

Adaptive Importance Sampling for Uniformly Recurrent Markov Chains

Paul Dupuis* and Hui Wang†

Lefschetz Center for Dynamical Systems
Brown University
Providence, R.I. 02912
USA

Abstract

Importance sampling is a variance reduction technique for efficient estimation of rare-event probabilities by Monte Carlo. In standard importance sampling schemes, the system is simulated using an a priori fixed change of measure suggested by a large deviation lower bound analysis. Recent work, however, has suggested that such schemes do not work well in many situations. In this paper, we consider adaptive importance sampling in the setting of uniformly recurrent Markov chains. By “adaptive,” we mean that the change of measure depends on the history of the samples. Based on a control-theoretic approach to large deviations, the existence of asymptotically optimal adaptive schemes is demonstrated in great generality. In this framework, the difference between a static change of measure and an adaptive change of measure amounts to the difference between an open-loop control and a feed-back control. The implementation of the adaptive schemes is carried out with the help of a limiting Bellman equation. Also presented are numerical examples contrasting the adaptive and standard schemes.

1 Introduction

Among variance reduction techniques for efficient Monte Carlo simulation is importance sampling, in which the data is generated using a probability

*Research of this author supported in part by the National Science Foundation (NSF-DMS-0072004, NSF-ECS-9979250) and the Army Research Office (DAAD19-00-1-0549, DAAD19-02-1-0425).

†Research of this author supported in part by the National Science Foundation (NSF-DMS-0103669).

Report Documentation Page				Form Approved OMB No. 0704-0188	
Public reporting burden for the collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to a penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number.					
1. REPORT DATE 2003		2. REPORT TYPE		3. DATES COVERED 00-00-2003 to 00-00-2003	
4. TITLE AND SUBTITLE Adaptive Importance Sampling for Uniformly Recurrent Markov Chains				5a. CONTRACT NUMBER	
				5b. GRANT NUMBER	
				5c. PROGRAM ELEMENT NUMBER	
6. AUTHOR(S)				5d. PROJECT NUMBER	
				5e. TASK NUMBER	
				5f. WORK UNIT NUMBER	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) Brown University, Division of Applied Mathematics, 182 George Street, Providence, RI, 02912				8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES)				10. SPONSOR/MONITOR'S ACRONYM(S)	
				11. SPONSOR/MONITOR'S REPORT NUMBER(S)	
12. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release; distribution unlimited					
13. SUPPLEMENTARY NOTES					
14. ABSTRACT					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT	18. NUMBER OF PAGES 38	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

distribution different from the true underlying distribution. It can be especially effective when applied to the estimation of expectations that are largely determined by rare events. To demonstrate the difficulty involved in simulating rare events by naive Monte Carlo, we consider a simple example. Let X be a random variable taking value in $\mathbb{R}^d$, and suppose we are interested in estimating $p = P\{X \in A\}$ for some Borel set $A \subset \mathbb{R}^d$. To this end, a sequence of independent and identically distributed (iid) copies $X_0, X_1, \dots$ of X are generated. With $I_k \doteq 1_{\{X_k \in A\}}$, an unbiased estimate for p based on the first K samples is just the sample mean: $Q_K \doteq (I_0 + I_1 + \dots + I_{K-1})/K$. The relative error associated with this estimator is

$$\text{relative error} \doteq \frac{\text{standard deviation of } Q_K}{\text{mean of } Q_K} = \frac{\sqrt{p - p^2}}{p} \cdot \frac{1}{\sqrt{K}}.$$

Since $\sqrt{p - p^2}/p \rightarrow \infty$ as p tends 0, a large sample size K is required for the estimator Q_K to achieve a reasonable relative error bound. For example, if $p = 10^{-8}$, ten billion samples are required to achieve a relative error bound of 10%.

The basic idea of importance sampling is as follows. Suppose that X has distribution θ , and consider an alternative sampling distribution τ . It is required that θ be absolutely continuous with respect to τ , so that the Radon-Nikodym derivative $f(x) \doteq (d\theta/d\tau)(x)$ exists. Independent and identically distributed samples $\bar{X}_0, \bar{X}_1, \dots$ with distribution τ are generated. Set $\bar{I}_k \doteq 1_{\{\bar{X}_k \in A\}}$, and form the estimate

$$\bar{Q}_K \doteq \frac{1}{K} \sum_{k=0}^{K-1} \bar{I}_k f(\bar{X}_k)$$

in lieu of Q_K . It is easy to check that $\bar{Q}_K$ is an unbiased estimate of p , with a rate of convergence determined by

$$\text{var} [\bar{I}_0 f(\bar{X}_0)] = \int_{\mathbb{R}} 1_{\{x \in A\}} f(x) \theta(dx) - p^2.$$

The optimization of this quantity over all possible τ is inappropriate. Indeed, taking $f(x) = p^{-1} 1_{\{x \in A\}}$ (i.e., τ is the conditional distribution of X given $X \in A$), the variance becomes 0, but this change of measure requires the knowledge of the unknown parameter p . Instead, one typically seeks to minimize over parameterized families of alternative sampling distributions.

When the distribution of X is connected to a large deviations problem, a standard heuristic is that the change of measure used to prove the large deviation lower bound should be a good (perhaps nearly optimal) distribution

to use for the purposes of importance sampling. The first result of this type was given by Siegmund [31]. The basic idea was subsequently investigated in many contexts, and a small selection of the literally hundreds of papers on the topics is [1, 2, 3, 7, 8, 9, 11, 12, 15, 17, 19, 23, 24, 26, 27, 30]. Necessary and sufficient conditions under which a prescribed scheme is asymptotically optimal are discussed in [10, 28, 29], while [20] gives a survey of rare-event simulation.

The validity of the heuristic, however, was challenged in [18]. Counterexamples were constructed to show that, under some very common settings, the change of measure suggested by large deviations leads to importance sampling scheme with very poor properties.

In order to explain these counterexamples, and more importantly, to find asymptotically optimal importance sampling algorithms in great generality, [16] introduces an adaptive importance sampling scheme and shows its asymptotic optimality in the setup of iid random variables (Cramér’s Theorem). The key observation is that *many* changes of measure are suggested by the large deviation lower bound analysis, and one must consider this larger class if one hopes to identify importance sampling schemes that work well in general. This leads to the development of schemes where the sampling distribution is “adaptive” in the sense that it depends on the history of the samples.

The present paper analyzes the estimation of rare-event probabilities associated with uniformly recurrent Markov chains. More precisely, let $\{Y_j, j \in \mathbb{N}_0\}$ be a uniformly recurrent Markov chain taking values in a Polish space $\mathcal{S}$, and let $g : \mathcal{S} \rightarrow \mathbb{R}^d$ be a bounded measurable function. Define $S_n \doteq g(Y_0) + g(Y_1) + \cdots + g(Y_{n-1})$. The probability of interest is $P\{S_n/n \in A\}$ for a Borel set $A \subset \mathbb{R}^d$ and n large. An asymptotic optimality result for traditional importance sampling is available in the one-dimensional case ($d = 1$), under the assumption which implies that the set A is within a half interval that does not contain the expectation of g under the invariant distribution [9]. A “dissection” approach was introduced for the high dimensional case [9]. This approach was later on applied to Markov additive sequences [11], and was also implicitly used in [18]. This dissection approach requires that one appropriately partition the set A into a finite number of subsets, and that a (possibly different) change of measure be applied to efficiently estimate the probability of each individual subset. However, there is no constructive way to obtain a suitable partition in general.

In this paper we develop adaptive importance sampling schemes for uniformly recurrent Markov chains. The existence of asymptotically optimal adaptive schemes is demonstrated for arbitrary dimension d , under very mild

conditions on the set A . It turns out that one must study the asymptotics of a small noise stochastic game in order to analyze the optimality of importance sampling schemes. The distinction between the change of measures used in traditional importance sampling and adaptive importance sampling amount, in control terminology, to the difference between “open-loop” and “feedback” controls. However, open loop controls are not nearly optimal in the setting of stochastic games, except for very special cases. For this reason, the traditional importance sampling will not be asymptotically optimal in general. Our analysis indicates that the adaptive scheme also works for estimating functionals (other than probabilities) largely determined by rare events.

The paper is organized as follows. The setting of the problem is introduced in Section 2, with a brief description of the large deviations principle for uniformly recurrent Markov chains. We also give the definition of asymptotic optimality in this section. In Section 3 we show that adaptive importance sampling schemes designed to minimize the second moment are asymptotically optimal. Section 4 discusses an alternative formal PDE approach to the adaptive scheme, and describes a method for the construction of an asymptotically optimal adaptive scheme that does not directly depend on the large deviation parameter n . Numerical examples are presented in Section 4.3. Certain technical proofs are deferred to the appendices to ease exposition.

2 Problem setup and background

2.1 Problem setup

Let $Y = \{Y_j, j \in \mathbb{N}_0\}$ be a time-homogeneous Markov chain taking values in a Polish space $\mathcal{S}$, with transition probability kernel

$$p(x, dy) = P\{Y_{j+1} \in dy \mid Y_j = x\}.$$

Let $g : \mathcal{S} \rightarrow \mathbb{R}^d$ be a bounded Borel-measurable function, and define

$$S_n \doteq g(Y_0) + g(Y_1) + \cdots + g(Y_{n-1}).$$

For an arbitrary Borel set $A \subset \mathbb{R}^d$, we wish to estimate

$$p_n \doteq P\{S_n/n \in A\}.$$

Throughout the paper, we will make use of the following *uniform recurrency* assumption.

Condition 2.1 *There exists a probability measure ν_p on $\mathcal{S}$, an integer $m_0 \in \mathbb{N}$, and a pair of strictly positive real numbers a, b such that*

$$a\nu_p(B) \leq p^{(m_0)}(x, B) \leq b\nu_p(B)$$

for all $x \in \mathcal{S}$ and Borel sets B . Here $p^{(m)}$ denotes the m -step probability transition kernel.

For example, an irreducible Markov chain with a finite state space is always uniformly recurrent.

The large deviation principle for a uniformly recurrent Markov chain is well known. It asserts that $\{S_n/n\}$ satisfies the large deviation principle with a convex rate function $L : \mathbb{R}^d \rightarrow [0, \infty]$. The identification of L is deferred to the next subsection. We will impose the following assumption throughout the paper.

Condition 2.2 *The Borel set $A \subset \mathbb{R}^d$ satisfies the condition*

$$\inf_{\beta \in \bar{A}} L(\beta) = \inf_{\beta \in A^\circ} L(\beta).$$

Under Conditions 2.1 and 2.2, we have the large deviations approximation

$$\lim_{n \rightarrow \infty} \frac{1}{n} \log P\{S_n/n \in A\} = - \inf_{\beta \in A} L(\beta).$$

2.2 LDP for a uniformly recurrent Markov chain

In this subsection we discuss two different approaches to the identification of the rate function L . The first approach suggests a parameterized family of change of measures (see Remark 2.1) that will be used later on to build importance sampling schemes. The second approach identifies the rate function L in terms of relative entropy, and will be used in the analysis of the asymptotic optimality of adaptive schemes.

The first approach is based on a generalized Perron-Frobenius theorem. Fix any $\alpha \in \mathbb{R}^d$. Then by [21], the non-negative kernel

$$\exp\{\langle \alpha, g(y) \rangle\} p(x, dy)$$

admits a unique real eigenvalue $\exp\{H(\alpha)\}$ and a unique (up to a multiplicative constant) eigenfunction $r(x; \alpha)$ in the sense that, for every $x \in \mathcal{S}$,

$$\int_{\mathcal{S}} e^{\langle \alpha, g(y) \rangle} r(y; \alpha) p(x, dy) = e^{H(\alpha)} r(x; \alpha), \quad (2.1)$$

and with the following properties. $H(\alpha)$ is an analytic, strictly convex function of $\alpha \in \mathbb{R}^d$ with $H(0) = 0$, and there exist $0 < c_\alpha < C_\alpha < \infty$ such that

$$c_\alpha \leq r(x; \alpha) \leq C_\alpha, \quad \forall x \in \mathcal{S}. \quad (2.2)$$

The paper [21] also shows that the rate function of the large deviation principle for $\{S_n/n\}$ is the convex conjugate of H , i.e.,

$$L(\beta) = \sup_{\alpha \in \mathbb{R}^d} [\langle \alpha, \beta \rangle - H(\alpha)]. \quad (2.3)$$

Note that in the special case when the Markov chain Y is an iid sequence, $H(\alpha)$ is the logarithm moment generating function of $g(Y_j)$ and $r(x; \alpha) \equiv 1$. Therefore, this result generalizes the classical Cramér's Theorem, at least for bounded iid random variables. For the case when Y is an irreducible Markov chain with finite state space, $\exp\{H(\alpha)\}$ is just the maximal eigenvalue of the irreducible non-negative matrix $\exp\{\langle \alpha, g(y) \rangle\} p(x, dy)$, and $r(\cdot; \alpha)$ is the associated right eigenvector.

