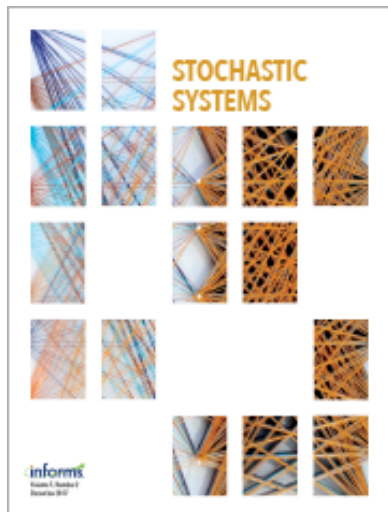




Stochastic Systems

Publication details, including instructions for authors and subscription information:
<http://pubsonline.informs.org>

Efficient Steady-State Simulation of High-Dimensional Stochastic Networks

Jose Blanchet, Xinyun Chen, Nian Si, Peter W. Glynn

To cite this article:

Jose Blanchet, Xinyun Chen, Nian Si, Peter W. Glynn (2021) Efficient Steady-State Simulation of High-Dimensional Stochastic Networks. *Stochastic Systems* 11(2):174-192. <https://doi.org/10.1287/stsy.2021.0077>

This work is licensed under a Creative Commons Attribution 4.0 International License. You are free to copy, distribute, transmit and adapt this work, but you must attribute this work as “*Stochastic Systems*. Copyright © 2021 The Author(s). <https://doi.org/10.1287/stsy.2021.0077>, used under a Creative Commons Attribution License: <https://creativecommons.org/licenses/by/4.0/>.”

Copyright © 2021, The Author(s)

Please scroll down for article—it is on subsequent pages



With 12,500 members from nearly 90 countries, INFORMS is the largest international association of operations research (O.R.) and analytics professionals and students. INFORMS provides unique networking and learning opportunities for individual professionals, and organizations of all types and sizes, to better understand and use O.R. and analytics tools and methods to transform strategic visions and achieve better outcomes. For more information on INFORMS, its publications, membership, or meetings visit <http://www.informs.org>

Efficient Steady-State Simulation of High-Dimensional Stochastic Networks

Jose Blanchet,^a Xinyun Chen,^b Nian Si,^a Peter W. Glynn^a

^aDepartment of Management Science and Engineering, Stanford University, Stanford, California 94305; ^bSchool of Data Science, Chinese University of Hong Kong, Shenzhen Institute of Artificial Intelligence and Robotics for Society, Shenzhen, Guangdong 518172, China

Contact: jose.blanchet@stanford.edu (JB); chenxinyun@cuhk.edu.cn,  <https://orcid.org/0000-0003-1727-0923> (XC); niansi@stanford.edu (NS); glynn@stanford.edu (PWG)

Received: February 19, 2020

Revised: October 31, 2020

Accepted: February 1, 2021

Published Online in Articles in Advance:
June 2, 2021

<https://doi.org/10.1287/stsy.2021.0077>

Copyright: © 2021 The Author(s)

Abstract. We propose and study an asymptotically optimal Monte Carlo estimator for steady-state expectations of a d -dimensional reflected Brownian motion (RBM). Our estimator is asymptotically optimal in the sense that it requires $\tilde{O}(d)$ (up to logarithmic factors in d) independent and identically distributed scalar Gaussian random variables in order to output an estimate with a controlled error. Our construction is based on the analysis of a suitable multilevel Monte Carlo strategy which, we believe, can be applied widely. This is the first algorithm with linear complexity (under suitable regularity conditions) for a steady-state estimation of RBM as the dimension increases.

 **Open Access Statement:** This work is licensed under a Creative Commons Attribution 4.0 International License. You are free to copy, distribute, transmit and adapt this work, but you must attribute this work as “Stochastic Systems. Copyright © 2021 The Author(s). <https://doi.org/10.1287/stsy.2021.0077>, used under a Creative Commons Attribution License: <https://creativecommons.org/licenses/by/4.0/>.”

Funding: J. Blanchet acknowledges support from the Air Force Office of Scientific Research [Award FA9550-20-1-0397] and the National Science Foundation [Grants 1915967, 1820942, and 1838676]. X. Chen acknowledges support from the National Natural Science Foundation of China [Grant 11901493].

Keywords: reflected Brownian motion • multilevel Monte Carlo • steady-state simulation • queueing networks

1. Introduction

Many complex stochastic systems can be modeled by a high-dimensional stochastic network, where applications include communication networks (Kushner 2013), cloud computing cluster (Maguluri et al. 2012), and patient flow in hospitals (Armony et al. 2015). Furthermore, the steady-state analysis of these systems is of interest because operators often focus on long-term average rewards/costs per unit of time. This motivates our focus in this paper, namely, the study of efficient Monte Carlo methods for computing steady-state performance measures for high-dimensional stochastic networks.

We consider a family of multidimensional reflected Brownian motions (RBMs) living on the positive orthant. We propose a steady-state simulation algorithm that is optimal under natural uniformity conditions (as the dimension increases), in the sense of requiring almost a linear number of independent and identically distributed (i.i.d.) Gaussian random variables to compute the steady-state expectation of the underlying RBM to given accuracy. We will provide an explicit description of the assumptions that we impose in Section 2. These conditions correspond basically to uniform stability and uniformly bounded variances. As far as we are aware, this paper provides the first class of optimal steady-state estimators for a reasonably general class of stochastic networks having a linear complexity in the dimension d .

RBM can be used to approximate the workload process for a wide range of stochastic networks in heavy traffic. In addition, RBM can be succinctly parameterized in terms of means, variances, and the routing architecture of the network. These properties make it an ideal vehicle for the study of Monte Carlo estimators for stochastic networks indexed by a set of parameters growing in the number of dimensions. In other words, RBM is a parsimonious, yet powerful, stylized model capturing the features that make steady-state analysis of stochastic networks challenging. We note that direct computation of the steady-state distribution for RBM is a very challenging problem, even in low dimensions. There is no closed-form expression for steady-state expectation in general, and even numerical methods are difficult to apply. These difficulties arise from the fact that RBM is defined in

terms of a system of constrained stochastic differential equations determined by the Skorokhod problem, which involves delicate local-time-like dynamics.

The steady-state distribution of RBM satisfies a partial differential equation (PDE) known as the basic adjoint relationship. Using this relationship, one can apply numerical methods such as finite differences to approximate the associated stationary distribution. These types of approaches are documented in, for instance, Dai and Harrison (1992) and Shen et al. (2002). Finite differences typically suffer from the curse of dimensionality. Moreover, in order to approximate the steady-state distribution, the methods in Dai and Harrison (1992) assume a certain degree of regularity that has not been established rigorously. Monte Carlo methods for steady-state analysis have also been studied, including those in Blanchet and Chen (2015), which apply coupling from the past technique to sample from the steady-state distribution of RBM with a controlled approximation error. However, the procedure in Blanchet and Chen (2015) exhibits exponential complexity in the dimension d .

Our analysis builds on recent work by Banerjee and Budhiraja (2020) and Blanchet and Chen (2020). The first results showing a polynomial rate in the dimension d of convergence to steady state for a high-dimensional RBM is given in Blanchet and Chen (2020). The proof technique used in Blanchet and Chen (2020) involves the following three ingredients: (a) the use of a coupling between a steady-state version of the RBM and one starting from a given initial condition driven by the same Brownian motion; (b) the application of results from Kella and Ramasubramanian (2012), which leads to a contraction factor as the product of certain random matrices when the process hits the constraint boundaries at a certain epochs; (c) a Lyapunov bound that estimates the return times of the contraction epochs—basically the return time to the constraint boundaries. By combining (a) through (c), Blanchet and Chen (2020) provided an estimate of the form $O(d^4 \log^2(d))$ for the relaxation time (measured in terms of the Wasserstein distance) between an RBM starting from the origin and its steady-state distribution. The work of Banerjee and Budhiraja (2020) introduced a weighted Lyapunov function (i.e., modifying step (c)), greatly improving these estimates and obtaining a relaxation time of $O(\log^2(d))$. This suggests simulating the RBM of interest for a time $O(\log^2(d))$ to control the size of the initial transient bias.

In addition to dealing with the initial transient bias, numerical simulation also involves discretization bias. In particular, discretizing a one-dimensional Brownian motion with a grid of size ε induces an error $\tilde{O}(\varepsilon^{1/2})$ in the uniform norm on compact intervals. (The tilde notation here means that we are ignoring poly-logarithmic factors in $1/\varepsilon$.) To control the simulation error, it is crucial to understand how the discretization error impacts the d -dimensional RBM as d grows. It is known, owing to the seminal work of Harrison and Reiman (1981), that RBM is a Lipschitz continuous functional of Brownian motion in the uniform metric. We show that under the uniform stability and uniform bounded variances assumptions, the Lipschitz constant is uniformly bounded in d . As a consequence, the discretization error is $\tilde{O}(1)$ in d . On the other hand, any simulation algorithms need to sample from each dimension, resulting in an $\Omega(d)$ computational cost lower bound measured by the number of i.i.d. standard Gaussian random variables simulated. Our algorithm in this paper shows that the estimation can actually achieve this linear lower bound (up to logarithmic factors in d).

We summarize the main contributions of this paper as follows:

I. First, we theoretically show that our simulation estimator approximates the steady-state distribution of the underlying RBM in $\tilde{O}(d)$ time (i.e., almost linear time), measured in terms of i.i.d. Gaussian random variables generated. Further, we show an $\Omega(d)$ lower bound for the computational cost, which demonstrates that our algorithm is optimal in terms of the dimension dependence (up to logarithmic factors).

II. Second, we provide an alternative method to deriving the contraction estimates for the initial transient bias, which is based on the derivative of the underlying RBM with respect to the initial condition; see Lemma 5 and Remark 4. The intuition is that the rate of convergence to stationarity is dictated by how fast the process forgets its initial condition, that is, how fast the derivative with respect to the initial condition converges to zero. Although the methods in Banerjee and Budhiraja (2020) and Blanchet and Chen (2020) could also lead to the estimates in the RBM case, we believe that our high-level approach could be applicable to broader settings, as we discuss in our conclusion section.

A key idea behind our first contribution is that for numerical simulation of a d -dimension RBM on finite time intervals, we analyze the contribution of discretization bias of the d Brownian motions altogether instead of separately, in order to obtain a finer bound on the simulation bias.

A crucial aspect in the development of contribution II is the use of derivative estimates of RBM with respect to the initial condition, using tools developed in Mandelbaum and Ramanan (2010). These derivatives, as it turns out, can be computed as the product of random matrices precisely arising in item (b) mentioned earlier in the analysis of Blanchet and Chen (2020). This is both reassuring and convenient because we can simply take advantage of the analysis both in Blanchet and Chen (2020) and Banerjee and Budhiraja (2020). However, studying the derivative process with respect to the initial condition is a type of strategy that can be applied in a wide range of settings of interest. So, we believe that the strategy deployed in this paper can be used as a blueprint for the

development of efficient Monte Carlo methods for high-dimensional steady-state analysis in many other settings. These developments will be studied in future research.

Our estimators are built using the multilevel Monte Carlo (MLMC) method (see Giles 2008) in conjunction with the key idea discussed earlier and also the contraction-estimating method mentioned in II). For a review of multilevel Monte Carlo, the reader is referred to Giles (2015). The MLMC method and its randomized variant, which can be used to remove bias under certain conditions (see Rhee and Glynn 2015), have been investigated both in the discretization of stochastic differential equations and, more recently, in the context of steady-state expectations; see Giles et al. (2020) and Glynn and Rhee (2014).

As in Giles et al. (2020), we are concerned both with the error in the numerical discretization of the underlying stochastic differential equation (SDE) and the time horizon contraction property. We use a synchronous coupling which, in our case, is motivated by the analysis in Blanchet and Chen (2020). A key difference, however, is that our goal is to study the complexity of the method as the dimension d increases to infinity and show that our estimator has essentially linear complexity in the dimension, as measured by the total number of generated random seeds. Indeed, we believe that this is also a key difference between our work and virtually every work to the date that uses multilevel Monte Carlo methods or steady-state Monte Carlo estimation in generic stochastic networks.

