

Optimal Allocation of Rates in Guaranteed Service Networks

Aniruddha Diwan

Motorola India Electronics Ltd

Bagmane Techpark, C. V. Raman Nagar, Bangalore 560093, India

E-mail: aniruddha@motorola.com

Joy Kuri

Department of Electronic Systems Engineering, Indian Institute of Science, Bangalore 560012, India

E-mail: kuri@cedt.iisc.ernet.in, <http://www.cedt.iisc.ernet.in/people/kuri/>

Sugata Sanyal

School of Technology & Computer Science

Tata Institute of Fundamental Research, Homi Bhabha Road, Mumbai 400005, India

E-mail: sanyals@gmail.com, <http://www.tifr.res.in/~sanyal>

Keywords: guaranteed service networks, IntServ, optimal rate allocation, rate reservation

Received: October 6, 2011

We examine the problem of rate allocation in Guaranteed Services networks by assigning a cost corresponding to a rate, and examining least cost allocations. We show that the common algorithm of allocating the same rate to a connection on all links along its path (called the *Identical Rates* algorithm henceforth) is, indeed, a least cost allocation in many situations of interest. This finding provides theoretical justification for a commonly adopted strategy, and is a contribution of this paper. However, it may happen that the single rate required is not available on some links of the route. The second contribution of this paper is an explicit expression for the optimal rate vector for the case where *Identical Rates* is infeasible. This leads to an algorithm called *General Rates* that can admit a connection by allocating possibly different rates on the links along a route. Finally, we simulate the *General Rates* algorithm in a dynamic scenario, and observe that it can provide, at best, marginally improved blocking probabilities. Our conclusion is that the performance benefits provided by *General Rates* are not compelling enough, and the simpler *Identical Rates* suffices in practice.

Povzetek: V članku je analiziran način alociranja v servisnih omrežjih.

1 Introduction

In this paper we consider the problem of rate allocation in the context of the Guaranteed (G) Service framework [1]. The G Service framework enables the receiver of an individual data flow to compute the rate (say R) that it needs to reserve, so that its QoS requirements are satisfied. Thus, a rate vector $(R, R, \dots, R)$ is reserved, with the vector having as many elements as there are links on the path between the source and the destination. R is a function of the traffic parameters and QoS requirements of the flow, as well as network characteristics such as the number of hops on the path, the scheduling policy employed at each hop and the propagation delay along each hop.

Now suppose that the rate R is not available on some links of the route. When this happens, the Connection Admission Control (CAC) module is usually programmed to block the incoming flow. But, in fact, sufficient bandwidth may not be available only on a few links. There might exist other rate vectors that satisfy the delay constraint and

the available link bandwidth constraints. Hence, the CAC module may end up blocking a connection request unnecessarily.

This scenario provides the starting point for the work described in this paper. We are motivated to examine the rate allocation problem in the hope that we may avoid unnecessary connection blocking and thereby achieve a bigger admission region than possible when a single rate (*viz.*, R) is reserved on all nodes along the path. To this end, we assign a *cost* to bandwidth allocation at each hop and look for minimum cost allocations.

Rate allocation in Guaranteed Services networks has a long and rich history [2, 3, 4, 5, 6, 7, 8, 9, 10, 11, 12, 13, 14, 15]. However, a common theme that runs through the vast literature is that of *identical* rate assignment at each node. Our approach to this problem is novel because we do not start by making this tacit assumption. We pose and answer the questions:

(a) Are there circumstances in which allocating unequal

rates is beneficial?

- (b) What criteria can one use to compare alternate rate allocations? What is the *best* rate allocation that can be done according to a chosen criterion?

To the best of our knowledge, the existing literature does not raise these questions.

We begin by listing in Section 3 a few important results in the areas of arrival and service curves, as well as end-to-end delays in Guaranteed Services. These results are used subsequently in our analysis.

2 Related Work

As we have mentioned above, the question that we address does not seem to have received attention in the literature. In this section, we discuss some papers that analyze questions that are related to, but distinct from, the object of study in our paper.

[4] surveys basic mechanisms to support QoS guarantees. Scheduling and buffer management are the basic techniques to achieve QoS. The authors discuss RSVP, the end-to-end resource reservation protocol that was suggested to ensure delay guarantees in the Internet. It is stated explicitly that resource reservation proceeds on the basis of a *single* reservation rate allocated at each node.

The general class of schedulers called Latency Rate (LR) servers is discussed in [6]. Latency and allocated rate are the two parameters that influence the delay experienced by a packet served by an LR server. To analyze the delay experienced by a connection passing through a network of servers, the authors assume that the *same* rate is allocated to a connection at every node on its path.

[10] studies a general network with multiple multihop connections sharing the nodes in the network, and focuses on obtaining end-to-end delay bounds, as well as understanding issues related to stability. Each session is associated with a minimum backlog clearing rate. With this rate given, the authors obtain formulae yielding end-to-end delay bounds. However, the issue of *how to allocate the service rate at each node* was not addressed; it was assumed that the rates are *given*. In this paper, we are concerned precisely with the question of how to allocate rates so that some pre-defined objective is met. Thus, our work addresses a novel aspect of the problem.

In the same way, [7] allows different allocated rates at different nodes, but the interest is in obtaining end-to-end delay bounds for a *given vector* of rates. On the other hand, our concern is with *how to allocate* the vector of rates, so that target delay guarantees can be provided while keeping aggregate resource consumption (*i.e.*, bandwidth consumption) low. Thus, the problem considered in our work is orthogonal to that considered in this paper.

Yet another paper considering related problems is [11]. Each node on the path of a connection uses a “rate-controlled service discipline,” but *traffic shapers* are employed at each hop. Thus, the network model considered

in this paper is different from the one in our work, because we do not have traffic shapers before each node.

To summarize, we find two threads in the literature. One group of papers assumes that the *same* rate is allocated at each hop on the path. The second group allows different rates to be allocated at each hop, but the focus is on finding formulae for end-to-end delay; the issue of assessing and comparing different possible rate vector allocations has received hardly any attention. In this paper, our concern is with precisely this question.

3 IntServ and Guaranteed Services

To ensure that the QoS requirements are guaranteed, the G Service requires each node on the route of the connection to dedicate an appropriate rate R and a buffer space B to the flow during its holding time. Also, a data flow is required to declare its flow characteristics at the time of connection setup and each node is required to declare its network characteristics. If the flow is admitted, only those packets of the flow that conform to the characteristics are guaranteed the QoS requirements. The rate R and the buffer space B required to guarantee the QoS requirements are a function of the flow characteristics and the network characteristics. In this paper, we consider only R as the resource requirement.

The flow characteristics are specified in terms of the token bucket parameters which are a triplet $(b, \rho, \hat{p})$, where b is the token bucket size, ρ is the token accumulation rate, and $\hat{p}$ is the peak rate of the flow. Node i specifies its network characteristics in terms of the parameters C_i and D_i ; typically, C_i and D_i approximate the departure of the service provided by node i from the “fluid” model of service, in which the flow effectively sees a dedicated channel of bandwidth R between source and receiver.

The arrival process has an envelope $A^{in}(\tau)$ that is given by

$$A^{in}(\tau) = \min\{L + \hat{p}\tau, b + \rho\tau\}, \tau \geq 0 \quad (1)$$

where L represents the packet size. When an identical rate R is reserved at each, the service curve for a tandem of nodes $1, 2, \dots, N$ is given by [16],

$$S_N(\tau) = \left[\left(\tau - \sum_{i=1}^N D_i \right) R - \sum_{i=1}^N C_i \right]^+ \quad (2)$$

C_i represents the Maximum Transmission Unit (MTU) on the link leaving the i^{th} node on the path, and D_i is the maximum non-preemption delay at the i^{th} node. The end-to-end delay bound is the maximum horizontal distance between the arrival curve $A^{in}(\tau)$ and the service curve $S_N(\tau)$. With this, the following delay bound can be eas-

ily obtained¹.

$$D^{\text{bound}} = \begin{cases} \frac{(b - L^{\max})(\hat{p} - R)}{R(\hat{p} - \rho)} + \frac{L^{\max}}{R} + \sum_{i=1}^N \left[\frac{C_i}{R} + D_i \right], & \hat{p} > R, \\ \frac{L^{\max}}{R} + \sum_{i=1}^N \left[\frac{C_i}{R} + D_i \right], & \hat{p} \leq R. \end{cases} \quad (3)$$

where $L^{\max}$ denotes the maximum packet size on this connection. If D^{reqd} is the worst-case end-to-end delay that is acceptable for the packets of a flow, $D^{\text{bound}} \leq D^{\text{reqd}}$ gives the minimum rate R that has to be allocated at each node on the route of the flow. We denote this minimum rate by R_{equal} .