Remark 2.1 It is not difficult to see that, thanks to (2.1), for each $\alpha \in \mathbb{R}^d$,

$$\exp\{\langle \alpha, g(y) \rangle - H(\alpha)\} \cdot \frac{r(y; \alpha)}{r(x; \alpha)} \cdot p(x, dy)$$

defines a probability transition kernel.

Another approach is the weak convergence methodology which utilizes a stochastic control representation for certain exponential integrals [14]. It first identifies the large deviations rate function for the empirical measure of the Markov chain in the τ -topology, then uses contraction principle to obtain the rate function for $\{S_n/n\}$. We will need the following definitions.

For an arbitrary Polish space $\mathcal{Z}$, we denote by $\mathcal{P}(\mathcal{Z})$ the collection of all probability measures on space $(\mathcal{Z}, \mathcal{B}(\mathcal{Z}))$. For a pair of probability measures $\gamma, \mu \in \mathcal{P}(\mathcal{Z})$, the *relative entropy* of γ with respect to μ is defined as

$$R(\gamma \parallel \mu) \doteq \int_{\mathcal{Z}} \log \frac{d\gamma}{d\mu} d\gamma$$

if $\gamma \ll \mu$ and $R(\gamma \parallel \mu) \doteq \infty$ otherwise. Given a probability transition kernel $q(x, dy)$ on space $\mathcal{Z}$, we define $\mu q \in \mathcal{P}(\mathcal{Z})$, $\mu \otimes q \in \mathcal{P}(\mathcal{Z} \times \mathcal{Z})$ by

$$\begin{aligned} \mu q(B) &\doteq \int_{\mathcal{Z}} q(x, B) \mu(dx) \\ (\mu \otimes q)(D \times B) &\doteq \int_{D \times B} \mu(dx) q(x, dy) = \int_D q(x, B) \mu(dx) \end{aligned}$$

for all Borel sets $D, B \subset \mathcal{Z}$. The collection of all probability transition kernels on $\mathcal{Z}$ is denoted by $\mathcal{T}(\mathcal{Z})$.

The weak convergence approach identifies the rate function for $\{S_n/n\}$ in terms of relative entropy:

$$L(\beta) = \inf \left\{ R(\mu \otimes q \| \mu \otimes p) : \mu \in \mathcal{P}(\mathcal{S}), q \in \mathcal{T}(\mathcal{S}), \mu q = \mu, \int_{\mathcal{S}} g d\mu = \beta \right\}. \quad (2.4)$$

The validity of the representation (2.4) is implied by the results in [14, Chapters 8 & 9], where the large deviation principle of the empirical measures associated with Markov chains are studied under weaker assumptions.

For future reference, we summarize the preceding discussion into the following proposition. The only part that has not been mentioned is the superlinearity of the rate function L , which is an easy consequence of (2.3) and the finiteness of H [14, Lemma 6.2.3(c)].

Proposition 2.1 *Under Condition 2.1, the sequence $\{S_n/n\}$ satisfies the large deviation principle with rate function L , which is given by equations (2.3) and (2.4). Moreover, the rate function L is convex, lower-semicontinuous, and superlinear in the sense that*

$$\lim_{N \rightarrow \infty} \inf_{\{\beta \in \mathbb{R}^d : \|\beta\| \geq N\}} \frac{L(\beta)}{\|\beta\|} = \infty.$$

In particular, L has compact level sets.

2.3 Asymptotic optimality

In this subsection we define *asymptotic optimality* for an importance sampling scheme.

Consider a probability space $(\Omega, \mathcal{F}, P)$ and a family of events $\{A_n\}$ with

$$\lim_{n \rightarrow \infty} \frac{1}{n} \log P\{A_n\} = -\gamma,$$

for some $\gamma \geq 0$. A general formulation of importance sampling for this problem can be described as follows. In order to estimate $P\{A_n\}$, a generic random variable $\bar{Z}_n$ is constructed such that $P\{A_n\} = E\bar{Z}_n$. Independent replications $(\bar{Z}_n^0, \bar{Z}_n^1, \dots, \bar{Z}_n^{K-1})$ of $\bar{Z}_n$ are then generated, and we obtain an estimator by averaging:

$$\bar{Q}_n^K \doteq \frac{\bar{Z}_n^0 + \bar{Z}_n^1 + \dots + \bar{Z}_n^{K-1}}{K}.$$

The estimator is unbiased, i.e., $E\bar{Q}_n^K = P\{A_n\}$. The rate of convergence associated with this estimator is determined by the variance of the summands, or equivalently, their second moment $E[(\bar{Z}_n)^2]$. The smaller the second moment, the faster the convergence, whence the smaller sample size K required. However, it follows from Jensen's inequality that

$$\limsup_{n \rightarrow \infty} -\frac{1}{n} \log E[(\bar{Z}_n)^2] \leq \lim_{n \rightarrow \infty} -\frac{1}{n} \log (E\bar{Z}_n)^2 = 2\gamma.$$

The estimator $\bar{Q}_n^K$ is said to be *asymptotically optimal* if

$$\lim_{n \rightarrow \infty} -\frac{1}{n} \log E[(\bar{Z}_n)^2] = 2\gamma.$$

Remark 2.2 Since the performance of the estimator $\bar{Q}_n^K$ is completely determined by the second moment of its generic, iid building block $\bar{Z}_n^k$, we will drop the superscript k hereafter. Note that n does *not* stand for sample size, but for the large deviation parameter.

3 Statement of the main result

The adaptive importance sampling scheme we consider dynamically selects the change of measure (or the parameter α) in the form suggested by Remark 2.1, according to the sample history. Naturally, the scheme is closely related to a control problem. Let the control $\alpha^n = \{\alpha_j^n(\cdot, \cdot), j = 1, \dots, n-1\}$ be given, where each $\alpha_j^n : \mathcal{S} \times \mathbb{R}^d \rightarrow \mathbb{R}^d$ is a Borel-measurable function. Then the state dynamics are governed by

$$\bar{S}_j^n \doteq \sum_{i=0}^{j-1} g(\bar{Y}_i^n), \quad j = 0, 1, \dots, n.$$

Here we set $\bar{Y}_0 = Y_0 \equiv y_0$, and for $j \geq 1$, $\bar{Y}_j^n$ is conditionally distributed, given $\{\bar{Y}_i^n, i = 0, 1, \dots, j-1\}$, according to

$$v_j^n(dy) = \exp \left\{ \langle \alpha_j^n, g(y) \rangle - H(\alpha_j^n) \right\} \cdot \frac{r(y; \alpha_j^n)}{r(\bar{Y}_{j-1}^n; \alpha_j^n)} \cdot p(\bar{Y}_{j-1}^n, dy)$$

with (abusing notation a bit) $\alpha_j^n = \alpha_j^n(\bar{Y}_{j-1}^n, \bar{S}_j^n/n)$.

An unbiased estimator of $P\{S_n/n \in A\}$ is defined as the average of independent copies of

$$\bar{X}_n = 1_{\{\bar{S}_n^n/n \in A\}} e^{\sum_{j=1}^{n-1} (-\langle \alpha_j^n, g(\bar{Y}_j^n) \rangle + H(\alpha_j^n))} \cdot \prod_{j=1}^{n-1} \frac{r(\bar{Y}_{j-1}^n; \alpha_j^n)}{r(\bar{Y}_j^n; \alpha_j^n)}.$$

Our goal is to minimize the second moment, hence the variance, of the summands $\bar{X}_n$ by judiciously choosing the control α^n . Thus we consider the value function defined by

$$\begin{aligned} V^n(y_0) &\doteq \inf_{\alpha^n} E[\bar{X}_n^2] \\ &= \inf_{\alpha^n} E \left[1_{\{\bar{S}_n^n/n \in A\}} e^{\sum_{j=1}^{n-1} (-2\langle \alpha_j^n, g(\bar{Y}_j^n) \rangle + 2H(\alpha_j^n))} \cdot \prod_{j=1}^{n-1} \frac{r^2(\bar{Y}_{j-1}^n; \alpha_j^n)}{r^2(\bar{Y}_j^n; \alpha_j^n)} \right] \end{aligned}$$

For convenience we write $V^n(y_0)$ as V^n when no confusion is incurred. We also consider the log transform

$$W^n = -\frac{1}{n} \log V^n.$$

We have the following result, which asserts the existence of asymptotically optimal adaptive importance sampling schemes.

Theorem 3.1 *Under Conditions 2.1 and 2.2, we have*

$$\lim_{n \rightarrow \infty} W^n = 2 \inf_{\beta \in A} L(\beta).$$

The detailed proof is deferred to Appendix A. It is worth to point out that the construction of asymptotically optimal or nearly optimal adaptive schemes (i.e. selection of the control α^n) is implied by a *dynamic programming equation* (DPE) appearing in the proof. Since the proof is rather lengthy and technical, it makes sense to give an outline and some intuitive discussion below, so that readers can proceed to the construction of the adaptive schemes (Section 4), without having to delve into the technical details of the proof.

Outline and intuition of the proof: Thanks to the discussion in Section 2.3, it suffices to show the lower bound

$$\liminf_n W^n \geq 2 \inf_{\beta \in A} L(\beta). \quad (3.1)$$

The proof will utilize the DPE that is satisfied by W^n . In order to do so, we first extend the dynamics. Abusing notation a bit, for $x \in \mathbb{R}^d$, $y \in \mathcal{S}$, and $i \in \{0, 1, \dots, n\}$, define the dynamics

$$\bar{S}_{i,j}^n = nx + \sum_{\ell=i}^{j-1} \bar{Y}_{i,\ell}^n, \quad j = i, \dots, n.$$

Here we set $\bar{Y}_{i,i} \equiv y$, and for $j \geq i+1$, $\bar{Y}_{i,j}^n$ is conditionally distributed, given $\{\bar{Y}_{i,\ell}, \ell = i, \dots, j-1\}$, according to

$$v_{i,j}^n(dz) = \exp \left\{ \langle \alpha_j^n, g(z) \rangle - H(\alpha_j^n) \right\} \cdot \frac{r(z; \alpha_j^n)}{r(\bar{Y}_{i,j-1}^n; \alpha_j^n)} p(\bar{Y}_{i,j-1}^n, dz),$$

where $\alpha_j^n = \alpha_j^n(\bar{Y}_{i,j-1}^n, \bar{S}_{i,j}^n/n)$. The original control problem corresponds to $x = 0, i = 0, y = y_0$. Define analogously

$$\begin{aligned} V^n(x, y; i) \\ \doteq \inf_{\alpha^n} E \left[1_{\{\bar{S}_{i,n}^n/n \in A\}} e^{\sum_{j=i+1}^{n-1} (-2\langle \alpha_j^n, g(\bar{Y}_{i,j}^n) \rangle + 2H(\alpha_j^n))} \prod_{j=i+1}^{n-1} \frac{r^2(\bar{Y}_{i,j-1}^n; \alpha_j^n)}{r^2(\bar{Y}_{i,j}^n; \alpha_j^n)} \right] \end{aligned}$$

and its log transform

$$W^n(x, y; i) = -\frac{1}{n} \log V^n(x, y; i).$$

The terminal conditions are

$$V^n(x, y; n) = 1_A(x), \quad W^n(x, y; n) = \infty \cdot 1_A(x).$$

Since it is inconvenient to study a problem with an ∞ terminal condition, we instead work with a mollified version of the control problem. Let $F : \mathbb{R}^d \rightarrow \mathbb{R}$ be an arbitrary bounded and Lipschitz continuous function. Suppose that V_F^n is defined as V^n , save that the indicator function $1_{\{\bar{S}_{i,n}^n/n \in A\}}$ is replaced by $\exp\{-2nF(\bar{S}_{i,n}^n/n)\}$. Similarly define

$$W_F^n(x, y; i) \doteq -\frac{1}{n} \log V_F^n(x, y; i). \quad (3.2)$$

Since V_F^n is the value function of a control problem, one can write down the DPE for V_F^n . Substituting (3.2) in this DPE, one obtains an equation for W_F^n ; see (A.1). The proof of the desired inequality (3.1) is based on the analysis of this recursive equation for W_F^n .

The relative entropy representation for exponential integrals [14, Proposition 1.4.2] states that

$$-\log \int_{\mathcal{S}} e^{-f(x)} \mu(dx) = \inf_{\gamma \in \mathcal{P}(\mathcal{S})} \left[R(\gamma \| \mu) + \int f d\gamma \right] \quad (3.3)$$

for all bounded and Borel measurable functions f . Applying this representation formula to the equation for W_F^n , one obtains

$$\begin{aligned} W_F^n(x, y; i) &= \sup_{\alpha \in \mathbb{R}^d} \inf_{\gamma \in \mathcal{P}(\mathcal{S})} \left[\int W_F^n \left(x + \frac{1}{n} g(y), z; i+1 \right) \gamma(dz) \right. \\ &\quad + \frac{1}{n} \left(R(\gamma(\cdot) \| p(y, \cdot)) + \int \langle \alpha, g(z) \rangle \gamma(dz) - H(\alpha) \right) \\ &\quad \left. + \frac{1}{n} \int \log \frac{r(z; \alpha)}{r(y; \alpha)} \gamma(dz) \right]. \end{aligned} \quad (3.4)$$

This equation suggests that W_F^n is the lower value of a discrete-time stochastic game. One of the two players of the game (the α -player) selects the parameter α , and is the weaker player. The other player (the γ -player) is the stronger player, and selects the distribution γ that determines the evolution of the state. The right hand side of (3.4) would take a simpler form if we could permute the sup and inf. However, this is not (in general) possible, since the last term

$$\frac{1}{n} \int \log \frac{r(z; \alpha)}{r(y; \alpha)} \gamma(dz) \quad (3.5)$$

may not be concave in α .