The rest of the paper is organized as follows. In Section 2, we review the definition of RBM and discuss the uniformity conditions we use to test the asymptotic optimality of our algorithm. The simulation algorithm is given in Section 3, together with the main result of this paper, Theorem 1. A numerical experiment that validates the theoretical performance of the algorithm, tested in the setting of networks of increasing size, is given in Section 4. Finally, the proofs of Theorems 1 and 2 are given in Section 5.

1.1. Notation

The boldface denotes multidimensional variables. All vector inequalities are assumed to be operated componentwise. We use $|x|$ for $x \in \mathbb{R}^d$ to denote the absolute value componentwise.

2. Model and Assumptions

2.1. Skorokhod Problem and RBM

A multidimensional RBM can be defined as the solution to a Skorokhod problem with Brownian input. In particular, let $\mathbf{X}(\cdot)$ be a multidimensional Brownian motion with drift vector $\boldsymbol{\mu}$, covariance matrix $\Sigma := CC^T$, and initial value $\mathbf{X}(0) = 0$. Let Q be a substochastic matrix, that is, $Q \geq 0$ and all its row sums ≤ 1 , and define $R = (I - Q)^T$. We assume R is an M -matrix, that is,

$$R^{-1} \text{ exists and it has nonnegative entries.} \quad (1)$$

The seminal paper by Harrison and Reiman (1981) shows that the following Skorokhod problem (2) is well posed, (i.e., it has a unique strong solution) in the case where the input $\mathbf{X}(\cdot)$ is continuous and R is an M -matrix.

Skorokhod Problem. Given a process $\mathbf{X}(\cdot)$ and a matrix R , we say that the pair $(\mathbf{Y}, \mathbf{L})$ solves the associated Skorokhod problem if

$$0 \leq \mathbf{Y}(t) = \mathbf{Y}(0) + \mathbf{X}(t) + R\mathbf{L}(t), \quad \mathbf{L}(0) = 0, \quad (2)$$

where the i th entry of $\mathbf{L}(\cdot)$ is nondecreasing and $\int_0^t Y_i(s) dL_i(s) = 0$.

When the input process $\mathbf{X}$ is a multidimensional Brownian motion with parameter $(\boldsymbol{\mu}, \Sigma)$, we call the process $\mathbf{Y}(\cdot)$ solved from (2) a $(\boldsymbol{\mu}, \Sigma, R)$ -RBM.

Remark 1. From the perspective that RBM $\mathbf{Y}(\cdot)$ is an approximation to the workload process of a stochastic network, the assumption that R is an M -matrix is equivalent to $Q^n \rightarrow 0$, that is, the network is open in the sense that all jobs will eventually leave the network.

For general Skorokhod problems, under the M -condition and some mild conditions on $\mathbf{X}(\cdot)$, the assumption that

$$R^{-1}E\mathbf{X}(1) = R^{-1}\boldsymbol{\mu} < 0, \quad (3)$$

implies that $\mathbf{Y}(t) \Rightarrow \mathbf{Y}(\infty)$ as $t \rightarrow \infty$, where $\mathbf{Y}(\infty)$ is a random variable with the (unique) stationary distribution of $\mathbf{Y}(\cdot)$. (We use $\Rightarrow$ to denote weak convergence.) In particular, according to Harrison and Williams (1987), condition (3) is necessary and sufficient for stability of the $(\boldsymbol{\mu}, \Sigma, R)$ -RBM (i.e., a unique stationary distribution exists) under the M -condition (1) (see also Kella and Ramasubramanian 2012, which studies necessary and sufficient conditions for more general types of input processes).

2.2. Assumptions

The goal of our simulation algorithm is to estimate the steady-state expectation of certain functions of a multidimensional RBM. In particular, let $(\boldsymbol{\mu}, \Sigma, R)$ be the parameters of the RBM and $f(\cdot)$ be the function to be evaluated. To study the complexity of the algorithm as the dimension grows, we shall consider a family of $(\boldsymbol{\mu}, \Sigma, R)$ -RBMs under certain uniformity assumptions for arbitrary dimension d , as in Blanchet and Chen (2020). Implicitly, R , $\boldsymbol{\mu}$, and Σ are indexed by their dimension. Now we state the uniformity conditions imposed throughout the paper.

A1. Uniform contraction: We let $R = I - Q^T$, where Q is substochastic and assume that there exists $\beta_0 \in (0, 1)$ and $\kappa_0 \in (0, \infty)$ independent of d such that

$$\|\mathbf{1}^T Q^n\|_\infty \leq \kappa_0 (1 - \beta_0)^n, \quad n \geq 1. \quad (4)$$

Under (4), we observe that

$$\|R^{-1} \mathbf{1}\|_\infty \leq b_1 := \kappa_0 / \beta_0 < \infty.$$

A2. Uniform stability: We write $\mathbf{X}(t) = \boldsymbol{\mu}t + C\mathbf{B}(t)$, where $\mathbf{B}(t) = (B_1(t), \dots, B_d(t))^T$ and the $B_i(\cdot)$'s are standard Brownian motions, and the matrix C satisfies $\Sigma = CC^T$. We assume that there exists $\delta_0 > 0$ independent of d such that

$$R^{-1} \boldsymbol{\mu} < -\delta_0 \mathbf{1}.$$

A3. Uniform marginal variability: Define $\sigma_i^2 = \Sigma_{i,i}$ (i.e., the variance of the i th coordinate of $\mathbf{X}$). We assume that there exists $b_0 \in (0, \infty)$, independent of $d \geq 1$, such that

$$b_0^{-1} \leq \sigma_i^2 \leq b_0.$$

A4. Lipschitz functions: Throughout the rest of the paper, we assume that the function $f: \mathbb{R}_+^d \rightarrow \mathbb{R}$, for which we shall estimate $E[f(\mathbf{Y}(\infty))]$, is Lipschitz continuous in l_∞ norm, that is, there exists a constant $\mathcal{L} > 0$ independent with d such that

$$|f(\mathbf{y}) - f(\mathbf{y}')| \leq \mathcal{L} \|\mathbf{y} - \mathbf{y}'\|_\infty, \quad \text{for all } \mathbf{y}, \mathbf{y}' \in \mathbb{R}_+^d.$$

Remark 2. A detailed discussion of Assumptions A1 to A3 is given in section 2.2 of Blanchet and Chen (2020). Assumption A4 holds if f is chosen to quantify the performance of a finite number of servers in the network, or when the performance measure of the system is scaled by d , for instance, the average workload at the servers.

3. Two-Parameter Multilevel Monte Carlo Algorithm

Any simulation estimator for stationary expectations of RBM is bound to contain two types of sources of bias. The first one is the discretization error, due to the fact that we can only simulate discrete approximation of continuous Brownian paths. The second source of bias is the initial transient bias or nonstationary error, due to the fact that we can only simulate the RBM during a finite time horizon. We call our simulation method a two-parameter MLMC algorithm because when constructing the MLMC estimator, we use two parameters, $\gamma \in (0, 1)$ and $T > 0$, to control the discretization and nonstationary errors, respectively.

As in the classic MLMC algorithm (Giles 2008), the precision of the MLMC estimator is controlled by the total number of levels L . Besides, we need to specify the initial state $\mathbf{y}_0$ to simulate the RBM paths. Given the parameter set $(\gamma, T, L, \mathbf{y}_0)$, plus the parameters $(\boldsymbol{\mu}, \Sigma, R)$ for the RBM and the function f to be evaluated, we now describe how to construct the two-parameter MLMC estimator for $E[f(\mathbf{Y}(\infty))]$ and we will summarize the whole procedure at the end of this section.

Let $\mathbf{B}(t) = (B_1(t), \dots, B_d(t))^T \in \mathbb{R}^d$ be a standard Brownian with drift $\mathbf{0}$ and covariance matrix I . Given parameter $\gamma \in (0, 1)$, we denote $\mathbb{D}_m = \{0, \gamma^m, 2\gamma^m, \dots\}$ for any integer $m \geq 0$. For every $t \geq 0$, we define $t_m^+ = \inf\{r \in \mathbb{D}_m : r > t\}$ and $t_m^- = \sup\{r \in \mathbb{D}_m : r \leq t\}$. Note that, following the definition, $t_m^- = t$ for $t \in \mathbb{D}_m$. Define a discretization of the standard Brownian motion of level m as $\mathbf{B}^m(t) = (B_1^m(t), \dots, B_d^m(t))^T$ such that

$$B_i^m(t) = B_i(t_m^-) + (t - t_m^-) \frac{B_i(t_m^+) - B_i(t_m^-)}{t_m^+ - t_m^-}, \quad \text{for all } t \geq 0 \text{ and } i = 1, 2, \dots, d.$$

It is easy to see that $\mathbf{B}^m(\cdot)$ is continuous and piecewise linear, and $\mathbf{B}^m(t) = \mathbf{B}(t)$ for all $t \in \mathbb{D}_m$. The corresponding discretization of the Brownian motion $\mathbf{X}(\cdot)$ driving the RBM (2) is defined as

$$\mathbf{X}^m(t) = \boldsymbol{\mu}t + C\mathbf{B}^m(t).$$

For any $0 \leq s \leq t < \infty$, we write $\mathbf{X}_{s:t}$ (respectively, $\mathbf{X}_{s:t}^m$) to denote the increment of $\mathbf{X}(\cdot)$ over $[s, t]$, that is, $\mathbf{X}_{s:t} = \{\mathbf{X}(s+u) - \mathbf{X}(s) : 0 \leq u \leq t-s\}$, (respectively, $\mathbf{X}_{s:t}^m = \{\mathbf{X}^m(s+u) - \mathbf{X}^m(s) : 0 \leq u \leq t-s\}$). We use $\mathbf{Y}(t-s; \mathbf{y}, \mathbf{X}_{s:t})$

(respectively, $\mathbf{Y}^m(t-s; \mathbf{y}, \mathbf{X}_{s:t}^m)$) to denote the value of RBM driven by $\mathbf{X}_{s:t}$ (respectively, $\mathbf{X}_{s:t}^m$) at time point $t-s$ given initial value $\mathbf{Y}(0) = \mathbf{y}$ (respectively, $\mathbf{Y}^m(0) = \mathbf{y}$). Following this notation, we have

$$\begin{aligned}\mathbf{Y}(t+s; \mathbf{y}, \mathbf{X}_{0:s+t}) &= \mathbf{Y}(t; \mathbf{Y}(s; \mathbf{y}, \mathbf{X}_{0:s}), \mathbf{X}_{s:s+t}), \\ \mathbf{Y}^m(t+s; \mathbf{y}, \mathbf{X}_{0:s+t}^m) &= \mathbf{Y}^m(t; \mathbf{Y}^m(s; \mathbf{y}, \mathbf{X}_{0:s}^m), \mathbf{X}_{s:s+t}^m).\end{aligned}\quad (5)$$

To construct the multilevel estimator, we introduce an integer-valued random variable $M \in \{0, 1, 2, \dots, L-1\}$, where L is the total number of levels. The random variable M is independent of the process $\mathbf{X}(\cdot)$ and follows probability distribution

$$P(M = m) = p(m) = \gamma^m(1-\gamma)/(1-\gamma^L) \triangleq K(\gamma)\gamma^m, \quad \text{for } 0 \leq m < L.$$

In the multilevel Monte Carlo method, to reduce the computational cost of estimating the target expected value, people use different levels of approximations $Z_0, Z_1, \dots, Z_L$ with increasing accuracy, but also increasing cost, which converges to the target random variable as $L \rightarrow \infty$. The basis of the multilevel method can thus be

$$E[Z_L] = E[Z_0] + \sum_{l=0}^{L-1} E[Z_{l+1} - Z_l] = E[Z_0] + \frac{1}{p(M)} E[Z_{M+1} - Z_M].$$

In the RBM setting, by choosing

$$Z_l = f(\mathbf{Y}^M(MT; \mathbf{y}_0, \mathbf{X}_{T:(M+1)T}^M)),$$

we give the formal definition of the two-parameter MLMC estimator Z for $E[f(\mathbf{Y}(\infty))]$ with input parameter set $(\gamma, T, L, \mathbf{y}_0)$ as

$$Z = \frac{1}{p(M)} \left\{ f(\mathbf{Y}^{M+1}(MT; \mathbf{Y}^{M+1}(T; \mathbf{y}_0, \mathbf{X}_{0:T}^{M+1}), \mathbf{X}_{T:(M+1)T}^{M+1})) - f(\mathbf{Y}^M(MT; \mathbf{y}_0, \mathbf{X}_{T:(M+1)T}^M)) \right\} + f(\mathbf{y}_0). \quad (6)$$

