

PREINTEGRATION VIA ACTIVE SUBSPACE*

SIFAN LIU[†] AND ART B. OWEN[†]

Abstract. Preintegration is an extension of conditional Monte Carlo to quasi-Monte Carlo and randomized quasi-Monte Carlo. Conditioning can reduce but not increase the variance in Monte Carlo. For quasi-Monte Carlo it can bring about improved regularity of the integrand with potentially greatly improved accuracy. We show theoretically that, just as in Monte Carlo, preintegration can reduce but not increase the variance when one uses scrambled net integration. Preintegration is ordinarily done by integrating out one of the input variables to a function. In the common case of a Gaussian integral one can also preintegrate over any linear combination of variables. For continuous functions that are differentiable almost everywhere, we propose to choose the linear combination by the first principal component in an active subspace decomposition. We show that the lead eigenvector in an active subspace decomposition is closely related to the vector that maximizes a computationally intractable criterion using a Sobol' index. A numerical example of Asian option pricing finds that this active subspace preintegration strategy is competitive with preintegrating the first principal component of the Brownian motion, which is known to be very effective. The new method outperforms others on some basket and rainbow options where there is no generally accepted counterpart to the principal components construction.

Key words. conditional Monte Carlo, option pricing, quasi-Monte Carlo, randomized quasi-Monte Carlo

MSC codes. 65C05, 65D30

DOI. 10.1137/22M1479129

1. Introduction. Preintegration [15] is a strategy for high dimensional numerical integration in which one variable is integrated out in a closed form (or by a very accurate quadrature rule) while the others are handled by quasi-Monte Carlo (QMC) sampling. This strategy has been used since the 1950s in Monte Carlo (MC) sampling [16, 45], where it is known as conditional Monte Carlo. It can reduce variance but not increase it. In the Markov chain Monte Carlo (MCMC) literature, such conditioning is called Rao-Blackwellization, although it does not generally bring the optimal variance reduction that results from the Rao-Blackwell theorem. In the MCMC setting it is possible for conditioning to increase variance [10]. See [40] for a survey of conditioning in MCMC.

The advantage of preintegration in QMC goes beyond the variance reduction that arises in MC. After preintegration, a d -dimensional integration problem with a discontinuity or a kink (discontinuity in the gradient) can be converted into a much smoother $(d - 1)$ -dimensional problem [15]. QMC exploits regularity of the integrand and this additional smoothness brings a benefit on top of the variance reduction that comes from conditioning. It is known that the resulting smoothness depends critically on a monotonicity property of the integrand with respect to the variable being integrated out [11, 14, 15]. Hoyt and Owen [19] give conditions where preintegration

*Received by the editors February 24, 2022; accepted for publication (in revised form) October 25, 2022; published electronically March 10, 2023.

<https://doi.org/10.1137/22M1479129>

Funding: The work of the authors was supported by National Science Foundation grant IIS-1837931. The work of the first author was also partially supported by the Stanford Data Science Scholars program.

[†]Department of Statistics, Stanford University, Stanford, CA 94305 USA (sfliu@stanford.edu, owen@stanford.edu).

reduces the mean dimension (that we will define later) of the integrand. It can reduce the mean dimension from proportional to $\sqrt{d}$ to $O(1)$ as $d \rightarrow \infty$ in a sequence of ridge functions with a discontinuity that the preintegration smooths out. He [17] studied the error rate of preintegration for scrambled nets applied to functions of the form $f(\mathbf{x}) = h(\mathbf{x})\mathbf{1}\{\phi(\mathbf{x}) \geq 0\}$ for a Gaussian variable $\mathbf{x}$. That work assumes that h and ϕ are smooth functions and ϕ is monotone in the variable to be preintegrated. Then preintegration has a smoothing effect that when combined with some boundary growth conditions yields a root mean squared error rate of $O(n^{-1+\varepsilon})$, where n is the number of samples. In some very early uses of preintegration in QMC, L'Ecuyer and Lemieux [23] integrate some variables out of a stochastic activity network problem to improve smoothness for applying a lattice rule.

