configurazione

Lo stato della rete GARR è descritto nel capitolo 4; i problemi a cui deve far fronte si scontrano con la necessità di avere alte prestazioni in termini di banda, affidabilità e servizi. Allo stato attuale delle tecnologie e della topologia è necessario rivedere le politiche di gestione, dato che furono impostate per una situazione molto diversa da quella odierna.

Innanzitutto, essendo aumentato il numero dei collegamenti punto-a-punto e delle sottoreti, si è complicata notevolmente la topologia; non è più sufficiente il protocollo distance vector IGRP per gestire le rapide variazioni dovute ai cambiamenti di stato dei collegamenti, ma è necessario un protocollo link state che sia in grado di garantire anche un rapido back-up.

L'aumento delle sottoreti connesse causa anche problemi dovuti al sovraccarico di informazioni di routing scambiate dai protocolli; è necessario prevedere forme di compattamento o gerarchizzazione degli indirizzi per ridurre la dimensione delle tabelle di routing e quindi accrescere l'efficienza dei router.

La presenza di tanti router in siti diversi implica una difficoltà nel gestire il routing in maniera uniforme e coerente; inoltre aumenta le probabilità che una modifica errata della configurazione di un router abbia gravi ripercussioni in tutto il routing domain. Può essere vantaggioso prevedere una gestione centralizzata del backbone indipendente dalle variazioni e dalla gestione dei diversi siti, anche se è molto complesso a livello politico.

Il nuovo protocollo IP è in fase di sviluppo e nel frattempo all'interno del

GARR si sta terminando la migrazione della rete DECnet a DECnet/OSI; la nuova configurazione IP deve tenere conto di entrambi questi aspetti: di come favorire una successiva migrazione a IPv6 e di come integrarsi con un protocollo di rete differente.

In definitiva i requisiti da rispettare sono:

1. Utilizzo di un protocollo in grado di gestire topologie complesse e rapidi cambiamenti della stessa.
2. Gerarchizzazione degli indirizzi.
3. Interazione con CLNP.
4. Predisposizione a IPv6.
5. Vari livelli di gestione.

6.2.2 I possibili approcci

1. I protocolli di routing in grado di gestire i rapidi cambiamenti di una topologia complessa sono quelli basati sull'algoritmo link state, ovvero OSPF e integrated IS-IS. La scelta dell'uno o dell'altro dipende molto dai punti successivi; è anche possibile considerare un uso di entrambi per separare la gestione dei diversi livelli di routing (aree e backbone).

Entrambi gestiscono un routing a due livelli; il backbone dovrà comprendere le linee ad alta velocità sulle lunghe distanze (figure 1.7 e 5.1). Le aree invece dovranno raccogliere reti connesse fra di loro e che faranno riferimento ad uno o più router sul backbone; nelle aree dovranno essere anche comprese le linee ad alta velocità su brevi distanze.

2. Per ridurre il numero di informazioni di routing che circolano nell'autonomous system, è necessario riassumere il più possibile gli indirizzi; questo significa, in pratica, mettere tutte le sottoreti delle reti di classe B nella stessa area in modo da annunciare all'esterno solo il prefisso che indica la rete ma non le sottoreti; inoltre, qualora siano stati assegnati degli indirizzi di classe C contigui con un prefisso comune tale che sia comune solo a loro, inserire anch'essi nella stessa area ed annunciarli con il loro indirizzo di supernetwork.

3. Sulla rete GARR si continuerà ad usare DECnet/OSI e quindi il proto-

collo IS-IS; questo influenza notevolmente la scelta del protocollo di routing per IP. Se si utilizzerà integrated IS-IS, le aree che verranno definite coincideranno con quelle definite dagli indirizzi OSI-NSAP dei router dei vari siti. Se invece si utilizzerà OSPF, il routing IP sarà completamente indipendente da quello DECnet/OSI. Questa soluzione implica che i protocolli IS-IS e OSPF dovranno essere utilizzati in parallelo (S.I.N.) e questo è possibile solo sui router più potenti e con molta memoria disponibile; tuttavia non tutti i router della rete sono tali, per cui questa soluzione richiederebbe un investimento economico per potenziarli o sostituirli; oltre a ciò si aggiunge una gestione più complessa.

4. IPv6 è caratterizzato da indirizzi gerarchici; quando verrà adottato dal GARR, sarà necessario assegnare gli indirizzi gerarchici in modo da contraddistinguere ogni area con un unico prefisso, come avviene già con gli indirizzi OSI-NSAP nelle aree DECnet/OSI. Anche in questo caso, se l'uso di integrated IS-IS viene definito in modo da fare corrispondere le aree IP e quelle CLNP, gli indirizzi IPv6 dovrebbero ricalcare quelli NSAP e quindi la transizione da OSI-NSAP a IPv6 sarebbe trasparente; altrimenti l'assegnamento degli indirizzi IPv6 dovrà solo rispettare le aree già definite.

Per quanto riguarda i protocolli di routing, per IPv6 dovrebbero essere sviluppati sia OSPF che integrated IS-IS; non dovrebbero esserci quindi molte variazioni nella loro configurazione.

5. Per separare il più possibile la gestione del routing sul backbone da quello delle aree, si dovrebbero definire dei routing domain diversi per ogni area e per il backbone; la divisione dei routing domain si ottiene utilizzando protocolli di routing differenti, ad esempio OSPF nelle aree e integrated IS-IS nel backbone (o viceversa).

Questa gestione però rischia di sovraccaricare molto i router di bordo che, oltre ai protocolli per CLNP, ne dovrebbero eseguire due per IP.

In seguito sono proposte alcune definizioni di aree e backbone ricavate dalla topologia attuale, sulle quali si sono poi ipotizzati alcuni scenari di utilizzo dei protocolli di routing link state.

6.2.3 Una configurazione per aree geografiche

In questa proposta, le sottoreti presenti nell'AS 137 sono state suddivise per aree geografiche coincidenti, più o meno, con le regioni italiane. La divisione è stata fatta raggruppando le reti che possono essere raggiunte efficientemente da un router di backbone; ad esempio, nell'area dell'Emilia sono state raggruppate buona parte delle reti dei siti connessi direttamente al CNAF e al CINECA, nell'area Piemonte i siti connessi direttamente ai router del Politecnico di Torino e così via.

È stato possibile raggruppare gli annunci di alcune reti di classe C contigue nelle loro supernetwork; con la divisione proposta, il numero di indirizzi annunciati all'esterno delle aree è stato ridotto da 309 a 176, che corrisponde ad un decremento di circa il 43%.

Denominazione delle aree

Le aree definite sono riportate qui di seguito; a fianco di ognuna sono elencati i siti che hanno connessioni interarea.

- *Emilia*: INFN/CNAF, Cineca.
- *Piemonte*: Università di Torino, Politecnico di Torino.
- *Lombardia*: Cilea, INFN Milano.
- *Veneto-Friuli*: Laboratorio Nazionale di Legnaro, Università di Padova, INFN Trieste.
- *Liguria*: INFN Genova, Università di Genova.
- *Toscana-Umbria*: CNR/Cnuce Pisa, INFN Firenze.
- *Lazio*: INFN Roma.
- *Abruzzo-Marche*: Laboratorio Nazionale del Gran Sasso.
- *Puglia-Basilicata*: Bari Csata, INFN Bari.
- *Campania-Calabria*: Università di Napoli.
- *Sicilia*: Università di Catania, Università di Palermo.
- *Sardegna*: CSR4 Cagliari, Sassari.

Gli indirizzi di rete che ognuna di queste aree comprende sono elencate nell'appendice C.

Backbone

Nella figura 6.8 sono riportate le schematizzazioni delle aree geografiche e le linee che le connettono; queste linee vengono a costituire il backbone dell'autonomous system. Lo schema dei collegamenti è semplificato rispetto alla realtà perché non si è distinto tra diversi router di una stessa città; ad esempio, tutti i link contrassegnati terminare al CNAF sono connessi in realtà a diversi router (come riportato in dettaglio nella figura 4.3).

Il backbone considerato è costituito dalle linee punto a punto su lunghe distanze e dai PVC della C-LAN.

6.2.4 Una configurazione per aree di affinità d'uso

In questa sezione si propone una suddivisione delle reti dell'autonomous 137 in aree che raggruppano gli enti affini, come ad esempio, l'INFN, le università, il CNR. Il vantaggio di una simile divisione è che all'interno delle singole aree, alcuni collegamenti su lunghe distanze possono essere dedicati all'utilizzo "privato" da parte dei siti dell'area; ad esempio, definendo un'area INFN a cui appartengono i siti INFN di Milano, Torino e Genova, essendo questi connessi dalla C-LAN potrebbero configurare i PVC come collegamenti solo intrarea, privilegiando in questo modo il proprio traffico.

In questa divisione però non è sempre stato possibile definire gruppi così netti, poiché sono presenti alcuni siti di determinati enti connessi solamente a siti di altri enti. Ad esempio, l'INFN di Trieste ha la linea a 2Mbps che gli permette di comunicare con il GARR collegata con il Cineca; sarebbe stato controproducente inserirlo in un'area INFN con Padova e Legnaro non essendovi connessa.

Denominazione delle aree

Le aree definite sono riportate qui di seguito; a fianco di ognuna sono elencati i siti che hanno connessioni interarea.

- *UNIVERSITÀ-1*: Università di Catania, Università di Palermo.
- *UNIVERSITÀ-2*: CSATA, Università di Lecce.