To see that Z is indeed a good estimator for $E[f(\mathbf{Y}(\infty))]$, we compute

$$\begin{aligned}E[Z] &= E[E[Z|M]] \\ &= \sum_{m=0}^{L-1} \left(E \left[f(\mathbf{Y}^{m+1}(mT; \mathbf{Y}^{m+1}(T; \mathbf{y}_0, \mathbf{X}_{0:T}^{m+1}), \mathbf{X}_{T:(m+1)T}^{m+1})) \right] - E \left[f(\mathbf{Y}^m(mT; \mathbf{y}_0, \mathbf{X}_{T:(m+1)T}^m)) \right] \right) \\ &\quad + f(\mathbf{y}_0) = \sum_{m=0}^{L-1} \left(E \left[f(\mathbf{Y}^{m+1}((m+1)T; \mathbf{y}_0, \mathbf{X}_{0:(m+1)T}^{m+1})) \right] - E[f(\mathbf{Y}^m(mT; \mathbf{y}_0, \mathbf{X}_{0:mT}^m))] \right) + f(\mathbf{y}_0) = E \left[f(\mathbf{Y}^L(LT; \mathbf{y}_0, \mathbf{X}_{0:LT}^L)) \right].\end{aligned}$$

The last equality holds because $\mathbf{Y}^m(mT; \mathbf{y}_0, \mathbf{X}_{0:mT}^m) = \mathbf{y}_0$ for $m = 0$. Consequently, we can split the estimation bias into two parts:

$$\begin{aligned}& E \left[f(\mathbf{Y}^L(LT; \mathbf{y}_0, \mathbf{X}_{0:LT}^L)) \right] - E[f(\mathbf{Y}(\infty))] \\ &= \left(E \left[f(\mathbf{Y}^L(LT; \mathbf{y}_0, \mathbf{X}_{0:LT}^L)) \right] - E[f(\mathbf{Y}(LT; \mathbf{y}_0, \mathbf{X}_{0:LT}))] \right) \\ &\quad + \left(E[f(\mathbf{Y}(LT; \mathbf{y}_0, \mathbf{X}_{0:LT}))] - E[f(\mathbf{Y}(\infty))] \right) \\ &= \text{discretization error} + \text{nonstationarity error}.\end{aligned}\quad (7)$$

Intuitively, as $L \rightarrow \infty$, the two errors will both go to 0, and as a consequence, we can obtain an accurate estimate of $E[f(\mathbf{Y}(\infty))]$ by taking L sufficiently large. In Sections 5.1 and 5.2, we shall provide theoretical upper bounds for those two errors in terms of L , and also analyze their dependence on the number of dimensions d . Then, we apply these theoretical error bounds to control the mean square error (MSE) of the simulation estimator, and obtain the main complexity analysis result for our simulation algorithm in Section 5.3.

The previous description of the two-parameter multilevel Monte Carlo method is summarized in Algorithm 1. The main result of the paper is as follows. We show that, under the proper choice of algorithm hyperparameters, as described in Algorithm 1, the computational cost for obtaining an estimator of a fixed accuracy level is almost linear in the dimension d . Here, we interchangeably use computational complexity and the total expected cost, which is measured by the number of scalar Gaussian random variable generated in the algorithm. Specifically, we generate one Gaussian random variable for each dimension, discretization point, and simulation path. The proof relies on an analysis at the dimension dependence of the discretization and the nonstationarity error in the simulation procedures, and will be given in Section 5.

Theorem 1. Suppose $\mathbf{Y}$ is a d -dimensional RBM satisfying Assumptions A1 to A4. By setting parameters as step size $0 < \gamma < 1$,¹ path length $T = O(\log(d)^2)$, the number of levels $L = \lceil (\log(\log(d)) + 2\log(1/\varepsilon) + k_1)/\log(1/\gamma) \rceil$, for a numerical constant k_1 , the initial point $\mathbf{y}_0 = 0$, and the number of sample paths $N = \lceil (1 - \gamma)^{-1}(1 - \gamma^L)\gamma^{-L}L \rceil$, the total expected cost, in terms of the number of scalar Gaussian random variables, for the two-parameter MLMC Algorithm 1 to produce an estimator $\bar{Z} = N^{-1} \sum_{i=1}^N Z_i$ (defined in (6)) of $E[f(\mathbf{Y}(\infty))]$ with mean square error (MSE) ε^2 is

$$O\left(\varepsilon^{-2} d \log(d)^3 (\log(\log(d)) + \log(1/\varepsilon))^3\right).$$

In fact, the following theorem shows that the linear dependence on the dimension d of the total expected cost is optimal (up to logarithmic factors in d). Before stating the theorem, we first define some useful notions.

Definition 1. Let $\Psi(\kappa_0, \beta_0, \delta_0, b_0)$ be the set of all possible RBMs $(\mathbf{X}; \mathbf{Y})$ that satisfy Assumptions A1 to A3, and let $\text{Lips}_{\mathcal{L}}$ be the set of all $\mathcal{L}$ -Lipschitz functions in l_∞ norm satisfying Assumption A4. Let Π^ε be the class of all algorithms π satisfying the following four requirements:

- can access the reflection matrix R ;
- is agnostic to the model parameters μ and Σ , but is able to sample from any given discrete skeleton of the process $\mathbf{X}$;
- is agnostic to the function $f \in \text{Lips}_{\mathcal{L}}$, but is able to query the function; and
- can achieve the target mean square error ε^2 when applied to each RBM $(\mathbf{X}, \mathbf{Y}) \in \Psi(\kappa_0, \beta_0, \delta_0, b_0)$ and the target function $f \in \text{Lips}_{\mathcal{L}}$.

Further, Let $\text{TC}(\pi; (\mathbf{X}, \mathbf{Y}), f)$ be the total expected cost of the algorithm $\pi \in \Pi^\varepsilon$, associated with the RBM $(\mathbf{X}, \mathbf{Y}) \in \Psi(\kappa_0, \beta_0, \delta_0, b_0)$ and the target function $f \in \text{Lips}_{\mathcal{L}}$.

Theorem 2. Algorithm 1 is inside the algorithm class Π^ε , and for any algorithm $\pi \in \Pi^\varepsilon$, we have a minimax lower bound for the total expected cost when $\varepsilon < \mathcal{L}/(16\delta_0)$,

$$\inf_{\pi \in \Pi^\varepsilon} \sup_{(\mathbf{X}, \mathbf{Y}) \in \Psi(\kappa_0, \beta_0, \delta_0, b_0), f \in \text{Lips}_{\mathcal{L}}} \text{TC}(\pi; (\mathbf{X}, \mathbf{Y}), f) \geq d.$$

4. Numerical Experiments

We test the theoretical performance guarantee (i.e., Theorem 1) of our algorithm using so-called symmetric RBMs. In this case, the true value of $E[Y_1(\infty)]$ has a closed-form expression, so that we can check the dimension dependence of the simulation MSE and complexity. To do this, we consider a sequence of symmetric RBMs of different dimensions from five up to 200. More precisely, for each $d \in \{5, 6, \dots, 200\}$, $\mu = -[1, 1, \dots, 1]^T$, the covariance matrix takes the form

$$\Sigma = \begin{bmatrix} 1 & \rho_\sigma & \cdots & \rho_\sigma \\ \rho_\sigma & 1 & \cdots & \rho_\sigma \\ \vdots & & 1 & \vdots \\ \rho_\sigma & \cdots & \rho_\sigma & 1 \end{bmatrix},$$

and the reflection matrix takes the form

$$R = \begin{bmatrix} 1 & -r & \cdots & -r \\ -r & 1 & \cdots & -r \\ \vdots & & 1 & \vdots \\ -r & \cdots & -r & 1 \end{bmatrix}.$$

Algorithm 1 (Two-Parameter MultilevelMonte Carlo for RBM)

Input:

The parameters of the RBM: (μ, Σ, R) ;

The function to evaluate: $f: \mathbb{R}_+^d \rightarrow \mathbb{R}$;

The target error level ε ;

Output:

An estimator for $E[\mathbf{Y}(\infty)]$, $\bar{Z}$;

Algorithm procedure:

1: for $i = 1$ to N do

- 2: Generate M with $P(M = m) = p(m) = K(\gamma)\gamma^m$;
- 3: Simulate a discrete Brownian path $\mathbf{B}^{M+1}(t)$ with step size γ^{M+1} on $[0, (M+1)T]$;
- 4: Compute $\mathbf{B}^M(t)$ as a discrete Brownian path such that $\mathbf{B}^M(t) = \mathbf{B}^{M+1}(t)$ for all $t \in \mathbb{D}_M$;
- 5: Solve the Skorokhod problem (Algorithm A.1 in the appendix) to compute

$$X^M(t) = \mu t + \mathbf{C}\mathbf{B}^M(t) \quad \text{and} \quad X^{M+1}(t) = \mu t + \mathbf{C}\mathbf{B}^{M+1}(t),$$

- 6: Compute

$$Z_i = \frac{1}{p(M)} \left(f(\mathbf{Y}^{M+1}((M+1)T, \mathbf{y}_0, \mathbf{X}_{0:(M+1)T})) - f(\mathbf{Y}^M(MT, \mathbf{y}_0, \mathbf{X}_{T:(M+1)T})) \right) + f(\mathbf{y}_0);$$

return $\bar{Z} = \frac{1}{N} \sum_{i=1}^N Z_i$.

To be consistent with Assumptions A1 to A3, we pick

$$\rho_\sigma = -\frac{1-\beta}{d-1} \quad \text{and} \quad r = \frac{1-\beta}{d-1},$$

for given $0 < \beta < 1$. According to Dai and Harrison (1992), the steady-state expectation of workload at each station equals to

$$E[Y_1(\infty)] = \frac{1 - (d-2)r + (d-1)r\rho_\sigma}{2(1+r)} = \frac{\beta}{2}.$$

For $\beta = 0.8$, the true value of $E[Y_1(\infty)] = 0.4$.

In the first group of numerical experiments, we compare the algorithm performance for different choices of parameter $\gamma \in \{0.01, 0.05, 0.1\}$ at a target error level $\varepsilon = 0.01$. The other parameters are as follows: $T = \log(d)^2/2$, $L = \lceil (\log(\log(d)) + 2\log(1/\varepsilon) - 2)/\log(1/\gamma) \rceil$, and $N = \lceil K(\gamma)^{-1}\gamma^{-L}L \rceil$, where $K(\gamma) = (1-\gamma)/(1-\gamma^L)$. Figure 1 shows the estimated mean and total complexity across dimensions from $d = 5$ to $d = 200$ for different choices of γ . It shows that most of the absolute error fluctuates around 0.01 and the total number of generated scalar Gaussian random variables (total complexity) grows approximately linearly in the number of dimension for all three values of γ . The simulation error is not sensitive to the choice of γ . Besides, the complexity is best when $\gamma = 0.05$, as indicated by our theoretical analysis (Lemma 7).

In our second group of numerical experiments, we aim to show that our choice of parameters achieves the target precision level tightly in the sense that, as the dimension increases, the estimation error remains stable around a value that is smaller than the target precision level. In particular, we estimate the MSE of the estimators for $\gamma = 0.05$ and target error level $\varepsilon = 0.05$ with the other parameters fixed. For each dimension in $\{10, 20, 30, \dots, 200\}$, we generate 250 estimators to estimate the MSE as well as the 95% confidence band of the MSE, and the results are reported in Figure 2. We see the MSE is stable around 5×10^{-4} across different dimensions, which is smaller than the target level $\varepsilon^2 = 0.0025$.

