

Quesito 1 (punti 5). Dato il sistema descritto dalla seguente rappresentazione insiemistica di assegnazione delle risorse:

$$\begin{aligned} P &= \{P_1, P_2, P_3, P_4\} \\ R &= \{R_1^1, R_2^1, R_3^1, R_4^2\} \\ E &= \{P_1 \rightarrow R_1, P_1 \rightarrow R_3, P_2 \rightarrow R_2, \\ &\quad R_1 \rightarrow P_2, R_2 \rightarrow P_3, R_3 \rightarrow P_2, R_4 \rightarrow P_3, R_4 \rightarrow P_4\} \end{aligned}$$

si verifichi, anche tramite analisi del relativo grafo di allocazione delle risorse, se il sistema si trovi in condizione di stallo. Successivamente si verifichi se l'eventuale ulteriore richiesta di accesso alla risorsa R_3 da parte del processo P_3 possa cambiare la situazione precedentemente rilevata.

Quesito 2 (punti 8). Si descrivano in dettaglio l'esecuzione e l'effetto complessivo del comando di *shell* UNIX e GNU/Linux:

```
sort < in | tail -5 > out
```

sapendo che:

- `sort` è un comando di *shell* che ordina in modo lessicografico righe (una stringa per riga) ricevute in ingresso
- `tail -n` è un comando di *shell* che restituisce le ultime n righe di un *file* ricevuto in ingresso
- `in` è un *file* contenente la prima riga del testo di questo quesito, scritto una parola per riga
- `out` non è un comando e non esiste un *file* con tale nome nella *directory* corrente.

Si ricordi che, nell'ordinamento lessicografico le maiuscole precedono le minuscole.

Quesito 3 (punti 3).

[3.A] Quale tra le seguenti caratteristiche costituisce un criterio valido di valutazione di una politica di ordinamento di processi:

- 1: la capacità di trattare anche processi di lunga durata
- 2: il numero di processi completati per unità di tempo
- 3: il numero di processi in esecuzione per unità di tempo
- 4: il numero di processi in attesa di essere eseguiti.

[3.B] Quale tra le seguenti affermazioni concernenti la politica di ordinamento Round-Robin è corretta:

- 1: il tempo di attesa è sempre maggiore del tempo di risposta
- 2: il tempo di attesa è sempre minore del tempo di risposta
- 3: il tempo di attesa è sempre uguale al tempo di risposta
- 4: il tempo di attesa e il tempo di risposta non hanno alcun legame prefissato.

[3.C] Quale tra le seguenti politiche di ordinamento in generale minimizza il tempo medio di attesa dei processi:

- 1: FCFS
- 2: Round-Robin con valutazione dell'attributo di priorità dei processi
- 3: Round-Robin senza valutazione dell'attributo di priorità dei processi
- 4: Shortest Job First.

Quesito 4 (punti 8). Si discutano le differenze che intercorrono tra gli algoritmi denominati leaky bucket e token bucket sul piano algoritmico e delle finalità. Si illustri il comportamento concreto di tali algoritmi, evidenziando le differenti scelte da essi eventualmente fatte, a fronte del seguente profilo d'ingresso:

<i>intervallo</i>	<i>durata (u.t.)</i>	<i>pacchetti/u.t.</i>
1	5	5
2	5	—
3	5	5

ipotizzando che i pacchetti in ingresso abbiano tutti dimensione costante e che per entrambi gli algoritmi valgano le seguenti condizioni: (1) stessa capacità di *bucket* di 15 unità, (2) stesso ciclo di 1 u.t., (3) stessa situazione iniziale di coda di ingresso e *bucket* vuoti.

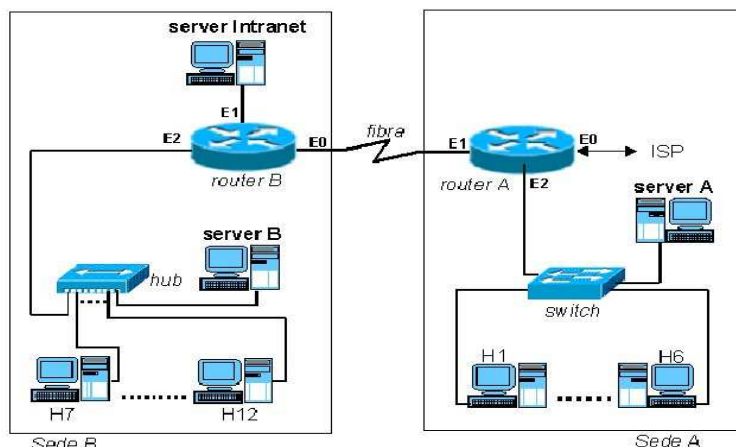


Figura 1: Architettura della rete interna dell'Azienda.