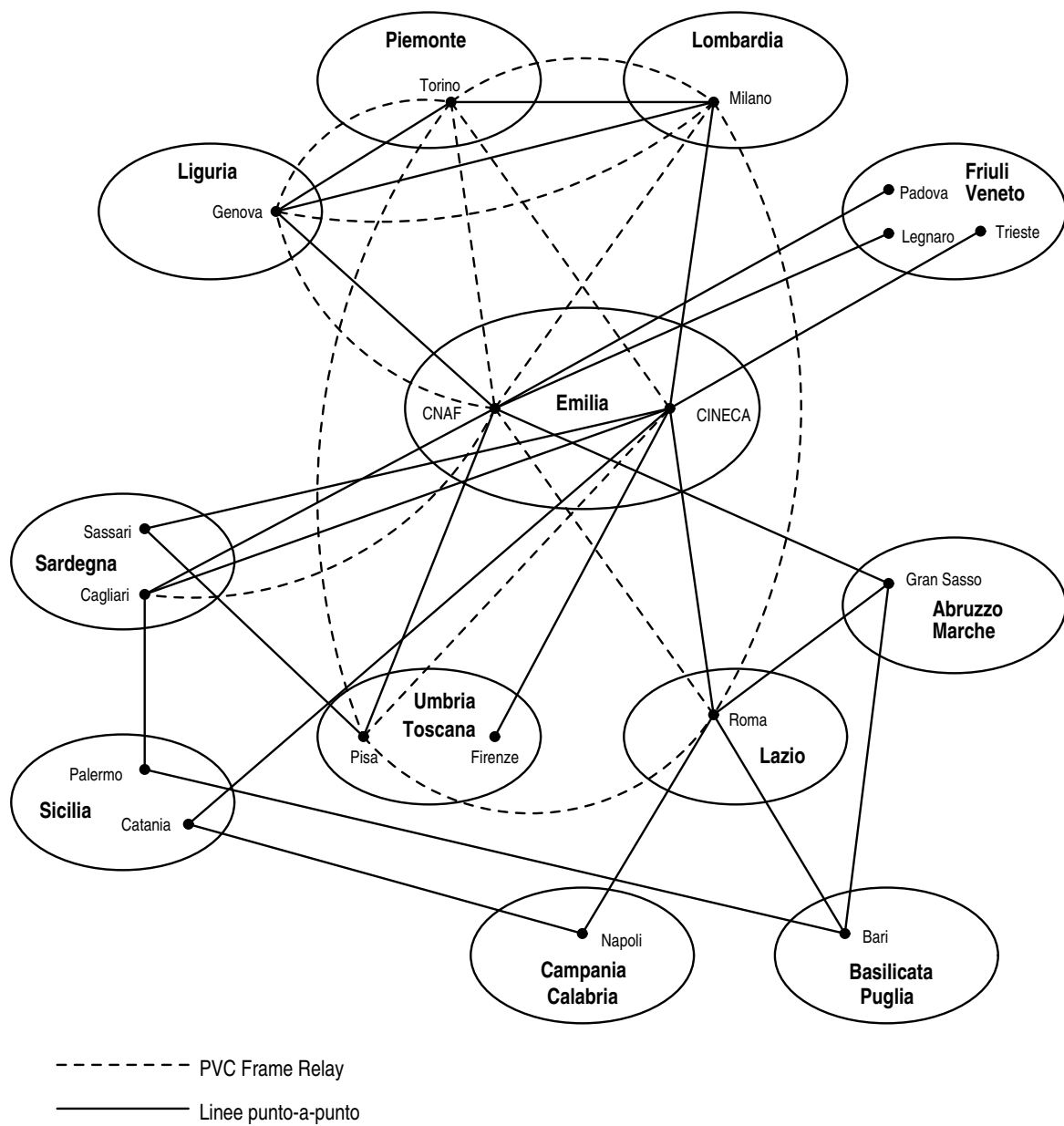


Figura 6.8. Aree geografiche e backbone.

- *UNIVERSITÀ-3*: Cineca, Università di Bologna (ALMAnet), INFN Trieste, Università di Cagliari, Università di Parma, Università di Ferrara.
- *UNIVERSITÀ-4*: Cilea Milano, Politecnico di Torino, Università di Pavia, Università di Genova, Università di Padova.
- *CNR-1*: Cnuce Pisa, Università di Firenze, Università di Pisa.
- *INFN-1*: INFN Napoli, INFN Catania.
- *INFN-2*: INFN Gran Sasso, INFN Aquila, INFN Bari.
- *INFN-3*: CNAF, INFN Bologna, INFN Parma, INFN Ferrara, INFN Cagliari, INFN Firenze.
- *INFN-4*: INFN Milano, INFN Torino, INFN Genova, INFN Pavia.
- *INFN-5*: INFN Padova, INFN Legnaro.
- *ROMA*: INFN Roma, INFN Frascati, Università di Roma.

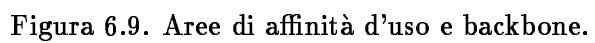
L'elenco delle reti IP appartenenti ad ogni singola area è riportato nell'appendice D. È stato possibile raggruppare gli annunci di alcune reti di classe C contigue nelle loro supernetwork; con la divisione proposta, il numero di indirizzi annunciati all'esterno delle aree è stato ridotto da 309 a 166 che corrisponde ad un decremento di circa il 46% (leggermente meglio del caso delle aree geografiche).

Backbone

La realizzazione del backbone per queste aree è più complessa che per la suddivisione proposta nella sezione 6.2.4, perché diverse aree hanno più di una intersezione tra di loro; ad esempio, l'area INFN-3 e UNIVERSITÀ-3 hanno entrambi siti a Bologna, Ferrara e Parma: quelli dell'INFN sono connessi al CNAF e gli altri al Cineca, ma nelle stesse città sono connessi anche tra di loro. È necessario allora configurare molti collegamenti di backbone interni alle aree per poter utilizzare tutti i link interarea; questo implicherà per OSPF la definizione di molti virtual link; per integrated IS-IS invece la necessità di avere molti IS di livello 2.

Questo fatto penalizza integrated IS-IS perché a livello 2 risulta un grafo con molti nodi; per cui l'algoritmo link state impiegherà più tempo per calcolare i cammini minimi.

La schematizzazione delle aree e del backbone è riportata nella figura 6.9.



6.2.5 Configurazione dei protocolli: possibili scenari

In questa sezione si vuole analizzare l'impatto che l'utilizzo di uno o più specifici protocolli link state avrebbero sulla gestione del routing nelle topologie proposte.

Sono possibili diversi scenari:

- *un protocollo unico per tutto l'autonomous system*: i vantaggi di una simile configurazione sono molteplici: innanzitutto si sfruttano appieno le potenzialità del protocollo link state soprattutto per quanto riguarda l'alta velocità di convergenza; in più, evitando i redistribute, se ne aumenta l'efficienza. La gestione di un unico protocollo è molto più semplice in quanto le configurazioni sarebbero di un unico tipo; inoltre eventuali upgrade del software riguarderebbero un solo protocollo.

- *un protocollo sulle aree e l'altro sul backbone*: in questo modo si determinano tanti routing domain quante sono le aree. Si ha quindi il vantaggio di separare nettamente la gestione delle aree da quella del backbone; inoltre all'interno delle aree è possibile definire di nuovo due livelli di routing, il che può rivelarsi utile in aree con molte reti. Per contro è necessario che il backbone sia connesso non potendo essere definiti virtual link all'interno di un routing domain integrated IS-IS oppure dei percorsi di livello 2 all'interno di un autonomous system OSPF. Altro svantaggio è che i redistribute richiedono ai border router di eseguire due protocolli IP.

- *alcune aree con un protocollo, alcune con l'altro e sul backbone tutti e due*: questo approccio ha il vantaggio di lasciare la libertà ad ogni area di adottare il protocollo più opportuno. In questo caso si verrebbero a costituire di fatto due routing domain distinti, ma integrati tra loro tramite dei redistribute effettuati su alcuni router; il problema di questa soluzione si evidenzia quando devono colloquiare host che appartengono a routing domain differenti, in quanto le comunicazioni interdomain sono svantaggiate rispetto a quelle intradomain dovendo passare comunque per determinati router.

1. Solo OSPF

OSPF deve essere configurato con le aree indicate nella figura 6.8 o nella figura 6.9; i collegamenti riportati sempre nelle stesse figure devono invece

appartenere al backbone; per fare questo occorre che i router a capo dei link siano configurati come area border router.

I collegamenti punto-a-punto ad alta velocità tra area border router della stessa area non devono essere configurati appartenenti al backbone altrimenti non sarebbero usati per comunicazioni intrarea veloci; è il caso della linea tra CNAF e Cineca, di quella tra Padova e Legnaro ed altre simili. È bene però che tra essi siano stabiliti dei virtual link: questi possono costituire delle “scorciatoie” alle comunicazioni fra coppie di aree separate da una terza.

Anche le linee punto-a-punto a bassa velocità (64 kbps) che collegano aree diverse devono essere configurate appartenere al backbone, altrimenti non verrebbero utilizzate (poiché le comunicazioni interarea avvengono solo sul backbone). In pratica è necessario configurare come area border router tutti quei router che hanno link interarea.

Nessuna delle aree previste può essere configurata come stub perché nessuna di esse ha un unico area border router; non ci si faccia trarre in inganno dal fatto che nelle figure i link terminano in un unico punto: come già detto, in realtà le connessioni sono su diversi router.

È necessario stabilire alcune norme per regolare una futura crescita del routing domain; in particolare, nel caso si debbano connettere nuove aree, bisogna che esse abbiano un proprio area border router da collegare ad un altro del backbone. Questo perché in OSPF i router eseguono un algoritmo per ogni area a cui appartengono: se uno stesso router diventa area border router per troppe aree subisce un calo (o un crollo) delle prestazioni (cfr. sezione 3.5.2).

L'utilizzo di OSPF è vantaggioso rispetto a quello di integrated IS-IS soprattutto nelle aree molto vaste e con molti area border router; questo perché i router interni sono in grado di scegliere il migliore per colloquiare con altre aree. Infatti, in un'area come quella dell'Emilia dove ci sono anche una decina di area border router su una stessa rete (es. la 131.154.1.0 al CNAF), è bene che gli internal router individuino subito il migliore per evitare salti inutili.

Lo svantaggio di OSPF è che in una rete come questa deve essere eseguito in parallelo a IS-IS per CLNP, per cui i router risultano sovraccaricati

dovendo eseguire due algoritmi di calcolo dei cammini minimi.

Dai risultati del test riportato nella sezione 5.2 bisogna dedurre però che attualmente una configurazione simile non è realizzabile in quanto nel backbone sono presenti router Cisco e DECnis.

2. Solo integrated IS-IS

Integrated IS-IS deve essere configurato a livello 1 con le aree di indicate nella figura 6.8 e nella figura 6.9; i collegamenti, riportati sempre nelle stesse figure, devono connettere IS di livello 2. Per definire queste nuove aree è necessario riassegnare gli NSAP alla maggior parte degli host; gli NSAP in uso attualmente per il CLNP delimitano infatti aree differenti.

Poiché gli IS di livello 1 non possono scegliere immediatamente l'IS di livello 2 per comunicare con altre aree, è meglio che, in ogni area, gli IS di livello 2 siano direttamente connessi; in questo modo si possono evitare percorsi non ottimali. Per fare questo, può rivelarsi necessario configurare dei router come IS di livello 2 anche se non hanno collegamenti interarea. Questo è un doppio svantaggio poiché oltre a dover fare almeno un salto inutile per trovare l'IS giusto che faccia uscire dall'area, ci devono essere molti più IS di livello 2 che area border router OSPF; ciò significa che ci saranno più router che eseguiranno due algoritmi; inoltre che il grafo di livello 2 risulterà avere più nodi per cui richiederà un maggior tempo di elaborazione.

In compenso integrated IS-IS pone meno limiti nella sua configurazione; infatti non è necessario stabilire quali collegamenti devono essere utilizzati per il routing intrarea e quali per quelli interarea: è sufficiente configurare ai loro capi i router o di livello 1 (routing intrarea) o di livello 2 (routing interarea) o di entrambi i livelli (sia intrarea che interarea); l'unico vincolo è che, per poter usare il collegamento, gli IS ai suoi capi devono essere dello stesso livello.

Il principale vantaggio dell'uso di integrated IS-IS rimane quello dell'aver un unico protocollo di routing tra IP e CLNP e quindi di impegnare meno la CPU dei router per calcolare i cammini minimi.

Il test di integrated IS-IS effettuato sulla C-LAN e riportato nella sezione 5.1 ha dimostrato che una configurazione di questo tipo è realizzabile in prati-

ca; rinarrebbe solo da verificare il funzionamento in topologie così complesse.

3. Backbone OSPF, aree integrated IS-IS

In una configurazione di questo tipo le aree risultano come routing domain indipendenti l'uno dall'altro e dal routing domain del backbone; i border router sono quelli dove vengono eseguiti i protocolli dell'area e del backbone e dove avviene una redistribuzione tra i due.

Il problema più grande che si verifica in questa situazione è che due backbone router appartenenti alla stessa area non possono essere collegati da virtual link poiché questi non attraversano nessuna area OSPF; solo nel caso che esista un collegamento diretto tra coppie di essi è possibile usarlo sia per integrated IS-IS che per OSPF. In più, all'interno dell'area gli IS non possono scegliere il border router migliore per uscire, ma useranno sempre il più vicino a loro; se i border router della stessa area non riescono a colloquiare insieme con OSPF, potrebbe essere seguito un percorso più tortuoso del minimo.