Now we describe the delay bound when *different* rates are allowed at the nodes. Let the rate allocated at node i be R_i . The arrival process envelope is $A^{\text{in}}(\tau)$ as before. Let $R_{\min} = \min_i R_i$ and $R_{\min} \geq \rho$. It was shown in [12, 6] that the service curve for a tandem of nodes $1, 2, \dots, N$ is given by (example 3.5 of [12])

$$S_N(\tau) = \left[\left(\tau - \sum_{i=1}^N D_i - \sum_{i=1}^N \frac{C_i}{R_i} \right) R_{\min} \right]^+ \quad (4)$$

As before, the end-to-end delay bound is the maximum horizontal distance between the arrival process envelope $A^{\text{in}}(\tau)$ and the service curve $S_N(\tau)$. Then the following delay bound can be easily obtained.

$$D^{\text{bound}} = \begin{cases} \frac{(b - L^{\max})(\hat{p} - R_{\min})}{R_{\min}(\hat{p} - \rho)} + \frac{L^{\max}}{R_{\min}} + \sum_{i=1}^N \left[\frac{C_i}{R_i} + D_i \right], & \hat{p} > R_{\min}, \\ \frac{L^{\max}}{R_{\min}} + \sum_{i=1}^N \left[\frac{C_i}{R_i} + D_i \right], & \hat{p} \leq R_{\min}. \end{cases} \quad (5)$$

It can be seen that when the rate allocated at all the nodes is identical and R , we recover the delay bound of Eqn. (3).

4 Minimum Cost Rate Allocation

Suppose that a rate R_i is allocated to the flow on link i , $1 \leq i \leq N$, in such a way that the QoS requirements of the flow are met. We call $\bar{R} = \{R_1, \dots, R_N\}$ the “rate vector” assigned to the flow. Then the flow can be admitted if there exists a rate vector assignment $\bar{R}$ that satisfies the following constraints:

1. $R_{\min} = \min_i R_i \geq \rho$, i.e., the minimum allocated rate along the route is greater than the average arrival rate of the flow.
2. $D^{\text{bound}} \leq D^{\text{reqd}}$, i.e., the end-to-end delay seen by packets of the flow is less than the end-to-end delay requirement specified by the flow.

3. $R_i \leq \gamma_i$, $1 \leq i \leq N$; γ_i is the available link bandwidth on link i (which may be less than the capacity of the link). This constraint simply says that the allocated rate should not exceed the available link bandwidth.

There may be many rate vector assignments that satisfy the above constraints. We associate a cost with allocating a rate on a link. The cost function $f_i(R_i)$ for link i is assumed to be strictly convex and increasing in the allocated rate R_i . This implies that low-cost rate allocations will tend to avoid assigning large rates. It is further assumed that the cost function is the same for each link and is denoted by $f(\cdot)$. For a rate vector $\bar{R} = \{R_1, \dots, R_N\}$, the total cost is defined as $t(\bar{R}) = \sum_{i=1}^N f(R_i)$. We would like to allocate *that* rate vector which minimizes the total cost and hence our optimization criterion is minimization of the total cost for the flow.

Without loss of generality, we order the N links in the tandem such that the link with the least available capacity is numbered 1 and so on, i.e., $\gamma_1 \leq \gamma_2 \leq \dots \leq \gamma_N$. The total cost minimization problem, denoted by **MinimumCost**, is as follows:

$$(\text{MinimumCost}) \quad \min \sum_{i=1}^N f(R_i)$$

subject to:

$$\frac{(b - L^{\max})(\hat{p} - R_{\min})}{R_{\min}(\hat{p} - \rho)} + \frac{L^{\max}}{R_{\min}} + \sum_{i=1}^N \left[\frac{C_i}{R_i} + D_i \right] \leq D^{\text{reqd}} \quad (6)$$

$$R_i \leq \gamma_i, 1 \leq i \leq N \quad (7)$$

$$R_{\min} \geq \rho \quad (8)$$

We note that if $R_1 = R_2 = \dots = R_N = \rho$ is a feasible solution to the **MinimumCost** problem, it is also the optimal solution. That is, the end-to-end delay requirement is met by simply allocating the average rate of arrivals at each link on the path. To avoid this trivial case, we consider only those connections for which allocating the average rate at each link violates the end-to-end delay requirement. Such connections are called “delay-constrained.” In other words, we assume that $(\rho, \rho, \dots, \rho)$ is *not* a feasible solution to the **MinimumCost** problem.

Substituting $\sigma = (b\hat{p} - \rho L^{\max})/(\hat{p} - \rho)$, we can rewrite the constraint in Eqn. (6) as follows.

$$\frac{\sigma}{R_{\min}} + \sum_{i=1}^N \frac{C_i}{R_i} \leq D^{\text{reqd}} - \sum_{i=1}^N D_i + \frac{(b - L^{\max})}{(\hat{p} - \rho)} \quad (9)$$

We denote the R.H.S. of Eqn. (9) by D and note that D is a positive quantity. Let $x_i = 1/R_i$, $1 \leq i \leq N$, $\theta_i = 1/\gamma_i$, $1 \leq i \leq N$ and $\delta = 1/\rho$. Let $x_{\max} = \max_i x_i$. Let $\bar{x} = (x_1, \dots, x_N)$. Let $h(x_i) = f(1/x_i)$. The two results that follow are stated without proof.

Lemma 1. *If $f(x)$ is a convex non-decreasing (concave non-increasing) function over $x \geq 0$, then $h(x) = f(1/x)$ is a convex non-increasing (concave non-decreasing) function over $x \geq 0$.*

¹We assume $R \geq \rho$ because when $R < \rho$, the delay is unbounded.

We now recast the **MinimumCost** problem in terms of $h(\cdot)$, x_i , θ_i and δ .

$$(\text{MinimumCost}) \quad \min \sum_{i=1}^N h(x_i)$$

subject to:

$$\sigma x_{\max} + \sum_{j=1}^N C_j x_j \leq D \quad (10)$$

$$x_i \geq \theta_i, \quad 1 \leq i \leq N \quad (11)$$

$$x_i \leq \delta, \quad 1 \leq i \leq N \quad (12)$$

Let $H(\bar{x}) = \sum_{i=1}^N h(x_i)$ where $\bar{x} = (x_1, \dots, x_N)$. The following can be easily proved.

Lemma 2. $H(\bar{x})$ is a convex function in $\bar{x}$.

It can be seen that the **MinimumCost** problem has a convex cost function with linear constraints. For such problems, there exist simple necessary and sufficient conditions (in terms of Lagrange multipliers) for a feasible solution to be optimal [17].

