

SHRiNK: A method for scaleable performance prediction and efficient network simulation

Rong Pan[†], Balaji Prabhakar[†], Konstantinos Psounis[†], Damon Wischik[‡]

[†] Stanford University, [‡] Cambridge University

Email: rong, balaji, kpsounis@stanford.edu, D.J.Wischik@statslab.cam.ac.uk

Abstract—In networks and in web-server farms, it is useful to collect performance measurements, to monitor the state of the system, and to perform simulations. However, the sheer volume of traffic in large high-speed network systems makes it hard to monitor their performance or to simulate them efficiently. And the heterogeneity of the Internet means it is time-consuming and difficult to devise the traffic models and analytic tools which would allow us to work with summary statistics.

We explore a method to side-step these problems by combining sampling, modeling and simulation. Our hypothesis is this: if we take a sample of the input traffic, and feed it into a suitably scaled version of the system, we can extrapolate from the performance of the scaled system to that of the original.

Our main findings are: When we scale an IP network which is shared by TCP-like, UDP and web flows; and which is controlled by a variety of active queue management schemes, then performance measures such as queueing delay and drop probability are left virtually unchanged. We show this in theory and in simulations. This makes it possible to capture the performance of large networks quite faithfully using smaller scale replicas.

I. INTRODUCTION

Measuring the performance of the Internet and predicting its behavior under novel protocols and architectures are important research problems. These problems are made difficult by the sheer size and heterogeneity of the Internet: it is very hard to simulate large networks and to pinpoint aspects of algorithms and protocols relevant to their behavior. This has prompted work on traffic sampling [6], [7]. Sampling certainly reduces the volume of data, although it can be hard to work backwards—to infer the performance of the original system.

A direct way to measure and predict performance is with exhaustive simulation: If we record the primitive inputs to the system, such as session arrival times and flow types, we can in principle compute the full state of the system. Further, through simulation we can test the behavior of the network under new protocols and architectures. But such large-scale simulation requires massive computing power.

Reduced-order models can go some way in reducing the burden of simulation. In some cases [11], [26] one can reduce the dimensionality of the data, for example by working with traffic matrices rather than full traces, while retaining enough information to estimate the state of the network. The trouble is that this requires careful traffic characterization and model-building. The heterogeneity of the Internet makes this time-consuming and difficult, since each scenario might potentially require a different new model.

In this paper we explore a way to reduce the computational requirements of simulations and the cost of experiments, and hence simplify network measurement and performance prediction. We do this by combining simulations with sampling and analysis. Our basic hypothesis, which we call SHRiNK (for Small-scale Hi-fidelity Reproduction of Network Kinetics), is this: if we take a *sample* of the traffic, and feed it into a *suitably scaled* version of the system, we can *extrapolate* from the performance of the scaled system to that of the original.

This has two benefits. First, by relying only on a sample of the traffic, SHRiNK reduces the amount of data we need to work with. Second, by using samples of actual traffic, it short-cuts the traffic characterization and model-building process while ensuring the relevance of the results.

This approach also presents challenges. At first sight, it appears optimistic. Might not the behavior of a large network with many users and higher link speeds be intrinsically different to that of a smaller network? Somewhat surprisingly we find that, in several essential ways, one can mimic a large network using a suitably scaled-down version. The key is to find suitable ways to scale down the network and extrapolate performance.

Our main results are: (i) For networks which carry long-lived TCP-like flows arriving in clusters, and which are controlled by a variety of active queue management schemes, performance measures such as queueing delay and drop probability are left virtually unchanged. In Section II we verify this using the differential-equation type models developed in [17]. Such models have been widely used in designing control algorithms and for conducting control-theoretic analyses of network behavior. (ii) For networks in which flows arrive at random times and whose sizes are heavy-tailed, we find a different scaling to that in Section II leaves the distribution of the number of active flows and of their normalized transfer times unchanged. These latter networks are representative of the Internet. A simple theoretical argument, using the M/GI-type models proposed in [14], reveals that the method we suggest for “SHRiNKing” networks in which flows arrive at random times will be widely applicable (i.e. for a variety of topologies, flow transfer protocols, and queue management schemes). By contrast, we find that the theoretical underpinning for SHRiNKing networks at which flows arrive in clusters depends on the type of queue management scheme used at the routers.

A motivating example: Before continuing, we consider a simple example which illustrates the key points: the $M/M/1$ queue. Suppose jobs arrive at a queue according to a Poisson

process of rate λ , and that service times are independent and exponential with rate $\mu > \lambda$. Let $Q(t)$ be the number of jobs in the system at time t .

Now scale the system as follows: Sample the arriving jobs, keeping each job with probability α , independent of the others, so that sampled arrivals form a Poisson process of rate $\alpha\lambda$. Consider feeding the sampled arrivals to a separate queue whose server runs slower than the first by a factor α . This is equivalent to multiplying the service times by a factor $1/\alpha$ (so that they are rate $\alpha\mu$ exponentials), and the second queue is also $M/M/1$. If $\tilde{Q}(t)$ is the number of jobs in the slower queue at time t , then it is not hard to see that $\tilde{Q}(t) = Q(\alpha t)$ in distribution. That is, the evolution of the slower queue is statistically equivalent to that of the original queue slowed down in time by a factor α . This is because the queue-size process in an $M/M/1$ queue is a birth-death chain. The birth and death rates in the original queue are λ and μ respectively; while they are $\alpha\lambda$ and $\alpha\mu$ in the slower queue.

As a consequence, in equilibrium, the marginal distributions of the two queues are equal: i.e. $P(Q \geq n) = (\lambda/\mu)^n = (\alpha\lambda/\alpha\mu)^n = P(\tilde{Q} \geq n)$. Thus, we have inferred the distribution of queue-size, and hence of delay, in the original high-speed system by looking at a smaller-scale version.

It is natural to be skeptical of the relevance of these results. After all, they assume Poisson input traffic, whereas Internet packet traffic exhibits long-range dependence. Even more, these are open networks (the rate of arrivals is independent of current network congestion), quite different from the window flow-controlled Internet.

Nevertheless we find in the coming sections that the SHRiNK approach can be applied to IP networks, because it relies on factors other than packet level statistics: we shall see that it relies on certain fundamental scalability properties of networks.

II. IP NETWORKS WITH LONG-LIVED FLOWS

In this section we explore how SHRiNK applies to IP networks used by long-lived TCP-like flows that arrive in clusters, and controlled by queue management schemes like RED.

First, we explain in general terms how we sample traffic, scale the network, and extrapolate performance.

Sampling is simple. We sample a proportion α of the flows, independently and without replacement.

We scale the network as follows: link speeds and buffer sizes are multiplied by α . The various AQM-specific parameters are also scaled, as we will explain in the following section II-A. The network topology is unchanged during scaling. In the cases we study, performance measures such as average queueing delay are virtually the same in the scaled and the unscaled system.