I vantaggi sono che all'interno delle aree non circolano le informazioni che riguardano il resto dell'autonomous system e che, sempre all'interno delle aree, il routing è integrato a quello CLNP, per cui si ha una maggior efficienza del routing intrarea; quest'ultima caratteristica è maggiormente utile nelle aree molto estese (ovvero con molti nodi).

4. Backbone integrated IS-IS, aree OSPF

Anche in una configurazione di questo tipo le aree risultano dei routing domain indipendenti l'uno dall'altro e dal routing domain del backbone; allo stesso modo i border router sono quelli dove vengono eseguiti i protocolli dell'area e del backbone e dove avviene una redistribuzione tra i due.

I border router appartenenti alla stessa area non possono essere connessi per integrated IS-IS a meno che non esista un collegamento diretto; in questa soluzione però, il problema è meno rilevante che nella precedente, in quanto nelle aree OSPF gli internal router sono in grado di selezionare il miglior area border router per le comunicazioni interarea.

Con questa configurazione i vantaggi del routing integrato si hanno sul backbone.

Rispetto alla topologia in esame, questa soluzione è migliore della precedente poiché tutte le aree hanno più di un router connesso al backbone e in tal caso la gestione delle aree di OSPF è più efficiente. In più, il backbone ha una topologia molto complessa e molti nodi, per cui l'impiego del routing integrato libera molte risorse di calcolo dei suoi router. Infatti, come riferito nella sezione 3.1.2, l'algoritmo di calcolo ha una complessità dell'ordine del quadrato del numero dei nodi del grafo; quindi, per il backbone in esame, è comunque necessario un certo tempo per eseguire il calcolo; se però si dovesse raddoppiare l'esecuzione, il tempo impiegato potrebbe diventare troppo.

5. Backbone integrated IS-IS e OSPF, aree integrated IS-IS o OSPF

In questa soluzione è possibile usare nelle aree uno dei due protocolli link state; sul backbone invece sono attivi entrambi e su alcuni suoi router sono effettuati dei redistribute tra i due protocolli per garantire la connettività totale. In questo modo vengono a costituirsi due routing domain distinti, quello OSPF e quello integrated IS-IS, uniti in un unico autonomous system tramite i redistribute.

I criteri di scelta del protocollo da adottare nelle singole aree possono essere diversi: innanzitutto i responsabili del routing di un'area potrebbero preferire mantenere la continuità con i protocolli attualmente in funzione, in modo da non disperdere le conoscenze e l'esperienza acquisite; poi, per quanto riguarda i motivi strettamente tecnici, l'uso di integrated IS-IS può essere vantaggioso dove è attivo il CLNP; oppure la necessità di colloquiare frequentemente con un'area che ha già scelto un determinato protocollo, suggerisce di scegliere lo stesso, in modo da determinare percorsi interni ad un routing domain.

L'uso di tutti e due i protocolli sul backbone serve appunto per non discriminare alcune aree rispetto ad altre; infatti, nel caso fosse utilizzato un solo protocollo sul backbone, le aree che impiegassero lo stesso protocollo sarebbero privilegiate rispetto alle altre. Per integrare pienamente i due routing domain, i redistribute devono avvenire sul maggior numero di router di backbone possibile; questo perché le comunicazioni interdomain passano

sempre per almeno un router dove è effettuato il redistribute tra i protocolli; se questi sono pochi, vengono sovraccaricati di lavoro ed i percorsi rischiano di non risultare ottimali, mentre se sono molti l'interazione tra i due routing domain risulta più trasparente.

Un apparente svantaggio di questa configurazione è che sul backbone si avrebbero due protocolli per IP e uno per CLNP; in realtà tra IS-IS e integrated IS-IS non c'è molta differenza di carico della CPU in quanto l'algoritmo di cammino minimo calcolato è sempre uno; si ha solo una maggiore occupazione di banda in quanto nei pacchetti integrated IS-IS sono aggiunti gli elenchi delle reti IP appartenenti alle singole aree. In questa situazione il carico dei backbone router è paragonabile ai casi precedenti con solo OSPF e IS-IS in parallelo. Resta comunque il fatto che in questa topologia dove il backbone ha molti nodi, i backbone router devono impiegare tempo per calcolare due algoritmi di cammino minimo.

Un inconveniente che potrebbe verificarsi con l'adozione di IPv6 è che molto probabilmente nei primi tempi sarà disponibile uno solo dei due protocolli; potrebbe dunque accadere di avere due differenti tempi di migrazione per i due routing domain oppure la necessità di adottare un unico protocollo di routing.

6.2.6 Conclusioni

La prima divisione in aree proposta (aree geografiche) è forse la più semplice concettualmente, poiché i confini tra le varie aree sono abbastanza netti; infatti i collegamenti interarea sono totalmente costituiti da link punto-a-punto. In più, il numero dei nodi del grafo di livello 2 è notevolmente inferiore all'altro caso, per cui è più rapido il calcolo dei cammini minimi.

È però la seconda proposta (aree di affinità) quella effettivamente più vantaggiosa; principalmente poiché sono più frequenti le comunicazioni tra siti degli stessi enti e quindi il raggrupparli in medesime aree implica una maggiore efficienza ed un miglior controllo del traffico; questo ha come conseguenza anche una minore interazione tra i responsabili delle reti dei vari enti, i quali risultano liberi di adottare le soluzioni più convenienti alla gestione delle proprie aree.

Per ottimizzare questa configurazione sarebbe meglio ridefinire alcuni link punto-a-punto in modo da riuscire a raggruppare effettivamente nella stessa area i siti affini; è il caso dell'INFN di Trieste “costretta” ad appartenere ad un'area dell'Università perché i suoi collegamenti principali non sono con altri siti INFN.

Per quanto riguarda i protocolli di routing da usare, gli scenari “mono-protocollo” sono sicuramente le soluzioni più efficienti e più facilmente gestibili, in quanto si avrebbe una situazione coerente in tutto l'autonomous system.

Nella realtà però sono soluzioni difficilmente realizzabili: la migrazione verso un unico protocollo sarebbe osteggiata da coloro che attualmente ne usano uno diverso. La via più praticabile risulta quindi il quinto scenario che permetterebbe ad ognuno di mantenere il protocollo preferito. Proprio per questo motivo, il quinto scenario è molto adatto alla divisione in aree di affinità: insieme garantiscono la massima indipendenza di gestione.

Le proposte avanzate hanno utilizzato l'infrastruttura dei collegamenti presenti al momento; avendo la possibilità di riconfigurare tutto, si potrebbero pensare molte altre soluzioni.

Ad esempio, volendo sfruttare i vantaggi offerti dalla tecnologia frame relay utilizzata dalla C-LAN, si potrebbe prevedere un accesso per ogni sito di ogni città e definire dei PVC tra le coppie di siti con hanno frequenti contatti. I siti sarebbero configurati come singole aree, mentre i PVC costituirebbero il backbone. In tale situazione si potrebbero usare i protocolli di routing come definito nei cinque scenari precedenti; anche qui la via più semplice è quella del protocollo unico, ma quella più realistica è invece quella del quinto scenario. In questo caso però, essendo il backbone molto articolato e con molti nodi, bisognerebbe provare la effettiva capacità dei router di bordo di eseguire due protocolli link state.

7

Conclusioni

Durante il periodo di preparazione di questa tesi, lo stato della rete e dei protocolli è mutato in continuazione. Rispetto alle linee punto-a-punto, si è aggiunta la possibilità di utilizzare i PVC della C-LAN; è risultata chiara la prospettiva di ottenere una topologia molto magliata che però i protocolli di routing distance vector non sarebbero in grado di gestire efficientemente. La necessità di avere rapidi back-up delle linee impone l'uso dei protocolli link state, veloci nel rilevare le interruzioni dei collegamenti e nel ricalcolare correttamente le tabelle di routing.

Per tali motivi si sono volute effettuare delle prove di funzionalità di questi protocolli utilizzando l'infrastruttura delle reti INFNet e GARR. È stato provato OSPF su una maglia tra Pisa e Bologna; inoltre si è verificata la funzionalità della C-LAN e quella del protocollo di routing link state integrated IS-IS su una topologia fortemente magliata.

Per quanto riguarda il protocollo IP, durante i primi mesi del 1994 sembrava che dovesse essere sostituito da TUBA, e quindi dal CLNP: in tale situazione, integrated IS-IS appariva molto più interessante di OSPF essendo, il primo, l'unico protocollo di routing per CLNP.

Poi è stata presentata una modifica agli indirizzi di SIPP e ai suoi meccanismi di transizione: SIPP ha così acquistato nuovo interesse e il 25 luglio 1994 è stato scelto per sostituire IP. Questo nuovo scenario ha ribaltato le prospettive: OSPF è tornato alla pari di integrated IS-IS e la transizione (in corso) di DECnet a DECnet/OSI (con il CLNP) ha perso una parte della

sua importanza.

Alla luce di tutto questo sono state formulate due proposte di evoluzione del routing del GARR: le singole reti locali sono state raggruppate in aree geografiche o di affinità d'uso e su queste topologie si è valutato l'impatto dell'utilizzo dell'uno o dell'altro protocollo link state. La proposta più convincente è quella che prevede di definire alcune aree di affinità d'uso, cioè che raggruppano le reti dei vari siti appartenenti agli stessi enti (aree INFN, aree università), e di usare entrambi i protocolli di routing (almeno sul backbone) per permettere ad ogni area di impiegare quello dei due più adatto a soddisfare le proprie necessità. Tale soluzione è però certamente più pesante per la gestione e per la capacità elaborativa richiesta ai router. Queste proposte verranno probabilmente prese in considerazione quando il GARR dovrà decidere le strategie di evoluzione della rete nazionale.

Appendice A

Configurazioni dei router per il test Integrated IS-IS su C-LAN

cnaf-gw-3.infn.it

Current configuration:

```
!  
version 10.0  
!  
hostname cnaf-gw-3  
!  
boot host cnaf-gw-3.conf_test 131.154.1.1  
boot network cnaf-gw-3.conf_test 131.154.1.1  
boot system gs3-k.100-2_7 131.154.1.1  
boot system gs3-k.100-2_5 131.154.1.1  
boot system gs3-k.921_bgp4_2_2 131.154.1.1  
boot system rom  
!  
ip accounting-threshold 32000  
frame-relay switching  
clns routing  
!  
interface Ethernet0  
description Ethernet  
ip address 192.135.23.252 255.255.255.0 secondary  
ip address 131.154.1.99 255.255.255.0  
ip accounting
```

```

no mop enabled
isis metric 1 level-1
isis metric 1 level-2
!
interface Ethernet1
no ip address
shutdown
!
interface Serial0
description frame-relay-C-LAN
no ip address
encapsulation frame-relay
bandwidth 2048
frame-relay route 21 interface Serial1 21
!
interface Serial0.1 point-to-point
description CRS4-PVC24
ip address 192.135.32.1 255.255.255.0
bandwidth 256
frame-relay interface-dlci 24 broadcast
!
interface Serial0.2 point-to-point
description ROMA-PVC23
ip unnumbered Ethernet0
ip router isis infn_cnaf
bandwidth 2048
frame-relay interface-dlci 23 broadcast
isis metric 1 level-1
isis metric 1 level-2
!
interface Serial0.3 point-to-point
description TORINO-PVC22
ip unnumbered Ethernet0
ip router isis infn_cnaf
bandwidth 512
frame-relay interface-dlci 22 broadcast
isis metric 1 level-1
isis metric 1 level-2
!
interface Serial0.4 point-to-point

```

```

description MILANO-PVC20
ip unnumbered Ethernet0
ip router isis infn_cnaf
bandwidth 512
frame-relay interface-dlci 20 broadcast
isis metric 1 level-1
isis metric 1 level-2
!
interface Serial1
no ip address
encapsulation frame-relay
bandwidth 2048
clock rate 2000000
frame-relay intf-type dce
frame-relay route 21 interface Serial0 21
!
autonomous-system 2875
!
router isis infn_cnaf
distance 90
distance 90 ip
passive-interface Ethernet0
net 39.380f.1001.0001.0000.0000.0500.0000.0c01.9df1.00
!
router igrp 137
network 131.154.0.0
network 192.135.32.0 passive
distribute-list 14 out Serial0.1
distance 19 192.135.32.2 0.0.0.0 1
distance 19 131.154.1.9 0.0.0.0 2
redistribute static metric 512 200 250 100 1500
neighbor 192.135.32.2
neighbor 131.154.1.9
neighbor 131.154.1.5
passive-interface Ethernet0
passive-interface Serial0
passive-interface Serial0.2
passive-interface Serial0.3
passive-interface Serial0.4
!