5. Proof of Theorems 1 and 2

In this section, we develop theoretical computational complexity bounds for the two-parameter MLMC estimator in terms of the number of dimensions d using the hyperparameters $(\gamma, T, L, \mathbf{y}_0)$ specified in Algorithm 1. As in (7), the estimation bias of the two-parameter MLMC estimator Z can be split into two parts corresponding to the discretization error and nonstationarity error. The sketch of the proof is as follows:

- 1: In Section 5.1, we derive an upper bound for the discretization error in Lemma 4, which is based on the discretization error for Brownian motion (Lemma 2) and an explicit upper bound for the Lipschitz constant of the Skorokhod mapping (Lemma 3).
- 2: In Section 5.2, we provide a bound for the nonstationarity error in Lemma 6 by analyzing the derivative of RBM with respect to its initial value (Lemma 5).
- 3: In Section 5.3, we derive an upper bound for the algorithm complexity based on the error bounds.
- 4: Finally, in Section 5.4, we derive a lower bound for the algorithm complexity based on the error bounds to prove Theorem 2.

5.1. Discretization Error Bounds

To bound the discretization error, we first bound the discretization error for multidimensional Brownian motion.

Lemma 1. Suppose $Z_1, Z_2, \dots, Z_n$ are Gaussian variables (not necessarily independent) with mean μ and variance 1. Then, we have $E[\max_{1 \leq i \leq n} (Z_i - \mu)^2] \leq 4(\log n + 1/2\log(2))$.

Figure 1. Simulation Results for Symmetric RBMs at Target Error Level $\epsilon = 0.01$

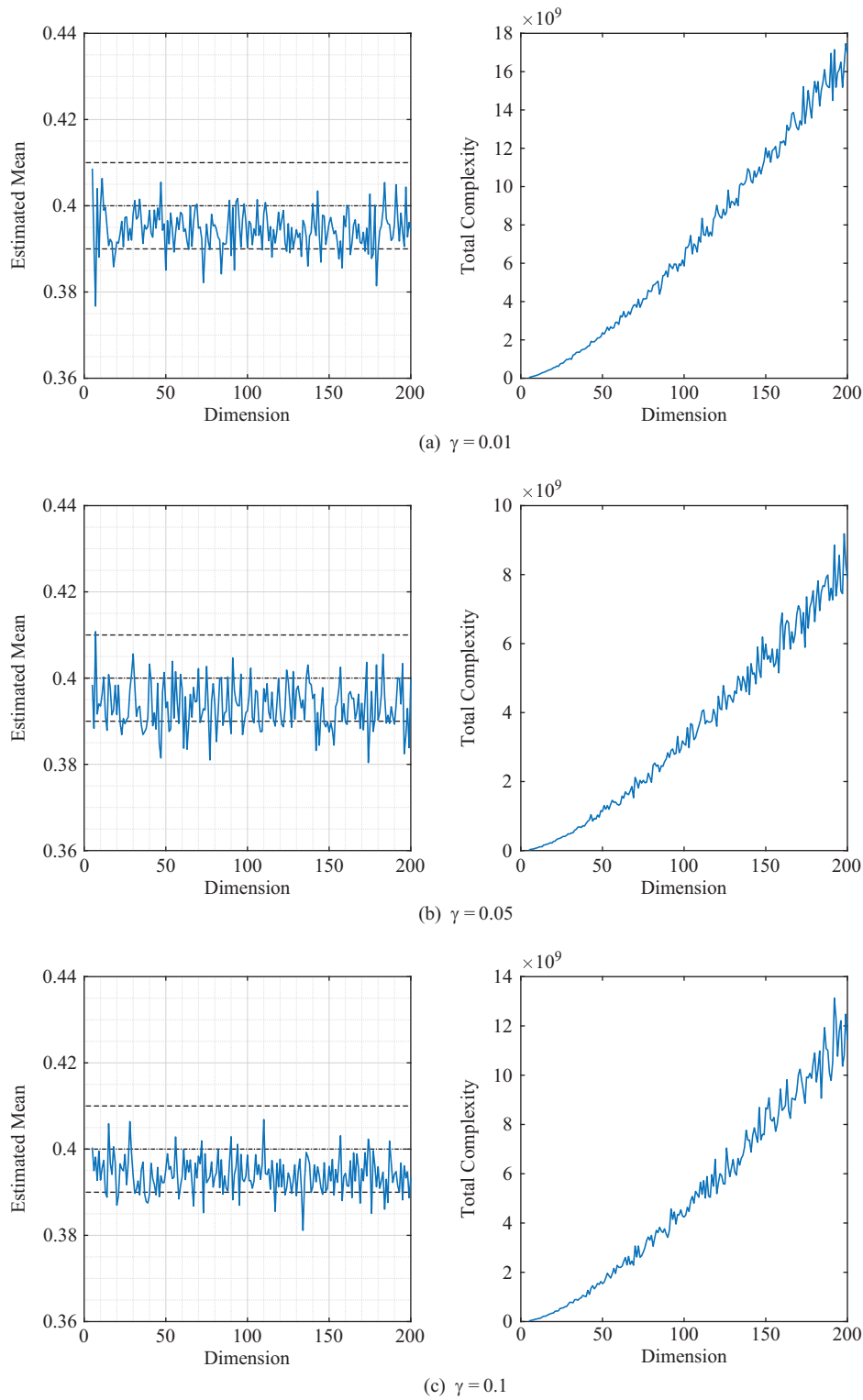
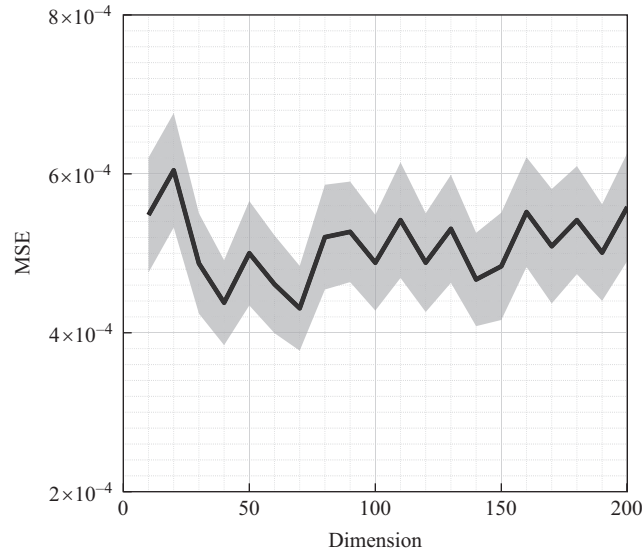


Figure 2. Mean Square Error of the Estimators at Target Error Level $\epsilon = 0.05$ for $\gamma = 0.05$ 

Note. The shaded area represents 95% confidence band for the MSE.

Proof of Lemma 1. For $\lambda \in (0, 1/2)$, we have

$$\begin{aligned} E\left[\max_{1 \leq i \leq n} (Z_i - \mu)^2\right] &= \frac{1}{\lambda} E\left[\log\left(\exp\left(\lambda \max_{1 \leq i \leq n} (Z_i - \mu)^2\right)\right)\right] \\ &\leq \frac{1}{\lambda} \log E\left[\exp\left(\lambda \max_{1 \leq i \leq n} (Z_i - \mu)^2\right)\right] \\ &\leq \frac{1}{\lambda} \log E\left[\sum_{i=1}^n \exp\left(\lambda (Z_i - \mu)^2\right)\right] \\ &= \frac{1}{\lambda} (\log n - 1/2 \log(1 - 2\lambda)). \end{aligned}$$

We can pick $\lambda = 1/4$ and then

$$E\left[\max_{1 \leq i \leq n} (Z_i - \mu)^2\right] \leq 4(\log n + 1/2 \log(2)). \quad \square$$

Lemma 2. For $0 < \gamma < 1$ and $m \geq 1$, let $\mathbf{X}^m(\cdot)$ be a discretized d -dimension Brownian path with step size γ^m . Then, there exists a positive constant C_0 , such that for any $d \geq 2$, $m \geq 1$, $t > \gamma$,

$$E\left[\max_{1 \leq i \leq d} \max_{0 \leq s \leq t} (X_i^m(s) - X_i(s))^2\right] \leq C_0 \gamma^m (\log(t) + \log(d) + m \log(1/\gamma)).$$

Proof of Lemma 2. Let $\tilde{\mathbf{X}}(t) = \mathbf{X}(t) - \mu t$ and $\tilde{\mathbf{X}}^m(t) = \mathbf{X}^m(t) - \mu t$. Note that

$$\begin{aligned} &\max_{1 \leq i \leq d} \max_{0 \leq s \leq t} (X_i^m(s) - X_i(s))^2 \\ &\leq \max_{1 \leq i \leq d} \max_{0 \leq s \leq \gamma^m \lceil t/\gamma^m \rceil} (X_i^m(s) - X_i(s))^2 \\ &= \max_{1 \leq i \leq d} \max_{0 \leq k \leq \lceil t/\gamma^m \rceil} \max_{0 \leq s \leq \gamma^m} (\tilde{X}_i(\gamma^m k + s) - \tilde{X}_i^m(\gamma^m k + s))^2. \end{aligned}$$

For $0 \leq s < \gamma^m$ and $0 \leq k \leq \lfloor t/\gamma^m \rfloor$, we have

$$\begin{aligned} & \left(\tilde{X}_i(\gamma^m k + s) - \tilde{X}_i^m(\gamma^m k + s) \right)^2 \\ & \leq \max \left\{ \left(\tilde{X}_i(\gamma^m k + s) - \tilde{X}_i(\gamma^m k) \right)^2, \left(\tilde{X}_i(\gamma^m k + \gamma^m) - \tilde{X}_i(\gamma^m k + s) \right)^2 \right\} \leq \left(\tilde{X}_i(\gamma^m k + s) - \tilde{X}_i(\gamma^m k) \right)^2 + \left(\tilde{X}_i(\gamma^m k + \gamma^m) - \tilde{X}_i(\gamma^m k + s) \right)^2. \end{aligned}$$

By time-reversibility of the Brownian process, we have

$$\left(\tilde{X}_i(\gamma^m k + s) - \tilde{X}_i(\gamma^m k) \right)^2 \stackrel{d}{=} \left(\tilde{X}_i(\gamma^m k + \gamma^m) - \tilde{X}_i(\gamma^m k + (\gamma^m - s)) \right)^2,$$

where $\stackrel{d}{=}$ indicates the two processes indexed by s follow the same probability law on the space of continuous functions. Then, using Brownian motion's independent increments and scaling properties, we have

$$\begin{aligned} & \max_{1 \leq i \leq d} \max_{0 \leq k \leq \lfloor t/\gamma^m \rfloor} \max_{0 \leq s \leq \gamma^m} \left(\tilde{X}_i(\gamma^m k + s) - \tilde{X}_i(\gamma^m k) \right)^2 \\ & \stackrel{d}{=} \gamma^m \max_{1 \leq i \leq d} \max_{0 \leq k \leq \lfloor t/\gamma^m \rfloor} \max_{0 \leq s \leq 1} \left(\tilde{X}_i^{(k)}(s) \right)^2, \end{aligned}$$

where $\tilde{X}^{(0)}, \tilde{X}^{(1)}, \dots$ are i.i.d. copies of $\tilde{X}$ and $\tilde{X}^{(k)} = \{\tilde{X}_1^{(k)}, \tilde{X}_2^{(k)}, \dots, \tilde{X}_d^{(k)}\}$. Recall that (e.g., Karlin and Taylor 1981, p. 346)

$$\max_{0 \leq s \leq 1} \left(\tilde{X}_i^{(k)}(s) \right)^2 \stackrel{d}{=} \left(\tilde{X}_i^{(k)}(1) \right)^2.$$

Then, by Lemma 1, we have for $d \geq 2$,

$$\begin{aligned} & E \left[\max_{1 \leq i \leq d} \max_{0 \leq s \leq \gamma^m \lfloor t/\gamma^m \rfloor} (X_i^m(s) - X_i(s))^2 \right] \\ & \leq 2\gamma^m E \left[\max_{1 \leq i \leq d} \max_{0 \leq k \leq \lfloor t/\gamma^m \rfloor} \left(\tilde{X}_i^{(k)}(1) \right)^2 \right] \\ & \leq 2b_0\gamma^m(4\log(d\lfloor t/\gamma^m \rfloor) + 2\log(2)) \\ & \leq 12b_0\gamma^m(\log(t) + \log(d) + m\log(1/\gamma)). \end{aligned}$$

The last inequality is due to $\lfloor t/\gamma^m \rfloor \leq 2t/\gamma^m$ when $t > \gamma, m \geq 1$ and $\log(d) \geq \log(2)$. $\square$

Lemma 3 shows that the Skorokhod mapping, from $\mathbf{X}$ to $\mathbf{Y}$, is Lipschitz continuous and provides a uniform upper bound for the Lipschitz constant, which is independent of the dimension d , the time s , and the input process $\mathbf{X}$. As a result, the discretization error $\sup_{0 \leq s \leq T} \|\mathbf{Y}(s) - \mathbf{Y}^m(s)\|_\infty$ can be bounded by $\sup_{0 \leq s \leq T} \|\mathbf{X}(s) - \mathbf{X}^m(s)\|_\infty$.