Quesito 5 (punti 8). Lo schema logico riportato in figura 1 rappresenta la rete dati di una piccola Agenzia di Servizi, suddivisa tra due sedi, indicate in figura come A e B, connesse tra loro tramite un collegamento in fibra ottica di proprietà aziendale. Tutti i dispositivi di rete operano a 100 Mbps. Nell'architettura logica della rete aziendale sono presenti tre server:

- il server intranet, che offre servizi di *application server* e di accreditamento agli utenti aziendali di entrambe le sedi
- il server A, che offre servizi di condivisione di *file* alle 6 postazioni di lavoro della Sede A
- il server B, che offre servizi di condivisione di *file* alle 6 postazioni di lavoro della Sede B.

Il traffico generato da ciascun utente dell'Agenzia è costituito essenzialmente da tre componenti: (1) componente Internet diretta genericamente verso lo spazio Web esterno, di valore massimo costante e pari a 200 Kbps per ogni utente; (2) componente Intranet, che fa capo all'apposito server Intranet; (3) componente utente-server che fa capo al server di condivisione *file* della sede di appartenenza dell'utente.

Sotto l'ipotesi che la componente (2) Intranet e la componente (3) utente-server assumano valori massimi uguali, si determini il flusso teorico massimo nel caso peggiore in cui tutti gli utenti vogliano operare simultaneamente.

L'Agenzia accede ad Internet mediante un unico indirizzo IP statico fornitole direttamente dal proprio ISP. Al proprio interno, invece, l'Agenzia decide di condividere gli indirizzi privati di una intera classe C (192.168.1.0), sfruttando la funzione di traduzione degli indirizzi (NAT, *Network Address Translation*) realizzata all'interno del router A. Sotto queste ipotesi, si proponga una suddivisione degli indirizzi utili in sottoreti di uguale ampiezza (dunque medesima *subnet mask*) e massima dimensione possibile. Si compili poi una tabella che indichi, per tutti i dispositivi della rete dell'Agenzia, l'indirizzo IP associato, la *subnet mask* corrispondente ed il *default gateway*.

Soluzione 1 (punti 5). La figura 2 riporta la versione grafica della rappresentazione insiemistica data, dalla quale non risulta alcun percorso chiuso, ciò che ci consente di concludere che nessun processo del sistema si trovi attualmente in situazione di stallo. Qualora, in questa situazione, il processo P_3 richiedesse però accesso alla risorsa R_3 , le cose si complicherebbero

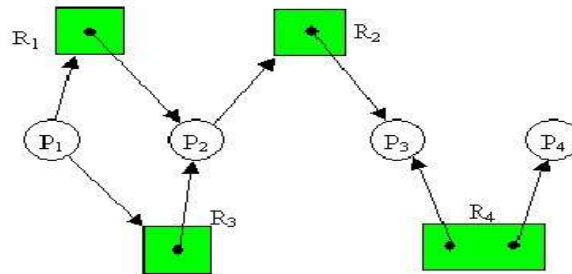


Figura 2: Grafo di allocazione delle risorse per il sistema in esame.

immediatamente:

- si produrrebbe il percorso chiuso $P_2 \rightarrow R_2 \rightarrow P_3 \rightarrow R_3$
- tale percorso coinvolgerebbe solo risorse unarie, ciò comportando stallo per i processi P_2 e P_3
- anche P_1 entrerebbe in stato di stallo trovandosi in attesa di una risorsa permanentemente indisponibile per stallo precedente di P_2
- il processo P_4 potrebbe invece completare normalmente la propria esecuzione, infine rilasciando una molteplicità della risorsa R_4 , senza però che questo produca alcun effetto positivo per il sistema in esame.