Our main theoretical tool is the recent work on fluid models for TCP networks [17]. While [17] shows these models to be reasonably accurate in most scenarios, the range of their applicability is not yet fully understood. However, in some cases the SHRiNK hypothesis holds even when the fluid model is not accurate, as shown in Section II-A.3.

A. RED

The key features of RED are the following two equations, which together specify the drop (or marking) probability. RED maintains a moving average q_a of the instantaneous queue size q ; and q_a is updated whenever a packet arrives, according to the rule

$$q_a := (1 - w)q_a + wq,$$

where the w parameter determines the averaging window. The average queue size determines the drop probability p , according to the equation

$$p_{\text{RED}}(q_a) = \begin{cases} 0 & \text{if } q_a < \min_{th} \\ p_{\text{max}} \left(\frac{q_a - \min_{th}}{\max_{th} - \min_{th}} \right) & \text{if } \min_{th} \leq q_a < \max_{th} \\ 1 & \text{if } q_a > \max_{th} \end{cases} \quad (1)$$

We now explain how we scale the parameters p_{max} , $\min_{th}$, $\max_{th}$ and w . We will multiply $\min_{th}$ and $\max_{th}$ by α . Recall that we are multiplying the buffer size by α : thus $\min_{th}$ and $\max_{th}$ are fixed to be a constant fraction of the buffer size. (This is in accord with the recommendations in [10].) We will keep p_{max} fixed at 10%, so that the drop probability is kept under 10% as long as the buffer is slightly congested. The averaging parameter w takes more thought. We shall multiply it by α^{-1} . The intuition is this: when the network is scaled down, packets arrive less frequently, so q_a is updated less often, so we make the updates larger in magnitude. Simulation and theory, described below, both indicate that this choice of scaling is natural for extrapolating performance.

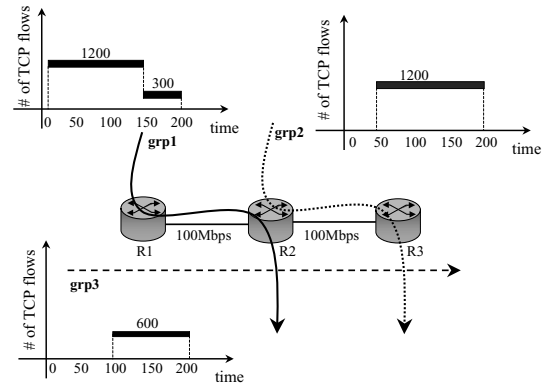


Fig. 1. Basic network topology and flow information

1) THE BASIC SETUP: We consider two congested links in tandem, as shown in Figure 1. There are three routers, $R1$, $R2$ and $R3$; and three groups of flows, $grp1$, $grp2$, and $grp3$, with group i connecting sources in src_i to receivers in rcv_i . The link speeds are 100Mbps and the buffers can hold 8000 packets. The RED parameters are $\min_{th} = 1000$, $\max_{th} = 2500$ and $w = 0.000005$. For the flows: $grp0$ consists of 1200 TCP flows each having a propagation delay of 150ms, $grp1$ consists of 1200 TCP flows each having a propagation delay of 200ms, and $grp2$

consists of 600 TCP flows each having a propagation delay of 250ms. The flows switch on and off as shown in the timing diagram in Figure 1. Note that 75% of *grp0* flows switch off at time 150s.

This network is scaled-down by factors $\alpha = 0.1$ and 0.02 , and the parameters are modified as described above.

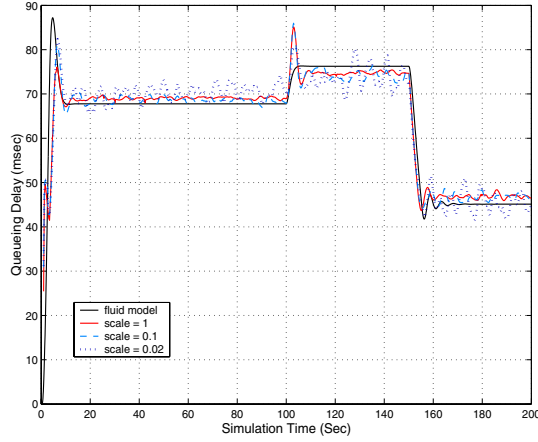


Fig. 2. Basic Setup: Average Queueing Delay at Q1

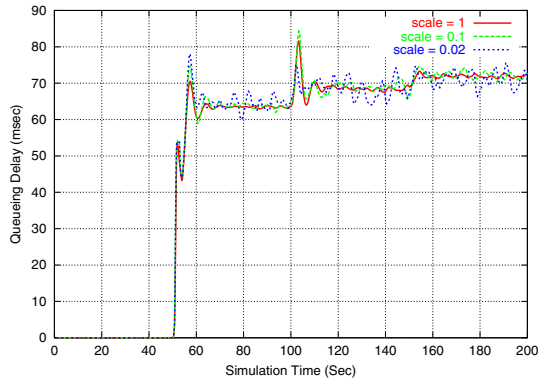


Fig. 3. Basic Setup: Average Queueing Delay at Q2

We plot the average queueing delay at Q1 and Q2 as a function of time in Figures 2 and 3. The drop probability at Q1 is shown in 4. Due to limited space, we omit the plot of drop probability for Q2 whose behavior is similar to that of Q1. We see that *the queueing delay is almost identical at different scales*. (It is worth noting that it is the queueing delay which is unchanged during scaling, whereas in the $M/M/1$ model it was the queue size distribution.)

Since the drop probability is also the same in the scaled and unscaled systems, the dynamics of the TCP flows are the same. In other words, an individual flow which survives the sampling process essentially cannot tell whether it is in the scaled or unscaled system.

2) THEORY: We now show that these simulation results are supported by the recently-proposed theoretical fluid model of TCP/RED [17].

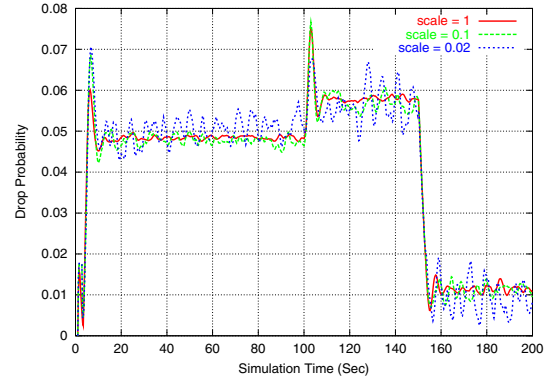


Fig. 4. Basic Setup: Drop Probability at Q1

Consider N flows sharing a link of capacity C . Let $W_i(t)$ and $R_i(t)$ be the window size and round-trip time of flow i at time t . Here $R_i(t) = (T_i + q(t))/C$, where T_i is the propagation delay and $q(t)$ is the queue size at time t . Let $p(t)$ be the drop probability at time t , and $q_a(t)$ the average queue size used by RED.

