

(2)

AD-A210 683

DISCRETE-TIME CONVERSION FOR
SIMULATING FINITE-HORIZON MARKOV
PROCESSES

by

Bennett L. Fox and Peter W. Glynn

DTIC
ELECTE
AUG 02 1989
S D

TECHNICAL REPORT No. 34

May 1989

Prepared under the Auspices
of

U.S. Army Research Contract
DAAL-88-K-0063

-03-

Approved for public release: distribution unlimited.

Reproduction in whole or in part is permitted for any
purpose of the United States government.

DEPARTMENT OF OPERATIONS RESEARCH
STANFORD UNIVERSITY
STANFORD, CALIFORNIA

89 7 31 077

DISCRETE-TIME CONVERSION FOR SIMULATING FINITE-HORIZON MARKOV PROCESSES

Bennett L. Fox* Peter W. Glynn†

May 1989

Abstract. We estimate via simulation the expectation of certain integrals of functionals of continuous-time Markov chains over a finite horizon, fixed or random. By computing conditional expectations given the sequence of states visited (and possibly other information), we reduce variance. This is discrete-time conversion. We further increase efficiency by combining discrete-time conversion with stratification and splitting.

Key words. simulation, finite-horizon, continuous-time Markov chains, variance reduction.

AMS (MOS) subject classification. 60J25, 62M05, 90B99.

1 Introduction

Estimating expected cumulative "reward", possibly continuously discounted, up to a finite horizon τ has practical interest. This and other examples in our paper are special cases of the following: estimating expectations of integrals of functionals f of a continuous-time Markov chain X with respect to a weight

*Department of Mathematics, Campus Box 170, University of Colorado, 1200 Larimer St., Denver 80204. Part of this research was carried out while the author was at Cornell University on leave from the University of Montreal. The research of this author was supported in part by the Natural Sciences and Engineering Research Council of Canada.

†Department of Operations Research, Stanford University, Stanford, California 94305-4022. Part of this research was carried out while the author was at the University of Wisconsin - Madison. The research of this author was supported by the National Science Foundation under Grant ECS-8404809 and by the U.S. Army under Contract No. DAAG29-84-K-0030.



Accession For	
NTIS	CRA&I ☒
ERIC	TAB ☐
Unannounced	☐
Justification	
By	
Distribution /	
Availability Codes	
Dist	1. Available for 2. Not available
A-1	

function G . Here $f(s)$ is the "reward" rate when in state s . We deal with two general classes of such integrals. Both have the form

$$\int_0^\tau f(X(t))G(dt)$$

and differ only in the definition of τ . By selecting f and G appropriately, many standard problems become special cases as Sections 3 and 5 detail. We estimate the expectations of these integrals via simulation. This is carried out *efficiently* via *discrete-time conversion*: computing conditional expectations given the sequence of states visited (and possibly other information). Section 2 shows that the less we condition on, the more variance σ^2 is reduced. However, the work W to compute conditional expectations depends on what we condition on. Section 2 shows that we want to minimize $\sigma^2 E[W]$, even when the conditional expectation and the work to compute it are correlated. Especially in sections 4 and 6, we discuss implementation and estimate the order of magnitude of the work involved.

Section 3 provides theory for random-time horizons, where τ is the hitting time of a specified subset of the state space. We relate this to regenerative steady-state simulations. Section 5 provides theory for fixed-time horizons τ . Proofs are deferred to the appendix.

The estimators in sections 3 and 5 are springboards to significant improvements in sections 4 and 6 respectively. New ideas (but not new theory) are introduced in sections 4 and 6; these ideas make our theoretical results more important and more practical.

2 Preliminaries

As indicated in the introduction, much of this paper develops variance reduction techniques for continuous-time Markov chains based on appropriately "conditioning out" holding times. Such methods are but special cases of the general variance reduction technique known as conditional Monte Carlo.

To set the stage for a discussion of conditional Monte Carlo, we first mathematically characterize the efficiency of an estimator. Suppose that we wish to estimate a parameter α that can be expressed as $\alpha = ER$ for some r.v. R . The parameter α can be calculated by generating iid copies $R_1, R_2, \dots$ of R and forming their sample mean. Let $\beta(R_i)$ be the amount of computer time required to generate R_i . We assume (reasonably) that the pairs $(R_i, \beta(R_i))$ are iid, though R_i and $\beta(R_i)$ may be correlated.

Given a computer budget t , the number of observations completed within the budget is

$$N(t) = \max\{n \geq 0 : \sum_{k=1}^n \beta(R_k) \leq t\}.$$

The sample mean $r(t)$ formed with the above budget constraint is

$$r(t) = \begin{cases} \frac{1}{N(t)} \sum_{k=1}^{N(t)} R_k & ; N(t) \geq 1 \\ 0 & ; N(t) = 0. \end{cases}$$

The (asymptotic) efficiency of the estimator $r(t)$ is determined by how quickly $r(t)$ converges to α as the budget t goes to infinity. This rate of convergence is characterized by the central limit theorem for $r(t)$.

THEOREM 1. Suppose that $0 < E[\beta(R_i)] < \infty$ and that $\sigma^2(R_i) \triangleq \text{var } R_i < \infty$. Then,

$$t^{1/2}(r(t) - \alpha) \Rightarrow (E[\beta(R_1)]\sigma^2(R_1))^{1/2}N(0, 1)$$

as $t \rightarrow \infty$.

This theorem is elementary only when $\beta(R_1)$ is deterministic.

Theorem 1 suggests defining the asymptotic efficiency of the estimator $r(t)$ as the reciprocal of $E[\beta(R_1)]\sigma^2(R_1)$; Hammersley and Handscomb [(1964), p. 51] suggest the same figure of merit without providing theoretical justification. Thus, the efficiency of a simulation algorithm increases when the product $E[\beta(R_1)]\sigma^2(R_1)$ decreases. This permits obtaining an improved estimator in which either $E[\beta(R_1)]$ or $\sigma^2(R_1)$ is increased, so long as the product decreases.

We now apply these ideas to our discussion of conditional Monte Carlo. Suppose that R is an $\mathcal{F}$ -measurable r.v. and let $\mathcal{G}$ and $\mathcal{H}$ be sub- σ -fields of $\mathcal{F}$ such that $\mathcal{G} \subseteq \mathcal{H}$. Set $R_{\mathcal{G}} = E[R|\mathcal{G}]$, $R_{\mathcal{H}} = E[R|\mathcal{H}]$. If $E|R| < \infty$, it is well known that $\alpha = ER = ER_{\mathcal{G}} = ER_{\mathcal{H}}$. Hence, competing estimators for α , based on averaging replicates of $R_{\mathcal{G}}$ and $R_{\mathcal{H}}$, can be considered; such estimators are known as conditional Monte Carlo estimators. Let $r_{\mathcal{G}}(t)$, $r_{\mathcal{H}}(t)$ be the estimators formed from sample means of iid copies of the r.v.'s $R_{\mathcal{G}}$ and $R_{\mathcal{H}}$, respectively. From Theorem 2.1, we find that the efficiencies of the estimators $r(t)$, $r_{\mathcal{G}}(t)$, and $r_{\mathcal{H}}(t)$ are the reciprocals of the products $E[\beta(R)]\sigma^2(R)$, $E[\beta(R_{\mathcal{G}})]\sigma^2(R_{\mathcal{G}})$, and $E[\beta(R_{\mathcal{H}})]\sigma^2(R_{\mathcal{H}})$, respectively. The quantities $E[\beta(R)]$, $E[\beta(R_{\mathcal{G}})]$, and $E[\beta(R_{\mathcal{H}})]$ measure the mean computation time to generate the three types of observations. These quantities are hard to quantify precisely, since the computer time to generate an observation is implementation dependent (although their respective orders of magnitude sometimes can be estimated). On the other hand, we can order *a priori* the variances $\sigma^2(R)$, $\sigma^2(R_{\mathcal{G}})$, and $\sigma^2(R_{\mathcal{H}})$, as the following well-known proposition shows (see Problem 6 on page 305 of Chung (1974) for an equivalent statement).

PROPOSITION 1. If $\text{var } R < \infty$, then $\text{var } R_{\mathcal{G}} \leq \text{var } R_{\mathcal{H}} \leq \text{var } R$.

Proposition 1 states that conditioning on less reduces variance (by "integrating out" more randomness). An extreme case is $\mathcal{G} = \{\Phi, \Omega\}$; here $R_{\mathcal{G}} = ER$ has zero variance, but can't be found exactly.