Soluzione 2 (punti 8). Sappiamo che in UNIX, e conseguentemente in GNU/Linux, ogni processo conosce 3 descrittori di *file* primordiali, denominati *standard input* (0), *standard output* (1) e *standard error* (2). Sappiamo poi che i simboli $<$ e $>$ denotano la redirezione dei descrittori 0 ed 1 rispettivamente verso le entità nominate successivamente. Sappiamo infine che il simbolo $|$ denota una *pipe*, ossia uno *pseudofile* dotato di descrittori distinti di scrittura e lettura, posto in comune tra due comandi di *shell*.

Il comando specificato dal quesito verrà conseguentemente interpretato dalla *shell* come 2 sottocomandi (2 processi), il primo dei quali (`sort < in`) scrive, per effetto dell'operatore *pipe* ($|$), la propria uscita nello *pseudofile* di ingresso del secondo (`tail -5 > out`).

Per effetto della redirezione $<$, il primo sottocomando troverà il proprio descrittore 0 associato al *file* `in`, di cui il quesito ci indica il contenuto. L'esecuzione di tale sottocomando scriverà nello *pseudofile* di ingresso al secondo sottocomando il riordinamento lessicografico delle stringhe del *file* `in`. Il secondo sottocomando preleverà le ultime 5 righe di tale *pseudofile* e le scriverà nel *file* `out`. Poichè tale *file* non esiste nella *directory* corrente del processo *shell*, esso verrà creato appositamente. Alla fine di tale esecuzione, un nuovo *file* di nome `out` sarà stato creato, contenente le seguenti 5 righe:

```
e
in
l'effetto
l'esecuzione
shell
```

Soluzione 3 (punti 3).

La risposta al quesito 3.A è: 2.

La risposta al quesito 3.C è: 4.

La risposta al quesito 3.B è: 4.

Nota Bene: Nel contesto del quesito 3.B, la risposta attesa era effettivamente la 4, anche se il testo del quesito potrebbe dare adito ad alcune perplessità. Ricordiamo le relative definizioni:

- **tempo di attesa:** somma degli intervalli di tempo in cui il processo, pur essendo noto all'ordinatore (*scheduler*), non progredisce la propria esecuzione;

- **tempo di risposta:** l'intervallo di tempo intercorrente tra l'istante di arrivo del processo e l'istante in cui esso inizia ad avanzare per la prima volta.

Indipendentemente dalla politica di ordinamento, tali definizioni indicano che il tempo di attesa è dato dalla somma del tempo di risposta più gli eventuali successivi tempi di inattività che il processo subisce fino al proprio completamento.

La formulazione della domanda e l'articolazione delle risposte portava ad escludere a priori le risposte 1, 2 e 3. La risposta 4 voleva invece escludere i legami assoluti citati nelle risposte precedenti, ed in tal senso era da ritenersi valida. Quanto sopra richiamato però potrebbe essere legittimamente interpretato come un "legame prefissato" tra le due componenti temporali, quindi inducendo a considerare errata anche la risposta 4. Per questo motivo, nella correzione dei compiti abbiamo ritenuto valide sia la risposta 4 che l'illustrazione dell'osservazione sopra riportata.

Soluzione 4 (punti 8). Gli algoritmi detti leaky bucket e token bucket sono entrambi utilizzati per il controllo di congestione del traffico su rete, e dunque sono di interesse del livello rete (*network, IP*) nella gerarchia dei protocolli ISO/OSI e TCP/IP. L'algoritmo leaky bucket produce un flusso di uscita costante o nullo a fronte di flussi di ingresso variabili ed irregolari. Concettualmente, può essere rappresentato come avente una coda di ingresso (*bucket*) a capacità finita fissata ed un'attività di estrazione dalla coda a periodo fissato e costante, di N pacchetti/periodo. (Nel caso posto dal quesito conosciamo l'ampiezza del periodo, 1 u.t., mentre, per la soluzione, possiamo assumere $N = 1$.) Il traffico in ingresso si accoda, ma solo finché vi trova spazio, altrimenti viene perduto, mentre il traffico di uscita viene prodotto a velocità costante finché vi siano pacchetti in coda.