The fluid model describes how these quantities evolve; or rather, since these quantities are random, the fluid model describes how their expected values evolve. Let $\bar{X}$ be the expected value of random variable X . Then the fluid model equations are these:

$$\frac{d\bar{W}_i(t)}{dt} = \frac{1}{R_i(\bar{q}(t))} - \frac{\bar{W}_i(t)\bar{W}_i(t - \tau_i)}{1.5R_i(\bar{q}(t - \tau_i))} \bar{p}(t - \tau_i) \quad (2)$$

$$\frac{d\bar{q}(t)}{dt} = \sum_{i=1}^N \frac{\bar{W}_i(t)}{R_i(\bar{q}(t - \tau_i))} - C \quad (3)$$

$$\frac{d\bar{q}_a(t)}{dt} = \frac{\log(1 - w)}{\delta} \bar{q}_a(t) - \frac{\log(1 - w)}{\delta} \bar{q}(t) \quad (4)$$

$$\bar{p}(t) = p_{\text{RED}}(\bar{q}_a(t)) \quad (5)$$

where $\tau_i = \tau_i(t)$ solves $\tau_i(t) = R_i(\bar{q}(t - \tau_i(t)))$, δ is the average packet inter-arrival time, and p_{RED} is as in (1).

Remarks: While the applicability of these equations is not yet fully understood, [17] indicates that empirically they are reasonably accurate. Also, note that we have the constant 1.5 in (2), not 2 as in [17]. This change improves the accuracy of the fluid model [21].¹ Finally, note that while these equations describe a single link, the extension to networks is straightforward, and is given in [17].

Returning to the differential equations, suppose we have a solution to these equations

$$(\bar{W}_i(\cdot), \bar{q}(\cdot), \bar{q}_a(\cdot), \bar{p}(\cdot)).$$

Now, suppose the network is scaled and denote by C' , N' , etc the parameters of the scaled system. When the network is scaled, the fluid model equations change, and so the solution

¹Due to space limitations, we omit the derivation here. The complete proof can be found in [21], and will be published in a longer version of the paper.

changes. Let $(\bar{W}'_i(\cdot), \bar{q}'(\cdot), \bar{q}'_a(\cdot), \bar{p}'(\cdot))$ be the solution of the scaled system. We claim that, in fact,

$$(\bar{W}'_i(\cdot), \bar{q}'(\cdot), \bar{q}'_a(\cdot), \bar{p}'(\cdot)) = (\bar{W}_i(\cdot), \alpha \bar{q}(\cdot), \alpha \bar{q}_a(\cdot), \bar{p}(\cdot)).$$

If our claim is established, we will obtain that the queueing delay $\bar{q}'/C' = \alpha \bar{q}/\alpha C$ is identical to that in the unscaled system. Note also that the drop probability is the same in each case ($\bar{p}(t) = \bar{p}'(t)$). Thus, we will have theoretical support for the observations in the previous section.

Establishing the claim. We will proceed through the fluid model equations one by one. Consider first (2). Note that $R'_i(\bar{q}'(t)) = T_i + \bar{q}'/C' = T_i + \alpha \bar{q}/\alpha C = R_i(\bar{q}(t))$, so that $\tau'(t) = \tau(t)$. Hence

$$\frac{d\bar{W}'_i(t)}{dt} = \frac{1}{R'_i(\bar{q}'(t))} - \frac{\bar{W}'_i(t)\bar{W}'_i(t-\tau')}{1.5R'_i(\bar{q}'(t-\tau'))} \bar{p}'(t-\tau').$$

Consider next (3). Suppose for simplicity that all flows have identical routes. Then the W_i are statistically identical, hence the expectations $\bar{W}_i$ are all equal. So we can rewrite the equation as

$$\frac{d\bar{q}(t)}{dt} = \frac{N\bar{W}_1(t)}{R_1(\bar{q}(t-\tau'))} - C.$$

It is then easy to see that

$$\begin{aligned} \frac{d\bar{q}'(t)}{dt} &= \alpha \frac{d\bar{q}(t)}{dt} \\ &= \frac{N'\bar{W}'_1(t)}{R'_1(\bar{q}'(t-\tau'))} - C'. \end{aligned}$$

This extends to the case of multiple groups of flows with different routes, provided we sample a proportion α from each group.

Consider next (4). Recall that $w' = w/\alpha$. Note that the average packet inter-arrival time increases as the number of flows and the capacity decrease, in proportion $\delta' = \delta/\alpha$. Making the approximation $\log(1 - w/\alpha) \approx \log(1 - w)/\alpha$, good for small w , we see that $\log(1 - w')/\delta' \approx \log(1 - w)/\delta$, and hence that

$$\frac{d\bar{q}'_a(t)}{dt} \approx \frac{\log(1 - w')}{\delta'} \bar{q}'_a(t) - \frac{\log(1 - w')}{\delta'} \bar{q}'(t).$$

In fact, we chose $w' = w/\alpha$ so that this equation would be satisfied, allowing us to scale properly.

Consider finally (5). Recall that $p'_{\max} = p_{\max}$, and that $\min'_{th} = \alpha \min_{th}$ and $\max'_{th} = \alpha \max_{th}$. It is then clear that

$$\bar{p}'(t) = p'_{\max} \left(\frac{\bar{q}'_a(t) - \min'_{th}}{\max'_{th} - \min'_{th}} \right).$$

This establishes the claim.

Figure 5 presents the solution of the fluid model for the queueing delay at Q1 under the scenario of Figure 1 for the scale parameters $\alpha = 1$ and 0.1. As can be seen, both the solutions are virtually identical, providing a numerical illustration of the scaling property of the differential equations established above.

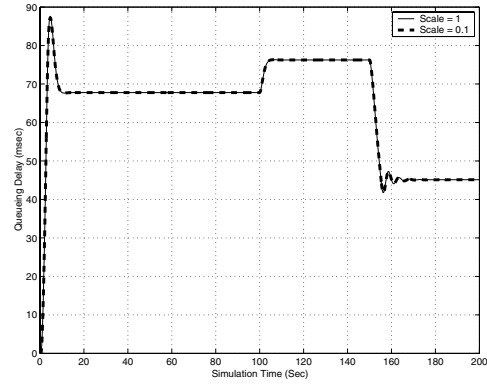


Fig. 5. Fluid model predicts scaling behavior

Remarks: It is worth remarking on a theoretical nicety related to the scaling property of these differential equations. If they had been derived from a limiting procedure in which the number of users, link capacities and buffer sizes all increase proportionally with N , then the scaling behavior would have been entirely expected (one has only to set N to equal αN before taking limits). However, they have been derived via a different route in [17]: by assuming that packet drops occur as a Poisson process. Therefore, the scaling property they exhibit is rather stunning. It strongly suggests that, in fact, they describe the behavior of the network in a large- N limit.