This last example illustrates the compromises in choosing the appropriate conditioning variables upon which to apply the method of conditional Monte Carlo. The choice is a tradeoff between the amount of variance reduction obtained (as measured by $\sigma^2(R_{\mathcal{G}})$) and the implementation difficulty inherent in computing the conditional expectation $R_{\mathcal{G}}$ (as measured by $ER_{\mathcal{G}}$). This need for compromise is a focus of our discussion in subsequent sections.

3 Discrete-time conversion for random-time horizons: theory

Let $X = (X(t) : t \geq 0)$ be a non-explosive continuous-time Markov chain living on state space S and let f be a real-valued function defined on S . For $B \in \mathcal{S}$, let $T(B) = \inf\{t \geq 0 : X(t) \in B, X(t-) \notin B\}$ be the first "hitting time" of the subset B . We further let $G : [0, \infty) \rightarrow [0, \infty)$ be a right-continuous (deterministic) non-decreasing function, which then acts as an "integrator". In this section we apply the method of conditional Monte Carlo to compute $\alpha = E[I]$, where

$$I = \int_{[0, T(B))} f(X(t))G(dt). \quad (1)$$

Assuming that $f(x)$ is interpreted as the rate at which "cost" is incurred while the process occupies state x , the above estimation problem arises in several different settings.

SETTING 1. If $G(t) = t$, then I corresponds to the total cost accumulated by the process X up to time $T(B)$.

SETTING 2. If $G(t) = r^{-1}(1 - e^{-rt})$ where $r > 0$, then I corresponds to the r -discounted cost accumulated by X up to time $T(B)$.

SETTING 3. Given $T \geq 0$, suppose that the system is charged a cost which depends on the state occupied at time T . We assume that if the process hits B before T , no cost is charged to the system. If $f(x)$ is the cost incurred when the system occupies state x at time T , then the expected cost $\alpha = E[I]$, where I takes the form (1) and $G(t) = I(t \geq T)$. If $f \equiv 1$, then $\alpha = P\{T(B) > T\}$, so that α is the right tail of the cumulative distribution function of $T(B)$.

Although it may appear that the estimation problem considered here pertains only to finite-horizon problems, it turns out to be also relevant to the infinite-horizon steady-state. If X is irreducible and positive recurrent on state space S , the process X is regenerative with respect to consecutive hitting times

of any state $z \in S$. Assuming suitable moment hypotheses are in force, regenerative process theory (Çinlar (1975), §9.2) shows that the steady-state rate α at which cost accrues is given by the well-known ratio formula

$$\alpha = \frac{E_z[\int_{[0, T'(z))} f(X(s)) ds]}{E_z T'(z)}$$

where $T'(z) = T(\{z\})$ and $E_z(\bullet)$ is the expectation operator conditional on $X(0) = z$. Both the numerator and denominator have the form (1), so variance reduction techniques developed for (1) are applicable to regenerative steady-state simulation of X ; e.g., see Fox and Glynn (1986). Via Little's law (e.g., see Glynn and Whitt (1989), this lets us estimate expected customer sojourn time in system in steady-state efficiently.

To apply conditional Monte Carlo to (1), we condition on the embedded chain $Z = (Z_0, Z_1, \dots)$, where $Z_i \neq Z_{i-1}$ and Z_i is the state visited by X just after jump i . Let τ_i be the time between jumps i and $i+1$, so that $S(n) = \tau_0 + \dots + \tau_n$ is the instant at which X makes its $(n+1)$ -st jump. Put $S(-1) = 0$ and let $H(B) = \inf\{n \geq 1 : Z_n \in B\}$. Note that

$$\begin{aligned} I &= \sum_{n=0}^{H(B)-1} \int_{[S(n-1), S(n))} f(X(t)) G(dt) \\ &= \sum_{n=0}^{H(B)-1} f(Z_n) [G(S(n)-) - G(S(n-1)-)]. \end{aligned} \quad (2)$$

Assuming that $E \int_{[0, T(B))} |f(X(t))| G(dt) < \infty$, we may take conditional expectations of both sides of (2) with respect to Z , yielding

$$E[I|Z] = \sum_{n=0}^{H(B)-1} f(Z_n) \{E[G(S(n)-)|Z] - E[G(S(n-1)-)|Z]\}. \quad (3)$$

Let $Q = (Q(x, y) : x, y \in S)$ be the generator of X and let $\lambda(x) = -Q(x, x)$. Continuous-time Markov chains have the convenient property that, conditional on Z , the holding times $\tau_0, \tau_1, \dots$ are conditionally independent with conditional distributions $P\{\tau_i \in dt|Z\} = f_{Z_i}(t)dt$, where $f_x(\bullet)$ is an exponential density with parameter $\lambda(x)$. Hence,

$$E[G(S(n)-)|Z] = \int_0^\infty G(t)(f_{Z_0} * \dots * f_{Z_n})(t)dt \quad (4)$$

for $n \geq 0$, where $*$ denotes convolution. We don't need $G(t^-)$ in the integrand because $G(t) = G(t^-)$ for all but at most countably many t . Combining (3) and (4) yields an expression for $I_Z \triangleq E[I|Z]$, from which a conditional Monte Carlo estimator for α can be obtained. We say that the resulting estimator

is obtained by "discrete-time conversion" since the new estimator depends on X only through a discrete-time chain, so that the holding times need not be simulated.

As discussed in Section 2, the efficiency of the estimator based on I_Z is determined by $\text{var } I_Z$ and by $E\beta(I_Z)$. Depending on the problem, the variance difference $\text{var } I - \text{var } I_Z$ can take on any value between zero and infinity, as the following example illustrates.

EXAMPLE 1. Let X be a pure birth process on the non-negative integers with constant birth rate equal to λ . Suppose $X(0) = 0$. If $G(t) = t$, then $I = T'(n)$ and $E[T'(n)|Z] = n/\lambda = ET'(n)$, so $\text{var } I_Z = 0$. On the other hand, $T'(n)$ is an Erlang- n r.v. with scale parameter λ , so $\text{var } I = n/\lambda^2$. Hence, the variance reduction can be made either arbitrarily large or small, depending on how one chooses n and λ .

Since we expect that the variance reduction will be (at least) moderate in most practical examples, the estimator based on I_Z is more efficient provided that the time required to compute the conditional expectation is at most moderately more than that to compute I . Note that I_Z requires simulation only of Z ; holding times need not be generated. Unfortunately, the convolution in (4) can be expensive to compute. However, for two important choices of G , the convolution (4) is relatively cheap to calculate. For such G 's, the discussion of Section 2 shows that the discrete-time estimator I_Z is a clear winner over I .

SETTING 1 (continued). If $G(t) = t$, then (4) is just the expected value of the sum of $n+1$ independent exponential r.v.'s with parameters $\lambda(Z_0), \dots, \lambda(Z_n)$. So, we obtain $E[G(S(n)-)|Z] \sum_{k=0}^n 1/\lambda(Z_k)$ and

$$E[I|Z] = \sum_{n=0}^{H(B)-1} f(Z_n)/\lambda(Z_n).$$

This estimator was first studied in a regenerative steady-state context by Hordijk, Iglehart, and Schassberger (1976), although it was not analyzed using the principle of conditional Monte Carlo. This estimator was further studied by Fox and Glynn (1986) in a steady-state context, where the ties to regenerative simulation were cut.

SETTING 2 (continued). Suppose $G(t) = r^{-1}(1 - e^{-rt})$. Note that $\int_0^\infty e^{-rt}(f_{Z_0} * \dots * f_{Z_n})(t)dt$ is the Laplace transform (evaluated at r) of the distribution of the sum of n independent exponential r.v.'s. Hence,

$$E[G(S(n)-)|Z] = \frac{1}{r} \left(1 - \prod_{k=0}^n \left(\frac{\lambda(Z_k)}{\lambda(Z_k) + r} \right) \right)$$

and

$$E[I|Z] = \sum_{n=0}^{H(B)-1} f(Z_n) \prod_{k=0}^{n-1} \left(\frac{\lambda(Z_k)}{\lambda(Z_k) + r} \right) \frac{1}{(\lambda(Z_n) + r)}.$$

An estimator similar to the above was first described by Fox and Glynn (1989a) in an infinite-horizon setting.

For general G , the integral (4) is typically (much) more computationally expensive to calculate, diminishing the attractiveness of I_Z as an estimator of α . This point is illustrated by our next example.

SETTING 3 (continued). For $G(t) = I(t \geq T)$, we find that (4) involves calculating the tail of the distribution function (evaluated at T) of the sum of $n+1$ independent exponential r.v.'s. Since the parameters of the $n+1$ exponentials typically differ from one another, this distribution function is neither available in closed form nor easy to calculate numerically.