This difficulty is also the main distinction from the setting of Cramér's Theorem where the Markov chain Y reduces to an iid sequence of random variables. The latter case gives $r(x; \alpha) \equiv 1$ and the unpleasant term (3.5) disappears, whence min/max theorem can be applied to convert the DPE of W_F^n into a DPE associated with a *control* problem, which is much simpler to analyze than a game [16]. However, the interchange of sup and inf is not possible with (3.4) as written.

The key idea to overcome this difficulty and to obtain a lower bound for W_F^n is as follows. Fix an integer m , and consider a variant of the game where the α -player is constrained to policies such that α must be constant over time intervals of length $1/m$. This new game is even more favorable to the γ -player, whence it will have a smaller lower-value. Letting n go to infinity, the lower value of the new game converges to a function U_F^m , and we expect

$$\liminf_{n \rightarrow \infty} W_F^n(x, y; \lfloor nk/m \rfloor) \geq U_F^m(x; k), \quad k = 0, 1, \dots, m.$$

A bonus of taking the limit is that the troubling terms (3.5), which can be interpreted as part of the running cost, cancel off, and it is not difficult to

guess that U_F^m should satisfy

$$U_F^m(x; k) = \sup_{\alpha \in \mathbb{R}^d} \inf_{\beta \in \mathbb{R}^d} \left[U_F^m \left(x + \frac{1}{m} \beta; k+1 \right) + \frac{1}{m} (L(\beta) + \langle \alpha, \beta \rangle - H(\alpha)) \right], \quad (3.6)$$

with terminal condition

$$U_F^m(x; m) \doteq 2F(x). \quad (3.7)$$

In the proof, U_F^m is in fact defined recursively through equations (3.6) and (3.7).

Equation (3.6) is much easier to analyze. Analogous to [16], one can show by a weak convergence argument that

$$\liminf_{m \rightarrow \infty} U_F^m(x, 0) \geq 2 \inf_{\beta \in \mathbb{R}^d} \{L(\beta) + F(x + \beta)\}, \quad (3.8)$$

which in turn implies

$$\liminf_{n \rightarrow \infty} W_F^n(x, y; 0) \geq 2 \inf_{\beta \in \mathbb{R}^d} \{L(\beta) + F(x + \beta)\}.$$

Letting $x = 0$ and the mollifier F tend to $\infty \cdot 1_A$, one arrives at the desired inequality (3.1). ■

The following result is useful in the identification of optimal adaptive importance sampling scheme in Section 4.

Corollary 3.2 *Fix an arbitrary $x \in \mathbb{R}^d$, and a bounded Lipschitz continuous function $F : \mathbb{R}^d \rightarrow \mathbb{R}$. Assume Condition 2.1, and define U_F^m recursively by (3.6) with the terminal condition (3.7). Then*

$$\lim_{m \rightarrow \infty} U_F^m(x; \lfloor tm \rfloor) = 2U_F(x, t), \quad \forall t \in [0, 1],$$

where

$$U_F(x, t) \doteq \inf_{\beta \in \mathbb{R}^d} \{(1-t)L(\beta) + F(x + (1-t)\beta)\}. \quad (3.9)$$

Proof. We will show the inequality for $t = 0$. The case with general $t \in [0, 1]$ is similar and thus omitted.

Thanks to (3.8), it suffices to prove

$$\limsup_{m \rightarrow \infty} U_F^m(x; 0) \leq 2U_F(x, 0) = 2 \inf_{\beta \in \mathbb{R}^d} \{L(\beta) + F(x + \beta)\}.$$

Fix an arbitrary $\beta \in \mathbb{R}^d$. The recursive definition of U_F^m (3.6) and equation (2.3) yield

$$\begin{aligned} U_F^m(x; k) &\leq \sup_{\alpha \in \mathbb{R}^d} \left[U_F^m \left(x + \frac{1}{m} \beta; k+1 \right) + \frac{1}{m} (L(\beta) + \langle \alpha, \beta \rangle - H(\alpha)) \right] \\ &= U_F^m \left(x + \frac{1}{m} \beta; k+1 \right) + \frac{2}{m} L(\beta). \end{aligned}$$

Repeatedly applying this inequality for $k = 0, 1, \dots, m-1$, we arrive at

$$U_F^m(x; 0) \leq U_F^m(x + \beta; m) + 2L(\beta) = 2F(x + \beta) + 2L(\beta),$$

thanks to (3.7). This completes the proof. $\blacksquare$

4 Implementation issues and examples

4.1 The limit control problem and implementation issues

Theorem 3.1 establishes the existence of asymptotically optimal adaptive sampling schemes. However, it does not explicitly discuss the construction of such schemes. On the other hand, the proof of the theorem implies that one approach of construction would be to solve, numerically if need be, the DPE (3.4) associated with W_F^n (W^n equals W_F^n when $F = \infty \cdot 1_A$). However, this approach may not only require a lot of computation effort, but the resulting adaptive sampling control (i.e., control α^n) will directly depend on n . In general, one would prefer schemes without this dependence.

An alternative approach is to consider the DPE associated with the limit problem of U_F^m as m tends to infinity. To this end, we rewrite equation (3.6) as

$$0 = \sup_{\alpha \in \mathbb{R}^d} \inf_{\beta \in \mathbb{R}^d} \left[\Delta U_F^m + \frac{1}{m} (L(\beta) + \langle \alpha, \beta \rangle - H(\alpha)) \right],$$

where

$$\Delta U_F^m \doteq U_F^m \left(x + \frac{1}{m} \beta; k+1 \right) - U_F^m(x; k).$$

Suppose that a t subscript denotes the partial derivative with respect to t , and that an x subscript denotes the vector of partials with respect to $x_i, i = 1, \dots, d$. Since Corollary 3.2 (for F bounded and Lipschitz continuous) asserts that

$$\lim_{m \rightarrow \infty} U_F^m(x; \lfloor tm \rfloor) = 2U_F(x; t),$$

we have formally the approximation

$$\Delta U_F^m \approx \left\langle \frac{1}{m} \beta, (2U_F)_x \right\rangle + \frac{1}{m} (2U_F)_t.$$

Substituting this back, we have

$$\begin{aligned} 0 &= \sup_{\alpha \in \mathbb{R}^d} \inf_{\beta \in \mathbb{R}^d} [\langle \beta, (2U_F)_x \rangle + (2U_F)_t + L(\beta) + \langle \alpha, \beta \rangle - H(\alpha)] \\ &= (2U_F)_t + \sup_{\alpha \in \mathbb{R}^d} \inf_{\beta \in \mathbb{R}^d} [L(\beta) + \langle \alpha + (2U_F)_x, \beta \rangle - H(\alpha)]. \end{aligned}$$

Representing the infimum in terms of the Legendre transform H of L gives

$$0 = (2U_F)_t + \sup_{\alpha \in \mathbb{R}^d} [-H(-\alpha - (2U_F)_x) - H(\alpha)].$$

The strict convexity of H implies that

$$\alpha^*(x, t) = -(U_F)_x(x, t), \quad (4.1)$$

and that

$$0 = (U_F)_t - H(-(U_F)_x). \quad (4.2)$$

The equation (4.1) identifies, at least formally, an optimal feedback control policy. However, this observation is not entirely satisfactory for several reasons. The first is that even if we have an exact formula for U_F , the partial derivatives may not be defined for all times and spatial points. A second reason is that U_F does not usually have an explicit solution, and so in many cases, numerical approximation is required. In order to obtain a formal characterization of α^* that is more useful, we observe that, thanks to the definition (3.9) of U_F and the convexity of L , U_F is the value function of the deterministic control problem

$$U_F(x, t) = \inf_{\phi} \left[\int_t^1 L(\dot{\phi}(s)) ds + F(\phi(1)) \right],$$

where the infimum is over all absolutely continuous ϕ which satisfy $\phi(t) = x$. It is straightforward to see from this control problem that an optimal control at (x, t) is the minimizer in (3.9), say $\beta^*(x, t)$, thanks to the convexity of L . Note that under very mild conditions $\beta^*(x, t)$ will always exist, e.g., if F is continuous and bounded. The standard dynamic programming argument implies that U_F (in a weak sense) satisfies the DPE

$$0 = (U_F)_t + \inf_{a \in \mathbb{R}^d} [L(a) + \langle a, (U_F)_x \rangle] = (U_F)_t - H(-(U_F)_x),$$

which, not surprisingly, is just equation (4.2). The optimal control $\beta^*(x, t)$ is, at least formally, the minimizer in the DPE, or, $\beta^*(x, t)$ and $-(U_F)_x(x, t)$ are conjugate. It follows that

$\alpha^*(x, t)$ is conjugate to the minimizer $\beta^*(x, t)$ in (3.9).

At points where $(U_F)_x(x, t)$ exists this definition gives $\alpha^*(x, t) = -(U_F)_x(x, t)$. At points where $(U_F)_x(x, t)$ does not exist there are multiple minimizing $\beta^*(x, t)$, and one should define $\alpha^*(x, t)$ through conjugacy in any Borel measurable way.

Remark 4.1 The original (unmollified) problem corresponds to $F = \infty \cdot 1_A$. In this case,

$$\beta^*(x, t) \in \operatorname{argmin}\{(1-t)L(\beta) : x + (1-t)\beta \in A\}, \quad (4.3)$$

and $\alpha^*(x, t)$ is its conjugate. Roughly speaking, this means that the controlled random walk always points to the “most likely” end point, given that the end point is in A and that we are currently at x with $1-t$ units of time to go. Note that we can assume without loss that A is closed. This suffices to guarantee the existence of $\beta^*(x, t)$ for all $x \in \mathbb{R}^d, t \in [0, 1]$.

4.2 A numerical example

Consider a simple finite-state Markov chain Y with state space $\mathcal{S} = \{1, -1\}$, and probability transition matrix

$$Q = \begin{bmatrix} p & 1-p \\ 1 & 0 \end{bmatrix} \doteq [Q(i, j)]_{2 \times 2},$$

for some constant $p \in (0, 1)$. Define $g : \mathcal{S} \rightarrow \mathbb{R}$ by $g(x) \doteq x$, and $S_n = g(Y_0) + g(Y_1) + \cdots + g(Y_{n-1}) = Y_0 + Y_1 + \cdots + Y_{n-1}$. Assume that the initial state is $Y_0 \equiv 1$. We wish to estimate $P\{S_n/n \in A\}$ for some Borel set $A \in \mathbb{R}$.

Since Y is an irreducible finite-state Markov chain, the eigenvalue $e^{H(\alpha)}$ and eigenfunction $r(\cdot; \alpha)$, as defined in (2.1) for $\alpha \in \mathbb{R}$, are just the maximal eigenvalue of the kernel

$$[e^{\alpha j} Q(i, j)] = \begin{bmatrix} pe^\alpha & (1-p)e^{-\alpha} \\ e^\alpha & 0 \end{bmatrix}$$

and the corresponding eigenvector, respectively. Simple algebra gives

$$H(\alpha) = \log \frac{pe^\alpha + \sqrt{p^2 e^{2\alpha} + 4(1-p)}}{2}, \quad \forall \alpha \in \mathbb{R},$$

which is a convex function with $H(0) = 0$, and an eigenvector

$$\begin{bmatrix} r(1; \alpha) \\ r(-1; \alpha) \end{bmatrix} = \begin{bmatrix} e^{H(\alpha)} \\ e^\alpha \end{bmatrix}.$$

Therefore, for any given $\alpha \in \mathbb{R}$, the corresponding change of measure is represented by the probability transition matrix

$$Q_\alpha = \left[e^{\alpha j - H(\alpha)} \frac{r(j; \alpha)}{r(i; \alpha)} Q(i, j) \right] = \begin{bmatrix} pe^{\alpha - H(\alpha)} & (1-p)e^{-2H(\alpha)} \\ 1 & 0 \end{bmatrix}. \quad (4.4)$$

In order to compute the convex conjugate of H , we first observe that $H'(-\infty) = 0$ and $H'(+\infty) = 1$. It follows immediately that $L(\beta) = \infty$ if $\beta < 0$ or $\beta > 1$. For $\beta \in (0, 1)$, some algebra yields

$$\begin{aligned} L(\beta) &= \sup_{\alpha \in \mathbb{R}} [\alpha\beta - H(\alpha)] \\ &= \frac{\beta}{2} \log \frac{4(1-p)\beta^2}{p^2(1-\beta^2)} + \frac{1}{2} \log \frac{1-\beta}{1+\beta} - \frac{1}{2} \log(1-p) \end{aligned}$$

with the minimizer

$$\alpha^* \doteq \alpha^*(\beta) = \frac{1}{2} \log \frac{4(1-p)\beta^2}{p^2(1-\beta^2)}, \quad (4.5)$$

and

$$L(0) = \lim_{\beta \downarrow 0} L(\beta) = -\frac{1}{2} \log(1-p), \quad L(1) = \lim_{\beta \uparrow 1} L(\beta) = -\log p.$$

Furthermore, $L(\beta) = 0$ if and only if $\beta = H'(0) = p/(2-p)$. Of course, $H'(0) = \int_{\mathcal{S}} g(y) \nu(dy)$ where ν is the stationary distribution of the Markov chain Y .