```

```

ip domain-name infn.it
ip name-server 131.154.2.49
ip name-server 131.154.2.3
ip route 131.154.2.0 255.255.255.0 131.154.1.55
ip route 131.154.10.0 255.255.255.0 131.154.1.55
ip route 127.0.0.0 255.0.0.0 Null0 254
ip route 159.93.0.0 255.255.0.0 192.84.135.20
snmp-server community anabasi R0
snmp-server host 131.154.1.52 public

```

garr-gw-mi.infn.it

Current configuration:

```

!
version 10.0
service config
!
hostname garr-gw-mi
!
boot host garr-gw-mi-config.07sep94 192.84.138.1
boot network garr-gw-mi-config.07sep94 192.84.138.1
boot system gs3-k100-2.7 192.84.138.1
boot system gs3-k921-2.4 192.84.138.1
boot system rom
!
ip accounting-threshold 32000
!
decnet routing 31.659
decnet node-type routing-iv
decnet max-area 46
decnet conversion 47.0020
!
clns routing
!
interface Ethernet0
ip address 192.135.8.254 255.255.255.0 secondary
ip address 192.84.139.254 255.255.255.0 secondary
ip address 192.167.36.254 255.255.255.0 secondary
ip address 192.135.14.253 255.255.255.0 secondary
ip address 192.84.138.254 255.255.255.0

```

```

decnet cost 1
no decnet route-cache
decnet router-priority 48
no mop enabled
no clns route-cache
isis priority 48 level-1
!
interface Ethernet1
ip address 192.135.21.254 255.255.255.0
no mop enabled
no clns route-cache
isis metric 1 level-1
isis metric 1 level-2
!
interface Serial0
description --->CERN
ip address 192.12.193.49 255.255.255.252
bandwidth 64
delay 5000
shutdown
!
interface Serial1
description --->CILEA
ip address 131.175.243.2 255.255.255.0
ip accounting
bandwidth 2048
no clns route-cache
!
interface Serial2
description C-LAN
no ip address
ip accounting
encapsulation frame-relay
bandwidth 512
!
interface Serial2.1 point-to-point
description PVC-CNAF
ip unnumbered Ethernet1
ip router isis infn_milano
frame-relay interface-dlci 20

```

```

isis metric 1 level-1
isis metric 1 level-2
!
interface Serial2.2 point-to-point
description PVC-ROMA
ip unnumbered Ethernet1
ip router isis infn_milano
frame-relay interface-dlci 23
isis metric 1 level-1
isis metric 1 level-2
!
interface Serial2.3 point-to-point
description PVC-TORINO
ip unnumbered Ethernet1
ip router isis infn_milano
frame-relay interface-dlci 22
isis metric 1 level-1
isis metric 1 level-2
!
interface Serial3
description --->INFN-PV
ip unnumbered Ethernet0
bandwidth 64
shutdown
hold-queue 60 out
!
autonomous-system 137
!
router isis infn_milano
distance 90 ip
passive-interface Ethernet0
passive-interface Ethernet1
net 39.380f.1001.0001.0000.0000.0200.0000.0c01.4226.00
!
router rip
network 192.84.138.0 passive
distribute-list 4 in
default-metric 1
neighbor 192.84.138.210
passive-interface Serial1

```

```

passive-interface Serial2
passive-interface Serial2.1
passive-interface Serial2.2
passive-interface Serial2.3
passive-interface Ethernet1
!
router igrp 137
network 192.84.138.0
network 192.84.139.0
network 192.12.193.0
network 192.135.8.0
network 192.167.36.0
network 192.84.142.0
network 131.175.0.0
network 192.135.14.0
distribute-list 20 out Serial3
default-metric 64 4000 255 64 1500
redistribute static
redistribute rip
redistribute egp 513
neighbor 192.84.140.254
neighbor 192.12.193.41
passive-interface Serial2
passive-interface Serial2.1
passive-interface Serial2.2
passive-interface Serial2.3
passive-interface Ethernet0
passive-interface Ethernet1
!
router egp 513
network 192.12.193.0 passive
distribute-list 5 out
distribute-list 10 in
default-metric 5
redistribute igrp 137
neighbor 192.12.193.50
!
ip domain-name mi.infn.it
ip name-server 192.84.138.1
ip name-server 192.84.138.160

```

```

ip route 192.84.142.0 255.255.255.0 Serial3 110
ip route 131.175.0.0 255.255.0.0 Serial1
logging 131.175.1.38
snmp-server community
snmp-server community public R0
snmp-server host 192.84.138.160 public snmp

```

garr-gw-rm.infn.it

Current configuration:

```

!
version 10.0
service config
service password-encryption
!
hostname mp2rm1
!
boot host mp2rm1-clns 141.108.5.3
boot network mp2rm1-clns 141.108.5.3
boot system gs3-k_921-2_4 141.108.5.3
boot system rom
!
!
decnet routing 38.681
decnet node-type area
decnet max-area 46
!
clns routing
!
interface Ethernet0
description lan roma1
ip address 192.12.193.69 255.255.255.252 secondary
ip address 141.108.4.4 255.255.255.0 secondary
ip address 141.108.5.4 255.255.255.0
decnet cost 1
decnet router-priority 50
!
interface Ethernet1
description LAN x TEST C-LAN
ip address 192.135.22.4 255.255.255.0

```

```

isis metric 1 level-1
isis metric 1 level-2
!
interface Serial0
description Frame-Relay (IP) : LNGS, CNAF
ip address 192.135.33.3 255.255.255.0
encapsulation frame-relay
bandwidth 2048
decnet cost 1
frame-relay map decnet 38.251 302 broadcast
frame-relay map ip 192.135.33.1 300 broadcast
frame-relay map ip 192.135.33.2 301 broadcast
frame-relay map decnet 38.350 303 broadcast
!
interface Serial1
description TEST C-LAN
no ip address
encapsulation frame-relay
!
interface Serial1.1 point-to-point
description PVC-CNAF
ip unnumbered Ethernet1
ip router isis infn_roma
frame-relay interface-dlci 20
isis metric 1 level-1
isis metric 1 level-2
!
interface Serial1.2 point-to-point
description PVC-MILANO
ip unnumbered Ethernet1
ip router isis infn_roma
frame-relay interface-dlci 21
isis metric 1 level-1
isis metric 1 level-2
!
interface Serial2
description Centro ric. sper. R.Elena
ip address 193.204.109.2 255.255.255.0
bandwidth 64
!

```

```

interface Serial3
description Cineca ex-bis 21/v/93
ip address 192.12.193.26 255.255.255.252
bandwidth 1024
!
interface Serial4
description csata
ip address 192.12.193.29 255.255.255.252
bandwidth 512
!
interface Serial5
description NAPOLI
ip address 192.12.193.37 255.255.255.252
bandwidth 512
!
interface Serial6
description Cassino 64k
ip address 193.205.62.2 255.255.255.0
bandwidth 64
!
interface Serial7
no ip address
shutdown
!
interface Serial8
no ip address
shutdown
!
interface Serial9
no ip address
shutdown
!
router isis infn_roma
distance 90 ip
passive-interface Ethernet0
passive-interface Ethernet1
net 39.380f.1001.0001.0000.0000.0400.0000.0c01.0fe7.00
!
router rip
network 141.108.0.0

```

```

default-metric 4
passive-interface Serial1
passive-interface Serial1.1
passive-interface Serial1.2
!
router igrp 137
network 141.108.0.0 passive
network 192.12.193.0 passive
network 193.204.109.0 passive
network 192.135.33.0 passive
network 193.205.62.0
distribute-list 1 out Serial0
default-metric 7 2000 200 9 128
distance 100 141.108.5.14 0.0.0.0
redistribute static
neighbor 193.205.62.1
neighbor 192.135.33.2
neighbor 192.135.33.1
neighbor 192.12.193.30
neighbor 141.108.5.143
neighbor 141.108.5.30
neighbor 141.108.5.14
neighbor 192.12.193.38
neighbor 192.12.193.25
neighbor 192.12.193.70
neighbor 193.204.109.1
passive-interface Ethernet0
passive-interface Serial1
passive-interface Serial1.1
passive-interface Serial1.2
!
ip name-server 141.108.5.3
ip route 193.204.108.0 255.255.255.0 193.204.109.1
ip route 141.108.7.0 255.255.255.0 141.108.5.143
ip route 193.204.160.0 255.255.255.0 141.108.5.143
ip route 193.204.102.0 255.255.255.0 141.108.5.143
ip route 141.108.15.0 255.255.255.0 141.108.5.143
ip route 141.108.8.0 255.255.255.0 141.108.5.143
ip route 151.100.0.0 255.255.0.0 141.108.5.14
ip route 193.204.4.0 255.255.255.0 141.108.5.14

```

```

ip route 193.204.5.0 255.255.255.0 141.108.5.14
ip route 193.204.6.0 255.255.255.0 141.108.5.14
ip route 141.108.11.0 255.255.255.0 141.108.5.143
ip route 141.108.12.0 255.255.255.0 141.108.5.97
ip route 141.108.13.0 255.255.255.0 141.108.5.143
ip route 192.167.34.0 255.255.255.0 141.108.5.14
arp 141.108.5.50 aa00.0400.479f ARPA
snmp-server community public RO

```

cisco.clan.to.infn.it

Current configuration:

```

!
version 10.0
!
hostname cisco
!
boot system /usr/cisco/xx-k.100-2.7 192.135.19.1
!
clns routing
!
interface Ethernet0
ip address 192.84.137.250 255.255.255.0
isis metric 1 level-2
isis priority 100 level-1
isis priority 100 level-2
!
interface Ethernet1
no ip address
shutdown
isis metric 1 level-1
isis metric 1 level-2
isis priority 100 level-1
isis priority 100 level-2
!
interface Ethernet2
ip address 192.135.19.254 255.255.255.0
isis metric 1 level-2
!
interface Serial0

```

```

no ip address
encapsulation frame-relay
bandwidth 512
!
interface Serial0.1 point-to-point
description pvc-cnaf
ip unnumbered Ethernet2
ip router isis infn_torino
bandwidth 512
isis metric 1 level-1
isis metric 1 level-2
frame-relay interface-dlci 20 broadcast
!
interface Serial0.2 point-to-point
description pvc-milano
ip unnumbered Ethernet2
ip router isis infn_torino
bandwidth 512
isis metric 1 level-1
isis metric 1 level-2
frame-relay interface-dlci 22
!
interface Serial1
no ip address
!
router isis infn_torino
net 39.380f.1001.0001.0000.0000.0300.0000.0c05.d512.00
passive-interface Ethernet2
!
ip domain-name clan.to.infn.it
ip name-server 192.84.137.1
ip route 0.0.0.0 0.0.0.0 192.84.137.253