Because of the above computational difficulty in calculating (4), a different approach to discrete-time conversion may be preferable. Suppose that X is now a uniformizable continuous-time Markov chain and let $\lambda = \sup\{\lambda(x) : x \in S\}$. For $\theta \geq \lambda$, X can be represented as $X(t) = Y_{N_\theta(t)}^\theta$, where $Y^\theta = (Y_0^\theta, Y_1^\theta, \dots)$ is a discrete-time Markov chain having transition matrix $P(\theta) = \theta^{-1}(Q + \theta I)$ and $N_\theta(\cdot)$ is a Poisson process, independent of Y^θ , having rate θ . We now develop a conditional Monte Carlo estimator for α based on the conditioning variables $Y^\theta = (Y_0^\theta, Y_1^\theta, \dots)$. Simulating the naive estimator I using the uniformized chain would waste work.

Let $\eta_0^\theta, \eta_1^\theta, \dots$ be the iid exponential (θ) interevent times of the Poisson process N_θ and let $T_n^\theta = \eta_0^\theta + \eta_1^\theta + \dots + \eta_n^\theta$ be the instant at which N_θ makes its $(n+1)$ -st jump. Put $J_\theta(B) = \inf\{n \geq 1 : Y_{n-1}^\theta \notin B, Y_n^\theta \in B\}$ and observe that

$$I = \sum_{n=0}^{J_\theta(B)-1} f(Y_n^\theta) [G(T_n^\theta-) - G(T_{n-1}^\theta-)].$$

Since the T_n^θ 's are independent of Y^θ , it follows that

$$E[I|Y^\theta] = \sum_{n=0}^{J_\theta(B)-1} f(Y_n^\theta) [EG(T_n^\theta-) - EG(T_{n-1}^\theta-)]. \quad (5)$$

Noting that T_n^θ is an Erlang r.v. with shape parameter $n+1$ and scale parameter θ , we find that

$$EG(T_n^\theta-) = \int_0^\infty G(t) \frac{\theta^{n+1} t^n e^{-\theta t}}{n!} dt. \quad (6)$$

Combining (5) and (6) yields an expression for $I_{Y^\bullet} \triangleq E[I|Y^\bullet]$.

The principal advantage of $I_{Y^\bullet}$ over I_Z is that the computation of (6) is typically (much) cheaper than that of (4). For example, in (4), the convolution $f_{Z_0} * \dots * f_{Z_n}$ must be evaluated numerically at many points before integrating numerically against $G(t)dt$, whereas in (6) the convolution is calculated analytically. In addition, for certain (practically important) choices of G , (6) can be calculated in closed form when (4) cannot.

SETTING 3 (continued). For $G(t) = I(t \geq T)$, we have that (6) equals $P\{T_n^\theta \geq T\} = P\{N_\theta(T) \leq n\}$. Hence,

$$I_{Y^\bullet} = \sum_{n=0}^{J_\bullet(B)-1} f(Y_n^\theta) e^{-\theta T} \frac{(\theta T)^n}{n!}.$$

We now choose θ optimally. Recall that Y^θ can be obtained from Z by adding "null jumps". More precisely, Y^θ can be represented as

$$Y_n^\theta = \sum_{k=0}^{\infty} Z_k I\left(\sum_{j=0}^{k-1} \nu_j^\theta \leq n < \sum_{j=0}^k \nu_j^\theta\right) \quad (7)$$

where $\nu_0^\theta, \nu_1^\theta, \dots$ are conditionally independent, given Z , and $P\{\nu_k^\theta = \ell | Z\} = \left(1 - \frac{\lambda(Z_k)}{\theta}\right)^{\ell-1} \frac{\lambda(Z_k)}{\theta}$.

Since $J_\theta(B) = \sum_{k=0}^{H(B)-1} \nu_k^\theta$ and the ν_k^θ 's are stochastically increasing in θ , clearly $J_\theta(B)$ is stochastically increasing in θ . Thus, the work required to compute (5) is minimized by taking θ as small as possible, namely $\theta = \lambda$.

As for the variance of $I_{Y^\bullet}$, Y^θ differs from Y^λ only in that it has more null jumps. In some sense, Y^θ contains more information than Y^λ and hence Proposition 1 ought to apply, yielding the conclusion that $\text{var } I_{Y^\bullet}$ is minimized by taking $\theta = \lambda$. Our next theorem confirms this assertion.

THEOREM 2. Suppose that $\text{var } I < \infty$. Then, $\text{var } I_{Y^\lambda} \leq \text{var } I_{Y^\bullet}$ for all $\theta \geq \lambda$.

Hence, the efficiency of the estimator based on replicating $I_{Y^\bullet}$ is maximized by taking $\theta = \lambda$.

This leaves us with the question of when I_{Y^λ} is more efficient than I_Z . First, note that $J_\lambda(B) \geq H(B)$ so that the number of summands in (3) is always less than the number in (5). Hence, the estimator based on I_{Y^λ} requires less work than does I_Z only when the integral (4) is significantly more expensive to compute than (6). For the variance comparison, it is clear that the sequence Z is a (deterministic) function of Y^λ , so that $\sigma(Z) \leq \sigma(Y^\lambda)$. Proposition 1 immediately yields the next result.

THEOREM 3. Suppose that $\text{var } I < \infty$. Then, $\text{var } I_Z \leq \text{var } I_{Y^*}$.

With this result in hand, it is clear that I_Z is the discrete-time estimator of choice whenever the computation of (4) is not too much more expensive than that of (6). If this condition is not satisfied, the choice is fuzzier; we defer discussion to the next section on implementation issues.

4 Discrete-time conversion for random-time horizons: implementation

In this section, we first consider, in subsection 4.1, how to efficiently calculate (4) [a key to computing $E[I|Z]$ in formula (3)] for general G . Next, in subsection 4.2, we consider how to implement formulas (5) and (6) which were based on conditioning on Y^θ . For finite horizons, whether or not it pays to uniformize depends on the problem (via Theorem 1). This carries over to infinite horizons. Subsection 4.3 considers steady-state sojourn time S . To estimate the expectation of a linear function of S , we don't uniformize. But for a nonlinear function, we do.

Fox (1989) and Fox and Young (1989) show how to generate Z (and hence Y^θ) quickly. The techniques developed there reduce $E\beta(I_Z)$ and $E\beta(I_{Y^*})$, increasing efficiency.

4.1 Without uniformization

It can be easily verified that when $\lambda(Z_0), \dots, \lambda(Z_n)$ are distinct, the convolutions in (4) take the form

$$c_{0,n} \exp(-\lambda(Z_0)t) + \dots + c_{n,n} \exp(-\lambda(Z_n)t). \quad (8)$$

Furthermore, the coefficients $c_{0,n}, \dots, c_{n,n}$ can be efficiently calculated from $c_{0,n-1}, \dots, c_{n-1,n-1}$ (the coefficients associated with $f_{Z_0} * \dots * f_{Z_{n-1}}$) in order n operations. With the initial condition $c_{00} = \lambda(Z_0)$, we recursively solve for the coefficients $c_{0,n}, \dots, c_{n,n}$ appearing in (7) in order n^2 operations. When the $\lambda(Z_i)$'s are not distinct, some terms in (7) get multiplied by (routinely-determined) powers of t and other terms vanish. Again, in this case, the coefficients can be calculated recursively.

To numerically evaluate (4), we use the recursive algorithm described above to find the symbolic representation of the convolution and then numerically evaluate the product of the convolution with $G(t)$ on a suitably defined grid of points.

In some cases, it may pay to look for (and exploit) common subsequences of $\lambda(Z_i)$'s across runs to avoid computing all convolutions for all runs from scratch.

4.2 With uniformization

Relation (5) can be written in the form

$$E[I|Y^\theta] = \sum_{k=0}^{H(B)-1} f(Z_k)[b(\chi_k^\theta, \theta) - b(\chi_{k-1}^\theta, \theta)] \quad (9)$$

where

$$\chi_k^\theta = \sum_{j=0}^k \nu_j^\theta,$$

and

$$b(n, \theta) = \int_0^\infty G(t) \theta^{n+1} \frac{e^{-\theta t} t^n}{n!} dt.$$

We check whether $b(n, \theta)$ has already been computed (and stored) from some previous run. If it has, we use it. If it hasn't, let $\bar{n}$ be the maximum j for which $b(j, \theta)$ has been previously computed. We can then recursively calculate $b(\bar{n} + 1, \theta)$ through $b(n, \theta)$ from $b(\bar{n}, \theta)$, by using the following ideas:

i) on a grid $t_1 < \dots < t_m$, we assume that the quantities

$$\Delta_{i\bar{n}} = G(t_i) \theta^{\bar{n}+1} \frac{e^{-\theta t_i} t_i^{\bar{n}}}{\bar{n}!} (t_i - t_{i-1})$$

($1 \leq i \leq m, t_0 = 0$) are already computed and stored. Choose t_m so that the integral to its right is negligible.

ii) for $1 \leq i \leq m$, we compute $\Delta_{i\bar{n}+1}, \dots, \Delta_{in}$ recursively, using

$$\Delta_{i,j+1} = \Delta_{ij} \frac{\theta t_i}{(j+1)}$$

and the initial condition $\Delta_{i\bar{n}}$ from (i)

iii) we use $\sum_{i=1}^{m-1} \Delta_{ij}$ as a numerical approximation to $b(j, \theta)$, $\bar{n} < j \leq n$.