L'algoritmo token bucket può essere visto come una variante del precedente, nella quale ciascun pacchetto in ingresso viene avviato all'uscita solo consumando un gettone (*token*) appositamente prelevato dal *bucket*. L'algoritmo produce tali gettoni con periodo fissato e costante, N gettoni/periodo, e li accumula nel *bucket* fino alla sua capacità massima. In assenza di traffico in ingresso l'algoritmo accumula dunque capacità di uscita, il cui andamento pertanto non è più costante, ma può assumere un profilo con picchi. Altra differenza importante dal leaky bucket è che questo algoritmo non ammette la perdita di pacchetti in ingresso, ma solo di gettoni in presenza di *bucket* pieno e traffico in ingresso nullo. (Il quesito ci ha indicato l'ampiezza del periodo di produzione gettoni, 1 u.t., mentre, per uniformità di confronto, assumeremo $N = 1$.)

Le figure 3 e 4 ci mostrano come, nella situazione proposta dal quesito, i due algoritmi si comportino in maniera analoga riguardo al flusso di uscita prodotto nell'intervallo 0 – 14 (i.e. flusso costante di 1 pacchetto/u.t.), ma in maniera assai diversa rispetto alla perdita di pacchetti: la versione leaky bucket perde 24 pacchetti complessivi sui 50 pervenuti, mentre la versione token bucket non ne perde alcuno, a condizione naturalmente di avere una coda di ampiezza ≥ 35 .

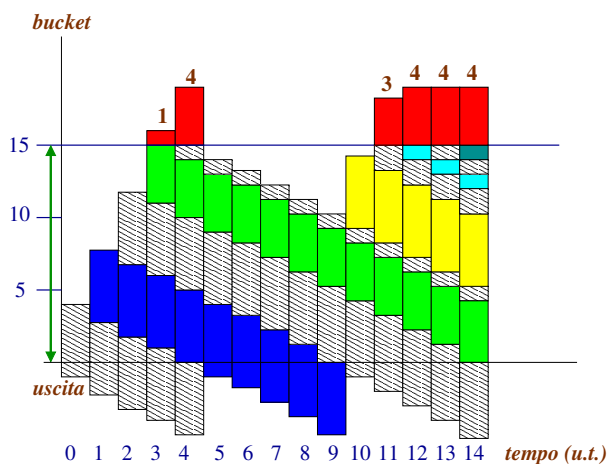


Figura 3: Comportamento dell'algoritmo leaky bucket.

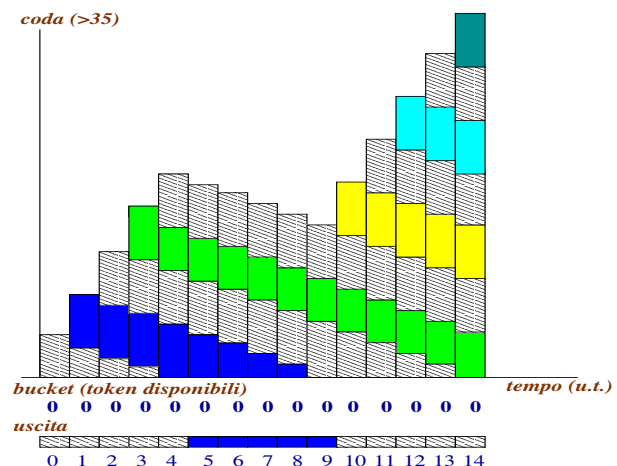


Figura 4: Comportamento dell'algoritmo token bucket.

Soluzione 5 (punti 8). Come mostrato in figura 5, la configurazione della rete aziendale consta chiaramente di 5 sottoreti interne così composte:

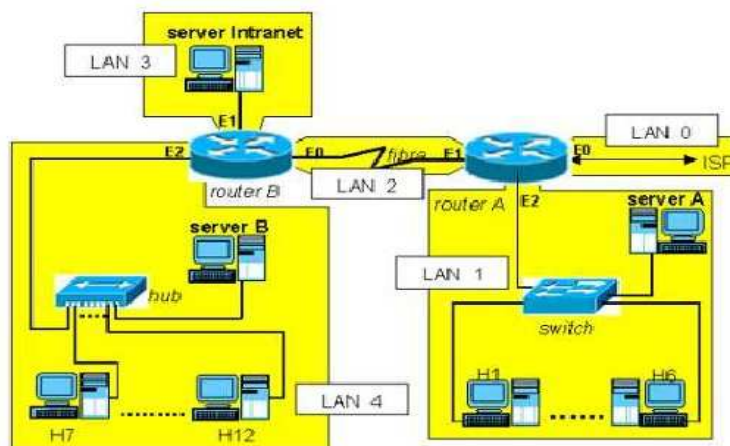


Figura 5: Articolazione delle sottoreti interne dell'Agenzia.