```

Niscn4.infn.it

Parte dello script per Niscn4

```

! This is an NCL script for the following DECnis
!
! Node:  NISCN4
! MOP Client Name:  NISCN4

```

```

! Hardware Type:  DECNIS-600
! Hardware Address:  08-00-2B-A2-0A-00
!
!
! Set up the Routing Module
!
create routing type L2Router , protocols { IP  ,  IS08473 }
set routing phaseiv address 31.11 , phaseiv prefix 47:0020:
set routing manual ip address 131.154.1.57
set routing manual L1 algorithm link state
set routing manual L2 algorithm link state
set routing Integrated ISIS protocols { IS08473  ,  IP }
!
! Create the routing network protocol ip
!
create routing network protocol ip
!
!
! Create Frame Relay Channel : W622-3-1
!
create modem connect line W622-3-1 communication port W622-3-1
set modem connect line W622-3-1 modem control -
    full, clock reflected, suppress test indicator TRUE
create frbs channel  W622-3-1 specification JOINT
set frbs channel  W622-3-1 physical line modem connect -
    line W622-3-1
!
!
!!
! Create a Frame Relay Connection: PVC-INFNGE
!
!
create frbs channel W622-3-1 connection PVC-INFNGE
set frbs channel W622-3-1 connection PVC-INFNGE preferred -
    dlci 20
create chdlc link PVC-INFNGE type synchronous
set chdlc link PVC-INFNGE lower layer entity frbs channel -
    W622-3-1 connection PVC-INFNGE, SLARP transmission timer 0
create routing circuit PVC-INFNGE  type  ppp
set routing circuit PVC-INFNGE network protocols {IP, IS08473}

```

```
set routing circuit PVC-INFNGE data link entity chdlc link -
    PVC-INFNGE
set routing circuit PVC-INFNGE L1 cost 1 , L2 cost 1
!
!
enable chdlc link PVC-INFNGE
enable routing circuit PVC-INFNGE
```

Appendice B

Configurazioni dei router per il test della maglia OSPF

Garr-gw.cnr.it

Garr-gw.cnr.it è configurato come backbone router, area border router dell'area 131.114.0.0, con un virtual link con cis1pi.pi.infn.it.

```
router ospf 1
network 192.65.131.0 0.0.0.255 area 131.114.0.0
network 192.12.193.0 0.0.0.255 area 0
area 131.114.0.0 virtual-link 192.84.133.3
distance 90
default-information originate metric 10 metric-type 1
redistribute igrp 137 metric 30 metric-type 1 subnets
```

Faeta.unipi.it

Faeta.unipi.it è configurato come internal router dell'area 131.114.0.0.

```
router OSPF 1
network 192.65.131.0 0.0.0.255 area 131.114.0.0
network 131.114.90.0 0.0.0.255 area 131.114.0.0
```

Garr-gw-bo.infn.it

Garr-gw-bo.infn.it è configurato come backbone router, area border router dell'area 131.154.0.0, con un virtual link con niscn4.infn.it.

```

router OSPF 1
network 131.154.1.0 0.0.0.255 area 131.154.0.0
network 192.12.193.0 0.0.0.255 area 0
distribute-list 42 in
distribute-list 43 out
redistribute igrp 137 metric n metric-type 2
redistribute bgp 137 metric m metric-type 2
default-information originate metric 20 metric-type 1
distance 90
area 131.154.0.0 virtual-link 131.154.1.57
ip ospf cost 48
!
router igrp 137
redistribute ospf 1 match internal external 1
!
access-list 42 deny 0.0.0.0
access-list 42 permit 0.0.0.0 255.255.255.255
!
access-list 43 deny 0.0.0.0 255.255.255.255
!

```

Cis1pi.pi.infn.it

Cis1pi.pi.infn.it è configurato backbone router e area border router dell'area 131.114.0.0, con un virtual link con garr-gw.cnr.it.

```

router OSPF 1
network 192.84.133.0 0.0.0.255 area 131.114.0.0
network 131.114.90.0 0.0.0.255 area 131.114.0.0
network 131.154.220.0 0.0.0.255 area 0
area 131.114.0.0 virtual-link 193.172.21.6
!
interface serial 1
ip ospf cost 48

```

Niscn4.infn.it

Niscn4.infn.it è configurato come backbone router, area border router dell'area 131.154.0.0, con un virtual link con garr-gw-bo.infn.it.

```
ncl>create routing control protocol OSPF-general type -
```

```

_ncl> OSPF general
ncl>create routing control protocol OSPF-area-0 type
_ncl> OSPF area
ncl>create routing control protocol OSPF-area-0 logical -
_ncl> circuit SERIAL-0 type point to point
ncl>create routing control protocol OSPF-area-B0 type
_ncl> OSPF area
ncl>create routing control protocol OSPF-area-B0 logical -
_ncl> circuit ETHERNET-0 type broadcast
ncl>create routing control protocol OSPF-area-0 logical -
_ncl> circuit VIRTUAL-1 type virtual

ncl>set routing control protocol OSPF-area-0 OSPF -
_ncl> general instance OSPF-general
ncl>set routing control protocol OSPF-area-0 OSPF area -
_ncl> id 0.0.0.0
ncl>set routing control protocol OSPF-area-0 logical -
_ncl> circuit SERIAL-0 circuit routing circuit SERIAL-0
ncl>set routing control protocol OSPF-area-B0 OSPF -
_ncl> general instance OSPF-general
ncl>set routing control protocol OSPF-area-B0 OSPF area -
_ncl> id 131.154.0.0
ncl>set routing control protocol OSPF-area-B0 logical -
_ncl> circuit ETHERNET-0 circuit routing circuit
_ncl> ETHERNET-0
ncl>set routing control protocol OSPF-area-0 logical -
_ncl> circuit VIRTUAL-1 ospf virtual neighbor router -
_ncl> id 192.135.33.1, ospf transit area OSPF-area-B0

ncl>enable routing control protocol OSPF-general
ncl>enable routing control protocol OSPF-area-0
ncl>enable routing control protocol OSPF-area-0 logical -
_ncl> circuit SERIAL-0
ncl>enable routing control protocol OSPF-area-B0
ncl>enable routing control protocol OSPF-area-B0 logical -
_ncl> circuit ETHERNET-0
ncl>enable routing control protocol OSPF-area-0 logical -
_ncl> circuit VIRTUAL-1

```

Appendice C

Reti dell'AS 137 divise per aree geografiche

In questa appendice sono elencate, suddivise per aree geografiche, tutte le reti IP dell'autonomous system 137. Quando possibile sono state indicate le supernet di alcuni indirizzi di classe C contigui con le maschere necessarie.

Area Emilia Romagna

Bologna:

130.136.0.0	255.255.0.0	Dip. Matematica
130.186.0.0	255.255.0.0	CINECA
131.154.0.0	255.255.0.0	INFN CNAF
137.204.0.0	255.255.0.0	BOLOGNA-ALMA-NET
192.94.70.0	255.255.254.0 (70,71)	CNR-BOLOGNA
192.107.60.0	255.255.252 (60-63, manca 60)	ENEA di Bologna
192.135.19.0	255.255.255.0	INFN CNAF Bologna
192.135.21.0	255.255.252.0 (20-23, manca 20)	INFN CNAF Bologna
192.135.32.0	255.255.254.0 (32-33)	INFN CNAF Bologna
192.167.161.0	255.255.224.0 (160-191) (manca 171-185,187,191)	CNR-BOLOGNA
193.42.96.0	255.255.255.0	Reg. Emilia Romagna
193.42.100.0	255.255.254.0	Reg. Emilia Romagna
193.204.120.0	255.255.255.0	CINECA
193.207.11.0	255.255.255.0	CINECA
193.207.12.0	255.255.252.0 (12-15, manca 15)	NETTUNO

Resto della regione:

155.185.0.0	255.255.0.0	Univ. Modena
192.167.3.0	255.255.255.0	IMGA CNR Modena
193.42.138.0	255.255.255.0	Ferrari (Maranello)
160.78.0.0	255.255.0.0	Univ. Parma
192.135.11.0	255.255.255.0	INFN Parma

192.84.144.0	255.255.255.0	INFN Ferrara
192.167.208.0	255.255.248.0 (208-215)	Univ. Ferrara

Le reti annunciate in questo modo sono 22 contro le 56 rilevate dalle tabelle di routing.

Area Piemonte

Torino:		
130.192.0.0	255.255.0.0	Politecnico Torino
163.162.0.0	255.255.0.0	IRI/STET
192.84.137.0	255.255.255.0	INFN Torino
192.167.203.0	255.255.255.0	CSI-Piemonte
193.42.224.0	255.255.252.0 (224-227)	Alenia Spazio
193.204.114.0	255.255.255.0	Ist. G. Ferraris
193.204.192.0	255.255.252.0 (192-195, manca 195)	CNR, CSTV

Le reti annunciate in questo modo sono 7 contro le 12 rilevate dalle tabelle di routing.

Area Lombardia

Milano:		
131.175.0.0	255.255.0.0	CILEA
146.133.0.0	255.255.0.0	CISE
149.132.0.0	255.255.0.0	Dip. Informatica
155.253.0.0	255.255.0.0	CNR-MILANO-NET
159.149.0.0	255.255.0.0	Univ. Milano
192.84.138.0	255.255.254.0 (138,139)	INFN Milano
192.84.140.0	255.255.254.0 (140,141)	INFN Milano
192.135.8.0	255.255.255.0	INFN Milano
192.135.14.0	255.255.255.0	INFN Milano
192.167.36.0	255.255.252.0 (36-39, manca 39)	Brera
193.205.16.0	255.255.255.0	Univ. Bocconi
193.205.40.0	255.255.248.0 (40-47, manca 41,44)	Univ. Sacro Cuore
193.205.48.0	255.255.255.0	Univ. Sacro Cuore
192.167.192.0	255.255.248.0 (192-199, manca 195-197)	Ist. San Raffaele
193.204.119.0	255.255.255.0	Cons. Citta' Ricerche
Resto della regione:		
192.84.142.0	255.255.255.0	INFN Pavia
193.204.32.0	255.255.252.0 (32-35, manca 33,35)	Univ. Pavia
192.167.16.0	255.255.240.0 (16-31, manca 19,22-27,31)	

		Univ. Brescia
193.204.248.0	255.255.252.0 (248-251, manca 251)	
		Univ. Bergamo

Le reti annunciate in questo modo sono 19 contro le 42 rilevate dalle tabelle di routing.

Area Veneto-Friuli

Friuli:		
140.105.0.0	255.255.0.0	Univ. Trieste
147.122.0.0	255.255.0.0	ISAS Trieste
158.110.0.0	255.255.0.0	Univ. Udine
Veneto:		
147.162.0.0	255.255.0.0	Univ. Padova
150.178.0.0	255.255.0.0	CNR LADSEB Padova
192.84.143.0	255.255.255.0	INFN Padova
192.84.148.0	255.255.255.0	INFN Legnaro
192.135.17.0	255.255.255.0	INFN Legnaro
192.135.29.0	255.255.255.0	INFN Legnaro
192.135.30.0	255.255.255.0	INFN Legnaro
157.138.0.0	255.255.0.0	Univ. Venezia
192.167.40.0	255.255.248.0 (40-47, manca 41,43-46)	CNR ISDGM Venezia
157.27.0.0	255.255.0.0	Univ. Verona

Le reti annunciate in questo modo sono 13 contro le 15 rilevate dalle tabelle di routing.