5 The Unbounded Link Capacities Case

First, we investigate the special case of the problem where we can allocate any amount of bandwidth on the links, i.e., $\gamma_i = \infty$ and $\theta_i = 0$, $1 \leq i \leq N$. We note that the constraint Eqn. (10) is actually equivalent to N constraints:

$$\sigma x_i + \sum_{j=1}^N C_j x_j \leq D, \quad 1 \leq i \leq N \quad (13)$$

We denote by **UnbddRates** this special case, i.e., the **MinimumCost** problem without the available link bandwidth constraints, viz.,

$$(\text{UnbddRates}) \quad \min \sum_{i=1}^N h(x_i)$$

subject to:

$$\sigma x_i + \sum_{j=1}^N C_j x_j \leq D, \quad 1 \leq i \leq N \quad (14)$$

$$x_i > 0, \quad 1 \leq i \leq N \quad (15)$$

$$x_i \leq \delta, \quad 1 \leq i \leq N \quad (16)$$

Now suppose that we decide to allocate an identical rate to the flow on each link along its route. It is easy to compute the minimum identical rate R_{equal} from the delay bound constraint of Eqn. (14); we have $x_{\text{equal}} = \frac{D}{\sigma + \sum_{i=1}^N C_i}$ and $R_{\text{equal}} = \frac{\sigma + \sum_{i=1}^N C_i}{D}$. Thus $\bar{x}_{\text{equal}} = (x_{\text{equal}}, \dots, x_{\text{equal}})$ is a feasible solution to the **UnbddRates** problem and among all identical rate vectors, $\bar{R}_{\text{equal}} = (R_{\text{equal}}, \dots, R_{\text{equal}})$ has the least total cost. This follows because reducing x_{equal} further will only push up the cost $\sum_{i=1}^N h(x_i)$, as $h(x_i)$ is a non-increasing function.

The approach of allocating of $\bar{R}_{\text{equal}}$ is widely used to provide end-to-end delay guarantees under the Guaranteed Services framework. We refer to this policy as

the **Identical Rates** policy. Our next theorem states that $\bar{x}_{\text{equal}}$ need not always be the optimal solution to the **UnbddRates** problem and gives an explicit condition for $\bar{x}_{\text{equal}}$ to be the optimal solution.

Theorem 1. $\bar{x}_{\text{equal}}$ is the optimal solution to the **UnbddRates** problem iff $\sigma + \sum_{j=1}^N C_j > NC_i$, $1 \leq i \leq N$.

Proof: We apply Proposition 3.4.2 of [17]. The constraints of the **UnbddRates** problem are ordered as follows.

- $\sigma x_i + \sum_{j=1}^N C_j x_j \leq D$ is the i th constraint, ($1 \leq i \leq N$).
- $x_i > 0$ is the $(N + i)$ th constraint, ($1 \leq i \leq N$).
- $x_i \leq \delta$ is the $(2N + i)$ th constraint, ($1 \leq i \leq N$).

As seen before, $\bar{x}_{\text{equal}}$ is a feasible solution to the **UnbddRates** problem. Let μ_i^* denote the multiplier for the i th constraint.

- If $\bar{x}_{\text{equal}}$ is also the optimal solution, we need to show the existence of appropriate scalars μ_i^* , $1 \leq i \leq 3N$ that satisfy the conditions of Proposition 3.4.2 of [17].
- If $\bar{x}_{\text{equal}}$ is *not* the optimal solution, we need to show that there exist no scalars μ_i^* , $1 \leq i \leq 3N$ that satisfy the conditions of Proposition 3.4.2 of [17].

Let J and $A(\cdot)$ be as defined in Proposition 3.4.2 of [17]. Therefore, let $J = \{1, \dots, 3N\}$ and at $\bar{x}_{\text{equal}}$, $A(\bar{x}_{\text{equal}}) = \{1, \dots, N\}$ as these are the only active constraints at $\bar{x}_{\text{equal}}$. If $\bar{x}_{\text{equal}}$ is the optimal solution to the **UnbddRates** problem, then $\mu_i^* = 0$, $(N + 1) \leq i \leq 3N$ and we need to show that there exist unique nonnegative $\mu_1^*, \dots, \mu_N^*$, not all zero, such that (using the notation of Proposition 3.4.2 of [17]), $g(x) = f(x) + \sum_{j \in J} \mu_j^* (a_j' x - b_j)$ is minimized at $\bar{x}_{\text{equal}}$. In our case,

$$g(\bar{x}) = h(x_1) + \dots + h(x_N) + \sum_{j=1}^N \mu_j^* \left[\sigma x_j + \sum_{i=1}^N C_i x_i - D \right] \quad (17)$$

Note that $g(\bar{x})$ is a convex function in $\bar{x}$. If $g(\bar{x})$ is minimized at $\bar{x}_{\text{equal}}$, the partial derivatives of $g(\bar{x})$ should vanish at $\bar{x}_{\text{equal}}$. If we take partial derivatives w.r.t. x_i , $1 \leq i \leq N$ and set each partial derivative equal to 0 at $\bar{x}_{\text{equal}}$, we get N linear equations for μ_j^* , $1 \leq j \leq N$ as follows,

$$\sigma \mu_i^* + C_i \sum_{j=1}^N \mu_j^* = -h'(x_{\text{equal}}), \quad 1 \leq i \leq N \quad (18)$$

where $h'(x_{\text{equal}})$ is the derivative of $h(x)$ at x_{equal} . Solving these N equations for μ_i^* gives

$$\mu_i^* = -\frac{h'(x_{\text{equal}}) \sigma + \sum_{j=1}^N C_j - NC_i}{\sigma + \sum_{j=1}^N C_j}, \quad i = 1, \dots, N \quad (19)$$

Note that $h'(x_{\text{equal}})$ is a negative quantity. This implies that if $\sigma + \sum_{j=1}^N C_j > NC_i \forall i$, then $\bar{x}_{\text{equal}}$ is the optimal solution. Otherwise, there do not exist non-negative $\mu_i \forall i \in \{1, \dots, N\}$ and hence $\bar{x}_{\text{equal}}$ is not the optimal solution. ■

There are two points to note here. Firstly, Theorem 1 is true for *any* convex non-decreasing cost function $f()$. Secondly, it can be seen from Theorem 1 that the identical rate vector $\bar{x}_{\text{equal}}$ need not always be optimal. We give a numerical example to show that a better rate vector indeed exists if $\sigma + \sum_{j=1}^N C_j > NC_i$ is not satisfied for every i .

Let us consider the simple cost function $f(R_i) = R_i$ which means $h(x_i) = 1/x_i$. Thus our optimization criterion is the minimization of the total allocated bandwidth. Consider the following choice of parameters.

- $N = 3$ and $\sigma = 30, C_1 = 25, C_2 = 5, C_3 = 5, \rho = 1, D = 6.5$.

Then it can be seen that $\sigma + \sum_{i=1}^3 C_i > 3C_2$ and $\sigma + \sum_{i=1}^3 C_i > 3C_3$, but $\sigma + \sum_{i=1}^3 C_i < 3C_1$.

- The identical rate vector is $\bar{x}_{\text{equal}} = (0.1, 0.1, 0.1)$ for which the total cost is $H(\bar{x}_{\text{equal}}) = 30$.
- Consider $\bar{x}' = (x'_1, x'_2, x'_3)$ where $x'_1 = 0.095, x'_2 = 0.103 + \frac{0.005}{35}, x'_3 = 0.103$. It can be easily verified that $\bar{x}'$ is a feasible solution and $H(\bar{x}') = 29.93 < H(\bar{x}_{\text{equal}})$.

This shows that $\bar{x}_{\text{equal}}$ is not the optimal solution as $\bar{x}'$ is a better rate vector.

Corollary 1. For the case $C_1 = C_2 = \dots = C_N$ and unbounded rates, $\bar{R}_{\text{equal}}$ is the optimal rate allocation vector.

C_i corresponds to the MTU at the interface i . When the same link technology is employed in a network backbone (a situation that may occur quite often in practice), we have the special case of equal MTU: $C_1 = C_2 = \dots = C_N (= C_{\text{equal}})$. In the rest of the paper we work with this assumption.