Lemma 3. Suppose $\mathbf{Y}(t)$ and $\mathbf{Y}'(t) \in \mathbb{R}^d$ are the solutions to two Skorokhod problems (2) with the same reflection matrix R satisfying Assumption A1, and input processes $\mathbf{X}(t)$ and $\mathbf{X}'(t)$, respectively, for $t \in [0, T]$. Then,

$$\sup_{0 \leq s \leq T} \|\mathbf{Y}(s) - \mathbf{Y}'(s)\| \leq 2R^{-1} \sup_{0 \leq s \leq T} \|\mathbf{X}(s) - \mathbf{X}'(s)\|,$$

where $\|\cdot\|$ and the inequality are assumed to be applied entry by entry. As a direct consequence, under Assumptions A1 to A3, we have

$$\|\mathbf{Y}(T) - \mathbf{Y}'(T)\|_\infty \leq \frac{2\kappa_0}{\beta} \sup_{0 \leq s \leq T} \|\mathbf{X}(s) - \mathbf{X}'(s)\|_\infty.$$

Proof of Lemma 3. The proof uses the fixed-point representation of the Skorokhod mapping as constructed in the proof of theorem 1 in Harrison and Reiman (1981). In detail, we first need to transform the space $\mathbb{R}^d$ using a diagonal matrix Θ with positive diagonal elements, which depends only on R , such that $(\Theta\mathbf{Y}, \Theta\mathbf{L})$ ($(\Theta\mathbf{Y}', \Theta\mathbf{L}')$) is the solution to a new Skorokhod problem of the form (2) with input process $\Theta\mathbf{X}$ ($\Theta\mathbf{X}'$) and reflection matrix $R^* = I - (\Theta^{-1}Q\Theta)^T$. (Note that in our notation, all vectors are column vectors, whereas in Harrison and Reiman (1981) they are treated as row vectors.)

Let $Q^* = \Theta^{-1}Q\Theta$. Then, according to Harrison and Reiman (1981), the amount of reflection $\Theta\mathbf{L}$ and $\Theta\mathbf{L}'$ solves the following fixed-point problem:

$$\begin{aligned}\Theta \mathbf{L}(t) &= \sup_{0 \leq s \leq t} (Q^{*T} \Theta \mathbf{L}(s) - \Theta \mathbf{X})^+ \\ \Theta \mathbf{L}'(t) &= \sup_{0 \leq s \leq t} (Q^{*T} \Theta \mathbf{L}'(s) - \Theta \mathbf{X}')^+ \quad \text{for all } 0 \leq t \leq T.\end{aligned}$$

Here the supremum is taken coordinate by coordinate. Since the elements of Q^{*T} are nonnegative, we have

$$\Theta(\mathbf{L}(t) - \mathbf{L}'(t)) \leq Q^{*T} \Theta \sup_{0 \leq s \leq t} |\mathbf{L}(s) - \mathbf{L}'(s)| + \sup_{0 \leq s \leq t} \Theta |\mathbf{X}(s) - \mathbf{X}'(s)|.$$

The inequality here also holds coordinate by coordinate. As Θ is a diagonal matrix with positive diagonal elements, we have

$$\begin{aligned}(\mathbf{L}(t) - \mathbf{L}'(t)) &\leq \Theta^{-1} Q^{*T} \Theta \sup_{0 \leq s \leq t} |\mathbf{L}(s) - \mathbf{L}'(s)| + \sup_{0 \leq s \leq t} |\mathbf{X}(s) - \mathbf{X}'(s)| \\ &= Q^T \sup_{0 \leq s \leq t} |\mathbf{L}(s) - \mathbf{L}'(s)| + \sup_{0 \leq s \leq t} |\mathbf{X}(s) - \mathbf{X}'(s)|.\end{aligned}$$

As a result,

$$\sup_{0 \leq s \leq T} |\mathbf{L}(s) - \mathbf{L}'(s)| \leq Q^T \sup_{0 \leq s \leq T} |\mathbf{L}(s) - \mathbf{L}'(s)| + \sup_{0 \leq s \leq T} |\mathbf{X}(s) - \mathbf{X}'(s)|.$$

Since $(I - Q^T)^{-1} = R^{-1}$ has nonnegative elements, we have

$$\sup_{0 \leq s \leq T} |\mathbf{L}(s) - \mathbf{L}'(s)| \leq R^{-1} \sup_{0 \leq s \leq T} |\mathbf{X}(s) - \mathbf{X}'(s)|.$$

In the end, we have

$$\begin{aligned}\sup_{0 \leq s \leq T} |\mathbf{Y}(s) - \mathbf{Y}'(s)| &\leq \sup_{0 \leq s \leq T} |\mathbf{X}(s) - \mathbf{X}'(s)| + |R| \sup_{0 \leq s \leq T} |\mathbf{L}(s) - \mathbf{L}'(s)| \\ &\leq \sup_{0 \leq s \leq T} |\mathbf{X}(s) - \mathbf{X}'(s)| + |R| R^{-1} \sup_{0 \leq s \leq T} |\mathbf{X}(s) - \mathbf{X}'(s)|.\end{aligned}$$

Let us denote R^{-1} by S , then $S_{ij} \geq 0$ for all $1 \leq i, j \leq d$. Based on the fact that $R_{ii} = 1$, $R_{ij} \leq 0$ for all $1 \leq i \neq j \leq d$ and $\sum_k R_{ik} S_{ki} = 1$ for all $1 \leq i \leq d$, we have

$$(|R|S)_{ii} = \sum_{k=1}^d |R_{ik}| S_{ki} = R_{ii} S_{ii} - \sum_{k \neq i} R_{ik} S_{ki} = R_{ii} S_{ii} + (-1 + R_{ii} S_{ii}) = 2S_{ii} - 1.$$

Note that $2S_{ii} - 1 > 0$ as the diagonal elements of R^{-1} are greater or equal to 1. Similarly, as $\sum_k R_{ik} S_{kj} = 0$ for all $1 \leq i \neq j \leq d$, we have

$$(|R|S)_{ij} = \sum_{k=1}^d |R_{ik}| S_{kj} = R_{ii} S_{ij} - \sum_{k \neq i} R_{ik} S_{kj} = R_{ii} S_{ij} + R_{ii} S_{ij} = 2S_{ij}.$$

Therefore, $|R|R^{-1} = 2R^{-1} - I$ where I is the identity matrix of dimension d , and we conclude

$$\begin{aligned}\sup_{0 \leq s \leq T} |\mathbf{Y}(s) - \mathbf{Y}'(s)| &\leq \sup_{0 \leq s \leq T} |\mathbf{X}(s) - \mathbf{X}'(s)| + |R| R^{-1} \sup_{0 \leq s \leq T} |\mathbf{X}(s) - \mathbf{X}'(s)| \\ &= 2R^{-1} \sup_{0 \leq s \leq T} |\mathbf{X}(s) - \mathbf{X}'(s)|.\end{aligned}$$

Recalling Assumption A1, $\|R^{-1}1\|_\infty \leq \kappa_0/\beta_0$, the desired result follows. $\square$

Given Lemma 2 and Lemma 3, we now are ready to provide an upper bound for the discretization error.

Lemma 4. For fixed $\gamma \in (0, 1)$, $t > 0$ and the number of dimensions d , there exists a positive constant C_1 such that

$$E[(\mathbf{Y}^m(t) - \mathbf{Y}(t))_\infty^2] \leq C_1 \gamma^m (\log(t) + \log(d) + m \log(1/\gamma)).$$

Proof of Lemma 4. By Lemma 3, we have

$$\begin{aligned} E[\|Y^m(t) - Y(t)\|_\infty^2] &\leq \frac{4\kappa_0^2}{\beta^2} E\left[\sup_{0 \leq s \leq t} \|X^m(s) - X(s)\|_\infty^2\right] \\ &= \frac{4\kappa_0^2}{\beta^2} E\left[\max_{1 \leq i \leq d} \max_{0 \leq s \leq t} |X_i^m(s) - X_i(s)|^2\right] \leq C_1 \gamma^m (\log(t) + \log(d) + m \log(1/\gamma)), \end{aligned}$$

the last inequality following from Lemma 2 with $C_1 = C_0 \cdot \frac{4\kappa_0^2}{\beta^2}$. $\square$

5.2. Nonstationary Error Bound

The convergence rate to stationarity of RBM has been analyzed in Banerjee and Budhiraja (2020) and Blanchet and Chen (2020) based on the synchronous coupling technique. Here we provide an alternative method based on the derivative of RBM with respect to the initial condition. Intuitively, the nonstationary error should have the same order as this derivative, as it reflects the impact of the initial condition on the RBM.

To do this, we first introduce the directional derivative of RBM as defined in Mandelbaum and Ramanan (2010). For every continuous input $X_{0:t}$, initial condition $\mathbf{y}$ and $\mathbf{h} \in \mathbb{R}^d$, Mandelbaum and Ramanan (2010) show that there exists a process $\mathfrak{D}_{\mathbf{h}}(t; \mathbf{y}, X_{0:t}) \in \mathbb{R}^d$ such that

$$\mathfrak{D}_{\mathbf{h}}(t; \mathbf{y}, X_{0:t}) = \lim_{\varepsilon \downarrow 0} \frac{Y(t; \mathbf{y} + \varepsilon \mathbf{h}, X_{0:t}) - Y(t; \mathbf{y}, X_{0:t})}{\varepsilon},$$

where the limit is taken componentwise. We first show that the derivative process can be bounded by $|\mathbf{h}|$ and the product of a series of matrices. Following the notation introduced in section 3 of Blanchet and Chen (2020), for the RBM $Y(\cdot; \mathbf{y}, \mathbf{X})$ starting from position $\mathbf{y}$ at time 0, we define a series of stopping times: $\eta_i^0(\mathbf{y}) = 0$, and for integer $k \geq 1$,

$$\begin{aligned} \eta_i^k(\mathbf{y}) &= \inf\{t > \eta_i^{k-1}(\mathbf{y}) + 1 : Y_i(t; \mathbf{y}) = 0\}, \\ \eta^k(\mathbf{y}) &= \sup\{\eta_i^k(\mathbf{y}) : 1 \leq i \leq d\}, \end{aligned} \tag{8}$$

and set-valued functions

$$\Gamma_i(t, \mathbf{y}) = \{\eta_i^k(\mathbf{y}) : \eta_i^k(\mathbf{y}) \leq t\}, \text{ and } \Gamma(t, \mathbf{y}) = \bigcup_{i=1}^d \Gamma_i(t, \mathbf{y}).$$

For any time point $t \geq 0$, define

$$\mathcal{C}(t) = \{1 \leq i \leq d : Y_i(t) = 0\} \text{ and } \bar{\mathcal{C}}(t) = \{1 \leq j \leq d : j \notin \mathcal{C}(t)\}.$$

Then, we introduce an auxiliary Markov chain $(W(n) : n \geq 0)$ living on the state-space $\{0, 1, \dots, d\}$ so that $P(W(n+1) = j | W(n) = i) = Q_{i,j}$ with $Q_{i,0} = 1 - \sum_{j=1}^d Q_{i,j}$ for $1 \leq i \leq d$, and $Q_{0,0} = 1$. We use P_i to refer to the probability law of $\{W(n) : n \geq 0\}$ given that $W(0) = i$. For any subset S of $\{1, 2, \dots, d\}$, we write

$$\begin{aligned} \tau(S) &= \inf\{n \geq 0 : W(n) \in S\}, \text{ and} \\ \tau(\{0\}) &= \inf\{n \geq 0 : W(n) = 0\}, \end{aligned}$$

and define the $d \times d$ matrix $\Lambda(S)$ as

$$\Lambda_{i,j}(S) = P_i(\tau(S) < \tau(\{0\}), W(\tau(S)) = j) \text{ for } i, j \in \{1, \dots, d\}.$$

The following result provides an explicit bound for the derivative matrix in terms of the product of Λ matrices.