Area Liguria

Genova:		
130.251.0.0	255.255.255.0	Univ. Genova
150.145.0.0	255.255.255.0	CNR Genova
192.84.136.0	255.255.255.0	INFN Genova
192.135.18.0	255.255.255.0	INFN Genova
192.107.66.0	255.255.255.0	ENEA La Spezia

Le reti annunciate in questo modo sono 5 come 5 sono quelle rilevate dalle tabelle di routing.

Area Toscana-Umbria

Toscana:		
131.114.0.0	255.255.0.0	CNR Pisa

158.47.0.0	255.255.0.0	ENEL Pisa
192.12.192.0	255.255.255.0	CNR CNUCE Pisa
192.65.131.0	255.255.255.0	CNR CNUCE Pisa
192.84.133.0	255.255.255.0	INFN Pisa
192.84.155.0	255.255.255.0	Scuola Norm. Pisa
192.135.9.0	255.255.255.0	INFN Pisa
192.167.204.0	255.255.255.0	Scuola Norm. Pisa
149.139.0.0	255.255.0.0	CNR Firenze
150.217.0.0	255.255.0.0	CeSIT Univ. Firenze
192.84.145.0	255.255.255.0	INFN Firenze
192.84.146.0	255.255.254.0 (146,147)	INFN Firenze
193.42.139.0	255.255.255.0	Soprintendenza Fi
193.42.140.0	255.255.255.0	Ist. Prog. Economica
193.204.168.0	255.255.252.0 (168-171)	CNR, ARNO-MAN

Umbria:

141.250.0.0	255.255.0.0	Univ. di Perugia
193.205.4.0	255.255.252.0 (4-7, manca 5,7)	Univ. Siena

Le reti annunciate in questo modo sono 17 contro le 21 rilevate dalle tabelle di routing.

Area Lazio

Roma:

141.108.0.0	255.255.0.0	INFN Roma
150.146.0.0	255.255.0.0	CNR Roma
151.100.0.0	255.255.0.0	Univ. "La Sapienza"
160.80.0.0	255.255.0.0	Univ. "Tor Vergata"
168.202.0.0	255.255.0.0	FAO Roma
192.106.252.0	255.255.255.0	ESRIN Roma
192.107.70.0	255.255.254.0 (70,71)	ENEA Casaccia Roma
192.107.72.0	255.255.254.0 (72,73)	ENEA Casaccia Roma
192.107.85.0	255.255.255.0	ENEA Centrale Roma
192.107.86.0	255.255.255.0	ENEA Centrale Roma
192.107.91.0	255.255.255.0	ENEA-DISP Roma
192.156.213.0	255.255.255.0	Oss. Astr. Roma
192.167.2.0	255.255.255.0	Ist.Biologia Molecolare
192.167.34.0	255.255.255.0	Ist.Sup. di Sanita'
192.167.224.0	255.255.255.0	CNR Roma-Montelibretti
193.43.36.0	255.255.255.0	FAO Headquarters, Roma
193.204.4.0	255.255.252.0 (4-7, manca 7)	Cons.App.Super calcolo
193.204.102.0	255.255.255.0	ASI Roma
193.204.108.0	255.255.254.0 (8,9)	Ist. Fisioterapici
193.204.116.0	255.255.255.0	Consorzio Roma Ricerche
193.204.160.0	255.255.252.0 (160-163, manca 163)	Fac.Scienze III Univ.
193.204.208.0	255.255.240.0 (208-215, manca 212)	

Fondazione Ugo Bordoni
CILEA Roma
ASI Roma

Frascati:

146.48.0.0	255.255.0.0	CNR Frascati
192.84.127.0	255.255.255.0	INFN Frascati
192.84.128.0	255.255.252.0 (128-131)	INFN Frascati
192.107.51.0	255.255.255.0	ENEA Frascati
192.107.52.0	255.255.254.0 (52,53)	ENEA Frascati
193.204.104.0	255.255.252.0 (104-107, manca 107)	ESRIN Frascati
		ESRIN Frascati
193.204.224.0	255.255.252.0 (224-227)	ESRIN Frascati
193.205.62.0	255.255.254.0 (62,63)	Univ. di Cassino

Le reti annunciate in questo modo sono 31 contro le 54 rilevate dalle tabelle di routing.

Area Abruzzo-Marche

Abruzzo:

192.84.135.0	255.255.255.0	INFN Gran Sasso
192.135.35.0	255.255.255.0	INFN Gran Sasso
192.84.154.0	255.255.255.0	INFN L'Aquila
192.150.194.0	255.255.254.0 (194,195)	Univ. L'Aquila
192.150.196.0	255.255.255.0	Univ. L'Aquila
192.167.12.0	255.255.252.0 (12-15)	Univ. Chieti
193.204.28.0	255.255.255.0	Univ. Chieti
193.204.100.0	255.255.254.0 (100,101)	Mario Negri Chieti
193.204.1.0	255.255.255.0	Collurania-Teramo

Marche:

193.204.8.0	255.255.252.0 (8-11, manca 11)	Univ. Camerino
193.205.2.0	255.255.255.0	Univ. Urbino

Le reti annunciate in questo modo sono 11 contro le 18 rilevate dalle tabelle di routing.

Area Puglia-Basilicata

Puglia:

138.66.0.0	255.255.0.0	CSATA Bari
192.106.246.0	255.255.255.0	CSATA Bari
193.204.48.0	255.255.252.0 (48-51)	Politecnico Bari
193.204.176.0	255.255.255.0	Univ. Bari
192.135.10.0	255.255.255.0	INFN Bari

192.195.110.0	255.255.255.0	CNR IESI Bari
192.84.152.0	255.255.255.0	INFN Lecce
193.204.64.0	255.255.255.0	Univ. Lecce
193.204.68.0	255.255.252.0 (68-71)	Univ. Lecce

Basilicata:

192.106.234.0	255.255.255.0	ASI-CGS Matera
193.204.16.0	255.255.254.0 (16,17)	Univ. Basilicata

Le reti annunciate in questo modo sono 11 contro le 18 rilevate dalle tabelle di routing.

Area Campania-Calabria

Campania:

140.164.0.0	255.255.0.0	CNR Napoli
143.225.0.0	255.255.0.0	CRIAI Napoli
192.132.34.0	255.255.255.0	CISED Univ. Napoli
192.133.28.0	255.255.255.0	CISED Univ. Napoli
192.135.165.0	255.255.255.0	CHEMNA Univ. Napoli
192.167.9.0	255.255.255.0	Ist. Navale Napoli
192.167.11.0	255.255.255.0	Univ. Napoli
192.167.33.0	255.255.255.0	Univ. Napoli
192.55.101.0	255.255.255.0	Univ. Napoli
192.84.134.0	255.255.255.0	INFN Napoli
192.84.149.0	255.255.255.0	INFN Napoli
156.14.0.0	255.255.0.0	CIRA Caserta
192.41.218.0	255.255.255.0	Univ. di Salerno
192.135.15.0	255.255.255.0	INFN CCSEM
138.41.0.0	255.255.0.0	CRAI

Calabria:

160.97.0.0	255.255.0.0	Univ. Calabria
192.167.201.0	255.255.255.0	Univ. Calabria

Le reti annunciate in questo modo sono 17 come 17 sono quelle rilevate dalle tabelle di routing.

Area Sicilia

147.163.0.0	255.255.0.0	Univ. Palermo
151.97.0.0	255.255.0.0	Univ. Catania
192.84.150.0	255.255.255.0	INFN Catania
193.42.136.0	255.255.254.0 (136,137)	CoRiMMe Catania
193.204.118.0	255.255.255.0	Catania Ricerche
193.204.240.0	255.255.252.0 (240-243, manca 241)	CNR Catania

193.204.244.0	255.255.254.0 (244,245)	CNR Catania
192.167.96.0	255.255.240.0 (96-11, manca 97,99,102-109)	Univ. Messina
193.204.96.0	255.255.255.0	CNR Messina
193.204.124.0	255.255.255.0	ASI Trapani
192.84.151.0	255.255.255.0	INFN LNS
192.106.247.0	255.255.255.0	Cent. Ric. Sicilia
193.204.24.0	255.255.252.0 (24-27)	CRES Sicilia

Le reti annunciate in questo modo sono 13 contro le 24 rilevate dalle tabelle di routing.

Area Sardegna

156.148.0.0	255.255.0.0	CRS4 Cagliari
192.84.132.0	255.255.255.0	INFN Cagliari
192.84.153.0	255.255.255.0	INFN Cagliari
192.146.242.0	255.255.255.0	Univ. Cagliari
192.160.156.0	255.255.255.0	Univ. Cagliari
192.167.8.0	255.255.255.0	Oss. Astr. Cagliari
192.167.128.0	255.255.224.0 (128-159, manca 137-140,149,150,152-154)	Univ. Cagliari
193.204.110.0	255.255.254.0 (110-111)	CNR Sassari
193.205.8.0	255.255.255.0	Univ. Sassari

Le reti annunciate in questo modo sono 9 contro le 27 rilevate dalle tabelle di routing.

Riduzione numero indirizzi annunciati

Con queste aree è stato possibile ridurre il numero di indirizzi da 309 a 175 con una riduzione del 43%.

Appendice D

Reti dell'AS 137 divise per aree di affinità d'uso

In questa appendice sono elencate, suddivise per aree di affinità d'uso, tutte le reti IP dell'autonomous system 137. Quando possibile sono state indicate le supernet di alcuni indirizzi di classe C contigui con le maschere necessarie.

Area UNIVERSITÀ-1

147.163.0.0	255.255.0.0	Univ. Palermo
151.97.0.0	255.255.0.0	Univ. Catania
192.167.96.0	255.255.240.0 (96-11, manca 97,99,102-109)	Univ. Messina
193.42.136.0	255.255.254.0 (136,137)	CoRimMe Catania
192.84.151.0	255.255.255.0	INFN LNS
192.106.247.0	255.255.255.0	Cent. Ric. Sicilia
193.204.24.0	255.255.252.0 (24-27)	CRES Sicilia
193.204.96.0	255.255.255.0	CNR Messina
193.204.118.0	255.255.255.0	Catania Ricerche
193.204.124.0	255.255.255.0	ASI Trapani
193.204.240.0	255.255.252.0 (240-243, manca 241)	CNR Catania
193.204.244.0	255.255.254.0 (244,245)	CNR Catania

Le reti annunciate in questo modo sono 12 contro le 24 rilevate dalle tabelle di routing.

Area UNIVERSITÀ-2

138.66.0.0	255.255.0.0	CSATA Bari
192.84.152.0	255.255.255.0	INFN Lecce
192.106.234.0	255.255.255.0	ASI-CGS Matera

192.106.246.0	255.255.255.0	CSATA Bari
192.195.110.0	255.255.255.0	CNR IESI Bari
193.204.48.0	255.255.252.0 (48-51)	Politecnico Bari
193.204.64.0	255.255.255.0	Univ. Lecce
193.204.68.0	255.255.252.0 (68-71)	Univ. Lecce
193.204.176.0	255.255.255.0	Univ. Bari

Le reti annunciate in questo modo sono 9 contro le 15 rilevate dalle tabelle di routing.