There is presently very little guidance in the literature about which variable to preintegrate over, beyond the monotonicity condition recently studied in [11] and a remark for ridge functions in [19]. Many of the use cases for preintegration involve integration with respect to the multivariate Gaussian distribution, especially for problems arising in finance. In the Gaussian context, we have more choices for the variable to preintegrate over. In addition to preintegration over any one of the d coordinates, preintegration over any linear combination of the variables remains integration with respect to a univariate Gaussian variable. Our proposal is to preintegrate over a linear combination of variables chosen to optimize a measure of variable importance derived from active subspaces [5].

When sampling from a multivariate Gaussian distribution by QMC, even without preintegration, one must choose a square root of the covariance matrix by which to multiply some sampled scalar Gaussian variables. While the choice of square root does not affect MC error, it does affect the QMC error. There are numerous choices for that square root. One can sample via the principal component matrix decomposition as [2] and many others do. For integrands defined with respect to Brownian motions, one can use the Brownian bridge construction studied by [26]. These choices have some potential disadvantages. It is always possible that the integrand is little affected by the first principal component. In a pessimistic scenario, the integrand could depend only on a principal component that is orthogonal to the first one. This is a well-known pitfall in principal components regression [22]. In a related phenomenon, Papageorgiou [36] exhibits an integrand where QMC via the standard construction is more effective than via the Brownian bridge.

Not only might a principal component direction perform poorly, but the first principal component is not necessarily well defined. Although the problem may be initially defined in terms of a specific Gaussian distribution, by a change of variable we can rewrite our integral as an expectation with respect to another Gaussian distribution with a different covariance matrix that has a different first principal component. Or, if the problem is posed with a covariance equal to the d -dimensional identity matrix, then every unit vector (i.e., L^2 -norm is 1) is a first principal component direction.

Some proposed methods take account of the specific integrand while formulating a sampling strategy. These include stratifying in a direction chosen from exponential tilting [13], exploiting a linearization of the integrand at $d+1$ special points starting with the center of the domain [20], and a gradient principal component analysis (GPCA) algorithm [47] that we describe in more detail later.

The problem we consider is to compute an approximation to $\mu = \mathbb{E}(f(\mathbf{x}))$, where $\mathbf{x} \in \mathbb{R}^d$ has the spherical Gaussian distribution denoted by $\mathcal{N}(0, I)$ and f is a continuous function with a gradient almost everywhere that is square integrable. Let

$C = \mathbb{E}(\nabla f(\mathbf{x})\nabla f(\mathbf{x})^\top) \in \mathbb{R}^{d \times d}$. The r -dimensional active subspace [5] is the space spanned by the r leading (i.e., top) eigenvectors of C . We always rank the eigenvalues and corresponding eigenvectors in descending order. Active subspace methods have been generalized to vector-valued functions [49] and nonlinear dimension reduction [3]. For other uses of the matrix C in numerical computation, see the references in [6]. For $r = 1$, let θ be the leading eigenvector of C normalized to be a unit vector in L^2 -norm. We propose to preintegrate f over $\theta^\top \mathbf{x} \sim \mathcal{N}(0, 1)$. Preintegration can smooth out both kinks and jumps. Because the active subspace direction is derived for continuous integrands it provides a principled preintegration choice for integrands with kinks, but its motivation does not extend to integrands with jumps.

The eigendecomposition of C is an uncentered principal components decomposition of the gradients, also known as the GPCA. The GPCA method [47] also uses the eigendecomposition of C to define a matrix square root for a QMC sampling strategy to reduce effective dimension, but it involves no preintegration. The algorithm in [46] preintegrates the first variable x_1 out of f . Then it applies GPCA to the remaining $d - 1$ variables in order to find a suitable $(d - 1) \times (d - 1)$ matrix square root for the remaining Gaussian variables. They preintegrate over a coordinate variable, while we always preintegrate over the leading eigenvector which is not generally one of the d coordinates. All of the algorithms that involve C typically take a sample to estimate it.