We are interested in estimating $p_n \doteq P\{S_n/n \in A\}$ for the Borel set

$$A = (-\infty, a] \cup [b, \infty), \quad 0 < a < H'(0) < b < 1.$$

In all the following discussion, we take $p = 1/2$, $a = 1/6$, $b = 1/2$, which implies

$$\inf_{\beta \in A} L(\beta) = L(b) < L(a). \quad (4.6)$$

We will compare the naive Monte Carlo simulation, traditional importance sampling, and adaptive importance sampling schemes below.

The naive Monte Carlo simulation will simulate the Markov chain under the original transition probability kernel Q . One can also regard this as a special change of measure with the corresponding $\alpha = 0$. In this case, the estimate is just the sample mean of K iid replications of

$$X_n = 1_{\{S_n/n \in A\}}.$$

Since the second moment of X_n satisfies

$$\lim_{n \rightarrow \infty} -\frac{1}{n} \log E \left[(X_n)^2 \right] = \lim_{n \rightarrow \infty} -\frac{1}{n} \log p_n = \inf_{\beta \in A} L(\beta) = L(b) < 2L(b),$$

the naive Monte Carlo sampling is not asymptotically optimal.

Thanks to (4.6), the traditional importance sampling will take $\beta^* = b$, and α^* is then defined by (4.5). The algorithm will generate a Markov chain $\tilde{Y}$ with probability transition matrix Q_{α^*} and $\tilde{Y}_0 \equiv 1$. Let $\tilde{S}_n \doteq \tilde{Y}_0 + \dots + \tilde{Y}_{n-1}$. The estimate is the sample mean of K iid replications of

$$\begin{aligned} \tilde{X}_n &= 1_{\{\tilde{S}_n/n \in A\}} \prod_{j=1}^{n-1} e^{-\alpha^* \tilde{Y}_j + H(\alpha^*)} \cdot \prod_{j=1}^{n-1} \frac{r(\tilde{Y}_{j-1}; \alpha^*)}{r(\tilde{Y}_j; \alpha^*)} \\ &= 1_{\{\tilde{S}_n/n \in A\}} e^{-\alpha^* \tilde{S}_n + nH(\alpha^*)} \cdot e^{\alpha^* \tilde{Y}_0 - H(\alpha^*)} \frac{r(\tilde{Y}_0; \alpha^*)}{r(\tilde{Y}_{n-1}; \alpha^*)}. \end{aligned}$$

Since $r(\cdot; \alpha^*)$ is clearly bounded from above and bounded away from zero, it is not difficult to see that

$$\lim_{n \rightarrow \infty} -\frac{1}{n} \log E \left[(\tilde{X}_n)^2 \right] = \lim_{n \rightarrow \infty} -\frac{1}{n} \log E \left[1_{\{\tilde{S}_n/n \in A\}} e^{-2n(\alpha^* \tilde{S}_n/n - H(\alpha^*))} \right].$$

Simple computation yields that $\{\tilde{S}_n/n\}$ satisfies the large deviation principle with rate function $\tilde{L}(\beta) = L(\beta) + H(\alpha^*) - \alpha^* \beta$. Now one can apply the Varadhan's Theorem [14, Theorem 1.3.4] (with slight modification) to show

$$\begin{aligned} \lim_{n \rightarrow \infty} -\frac{1}{n} \log E \left[(\tilde{X}_n)^2 \right] &= \inf_{\beta \in A} \left[2\alpha^* \beta - 2H(\alpha^*) + \tilde{L}(\beta) \right] \\ &= \inf_{\beta \in A} \left[\alpha^* \beta - H(\alpha^*) + L(\beta) \right]. \end{aligned}$$

In the configuration of this example, the infimum in the right-hand-side is achieved at $\beta = a$, and

$$\lim_{n \rightarrow \infty} -\frac{1}{n} \log E \left[(\tilde{X}_n)^2 \right] = a\alpha^* - H(\alpha^*) + L(a) < 2L(b).$$

Therefore, the traditional importance sampling scheme is not asymptotically optimal.

In Section 3, we argue the existence of asymptotically optimal adaptive importance sampling schemes in general. The construction of such adaptive schemes involves the selection of a nearly optimal control $\alpha^n = \{\alpha_j^n(\cdot, \cdot) : j = 0, 1, \dots, n-1\}$. It was formally suggested in Section 4.1 that a good choice is to sample $\bar{Y}_j^n$, conditional on $\{\bar{Y}_i^n, i = 0, \dots, j-1\}$, according to the transition probability matrix Q_α as in (4.4) with α being the conjugate of $\beta^*(x, t)$ given in (4.3) where $x = \bar{S}_n^n/n = (\bar{Y}_0^n + \dots + \bar{Y}_{j-1}^n)/n$ and $t = 1 - j/n$. In case the conjugate of $\beta^*(x, t)$ is ∞ or $-\infty$, α is taken as a large positive or negative number; see Remark 4.3 for more details. The estimate is the sample mean of K iid replications of

$$\bar{X}_n = 1_{\{\bar{S}_n^n/n \in A\}} e^{\sum_{j=1}^{n-1} (-\langle \alpha_j^n, g(\bar{Y}_j^n) \rangle + H(\alpha_j^n))} \cdot \prod_{j=1}^{n-1} \frac{r(\bar{Y}_{j-1}^n; \alpha_j^n)}{r(\bar{Y}_j^n; \alpha_j^n)}.$$

The numerical results show that the controls constructed in this way have asymptotically optimal performance (Table 6).

The numerical results are reported below for $n = 60$. The theoretical value of p_n is

$$p_n = P\{S_n/n \leq a\} + P\{S_n/n \geq b\} = 0.83\% + 2.44\% = 3.27\%.$$

See Remark 4.2 for the computation of this theoretical value. For each tableau, we run four simulations each with sample size $K = 10000$.

	No. 1	No. 2	No. 3	No. 4
Estimate $\hat{p}_n$ (%)	3.11	3.20	3.23	3.09
Standard Error (%)	0.17	0.18	0.18	0.17
95% Confidence Interval (%)	[2.76, 3.46]	[2.85, 3.55]	[2.88, 3.58]	[2.74, 3.44]

Table 1. Naive Monte Carlo Scheme

	No. 1	No. 2	No. 3	No. 4
Estimate $\hat{p}_n$ (%)	2.41	2.48	2.44	16.71
Standard Error (%)	0.04	0.04	0.04	14.22
95% Confidence Interval (%)	[2.34, 2.48]	[2.41, 2.56]	[2.37, 2.51]	[-11.73, 45.15]

Table 2. Traditional Importance Sampling Scheme

	No. 1	No. 2	No. 3	No. 4
Estimate $\hat{p}_n$ (%)	3.17	3.21	3.35	3.33
Standard Error (%)	0.15	0.13	0.17	0.18
95% Confidence Interval (%)	[2.86, 3.47]	[2.85, 3.47]	[3.00, 3.69]	[2.96, 3.70]

Table 3. Adaptive Importance Sampling Scheme

An interesting observation is that the traditional importance sampling scheme exhibits seemingly bizarre and inconsistent simulation results (Table 2). The explanation is as follows. Under the alternative sampling distribution Q_{α^*} , most of the sample means $\tilde{S}_n/n$ will end up near the point b . However, a few samples (“rogue” trajectories) have means that fall into the interval $(-\infty, a]$. Even though the “rogue” trajectories are rare, the Radon-Nikodym derivatives associated with them are so large that they dominate the variance. In simulation No. 4, the presence of a single “rogue” trajectory greatly raises the standard error associated with the estimate. Indeed, the proportion of the contribution to the second moment from this single “rogue” trajectories is more than 99%. In simulations No. 1, No. 2, and No. 3, however, there are no “rogue” trajectories, and the standard error associated with the estimate is deceptively small. The reason is that the standard error is itself estimated from the sample. Without “rogue” trajectories, we actually underestimate the standard error. Therefore, we cannot put much confidence in the standard errors thus obtained, or in the “tight” confidence intervals that follow. Note that the confidence intervals from these three simulations do *not* contain the true value. Similar phenomenon also occurs in the setting of Cramér’s Theorem, that is, where the Markov chain Y reduces to a sequence of iid random variables; see [18, 16].

In contrast, the adaptive importance sampling, on the other hand, yields more accurate estimates and its performance is much more stable. Even though it does not show great advantage over naive Monte Carlo simulation for $n = 60$, it quickly does so when n gets larger. The numerical results for different n (with $K = 10000$ fixed as before) are reported in Table 4–6.

The naive Monte Carlo does not work well for bigger n . For $n = 120$ and $n = 180$, it yields estimates with large standard errors, and for $n = 240$, the simulation yields an estimate 0, i.e., no sample mean reaches the target set A . As for the traditional importance sampling, each simulation gives a very “tight” confidence interval, due to the absence of “rogue” trajectories. However, as discussed before, we cannot put much belief into these estimates. Indeed, none of these confidence intervals cover the true value of p_n .

On the other hand, the adaptive importance scheme yields much more accurate estimates. In Table 6, the variable $\hat{V}^n$ denotes the sample estimate of the second moment $E[(\bar{X}_n)^2]$. Observe that as n gets larger, the ratio

$$\frac{-\frac{1}{n} \log E[(\bar{X}_n)^2]}{-\frac{1}{n} \log p_n} = \frac{-\log E[(\bar{X}_n)^2]}{-\log p_n} \approx \frac{-\log \hat{V}^n}{-\log \hat{p}^n}$$

approaches 2. In other words, the adaptive importance sampling scheme is approaching optimality.

	$n = 120$	$n = 180$	$n = 240$
Theoretical p_n	1.61×10^{-3}	9.66×10^{-5}	6.35×10^{-6}
Estimate $\hat{p}_n$	1.80×10^{-3}	20.00×10^{-5}	0
Standard Error	0.42×10^{-3}	14.14×10^{-5}	NA
95% Confidence Interval	$[0.95, 2.65] \times 10^{-3}$	$[-8.28, 48.28] \times 10^{-5}$	NA

Table 4. Naive Monte Carlo Simulation

	$n = 120$	$n = 180$	$n = 240$
Theoretical p_n	1.61×10^{-3}	9.66×10^{-5}	6.35×10^{-6}
Estimate $\hat{p}_n$	1.40×10^{-3}	8.76×10^{-5}	6.01×10^{-6}
Standard Error	0.02×10^{-3}	0.18×10^{-5}	0.13×10^{-6}
95% Confidence Interval	$[1.35, 1.45] \times 10^{-3}$	$[8.41, 9.12] \times 10^{-5}$	$[5.74, 6.28] \times 10^{-6}$

Table 5. Traditional Importance Sampling Scheme

	$n = 120$	$n = 180$	$n = 240$
Theoretical p_n	1.61×10^{-3}	9.66×10^{-5}	6.35×10^{-6}
Estimate $\hat{p}_n$	1.56×10^{-3}	9.73×10^{-5}	6.29×10^{-6}
Standard Error	0.04×10^{-3}	0.15×10^{-5}	0.07×10^{-6}
95% Confidence Interval	$[1.49, 1.63] \times 10^{-3}$	$[9.44, 10.02] \times 10^{-6}$	$[6.15, 6.43] \times 10^{-6}$
$(-\log \hat{V}^n)/(-\log \hat{p}_n)$	1.72	1.87	1.93

Table 6. Adaptive Importance Sampling Scheme: Asymptotic Optimality

Remark 4.2 The theoretical value of p_n can be computed as follows. Let X_n be the number of -1 's in a trajectory, i.e. $X_n \doteq \sum_{j=0}^{n-1} 1_{\{Y_j = -1\}}$. Since $Y_0 \equiv 1$ and $Q(-1, 1) = 1$, we have $0 \leq X_n \leq n/2$ with probability one. Clearly $S_n = n - 2X_n$, whence it suffices to compute $P(X_n \geq m)$ for all non-negative integers m such that $2m \leq n$. But if we define $T_1 \doteq \inf\{j \geq 0 : Y_j = -1\}$, then $T_1 \geq 1$ and $T_1 - 1$ is geometrically distributed with parameter $(1 - p)$. Moreover, $Y_{T_1} = -1$, and $Y_{1+T_1} = 1$. Now recursively define for $i \geq 2$,

$$T_i \doteq \inf\{j \geq 1 + T_{i-1} : Y_j = -1\}.$$

Then $\{T_1 - 1, T_2 - T_1 - 2, T_3 - T_2 - 2, \dots\}$ is clearly a sequence of iid geometrically distributed random variables with parameter $(1 - p)$, and

$$P(X_n \geq m) = P(T_m \leq n) = P(T_m - 2m + 1 \leq n - 2m + 1).$$

But

$$T_m - 2m + 1 = (T_1 - 1) + (T_2 - T_1 - 2) + \dots + (T_m - T_{m-1} - 2),$$

is the sum of iid geometrically distributed random variables, whence has a negative binomial distribution with parameter m and $(1-p)$. Standard softwares such as Splus contain the cumulative distribution functions of negative binomial distributions, and can easily yield the desired probabilities.