6 The Bounded Link Capacities Case

Now suppose that the optimal rate R_{equal} is not available on some links of the route. When this happens, the Connection Admission Control (CAC) module is usually programmed to block the incoming flow. But, in fact, sufficient bandwidth might not be available only on a few links. There might exist another rate vector which satisfies the delay constraint and the available link bandwidth constraints. We note that if such a rate vector exists, the total cost of this rate vector is greater than $H(\bar{x}_{\text{equal}})$. Thus, this is the price to be paid for admitting the flow: the cost is higher.

When $\bar{x}_{\text{equal}}$ is infeasible, there is at least one link which the available bandwidth is less than R_{equal} . By our numbering convention, link 1 has insufficient bandwidth; i.e.,

$\gamma_1 < R_{\text{equal}}$. Let $\bar{x}^* = (x_1^*, x_2^*, \dots, x_N^*)$ be the optimal solution when $\bar{x}_{\text{equal}}$ is infeasible.

Lemma 3. When $\bar{x}_{\text{equal}}$ is infeasible, the optimal vector $\bar{x}^*$ is characterized by $x_1^* = \theta_1$ and $x_i^* \leq \theta_1, 2 \leq i \leq N$.

The result above says that when $\bar{x}_{\text{equal}}$ is infeasible, it is optimal to allocate the *entire* available bandwidth on link 1 (reciprocal of θ_1) to the incoming connection. *Proof:* The proof relies on establishing the three claims that follow.

Claim 1: There exists $i \in \{2, \dots, N\}$ such that $x_i^* < \theta_1$. *Proof:* We have

$$\sigma x_{\text{equal}} + \sum_{i=1}^N C_{\text{equal}} x_{\text{equal}} = D$$

Let $x_{\text{max}}^* = \max_i x_i^*$. As $\bar{x}^*$ is the optimal solution,

$$\sigma x_{\text{max}}^* + C_{\text{equal}} \sum_{i=1}^N x_i^* \leq D$$

But $x_1^* \geq \theta_1 > x_{\text{equal}}$. This implies

$$\{\sigma + (N-1)C_{\text{equal}}\}x_{\text{equal}} \geq \sigma x_{\text{max}}^* + C_{\text{equal}} \sum_{i=2}^N x_i^*$$

But $x_{\text{max}}^* \geq x_1^* \Rightarrow x_{\text{max}}^* > x_{\text{equal}}$. This implies

$$\{\sigma + (N-1)C_{\text{equal}}\}x_{\text{equal}} > \sigma x_{\text{equal}} + C_{\text{equal}} \sum_{i=2}^N x_i^*$$

which gives

$$\frac{1}{(N-1)} \sum_{i=2}^N x_i^* < x_{\text{equal}}$$

Hence there exists an $i \in \{2, \dots, N\}$ such that $x_i^* < x_{\text{equal}} < \theta_1$. Denote it by x_k^* .

Claim 2: $x_1^* = \theta_1$. *Proof:* We prove this by contradiction. Assume that $x_1^* > \theta_1$. Recalling that $x_k^* < \theta_1$, consider a new rate vector $\bar{x}'$ such that, $x'_1 = x_1^* - \epsilon > \theta_1$, $x'_k = x_k^* + \epsilon < \theta_1$ and $x'_i = x_i^*, i \neq \{1, k\}$. Let $x'_{\text{max}} = \max_i x'_i$. We choose an ϵ that preserves the structure of $\bar{x}^*$, i.e., if $x_{\text{max}}^* = x_j^*$, then $x'_{\text{max}} = x'_j$. It is easy to see that such an ϵ exists. Note that $\bar{x}'$ satisfies the delay constraint. Also note that

$$h(x_1^*) + h(x_k^*) \geq h(x'_1) + h(x'_k)$$

as $h()$ is a convex function. This implies $\bar{x}'$ is the optimal rate vector and not $\bar{x}^*$. This is a contradiction. Thus $x_1^* = \theta_1$.

Claim 3: $x_i^* \leq \theta_1, 1 \leq i \leq N$. *Proof:* We know that $x_1^* = \theta_1$ and $x_k^* < \theta_1$. We need to prove the claim for the rest. We do this by contradiction. Let there exist a j (among the rest) such that $x_j^* > \theta_1$. Consider $x'_1 = x_1^* + \epsilon$ and $x'_j = x_j^* - \epsilon$. With arguments similar to those of claim 2, we get a better cost at $\bar{x}'$ which contradicts the assumption that $\bar{x}^*$ was the optimum. Thus $x_i^* \leq \theta_1$, as required. ■

6.1 The OneLink Problem

Now we consider the case when $\gamma_1 < R_{\text{equal}}$, but there are no restrictions on the available link capacities on other links. We refer to this as the **OneLink** problem and want to obtain the corresponding optimal rate vector. Let $\bar{x}^* = \{x_1^*, \dots, x_N^*\}$ denote the optimal rate vector. From Lemma 3, we know that $x_i^* \leq \theta_1, i = \{2, \dots, N\}$ and $x_1^* = \theta_1$. Then the **OneLink** problem is as follows:

$$(\text{OneLink}) \quad \min \sum_{i=2}^N h(x_i)$$

subject to:

$$C_{\text{equal}} \sum_{j=2}^N x_j \leq D - (\sigma + C_{\text{equal}})\theta_1 \quad (20)$$

$$\sigma x_i + C_{\text{equal}} \sum_{j=2}^N x_j \leq D - C_{\text{equal}}\theta_1, 2 \leq i \leq N \quad (21)$$

$$x_i > 0, \quad 2 \leq i \leq N \quad (22)$$

$$x_i \leq \delta, \quad 2 \leq i \leq N \quad (23)$$

Lemma 4. For the **OneLink** problem, the optimal rate vector is such that $x_2^* = \dots = x_N^*$.

Proof: By contradiction. Assume there exist $x_j^* \neq x_k^*$. Without loss of generality $x_j^* < x_k^*$. Let $x' = \frac{x_j^* + x_k^*}{2}$. Then by convexity of $h()$,

$$h\left(\frac{x_j^* + x_k^*}{2}\right) \leq \frac{1}{2} \{h(x_j^*) + h(x_k^*)\}$$

which gives

$$h(x') + h(x') \leq h(x_j^*) + h(x_k^*)$$

contradicting the assumption that $\bar{x}^*$ is the optimal rate vector. ■

It is easy to see that the delay constraint inequality is tight at the optimal rate vector. This gives $x_2^* = \dots = x_N^* = \frac{D - (\sigma + C_{\text{equal}})\theta_1}{(N-1)C_{\text{equal}}}$. Let $x_{\text{equal}}^1 = \frac{D - (\sigma + C_{\text{equal}})\theta_1}{(N-1)C_{\text{equal}}}$ and $R_{\text{equal}}^1 = \frac{1}{x_{\text{equal}}^1}$. Let $\bar{x}_{\text{OneLink}} = (\theta_1, x_{\text{equal}}^1, \dots, x_{\text{equal}}^1)$ and $\bar{R}_{\text{OneLink}} = (\gamma_1, R_{\text{equal}}^1, \dots, R_{\text{equal}}^1)$. Summarizing, we have the following:

Theorem 2. $\bar{x}_{\text{OneLink}}$ is the optimal solution to the **OneLink** problem and therefore $\bar{R}_{\text{OneLink}}$ the optimal rate allocation vector for the **OneLink** problem.

6.2 The TwoLinks Problem

It is possible that $\bar{x}_{\text{equal}}$ and $\bar{x}_{\text{OneLink}}$ are both not feasible. Then we know that $\theta_1 > x_{\text{equal}}$ and $\theta_2 > x_{\text{equal}}^1$. Let $\bar{x}^* = \{x_1^*, \dots, x_N^*\}$ be the optimal solution for this case.

Lemma 5. When $\bar{x}_{\text{equal}}$ and $\bar{x}_{\text{OneLink}}$ are both infeasible, the optimal rate vector $\bar{x}^*$ is characterized by $x_1^* = \theta_1, x_2^* = \theta_2$ and $x_i^* \leq \theta_2, 3 \leq i \leq N$.