This paper is organized as follows. Section 2 provides some background on randomized QMC (RQMC) and preintegration. Section 3 shows that preintegration never increases the variance of scrambled net integration, regardless of the smoothness of the integrand. Thus this well-known property of conditional MC, which does not always extend to conditional MCMC, does extend to conditional RQMC. In section 4, we describe using the unit vector θ which maximizes the Sobol' index for $\theta^\top \mathbf{x}$. This strategy is to find a matrix square root for which the first column maximizes the criterion from section 3. We show that this Sobol' index $\bar{\tau}_\theta^2$ is well defined in that it does not depend on how we parameterize the space orthogonal to θ . We apply active subspace preintegration to option pricing examples in section 5. These include an Asian call option, a basket option, and call on max and call on min options. Section 6 has our conclusions.

2. Background. In this section, we introduce some background of RQMC, preintegration, and active subspaces. First we introduce some notation.

For a positive integer d , we let $1:d = \{1, 2, \dots, d\}$. For a subset $u \subseteq 1:d$, we let $-u = 1:d \setminus u$. For an integer $j \in 1:d$, we use j to represent $\{j\}$ when the context is clear and $-j = 1:d \setminus \{j\}$. Let $|u|$ denote the cardinality of u . For $\mathbf{x} \in \mathbb{R}^d$, we let $\mathbf{x}_u$ be the $|u|$ -dimensional vector containing only the x_j with $j \in u$. For $\mathbf{x}, \mathbf{z} \in \mathbb{R}^d$ and $u \subseteq 1:d$, we let $\mathbf{x}_u : \mathbf{z}_{-u}$ be the d -dimensional vector, whose j th entry is x_j if $j \in u$ and z_j if $j \notin u$. We use $\mathbb{N} = \{1, 2, \dots\}$ for the natural numbers and $\mathbb{N}_0 = \mathbb{N} \cup \{0\}$. We denote the density and the cumulative distribution function (CDF) of the standard Gaussian distribution $\mathcal{N}(0, 1)$ as φ and Φ , respectively. We let Φ^{-1} denote the inverse CDF of $\mathcal{N}(0, 1)$. We also use φ to denote the density of the d -dimensional standard Gaussian distribution $\mathcal{N}(0, I_d)$, $\varphi(\mathbf{x}) = (2\pi)^{-d/2} \exp(-\|\mathbf{x}\|^2/2)$, where $\|\mathbf{x}\|$ always denotes the L^2 -norm unless otherwise stated. We use $\mathcal{N}(0, I)$ when the dimension of the random variable is clear from context. For a matrix $\Theta \in \mathbb{R}^{d \times d}$, we let $\Theta[u]$ denote the columns of Θ whose indices are in u . For an event A , we define the indicator function $\mathbf{1}_A$ to be 1 if A happens and 0 otherwise. When A contains many levels of subscripts and/or superscripts we use $\mathbf{1}\{A\}$ instead.

2.1. QMC and RQMC. QMC provides a way to estimate $\mu = \int_{[0,1]^d} f(\mathbf{x}) d\mathbf{x}$ with typically greater accuracy than can be done by MC, while sharing with MC the ability to handle larger dimensions d than can be handled well by classical quadrature methods such as those in [7]. The QMC estimate, like the MC one, takes the form $\hat{\mu}_n = (1/n) \sum_{i=0}^{n-1} f(\mathbf{x}_i)$, except that instead of $\mathbf{x}_i \stackrel{\text{iid}}{\sim} \mathbb{U}[0,1]^d$ the sample points are chosen strategically and deterministically to cover the unit cube more evenly as quantified by a discrepancy measure (see [4]) such as the star discrepancy D_n^* (defined in the supplementary material (supplement.pdf [local/web 411KB])).