Remark 4.3 If $\beta^*(x, t) \geq 1$, then its conjugate is $\alpha^*(x, t) = +\infty$, in the sense that

$$L(\beta^*(x, t)) = \sup_{\alpha} [\alpha \beta^*(x, t) - H(\alpha)] = \lim_{\alpha \rightarrow +\infty} [\alpha \beta^*(x, t) - H(\alpha)].$$

The change of measure corresponding to $\alpha^*(x, t) = +\infty$ (at least formally) is

$$Q_{+\infty} \doteq \lim_{\alpha \rightarrow +\infty} Q_{\alpha} = \begin{bmatrix} 1 & 0 \\ 1 & 0 \end{bmatrix}.$$

Similarly, if $\beta^*(x, t) \leq 0$, then its conjugate is $\alpha^*(x, t) = -\infty$, with the corresponding change of measure

$$Q_{-\infty} \doteq \lim_{\alpha \rightarrow -\infty} Q_{\alpha} = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}.$$

However, neither of these two probability transition kernels is suitable for the purpose of importance sampling, since the probability measure induced by the original probability transition kernel Q is *not* absolutely continuous with respect to the probability measure induced by $Q_{+\infty}$ or $Q_{-\infty}$.

To overcome this difficulty, we just take α to be a large positive or negative number whenever $\alpha^*(x, t) = +\infty$ or $\alpha^*(x, t) = -\infty$. In our numerical simulation, α is taken to be 5 if $\alpha^*(x, t) = +\infty$ and -5 if $\alpha^*(x, t) = -\infty$. The probability transition kernels corresponding to $\alpha = \pm 5$ are

$$Q_{+5} = \begin{bmatrix} 0.9999 & 0.0001 \\ 1 & 0 \end{bmatrix}, \quad Q_{-5} = \begin{bmatrix} 0.0047 & 0.9953 \\ 1 & 0 \end{bmatrix},$$

which are very close to $Q_{\pm\infty}$.

Remark 4.4 In the high dimensional case or when we consider functionals more general than probabilities, there may be no clear-cut definition of “rogue” trajectories. This is because the non-optimal behavior (from the point of view of variance reduction) of the traditional importance sampling can creep into the problems in more subtle and insidious ways. See [16] for a two dimensional example where Y is a sequence of iid normal random variables.

Appendix A: Proof of Theorem 3.1

Proof of Theorem 3.1. Proposition 2.1 and Condition 2.2 imply that

$$\lim_{n \rightarrow \infty} \frac{1}{n} \log P\{S_n/n \in A\} = - \inf_{\beta \in A} L(\beta).$$

Thanks to the discussion in Section 2.3, it suffices to show the lower bound (3.1), or

$$\liminf_n W^n \geq 2 \inf_{\beta \in A} L(\beta).$$

To this end, we extend the dynamics as in Section 3, and consider a mollified version of the original control problem. In other words, let $F : \mathbb{R}^d \rightarrow \mathbb{R}$ be an arbitrary bounded and Lipschitz continuous function, and define V_F^n, W_F^n correspondingly; see the discussion from equation (3.1) to equation (3.2).

Since V_F^n is the value function of a control problem, it satisfies the Bellman equation [4]

$$\begin{aligned} V_F^n(x, y; i) &= \inf_{\alpha \in \mathbb{R}^d} \int_{\mathcal{S}} e^{-2\langle \alpha, g(z) \rangle + 2H(\alpha)} \cdot \frac{r^2(y; \alpha)}{r^2(z; \alpha)} V_F^n \left(x + \frac{1}{n} g(z), z; i+1 \right) \\ &\quad \cdot e^{\langle \alpha, g(z) \rangle - H(\alpha)} \cdot \frac{r(z; \alpha)}{r(y; \alpha)} p(y, dz) \\ &= \inf_{\alpha \in \mathbb{R}^d} \int_{\mathcal{S}} e^{-\langle \alpha, g(z) \rangle + H(\alpha)} \cdot \frac{r(y; \alpha)}{r(z; \alpha)} V_F^n \left(x + \frac{1}{n} g(y), z; i+1 \right) p(y, dz), \end{aligned}$$

together with terminal condition $V_F^n(x, y; n) = \exp\{-2nF(x)\}$. It follows from (3.2) that

$$\begin{aligned} W_F^n(x, y; i) &= -\frac{1}{n} \log \inf_{\alpha \in \mathbb{R}^d} \int_{\mathcal{S}} e^{-\langle \alpha, g(z) \rangle + H(\alpha)} \cdot \frac{r(y; \alpha)}{r(z; \alpha)} e^{-nW_F^n(x + \frac{1}{n} g(y), z; i+1)} p(y, dz) \end{aligned} \quad (\text{A.1})$$

and that $W_F^n(x, y; n) = 2F(x)$.

The discussion in Section 3 now prompts the following definition. Fixing an arbitrary $m \in \mathbb{N}$, for $0 \leq k \leq m-1$ define recursively,

$$U_F^m(x; k) = \sup_{\alpha \in \mathbb{R}^d} \inf_{\beta \in \mathbb{R}^d} \left[U_F^m \left(x + \frac{1}{m} \beta; k+1 \right) + \frac{1}{m} (L(\beta) + \langle \alpha, \beta \rangle - H(\alpha)) \right], \quad (\text{A.2})$$

given the terminal condition

$$U_F^m(x; m) \doteq 2F(x), \quad \forall x \in \mathbb{R}^d. \quad (\text{A.3})$$

See Section 3 for the interpretation of W_F^n and U_F^m as lower values of games. The key observation is the following lemma, whose proof is deferred to Appendix C.

Lemma A.1 *For an arbitrary sequence $x^n \rightarrow x \in \mathbb{R}^d$, we have*

$$\liminf_{n \rightarrow \infty} \inf_{y \in S} W_F^n(x^n, y; \lfloor nk/m \rfloor) \geq U_F^m(x; k), \quad k = 0, 1, \dots, m.$$

Assume Lemma A.1 holds for the moment. All that remains to show is the inequality

$$\liminf_{m \rightarrow \infty} U_F^m(x; 0) \geq 2 \inf_{\beta \in \mathbb{R}^d} \{L(\beta) + F(x + \beta)\}. \quad (\text{A.4})$$

Indeed, suppose (A.4) is true. Fix an arbitrary $j \in \mathbb{N}$, and define $F_j(y) \doteq j(d(y, \bar{A}) \wedge 1)$, which is bounded and Lipschitz continuous. Since $1_A(y) \leq \exp\{-2nF_j(y)\}$, we have

$$\begin{aligned} \liminf_{n \rightarrow \infty} W^n &\geq \liminf_{n \rightarrow \infty} W_{F_j}^n(0, y_0; 0) \\ &\geq \liminf_{m \rightarrow \infty} U_{F_j}^m(0; 0) \\ &\geq 2 \inf_{\beta \in \mathbb{R}^d} [L(\beta) + F_j(\beta)]. \end{aligned}$$

Exactly as on [14, pages 10-11], a compactness argument shows that

$$\lim_{j \rightarrow \infty} \inf_{\beta \in \mathbb{R}^d} \{L(\beta) + F_j(\beta)\} = \inf_{\beta \in \bar{A}} L(\beta),$$

and we complete the proof.

Now we show inequality (A.4). The idea is to represent U_F^m as the value function of a control problem with the help of the min/max theorem. To this end, define

$$\mathcal{C} \doteq \left\{ \theta \in \mathcal{P}(\mathbb{R}^d) : \int L(\beta) \theta(d\beta) < \infty \right\}$$

and rewrite (A.2) as

$$\begin{aligned} U_F^m(x; k) &= \sup_{\alpha \in \mathbb{R}^d} \inf_{\theta \in \mathcal{C}} \left[\int U_F^m \left(x + \frac{1}{m} \beta; k+1 \right) \theta(d\beta) \right. \\ &\quad \left. + \frac{1}{m} \left(\int L(\beta) \theta(d\beta) + \left\langle \alpha, \int \beta \theta(d\beta) \right\rangle - H(\alpha) \right) \right]. \end{aligned}$$

We make the following useful observation, whose proof is deferred to Appendix C.

Lemma A.2 $U_F^m(\cdot; k)$ is bounded and Lipschitz continuous for every k . Indeed

$$\|U_F^m(x; k)\| \leq 2\|F\|_\infty, \quad \forall x \in \mathbb{R}^d, \quad k = 0, 1, \dots, m,$$

and $U_F^m(\cdot; k)$ is Lipschitz continuous with Lipschitz constant $2L_F$, where L_F is the Lipschitz constant for the mollifier F .

The next lemma is a version of min/max theorem, whose proof is almost identical to [16, Lemma 2.2] and whence omitted.

Lemma A.3 For any bounded and lower semi-continuous function $f : \mathbb{R}^d \rightarrow \mathbb{R}$, we have

$$\begin{aligned} & \sup_{\alpha \in \mathbb{R}^d} \inf_{\theta \in \mathcal{C}} \left[\int f(\beta) d\theta + \int L(\beta) d\theta + \left\langle \alpha, \int \beta d\theta \right\rangle - H(\alpha) \right] \\ &= \inf_{\theta \in \mathcal{C}} \sup_{\alpha \in \mathbb{R}^d} \left[\int f(\beta) d\theta + \int L(\beta) d\theta + \left\langle \alpha, \int \beta d\theta \right\rangle - H(\alpha) \right] \end{aligned}$$

Thanks to Lemma A.2 and Lemma A.3, we obtain

$$\begin{aligned} U_F^m(x; k) &= \inf_{\theta \in \mathcal{C}} \sup_{\alpha \in \mathbb{R}^d} \left[\int U_F^m \left(x + \frac{1}{m} \beta; k+1 \right) \theta(d\beta) \right. \\ &\quad \left. + \frac{1}{m} \left(\int L(\beta) \theta(d\beta) + \left\langle \alpha, \int \beta \theta(d\beta) \right\rangle - H(\alpha) \right) \right] \\ &= \inf_{\theta \in \mathcal{C}} \left[\int U_F^m \left(x + \frac{1}{m} \beta; k+1 \right) \theta(d\beta) \right. \\ &\quad \left. + \frac{1}{m} \left(\int L(\beta) \theta(d\beta) + L \left(\int \beta \theta(d\beta) \right) \right) \right]. \quad (\text{A.5}) \end{aligned}$$

This last display implies that U_F^m has an interpretation as the minimal cost of a stochastic control problem. To simplify the notation, we state the control problem only for the case $k = 0$. The control problem will be defined on a probability $(\tilde{\Omega}, \tilde{\mathcal{F}}, \tilde{P})$, and $\tilde{E}_x$ will denote that the initial condition of the state process is x . An admissible control is a sequence $\{\nu_j^n, j = 0, 1, \dots, m-1\}$, with each ν_j^n be a stochastic kernel on $\mathbb{R}^d$ given $\mathbb{R}^d$. Given an admissible control sequence, the state dynamics are defined by $\tilde{S}_0^m = mx$ and

$$\tilde{S}_{j+1}^m \doteq \tilde{S}_j^m + \tilde{Y}_j^m,$$

where

$$\tilde{P} \left\{ \tilde{Y}_j^m \in dy \mid \tilde{Y}_i^m, 0 \leq i < j \right\} = \tilde{P} \left\{ \tilde{Y}_j^m \in dy \mid \tilde{S}_j^m/m \right\} = \nu_j^n(dy \mid \tilde{S}_j^m/m).$$

We then define the value function

$$\begin{aligned} & \tilde{v}_F^m(x; 0) \\ & \doteq \inf_{\{\nu_j^m\}} \tilde{E}_x \left[\sum_{j=0}^{m-1} \frac{1}{m} \left[\int L(y) \nu_j^m(dy) + L \left(\int y \nu_j^m(dy) \right) \right] + 2F(\tilde{S}_m^m/m) \right], \end{aligned}$$

where the infimum is taken over all controls $\{\nu_j^m\}$ and resulting controlled processes $\{\tilde{S}_j^m/m\}$ that start at x at time 0. Since $\tilde{v}_F^m$ also satisfies the DPE (A.5) [4, Chapter 8] and terminal condition $\tilde{v}_F^m(x; m) = U_F^m(x; m) = 2F(x)$, we obtain by induction that $U_F^m(x; k) = \tilde{v}_F^m(x; k)$ for all $x \in \mathbb{R}^d$ and $k \in \{0, \dots, m\}$.