Similar ideas are compatible with more sophisticated numerical integration methods.

Returning to (9), we note that the ν_i^θ 's can be generated (by "inversion") as geometric variates with parameter $\lambda(Z_i)/\theta$, in computation time independent of $\lambda(Z_i)/\theta$; e.g., see Bratley, Fox, and Schrage [(1987), §5.4.5]. Sometimes null-jump sequences can start only in certain states; for example, in the M/M/1 queue only from the state corresponding to an empty system. This condition is detected automatically when $\lambda(Z_i) = \theta$.

In certain applications (particularly, those in which $\inf\{\lambda(x) : x \in S\} \ll \theta$), the ν_i^θ 's may well account for most of the variance of $E[I|Y^\theta]$. This suggests

that it may pay to generate multiple Y^θ sequences from a single Z by inserting (say k) conditionally independent copies of $\nu_0^\theta, \dots, \nu_{H(B)-1}^\theta$ into Z . Generating such a copy is usually trivial relative to the work to generate Z , because each null-jump sequence is generated in $O(1)$ time (with a small implicit constant). We then use

$$\sum_{r=0}^{H(B)-1} f(Z_r) \frac{1}{k} \sum_{i=1}^k [b(\chi_{ri}^\theta, \theta) - b(\chi_{r-1,i}^\theta, \theta)]$$

to estimate α , where $\chi_{ri}^\theta = \sum_{j=0}^r \nu_{ji}^\theta$ and ν_{ji}^θ is the i -th copy of ν_j^θ . Our (overall) estimator of α is the average of such iid estimators generated for a given computer budget. Fox and Glynn (1989b) show how to choose k to maximize efficiency as a special case of their results on splitting.

4.3 Steady-state sojourn times

Consider estimating the expectation of a function g of system sojourn time S for customers in steady state. When g is linear, we use Little's law as detailed in Section 3; in this case, it does not pay to uniformize. When g is nonlinear, however, this problem does not have a form analogous to (1), even via Little's law, unless the state space is expanded to keep track of customer entry and exit times.

Nonetheless, we can still convert to discrete time (and not expand the state space). We note for each customer the difference d between the transition numbers when he leaves the system and when he enters it. When the chain is *uniformized*, clearly $E[g(S)|d] = E[g(e(d))]$ where e is an Erlang variate with shape parameter d . Without uniformization, finding $E[g(S)|Z]$ requires recording each customer's path through the system, numerically computing the density of each customer's sojourn time (via a convolution) at many grid points, and numerically computing the corresponding expectations; this is much harder.

5 Discrete-time conversion for deterministic-time horizons: theory

Our goal in this section is to use discrete-time conversion to estimate $\alpha = E[I]$, where

$$I = \int_{[0,T]} f(X(t))G(dt). \quad (10)$$

We assume that X is a non-explosive continuous-time Markov chain on state space S , f is a real-valued function defined on S , T is deterministic, and G is a (deterministic) non-decreasing right-continuous function on $[0, \infty]$. Such estimation problems arise in a number of different settings.

SETTING 4. Suppose that we interpret $f(x)$ as the rate at which cost accrues when the process occupies state x . If $G(t) = t$, we find that I is the total cost accumulated over the horizon $[0, T]$.

SETTING 5. Suppose $G(t) = r^{-1}(1 - e^{-rt})(r > 0)$ and that f is interpreted as in Setting 4. Then, I is the r -discounted cost over the horizon $[0, T]$.

SETTING 6. Suppose that the system is charged an amount $f(x)$ if the process X occupies state x at time T . Then, the expected cost α charged to the system is given by $\alpha = E[I]$, where $G(t) = I(t \geq T)$.

Based on the discussion of Section 3, we would ideally like to obtain our discrete-time estimator for α by conditioning on Z , at least if $E[I|Z]$ is not hard to compute relative to alternative estimators. The following proposition will help us compute our conditional expectations.

PROPOSITION 2. Let $\mathcal{G}$ be a σ -field and suppose that

$$E \int_{[0, T]} |f(X(t))| G(dt) < \infty.$$

Then,

$$E[I|\mathcal{G}] = \int_{[0, T]} E[f(X(t))|\mathcal{G}] G(dt).$$

To apply this proposition to the calculation of $E[I|Z]$, observe that

$$f(X(t)) = \sum_{n=0}^{\infty} f(Z_n) I\{S_{n-1} \leq t < S_n\}$$

so that

$$E[f(X(t))|Z] = \sum_{n=0}^{\infty} f(Z_n) \int_0^t \int_{t-u}^{\infty} f_{Z_n}(v) dv (f_{Z_0} * \cdots * f_{Z_{n-1}})(u) du.$$

Hence, by Proposition 2,

$$E[I|Z] = \sum_{n=0}^{\infty} f(Z_n) \int_{[0, T]} \int_0^t \exp(-\lambda(Z_n)(t-u)) (f_{Z_0} * \cdots * f_{Z_{n-1}})(u) du G(dt). \quad (11)$$

Several difficulties with $I_Z \triangleq E[I|Z]$ are apparent. First, I_Z depends on the entire history of Z and therefore requires an infinite amount of time to compute (exactly). Second, the summands that define Z involve difficult double integrals that do not simplify significantly, even for well-chosen G .

SETTING 4 (continued). Here $G(t) = t$ so that the double integral in (11) becomes

$$\begin{aligned}
& \int_0^T \int_0^t \int_{t-u}^\infty f_{Z_n}(v) dv (f_{Z_0} * \cdots * f_{Z_{n-1}})(u) du dt \\
&= \int_0^T \int_0^T \int_0^\infty I(u \leq t \leq u+v) f_{Z_n}(v) dv (f_{Z_0} * \cdots * f_{Z_{n-1}})(u) du dt \\
&= \int_0^T \int_0^\infty \int_0^T I(u \leq t \leq u+v) dt f_{Z_n}(v) dv (f_{Z_0} * \cdots * f_{Z_{n-1}})(u) du \\
&= \int_0^T \int_0^{T-u} v f_{Z_n}(v) dv (f_{Z_0} * \cdots * f_{Z_{n-1}})(u) du \\
&\quad + \int_0^T \int_{T-u}^\infty (T-u) f_{Z_n}(v) dv (f_{Z_0} * \cdots * f_{Z_{n-1}})(u) du \\
&= \int_0^T \int_0^\infty v f_{Z_n}(v) dv (f_{Z_0} * \cdots * f_{Z_{n-1}})(u) du \\
&\quad - \int_0^T \int_{T-u}^\infty (v-T+u) f_{Z_n}(v) dv (f_{Z_0} * \cdots * f_{Z_{n-1}})(u) du \\
&= \lambda(Z_n)^{-1} \int_0^T (f_{Z_0} * \cdots * f_{Z_{n-1}})(u) du \\
&\quad - \lambda(Z_n)^{-1} \int_0^T \exp(-\lambda(Z_n)(T-u)) (f_{Z_0} * \cdots * f_{Z_{n-1}})(u) du.
\end{aligned} \tag{12}$$

Hence, unless the convolutions can be rapidly calculated (perhaps as in Section 4.1), I_Z appears impractical.

Similar difficulties arise when we wish to calculate I_Z for the estimators that arise in Settings 5 and 6. Therefore, we now consider alternatives. We now assume that X is uniformizable with $\lambda = \sup\{\lambda(x) : x \in S\} < \infty$. The analog of (5) is given by

$$E[I|Y^\theta] = f(Y_0^\theta) \cdot \int_{[0,T]} e^{-\theta t} G(dt) + \frac{1}{\theta} \sum_{n=1}^{\infty} f(Y_n^\theta) \int_{[0,T]} \frac{\theta^{n+1} t^n e^{-\theta t}}{n!} G(dt). \tag{13}$$

The inner (convolution) integral that appears in (12) is analytically calculated here: it becomes Erlang. Also, for certain G , the conditional expectation $I_{Y^\theta} = E[I|Y^\theta]$ simplifies further.