Proof: From Lemma 3, $x_1^* = \theta_1$ and $x_i^* \leq \theta_1$.

Claim 1: There exists $i \in \{3, \dots, N\}$ such that $x_i^* < \theta_2$. *Proof:* We have

$$\sigma\theta_1 + C_{\text{equal}}\theta_1 + \sum_{i=2}^N C_{\text{equal}}x_{\text{equal}}^1 = D$$

Let $x_{\text{max}}^* = \max_i x_i^*$, then $x_{\text{max}}^* = \theta_1$. As $\bar{x}^*$ is the optimal solution,

$$\sigma x_{\text{max}}^* + C_{\text{equal}} \sum_{i=1}^N x_i^* \leq D$$

Therefore,

$$\sigma\theta_1 + C_{\text{equal}}\theta_1 + C_{\text{equal}} \sum_{i=2}^N x_i^* \leq D$$

But $x_2^* \geq \theta_2 > x_{\text{equal}}^1$. This implies that

$$(N-2)C_{\text{equal}}x_{\text{equal}}^1 \geq C_{\text{equal}} \sum_{i=3}^N x_i^*$$

which gives

$$\frac{1}{(N-2)} \sum_{i=3}^N x_i^* < x_{\text{equal}}^1$$

Hence there exists a $i \in \{3, \dots, N\}$ such that $x_i^* < x_{\text{equal}} < \theta_2$. Denote it by x_k^* .

Claim 2: $x_2^* = \theta_2$. *Proof:* Similar to that of Claim 2 of Lemma 3.

Claim 3: $x_i^* \leq \theta_2, 3 \leq i \leq N$. *Proof:* Similar to that of Claim 3 of Lemma 3. ■

Next, we consider the problem where $\gamma_1 < R_{\text{equal}}$ and $\gamma_2 < R_{\text{equal}}^1$, but there are no restrictions on the available link capacities on other links. This is the **TwoLinks** problem. Let $\bar{x}^* = \{x_1^*, \dots, x_N^*\}$ denote the optimal rate vector for this problem. From Lemma 5, we know that $x_1^* = \theta_1$ and $x_2^* = \theta_2$. Then the problem is:

$$(\text{TwoLinks}) \quad \min \sum_{i=3}^N h(x_i)$$

subject to:

$$C_{\text{equal}} \sum_{j=3}^N x_j \leq D - (\sigma + C_{\text{equal}})\theta_1 - C_{\text{equal}}\theta_2,$$

$$C_{\text{equal}} \sum_{j=3}^N x_j \leq D - C_{\text{equal}}\theta_1 - (\sigma + C_{\text{equal}})\theta_2,$$

$$\sigma x_i + C_{\text{equal}} \sum_{j=3}^N x_j \leq D - C_{\text{equal}}(\theta_1 + \theta_2) \quad 3 \leq i \leq N$$

$$x_i > 0, \quad 3 \leq i \leq N$$

$$x_i \leq \delta, \quad 3 \leq i \leq N$$

Lemma 6. $x_3^* = \dots = x_N^*$.

Proof: Same as that of Lemma 4. ■

It is easy to see that the delay constraint inequality is tight at the optimal rate vector. This gives $x_3^* = \dots = x_N^* = \frac{D - (\sigma + C_{\text{equal}})\theta_1 - C_{\text{equal}}\theta_2}{(N-2)C_{\text{equal}}}$.

$$\text{Let } x_{\text{equal}}^2 = \frac{D - (\sigma + C_{\text{equal}})\theta_1 - C_{\text{equal}}\theta_2}{(N-2)C_{\text{equal}}} \text{ and } R_{\text{equal}}^2 = \frac{1}{x_{\text{equal}}^2}.$$

Let $\bar{x}_{\text{TwoLinks}} = (\theta_1, \theta_2, x_{\text{equal}}^2, \dots, x_{\text{equal}}^2)$ and $\bar{R}_{\text{TwoLinks}} = (\gamma_1, \gamma_2, R_{\text{equal}}^2, \dots, R_{\text{equal}}^2)$. Then we have:

Theorem 3. $\bar{x}_{\text{TwoLinks}}$ is the optimal solution to the **TwoLinks** problem and therefore $\bar{R}_{\text{TwoLinks}}$ is the optimal rate allocation vector for the **TwoLinks** problem.

6.3 The K -Links Problem

It is clear that in a similar way, one can have problems where 3, 4, 5, $\dots$ links have limited capacities. The pattern of the optimal solution for the **K-Links** problem, $K \geq 3$, is already clear from the optimal solutions for the **OneLink** and the **TwoLinks** problems. Let

$$R_{\text{equal}}^I = \frac{(N-I)C_{\text{equal}}}{D - \frac{\sigma + C_{\text{equal}}}{\gamma_1} - \frac{C_{\text{equal}}}{\gamma_2} - \dots - \frac{C_{\text{equal}}}{\gamma_I}}, I \geq 1$$

and $x_{\text{equal}}^I = 1/R_{\text{equal}}^I$.

Let $\bar{R}_{I\text{Links}} = (\gamma_1, \dots, \gamma_I, R_{\text{equal}}^I, \dots, R_{\text{equal}}^I)$ and

$\bar{x}_{I\text{Links}} = (\theta_1, \dots, \theta_I, x_{\text{equal}}^I, \dots, x_{\text{equal}}^I), I \geq 1$.

The **K-Links** problem is as follows.

- The available link capacity on link $I, 1 \leq I \leq K$, is γ_I , where $\gamma_I < R_{\text{equal}}^{(I-1)}$ and unlimited capacity is available on links $(K+1), \dots, N$. What is the optimal rate allocation vector that minimizes the total cost and obeys the delay constraint and the available link capacity constraints?

As the technique of obtaining the optimal solution is essentially the same as that in the **TwoLinks** case, we state the results directly, without proof. Let $\bar{x}^* = \{x_1^*, \dots, x_N^*\}$ be the optimal solution to the **K-Links** problem.

Lemma 7. $x_i^* = \theta_i, 1 \leq i \leq K$ and $x_i^* \leq \theta_K, (K+1) \leq i \leq N$.

Lemma 8. $x_{(K+1)}^* = \dots = x_N^*$.

Theorem 4. $\bar{x}_{K\text{Links}}$ is the optimal solution to the **K-Links** problem and therefore $\bar{R}_{K\text{Links}}$ is the optimal rate allocation vector for the **K-Links** problem.

In practice, the available capacities of the links are always bounded. The previous results lead directly to a **General Rates** allocation algorithm that can be used to decide whether a flow can be admitted and, if so, the optimal rate allocation vector.

7 Simulation Results

General Rates allows an incoming connection to be possibly accepted even when **Identical Rates** blocks

it. But, as we have remarked earlier, the total bandwidth required to be allocated is more for **General Rates**. This means that less bandwidth may be available for calls in the future. So even though **General Rates** looks appealing in the short term, the long-term consequences of following it need to be examined. In this section, we simulate the **General Rates** algorithm to study this aspect.

We provide some details of our simulation scenario. At each link, the packets are scheduled according to the **Weighted Fair Queueing (WFQ)** policy [18, 9, 10]. WFQ falls under the **Guaranteed Services** framework with the following parameters: For a flow j at link i , $C_i = L_j^{\max}$ and $D_i = \frac{L_j^{\max}}{\gamma_i} + d_i^{\text{prop}}$, where $L_j^{\max}$ is the maximum packet size of flow j , $L_j^{\max}$ is the maximum size of the packets at link i , γ_i is the capacity of link i and d_i^{prop} is the propagation delay of link i .

We assume that all the connections to be routed are full-duplex, that all links are bidirectional and the two halves of a full-duplex connection are to be routed on the same path. Connection requests are assumed to arrive according to a Poisson process and last for a duration that is exponentially distributed. We further assume that the traffic specifications and delay requirements of the connection requests are as specified in Table 1. Our performance criterion is the blocking probability of connections.