Define a stochastic kernel ν^m on $\mathbb{R}^d$ given $[0, 1]$ by

$$\nu^m(dy|t) \doteq \begin{cases} \nu_j^m(dy) & \text{if } t \in [j/m, (j+1)/m), \ j = 0, 1, \dots, m-2 \\ \nu_{m-1}^m(dy) & \text{if } t \in [(m-1)/m, 1] \end{cases}.$$

Let λ denote Lebesgue measure. Then the definition of $\nu^m(dy|t)$ and the convexity of L imply that

$$\begin{aligned} & U_F^m(x; 0) \\ & = \inf_{\{\nu_j^m\}} \tilde{E}_x \left[\int_0^1 \int_{\mathbb{R}^d} L(y) \nu^m(dy|t) dt + \sum_{j=0}^{m-1} \frac{1}{m} L \left(\int_{\mathbb{R}^d} y \nu_j^m(dy) \right) + 2F(\tilde{S}_m^m/m) \right] \\ & \geq \inf_{\{\nu_j^m\}} \tilde{E}_x \left[\int_0^1 \int_{\mathbb{R}^d} L(y) \nu^m(dy|t) dt + L \left(\sum_{j=0}^{m-1} \frac{1}{m} \int_{\mathbb{R}^d} y \nu_j^m(dy) \right) + 2F(\tilde{S}_m^m/m) \right] \\ & = \inf_{\{\nu_j^m\}} \tilde{E}_x \left[\int_{\mathbb{R}^d \times [0, 1]} L(y) \nu^m(dy \times dt) + L \left(\int_{\mathbb{R}^d \times [0, 1]} y \nu^m(dy \times dt) \right) \right. \\ & \quad \left. + 2F(\tilde{S}_m^m/m) \right], \end{aligned}$$

where $\nu^m(dy \times dt) \doteq \nu^m(dy|t)dt$. A straightforward weak convergence approach will be adopted to derive the desired inequality (A.4). Since the proof is essentially the same as [14, Theorem 5.3.5], we only give a sketch.

For each $\varepsilon > 0$, there exist a sequence of controls $\{\nu^m, m \in \mathbb{N}\}$ such that, for every m , we have

$$U_F^m(x; 0) + \varepsilon \geq \tilde{E}_x \left[\int_{\mathbb{R}^d \times [0, 1]} L(y) d\nu^m + L \left(\int_{\mathbb{R}^d \times [0, 1]} y d\nu^m \right) + 2F(\tilde{S}_m^m/m) \right].$$

Furthermore, since L is non-negative and F is bounded, we have

$$\sup_{m \in \mathbb{N}} \tilde{E}_x \int_{\mathbb{R}^d \times [0,1]} L(y) \nu^m(dy \times dt) < \infty$$

However, since function L is superlinear (Proposition 2.1), it is not difficult to check that $\{\nu^m\}$ is uniformly integrable in the sense that

$$\lim_{C \rightarrow \infty} \sup_{m \in \mathbb{N}} \tilde{E}_x \int_{\{y: \|y\| > C\} \times [0,1]} \|y\| \nu^m(dy \times dt) = 0.$$

It follows from that proof of [14, Proposition 5.3.2] that $\{\nu^m\}$ is indeed tight. Therefore we can extract a weakly convergent sub-subsequence, still denoted by $\{\nu^m\}$, such that $\nu^m \Rightarrow \nu$ for some stochastic kernel ν whose second marginal is Lebesgue measure [14, Lemma 5.3.4]. We utilize the Skorokhod representation [6], which allows us to assume (when calculating the limits of the integrals) that the convergence is actually w.p.1. It follows from the uniform integrability of $\{\nu^m\}$ and the proof of [14, Proposition 5.3.5] that

$$\int_{\mathbb{R}^d \times [0,1]} y \nu^m(dy \times dt) \xrightarrow{P} \int_{\mathbb{R}^d \times [0,1]} y \nu(dy \times dt)$$

and

$$\tilde{S}_m^m/m \xrightarrow{P} Z \doteq x + \int_{\mathbb{R}^d \times [0,1]} y \nu(dy \times dt).$$

Furthermore, it follows from the lower semicontinuity and non-negativity of L [14, Lemma A.3.12] that, with probability one,

$$\liminf_m \int_{\mathbb{R}^d \times [0,1]} L(y) \nu^m(dy \times dt) \geq \int_{\mathbb{R}^d \times [0,1]} L(y) \nu(dy \times dt).$$

Thanks to convexity of L and Jensen's inequality, we have

$$\int_{\mathbb{R}^d \times [0,1]} L(y) \nu(dy \times dt) \geq L \left(\int_{\mathbb{R}^d \times [0,1]} y \nu(dy \times dt) \right)$$

By Fatou's Lemma and the lower-semicontinuity of L [25], we have

$$\liminf_n U_F^m(x; 0) + \varepsilon \geq \tilde{E}_x \left[2L \left(\int_{\mathbb{R}^d \times [0,1]} y \nu(dy \times dt) \right) + 2F(Z) \right].$$

It is now trivial that the right hand side of the last inequality is bounded below by

$$2 \inf_{\beta \in \mathbb{R}^d} [L(\beta) + F(x + \beta)].$$

Since $\varepsilon > 0$ is arbitrary, (A.4) follows readily, which completes the proof. ■

Appendix B: A large deviation upper bound

In this section, we study a uniform large deviation principle upper bound, which is essential for proving the key Lemma A.1. We present a proof based on the weak convergence approach [14]. Alternatively, one can adapt the methodology in [13].

The following two lemmas will be useful.

Lemma B.1 *Suppose $\mathcal{S}$ is a Polish space and $\mathcal{P}(\mathcal{S})$ is the space of probability measures on $\mathcal{S}$ endowed with the weak convergence topology. Consider a sequence of random variables $\mu_n : (\Omega^n, \mathcal{F}^n, P^n) \rightarrow \mathcal{P}(\mathcal{S})$. In other words, $\{\mu^n\}$ is a sequence of random probability measures. Then $\{\mu^n\}$ is tight if and only if the sequence $\{E^n \mu^n\}$ is tight. Here $E^n \mu^n \in \mathcal{P}(\mathcal{S})$ is defined by*

$$(E^n \mu^n)(A) \doteq \int_{\Omega^n} \mu^n(\omega)(A) P^n(d\omega)$$

for every Borel set A in $\mathcal{P}(\mathcal{S})$.

Proof. See [22, Theorem 6.1, Chapter 1]. ■

Lemma B.2 *Suppose $\mathcal{S}$ is a Polish space, $\{\mu^n\} \subset \mathcal{P}(\mathcal{S})$, and $p(\cdot, \cdot)$ a probability transition kernel. If $\mu^n \rightarrow \mu$ in the τ -topology for some $\mu \in \mathcal{P}(\mathcal{S})$, then*

$$\mu^n \otimes p \rightarrow \mu \otimes p$$

in the τ -topology. Here $\mu \otimes p$ denotes the probability measure on $\mathcal{S} \times \mathcal{S}$ given by

$$(\mu \otimes p)(B) \doteq \int_B \mu(dx) p(x, dy)$$

for every Borel set $B \subseteq \mathcal{S} \times \mathcal{S}$.

Proof. It suffices to show that

$$\int_{\mathcal{S} \times \mathcal{S}} f(x, y) \mu^n(dx) p(x, dy) \rightarrow \int_{\mathcal{S} \times \mathcal{S}} f(x, y) \mu(dx) p(x, dy)$$

for every bounded, measurable function f . Since $\mu^n \rightarrow \mu$ in the τ -topology, it remains to show that

$$\int_{\mathcal{S}} f(x, y) p(x, dy)$$

is a bounded and measurable function (over x). The boundedness is trivial, and the measurability follows from Fubini's Theorem; c.f. [5, Exercise 18.20].

■

Proposition B.3 Suppose $Y = \{Y_j, j \in \mathbb{N}\}$ is a Markov chain that takes values in a Polish space $\mathcal{S}$. Let p denote the probability transition kernel of Y , and assume Condition 2.1 holds. Suppose $g : \mathcal{S} \rightarrow \mathbb{R}^d$ is a bounded measurable function, and define H and L as in (2.1) through (2.3). Then for any fixed $\alpha \in \mathbb{R}^d$, bounded and continuous function $f : \mathbb{R}^d \rightarrow \mathbb{R}$, and sequence $x^n \rightarrow x \in \mathbb{R}^d$, we have

$$\liminf_{n \rightarrow \infty} \inf_{y \in \mathcal{S}} -\frac{1}{n} \log E_y \left[e^{-\langle \alpha, \sum_{j=0}^{n-1} g(Y_j) \rangle} e^{-nf\left(x^n + \frac{1}{n} \sum_{j=0}^{n-1} g(Y_j)\right)} \right] \geq I_f(x),$$

where

$$I_f(x) \doteq \inf_{\beta} [f(x + \beta) + L(\beta) + \langle \alpha, \beta \rangle].$$

Proof. Let

$$\begin{aligned} v^n(x, y) &\doteq -\frac{1}{n} \log E_y \left[e^{-\langle \alpha, \sum_{j=0}^{n-1} g(Y_j) \rangle} e^{-nf\left(x + \frac{1}{n} \sum_{j=0}^{n-1} g(Y_j)\right)} \right] \\ &= -\frac{1}{n} \log \int e^{-\langle \alpha, \sum_{j=0}^{n-1} g(y_j) \rangle} e^{-nf\left(x + \frac{1}{n} \sum_{j=0}^{n-1} g(y_j)\right)} d\pi_y^n. \end{aligned}$$

Here π_y^n is the joint distribution of $(Y_0, Y_1, \dots, Y_{n-1})$, or

$$\pi_y^n(dy_0, dy_1, \dots, dy_n) \doteq \delta_y(dy_0) p(y_0, dy_1) p(y_1, dy_2) \cdots p(y_{n-1}, dy_n).$$

Clearly v^n is bounded, thanks to the boundedness of g and f . It suffices to show that for every sequence $x^n \rightarrow x$ and $\{y^n\} \subseteq \mathcal{S}$,

$$\liminf_{n \rightarrow \infty} v^n(x^n, y^n) \geq I_f(x). \quad (\text{B.1})$$

For an arbitrary $\varepsilon > 0$, the relative entropy representation of exponential integrals (3.3) [14, Proposition 1.4.2] yields the existence of a probability measure μ^n on $\mathcal{S}^{n+1}$ such that

$$\begin{aligned} v^n(x^n, y^n) + \varepsilon & \\ &\geq \frac{1}{n} R(\mu^n \| \pi_{y^n}^n) + \left\langle \alpha, \int \frac{1}{n} \sum_{j=0}^{n-1} g(y_j) d\mu^n \right\rangle + \int f\left(x^n + \frac{1}{n} \sum_{j=0}^{n-1} g(y_j)\right) d\mu^n. \end{aligned} \quad (\text{B.2})$$

In particular, it is not hard to see that

$$\sup_{n \in \mathbb{N}} \frac{1}{n} R(\mu^n \| \pi_{y^n}^n) < \infty. \quad (\text{B.3})$$

We can factor μ^n as

$$\mu^n(dy_0, dy_1, \dots, dy_n) = \mu_0^n(dy_0) \mu_1^n(dy_1|y_0) \cdots \mu_n^n(dy_n|y_{n-1}, y_{n-2}, \dots, y_0).$$

Now consider a probability space $(\tilde{\Omega}, \tilde{\mathcal{F}}, \tilde{P})$, on which we define a stochastic process given by

$$\begin{aligned} \tilde{P}(\tilde{Y}_0^n \in dy) &= \mu_0^n(dy_0) \\ \tilde{P}(\tilde{Y}_{j+1}^n \in dy | \tilde{Y}_i^n, i = 0, 1, \dots, j) &= \mu_{j+1}^n(dy | \tilde{Y}_i^n, i = 0, 1, \dots, j) \end{aligned}$$

for $j = 0, 1, \dots, n-1$. To ease exposition, let

$$\bar{\mu}_{j+1}^n(dy) \doteq \mu_{j+1}^n(dy | \tilde{Y}_i^n, i = 0, 1, \dots, j),$$

which is a random probability measure on $\mathcal{S}$. Also define a random probability measure on $\mathcal{S} \times \mathcal{S}$ by

$$\gamma^n(dx \times dy) = \frac{1}{n} \sum_{j=0}^{n-1} \delta_{\tilde{Y}_j^n}(dx) \times \bar{\mu}_{j+1}^n(dy),$$

whose marginals are

$$(\gamma^n)_1 = \frac{1}{n} \sum_{j=0}^{n-1} \delta_{\tilde{Y}_j^n} \doteq \tilde{L}^n, \quad (\gamma^n)_2 = \frac{1}{n} \sum_{j=0}^{n-1} \bar{\mu}_{j+1}^n.$$