SETTING 4 (continued). For $G(t) = t$, we get

$$\int_0^T \frac{\theta^{n+1} t^n e^{-\theta t}}{n!} dt = \sum_{k=n+1}^{\infty} e^{-\theta T} \frac{(\theta T)^k}{k!}$$

and so

$$\begin{aligned}
E[I|Y^\theta] &= \frac{1}{\theta} \sum_{n=0}^{\infty} f(Y_n^\theta) \sum_{k=n+1}^{\infty} e^{-\theta T} \frac{(\theta T)^k}{k!} \\
&= T \sum_{k=0}^{\infty} \sum_{n=0}^k f(Y_n^\theta) e^{-\theta T} \frac{(\theta T)^k}{(k+1)!}
\end{aligned} \tag{14}$$

Gross and Miller (1984) find (14) by a different route.

SETTING 5 (continued). For $G(t) = r^{-1}(1 - e^{-rt})$, we find that

$$\begin{aligned}
\int_0^T \frac{\theta^{n+1} t^n e^{-\theta t}}{n!} e^{-rt} dt &= \left(\frac{\theta}{\theta+r} \right)^{n+1} \int_0^T \frac{(\theta+r)^{n+1} t^n e^{-(\theta+r)t}}{n!} dt \\
&= \left(\frac{\theta}{\theta+r} \right)^{n+1} \sum_{k=n+1}^{\infty} e^{-(\theta+r)T} \frac{((\theta+r)T)^k}{k!}
\end{aligned}$$

so

$$E[I|Y^\theta] = \sum_{n=0}^{\infty} f(Y_n^\theta) \frac{\theta^n}{(\theta+r)^{n+1}} \sum_{k=n+1}^{\infty} e^{-(\theta+r)T} \frac{((\theta+r)T)^k}{k!}$$

SETTING 6 (continued). For $G(t) = I(t \geq T)$, we obtain

$$E[I|Y^\theta] = \sum_{n=0}^{\infty} f(Y_n^\theta) e^{-\theta T} \frac{(\theta T)^n}{n!}.$$

Just as in Section 3, the choice $\theta = \lambda$ minimizes both the average work $E[\beta(I_{Y^\bullet})]$ and the variance $\text{var } I_{Y^\bullet}$. The arguments are identical to those of Section 3.

A major difficulty with the estimators is that, as in the case of I_Z , they depend on the entire history of Y^θ . One (obvious) solution terminates the simulation of Y^θ after a certain fixed number of transitions (say m) and deletes those terms from the estimators which depend on Y_n^θ for $n \geq m$. For several (practically-important) G 's, the resulting (truncation) error can be bounded if f is bounded using bounds on Poisson tails in Fox and Glynn (1988). If f is not bounded, it seems hard or impossible to get useful error bounds. In any case, Fox and Glynn (1989c) show that such termination strategies are also undesirable in terms of their (theoretical) large-sample convergence rate. It is shown there that even if the termination index is chosen to depend optimally on the budget t , the resulting estimator will converge (slightly) more slowly than $t^{-1/2}$ in the budget t . This contrasts with the estimators discussed in Section 2, where Theorem 1 establishes a (canonical) convergence rate of $t^{-1/2}$ for the

estimator $r(t)$. Consequently, we now discuss variants of the above estimators that avoid this difficulty.

Suppose that we enrich the conditioning σ -field so that it also contains information on the number of events of N_θ completed by time $V \geq T$. While this will clearly increase the variance of the estimator (see Proposition 1), we would expect that the resulting estimator will depend on the history of Y^θ only up through the $N_\theta(V)$ -th transition. This decreases the work required to calculate the estimator. In particular, the time required to (exactly) calculate the estimator will now be finite, whereas the time to calculate I_{Y^θ} is infinite.

Let $\mathcal{G}_{\theta,V}$ be the σ -field generated by Y^θ and $N_\theta(V)$. Then, for $t \leq V$,

$$\begin{aligned} E[f(X(t))|Y^\theta, N_\theta(V) = m] &= \sum_{j=0}^m E[f(X(t))I(N_\theta(t) = j)|Y^\theta, N_\theta(V) = m] \\ &= \sum_{j=0}^m f(Y_j^\theta)P\{N_\theta(t) = j|Y^\theta, N_\theta(V) = m\} \\ &= \sum_{j=0}^m f(Y_j^\theta)P\{N_\theta(t) = j|N_\theta(V) = m\}. \end{aligned}$$

We used the independence of Y^θ and N_θ in the last step. But

$$P\{N_\theta(t) = j|N_\theta(V) = m\} = \binom{m}{j} \left(\frac{t}{V}\right)^j \left(\frac{V-t}{V}\right)^{m-j}.$$

Hence, in light of Proposition 2, we obtain

$$E[I|\mathcal{G}_{\theta,V}] = \sum_{j=0}^{N_\theta(V)} f(X_j^\theta) \int_{[0,T]} \binom{N_\theta(V)}{j} \left(\frac{t}{V}\right)^j \left(\frac{V-t}{V}\right)^{N_\theta(V)-j} G(dt). \quad (15)$$

The estimator $I_{\mathcal{G}_{\theta,V}} \triangleq E[I|\mathcal{G}_{\theta,V}]$ can be calculated in finite time. The question naturally arises as to which choice of θ and V minimizes the variance of $I_{\mathcal{G}_{\theta,V}}$. As in Section 3, the choice $\theta = \lambda$ minimizes the variance for any (fixed) value of V . The next proposition states that the variance of $I_{\mathcal{G}_{\theta,V}}$ is non-increasing in V .

PROPOSITION 3. If $\text{var } I < \infty$, then $\text{var } E[I|\mathcal{G}_{\theta,V}]$ is non-increasing in $V \geq T$.

This result seems to suggest that we should choose V as large as possible. This, however, has the effect of increasing the expected time required to compute $I_{\mathcal{G}_{\theta,V}}$. Assuming (reasonably) that the average work increases proportionately to V , this suggests choosing V to minimize $Vc(V)$ over $V \geq T$, where $c(V) =$

$\text{var } I_{\mathcal{G}_{\theta,V}}$. Noting that the σ -field generated by Y^θ is contained in $\mathcal{G}_{\theta,V}$, we may conclude that $c(V) \geq \text{var } I_{Y^\theta}$ for all V . Hence, $Vc(V) \rightarrow \infty$ as $V \rightarrow \infty$, so that we are unlikely to find a minimizing value of V (much) larger than T . Since the minimizing V is (probably) close to T , we therefore suggest setting $V = T$ to avoid the (expensive) trial runs necessary to find the optimal choice of V . Henceforth, we take $V = T$ and $\theta = \lambda$.

Given this choice of V , (15) simplifies further when G is appropriately chosen.

SETTING 4 (continued). For $G(t) = t$, consider

$$\int_0^T \binom{m}{j} \left(\frac{t}{T}\right)^j \left(\frac{T-t}{T}\right)^{m-j} dt.$$

The integrand is a beta density (up to a proportionality constant) over $[0, T]$, from which we conclude that the integral is $T/(m+1)$. So,

$$E[I|\mathcal{G}_{\lambda,T}] = T \sum_{j=0}^{N_\lambda(T)} f(Y_j^\lambda)/(N_\lambda(T) + 1).$$

SETTING 5 (continued). If $G(t) = r^{-1}(1 - e^{rt})$, then Section 6.4 shows how to compute the set of required integrals with no numerical integration in $O(L^2)$ time, where L is the largest Poisson variate generated over all runs.

SETTING 6 (continued). Since $X(T) = Y^\lambda(N_\lambda(T))$, it follows that when $G(t) = I(t \geq T)$, I is $\mathcal{G}_{\lambda,T}$ measurable so that $E[I|\mathcal{G}_{\lambda,T}] = I$. Hence, conditioning on $\mathcal{G}_{\lambda,T}$ just leads us back to the naive estimator I .

We conclude this section with the introduction of an estimator that replaces holding time r.v.'s with their means, yet can not be represented as a conditional expectation of the form $E[I|\mathcal{G}]$. Nevertheless, we classify the following estimator as a discrete-time estimator, because continuous holding time r.v.'s are integrated out.