We also draw attention to some interesting features of all of the performance-related figures in this section. Note that transients are pretty well mimicked at the smaller scales. Also note that the smaller scale plots look more jagged, as if they are a noisy version of the original plots. The last point would be an easy consequence of a limit theorem: If in the large- N limit the behavior of the network is describable using deterministic differential equations, then away from the limit (at smaller and smaller scales) a corresponding Central Limit Theorem would suggest that the noise would be proportional to $1/\sqrt{\alpha}$.

3) WITH FASTER AND SLOWER LINKS: Suppose we alter the basic setup, by increasing the link speeds to 500Mbps, while keeping all other parameters the same. Figure 6 (zoomed in to emphasize the point) illustrates that, once again, scaling the network does not alter the queueing delay. Note that under these conditions the queue oscillates. There have been various proposals for stabilizing RED [15], [20]. We are not concerned with stabilizing RED here: we mention this case to show that SHRINK can work whether or not the queue oscillates.

Suppose we instead alter the basic setup, by decreasing the link speeds to 50Mbps, while keeping all other parameters the same. Once again, scaling the network does not alter the queueing delay. For such a simulation scenario, especially in the time frame 100sec-150sec, the fluid model is not a good fit (see Figure 7). This is not unexpected [25]: actual window and queue sizes are integer-valued whereas fluid solutions are real-valued; rounding errors are non-negligible when window sizes are small as is the case here. The range of applicability of the fluid model is not our primary concern in this paper: we mention this case

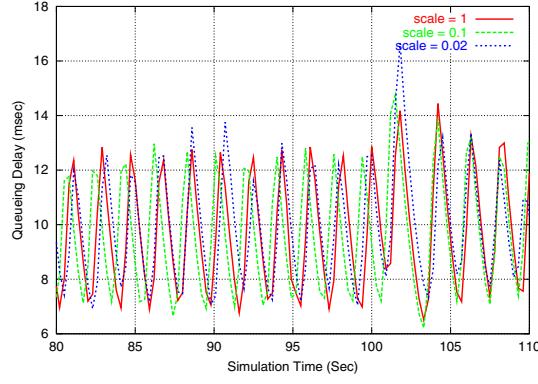


Fig. 6. With faster links: Average queueing delay at Q1 (zoomed in)

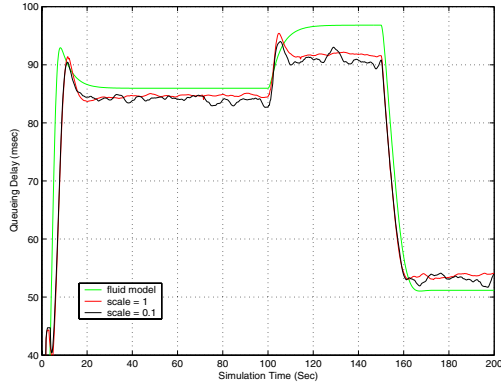


Fig. 7. With slower links: Average queueing delay at Q1

to show that SHRiNK can work whether or not the fluid model is appropriate.

4) IN A MORE COMPLEX NETWORK: As a further validation, we test SHRiNK in a more complex network, shown in Figure 8. There are seven routers R1 to R7. Links R1–R2, R2–R3, R1–R5, R3–R5 and R4–R5 run at 150 Mbps, links R1–R4 and R5–R6 run at 100 Mbps, and all other links run at 50 Mbps. The traffic is a mixture of UDP and web flows; and long-lived TCP, AIMD and Binomial [2] flows. These last types have the following common form: on receiving an acknowledgement, increase the congestion window w by aw^{n-1} (TCP uses $a = 1$, $n = 0$), and on incurring a mark/drop, decrease w by bw^m (TCP uses $b = w/2$, $m = 1$). The parameters $(a, n; b, m)$ describe each class.

We omit a detailed description of all the flows, except those traversing link R1–R5 whose queueing dynamics are shown in Figure 9. Link R1→R5 carries 1000 long-lived flows, divided into five groups: 200 normal TCP, 200 AIMD (1, 0; .1, 1), 200 AIMD (2, 0; .5, 1), 200 Binomial (1, 1; .5, 1) and 200 Binomial (1.5, –1; .5, 1). The links are controlled by RED with $min_{th} = 1000$, $max_{th} = 2000$ and $w = 0.000005$. As before, we see that scaling the network does not affect the queueing delay.

To illustrate the potential savings in resources, we report the CPU time to run each of the simulations. We simulated the original system, and logged the start times of flows, the sizes of

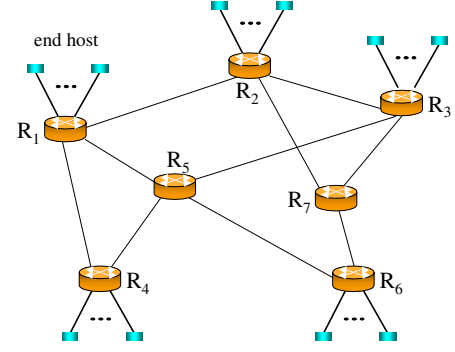


Fig. 8. A more complex topology

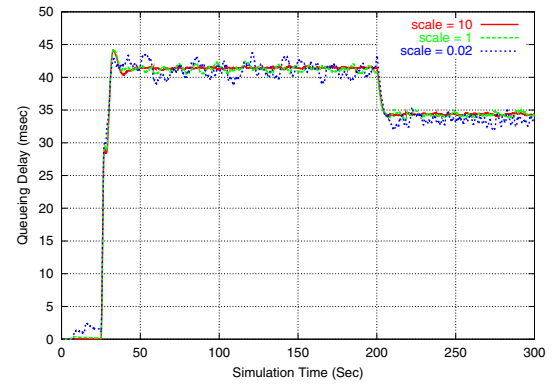


Fig. 9. In a more complex network: Average queueing delay at R1–R5

sessions, and so on. This data was sampled, and used to drive simulations at various scales α . The CPU times are: 3752.50 secs at $\alpha = 1$, 172.85 secs at $\alpha = .1$, and 49.13 secs at $\alpha = 0.02$.

B. Proportional-Integral (PI) Controller

A different AQM scheme is the PI controller [16], which attempts to stabilize the queue size around a given target value. The PI controller drops/marks packets with a probability p which is updated periodically by

$$p(t + \delta t) = p(t) + a(q(t + \delta t) - q_{target}) - b(q(t) - q_{target}). \quad (6)$$

Here, q is the instantaneous queue size, q_{target} is the target queue size, δt is the update timestep (here fixed at 0.01s), and a and b are arbitrary parameters.

We first explain how we will scale the network. As usual, let a' etc. be the scaled parameters. We will sample a fraction α of the flows, and set $a' = a/\alpha$, $b' = b/\alpha$ and $q'_{target} = \alpha q_{target}$. (This is in accordance with the design rules in [16].)

We simulated the basic setup of Section II-A, replacing RED by the PI controller. We use $a = 8.8681 \times 10^{-7}$ and $b = 8.7427 \times 10^{-7}$, as suggested in [16]. We set q_{target} to be 1750 packets, which is half-way between our min_{th} and max_{th} parameters from the last section.