Thanks to the chain rule [14, Theorem B.2.1], we have

$$\frac{1}{n} R(\mu^n \| \pi_{y^n}^n) = \frac{1}{n} \tilde{E} \left[R(\mu_0^n \| \delta_{y^n}) + \sum_{j=0}^{n-1} R(\bar{\mu}_{j+1}^n(\cdot) \| p(\tilde{Y}_j, \cdot)) \right]. \quad (\text{B.4})$$

However,

$$\begin{aligned} & \frac{1}{n} \sum_{j=0}^{n-1} R(\bar{\mu}_{j+1}^n(\cdot) \| p(\tilde{Y}_j^n, \cdot)) \\ &= \frac{1}{n} \sum_{j=0}^{n-1} R(\delta_{\tilde{Y}_j^n}(dy) \times \bar{\mu}_{j+1}^n(dz) \| \delta_{\tilde{Y}_j^n}(dy) \times p(\tilde{Y}_j^n, dz)) \\ &= \frac{1}{n} \sum_{j=0}^{n-1} R(\delta_{\tilde{Y}_j^n}(dy) \times \bar{\mu}_{j+1}^n(dz) \| \delta_{\tilde{Y}_j^n}(dy) \otimes p(y, dz)) \\ &\geq R \left(\frac{1}{n} \sum_{j=0}^{n-1} \delta_{\tilde{Y}_j^n}(dy) \times \bar{\mu}_{j+1}^n(dz) \left\| \frac{1}{n} \sum_{j=0}^{n-1} \delta_{\tilde{Y}_j^n}(dy) \otimes p(y, dz) \right\| \right) \\ &= R(\gamma^n \| \tilde{L}^n \otimes p), \end{aligned}$$

where the inequality follows from the convexity of relative entropy $R(\cdot\|\cdot)$. Thanks to (B.2), and observing that

$$\begin{aligned}\int \frac{1}{n} \sum_{j=0}^{n-1} g(y_j) d\mu^n &= \tilde{E} \int g d\tilde{L}^n \\ \int f \left(x^n + \frac{1}{n} \sum_{j=0}^{n-1} g(y_j) \right) d\mu^n &= \tilde{E} f \left(x^n + \int g d\tilde{L}^n \right)\end{aligned}$$

we arrive at

$$v^n(x^n, y^n) + \varepsilon \geq \tilde{E} \left[R(\gamma^n \|\tilde{L}^n \otimes p) + \left\langle \alpha, \int g d\tilde{L}^n \right\rangle + f \left(x^n + \int g d\tilde{L}^n \right) \right].$$

It suffices to show $\{\gamma^n\}$ is tight. Indeed, if this is true, the same argument as in [14, Theorem 8.2.8] allows us to extract a weak convergent subsequence of $(\gamma^n, \tilde{L}^n)$, still indexed by n , such that

$$(\gamma^n, \tilde{L}^n) \Rightarrow (\gamma, \tilde{L})$$

for some stochastic kernel γ on $\mathcal{S} \times \mathcal{S}$ and some stochastic kernel on $\mathcal{S}$, and a (random) transition probability function q such that

$$\gamma(dy \times dz) = \tilde{L}(dy) \otimes q(y, dz)$$

and

$$\tilde{L}q = \tilde{L} \tag{B.5}$$

hold almost surely. In particular, we have

$$(\gamma^n)_2 \Rightarrow (\gamma)_2 = \tilde{L}q = \tilde{L}$$

Note equation (B.5) says that $\tilde{L}$ is indeed the invariant measure for the transition probability function q . Also observe that

$$\sup_{n \in \mathbb{N}} \tilde{E} R(\gamma^n \|\tilde{L}^n \otimes p) < \infty.$$

This implies the existence of a subsequence, still indexed by n , such that

$$\tilde{L}^n \rightarrow \tilde{L}, \quad (\gamma^n)_2 \rightarrow \tilde{L}$$

in the τ -topology; see the proof of [14, Lemma 9.3.3]. Therefore,

$$\int g d\tilde{L}^n \rightarrow \int g d\tilde{L}$$

almost surely. Furthermore, thanks to Lemma B.2, $\tilde{L}^n \otimes p \rightarrow \tilde{L} \otimes p$ in the τ -topology (hence in the weak-topology) almost surely. The lower semi-continuity of $R(\cdot\|\cdot)$ implies

$$\liminf_{n \rightarrow \infty} R(\gamma^n \|\tilde{L}^n \otimes p) \geq R(\gamma \|\tilde{L} \otimes p).$$

It follows readily from Fatou's Lemma that

$$\begin{aligned} & \liminf_{n \rightarrow \infty} v^n(x^n, y^n) + \varepsilon \\ & \geq \tilde{E} \left[R(\gamma \|\tilde{L} \otimes p) + \left\langle \alpha, \int g d\tilde{L} \right\rangle + f \left(x + \int g d\tilde{L} \right) \right] \\ & = \tilde{E} \left[R(\tilde{L} \otimes q \|\tilde{L} \otimes p) + \left\langle \alpha, \int g d\tilde{L} \right\rangle + f \left(x + \int g d\tilde{L} \right) \right] \\ & \geq \inf_{\{\mu q = \mu\}} \left[R(\mu \otimes q \|\mu \otimes p) + \left\langle \alpha, \int g d\mu \right\rangle + f \left(x + \int g d\mu \right) \right]. \end{aligned}$$

Recalling equation (2.4) and letting $\varepsilon \rightarrow 0$, we obtain

$$\liminf_{n \rightarrow \infty} v^n(x^n, y^n) \geq \inf_{\beta} [L(\beta) + \langle \alpha, \beta \rangle + f(x + \beta)].$$

which is the desired inequality (B.1).

It remains to show the tightness of $\{\gamma^n\}$. All we need is the tightness of the two marginals, $\{(\gamma^n)_1\}$ and $\{(\gamma^n)_2\}$. However, it is not difficult to observe that

$$\tilde{E}(\gamma^n)_1 = \tilde{E}\tilde{L}^n = \frac{1}{n} \sum_{j=0}^{n-1} \tilde{E}\delta_{\tilde{Y}_j^n} = \frac{1}{n} \sum_{j=0}^{n-1} \mu^{n,j},$$

where $\mu^{n,j}$ denotes the j -th marginal of the probability μ^n , and similarly

$$\tilde{E}(\gamma^n)_2 = \frac{1}{n} \sum_{j=0}^{n-1} \tilde{E}\bar{\mu}_{j+1}^n = \frac{1}{n} \sum_{j=0}^{n-1} \mu^{n,j+1}.$$

Letting $\|\cdot\|_v$ denote the total variation metric, we have

$$\|\tilde{E}(\gamma^n)_1 - \tilde{E}(\gamma^n)_2\|_v = \frac{1}{n} \|\mu^{n,0} - \mu^{n,n}\|_v \leq \frac{2}{n}. \quad (\text{B.6})$$

If we can show $\{(\gamma^n)_2\}$ is tight, then Lemma B.1 implies $\{\tilde{E}(\gamma^n)_2\}$ is tight, which in turns yields the tightness of $\{\tilde{E}(\gamma^n)_1\}$, thanks to (B.6). Applying Lemma B.1 once again, we have the tightness of $\{(\gamma^n)_1\}$. Therefore, it is

sufficient to show that $\{(\gamma^n)_2\}$ is tight. The proof will distinguish two cases: $m_0 = 1$ and $m_0 > 1$.

Suppose that $m_0 = 1$. Note that the non-negativity of relative entropy, (B.3), and (B.4) imply

$$\sup_{n \in \mathbb{N}} \tilde{E} \frac{1}{n} \sum_{j=0}^{n-1} R(\bar{\mu}_{j+1}^n(\cdot) \| p(\tilde{Y}_j^n, \cdot)) < \infty$$

It follows from the assumption of uniform recurrency

$$a\nu_p(\cdot) \leq p(y, \cdot) \leq b\nu_p(\cdot), \quad \forall y \in \mathcal{S},$$

that

$$R(\bar{\mu}_{j+1}^n(\cdot) \| p(\tilde{Y}_j^n, \cdot)) \geq cR(\bar{\mu}_{j+1}^n \| \nu_p)$$

for some constant $c > 0$. It is now easy to derive from the convexity of relative entropy that

$$\sup_{n \in \mathbb{N}} \tilde{E} R((\gamma^n)_2 \| \nu_p) \leq \sup_{n \in \mathbb{N}} \tilde{E} \frac{1}{n} \sum_{j=0}^{n-1} R(\bar{\mu}_{j+1}^n \| \nu_p) < \infty,$$

which further implies the tightness of $\{(\gamma^n)_2\}$ since $R(\cdot \| \nu_p)$ is a tightness function on $\mathcal{P}(\mathcal{S})$. Note that

$$\tilde{E}(\gamma^n)_2 = \frac{1}{n} \sum_{j=0}^{n-1} \mu^{n,j+1}$$

is also tight, thanks to Lemma B.1.

The general case with $m_0 > 1$ is slightly more complicated. We will give a proof with $m_0 = 2$, and observe that the proof for $m_0 > 2$ is essentially the same and thus omitted. Without loss of generality, we show $\{(\gamma^n)_2 : n \text{ even}\}$ to be tight. The tightness for $\{(\gamma^n)_2 : n \text{ odd}\}$ is similar.

To ease notation, let $\pi^n \doteq \pi_{y^n}^n$, and $\pi^{n,e}$ be the marginal distribution of π^n over even coordinates; i.e.

$$\pi^{n,e}(dy_0, dy_2, \dots, dy_{n-2}, dy_n) = \delta_y(dy_0) p^{(2)}(y_0, dy_2) \cdots p^{(2)}(y_{n-2}, dy_n).$$

One can similarly define $\mu^{n,e}$, or

$$\mu^{n,e}(dy_0, dy_2, \dots, dy_n) = \mu_0^n(dy_0) \mu_2^{n,e}(dy_2 | y_0) \cdots \mu_n^{n,e}(dy_n | y_{n-2}, \dots, y_2, y_0).$$

Thanks to the chain rule and non-negativity of the relative entropy, we have

$$R(\mu^{n,e} \| \pi^{n,e}) \leq R(\mu^n \| \pi^n),$$

and thus $\sup_n \frac{1}{n} R(\mu^{n,e} \|\pi^{n,e}) < \infty$. With the same proof as for the case $m_0 = 1$, we have that

$$\frac{2}{n} \sum_{j=0}^{\frac{n}{2}-1} \mu^{n,e,j+1}$$

is tight; here $\mu^{n,e,j}$ is the j -th marginal of $\mu^{n,e}$; i.e.

$$\mu^{n,e,j}(dy_{2j}) = \mu^{n,e}(\mathcal{S}, \dots, \mathcal{S}, dy_{2j}, \mathcal{S}, \dots, \mathcal{S}).$$

One can similarly define $\mu^{n,o}$ as the marginal distribution of μ^n over odd coordinates, and the same argument can be carried over to prove the tightness of

$$\frac{2}{n} \sum_{j=0}^{\frac{n}{2}-1} \mu^{n,o,j+1}.$$

However, observe that

$$\mu^{n,e,j} = \mu^{n,2j}, \quad \mu^{n,o,j} = \mu^{n,2j+1}.$$

We have

$$\frac{1}{2} \left(\frac{2}{n} \sum_{j=0}^{\frac{n}{2}-1} \mu^{n,e,j+1} + \frac{2}{n} \sum_{j=0}^{\frac{n}{2}-1} \mu^{n,o,j+1} \right) = \frac{1}{2} \sum_{j=0}^{n-1} \mu^{n,j+1} = \tilde{E}(\gamma^n)_2.$$

This implies the tightness of $\{\tilde{E}(\gamma^n)_2 : n \text{ even}\}$, which is equivalent to the tightness of $\{(\gamma^n)_2 : n \text{ even}\}$, thanks to Lemma B.1. ■

Appendix C: Proof of Lemma A.1 and Lemma A.2

Proof of Lemma A.1. The proof is by induction. For $k = m$, we have $\lfloor nk/m \rfloor = n$. By definition,

$$\liminf_{n \rightarrow \infty} \inf_{y \in \mathcal{S}} W_F^n(x^n, y; n) = \liminf_{n \rightarrow \infty} \inf_{y \in \mathcal{S}} 2F(x^n) = \liminf_{n \rightarrow \infty} 2F(x^n),$$

and Lemma A.1 follows trivially from the continuity of F .

Assume now the claim holds for $k+1$. Let $\ell(n) \doteq \lfloor n(k+1)/m \rfloor - \lfloor nk/m \rfloor$. Also, let π_y^j be the probability measure on $\mathcal{S}^{j+1}$ defined by

$$\pi_y^j(dy_0, dy_1, \dots, dy_j) \doteq \delta_y(dy_0) p(y_0, dy_1) p(y_1, dy_2) \cdots p(y_{j-1}, dy_j)$$

for every $y \in \mathcal{S}$ and every $j \in \mathbb{N}$.