Area UNIVERSITÀ-3

130.136.0.0	255.255.0.0	Dip. Matematica Bologna
130.186.0.0	255.255.0.0	CINECA
137.204.0.0	255.255.0.0	BOLOGNA-ALMA-NET
140.105.0.0	255.255.0.0	Univ. Trieste
147.122.0.0	255.255.0.0	ISAS Trieste
147.162.0.0	255.255.0.0	Univ. Padova
150.178.0.0	255.255.0.0	CNR LADSEB Padova
155.185.0.0	255.255.0.0	Univ. Modena
157.27.0.0	255.255.0.0	Univ. Verona
157.138.0.0	255.255.0.0	Univ. Venezia
158.110.0.0	255.255.0.0	Univ. Udine
160.78.0.0	255.255.0.0	Univ. Parma
192.94.70.0	255.255.254.0 (70,71)	CNR-BOLOGNA
192.107.60.0	255.255.252 (60-63, manca 60)	ENEA di Bologna
192.146.242.0	255.255.255.0	Univ. Cagliari
192.160.156.0	255.255.255.0	Univ. Cagliari
192.167.3.0	255.255.255.0	IMGA CNR Modena
192.167.8.0	255.255.255.0	Oss. Astr. Cagliari
192.167.40.0	255.255.248.0 (40-47, manca 41,43-46)	CNR ISDGM Venezia
192.167.128.0	255.255.192.0 (128-159, manca 137-140,149,150, 152-154,171-185,187,191)	128-159 Univ. Cagliari 160-191 CNR-BOLOGNA
192.167.208.0	255.255.248.0 (208-215)	Univ. Ferrara
193.42.96.0	255.255.255.0	Reg. Emilia Romagna
193.42.100.0	255.255.254.0	Reg. Emilia Romagna
193.42.138.0	255.255.255.0	Ferrari (Maranello)
193.204.8.0	255.255.252.0 (8-11, manca 11)	Univ. Camerino
193.204.120.0	255.255.255.0	CINECA
193.205.2.0	255.255.255.0	Univ. Urbino
193.207.11.0	255.255.255.0	CINECA
193.207.12.0	255.255.252.0 (12-15, manca 15)	NETTUNO

Le reti annunciate in questo modo sono 29 contro le 83 rilevate dalle tabelle di routing.

Area UNIVERSITÀ-4

130.192.0.0	255.255.0.0	Politecnico Torino
130.251.0.0	255.255.255.0	Univ. Genova
131.175.0.0	255.255.0.0	CILEA
146.133.0.0	255.255.0.0	CISE
149.132.0.0	255.255.0.0	Dip. Informatica
150.145.0.0	255.255.255.0	CNR Genova
155.253.0.0	255.255.0.0	CNR-MILANO-NET
159.149.0.0	255.255.0.0	Univ. Milano
163.162.0.0	255.255.0.0	IRI/STET
192.107.66.0	255.255.255.0	ENEA La Spezia
192.167.16.0	255.255.240.0 (16-31, manca 19,22-27,31)	Univ. Brescia
192.167.192.0	255.255.248.0 (192-199, manca 195-197)	Ist. San Raffaele
192.167.203.0	255.255.255.0	CSI-Piemonte
193.42.224.0	255.255.252.0 (224-227)	Alenia Spazio
193.204.32.0	255.255.252.0 (32-35, manca 33,35)	Univ. Pavia
193.204.114.0	255.255.255.0	Ist. G. Ferraris
193.204.119.0	255.255.255.0	Cons. Citta' Ricerche
193.204.192.0	255.255.252.0 (192-195, manca 195)	CNR, CSTV
193.204.248.0	255.255.252.0 (248-251, manca 251)	Univ. Bergamo
193.205.16.0	255.255.255.0	Univ. Bocconi
193.205.40.0	255.255.248.0 (40-47, manca 41,44)	Univ. Sacro Cuore
193.205.48.0	255.255.255.0	Univ. Sacro Cuore

Le reti annunciate in questo modo sono 22 contro le 46 rilevate dalle tabelle di routing.

Area CNR-1

131.114.0.0	255.255.0.0	CNR Pisa
149.139.0.0	255.255.0.0	CNR Firenze
150.217.0.0	255.255.0.0	CeSIT Univ. Firenze
158.47.0.0	255.255.0.0	ENEL Pisa
192.12.192.0	255.255.255.0	CNR CNUCE Pisa
192.65.131.0	255.255.255.0	CNR CNUCE Pisa
192.84.133.0	255.255.255.0	INFN Pisa
192.84.155.0	255.255.255.0	Scuola Norm. Pisa
192.135.9.0	255.255.255.0	INFN Pisa
192.167.204.0	255.255.255.0	Scuola Norm. Pisa
193.42.139.0	255.255.255.0	Soprintendenza Fi

193.42.140.0	255.255.255.0	Ist. Prog. Economica
193.204.110.0	255.255.254.0 (110-111)	CNR Sassari
193.204.168.0	255.255.252.0 (168-171)	CNR, ARNO-MAN
193.205.4.0	255.255.252.0 (4-7, manca 5,7)	Univ. Siena
193.205.8.0	255.255.255.0	Univ. Sassari

Le reti annunciate in questo modo sono 16 contro le 21 rilevate dalle tabelle di routing.

Area INFN-1

138.41.0.0	255.255.0.0	CRAI
140.164.0.0	255.255.0.0	CNR Napoli
143.225.0.0	255.255.0.0	CRAI Napoli
156.14.0.0	255.255.0.0	CIRA Caserta
160.97.0.0	255.255.0.0	Univ. Calabria
192.41.218.0	255.255.255.0	Univ. di Salerno
192.55.101.0	255.255.255.0	Univ. Napoli
192.84.134.0	255.255.255.0	INFN Napoli
192.84.149.0	255.255.255.0	INFN Napoli
192.84.150.0	255.255.255.0	INFN Catania
192.132.34.0	255.255.255.0	CISED Univ. Napoli
192.133.28.0	255.255.255.0	CISED Univ. Napoli
192.135.15.0	255.255.255.0	INFN CCSEM
192.135.165.0	255.255.255.0	CHEMNA Univ. Napoli
192.167.9.0	255.255.255.0	Ist. Navale Napoli
192.167.11.0	255.255.255.0	Univ. Napoli
192.167.33.0	255.255.255.0	Univ. Napoli
192.167.201.0	255.255.255.0	Univ. Calabria
193.204.16.0	255.255.254.0 (16,17)	Univ. Basilicata

Le reti annunciate in questo modo sono 18 contro le 19 rilevate dalle tabelle di routing.

Area INFN-2

192.84.135.0	255.255.255.0	INFN Gran Sasso
192.84.154.0	255.255.255.0	INFN L'Aquila
192.135.10.0	255.255.255.0	INFN Bari
192.135.35.0	255.255.255.0	INFN Gran Sasso
192.150.194.0	255.255.254.0 (194,195)	Univ. L'Aquila
192.150.196.0	255.255.255.0	Univ. L'Aquila
192.167.12.0	255.255.252.0 (12-15)	Univ. Chieti
193.204.1.0	255.255.255.0	Collurania-Teramo
193.204.28.0	255.255.255.0	Univ. Chieti
193.204.100.0	255.255.254.0 (100,101)	Mario Negri Chieti

Le reti annunciate in questo modo sono 10 contro le 15 rilevate dalle tabelle di routing.

Area INFN-3

131.154.0.0	255.255.0.0	INFN CNAF
156.148.0.0	255.255.0.0	CRS4 Cagliari
192.84.132.0	255.255.255.0	INFN Cagliari
192.84.144.0	255.255.252.0 (144-147)	144 INFN Ferrara
		145-147 INFN Firenze
192.84.153.0	255.255.255.0	INFN Cagliari
192.135.11.0	255.255.255.0	INFN Parma
192.135.19.0	255.255.255.0	INFN CNAF Bologna
192.135.21.0	255.255.252.0 (20-23, manca 20)	INFN CNAF Bologna
192.135.32.0	255.255.254.0 (32-33)	INFN CNAF Bologna

Le reti annunciate in questo modo sono 9 contro le 15 rilevate dalle tabelle di routing.

Area INFN-4

192.84.136.0	255.255.248.0 (136-143, manca 143)	
	.136 INFN Genova, .137 INFN Torino,	
	.138 .139 INFN Milano .140 .141 INFN Milano	
	.142 INFN Pavia	
192.135.8.0	255.255.255.0	INFN Milano
192.135.14.0	255.255.255.0	INFN Milano
192.135.18.0	255.255.255.0	INFN Genova
192.167.36.0	255.255.252.0 (36-39, manca 39)	Brera Milano

Le reti annunciate in questo modo sono 5 contro le 13 rilevate dalle tabelle di routing.

Area INFN-5

192.84.143.0	255.255.255.0	INFN Padova
192.84.148.0	255.255.255.0	INFN Legnaro
192.135.17.0	255.255.255.0	INFN Legnaro
192.135.29.0	255.255.255.0	INFN Legnaro
192.135.30.0	255.255.255.0	INFN Legnaro

Le reti annunciate in questo modo sono 5 come 5 sono quelle rilevate dalle tabelle di routing.

Area ROMA

Questa area non è stata spezzata perché tutte queste reti utilizzano il router dell'INFN Roma per connettersi al GARR.

141.108.0.0	255.255.0.0	INFN Roma
141.250.0.0	255.255.0.0	Univ. di Perugia
146.48.0.0	255.255.0.0	CNR Frascati
150.146.0.0	255.255.0.0	CNR Roma
151.100.0.0	255.255.0.0	Univ. "La Sapienza"
160.80.0.0	255.255.0.0	Univ. "Tor Vergata"
168.202.0.0	255.255.0.0	FAO Roma
192.84.127.0	255.255.255.0	INFN Frascati
192.84.128.0	255.255.252.0 (128-131)	INFN Frascati
192.106.252.0	255.255.255.0	ESRIN Roma
192.107.51.0	255.255.255.0	ENEA Frascati
192.107.52.0	255.255.254.0 (52,53)	ENEA Frascati
192.107.70.0	255.255.254.0 (70,71)	ENEA Casaccia Roma
192.107.72.0	255.255.254.0 (72,73)	ENEA Casaccia Roma
192.107.85.0	255.255.255.0	ENEA Centrale Roma
192.107.86.0	255.255.255.0	ENEA Centrale Roma
192.107.91.0	255.255.255.0	ENEA-DISP Roma
192.156.213.0	255.255.255.0	Oss. Astr. Roma
192.167.2.0	255.255.255.0	Ist.Biologia Molecolare
192.167.34.0	255.255.255.0	Ist.Sup. di Sanita'
192.167.224.0	255.255.255.0	CNR Roma-Montelibretti
193.43.36.0	255.255.255.0	FAO Headquarters, Roma
193.204.4.0	255.255.252.0 (4-7, manca 7)	Cons.App.Super calcolo
193.204.102.0	255.255.255.0	ASI Roma
193.204.104.0	255.255.252.0 (104-107, manca 107)	ESRIN Frascati
193.204.108.0	255.255.254.0 (8,9)	Ist. Fisioterapici
193.204.116.0	255.255.255.0	Consorzio Roma Ricerche
193.204.160.0	255.255.252.0 (160-163, manca 163)	Fac.Sienze III Univ.
193.204.208.0	255.255.240.0 (208-215, manca 212)	Fondazione Ugo Bordoni
		CILEA Roma
		ASI Roma
193.204.224.0	255.255.252.0 (224-227)	ESRIN Frascati
193.205.62.0	255.255.254.0 (62,63)	Univ. di Cassino

Le reti annunciate in questo modo sono 31 contro le 54 rilevate dalle tabelle di routing.

Riduzione numero indirizzi annunciati

Con queste aree è stato possibile ridurre il numero di indirizzi da 309 a 166 con una riduzione del 46%.