Our discussion is specific to $G(t) = t$. We return to formula (12) and note that it equals

$$\lambda(Z_n)^{-1}(P\{S_{n-1} \leq T|Z\} - P\{S_{n-1} \leq T, S_n > T|Z\}) = \lambda(Z_n)^{-1}P\{S_n \leq T|Z\}.$$

By (11), we get

$$E[I|Z] = \sum_{n=0}^{\infty} \frac{f(Z_n)}{\lambda(Z_n)} P\{S_n \leq T|Z\}.$$

Taking expectations of both sides of the above equation gives

$$E[I] = E \sum_{n=0}^{\tilde{N}(T)-1} \frac{f(Z_n)}{\lambda(Z_n)},$$

where $\tilde{N}(T) = \min\{n \geq 0 : S_n > T\}$. This suggests considering an estimator for the expected total reward over $[0, T]$, based on replicates of the r.v. I_d , where

$$I_d = \sum_{n=0}^{\tilde{N}(T)-1} \frac{f(Z_n)}{\lambda(Z_n)}. \quad (16)$$

We claim that I_d typically can't be represented as $E[I|\mathcal{G}]$ for some σ -field $\mathcal{G}$. This necessarily implies that I_d is distinct from $E[I|\mathcal{G}_{\lambda,T}]$. We prove this claim by providing an example in which $\text{var } I_d > \text{var } I$. If I_d were a conditional expectation of I , this would violate Proposition 1. To obtain the example, just take $f \equiv 1$ so that $I = T$. Then, $\text{var } I = 0$ but $\text{var } I_d > 0$.

We include the estimator I_d in our paper because it turns out that I_d is sometimes more efficient than $I_{\mathcal{G}_{\lambda,T}}$, though it is not a conditional expectation. To calculate $I_{\mathcal{G}_{\lambda,T}}$ requires generating $Y_0^\lambda, Y_1^\lambda, \dots, Y_{N_\lambda(T)}^\lambda, N_\lambda(T)$, whereas I_d is a (deterministic) function of $Z_0, Z_1, \dots, Z_{\tilde{N}(T)}, \tilde{N}(T)$. Since Y^λ includes null jumps not present in Z , $N_\lambda(T) \geq \tilde{N}(T)$. Hence, $E[\beta(I_d)] \geq E[\beta(I_{\mathcal{G}_{\lambda,T}})]$, so that I_d always beats $I_{\mathcal{G}_{\lambda,T}}$ in terms of work.

Furthermore, I_d sometimes beats $I_{\mathcal{G}_{\lambda,T}}$ in terms of variance, as our next example illustrates. Suppose $S = \{0, 1\}$ and $\lambda(0) = 1, \lambda(1) = 0$, so that state 1 is absorbing. Clearly, $f(0) = 1, f(1) = 0$, and note that $I_d = I(\tilde{N}(T) > 0), I_{\mathcal{G}_{\lambda,T}} = T/(N_1(T) + 1)$. Then, $\text{var } I_d = e^{-t} - e^{-2t}$. Note that as $T \rightarrow \infty$,

$$T^{1/2} \left(\frac{T}{N_1(T) + 1} - 1 \right) = \left(\frac{T}{N_1(T) + 1} \right) \left(\frac{T - N(T)}{\sqrt{T}} \right) \Rightarrow N(0, 1).$$

Establishing uniform integrability is easy, from which we conclude that $E(T/(N_1(T) + 1) - 1)^2 \sim 1/T$ as $T \rightarrow \infty$. Since $E[I_{\mathcal{G}_{\lambda,T}}] = 1 - e^{-t}$, it follows that $\text{var } I_{\mathcal{G}_{\lambda,T}} \sim 1/T$ as $T \rightarrow \infty$. Hence, $\text{var } I_d \ll \text{var } I_{\mathcal{G}_{\lambda,T}}$ for large T .

6 Discrete-time conversion for fixed-time horizons: implementation

Whereas for random-time horizons the estimator $E[I|Z]$ is sometimes competitive, for fixed-time horizons $E[I|Z]$ seems very unlikely to be competitive with the improvements of $E[I|\mathcal{G}_{\lambda,T}]$ presented here. For the case $G(t) = t$, the estimator I_d of Section 5 might be competitive with these improvements; that estimator is straightforward to implement.

Our first improvement of $E[I|\mathcal{G}_{\theta,T}]$ stratifies $N_{\theta}(T)$. Let $\delta = \theta T$. Section 5 indicates that picking $\theta = \lambda$ is good, but here we allow any $\theta \geq \lambda$ for flexibility. Stratum one is the integers $0, \dots, K_{\delta}$, and stratum two is the remaining positive integers. We choose K_{δ} so that $P\{N_{\theta}(T) \leq K_{\delta}\}$ is near one. Next, we integrate $N_{\theta}(T)$ out of stratum one. Although the second stratum has an estimator conditioned on $N_{\theta}(T)$ as well as Y^{θ} , the variance of the overall estimator relative to that of $E[I|\mathcal{G}_{\theta,T}]$ is small. Our overall estimator $\tilde{\alpha}$ is unbiased. Unlike $E[I|Y^{\theta}]$, it can be computed exactly with a finite amount of work. Let $\hat{\alpha}_1(t)$ average iid copies of $\tilde{\alpha}$ generated with computer budget t . Clearly, $\hat{\alpha}_1(t)$ converges at the canonical rate $t^{-1/2}$ to α . Let $\hat{\alpha}_2(t)$ average iid copies of $E[I|\sigma(Y^{\theta}, N_{\theta}(T))]$ and assume (reasonably) that the expected work to generate $\tilde{\alpha}$ is at most the expected work to generate $E[I|\sigma(Y^{\theta}, N_{\theta}(T))]$. Neglecting rounding in stratification, $\text{var } \hat{\alpha}_1(t) \leq \text{var } \hat{\alpha}_2(t)$ as holds whenever we stratify proportionally; e.g., see Bratley, Fox, and Schrage [(1987), §2.4].

In Section 6.3, we increase the efficiency of the first-stratum estimator above via *splitting*. Though the expected work per run increases, the product of expected work per run and output variance per run decreases – just what we want according to Theorem 1. Call $\hat{\alpha}_3(t)$ the resulting (unbiased) overall estimator of α . We get $\text{var } \hat{\alpha}_3(t) \leq \text{var } \hat{\alpha}_1(t)$. The variance decrease can be dramatic. We already introduced a version of splitting in Section 4.2. Here is another version. First, we partition stratum one into $\Phi = \{0, 1, \dots, K_{\varphi}\}$ and $\Delta = \{K_{\varphi} + 1, \dots, K_{\delta}\}$ where $|\Phi| \gg |\Delta|$ but $P\{N_{\theta}(T) \in \Delta\}$ is (still) near one. If δ is large, we pick K_{φ} a few standard deviations to the left of the mean (i.e., $K_{\varphi} = \delta - c_1 \delta^{1/2}$ where $c_1 = 4$ say) and K_{δ} a few standard deviations right of the mean (i.e., $K_{\delta} = \delta + c_2 \delta^{1/2}$ where $c_2 = 4$ say); thus, $|\Phi| = O(\delta)$ but $|\Delta| = O(\delta^{1/2})$. When the simulation reaches $Y_{K_{\varphi}}$, we split it into (say) k subruns all starting at $Y_{K_{\varphi}}$ and going on to K_{δ} . We can nest the version of splitting described in Section 4.2 inside this procedure.

Section 6.4 shows how to compute the integrals in (15) for general G . When $G(t) = t$ or $G(t) = I\{T \geq t\}$, then these integrals are available in simple closed form as Section 5 shows.

If the chain's generator is piecewise constant, then an estimator of the form $E[I|\mathcal{G}_{\theta,T}]$ applies to each piece, with the final state for piece i becoming the initial state for piece $i + 1$. Since except for the last piece we need to know the final state, integrating $N_{\theta}(T)$ out of stratum one works only for the last piece. In each piece we simulate a stationary process, whereas to get I one must simulate a nonstationary process. This makes I a less efficient estimator relative to $E[I|\mathcal{G}_{\theta,T}]$ than in the stationary case. Poisson arrivals with a piecewise-constant intensity to a system which is otherwise a stationary continuous-time Markov chain produce a piecewise-constant generator.

6.1 Telescoping null-jump sequences

Analogously to what was done in Section 4.2, we reexpress Y^θ as $\{Z_0, \nu_0^\theta - 1, Z_1, \nu_1^\theta - 1, \dots\}$ and rewrite $E[I|Y^\theta]$ and $E[I|\sigma(Y^\theta, N_\theta(T))]$ accordingly.

$$E[I|Y^\theta] = f(Z_0) \int_{[0,T]} e^{-\theta t} G(dt) + \theta^{-1} f(Z_0) [\nu_0^\theta - 1] c(0, \theta) + \theta^{-1} \sum_{n=1}^{\infty} f(Z_n) \nu_n^\theta c(n, \theta) \quad (17)$$

where

$$c(n, \theta) = \int_{[0,T]} \frac{\theta^{n+1} t^n e^{-\theta t}}{n!} G(dt). \quad (18)$$

$$\begin{aligned} E[I|\sigma(Y^\theta, N_\theta(T))] \\ = \sum_{j=0}^{\infty} f(Z_j) \nu_j^\theta d(j, N_\theta(T)) I\{\chi_{j+1}^\theta \leq N_\theta(T)\} \\ + f(Z_m) [N_\theta(T) - \chi_m^\theta] d(m, N_\theta(T)) \bullet I\{\chi_m^\theta \leq N_\theta(T) < \chi_{m+1}^\theta\} \end{aligned} \quad (19)$$

where

$$d(j, \ell) = \int_{[0,T]} \binom{\ell}{j} \left(\frac{t}{T}\right)^j \left(\frac{T-t}{T}\right)^{\ell-j} G(dt) \quad (20)$$

and $\chi_j^\theta = \sum_{k=0}^j \nu_k^\theta$. Section 4.2 discusses computation of integrals almost the same as $c(n, \theta)$, after an integration by parts. Section 6.4 shows how to compute the $d(j, \ell)$'s recursively for general G . Section 5 shows, for example, that $d(j, \ell) = T/(\ell+1)$ when $G(t) = t$.