LAN_0	rete esterna, di collegamento tra l'ISP e la rete aziendale
LAN_1	rete della Sede A, facente capo all'interfaccia $E2$ del router A, composta da un unico dominio di diffusione e da 8 domini di collisione, uno per ogni porta dello switch
LAN_2	rete di interconnessione tra le due sedi, facente capo rispettivamente all'interfaccia $E1$ del router A ed all'interfaccia $E0$ del router B, composta da un unico dominio di diffusione ed un unico dominio di collisione
LAN_3	settore protetto della rete contenente i servizi Web interni, accessibili solo agli utenti interni dell'Agenzia, composto da un unico dominio di diffusione ed un unico dominio di collisione
LAN_4	rete della Sede B, facente capo all'interfaccia $E2$ del router B, composta da un unico dominio di diffusione ed un unico dominio di collisione.

Il calcolo dei flussi nel caso peggiore può essere determinato considerando la situazione delle sole reti interne ($LAN_1 - LAN_4$). Per prima cosa occorre individuare i flussi utili descritti dal quesito. Detti:

- X flusso di dati prodotto dal generico utente della Sede A ($H1 - H6$) per la componente facente capo al server Intranet, il cui valore è numericamente uguale alla componente del traffico utente-server verso il server A
- Y flusso di dati prodotto dal generico utente della Sede B ($H7 - H12$) per la componente facente capo al server Intranet, il cui valore è numericamente uguale alla componente del traffico utente-server verso il server B

si ottiene facilmente la distribuzione rappresentata in figura 6 in cui il traffico utile è indicato in grassetto, ed il traffico effettivo (ove diverso da quello utile) è racchiuso in un ovale.

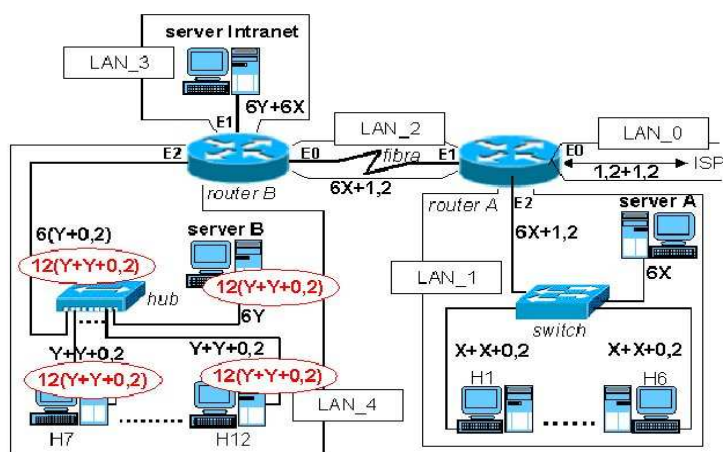


Figura 6: Determinazione dei flussi utili di caso peggiore.

Ogni ramo della rete può trasferire al massimo 100 Mbps. Deriviamo pertanto le seguenti condizioni:

$$X + X + 0,2 = 100 \text{ Mbps} \quad (H1 - H6) \quad (1)$$

$$6X = 100 \text{ Mbps} \quad (\text{collegamento al server A}) \quad (2)$$

$$6 \times (X + 0,2) = 100 \text{ Mbps} \quad (\text{collegamento tra switch ed interfaccia E2 del router A}) \quad (3)$$

$$12 \times (Y + Y + 0,2) = 100 \text{ Mbps} \quad (H7 - H12, \text{ tutte le porte dell'hub}) \quad (4)$$

$$6X + 6Y = 100 \text{ Mbps} \quad (\text{collegamento tra interfaccia E1 del router B e server Intranet}) \quad (5)$$

$$6X + 1,2 = 100 \text{ Mbps} \quad (\text{collegamento tra le sedi}) \quad (6)$$

Dalla condizione 1 ricaviamo il valore $X = 49,9 \text{ Mbps}$.