The QMC methods we study are digital nets and sequences. To define them, for an integer base $b \geq 2$ let $E(\mathbf{k}, \mathbf{c}) = \prod_{j=1}^d [\frac{c_j}{b^{k_j}}, \frac{c_j+1}{b^{k_j}})$ for $\mathbf{k} = (k_1, \dots, k_d)$ and $\mathbf{c} = (c_1, \dots, c_d)$, where $k_j \in \mathbb{N}_0$ and $0 \leq c_j < b^{k_j}$. The sets $E(\mathbf{k}, \mathbf{c})$ are called elementary intervals in base b . They have volume $b^{-|\mathbf{k}|}$ where $|\mathbf{k}| = \sum_{j=1}^d k_j$. For integers $m \geq t \geq 0$, the points $\mathbf{x}_0, \dots, \mathbf{x}_{n-1}$ with $n = b^m$ are a (t, m, d) -net in base b if every elementary interval $E(\mathbf{k}, \mathbf{c})$ with $|\mathbf{k}| \leq m - t$ contains $b^{m-|\mathbf{k}|}$ of those points, which is exactly n times the volume of $E(\mathbf{k}, \mathbf{c})$. Here, t is known as the quality parameter and smaller values are better. It is not always possible to get $t = 0$. The power of digital nets is that the points $\mathbf{x}_i$ satisfy $\binom{m-t+d-1}{d-1}$ such stratifications simultaneously. They attain a star discrepancy [27] of $D_n^* = O((\log n)^{d-1}/n)$ after approximating each $[\mathbf{0}, \mathbf{a}]$ by sets $E(\mathbf{k}, \mathbf{c})$ [27, Theorem 4.10]. An infinite sequence $\{\mathbf{x}_i\}_{i \geq 0}$ is called a (t, d) -sequence if for all $k \geq 0$ and $m \geq t$ the points $\mathbf{x}_{kb^m}, \dots, \mathbf{x}_{(k+1)b^m-1}$ are a (t, m, d) -net in base b .

Randomization techniques can be applied so that one can estimate the error by multiple independent replicates. In RQMC, the points $\mathbf{x}_i \sim \mathbb{U}([0,1]^d)$ individually while collectively they form a (t, m, d) -net or sequence with probability 1. One randomization technique that has this property is nested uniform scrambling introduced by [29]. The estimate $\hat{\mu}_n$ taken over a scrambled (t, d) -sequence satisfies a strong law of large numbers if $f \in L^2[0,1]$ [35]. If $f \in L^2[0,1]^d$, then $\text{Var}(\hat{\mu}_n) = o(1/n)$, giving the method asymptotically unbounded efficiency versus MC which has variance σ^2/n for $\sigma^2 = \text{Var}(f(\mathbf{x}))$. For smooth enough f , $\text{Var}(\hat{\mu}_n) = O(n^{-3}(\log n)^{d-1})$ [31, 33] with the sharpest sufficient condition in [48]. Also there exists $\Gamma < \infty$ with $\text{Var}(\hat{\mu}_n) \leq \Gamma \sigma^2/n$ for all $f \in L^2[0,1]^d$ [32]. This bound involves no powers of $\log(n)$.

For more background about QMC and RQMC, see the supplement (SM1).

2.2. Preintegration. For $j \in 1:d$, $\int_{[0,1]^d} f(\mathbf{x}) d\mathbf{x} = \int_{[0,1]^{d-1}} \int_0^1 f(\mathbf{x}) dx_j d\mathbf{x}_{-j}$, which we can also write as $\mathbb{E}(f(\mathbf{x})) = \mathbb{E}(\mathbb{E}(f(\mathbf{x})|\mathbf{x}_{-j}))$ for $\mathbf{x} \sim \mathbb{U}[0,1]^d$. For $\mathbf{x} \in [0,1]^d$, define $g(\mathbf{x}) = g_j(\mathbf{x}) = \int_0^1 f(\mathbf{x}) dx_j$. It simplifies the presentation of some of our results, especially Theorem 3.2 on variance reduction, to keep g defined as above on $[0,1]^d$ even though it only depends on $\mathbf{x}_{-j}$. In preintegration,

$$\hat{\mu}_n = \hat{\mu}_{n,j} = \frac{1}{n} \sum_{i=0}^{n-1} g_j(\mathbf{x}_i)$$

which, as we noted in the introduction, is conditional MC except that we now use RQMC inputs.