Bibliografia

- [Bgp4ospf] Varadhan K., Hares S., Rekhter Y., *draft-ietf-idrp-bgp4ospf-interact-08: BGP4/IDRP for IP—OSPF interaction*. OARnet, Merit, NSFnet, T.J. Watson Research Center, IBM corp., Ottobre 1994.
- [Black] Black Uyless, *Data link protocols*. Prentice Hall, New Jersey, 1993.
- [Catnip] Ullman Robert, *draft-ietf-catnip-base-02: CATNIP Common Architecture for Next-generation Internet Protocol*. Lotus Development Corporation, Gennaio 1994.
- [Cisco] Cisco, *Internetworking technology overview*. Cisco system, Inc. Menlo Park, California, 1993.
- [Comer] Comer Douglas, *Internetworking con TCP/IP*. Gruppo editoriale Jackson - Prentice Hall International, Italia, 1992.
- [DECnetV] Martin James, Leben Joe, *DECnet phase V: an OSI implementation*. Digital Press, Bedford, Massachusetts, 1992.
- [Hedrick] Hedrick Charles, *An introduction to IGRP*. Rutgers, the State University of New Jersey, 1989.
- [IPng-crit] Kastenholz F., Partridge C., *draft-kastenholz-ipng-criteria-01: Technical criteria for choosing IP: the next generation (IPng)*. Marzo 1994.
- [IPng-rec] Bradner S., Mankin A., *draft-ipng-recommendation-01: The recommendation for the IP next generation protocol*. Harvard University, NRL, Ottobre 1994.
- [IPng-view] Hinden Robert, *draft-hinden-ipng-overview-00: IP next generation overview*. Sun Microsystems Inc., Ottobre 1994.

- [IPv6-spec] Hinden Robert, *draft-ietf-ipng-ipv6-spec-01: Internet Protocol, version 6 (IPv6)*. Sun Microsystems Inc., Ottobre 1994.
- [IPv6-trans] Gilligan Robert, *draft-gilligan-ipv6-sit-overview-00: Simple Internet Transition (SIT) overview*. Sun Microsystems Inc., Ottobre 1994.
- [ISO 8473] *ISO/IEC 8473: Information technology - Protocol for providing the Connectionless-mode Network Service*. ISO International Organization for Standardization, 1988, 1994 (seconda edizione).
- [ISO 9542] *ISO/IEC 9542: Information technology - End System to Intermediate System routing Exchange Protocol for use with ISO 8473*. ISO International Organization for Standardization, 1988.
- [ISO 10589] *ISO/IEC 10589: Information technology - Telecommunications and information exchange between systems - Intermediate system to Intermediate system intradomain routing information exchange protocol for use in conjunction with the protocol for providing the Connectionless-mode Network Service (ISO 8473)*. ISO International Organization for Standardization, Gennaio 1992.
- [ISO 10747] *ISO/IEC 10747: Information technology - Telecommunications and information exchange between systems - Protocol for exchange of inter-domain routing information among intermediate system to support forwarding of ISO 8473 PDUs*. ISO International Organization for Standardization, Ottobre 1993.
- [Perlman] Perlman Radia, *Interconnections bridges and routers*. Addison Wesley Professional computing series, Reading, Massachusetts, Luglio 1992.
- [RFC 768] Postel J.B., *RFC 768: User Datagram Protocol*. ISI, Agosto 1980.
- [RFC 791] Postel J.B., *RFC 791: Internet Protocol*. Settembre 1981.
- [RFC 793] Information Sciences Institute - University of Southern California, *RFC 793 Transmission Control Protocol*. Settembre 1981.
- [RFC 826] Plummer D.C., *RFC 826: Ethernet Address Resolution Protocol: Or converting network protocol addresses to 48.bit Ethernet address for transmission on Ethernet hardware*. Novembre 1982.

- [RFC 904] Mills D.L., *RFC 904: Exterior Gateway Protocol formal specification*. Aprile 1984.
- [RFC 950] Mogul J.C. Postel J.B., *RFC 950: Internet Standard Subnetting Procedure*. Agosto 1985.
- [RFC 1058] Hedrick Charles, *RFC 1058: Routing Information Protocol*. Rutgers University, Giugno 1988.
- [RFC 1060] Reynolds J.K. Postel J.B., *RFC 1060 Assigned numbers*. Marzo 1990.
- [RFC 1195] Callon Ross W., *RFC 1195: Use of OSI IS-IS for Routing in TCP/IP and Dual Environments*. Digital Equipment Corporation, Dicembre 1990.
- [RFC 1347] Callon Ross, *RFC 1347: TCP and UDP with Bigger Addresses*. Digital Equipment Corporation, Giugno 1992.
- [RFC 1403] Varadhan K., *RFC 1403: BGP OSPF Interaction*. OARnet, Gennaio 1993.
- [RFC 1519] Fuller V., Li T., Yu J., Varadhan K., *RFC 1519: Classless Inter-Domain Routing (CIDR): an Address Assignment and Aggregation Strategy*. BARRnet, Cisco, MERIT, OARnet, Settembre 1993.
- [RFC 1561] Piscitello David, *RFC 1561: Use of ISO CLNP in TUBA environments*. Core Competence Inc., Dicembre 1993.
- [RFC 1583] Moy John, *RFC 1583: OSPF version 2*. Proteon Inc., Marzo 1994.
- [RFC 1654] Rekhter Y., Li T., *RFC 1654: A Border Gateway Protocol 4 (BGP-4)*. T.J. Watson Research Center, IBM Corp., cisco Systems, Luglio 1994.
- [RFC 1676] Ghiselli A., Salomoni D., Vistoli C., *RFC 1676: INFN Requirements for an IPng*. INFN/CNAF, Agosto 1994.
- [Sipp-ipae] Gilligan R., Nordmark E., Hinden B., *draft-ietf-sipp-ipae-transition-00: IPAE: The SIPP interoperability and transition mechanism*. Sun Microsystems Inc., Novembre 1993.

- [Sipp-spec] Deering Steve, *draft-ietf-sipp-spec-01: Simple Internet Protocol Plus (SIPP) specification (128 bit address version)*. Xerox PARC, Luglio 1994.
- [Sipp-view] Deering S., Francis P., Govindan R., *draft-ietf-sip-overview-00: Simple Internet Protocol Plus: overview of routing and addressing extension to SIP*. Xerox PARC-Bellcore, Ottobre 1993.
- [Tuba-trans] Piscitello David, *draft-ietf-tuba-transition-00: Transition plan for TUBA/CLNP*. Core Competence Inc., Marzo 1994.
- [Wilder] Wilder Floyd, *A guide to the TCP/IP protocol suite*. Artech House London, 1993.

Indice analitico

A

ARP: 31
AS boundary router: 48, 78
AS external link advertisement: 57
Address: 20
Adiacenza: 50
Algoritmo: 45
Application layer: 9, 12
Area border router: 47, 48
Aree OSPF: 46
Aree integrated IS-IS: 59
Aree: 144
Autonomous system 137: 13, 78, 89, 143
Autonomous system: 46, 78, 89

B

BGP: 41, 80, 89
Backbone OSPF: 47
Backbone router: 47, 48
Backbone: 144, 147
Boundary intermediate system: 41

C

C-LAN: 93
CLNP: 19, 20, 135
CNAF: 79
Classi IP: 22
Client-server: 7
Cluster address: 138
Complessità: 45

Complete sequence number PDU: 67
CompleteSNPInterval: 67

D

DECnet/OSI: 144
DLCI: 95
DNS: 136
Data link layer: 11
Database description packet: 52
Datagramma IP: 26
Designated IS: 62
Designated router: 49, 75
Dijkstra: 45
Distance vector: 34
Dual stack: 136

E

EGP: 40
EIGRP: 38
ES-IS: 32

F

File transfer: 6
Flooding: 44
Forwarding database: 61, 76
Frame relay switch: 100
Frame relay: 94
Frammentazione: 27
Ftp site: 6
Ftp: 6

G

GARR: 13,78,143

GIX: 81

Garr-gw: 84

H

HDLC: 94

Hardware layer: 8

Hello packet OSPF: 51

HelloInterval: 51, 118

HelloTimer: 64, 65, 118

Holddowns: 37

Hop: 34

I

IDRP: 41

IETF: 132

IGRP: 37, 83

INFNet: 12, 78

INFN: 12, 93

Indirizzamento gerarchico: 24

Indirizzamento piatto: 24

Indirizzi CLNP: 23

Indirizzi IP: 21

Indirizzo: 20

Integrated IS-IS, test: 93

Integrated IS-IS: 39, 58, 84, 144

Integrated routing: 68, 117

Interdomain: 40, 78

Internal router: 47, 48

Internet: 3, 131

internet: 3

Internet layer: 8

Internetting: 4

Internetworking: 3

Intestazione SIPP: 139

Intradomain: 35, 78

IP: 9, 19

IPng: 18, 131, 136

IPv6: 136, 141, 145

IS-IS: 39

ISO: 10

L

LAN IS to IS hello PDU: 63, 64

LSPGenerationInterval: 62, 65, 118

LSRefreshTimer: 53, 118

Layer ISO: 11

Layer TCP/IP: 8

Level 2 subdomain: 60

Link state PDU: 65, 66

Link state acknowledgments packet:
54

Link state database: 60

Link state request packet OSPF:
53

Link state routing protocol: 43

Link state update packet OSPF: 53

Link state, algoritmo: 45

Link state: 35

Livello 1: 59

Livello 2: 60

Login remoto: 7

Loop: 34

M

Mail: 4

Manual area addresses: 59

Multicast address: 138

Multicast: 7

N

NET: 23

NSAP address: 23

Neighbor greeting: 31

Network interface layer: 8

Network layer protocol: 19
Network layer: 11
Network link advertisement: 55

O

OSI reference model: 11
OSPF, test: 109
OSPF: 38, 46, 86, 144

P

PPP: 94
PVC: 95
Pacchetti IS-IS: 123
Pacchetti integrated IS-IS: 63, 126
Pacchetti OSPF: 51, 118
Pacchetto CLNP: 28
Partial sequence number PDU: 67, 68
Partition designated level 2 IS: 62
Percorso asimmetrico: 91
Physical layer: 11
Ping: 107
Point-to-point IS to IS hello PDU: 65
Poison reverse: 37
Posta: 4
Presentation layer: 12
Produzione: 96
Pseudonodo: 49, 75

Q

Quality of Service (QoS): 30

R

RFC: 77
RIP: 36, 82
Redistribute: 89

Remote login: 7
Route poisoning: 38
Routeing function: 61
Router link advertisement: 55
Router: 81
Routing: 30
Routing SIPP: 140
Routing domain: 59, 89
Routing table: 31, 49, 76
RxmtInterval: 52

S

S.I.N.: 145
SIPP: 18, 136
SIT: 140
Sequence number: 44
Session layer: 11
Ship In the Night (S.I.N.): 68, 117
Source routing: 140
Split horizon: 37
Stub area: 47, 50
Subnetting: 22
Subnetwork: 22
Summary link advertisement: 57
Supernetwork: 23

T

TCP/IP: 4
TCP: 10
TUBA: 18, 135
Time To Live (TTL): 27
Timestamp: 44
Topologia OSPF: 46
Topologia integrated IS-IS: 59
Topological database: 49, 76
Transit LAN: 55
Transizione SIPP: 140
Transizione TUBA: 136
Transport layer: 9, 11

Trasferimento file: 6
Triggered update: 38
Tunneling: 141
Type of Service (TOS): 27

U
Unicast address: 137
Virtual link: 47, 72