Class	b (kB)	ρ (Mbps)	$\hat{\rho}$ (Mbps)	D^{reqd} (ms)	$L^{\max}$ (kB)
Voice	0.1	0.064	0.064	50	0.1
Vid conf	10	0.5	10	75	1.5
Stored vid	100	3	10	100	1.5

Table 1: Traffic Parameters

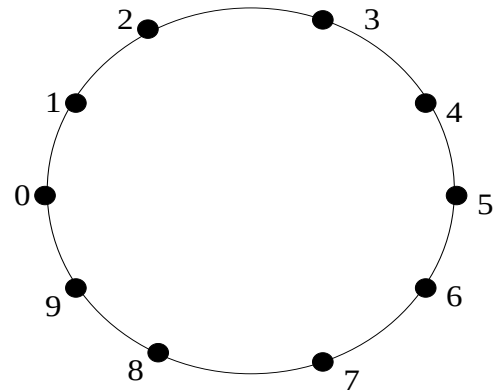


Figure 1: The homogeneous ring topology, with 155 Mbps links and propagation delay of 4 ms on each link.

To study the impact of the **General Rates** policy, we consider source-destination (SD) pairs with at least 2 hops on the path. We consider the homogeneous ring topology of Fig 1 and the NSFNET topology of Fig. 2. Both 2 hops and 3 hops away SD pairs are considered.

To study the effect of non-identical link characteristics, we also consider the ring and NSFNET topologies with link

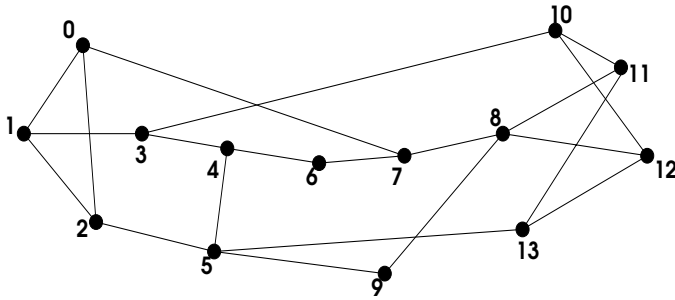


Figure 2: The NSFNET topology, with 155 Mbps links and propagation delay of 4 ms on each link.

link	rate (Mbps)	delay (ms)
0 1	155	4
0 9	310	4
1 2	155	4
2 3	310	4
3 4	155	4
4 5	310	4
5 6	155	4
6 7	310	4
7 8	155	4
8 9	310	4

Table 2: Link parameters for the nonhomogeneous ring topology.

link	rate (Mbps)	delay (ms)
0 1	155	4
0 2	310	4
0 7	620	4
1 2	155	4
1 3	310	4
2 5	620	4
3 4	155	4
3 10	310	4
4 5	620	4
4 6	310	4
5 9	620	4
5 13	155	4
6 7	310	4
7 8	620	4
8 9	155	4
8 11	310	4
8 12	620	4
10 11	155	4
10 12	310	4
11 13	620	4
12 13	155	4

Table 3: Link parameters for the nonhomogeneous NSFNET topology.

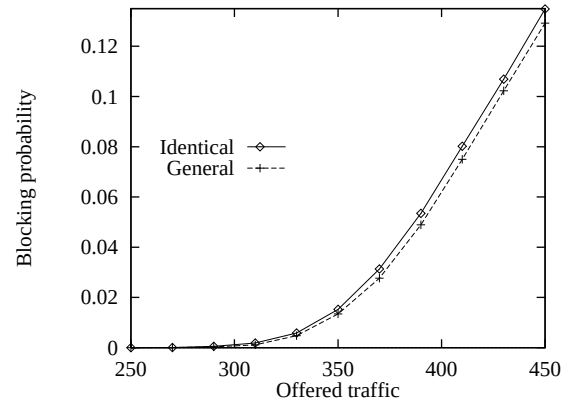


Figure 3: Blocking probabilities for the ring topology of Fig. 1; uniform load, 2 hops away SD pairs.

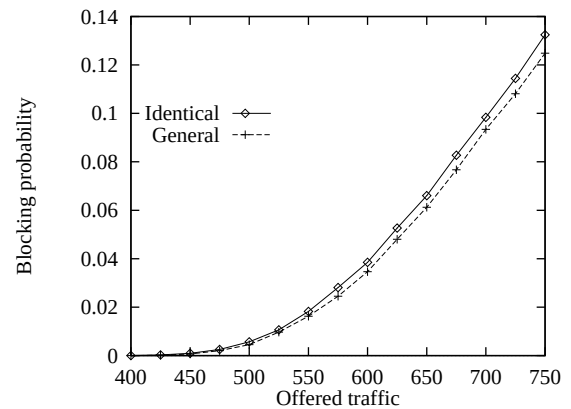


Figure 4: Blocking probabilities for the NSFNET topology of Fig. 2; uniform traffic distribution, 2 hops away SD pairs.

bandwidths taking values 155 and 310 Mbps, and 155, 310 and 620 Mbps, respectively, as shown in Tables 2 and 3.

7.1 Uniform Load

First, we consider the situation when traffic load is *uniform* across all SD pairs. In the next subsection, we present results when traffic load is distributed among SD pairs in an uneven manner.

In Figs 3 to 6, we consider topologies that are “homogeneous” in the sense that link capacities are equal. In Figs 7 to 9, link capacities are unequal. Each of the figures shows a plot of connection blocking probability as traffic load increases.

Fig 3 shows that for the homogeneous ring topology and 2-hop away SD pairs, the General Rates policy achieves slightly better (lower) connection blocking probability than the Identical Rates policy. The same is true for the homogeneous NSFNET topology (Fig 4). However, when we consider 3-hop away SD pairs, the General Rate policy is unable to show any appreciable improvement (Figs 5 and 6).

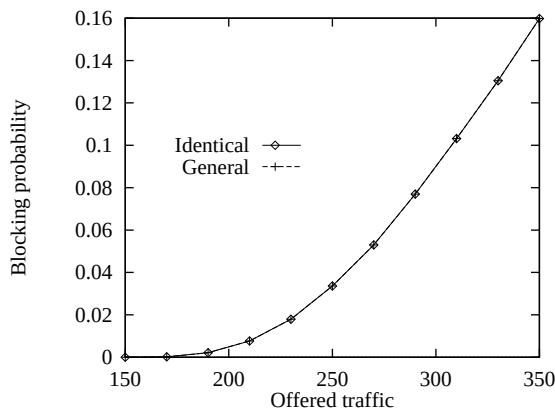


Figure 5: Blocking probabilities for the NSFNET topology of Fig. 2; uniform traffic distribution, 3 hops away SD pairs.

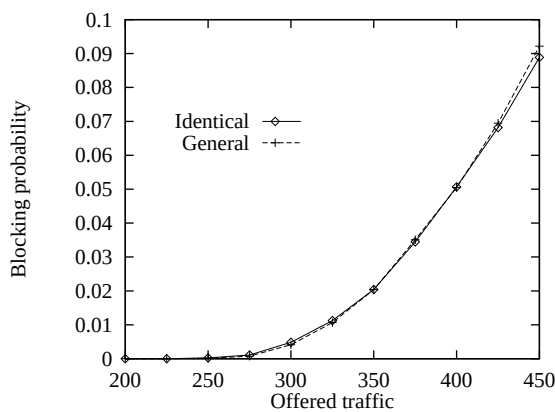


Figure 6: Blocking probabilities for the NSFNET topology of Fig. 2; uniform traffic distribution, 2 & 3 hops away SD pairs.