Dalla condizione 2 ricaviamo il valore $X = 16,6 \text{ Mbps}$, più restrittivo del precedente e che dunque lo sostituisce.

Dalla condizione 3 ricaviamo il valore $X = 16,46 \text{ Mbps}$, ancora più restrittivo del precedente, e che dunque lo sostituisce.

Dalla condizione 4 ricaviamo il valore $Y = 4,06 \text{ Mbps}$.

La condizione 5 ci serve da verifica. Sostituendo X ed Y in essa otteniamo: $6X + 6Y = 6 \times 16,46 + 6 \times 4,06 = 132,3$ valore maggiore della banda massima e dunque inammissibile. Dobbiamo pertanto ritenere che il collo di bottiglia della rete aziendale risieda proprio nel segmento che induce la condizione 5, ossia il collegamento tra l'interfaccia $E1$ del router B ed il server Intranet.

Mantenendo inalterato il rapporto tra X e Y , ed imponendo la condizione 5, possiamo ricavare i valori ammissibili di X ed Y procedendo, ad esempio, come segue:

$$X' : X = (X' + Y') : (X + Y) \quad (7)$$

$$Y' : Y = (X' + Y') : (X + Y) \quad (8)$$

$$X' + Y' = \frac{100}{6} = 16,66 \quad (\text{dalla condizione 5}) \quad (9)$$

$$X + Y = 16,46 + 4,06 = 20,53 \quad (10)$$

$$X' = 16,46 \times \frac{16,66}{20,53} = 13,36 \quad (11)$$

$$Y' = 4,06 \times \frac{16,66}{20,53} = 3,3 \quad (12)$$

dove X' ed Y' soddisfano la condizione 5 per costruzione. La condizione 6 di fatto duplica la condizione 3, per cui può essere trascurata. Ne segue che i valori massimi teorici di caso peggiore del traffico nei vari punti della rete valgono rispettivamente:

$$X_{MAX} = 13,36 \text{ Mbps} \quad Y_{MAX} = 3,3 \text{ Mbps}$$

Per quanto riguarda la ripartizione degli indirizzi IP della rete aziendale, per la rete LAN_0 non sono necessarie considerazioni particolari, dato che la sua configurazione di natura statica è in generale a carico dell'ISP. Per gli altri dispositivi abbiamo a disposizione un'intera rete di classe C di indirizzi riservati 192.168.1.0. L'Agenzia, come sappiamo è composta da 4 sottoreti interne, $LAN_1 - LAN_4$. Il dato di progetto da cui partire per l'individuazione degli indirizzi delle sottoreti è che esse debbano avere tutte la medesima dimensione, e dunque la stessa *subnet mask*. Dovendo creare 4 sottoreti utili e sapendo che il numero di sottoreti utili S è dato dall'equazione $S = 2^N - 2$, dove N è il numero di *bit* impiegati per designarle, otteniamo che il primo valore intero appropriato è $S = 6 \geq 4$, da cui ricaviamo $N = 3$. La parte nodo dell'indirizzo IP sarà allora ampia $8 - 3 = 5 \text{ bit}$, con la quale possiamo denotare fino a $2^5 - 2 = 30$ nodi. Questo valore è più che sufficiente per le sottoreti aziendali, per cui possiamo considerare risolto il problema. La *subnet mask* interna all'Agenzia, comune a tutte le sue sottoreti interne, varrà allora:

255.255.255.224 in binario: 11111111.11111111.11111111.11100000 ed in forma compatta: /27

192	168	1	0	
11000000	10101000	00000001	000	000000
parte di rete (classe C)			sottoreti interne	nodi di sottorete

Tabella 1: Base dell'insieme di indirizzi di classe C riservati ad uso interno dell'Agenzia.