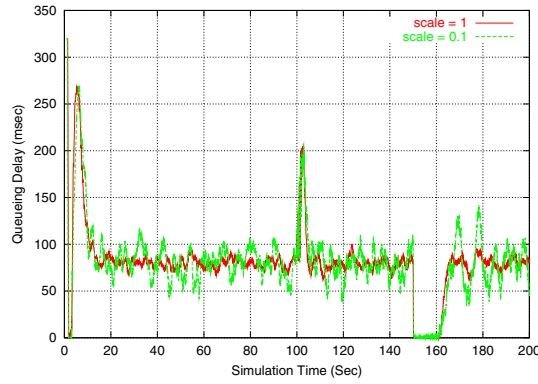


Fig. 10. PI Controller: Average queueing delay at Q1

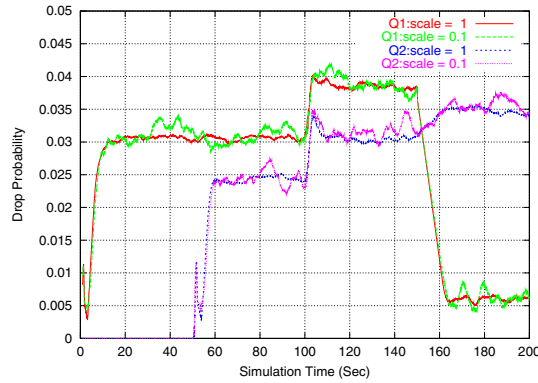


Fig. 11. PI Controller: Drop probabilities at Q1 and Q2

Figure 10 shows the average queueing delay at different scales for Q1. We see that scaling the network does not affect queueing delay, at least in steady state. There are some spikes when the load changes abruptly, and the small-scale network shows slightly larger spikes. Figure 11 shows that neither is the drop probability affected by scaling the network.

We can again use the fluid model to understand this behaviour. To obtain the fluid model for the PI controller, we simply replace (4) and (5) in the fluid model by the fluid analog of (6): the expected drop probability $\bar{p}$ evolves according to

$$\frac{d\bar{p}}{dt} = -b\frac{d\bar{q}}{dt} + (b-a)(\bar{q}(t) - q_{\text{target}}).$$

As before, by our choice of scaling,

$$\frac{d\bar{p}}{dt} = \frac{d\bar{p}'}{dt} = -b'\frac{d\bar{q}'}{dt} + (b' - a')(\bar{q}'(t) - q'_{\text{target}}).$$

Thus the fluid model also scales.

C. Summary

In all the examples we have studied in this section—with heterogeneous end-systems, with different of active queue management policies, and with a range of system parameters—we have found that basic performance measures such as queueing delay are left unchanged, when we sample the input traffic and

scale the network parameters in proportion. This conclusion is supported by the theory of fluid models, and even holds where the fluid models fail. A notable exception is provided by the queue management scheme DropTail, as described next.

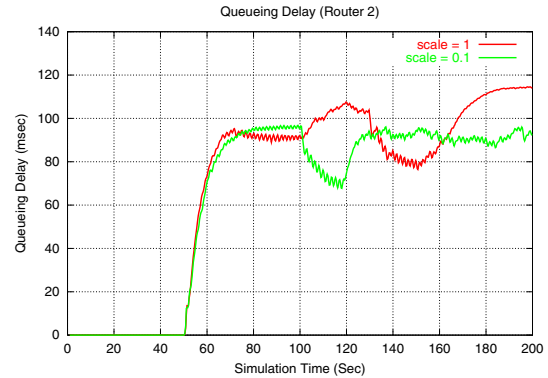


Fig. 12. DropTail: Average queueing delay at Q2

DropTail: Consider the basic network setup of Section II-A, and suppose that the routers use DropTail instead of RED. Figure 12 shows the average queueing delay at Q2. Clearly, the queueing delays at different scales do not match. DropTail drops all the packets that arrive at a full buffer. As a result, it could cause a number of consecutive packets to be lost. These bursty drops underlie the failure of the scaling hypothesis in this case, as explained in [22]. Separately, note that when packet drops are bursty and correlated, the assumption that packet drops occur as a Poisson process (see [17]) is violated and the differential equations become invalid. The connection between these two phenomena (the failure of the scaling hypothesis and the invalidation of the differential equation models) is explored in [22].

III. IP NETWORKS WITH SHORT AND LONG FLOWS

It has been shown that the size distribution of flows on the Internet is heavy-tailed [27]. Hence, Internet traffic consists of a large fraction of short flows, and a small fraction of long flows that carry most of the traffic. Also, it has been recently argued that since network sessions arrive as a Poisson process [9], [19], [23],² network flows are *as if* they were Poisson [14]. (In particular, the equilibrium distribution of the number of flows in progress at any time can be obtained by assuming that flows arrive as a Poisson process.) We take these observations into account and study the scaling behavior of IP networks carrying heavy-tail distributed, Poisson flows. Such networks are a plausible representation of today's Internet.

A. Sampling and Scaling

We start with sampling: Due to the tremendous increase in the volume and speed of network traffic, it is very expensive

²That network sessions are Poisson is not surprising since a Poisson process is known to result from the superposition of a large number of independent user processes.

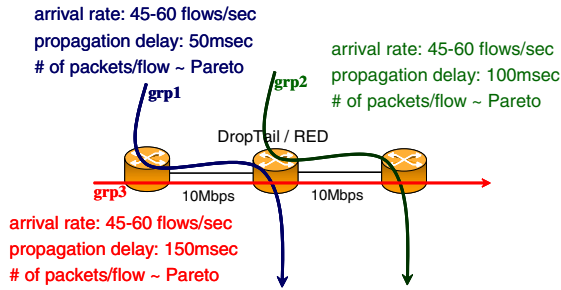


Fig. 13. Basic network topology and flow information.

to sample packets. At the other end of the spectrum, one may sample network sessions, e.g. modem calls, ftp, telnet, or web sessions. However, sampling sessions is hard in practice, because only end users have enough information to distinguish between different sessions. Hence, we choose to sample network flows.³ This reduces the traffic we have to deal with, and is easy to implement in practice.

A second issue related to sampling is: How are the network flows sampled? The method samples the flows, choosing each one with probability α , all choices being independent, in an i.i.d. fashion with some probability α . The last issue with sampling is: Where are flows sampled? The method samples at network entry points, e.g. at edge routers.

What is left is to describe how to obtain the small replica of the original network. This is done as follows: (i) link capacities are reduced by a factor α , (ii) propagation delays are scaled up by a factor $1/\alpha$, and (iii) protocol timeouts are also scaled up by the same factor. Intuitively, these steps aim to slow down the speed of the network. This will become more clear in Section III-C.