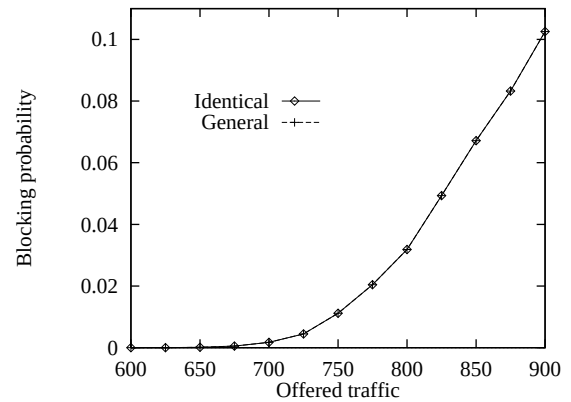


Figure 7: Blocking probabilities for the ring topology of Fig. 1 with link parameters of Table 2; uniform load, 2 hops away SD pairs.

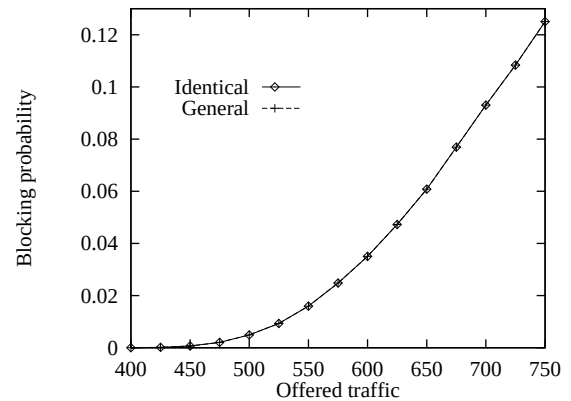


Figure 8: Blocking probabilities for the NSFNET topology of Fig. 2 with non-identical link parameters of Table 3; uniform traffic distribution, 2 hops away SD pairs.

Next, we consider nonhomogeneous topologies. Figs 7, 8 and 9 show that once again, the *General Rates* and *Identical Rates* policies perform similarly as far as connection blocking probability is concerned.

7.2 Nonuniform Load

With nonuniform loads, the general trends are similar. Figs 10 and 11 show that for the homogeneous case (link capacities equal), as long as the SD pairs are 2 hops away, the *General Rates* policy shows a slight improvement in connection blocking. However, as we include SD pairs that are 3 hops away, the advantage disappears (Fig 12 and Fig 13).

Similarly, for nonhomogeneous topologies, there is little to distinguish between the performances of the *General Rates* and the *Identical Rates* policies. We show the 2-hop away SD pair case for the Ring topology in Fig 14, the 2-hop away SD pair case for the NSFNET topology in Fig 15. The case with a mixture of 2-hop away and 3-hop away SD pairs is shown in Fig 16.

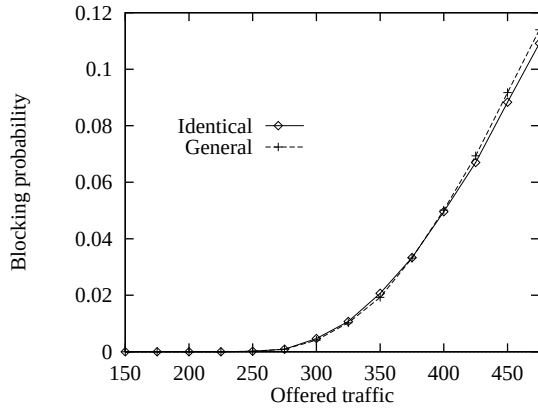


Figure 9: Blocking probabilities for the NSFNET topology of Fig. 2 with non-identical link parameters Table 3; uniform traffic distribution, 2 & 3 hops away SD pairs.

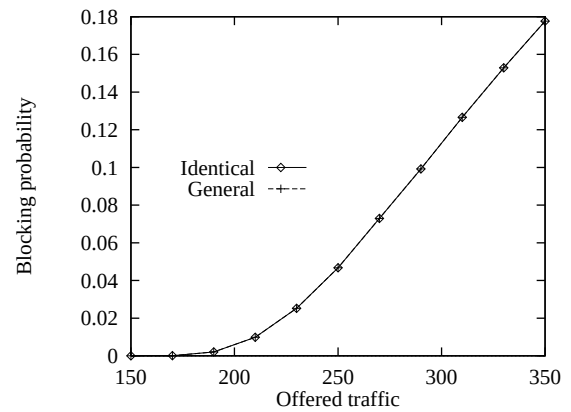


Figure 12: Blocking probabilities for the NSFNET topology of Fig. 2; nonuniform load, 3 hops away SD pairs.

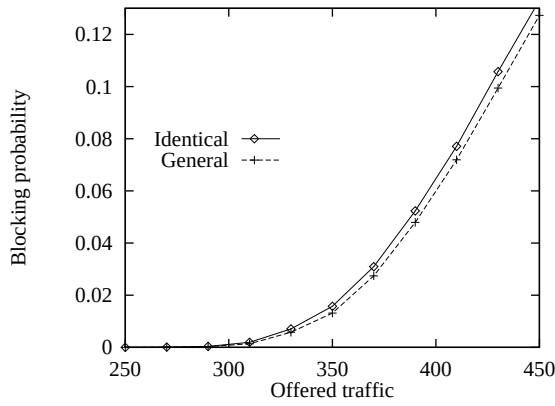


Figure 10: Blocking probabilities for the ring topology of Fig. 1; nonuniform load, 2 hops away SD pairs.

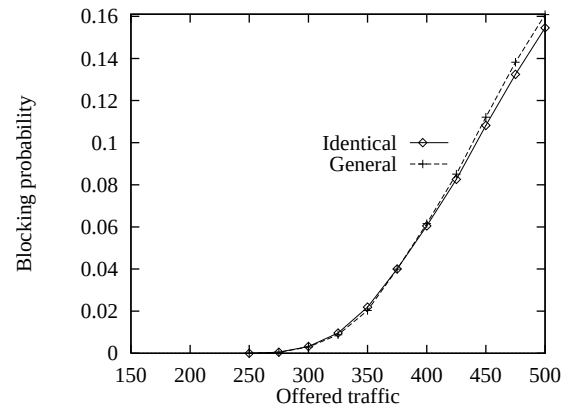


Figure 13: Blocking probabilities for the NSFNET topology of Fig. 2; nonuniform load, 2 & 3 hops away SD pairs.

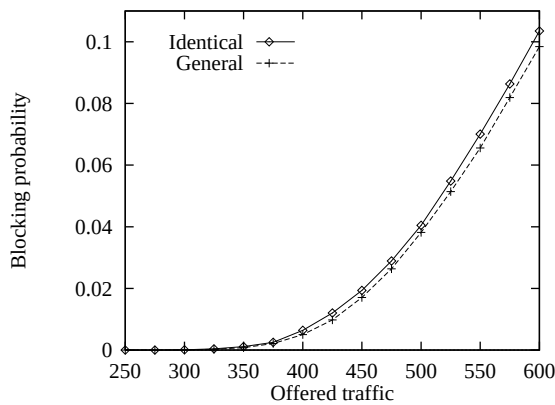


Figure 11: Blocking probabilities for the NSFNET topology of Fig. 2; nonuniform load, 2 hops away SD pairs.

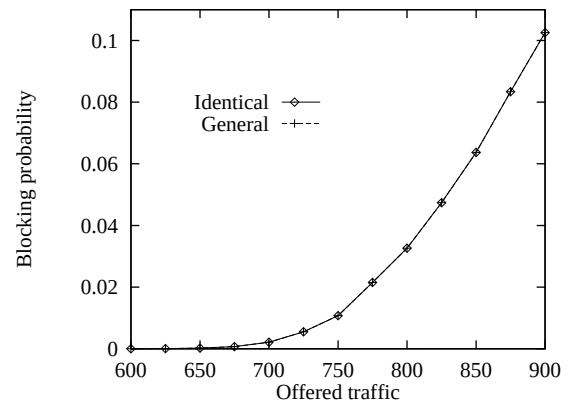


Figure 14: Blocking probabilities for the ring topology of Fig. 1 with link parameters of Table 2; nonuniform load, 2 hops away SD pairs.