In tabella 2 ricapitoliamo le caratteristiche degli indirizzi IP delle 4 sottoreti utili, mentre in tabella 3 riportiamo una possibile assegnazione di indirizzi IP per i dispositivi aziendali interni.

parte di rete	sottorete interna	nodo	destinazione indirizzo		
11000000. 10101000. 00000001.	001	00000	192.168.1.32	riservato per la I sottorete	
11000000. 10101000. 00000001.	001	00001	192.168.1.33	utilizzabile per 1° nodo nella I sottorete	
...					
11000000. 10101000. 00000001.	001	11110	192.168.1.62	utilizzabile per ultimo nodo nella I sottorete	
11000000. 10101000. 00000001.	001	11111	192.168.1.63	riservato al <i>broadcast</i> nella I sottorete	
11000000. 10101000. 00000001.	010	00000	192.168.1.64	riservato per la II sottorete	
11000000. 10101000. 00000001.	010	00001	192.168.1.65	utilizzabile per 1° nodo nella II sottorete	
...					
11000000. 10101000. 00000001.	010	11110	192.168.1.94	utilizzabile per ultimo nodo nella II sottorete	
11000000. 10101000. 00000001.	010	11111	192.168.1.95	riservato al <i>broadcast</i> nella II sottorete	
11000000. 10101000. 00000001.	011	00000	192.168.1.96	riservato per la III sottorete	
11000000. 10101000. 00000001.	011	00001	192.168.1.97	utilizzabile per 1° nodo nella III sottorete	
...					
11000000. 10101000. 00000001.	011	11110	192.168.1.126	utilizzabile per ultimo nodo nella III sottorete	
11000000. 10101000. 00000001.	011	11111	192.168.1.127	riservato al <i>broadcast</i> nella III sottorete	
11000000. 10101000. 00000001.	100	00000	192.168.1.128	riservato per la IV sottorete	
11000000. 10101000. 00000001.	100	00001	192.168.1.129	utilizzabile per 1° nodo nella IV sottorete	
...					
11000000. 10101000. 00000001.	100	11110	192.168.1.158	utilizzabile per ultimo nodo nella IV sottorete	
11000000. 10101000. 00000001.	100	11111	192.168.1.159	riservato al <i>broadcast</i> nella IV sottorete	
11000000. 10101000. 00000001.	101	00000	192.168.1.160	riservato per la V sottorete	
11000000. 10101000. 00000001.	101	00001	192.168.1.161	utilizzabile per 1° nodo nella V sottorete	
...					
11000000. 10101000. 00000001.	101	11110	192.168.1.190	utilizzabile per ultimo nodo nella V sottorete	
11000000. 10101000. 00000001.	101	11111	192.168.1.191	riservato al <i>broadcast</i> nella V sottorete	
11000000. 10101000. 00000001.	110	00000	192.168.1.192	riservato per la V sottorete	
11000000. 10101000. 00000001.	110	00001	192.168.1.193	utilizzabile per 1° nodo nella V sottorete	
...					
11000000. 10101000. 00000001.	110	11110	192.168.1.222	utilizzabile per ultimo nodo nella V sottorete	
11000000. 10101000. 00000001.	110	11111	192.168.1.223	riservato al <i>broadcast</i> nella V sottorete	

Tabella 2: Ripartizione degli indirizzi nelle 6 sottoreti interne disponibili.

rete	dispositivo	IP address	subnet mask	default gateway
LAN ₁	H1	192.168.1.33	255.255.255.224	192.168.1.62
	H2	192.168.1.34	/27	192.168.1.62
	H3	192.168.1.35	/27	192.168.1.62
	H4	192.168.1.36	/27	192.168.1.62
	H5	192.168.1.37	/27	192.168.1.62
	H6	192.168.1.38	/27	192.168.1.62
	server A	192.168.1.61	/27	192.168.1.62
	router A porta E2	192.168.1.62	/27	—
LAN ₂	router A porta E1	192.168.1.93	/27	—
	router B porta E0	192.168.1.94	/27	—
LAN ₃	server Intranet	192.168.1.125	/27	192.168.1.126
	router B porta E1	192.168.1.126	/27	—
LAN ₄	H7	192.168.1.129	/27	192.168.1.158
	H8	192.168.1.130	/27	192.168.1.158
	H9	192.168.1.131	/27	192.168.1.158
	H10	192.168.1.132	/27	192.168.1.158
	H11	192.168.1.133	/27	192.168.1.158
	H12	192.168.1.134	/27	192.168.1.158
	server B	192.168.1.157	/27	192.168.1.158
	router B porta E2	192.168.1.158	/27	—

Tabella 3: Una possibile attribuzione di indirizzi IP interni alle postazioni della rete Aziendale.