Lemma 5. For any $\mathbf{h} \in \mathbb{R}_+^d$, the derivative process

$$\mathfrak{D}_{\mathbf{h}}(t; \mathbf{y}, X_{0:t}) \leq R^{-1} \prod_{s \in \Gamma(t, \mathbf{y})} \Lambda^T(\bar{\mathcal{C}}(s)) \cdot \mathbf{h},$$

where the inequality holds componentwise.

Remark 3. Under the uniformity assumptions, $(R^{-1}1(\leq b_1)$. As a result, for any $\mathbf{h} \in \mathbb{R}_+^d$,

$$\|\mathfrak{D}_{\mathbf{h}}(t; \mathbf{y}, X_{0:t})\|_1 \leq b_1 \|\prod_{s \in \Gamma(t, \mathbf{y})} \Lambda^T(\bar{\mathcal{C}}(s))\|_\infty \|\mathbf{h}\|_1.$$

Proof of Lemma 5. For simplicity of notation, we shall write $\mathfrak{D}_{\mathbf{h}}(t; \mathbf{y}, \mathbf{X}_{0:t}) = \mathfrak{D}_{\mathbf{h}}(t)$ and define $\gamma(t) = R^{-1}(\mathfrak{D}_{\mathbf{h}}(t) - I)$, that is, $\gamma(t)$ is the directional derivative of $\mathbf{L}(t)$ with respect to initial value $\mathbf{y}$ in the direction $\mathbf{h}$ (see Mandelbaum and Ramanan 2010).

According to theorem 1.1 of Kella and Ramasubramanian (2012), the process $R^{-1}(\mathbf{Y}(t, \mathbf{y}_1, \mathbf{X}_{0:t}) - \mathbf{Y}(t, \mathbf{y}_2, \mathbf{X}_{0:t}))$ is nonincreasing in t , for any $\mathbf{y}_1 \geq \mathbf{y}_2$. As a direct consequence, we can conclude that $\gamma(t)$ is nonincreasing in t component by component.

Suppose $\Gamma(t, \mathbf{y}) = \{\tau_1, \tau_2, \dots\}$ with $\tau_1 < \tau_2 < \dots$ in order. We define

$$D_n = \prod_{k \leq n} \Lambda^T(\bar{\mathcal{C}}(\tau_k)), \text{ and } \gamma_n = R^{-1}(D_n - I) \cdot \mathbf{h}.$$

In particular, $D_0 = I$ and $\gamma_0 = 0$. We shall prove by induction that for any $\tau_n < t \leq \tau_{n+1}$,

$$\gamma(t) \leq \gamma_n \text{ and hence } R^{-1}\mathfrak{D}_{\mathbf{h}}(t) \leq R^{-1} \prod_{k \leq N(t)} \Lambda^T(\bar{\mathcal{C}}(\tau_k)) \cdot \mathbf{h}, \quad (9)$$

where $N(t) = \sup\{k : \tau_k < t\}$. First, when $t < \tau_1$, by definition $\gamma(t) = \gamma_0 = 0$. Now suppose (9) holds for all $n \leq m-1$ and we consider a fixed time $\tau_m < t \leq \tau_{m+1}$. According to Mandelbaum and Ramanan (2010), the derivative processes $\gamma(t)$ is the unique solution to the following system of equations:

$$\gamma_i(t) = \sup_{s \in \Phi^i(t)} [-h_i + (P\gamma(s))_i],$$

where $\Phi^i(t) = \{s \leq t : L_i(s) = L_i(t)\}$ and $P = I - R \geq 0$. For any $i \in \mathcal{C}(\tau_m)$, $L_i(t) > L_i(\tau_m)$ with probability 1. By the fact that $\gamma(t)$ is nonincreasing in t and $P \geq 0$, we have

$$\gamma_i(t) \leq -h_i + (P\gamma(\tau_m))_i \leq -h_i + (P\gamma_{m-1})_i,$$

where the last inequality holds by the induction assumption. For any $i \in \bar{\mathcal{C}}(\tau_m)$, we have $\gamma_i(t) \leq \gamma_i(\tau_m) \leq \gamma_{m-1,i}$, where $\gamma_{m-1,i}$ denotes the i th element of γ_{m-1} . Suppose $\bar{\gamma}$ is the solution to the following systems of linear equations:

$$\bar{\gamma}_i = \begin{cases} -h_i + (P\gamma_{m-1})_i & \text{if } i \in \mathcal{C}(\tau_m), \\ \gamma_{m-1,i} & \text{if } i \in \bar{\mathcal{C}}(\tau_m). \end{cases}$$

Then, $\gamma(t) \leq \bar{\gamma}$ component by component. For the simplicity of notation, we write $\mathcal{C} = \mathcal{C}(\tau_m)$. Then, $\bar{\gamma}_i$ can be solved explicitly as

$$\bar{\gamma}_{\bar{\mathcal{C}}} = \gamma_{m-1, \bar{\mathcal{C}}}; \bar{\gamma}_{\mathcal{C}} = R_{\mathcal{C}\mathcal{C}}^{-1}(-\mathbf{h}_{\mathcal{C}} + P_{\mathcal{C}\bar{\mathcal{C}}}\gamma_{m-1, \bar{\mathcal{C}}}).$$

More precisely, we write

$$\bar{\gamma} = \begin{pmatrix} -R_{\mathcal{C}\mathcal{C}}^{-1}I_{\mathcal{C}} \\ 0 \end{pmatrix} \mathbf{h} + \begin{pmatrix} 0 & R_{\mathcal{C}\mathcal{C}}^{-1}P_{\mathcal{C}\bar{\mathcal{C}}} \\ 0 & I_{\bar{\mathcal{C}}} \end{pmatrix} \gamma_{m-1}.$$

One can check that

$$\Lambda_m^T \triangleq \Lambda^T(\bar{\mathcal{C}}(\tau_m)) = I + R \begin{pmatrix} -R_{\mathcal{C}\mathcal{C}}^{-1}I_{\mathcal{C}} \\ 0 \end{pmatrix},$$

and

$$\begin{aligned} R^{-1}\Lambda_m^T R &= R^{-1} \left[I + R \begin{pmatrix} -R_{\mathcal{C}\mathcal{C}}^{-1}I_{\mathcal{C}} \\ 0 \end{pmatrix} \right] R = I + \begin{pmatrix} -R_{\mathcal{C}\mathcal{C}}^{-1}I_{\mathcal{C}} & 0 \\ 0 & 0 \end{pmatrix} \begin{pmatrix} R_{\mathcal{C}\mathcal{C}} & R_{\mathcal{C}\bar{\mathcal{C}}} \\ R_{\bar{\mathcal{C}}\mathcal{C}} & R_{\bar{\mathcal{C}}\bar{\mathcal{C}}} \end{pmatrix} \\ &= I + \begin{pmatrix} -I_{\mathcal{C}} & -R_{\mathcal{C}\mathcal{C}}^{-1}R_{\mathcal{C}\bar{\mathcal{C}}} \\ 0 & 0 \end{pmatrix} = \begin{pmatrix} 0 & R_{\mathcal{C}\mathcal{C}}^{-1}P_{\mathcal{C}\bar{\mathcal{C}}} \\ 0 & I_{\bar{\mathcal{C}}} \end{pmatrix}. \end{aligned}$$

Therefore, we have

$$\begin{aligned}\bar{\gamma} &= R^{-1}(\Lambda_m^T - I)\mathbf{h} + R^{-1}\Lambda_m^T R\gamma_{m-1} \\ &= R^{-1}(\Lambda_m^T - I)\mathbf{h} + R^{-1}\Lambda_m^T R \cdot R^{-1}\left(\prod_{k \leq m-1} \Lambda_k^T - I\right)\mathbf{h} \\ &= R^{-1}\left(\prod_{k \leq m} \Lambda_k^T - I\right)\mathbf{h} = \gamma_m.\end{aligned}$$

As a result, (9) holds by induction and we have

$$R^{-1}\mathfrak{D}_{\mathbf{h}}(t) \leq R^{-1} \prod_{k \leq N(t)} \Lambda^T(\bar{\mathcal{C}}(\tau_k)) \cdot \mathbf{h}.$$

Since all the components of R^{-1} are nonnegative and all its diagonal entries are greater or equal to 1, we can conclude that, component by component

$$\mathfrak{D}_{\mathbf{h}}(t) \leq R^{-1} \prod_{k \leq N(t)} \Lambda^T(\bar{\mathcal{C}}(\tau_k)) \cdot \mathbf{h}. \quad \square$$

Remark 4. During the revision of the paper, we learned that Lipshutz and Ramanan (2021, lemma 7.5) produce a similar result that the derivative with respect to the initial condition contracts when the RBM hits all the faces. Although the result in Lipshutz and Ramanan (2021) holds in a more general setting, that is, RBM in a convex polyhedral cone, a quantitative bound on the contraction size is not directly provided in Lipshutz and Ramanan (2021). In contrast, Lemma 5 provides an explicit expression of the contraction whenever RBM hits a face, and this is necessary for high-dimensional analysis.

Now, we are ready to derive the upper bound for the error due to nonstationarity.

Lemma 6. *There exist constants C_2 and $\xi_1 > 0$ such that*

$$E[\|\mathbf{Y}(t; \mathbf{Y}(\infty), \mathbf{X}_{0:t}) - \mathbf{Y}(t; 0, \mathbf{X}_{0:t})\|_{\infty}^2] \leq C_2 d^3 \exp\left(-\xi_1 \frac{t}{\log(d)}\right).$$

Proof of Lemma 6. By the definition of directional derivative of RBM, for any $\mathbf{y} \in \mathbb{R}_+^d$,

$$\mathbf{Y}(t; \mathbf{y}, \mathbf{X}_{0:t}) - \mathbf{Y}(t; 0, \mathbf{X}_{0:t}) = \int_0^1 \mathfrak{D}_{\mathbf{y}}(t; u \cdot \mathbf{y}, \mathbf{X}_{0:t}) du.$$

Then, following Lemma 5,

$$\|\mathbf{Y}(t; \mathbf{y}, \mathbf{X}_{0:t}) - \mathbf{Y}(t; 0, \mathbf{X}_{0:t})\|_{\infty} \leq b_1 \int_0^1 \left\| \prod_{s \in \Gamma(t, u \cdot \mathbf{y})} \Lambda^T(\bar{\mathcal{C}}(s)) \right\|_{\infty} du \cdot \|\mathbf{y}\|_1.$$

Let's denote $\left\| \prod_{s \in \Gamma(t, u \cdot \mathbf{y})} \Lambda^T(\bar{\mathcal{C}}(s)) \right\|_{\infty} = \Theta(u)$. Then we have

$$\|\mathbf{Y}(t; \mathbf{y}, \mathbf{X}_{0:t}) - \mathbf{Y}(t; 0, \mathbf{X}_{0:t})\|_{\infty}^2 \leq b_1^2 \|\mathbf{y}\|_1^2 \left(\int_0^1 \Theta(u) du \right)^2 \leq b_1^2 \|\mathbf{y}\|_1^2 \int_0^1 \Theta(u) du.$$

The last equality holds as $\Theta(u) \leq 1$ for all $0 \leq u \leq 1$.