B. Simulation Results

In this section we investigate how accurately SHRiNK can predict the performance of IP networks from small-scale replicas, using the network simulator ns-2 [18].

For simplicity, consider the topology illustrated in Figure 13. In Section III-C we establish that the results are independent of the particular topology. There are three routers, R_1 , R_2 and R_3 , two links in tandem, and three groups of flows, $grp1$, $grp2$, and $grp3$. The link speeds for are 10Mbps.

Routers use either the Random Early Detection (RED) or the DropTail queue management schemes. The RED parameters are $min_{th} = 100$, $max_{th} = 250$ and $w = 0.00005$. When using DropTail, the buffer can hold 200 packets.

Within each group, flows arrive as a Poisson process with rate λ . We vary λ to study both uncongested and congested scenarios. (We use the ns-2 built-in routines to generate web sessions consisting of a single object each. This is what we

³Notice that in accordance with the usual practice [8], [12], [13], packets are said to belong to the same flow if they have the same source and destination IP address, and source and destination port number. A flow is "on" if its packets arrive more frequently than a timeout of some seconds. This timeout is usually set to something less than 60 seconds.

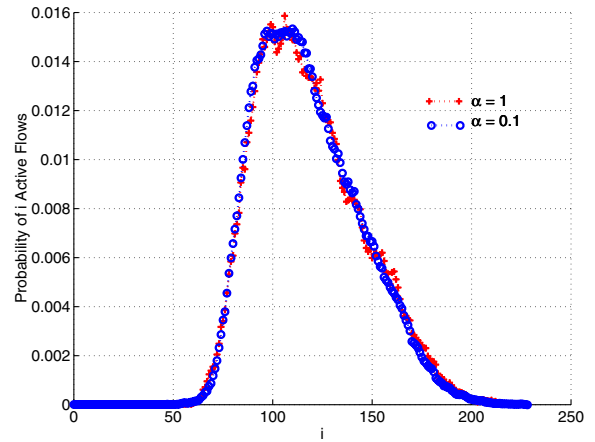


Fig. 14. Distribution of number of active flows on the first link (RED).

call a flow in the simulations.) Each flow consists of a Pareto-distributed number of packets with average size 12 packets and shape parameter equal to 1.2. The packet size is set to 1000 bytes. The propagation delay of each flow of $grp1$, $grp2$, and $grp3$, is 50msec, 100msec, and 150msec respectively.

We run the experiments for scale factors $\alpha = 1$ and 0.1, and compare the distribution of the number of active flows as well as the histogram of the normalized delays of the flows in the original and the scaled system. (The normalized delays are the flow transfer times multiplied by α .) We also compare more detailed performance measures such as the distribution of active flows that are less than some size and belong to a particular group, and the distribution of the packet buffer occupancies. As will be shown in Section III-C, the method can predict the marginal and joint distributions of a large number of performance measures.

Due to limitations of space, we do not present results when the links are uncongested, but only compare distributions for the more challenging and realistic case of congested networks. For a full exposition of the simulation results, interested readers are referred to [24].

Accordingly, to induce congestion, flow arrival rates are set to 60 flows/sec within each group. Flows experience drops that account for up to 5% of the total traffic. We first present simulations where all three routers use RED.

Figure 14 plots the distribution of the number of active flows in the first link. The two distributions match. A similar conclusion is obtained at the second link.

Figure 15 plots the histogram of the normalized delays of the flows of $grp1$. To generate the histogram, we use normalized delay chunks of 10msec each. There are 150 such delay chunks in the plot, corresponding to flows having a normalized delay of 0 to 10msec, 10msec to 20msec, and so on. The last delay chunk is for flows that have a normalized delay of at least 1.5sec. The plot reveals that the distribution of the normalized delays match. The results for the other two groups of flows are similar.

The peaks in the delay plots are due to the TCP slow-start

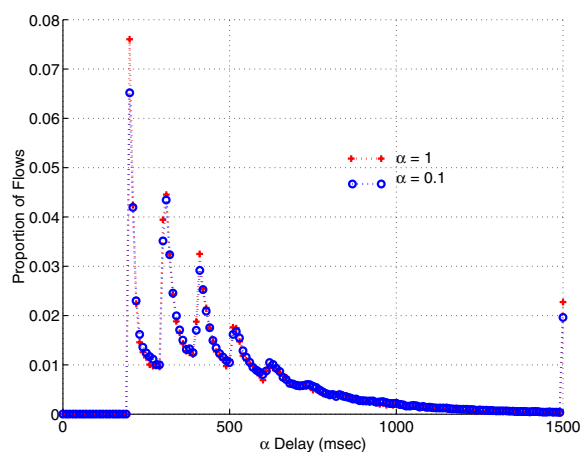


Fig. 15. Histogram of normalized delays of *grp1* flows (RED).

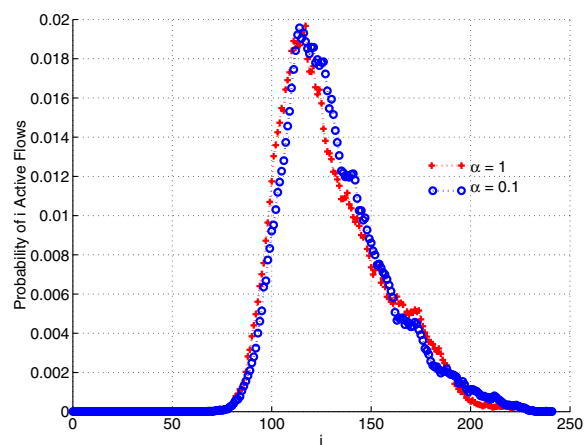


Fig. 17. Distribution of number of active flows on the second link (DropTail).

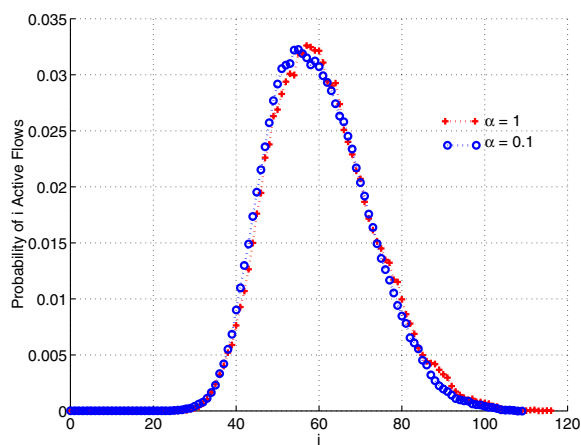


Fig. 16. Distribution of number of active *grp3* flows with size less than 12 packets (RED).