Preintegration can bring some advantages for RQMC. The plain MC variance of $g(\mathbf{x})$ is no larger than that of $f(\mathbf{x})$ and is generally smaller unless f does not depend at all on x_j . Thus the bound $\Gamma \sigma^2/n$ is reduced. Moreover, the preintegrated integrand g can be much smoother than f and (R)QMC improves on MC by exploiting this additional smoothness. For example, it has been observed that for some option pricing integrands, preintegrating certain variables can remove the discontinuities in the integrand or its gradient [15, 46].

The integrands we consider here are defined with respect to a Gaussian random variable. We are interested in $\mu = \mathbb{E}(f(\mathbf{y}))$ for $\mathbf{y} \sim \mathcal{N}(\mathbf{0}, \Sigma)$ with a positive definite covariance $\Sigma \in \mathbb{R}^{d \times d}$. Letting $R_0 \in \mathbb{R}^{d \times d}$ with $R_0 R_0^\top = \Sigma$ we can write $\mu = \mathbb{E}(f(R_0 \mathbf{z}))$ for $\mathbf{z} \sim \mathcal{N}(\mathbf{0}, I)$. For an orthogonal matrix $Q \in \mathbb{R}^{d \times d}$ we also have $Q\mathbf{z} \sim \mathcal{N}(\mathbf{0}, I)$. Then taking $\mathbf{z} = \Phi^{-1}(\mathbf{x})$ componentwise leads us to the estimate $\hat{\mu} = \frac{1}{n} \sum_{i=0}^{n-1} f(R\Phi^{-1}(\mathbf{x}_i))$ with $R = R_0 Q$ for RQMC points $\mathbf{x}_i$. The choice of Q or equivalently R does not affect the MC variance of $\hat{\mu}$ but it can change the RQMC variance. We will consider some examples later. The mapping Φ^{-1} from $\mathbb{U}[0, 1]^d$ to $\mathcal{N}(\mathbf{0}, I)$ can be replaced by another one such as the Box–Muller transformation. The choice of transformation does not affect the MC variance but does affect the RQMC variance. Most researchers use Φ^{-1} but [28] advocates for Box–Muller.

When we are using preintegration for a problem defined with respect to an $\mathcal{N}(\mathbf{0}, \Sigma)$ random variable we must choose R and then the coordinate j over which to preintegrate. Our approach is to fix $j = 1$ while choosing R so that x_1 is the most important linear combination of $\mathbf{x}$ in an active subspace approximation as described in section 4.

2.3. The ANOVA decomposition. For $f \in L^2[0, 1]^d$ we can define an analysis of variance (ANOVA) decomposition from [9, 18, 42]. For details see [34, Appendix A.6]. This decomposition reads

$$f(\mathbf{x}) = \sum_{u \subseteq 1:d} f_u(\mathbf{x})$$

where f_u depends on $\mathbf{x}$ only through x_j with $j \in u$ and also $\int_0^1 f_u(\mathbf{x}) dx_j = 0$ whenever $j \in u$. The term $f_\emptyset$ is the constant function everywhere equal to $\mu = \int_{[0,1]^d} f(\mathbf{x}) d\mathbf{x}$. The decomposition is orthogonal in that $\int_{[0,1]^d} f_u(\mathbf{x}) f_v(\mathbf{x}) d\mathbf{x} = 0$ if $u \neq v$. To each effect f_u there corresponds a variance component

$$\sigma_u^2 = \text{Var}(f_u(\mathbf{x})) = \begin{cases} \int_{[0,1]^d} f_u(\mathbf{x})^2 d\mathbf{x}, & |u| > 0 \\ 0 & \text{else.} \end{cases}$$

The variance components sum to $\sigma^2 = \text{Var}(f(\mathbf{x}))$.