The rest of proof follows the same argument as in Banerjee and Budhiraja (2020). By lemma 2 and lemma 3 of Blanchet and Chen (2020), all $0 \leq u \leq 1$,

$$\Theta(u) \leq \|Q^{\mathcal{N}(t, u \cdot \mathbf{y})} \mathbf{1}\|_{\infty},$$

where $\mathcal{N}(t, \mathbf{y}) = \sup\{k \geq 0 : \eta^k(\mathbf{y}) \leq t\}$. Then, we have

$$\begin{aligned}E[\|\mathbf{Y}(t; \mathbf{Y}(\infty), \mathbf{X}_{0:t}) - \mathbf{Y}(t; 0, \mathbf{X}_{0:t})\|_{\infty}^2] &\leq b_1^2 E[\|\mathbf{Y}(\infty)\|_1^2 \|Q^{\mathcal{N}(t, \mathbf{Y}(\infty))} \mathbf{1}\|_{\infty}^2] \\ &\leq b_1^2 E[\|\mathbf{Y}(\infty)\|_1^4]^{1/2} E[\|Q^{\mathcal{N}(t, \mathbf{Y}(\infty))} \mathbf{1}\|_{\infty}^2]^{1/2} \leq b_1^2 E[\|\mathbf{Y}(\infty)\|_1^4]^{1/2} E[\|Q^{\mathcal{N}(t, \mathbf{Y}(\infty))} \mathbf{1}\|_{\infty}^2]^{1/2}.\end{aligned}$$

The proof of theorem 1 of Banerjee and Budhiraja (2020, p. 20) shows that, under Assumptions A1 to A3,

$$E[\|\mathbf{Y}(\infty)\|_1^4]^{1/2} \leq \frac{4b_0^2}{\delta_0^2} d^2,$$

$$E[\|Q^{N(t, \mathbf{Y}(\infty))} \mathbf{1}\|_\infty]^{1/2} \leq C_0 d \left(\exp\left(-\xi_1 \frac{t}{\log(d)}\right) \right).$$

Therefore, letting $C_2 = \frac{4b_0^2 b^2}{\delta_0^2} C_0$, we get

$$E[\|\mathbf{Y}(t; \mathbf{Y}(\infty), \mathbf{X}_{0:t}) - \mathbf{Y}(t; 0, \mathbf{X}_{0:t})\|_\infty^2] \leq C_2 d^3 \exp\left(-\xi_1 \frac{t}{\log(d)}\right). \quad \square$$

5.3. Complexity Analysis

Given the error bounds Lemma 4 and Lemma 6, we are ready to show that, for the two-parameter multilevel Monte Carlo Algorithm 1, the computational budget to obtain an estimator at a fixed accuracy level is almost linear in dimension d .

Proof of Theorem 1. Recall that for a given sequence of RBMs and a given function f to be evaluated, the algorithm has five input parameters, $(\gamma, T, L, \mathbf{y}_0, N)$. In the following analysis, we choose $\mathbf{y}_0 = 0$.

For fixed d and ε , the mean square error of the estimator $\bar{Z}$ can be expressed as

$$\begin{aligned} E[(\bar{Z} - E[f(\mathbf{Y}(\infty))])^2] &= \text{Var}[\bar{Z}] + (E[\bar{Z}] - E[f(\mathbf{Y}(\infty))])^2 \\ &\leq \frac{1}{N} E[(Z - f(\mathbf{y}_0))^2] + (E[Z] - E[f(\mathbf{Y}(\infty))])^2 \\ &= \frac{1}{N} \sum_{m=0}^{L-1} p(m)^{-1} E\left[\left(f\left(\mathbf{Y}^{m+1}\left((m+1)T; \mathbf{y}_0, \mathbf{X}_{0:(m+1)T}^{m+1}\right)\right) - f\left(\mathbf{Y}^m\left(mT; \mathbf{y}_0, \mathbf{X}_{T:(m+1)T}^m\right)\right)\right)^2\right] \\ &\quad + (E[Z] - E[f(\mathbf{Y}(\infty))])^2 \\ &\triangleq \frac{1}{N} \sum_{m=0}^{L-1} K(\gamma)^{-1} \gamma^{-m} V_m + \text{Bias}^2. \end{aligned} \quad (10)$$

We first analyze the variance terms V_m for each $m = 0, 1, \dots, L-1$. Following Assumption A4,

$$\begin{aligned} &f\left(\mathbf{Y}^{m+1}\left((m+1)T; \mathbf{y}_0, \mathbf{X}_{0:(m+1)T}^{m+1}\right)\right) - f\left(\mathbf{Y}^m\left(mT; \mathbf{y}_0, \mathbf{X}_{T:(m+1)T}^m\right)\right) \\ &\leq \mathcal{L} \|\mathbf{Y}^{m+1}\left((m+1)T; \mathbf{y}_0, \mathbf{X}_{0:(m+1)T}^{m+1}\right) - \mathbf{Y}^m\left(mT; \mathbf{y}_0, \mathbf{X}_{T:(m+1)T}^m\right)\|_\infty, \end{aligned}$$

and

$$\|\mathbf{Y}^{m+1}\left((m+1)T; \mathbf{y}_0, \mathbf{X}_{0:(m+1)T}^{m+1}\right) - \mathbf{Y}^m\left(mT; \mathbf{y}_0, \mathbf{X}_{T:(m+1)T}^m\right)\|_\infty \quad (11)$$

$$\begin{aligned} &\leq \|\mathbf{Y}^{m+1}\left((m+1)T; \mathbf{y}_0, \mathbf{X}_{0:(m+1)T}^{m+1}\right) - \mathbf{Y}\left((m+1)T; \mathbf{y}_0, \mathbf{X}_{0:(m+1)T}\right)\|_\infty \\ &\quad + \|\mathbf{Y}^m\left(mT; \mathbf{y}_0, \mathbf{X}_{T:(m+1)T}^m\right) - \mathbf{Y}\left(mT; \mathbf{y}_0, \mathbf{X}_{T:(m+1)T}\right)\|_\infty \end{aligned} \quad (12)$$

$$+ \|\mathbf{Y}\left((m+1)T; \mathbf{y}_0, \mathbf{X}_{0:(m+1)T}\right) - \mathbf{Y}\left(mT; \mathbf{y}_0, \mathbf{X}_{T:(m+1)T}\right)\|_\infty. \quad (13)$$

For (11) and (12), by following Lemma 4, we have

$$\begin{aligned} &E[\|\mathbf{Y}^{m+1}\left((m+1)T; \mathbf{y}_0, \mathbf{X}_{0:(m+1)T}^{m+1}\right) - \mathbf{Y}\left((m+1)T; \mathbf{y}_0, \mathbf{X}_{0:(m+1)T}\right)\|_\infty^2] \\ &\leq C_1 \gamma^{m+1} (\log((m+1)T) + \log(d) + (m+1)\log(1/\gamma)), \end{aligned}$$

and

$$\begin{aligned} &E[\|\mathbf{Y}^m\left(mT; \mathbf{y}_0, \mathbf{X}_{T:(m+1)T}^m\right) - \mathbf{Y}\left(mT; \mathbf{y}_0, \mathbf{X}_{T:(m+1)T}\right)\|_\infty^2] \\ &\leq C_1 \gamma^m (\log(mT) + \log(d) + m\log(1/\gamma)). \end{aligned}$$

For (13), by Kella and Ramasubramanian (2012, theorem 1.1), we have

$$\mathbf{Y}(T; \mathbf{Y}(\infty), \mathbf{X}_{0:T}) \geq \mathbf{Y}(T; 0, \mathbf{X}_{0:T}) \geq 0,$$

and by using Kella and Ramasubramanian (2012, theorem 1.1) again, we have

$$\begin{aligned} \mathbf{Y}(mT; \mathbf{Y}(T; \mathbf{Y}(\infty), \mathbf{X}_{0:T}), \mathbf{X}_{T:(m+1)T}) &\geq \mathbf{Y}(mT; \mathbf{Y}(T; 0, \mathbf{X}_{0:T}), \mathbf{X}_{T:(m+1)T}) \\ &\geq \mathbf{Y}(mT; 0, \mathbf{X}_{T:(m+1)T}) \end{aligned}$$

Therefore, by following Lemma 6 and $\mathbf{y}_0 = 0$, we have

$$\begin{aligned} &E\left[\|\mathbf{Y}((m+1)T; \mathbf{y}_0, \mathbf{X}_{0:(m+1)T}) - \mathbf{Y}(mT; \mathbf{y}_0, \mathbf{X}_{T:(m+1)T})\|_\infty^2\right] \\ &= E\left[\|\mathbf{Y}(mT; \mathbf{Y}(T; 0, \mathbf{X}_{0:T}), \mathbf{X}_{T:(m+1)T}) - \mathbf{Y}(mT; 0, \mathbf{X}_{T:(m+1)T})\|_\infty^2\right] \\ &\leq E\left[\|\mathbf{Y}(mT; \mathbf{Y}(mT; \mathbf{Y}(T; \mathbf{Y}(\infty), \mathbf{X}_{0:T}), \mathbf{X}_{T:(m+1)T}) - \mathbf{Y}(mT; 0, \mathbf{X}_{T:(m+1)T})\|_\infty^2\right] \\ &= E\left[\|\mathbf{Y}(mT; \mathbf{Y}(\infty), \mathbf{X}_{T:(m+1)T}) - \mathbf{Y}(mT; 0, \mathbf{X}_{T:(m+1)T})\|_\infty^2\right] \\ &\leq C_2 \cdot d^3 \exp\left(-\xi_1 \frac{mT}{\log d}\right). \end{aligned}$$

Recalling that $(a + b + c)^2 \leq 3(a^2 + b^2 + c^2)$, we therefore have that

$$V_m \leq 3\mathcal{L}^2 \left(2C_1 \gamma^m (\log((m+1)T) + \log(d) + (m+1)\log(1/\gamma)) + C_2 \cdot d^3 \exp\left(-\xi_1 \frac{mT}{\log d}\right) \right).$$

Let $T = \lceil (3\log(d)^2 + \log(1/\gamma)\log(d)) / \xi_1 \rceil$ and $C_3 = 3\mathcal{L}^2(2C_1 + C_2)$. We have

$$\begin{aligned} V_m &\leq 3\mathcal{L}^2(2C_1 \gamma^m (\log((m+1)T) + \log(d) + (m+1)\log(1/\gamma)) + C_2 \gamma^m) \\ &\leq C_3 \gamma^m (\log((m+1)T) + \log(d) + (m+1)\log(1/\gamma)). \end{aligned}$$

Therefore, the total variance of our estimator is

$$\begin{aligned} V_{total} &= \frac{1}{N} \sum_{m=0}^{L-1} K(\gamma)^{-1} \gamma^{-m} V_m \\ &\leq \frac{1}{N} K(\gamma)^{-1} \sum_{m=0}^{L-1} C_3 (\log((m+1)T) + \log(d) + (m+1)\log(1/\gamma)) \\ &\leq \frac{1}{N} C_3 K(\gamma)^{-1} L (\log(LT) + \log(d) + L\log(1/\gamma)). \end{aligned} \tag{14}$$

Now we turn to the term of bias in (10). Following Assumption A4, we have

$$\begin{aligned} &\text{Bias}^2 \\ &= (E[Z] - E[f(\mathbf{Y}(\infty))])^2 = (E[f(\mathbf{Y}^L(TL; \mathbf{y}_0, \mathbf{X}_{0:LT}^L)) - f(\mathbf{Y}(\infty))])^2 \\ &\leq 2\left((E[f(\mathbf{Y}^L(TL; \mathbf{y}_0, \mathbf{X}_{0:LT}^L)) - f(\mathbf{Y}(TL; \mathbf{y}_0, \mathbf{X}_{0:LT}))])^2 + (E[f(\mathbf{Y}(TL; \mathbf{y}_0, \mathbf{X}_{0:LT})) - f(\mathbf{Y}(\infty))])^2\right) \\ &\leq 2\left(E\left[\left(f(\mathbf{Y}^L(TL; \mathbf{y}_0, \mathbf{X}_{0:LT}^L)) - f(\mathbf{Y}(TL; \mathbf{y}_0, \mathbf{X}_{0:LT}))\right)^2\right] + E\left[\left(f(\mathbf{Y}(TL; \mathbf{y}_0, \mathbf{X}_{0:LT})) - f(\mathbf{Y}(TL; \mathbf{Y}(\infty), \mathbf{X}_{0:LT}))\right)^2\right]\right) \\ &\leq 2\mathcal{L}^2\left(E[\|\mathbf{Y}^L(TL; \mathbf{y}_0, \mathbf{X}_{0:LT}^L) - \mathbf{Y}(TL; \mathbf{y}_0, \mathbf{X}_{0:LT})\| + E[\|\mathbf{Y}(TL; \mathbf{y}_0, \mathbf{X}_{0:LT}) - \mathbf{Y}(TL; \mathbf{Y}(\infty), \mathbf{X}_{0:LT})\|_\infty^2]]\right). \end{aligned}$$

Following Lemma 4, we have

$$E[\|\mathbf{Y}^L(TL; \mathbf{y}_0, \mathbf{X}_{0:LT}) - \mathbf{Y}(TL; \mathbf{y}_0, \mathbf{X}_{0:LT})\|_\infty^2] \leq C_1 \left(\gamma^L (\log(LT) + \log(d) + L\log(1/\gamma)) \right).$$

Following Lemma 6, we have

$$E[\|\mathbf{Y}(TL; \mathbf{y}_0, \mathbf{X}_{0:LT}) - \mathbf{Y}(TL; \mathbf{Y}(\infty), \mathbf{X}_{0:LT})\|_\infty^2] \leq C_2 \cdot d^3 \exp\left(-\xi_1 \frac{LT}{\log(d)}\right) \leq C_2 \cdot \gamma^L,$$

for $T = \lceil (3\log(d)^2 + \log(1/\gamma)\log(d)) / \xi_1 \rceil$.