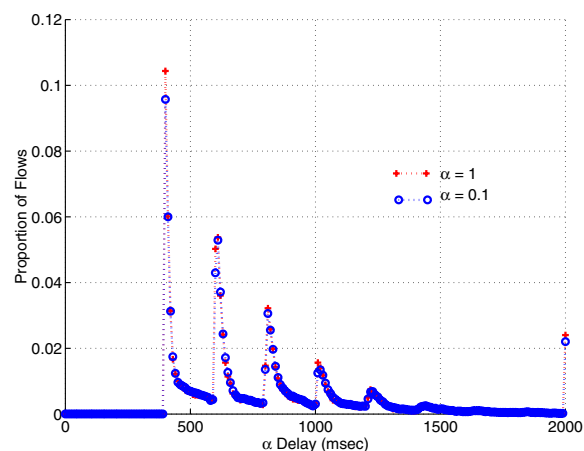


Fig. 18. Histogram of normalized delays of *grp2* flows (DropTail).

mechanism. The left-most peak corresponds to flows which send only one packet and face no congestion. These flows only have to wait for the setup of the TCP connection. (Hence, for example, in Figure 15 where propagation delays are 50msec, the normalized delay for these flows is a bit more than 200msec accounting for SYN, SYN-ACK, the data packet, the ACK for the packet, and insignificant transmission and queueing delays.) The portion of the curve between the first and second peaks corresponds to flows which send only one packet and face congestion (but no drops). The next peak corresponds to flows which send two or three packets and face no congestion. These flows have to wait for an additional round trip time for the acknowledgment for the first packet to arrive. The third peak corresponds to flows which send between four and seven packets and face no congestion, and so on.⁴

What about more detailed performance measures? As an example, we compare the distribution of active flows belonging to *grp3* that are less than 12 packets long. Figure 16 compares the

⁴Recall: Whenever an acknowledgment arrives, TCP senders double their window sizes.

two distributions from the original and scaled system. Again, the plots match.

We will now investigate if distributions scale when DropTail is used. Figure 17 plots the distribution of the number of concurrently active flows in the second link between routers *R2* and *R3* when all routers use DropTail. It is evident from the plot that the two distributions match as before. A similar scaling holds for the other link.

Figure 18 plots the histogram of the normalized delays of the flows of *grp2* when DropTail is employed. The distributions match as before. A similar scaling holds for the other two groups of flows.

So far, the method has successfully predicted the distribution of various performance measures at the *flow* level. Figure 19 compares the distribution of the number of *packets* at the first queue, which uses RED, in the original and scaled network. As evident from the plot, the method can also predict the distribution of the queue occupancies.

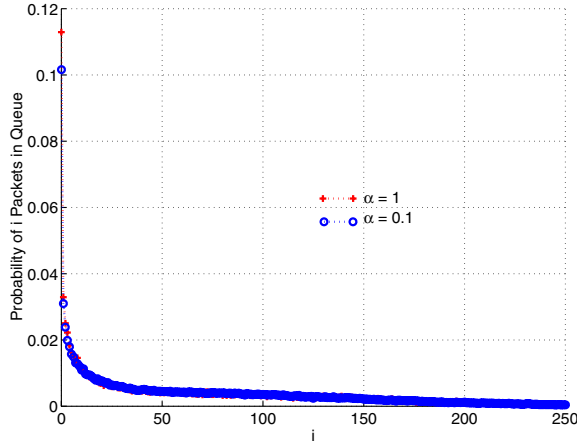


Fig. 19. Distribution of number of packets in $R1$.

C. Understanding SHRiNK

Recall that flows arrive as a Poisson process, bearing sizes drawn independently from a common (Pareto) distribution.⁵ The above observation allow the use of M/GI -type queueing models to analyze SHRiNK.

By the state of the network at time t we mean the total information that is needed to resume the evolution of the network from time t onwards, given input data (flow arrival times and sizes) after time t . For example, the state consists of information about currently active flows: their transferred packets, their packets in transit, and where they are at time t , etc. Write $S(t)$ for the state at time t . If $I(t)$ denotes the input data to the system then $S(t)$ is some function, $\mathcal{F}$, of the input until time t . Symbolically, $S(t) = \mathcal{F}[I(s), s < t]$. We shall abbreviate this to $S(t) = \mathcal{F}[I(\cdot)]$. Note that $\mathcal{F}$ is some complicated function depending on transport protocols, queue management schemes, and other network- and user-specific details.

Theorem 1: Let $S(t)$ be the state of the original network at time t , and $\tilde{S}(t)$ be the state of the scaled network at time t . Then $S(\alpha t) \stackrel{d}{=} \tilde{S}(t)$, i.e. the same in distribution.

Proof: Let $I(\cdot)$ and $\tilde{I}(\cdot)$ be the inputs to the original and scaled systems, respectively. Let $\mathcal{F}^o$ and $\mathcal{F}^s$ denote the functions corresponding to the original and scaled (slowed-down) networks. This gives $S(t) = \mathcal{F}^o[I(\cdot)]$ and $\tilde{S}(t) = \mathcal{F}^s[\tilde{I}(\cdot)]$. Our method of proof consists of constructing a third system, the “time-stretched system”, which is obtained by applying the input $\hat{I}(t) \equiv I(\alpha t)$ to the scaled system. Thus, the time-stretched system has as input the input of the original system stretched out in time by a factor α . That is, flow f , of size s , arrives to the original system at time t iff it arrives, again with size s , to the time-stretched system at time t/α . It is a simple, but far-reaching, property of the Poisson process that $\tilde{I}(\cdot) \stackrel{d}{=} \hat{I}(\cdot)$, since sampling a proportion α of the points of a rate λ Poisson process will yield a rate $\alpha\lambda$ Poisson process. And the independent nature of the sampling process does not destroy the i.i.d. nature

⁵Note that whereas flow sizes are independent, their delays (equal to their total transfer times) are usually dependent.

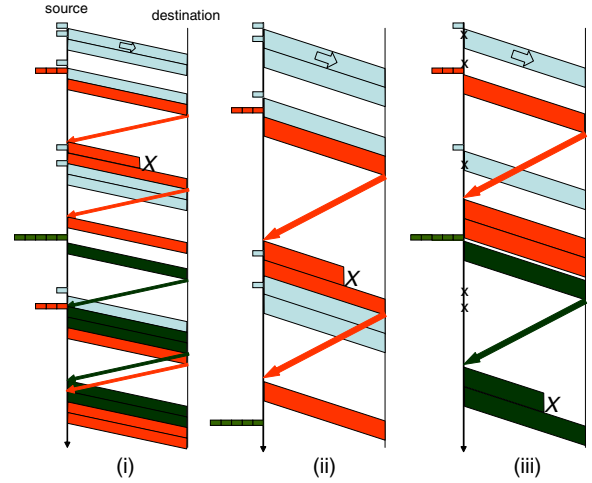


Fig. 20. Time evolution of: (i) the original, (ii) the time-stretched, and (iii) the scaled system.

of the flow sizes.

Let $\hat{S}(t) = \mathcal{F}^s[\hat{I}(\cdot)]$ denote the state of the time-stretched system at time t . We shall show that the following identity is satisfied at every time t :

$$\hat{S}(t) = S(\alpha t). \quad (7)$$

Establishing this will complete our proof since $S(\alpha t) = \hat{S}(t) = \mathcal{F}^s[\hat{I}(\cdot)] \stackrel{d}{=} \mathcal{F}^s[\tilde{I}(\cdot)] = \tilde{S}(t)$.