6.2 Stratification

Integrating the Poisson variate out of the first stratum gives

$$\begin{aligned} \sum_{i=0}^{K_\delta} E[I|Y, N_\theta(T) = i] P\{N_\theta(T) = i | N_\theta(T) \leq K_\delta\} \\ = q^{-1} \sum_{i=0}^{K_\delta} \sum_{j=0}^i f(Y_j^\theta) d(j, i) e^{-\delta} \delta^i / i! \end{aligned} \quad (21)$$

$$= q^{-1} \sum_{j=0}^{K_\delta} f(Y_j^\theta) e_j \quad (22)$$

where

$$q = P\{N_\theta(T) \leq K_\delta\} \quad (23)$$

$$e_j = \sum_{\ell=j}^{K_\delta} d(j, \ell) e^{-\delta} \delta^\ell / \ell! \quad (24)$$

$$= e_{j+1} + d(j, j) [e^{-\delta} \delta^j / j!], j < K_\delta. \quad (25)$$

We compute the e_j 's recursively starting from $j = K_\delta$. The (bracketed) "Poisson" terms are computed recursively starting from $j = \lfloor \delta \rfloor$ as in Fox and Glynn (1988). To simulate (22), use (19) with $d(j, N_\theta(T))$ replaced by e_j and $N_\theta(T)$ replaced by K_δ . This handles stratum one.

To handle stratum two, we force $N_\theta(T)$ to fall there and use (19) multiplied by $1/(1-q)$. Devroye (1986) gives an $O(1)$ average-time rejection algorithm to generate variates from Poisson right tails that, with (hypothetical) infinite-precision computers, requires no right-hand truncation of the tail; its worst-case time is unbounded. Now we suggest an alternative which is sometimes more attractive. For several important G 's and bounded f , the (close) upper bounds on Poisson-tail masses in Fox and Glynn (1988) let us find a truncation point M_δ such that the mass to its right is negligible. This suggests forcing $N_\theta(T)$ to fall between $K_\delta + 1$ and M_δ and using "inversion" to generate $N_\theta(T)$ thus conditioned. Most of the Poisson probabilities required for "inversion" are already computed in (25), since $K_\delta/M_\delta \approx 1$. Since most of the mass between $K_\delta + 1$ and M_δ is concentrated at $K_\delta + 1$, straightforward "inversion" appears competitive even for average time. Correspondence on this point with Bruce Schmeiser was helpful. Alternatively, "inversion" can be implemented with the alias method; e.g., see Bratley, Fox, and Schrage [(1987), §5.2.8]. This takes $O(1)$ marginal time per variate after a one-time $O(M_\delta - K_\delta)$ setup. Clearly $M_\delta - K_\delta = O(\delta^{1/2})$.

With S runs altogether, proportional stratification uses (22) on $\lfloor qS \rfloor$ runs and stratum two on $\lfloor (1-q)S \rfloor$ runs. On the remaining runs, it selects stratum one with probability proportional to $qs - \lfloor qS \rfloor$ and stratum two with probability proportional to $(1-q)S - \lfloor (1-q)S \rfloor$. If an exact algorithm to generate from Poisson tails is used, no bias results. Ignoring rounding, we get lower variance than by averaging iid copies of $E[I|\sigma(Y^\theta, N_\theta(T))]$ – even without the additional variance reduction obtained by integrating $N_\theta(T)$ out of stratum one. If the output variances and expected unit sampling costs for each stratum were known, we could get even higher efficiency by modifying the strata sampling allocations accordingly; e.g., see Bratley, Fox, and Schrage [(1987), §2.4]. One could estimate these quantities after every run and adaptive allocate subsequent runs to strata accordingly; this is a topic for future research.

6.3 Splitting

We separate the sum (22) into two terms. In the first term (say V_1) the summation goes from 0 to K_φ . Call the other term V_2 . The weights e_j in V_1 are

collectively small for many important G such as $G(t) = t$. When this occurs, the main function of V_1 is to supply an initial state Y_{K_v+1} for V_2 . It may be that V_2 is insensitive to Y_{K_v+1} . More importantly, it takes $O(\delta)$ work to get Y_{K_v+1} but only $O(\delta^{1/2})$ to compute V_2 thereafter.

We now *split*. For each V_1 we take (say) b copies of V_2 , conditionally independent given Y_{K_v+1} . Call these replicates $V_{21}, V_{22}, \dots, V_{2b}$. Thus (22) and

$$\tilde{V}_b = V_1 + b^{-1} \sum_{j=1}^b V_{2j}$$

have the same expectation. Our overall estimator of that expectation averages of iid copies of $\tilde{V}_b$. Fox and Glynn (1989b) show, in a more general setting, how to pick b to maximize efficiency in the sense of Theorem 1.

Because stratum one gets almost all the weight, it is probably not worth extending splitting to stratum two.

6.4 Computing the $d(j, \ell)$'s

Let

$$I(j, \ell) = \int_{[0, T]} (t/T)^j (1 - t/T)^\ell G(dt).$$

Since

$$I(j, \ell + 1) = I(j, \ell) - I(j + 1, \ell),$$

we get

$$d(j, \ell + 1) = [(\ell + 1)/(\ell + 1 - j)]d(j, \ell) - [(j + 1)/(\ell + 1 - k)]d(j + 1, \ell + 1) \quad (26)$$

where

$$d(\ell, \ell) = I(\ell, 0)$$

$$I(0, 0) = G(T-) - G(0-).$$

Thus, given $d(\ell, \ell)$ for $\ell = 0, 1, \dots, L$, we compute $d(j, \ell)$ for $j \neq \ell$ and $0 \leq j, \ell \leq L$ from (26) recursively with $O(L^2)$ work and memory. So at most $L + 1$ numerical integrations are needed. Since the integrand for $I(\ell, \ell)$ is just $(t/T)(1 - t/T)$ times the integrand for $I(\ell - 1, \ell - 1)$, ideas similar to those in Section 4.2 expedite computation of $\{I(\ell, \ell) : \ell = 0, \dots, L\}$.

When $G(dt) = e^{-rt}dt$, then

$$d(0, 0) = (1 - e^{-rT})/r$$

and integrating by parts gives $d(\ell, \ell) = (\ell/rT)d(\ell - 1, \ell - 1) - e^{-rT}/r$. So no numerical integration is needed for continuous discounting.

If the accurate but approximate "inversion" method for generating variates from Poisson tails suggested in Section 6.2 is used, pick $L = M_\delta$. If an exact method is used, generate all the Poisson variates needed over all runs at the start; the largest of these is L .

APPENDIX

Proof of Theorem 1. If $\sigma^2(R_i) = 0$, the result is immediate since $r(t) = \alpha$ a.s. If $\sigma^2(R_i) > 0$, then the classical central limit theorem implies that

$$n^{1/2} \left(\frac{1}{n} \sum_{k=1}^n R_k - \alpha \right) \Rightarrow \sigma(R_1)N(0, 1)$$

as $n \rightarrow \infty$. Classical renewal theory shows that $N(t)/t \rightarrow 1/E[\beta(R_1)]$ a.s. as $t \rightarrow \infty$. Apply a random time change theorem [Billingsley (1968), p. 146] to get

$$N(t)^{1/2}(r(t) - \alpha) \Rightarrow \sigma(R_1)N(0, 1)$$

as $t \rightarrow \infty$. A converging-together argument [e.g., see Billingsley (1968), p. 25] yields the theorem.

Proof of Theorem 2. For $\theta \geq \lambda$, we shall find a σ -field $\mathcal{K}_\theta$ such that $\sigma(Y^\lambda) \subseteq \mathcal{K}_\theta$ and

$$E[I|\mathcal{K}_\theta] = E[I|Y^\theta].$$

Then, Proposition 1 implies that $\text{var}(E[I|\mathcal{K}_\theta]) \geq \text{var} I_{Y^\lambda}$ and so, once we construct a suitable $\mathcal{K}_\theta$, we have proved that setting $\theta = \lambda$ minimizes variance.