For an arbitrary $\alpha \in \mathbb{R}^d$, let

$$U_{\alpha,F}^m(x; k) \doteq \inf_{\beta \in \mathbb{R}^d} \left[U_F^m \left(x + \frac{1}{m} \beta; k+1 \right) + \frac{1}{m} (L(\beta) + \langle \alpha, \beta \rangle - H(\alpha)) \right].$$

It follows from the definition that $U_F^m(x; k) = \sup_{\alpha} U_{\alpha,F}^m(x; k)$. Therefore, all we need to show is that, for every $\alpha \in \mathbb{R}^d$ and any sequence $x^n \rightarrow x$,

$$\liminf_{n \rightarrow \infty} \inf_{y \in \mathcal{S}} W_F^n(x^n, y; \lfloor nk/m \rfloor) \geq U_{\alpha,F}^m(x; k).$$

However, for an arbitrary fixed $\alpha \in \mathbb{R}^d$, the dynamic programming principle implies that

$$\begin{aligned} W_F^n(x, y; \lfloor nk/m \rfloor) &\geq -\frac{1}{n} \log \int e^{-\langle \alpha, \sum_{j=1}^{\ell(n)} g(y_j) \rangle + \ell(n) H(\alpha)} \cdot \prod_{j=1}^{\ell(n)} \frac{r(y_{j-1}; \alpha)}{r(y_j; \alpha)} \\ &\quad \cdot e^{-n W_F^n \left(x + \frac{1}{n} \sum_{j=0}^{\ell(n)-1} g(y_j), y_{\ell(n)}; \lfloor n(k+1)/m \rfloor \right)} d\pi_y^{\ell(n)} \\ &= -\frac{1}{n} \log \int e^{-\langle \alpha, \sum_{j=1}^{\ell(n)} g(y_j) \rangle + \ell(n) H(\alpha)} \cdot \frac{r(y_0; \alpha)}{r(y_{\ell(n)}; \alpha)} \\ &\quad \cdot e^{-n W_F^n \left(x + \frac{1}{n} \sum_{j=0}^{\ell(n)-1} g(y_j), y_{\ell(n)}; \lfloor n(k+1)/m \rfloor \right)} d\pi_y^{\ell(n)}. \end{aligned}$$

Since g is bounded and $r(\cdot; \alpha)$ is both bounded from above and bounded away from zero by (2.2), it suffices to show

$$\liminf_{n \rightarrow \infty} \inf_{y \in \mathcal{S}} \bar{v}_F^n(x^n, y; 0) \geq U_{\alpha,F}^m(x; k), \quad (\text{C.1})$$

where

$$\begin{aligned} \bar{v}_F^n(x, y; 0) &\doteq -\frac{1}{n} \log \int e^{-\langle \alpha, \sum_{j=0}^{\ell(n)-1} g(y_j) \rangle + \ell(n) H(\alpha)} \\ &\quad \cdot e^{-n W_F^n \left(x + \frac{1}{n} \sum_{j=0}^{\ell(n)-1} g(y_j), y_{\ell(n)}; \lfloor n(k+1)/m \rfloor \right)} d\pi_y^{\ell(n)}. \end{aligned}$$

We claim that inequality (C.1) is a direct consequence of

$$\liminf_{n \rightarrow \infty} \inf_{y \in \mathcal{S}} v_F^n(x^n, y; 0) \geq U_{\alpha,F}^m(x; k), \quad (\text{C.2})$$

where

$$\begin{aligned} v_F^n(x, y; 0) &\doteq -\frac{1}{n} \log \int e^{-\langle \alpha, \sum_{j=0}^{\ell(n)-1} g(y_j) \rangle + \ell(n) H(\alpha)} \\ &\quad \cdot e^{-n U_F^m \left(x + \frac{1}{n} \sum_{j=0}^{\ell(n)-1} g(y_j); k+1 \right)} d\pi_y^{\ell(n)}. \end{aligned}$$

Indeed, since $\ell(n) \leq n$, one can always find a compact set $K \subseteq \mathbb{R}^d$ such that

$$x^n + \frac{1}{n} \sum_{j=0}^{\ell(n)-1} g(y_j) \in K, \quad \forall (y_0, y_1, \dots, y_{\ell(n)}), \forall n \in \mathbb{N},$$

thanks to the boundedness of g and the assumption $x^n \rightarrow x$. It is also not hard to show by contradiction from the induction hypothesis and the continuity of U_F^m (Lemma A.2) that, for any $\varepsilon > 0$, there exists $N(\varepsilon) \in \mathbb{N}$ such that for all $x \in K$ and $n \geq N(\varepsilon)$,

$$\inf_{y \in \mathcal{S}} W_F^n(x, y; \lfloor n(k+1)/m \rfloor) - U_F^m(x; k+1) \geq -\varepsilon.$$

We arrive at

$$\liminf_{n \rightarrow \infty} \inf_{y \in \mathcal{S}} \bar{v}_F^n(x^n, y; 0) \geq \liminf_{n \rightarrow \infty} \inf_{y \in \mathcal{S}} v_F^n(x^n, y; 0) - \varepsilon$$

for every $\varepsilon > 0$. It follows that (C.1) is implied by (C.2).

It remains to show (C.2), which is an easy consequence of the uniform large deviation bound Proposition B.3, Lemma A.2, boundedness of g , and that

$$\left| \frac{\ell(n)}{n} - \frac{1}{m} \right| \rightarrow 0$$

This completes the proof. ■

Proof of Lemma A.2. That $U_F^m(\cdot; k)$ is Lipschitz continuous with Lipschitz constant $2L_F$ follows trivially by induction, and the terminal condition (A.3).

As for the boundedness of $U_F^m(\cdot; k)$, we first show it is bounded from below. Since $H(0) = 0$ and L is non-negative, definition (A.2) gives

$$\begin{aligned} U_F^m(x; k) &\geq \inf_{\beta \in \mathbb{R}^d} \left[U_F^m \left(x + \frac{1}{m} \beta; k+1 \right) + \frac{1}{m} L(\beta) \right] \\ &\geq \inf_{\beta \in \mathbb{R}^d} U_F^m \left(x + \frac{1}{m} \beta; k+1 \right) \\ &= \inf_{z \in \mathbb{R}^d} U_F^m(z; k+1) \end{aligned}$$

for every x . It follows that, for every k ,

$$\inf_{x \in \mathbb{R}^d} U_F^m(x; k) \geq \inf_{x \in \mathbb{R}^d} U_F^m(x; k+1) \geq \dots \geq \inf_{x \in \mathbb{R}^d} U_F^m(x; m) \geq -2\|F\|_\infty.$$

It remains to show that U_F^m is bounded from above. Let $\bar{\beta}$ be a subdifferential of the convex function H at $\alpha = 0$. Then

$$L(\bar{\beta}) = \sup_{\alpha \in \mathbb{R}^d} [\langle \alpha, \bar{\beta} \rangle - H(\alpha)] = 0$$

and the supremum is achieved at $\alpha = 0$. By definition (A.2) again, we have

$$\begin{aligned} U_F^m(x; k) &\leq \sup_{\alpha \in \mathbb{R}^d} \left[U_F^m \left(x + \frac{1}{m} \bar{\beta}; k+1 \right) + \frac{1}{m} (\langle \alpha, \bar{\beta} \rangle - H(\alpha)) \right] \\ &= U_F^m \left(x + \frac{1}{m} \bar{\beta}; k+1 \right) \\ &\leq \sup_{z \in \mathbb{R}^d} U_F^m(z; k+1) \end{aligned}$$

for every x . It follows that, for every k ,

$$\sup_{x \in \mathbb{R}^d} U_F^m(x; k) \leq \sup_{x \in \mathbb{R}^d} U_F^m(x; k+1) \leq \cdots \leq \sup_{x \in \mathbb{R}^d} U_F^m(x; m) \leq 2\|F\|_\infty.$$

This completes the proof. ■

References

- [1] S. Asmussen. Conjugate processes and the simulation of ruin probability. *Stochastic Process. Appl.*, 20:213–229, 1985.
- [2] S. Asmussen. Risk theory in a Markovian environment. *Scand. Actuarial J.*, pages 69–100, 1989.
- [3] S. Asmussen, R. Rubinstein, and C.L. Wang. Regenerative rare event simulation via likelihood ratios. *J. Appl. Probab.*, 31:797–815, 1994.
- [4] D. Bertsekas and S. Shreve. *Stochastic Optimal Control: The Discrete Time Case*. Academic Press, San Diego, California, 1978.
- [5] P. Billingsley. *Probability and Measure*. John Wiley, New York, third edition, 1995.
- [6] L. Breiman. *Probability Theory*. Addison-Wesley, Reading, Mass., 1968.
- [7] J.A. Bucklew. *Large Deviations Techniques in Decision, Simulation and Estimation*. Wiley, New York, 1990.

- [8] J.A. Bucklew. The blind simulation problem and regenerative processes. *IEEE Trans. Inform. Theory*, 44:2877–2891, 1998.
- [9] J.A. Bucklew, P. Ney, and J.S. Sadowsky. Monte Carlo simulation and large deviations theory for uniformly recurrent Markov chains. *J. Appl. Probab.*, 27:44–59, 1990.
- [10] J. Chen, D. Lu, J.S. Sadowsky, and K. Yao. On importance sampling in digital communications. Part I: Fundamentals. *IEEE J. Selected Areas in. Comm.*, 11:289–307, 1993.
- [11] J.F. Collamore. Importance sampling techniques for the multidimensional ruin problem for general Markov additive sequences of random vectors. *Ann. Appl. Prob.*, 12:382–421, 2002.
- [12] M. Cottrel, J.C. Fort, and G. Malgouyres. Large deviations and rare events in the study of stochastic algorithms. *IEEE Trans Automat. Control*, AC-28:907–920, 1983.
- [13] A. de Acosta. Large deviations for empirical measures of Markov chains. *Journal of Theoretical Probability*, 3, 1990.
- [14] P. Dupuis and R. S. Ellis. *A Weak Convergence Approach to the Theory of Large Deviations*. John Wiley & Sons, New York, 1997.
- [15] P. Dupuis and H.J. Kushner. Stochastic systems with small noise, analysis and simulation; a phase locked loop example. *SIAM J. Appl. Math.*, 47:643–661, 1987.
- [16] P. Dupuis and H. Wang. Importance sampling, large deviations, and differential games. *Annals of Probab.*, 2002. Submitted.
- [17] P. Glasserman and S. Kou. Analysis of an importance sampling estimator for tandem queues. *ACM Trans. Modeling Comp. Simulation*, 4:22–42, 1995.
- [18] P. Glasserman and Y. Wang. Counter examples in importance sampling for large deviations probabilities. *Ann. Appl. Prob.*, 7:731–746, 1997.
- [19] P.W. Glynn. Large deviations for the infinite server queue in heavy traffic. *IMA Vol. Math. Appl.*, 71:387–394, 1995.
- [20] P. Heidelberger. Fast simulation of rare events in queueing and reliability models. *ACM Trans. Modeling Comp. Simulation*, 4:43–85, 1995.

- [21] I. Iscoe, P. Ney, and E. Nummelin. Large deviations of uniformly recurrent Markov additive processes. *Adv. Appl. Math.*, 6:373–412, 1985.
- [22] H.J. Kushner. *Weak Convergence Methods and Singularly Perturbed Stochastic Control and Filtering Problems*, volume 3 of *Systems and control*. Birkhaeuser, Boston, 1990.
- [23] T. Lehtonen and H. Nyhrinen. Simulating level crossing probabilities by importance sampling. *Adv. Appl. Probab.*, 24:858–874, 1992a.
- [24] T. Lehtonen and H. Nyhrinen. On asymptotically efficient simulating of ruin probabilities in a Markovian environment. *Scand. Acturial. J.*, 1:60–75, 1992b.
- [25] R.T. Rockafellar. *Convex Analysis*. Princeton University Press, Princeton, 1970.
- [26] J.S. Sadowsky. Large deviations and efficient simulation of excessive backlogs in a $GI/G/m$ queue. *IEEE. Trans. Automat. Control*, 36:1383–1394, 1991.
- [27] J.S. Sadowsky. On the optimality and stability of exponential twisting in Monte Carlo estimation. *IEEE. Trans. Inform. Theory*, 39:119–128, 1993.
- [28] J.S. Sadowsky. On Monte Carlo estimation of large deviations probabilities. *Ann. Appl. Prob.*, 6:399–422, 1996.
- [29] J.S. Sadowsky and J.A. Bucklew. On large deviations theory and asymptotically efficient Monte Carlo estimation. *IEEE Trans. Inform. Theory*, 36:579–588, 1990.
- [30] P. Shahsbuddin. Importance sampling for the simulation of highly reliable Markovian systems. *Management Sciences*, 40:333–352, 1994.
- [31] D. Siegmund. Importance sampling in the Monte Carlo study of sequential tests. *Ann. Statist.*, 4:673–684, 1976.