The formula for f_u can be written recursively as $f_\emptyset = \mu$ and then for $|u| > 0$,

$$f_u(\mathbf{x}) = \int_{[0,1]^{d-|u|}} \left(f(\mathbf{x}) - \sum_{v \subsetneq u} f_v(\mathbf{x}) \right) d\mathbf{x}_{-u}.$$

The integral above is an expectation with respect to independent random x_j for $j \notin u$. For Gaussian or indeed any other distribution on components x_j one replaces the integral by the analogous expectation. We will use the ANOVA decomposition below when describing how to choose a preintegration variable.

Sobol' indices [43] are derived from the ANOVA decomposition. For $u \subseteq 1:d$ these are

$$\tau_u^2 = \sum_{v \subseteq u} \sigma_v^2 \quad \text{and} \quad \bar{\tau}_u^2 = \sum_{v \subseteq 1:d} \mathbf{1}_{\{u \cap v \neq \emptyset\}} \sigma_v^2.$$

They provide two ways to judge the importance of the set of variables x_j for $j \in u$. They are usually normalized by σ^2 to get an interpretation as a proportion of variance explained.

The mean dimension of f is $\nu(f) = \sum_{u \in \{1,d\}} |u| \sigma_u^2 / \sigma^2$. It satisfies $\nu(f) = \sum_{j=1}^d \bar{\tau}_j^2 / \sigma^2$. A ridge function takes the form $f(\mathbf{x}) = h(\Theta^\top \mathbf{x})$ for $\Theta \in \mathbb{R}^{d \times r}$ with $\Theta^\top \Theta = I_r$ and a function $h: \mathbb{R}^r \rightarrow \mathbb{R}$. For $\mathbf{x} \sim \mathcal{N}(0, I_d)$, the distribution of $f(\mathbf{x})$ does not depend on d and the mean dimension is $O(1)$ as $d \rightarrow \infty$ if h is Lipschitz [19]. If h is an indicator function and $r = 1$, then it is possible to have $\nu(f) = \Omega(\sqrt{d})$ reduced to $O(1)$ by preintegration over a component variable x_j with Θ_{j1} bounded away from zero as $d \rightarrow \infty$. See [19] for an example.

3. Preintegration and scrambled net variance. Conditional MC can reduce but not increase the variance of plain MC integration. Here we show that the same thing holds for scrambled nets using the nested uniform scrambling of [29]. The affine linear scrambling of [25] has the same variance and hence the same result. We assume that $f \in L^2[0, 1]^d$. For any $f \in L^2[0, 1]^d$ we could set $f(\mathbf{x}) = 0$ for any $\mathbf{x} \in D = [0, 1]^d \setminus [0, 1]^d$ and get an equivalent function with the same integral and, almost surely, the same RQMC estimate because all $\mathbf{x}_i \sim \mathcal{U}[0, 1]^d$ avoid D with probability one.

We will preintegrate over one of the d components of $\mathbf{x} \in [0, 1]^d$. It is also possible to preintegrate over multiple components and reduce the RQMC variance each time, though the utility of that strategy is limited by the availability of suitable closed forms or effective quadratures.

3.1. Walsh function expansions. To get variance formulas for scrambled nets we follow Dick and Pillichshammer [8], who work with Walsh function expansion in $L^2[0, 1]^d$, for which they credit Pirsic [38]. Let $\omega_b = e^{2\pi i/b}$ with i being the imaginary unit. For $k \in \mathbb{N}_0$ write $k = \sum_{j \geq 0} \kappa_j b^j$ with base b digits $\kappa_j \in \{0, 1, \dots, b-1\}$. For $x \in [0, 1]$ write $x = \sum_{j \geq 1} \xi_j b^{-j}$ with base b digits $\xi_j \in \{0, 1, \dots, b-1\}$. This is unique in the sense that infinitely many ξ_j 's must be different from $b-1$.