That construction begins with an iid sequence $(\eta_i : i \geq 1)$ of Bernoulli variates, independent of Y^θ and N_θ , each with $P\{\eta_i = 1\} = \lambda/\theta$. We use it to thin null jumps of Y^θ to obtain Y^λ . The i -th null jump is retained if and only if $\eta_i = 1$.

Now, we set $\mathcal{K}_\theta = \sigma(Y^\theta, \eta_1, \eta_2, \dots)$. Since I is a (measurable) function of Y^θ and N_θ , we get

$$E[I|Y^\theta, \eta_1, \eta_2, \dots] = E[I|Y^\theta]$$

using (3) on p. 308 of Chung (1974), with $\mathcal{F}_1 = \sigma(Y^\theta, N_\theta)$, $\mathcal{F}_2 = \sigma(Y^\theta)$, $\mathcal{F}_3 = \sigma(\eta_1, \eta_2, \dots)$, completing the proof.

Proof of Proposition 2. The right-hand side, namely $\int_{[0, T]} E[f(X(t))|\mathcal{G}]G(dt)$, is a $\mathcal{G}$ -measurable r.v. Also, if $A \in \mathcal{G}$, then Fubini implies that

$$\begin{aligned} \int_A \int_{[0, T]} E[f(X(t))|\mathcal{G}]G(dt)P(dw) &= \int_{[0, T]} \int_A E[f(X(t))|\mathcal{G}]P(dw)G(dt) \\ &= \int_{[0, T]} \int_A f(X(t))P(dw)G(dt) \\ &= \int_A \int_{[0, T]} f(X(t))G(dt)P(dw). \end{aligned}$$

Using the defining relation for conditional expectations, we see that the proposition is proved.

Proof of Proposition 3. Suppose that $t_2 > t_1 \geq T$. Note that $N_\theta(t_2) - N_\theta(t_1)$ is independent of $\mathcal{G}_{\theta,t_1}$ and $\sigma(N_\theta(s) : s \leq t_1)$. Hence, we get

$$E[I|\mathcal{G}_{\theta,t_1}] = E[I|\mathcal{G}_{\theta,t_1}, N_\theta(t_2) - N_\theta(t_1)]$$

from (3) on p. 308 of Chung (1974), with $\mathcal{F}_1 = \sigma(N_\theta(s) : s \leq t_1) \vee \mathcal{G}_{\theta,t_1}$, $\mathcal{F}_2 = \mathcal{G}_{\theta,t_1}$, $\mathcal{F}_3 = \sigma(N_\theta(t_2) - N_\theta(t_1))$, noting that I is $\mathcal{F}_1$ measurable. But $\mathcal{G}_{\theta,t_2} \leq \mathcal{G}_{\theta,t_1} \vee \sigma(N_\theta(t_2) - N_\theta(t_1))$, so $\text{var } E[I|\mathcal{G}_{\theta,t_2}] \leq \text{var } E[I|\mathcal{G}_{\theta,t_1}, N_\theta(t_2) - N_\theta(t_1)] = \text{var } E[I|\mathcal{G}_{\theta,t_1}]$, by Proposition 1.

References

- [1] P. BILLINGSLEY, *Convergence of Probability Measures*, Wiley, New York, 1968.
- [2] P. BRATLEY, B.L. FOX, and L.E. SCHRAGE, *A Guide to Simulation*, Springer-Verlag, New York, second edition 1987.
- [3] K.L. CHUNG, *A Course in Probability Theory*, Academic Press, New York, 1974.
- [4] E. ÇINLAR, *Introduction to Stochastic Processes*, Prentice-Hall, Englewood Cliffs, New Jersey, 1975.
- [5] L. DEVROYE, *A simple tail-of-Poisson random variate generator*, unpublished manuscript, 1986.
- [6] B.L. FOX, *Generating Markov-chain transitions quickly*, Technical report, University of Colorado at Denver, 1989.
- [7] B.L. FOX and A.R. YOUNG, *Generating Markov-chain transitions quickly: II*, Technical report, University of Colorado at Denver, 1989.
- [8] B.L. FOX and P.W. GLYNN, *Discrete-time conversion for simulating semi-Markov processes*, Operations Research Letters, 5(1986), 191-196.
- [9] B.L. FOX and P.W. GLYNN, *Computing Poisson probabilities*, Comm. ACM, 31 (1988), 440-445.
- [10] B.L. FOX and P.W. GLYNN, *Simulating discounted costs*, Management Science (1989a), to appear.
- [11] B.L. FOX and P.W. GLYNN, *Splitting as conditional Monte Carlo*, Technical report in preparation, (1989b).
- [12] B.L. FOX and P.W. GLYNN, *Replication schemes for limiting expectations*, Probability in the Engineering and Informational Sciences, 3 (1989c), 299-318 (to appear).

- [13] P.W. GLYNN and W. WHITT, *Indirect estimation via $L = \lambda W$* . Operations Res., 37 (1989), 82-103.
- [14] D. GROSS and D.R. MILLER, *The randomization technique as a modeling tool and solution procedure for transient Markov processes*, Operations Res., 32 (1984), 343-361.
- [15] J.M. HAMMERSLEY and D.C. HANDSCOMB, *Monte Carlo Methods*, Methuen, London, 1984.
- [16] A. HORDIJK, D.L. IGLEHART, and R. SCHASSBERGER, *Discrete-time methods of simulating continuous-time Markov chains*, Adv. Appl. Prob., 8 (1976), 772-788.

SECURITY CLASSIFICATION OF THIS PAGE

REPORT DOCUMENTATION PAGE

1a REPORT SECURITY CLASSIFICATION Unclassified		1b RESTRICTIVE MARKINGS	
2a SECURITY CLASSIFICATION AUTHORITY		3 DISTRIBUTION/AVAILABILITY OF REPORT Approved for public release; distribution unlimited.	
2b DECLASSIFICATION/DOWNGRADING SCHEDULE		5 MONITORING ORGANIZATION REPORT NUMBER(S) APO 25839-7-MA	
4 PERFORMING ORGANIZATION REPORT NUMBER(S) Technical Report NO. 34		7a NAME OF MONITORING ORGANIZATION U. S. Army Research Office	
6a NAME OF PERFORMING ORGANIZATION Dept. of Operations Research	6b OFFICE SYMBOL (if applicable)	7b ADDRESS (City, State, and ZIP Code) P. O. Box 12211 Research Triangle Park, NC 27709-2211	
8a NAME OF FUNDING/SPONSORING ORGANIZATION U. S. Army Research Office	8b OFFICE SYMBOL (if applicable)	9 PROCUREMENT INSTRUMENT IDENTIFICATION NUMBER DAA03-88-K-0063	
8c ADDRESS (City, State, and ZIP Code) P. O. Box 12211 Research Triangle Park, NC 27709-2211		10 SOURCE OF FUNDING NUMBERS PROGRAM ELEMENT NO PROJECT NO TASK NO WORK UNIT ACCESSION NO	
11 TITLE (Include Security Classification) Discrete-Time Conversion for Simulating Finite-Horizon Markov Processes			
12 PERSONAL AUTHOR(S) Bennett L. Fox and Peter W. Glynn			
13a TYPE OF REPORT Technical	13b TIME COVERED FROM TO	14 DATE OF REPORT (Year, Month, Day) May 1989	15 PAGE COUNT 24
16 SUPPLEMENTARY NOTATION The view, opinions and/or findings contained in this report are those of the author(s) and should not be construed as an official Department of the Army position, policy, or decision, unless so designated by other documentation.			
17 COSATI CODES FIELD GROUP SUB-GROUP		18. SUBJECT TERMS (Continue on reverse if necessary and identify by block number) simulation; finite-horizon, continuous-time Markov chains, variance reduction.	
19 ABSTRACT (Continue on reverse if necessary and identify by block number) Abstract. We estimate via simulation the expectation of certain integrals of functionals of continuous-time Markov chains over a finite horizon, fixed or random. By computing conditional expectations given the sequence of states visited (and possibly other information), we reduce variance. This is discrete-time conversion. We further increase efficiency by combining discrete-time conversion with stratification and splitting.			
20 DISTRIBUTION/AVAILABILITY OF ABSTRACT ☐ UNCLASSIFIED/UNLIMITED ☐ SAME AS RPT ☐ DTIC USERS		21 ABSTRACT SECURITY CLASSIFICATION Unclassified	
22a NAME OF RESPONSIBLE INDIVIDUAL		22b TELEPHONE (Include Area Code)	22c OFFICE SYMBOL