Therefore, we have

$$\begin{aligned} \text{Bias}^2 &\leq C_3 \left(\gamma^L (\log(LT) + \log(d) + L \log(1/\gamma)) \right) \leq \varepsilon^2/2, \text{ for} \\ T &= \lceil (3 \log(d)^2 + \log(1/\gamma) \log(d)) / \xi_1 \rceil, \text{ and} \\ L &= \lceil (\log(\log(d)) + 2 \log(1/\varepsilon) + k_1) / \log(1/\gamma) \rceil, \end{aligned}$$

where k_1 is a numerical constant.

To equalize the variance and bias of our estimator, we enforce

$$C_3 \left(\gamma^L (\log(LT) + \log(d) + L \log(1/\gamma)) \right) = \frac{1}{N} C_3 K(\gamma)^{-1} L (\log(LT) + \log(d) + L \log(1/\gamma)). \quad (15)$$

Hence, $N = \lceil K(\gamma)^{-1} L / \gamma^L \rceil = O(\varepsilon^{-2} K(\gamma)^{-1} L \log(d))$. Then, by plugging the choice of T, L, N in (14), we have $V_{\text{total}} \leq \varepsilon^2/2$.

Note that the complexity, in terms of expected number of scalar Gaussian random variables generated to simulate one sample of Z , should be

$$\mathcal{C} = \sum_{m=0}^{L-1} p(m) \gamma^{-(m+1)} T(m+1) d = \frac{1}{2} K(\gamma) \gamma^{-1} d T L (L+1).$$

Then, the total complexity to compute $\bar{Z}$ by N numbers of samples, with our choice of (γ, T, L, N) , is

$$\begin{aligned} N \times \mathcal{C} &= O(\varepsilon^{-2} K(\gamma)^{-1} L \log(d)) \times \left(\frac{1}{2} K(\gamma) \gamma^{-1} d T L (L+1) \right) \\ &= O(\varepsilon^{-2} d T \log(d) L^3) = O(\varepsilon^{-2} d \log(d)^3 (\log(\log(d)) + \log(1/\varepsilon))^3). \quad \square \end{aligned} \quad (16)$$

Lemma 7. The optimal $\gamma^* = 0.05$.

Proof of Lemma 7. According to (16), we have the dependence of the total complexity on γ is approximately $\gamma^{-1} (\log(1/\gamma))^{-3}$. Optimizing γ to obtain the optimal complexity, we find that the optimal γ is

$$\gamma^* = \arg \min_{0 < \gamma < 1} \gamma^{-1} (\log(1/\gamma))^{-3} = 0.05. \quad \square$$

5.4. Proof of the Lower Bound in Theorem 2

Proof of Theorem 2. The fact that Algorithm 1 belongs to the specified class is direct from its construction. Now, let π be any algorithm in the class. It suffices to show that each dimension must be sampled in the algorithm. If this is not the case, without loss of generality, we may assume that the algorithm does not sample any observation in the first coordinate. Consider the case $R = \Sigma = I$ and $f(\mathbf{y}) = \mathcal{L} \mathbf{y}_1$. We consider two different specifications of μ : $\mu^{(1)} = -[2\delta_0, 1, \dots, 1]^T$ for RBM1 and $\mu^{(2)} = -[4\delta_0, 1, \dots, 1]^T$ for RBM2. Then, since the dimensions are mutually independent, computing the steady-state distribution of the RBM reduces to studying one-dimensional RBMs. It is well known (see, for example, Harrison 1985) that the stationary distribution is exponential distribution with mean $-1/(2\mu_i)$. Specifically, $E[f(\mathbf{Y}^{(1)}(\infty))] = \mathcal{L}/(4\delta_0)$ and $E[f(\mathbf{Y}^{(2)}(\infty))] = \mathcal{L}/(8\delta_0)$.

Since the algorithm never observes the first coordinate, the algorithm should give an identical output for RBM1 and RBM2. Suppose the output is the random variable $\tilde{Z}$. The sum of MSEs of RBM1 and RBM2 becomes

$$E \left[\left(\frac{\mathcal{L}}{4\delta_0} - \tilde{Z} \right)^2 \right] + E \left[\left(\frac{\mathcal{L}}{8\delta_0} - \tilde{Z} \right)^2 \right] \geq \frac{\mathcal{L}^2}{128\delta_0^2}.$$

Therefore, the target error of order MSE $\varepsilon^2 < (\mathcal{L}/(16\delta_0))^2$ cannot be achieved simultaneously for both RBMs. $\square$

6. Conclusion

We have presented and analyzed a Monte Carlo algorithm that constructs asymptotically optimal estimators for steady-state expectations of high-dimensional RBM. We believe that the strategy that we present can be applied

to more general networks. A key idea is to consider the so-called synchronous coupling in combination with multilevel Monte Carlo. Although this idea is not new (see, for example, Glynn and Rhee 2014), the analysis, which is based on the rate of decay to zero of the product of substochastic random matrices is, we believe, applicable to other settings. In particular, the sensitivity to the initial condition in every stochastic flow naturally leads to the study of products of random matrices and the analysis of the so-called top Lyapunov exponent. In this paper, we are able to use implicit estimates for this product from Banerjee and Budhiraja (2020) and Blanchet and Chen (2020). This, we expect, will provide a blueprint that can be used in other settings, as we expect to report in future research.

Acknowledgments

The authors thank the two referees and the editor for their careful reading and helpful comments.

Appendix A. Routine to Solve Skorokhod Problem in Algorithm 1

In step 5 of Algorithm 1, once the piecewise linear approximation is obtained for the underlying Brownian motion, we obtain the solution to the Skorokhod problem by solving, at each time step, a static linear complementarity problem (see, for example, Cottle et al. 2009). Since R is an M -matrix, we here provide a simple yet numerical stable algorithm to solve the linear complementarity problem in Algorithm A.1.

Algorithm A.1 (Algorithm for the Linear Complementarity Problem)

Input:

The reflection matrix: R ;

The initial vector: $\mathbf{x}$;

Output:

The solution of the linear complementarity problem: $\mathbf{y} \geq 0$, where $\mathbf{y} = \mathbf{x} + R\mathbf{L}$ for $\mathbf{L} \geq 0$.

1: Set $\epsilon = 10^{-8}$;

2: $\mathbf{y} = \mathbf{x}$;

3: **while** Exists $\mathbf{y}_i < -\epsilon$ **do**

4: Compute the set $B = \{i : \mathbf{y}_i < \epsilon\}$;

5: Compute $\mathbf{L}_B = -R_{B,B}^{-1}\mathbf{x}_B$;

6: Compute $\mathbf{y} = \mathbf{x} + R_{:,B} \times \mathbf{L}_B$;

return $\mathbf{y}$.

Appendix B. Lower Bound on Constant ξ_1

We also provide a lower bound for the constant ξ_1 , which is not given explicitly in either Banerjee and Budhiraja (2020) or Blanchet and Chen (2020). The lower bound is computed based on a worst-case analysis in Banerjee and Budhiraja (2020). We believe that it is far from tight, as shown in the numerical experiments in Section 4. We provide this, nevertheless, for completeness.

Lemma A.1. *The constant ξ_1 satisfies*

$$\xi_1 \geq D_1 \left(\frac{\log(2)}{\log(1 - \beta_0)^{-1}} + 1 \right)^{-1} \left(2 + \frac{\kappa_0^2 b_0}{\beta_0^2 \delta_0^2} \right)^{-1}$$

with $D_1 = 1/557065$.

Proof of Lemma A.1. Our ξ_1 is equivalent to E_2 as defined in theorem 3 in Banerjee and Budhiraja (2020), that is,

$$E_2 = D_1 \left(\frac{\log(2)}{\log(1 - \beta_0)^{-1}} + 1 \right)^{-1} \left(2 + \frac{\kappa_0^2 b_0}{\beta_0^2 \delta_0^2} \right)^{-1},$$

with $D_1 = \delta'/128$, $\delta' = (64C_1)^{-1}$ and $C_1 = C_0 = A_0 = 68$ according to lemmas 7 and 8 in Banerjee and Budhiraja (2020). $\square$

Endnote

¹ We recommend choosing step size γ around 0.05 (please check Lemma 7 for the reason), but our algorithm is not sensitive to the specific choice of γ .

References

Armony M, Israelit S, Mandelbaum A, Marmor YN, Tseytlin Y, Yom-Tov GB (2015) On patient flow in hospitals: A data-based queueing-science perspective. *Stochastic Systems* 5(1):146–194.

- Banerjee S, Budhiraja A (2020) Parameter and dimension dependence of convergence rates to stationarity for reflecting Brownian motions. *Ann. Appl. Probab.* 30(5):2005–2029.
- Blanchet J, Chen X (2015) Steady-state simulation of reflected Brownian motion and related stochastic networks. *Ann. Appl. Probab.* 25(6):3209–3250.
- Blanchet J, Chen X (2020) Rates of convergence to stationarity for reflected Brownian motion. *Math. Oper. Res.* 45(2):660–681.
- Cottle RW, Pang JS, Stone RE (2009) *The Linear Complementarity Problem* (Society for Industrial and Applied Mathematics, Philadelphia, PA).
- Dai JG, Harrison JM (1992) Reflected Brownian motion in an orthant: Numerical methods for steady-state analysis. *Ann. Appl. Probab.* 2:65–86.
- Giles MB (2008) Multilevel Monte Carlo path simulation. *Oper. Res.* 56(3):607–617.
- Giles MB (2015) Multilevel Monte Carlo methods. *Acta Numerica* 24:259–328.
- Giles MB, Majka M, Szpruch L, Vollmer S, Zygalkis K (2020) Multi-level Monte Carlo methods for the approximation of invariant measures of stochastic differential equations. *Statistics and Computing* 30:507–524.
- Glynn PW, Rhee CH (2014) Exact estimation for Markov chain equilibrium expectations. *J. Appl. Probab.* 51A:377–389.
- Harrison JM (1985) *Brownian Motion and Stochastic Flow Systems* (Wiley, New York).
- Harrison JM, Reiman MI (1981) Reflected Brownian motion on an orthant. *Ann. Appl. Probab.* 9(2):302–308.
- Harrison JM, Williams RJ (1987) Brownian models of open queueing networks with homogeneous customer populations. *Stochastics* 22(2):77–115.
- Karlin S, Taylor HE (1981) *A Second Course in Stochastic Processes* (Academic Press, New York).
- Kella O, Ramasubramanian S (2012) Asymptotic irrelevance of initial conditions for Skorokhod reflection mapping on the nonnegative orthant. *Math. Oper. Res.* 37(2):301–312.
- Kushner H (2013) *Heavy Traffic Analysis of Controlled Queueing and Communication Networks*, vol. 47 (Springer-Verlag, New York).
- Lipshutz D, Ramanan K (2021) Sensitivity analysis for the stationary distribution of reflected Brownian motion in a convex polyhedral cone. *Math. Oper. Res.* Forthcoming.
- Maguluri ST, Srikant R, Ying L (2012) Stochastic models of load balancing and scheduling in cloud computing clusters. *Proc. IEEE INFOCOM (IEEE)*, 702–710.
- Mandelbaum A, Ramanan K (2010) Directional derivatives of oblique reflection maps. *Math. Oper. Res.* 35(3):527–558.
- Rhee CH, Glynn PW (2015) Unbiased estimation with square root convergence for SDE models. *Oper. Res.* 63(5):1026–1043.
- Shen X, Chen H, Dai J, Dai W (2002) The finite element method for computing the stationary distribution of an SRBM in a hypercube with applications to finite buffer queueing networks. *Queueing Systems* 42(1):33–62.