We now establish the identity at (7). Consider the consequences of our method of scaling (slowing down) the original network: reducing link speeds by a factor α will increase queueing delays by factor $1/\alpha$, increasing propagation delays by a factor $1/\alpha$ will increase transmission times by $1/\alpha$. Since the total delay of a packet is the sum of its queueing and transmission times, we have effectively increased the delay of every packet by $1/\alpha$. This in turn increases the delay of every flow transfer time by a factor $1/\alpha$. It is now quite easy to see that much more is true: Since the networks are all discrete-event systems, clocked by transmissions and acknowledgements of packets, every event that occurred in the original system at time t will occur in the time-stretched system at time t/α . Therefore $S(\alpha t) = \hat{S}(t)$, and the theorem is proved.

Remark 1: It is instructive to consider an illustration of the three systems, as in Figure 20. The time evolution of each of the three systems is shown between some one source-destination pair. In each sub-figure, the corresponding input process is shown on the left axis. The graph of an input process denotes flow arrival times and their corresponding sizes. The lines going from right to left denote acknowledgments. Finally, the big “X”’s denote packet drops. The original system has an input process of $I(t)$. For the time-stretched system, packets have larger transmission and propagation delays, denoted by “fatter” parallelograms going from left to right and larger slopes respectively; and the input process, $\hat{I}(t)$, is a time-stretched version of $I(t)$. Notice that the time-stretched system is just a device for

the proof, it does not exist. The input of the scaled system, $\tilde{I}(t)$, is just a subsample of the flows of $I(t)$. The unsampled flows of $I(t)$ are denoted by tiny “x” ’s on the vertical axis of Figure 20(iii).

Remark 2: The theorem explains why the distributions of various performance measures match in distribution. Further, it shows that performance scaling involves speeding up the time, and this is why we compare normalized delays rather than delays. The proof of the theorem only relies on the assumptions about inputs (Poisson flow arrivals and i.i.d. sizes) and the fact that the network evolves as a discrete-event system. Therefore, when these assumptions are met, ⁶ SHRiNK is widely applicable for marginal, joint, steady-state and transient distributions of a large family of performance measures, for any network topology, transport protocol, and queue mechanism. Another consequence of Theorem 1 is that SHRiNK works for any value of α . Thus, networks can be slowed down arbitrarily. However, the smaller the α , the slower the network is, and the longer it takes for distributions to converge.

D. Applications

Since the method provides a way to deduce the performance of a fast network from a slowed-down replica, it can be used to reduce the cost of experimentation: Imagine a test-network with slow network interfaces, slow switches and routers, and cheap links, that is fed with a sample of the actual network traffic. ⁷ In this network one may experiment with new algorithms, protocols, and architectures, and extrapolate performance.

Another use of the method is the following: There has been a recent development of research prototypes and products [5] that record partial information about the network by sampling incoming traffic. SHRiNK offers a systematic way to reproduce offline the whole behavior of the network using this sample.

IV. WEB SERVER FARMS

In this section we briefly outline how SHRiNK may apply to web server farms. Since a rapid growth in the size and capacity of web server farms makes it increasingly difficult to take performance measurements and to evaluate new algorithms and architectures, if SHRiNK applies to web server farms it would help reduce this difficulty significantly.

How should server farms be scaled? Consider a web server farm with N servers each having speed s , as in Figure 21. Sample the requests for the original farm, retaining each independently with probability α . Feed the sampled traffic into a scaled-down farm consisting of either (i) a fraction α of the original web servers, or (ii) the same number of servers each having speed αs (see (i) and (ii) of Figure 21). Of interest is the closeness of the average response time, and the server throughput and capacity (maximum throughput) in the scaled system to those in the original system.

⁶We refer the reader [14] and [4] for an interesting discussion of the M/GI models and their role in generating the well-documented self-similar nature of network traffic.

⁷This network should also have larger propagation delay than the original. This can be achieved in software, or with delay-loops.

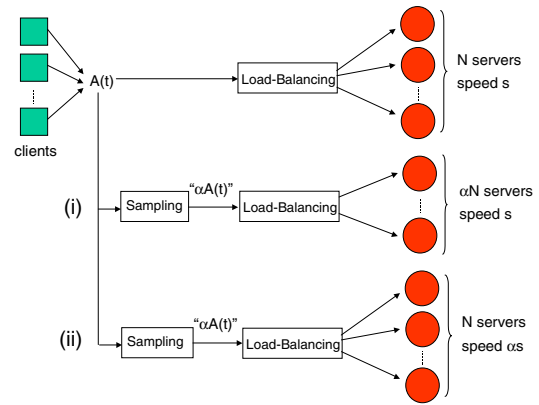


Fig. 21. Scaling a web server farm⁸: (i) scaling the number of servers, (ii) scaling the speed of the servers.

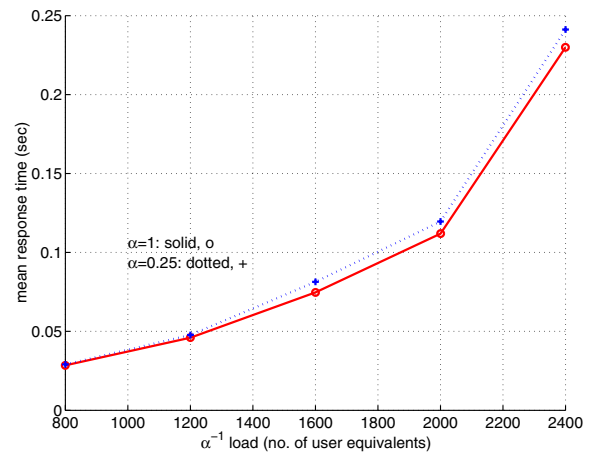


Fig. 22. Average response time, when sampling user equivalents.

We conducted some preliminary experiments using eight Linux machines configured with a Pentium III at 550MHz and 384MB of RAM, connected to a 100Mbps/sec switch. Four machines constitute the original farm and act as servers, each of which host one Apache 1.3.9 [1] web server. The other four machines act as clients, each of which run Surge [3] to generate HTTP requests.

Due to lack of space we only present results for the case where $\alpha = 0.25$, one scales the number of servers, the clients use HTTP1.1, load-balancing is a simple round-robin scheme, and both load-balancing and sampling take place at the user-equivalent level. (Surge uses the notion of user equivalents to generate sequences of requests similar to those generated by web sessions that stays “on” throughout the experiment.) Please refer to [24] for more experimental results.

Figures 22 and 23 show the average response time and the normalized server throughput as a function of the normalized

⁸This is a simplified picture of a farm, since the application-servers, the databases, and the switches used to interconnect the various components are absent.