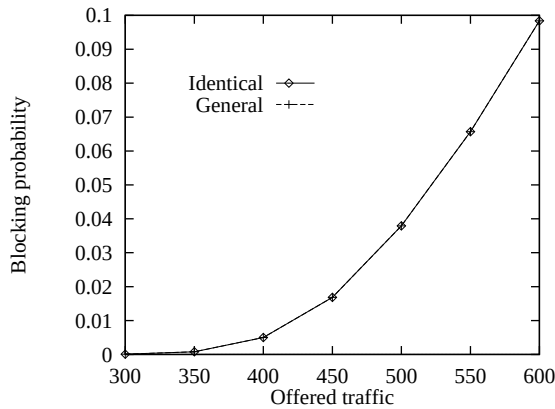


Figure 15: Blocking probabilities for the NSFNET topology of Fig. 2 with link parameters of Table 3; nonuniform load, 2 hops away SD pairs.

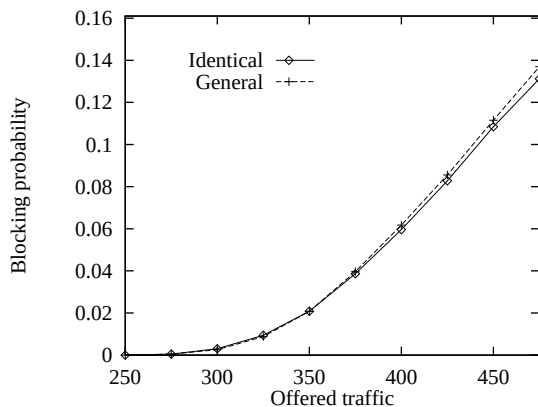


Figure 16: Blocking probabilities for the NSFNET topology of Fig. 2 with link parameters of Table 3; nonuniform load, 2 & 3 hops away SD pairs.

7.3 Summary

The series of plots indicates that the blocking performances achieved by the *Identical* and *General* policies remain close, with slight improvement in one or the other in some cases. The results indicate that there is not much to choose between the two policies as far as blocking performance is concerned.

8 Connection Blocking

We saw in the previous section that accepting a connection whenever possible (what the *General Rates* policy does) may lead to significant blocking of future connections. In this section, we attempt to obtain some insight into this aspect by using very simple arguments.

Suppose that we have a very simple “network” with two links in tandem. Let the capacities of the links be 2.5 and 10 units respectively. Type *A* connections traverse both the links while type *B* connections pass over the second link only. The average rate of traffic from a connection of either type is 0.45. The delay requirement of a type *A* connection can be met by allocating 2 units of bandwidth on the two links. Similarly, the delay requirement of the Type *B* connection can be met by allocating 2 units on the second link.

For the network and traffic types above, the *Identical Rates* policy would admit only one Type *A* connection. Now let us consider the *General Rates* policy. Let us assume that it is possible to accommodate a *second* Type *A* connection by allocating 0.5 and 6 units on the first and second links, respectively. Let us now suppose that the network is empty to begin with, and two Type *A* connections arrive in quick succession. Also, we assume that these connections have infinite lifetimes.

The *Identical Rates* policy will admit the first Type *A* connection and block the second, as well as all future Type *A* connections. The fraction of Type *B* connections that will be blocked can now be obtained using the Erlang-B formula. The bandwidth available on the second link is 8 units and, therefore, 4 Type *B* connections can be accommodated; the $M/M/4$ model applies.

The *General Rates* policy will admit both the Type *A* connections. As a result, the bandwidth available on the second link is only 2 units. For this case, the $M/M/1$ model applies. Clearly, the overall blocking probability in this case will be much higher: the fraction of Type *A* connections blocked is the same as for *Identical Rates*, while the fraction of Type *B* connections blocked is much more.

Admittedly, our example network and traffic scenario are very simple and contrived. Nevertheless, they do highlight the problem that can result when the *General Rates* policy is followed. We are currently working on a complete analysis of this problem to extend the intuitive understanding that can be obtained from this simple model.

9 Conclusions

In this work, we study the problem of rate allocation from the perspective of total cost minimization. It is shown that the *Identical Rates* policy need not always minimize the total cost, and a specific condition is derived which, if satisfied, makes the *Identical Rates* policy optimal. However, if the computed identical rate is not available on every link along the path of a connection, *Identical Rates* blocks a request without searching for other rate allocations. To admit a connection whenever it is possible at all, a *General Rates* algorithm is developed; it computes an optimal rate vector with possibly different rates allocated on different links. However, *General Rates* is forced to allocate more bandwidth on the end-to-end path, and this leaves less resources for future connections. Simulations show that the increased bandwidth consumption of *General Rates* can become significant in the long run, as a result of which the algorithm is unable to show appreciably improved blocking probability. Hence, the simpler *Identical Rates* algorithm is sufficient for all practical purposes.

References

- [1] S. Shenker, C. Partridge, and R. Guérin. Specification of guaranteed quality of service. RFC 2212, September 1997.
- [2] R. Nagarajan, J. Kurose, and D. Towsley. Allocation of local Quality of Service constraints to meet end-to-end requirements. In *IFIP Workshop on the Performance Analysis of ATM Systems*, Martinique, January 1993.
- [3] D. Raz and Y. Shavitt. Optimal partition of QoS requirements with discrete cost functions. In *IEEE INFOCOM*, 2000.
- [4] R. Guérin and V. Peris. Quality-of-Service in packet networks: basic mechanisms and directions. *Computer Networks*, 31(3):169–179, February 1999.
- [5] P. White. RSVP and integrated services in the Internet. *IEEE Communications Magazine*, May 1997.
- [6] D. Stiliadis and A. Varma. Latency-rate servers: a general model for analysis of traffic scheduling algorithms. *IEEE/ACM Transactions on Networking*, 6(5):611–624, October 1998.
- [7] P. Goyal and H. Vin. Generalized guaranteed rate scheduling algorithms: a framework. *IEEE/ACM Transactions on Networking*, 5(4):561–571, August 1997.
- [8] H. Zhang. Service disciplines for guaranteed performance service in packet-switching networks. *Proceedings of the IEEE*, pages 1374–1396, October 1995.
- [9] A. Parekh and R. Gallager. A generalized processor sharing approach to flow control in integrated services networks: the single node case. *IEEE/ACM Transactions on Networking*, 1(3):137–150, 1993.
- [10] A. Parekh and R. Gallager. A generalized processor sharing approach to flow control in integrated services networks: the multiple node case. *IEEE/ACM Transactions on Networking*, 2(2):347–357, 1993.
- [11] V. Peris, L. Georgiadis, R. Guérin, and K. Sivaranjan. Efficient network QoS provisioning based on per node traffic shaping. *IEEE/ACM Transactions on Networking*, 4(4):482–501, August 1996.
- [12] R. Agrawal and R. Rajan. Performance bound for guaranteed and adaptive services. Technical Report RC20649, IBM Research Report, December 1996.
- [13] L. Georgiadis, R. Guérin, and A. Parekh. Optimal multiplexing on a single link: delay and buffer requirements. In *IEEE INFOCOM*, 1994.
- [14] S. Shenker and L. Breslau. Two issues in reservation establishment. In *SIGCOMM Symposium on Communications Architectures and Protocols*, Cambridge, Massachusetts, September 1995.
- [15] H. Ahmadi, J. Chen, and R. Guérin. Dynamic routing and call control in high-speed integrated networks. In *13th International Teletraffic Congress*, pages 19–26, Copenhagen, Denmark, 1991.
- [16] L. Georgiadis, R. Guérin, V. Peris, and R. Rajan. Efficient support of delay and rate guarantees in an Internet. In *ACM SIGCOMM*, pages 106–116, August 1996.
- [17] D. Bertsekas. *Nonlinear Programming*. Athena Scientific, Belmont, Massachusetts, 1995.
- [18] A. Demers, S. Keshav, and S. Shenker. Analysis and simulation of a fair queueing algorithm. In *ACM SIGCOMM*, pages 3–12, 1989.