Using the above notation we can define the k th b -adic Walsh function ${}_b\text{wal}_k: [0, 1] \rightarrow \mathbb{C}$ as ${}_b\text{wal}_k(x) = \omega_b^{\sum_{j \geq 1} \xi_j \kappa_{j-1}}$. The summation in the exponent is finite because $k < \infty$. Note that ${}_b\text{wal}_0(x) = 1$ for all $x \in [0, 1]$. For $\mathbf{x} = (x_1, \dots, x_d) \in [0, 1]^d$ and $\mathbf{k} = (k_1, \dots, k_d) \in \mathbb{N}_0^d$, the d -dimensional Walsh functions are defined as ${}_b\text{wal}_{\mathbf{k}}(\mathbf{x}) = \prod_{j=1}^d {}_b\text{wal}_{k_j}(x_j)$. The Walsh series expansion of $f(\mathbf{x})$ is

$$f(\mathbf{x}) \sim \sum_{\mathbf{k} \in \mathbb{N}_0^d} \hat{f}(\mathbf{k}) {}_b\text{wal}_{\mathbf{k}}(\mathbf{x}), \text{ where } \hat{f}(\mathbf{k}) = \int_{[0,1]^d} f(\mathbf{x}) \overline{{}_b\text{wal}_{\mathbf{k}}(\mathbf{x})} d\mathbf{x}.$$

The d -dimensional b -adic Walsh function system is a complete orthonormal basis in $L_2([0, 1]^d)$ [8, Theorem A.11] and the series expansion converges to f in L^2 .

For $\mathbf{x} \sim \mathcal{U}([0, 1]^d)$, the variance of $f(\mathbf{x})$ has the decomposition

$$\text{Var}(f(\mathbf{x})) = \sum_{\mathbf{k} \in \mathbb{N}_0^d \setminus \{\mathbf{0}\}} |\hat{f}(\mathbf{k})|^2.$$

The variance under scrambled nets is different. To study it we group the Walsh coefficients. For $\ell \in \mathbb{N}_0^d$ let $C_\ell = \{\mathbf{k} \in \mathbb{N}_0^d \mid \lfloor b^{\ell_j-1} \rfloor \leq k_j < b^{\ell_j}, 1 \leq j \leq d\}$. Then define $\beta_\ell(\mathbf{x}) = \sum_{\mathbf{k} \in C_\ell} \hat{f}(\mathbf{k}) {}_b\text{wal}_{\mathbf{k}}(\mathbf{x})$. The functions β_ℓ are orthogonal in that $\int_{[0,1]^d} \beta_\ell(\mathbf{x}) \overline{\beta_{\ell'}(\mathbf{x})} d\mathbf{x} = 0$ when $\ell' \neq \ell$. For $\ell \neq \mathbf{0}$, $\beta_\ell(\mathbf{x})$ has variance

$$\sigma_\ell^2(f) = \int_{[0,1]^d} |\beta_\ell(\mathbf{x})|^2 d\mathbf{x} = \sum_{\mathbf{k} \in C_\ell} |\hat{f}(\mathbf{k})|^2.$$

Let $\mathbf{a}_0, \dots, \mathbf{a}_{n-1}$ be a point set in $[0, 1]^d$ and let $\mathbf{x}_0, \dots, \mathbf{x}_{n-1}$ be the scrambled version of $\mathbf{a}_i$'s. Then for the estimator $\hat{\mu} = \hat{\mu}_n = \frac{1}{n} \sum_{i=0}^{n-1} f(\mathbf{x}_i)$ we have

$$(3.1) \quad \text{Var}(\hat{\mu}_n) = \frac{1}{n} \sum_{\ell \in \mathbb{N}_0^d \setminus \{\mathbf{0}\}} \Gamma_\ell \sigma_\ell^2(f)$$

for a collection of gain coefficients $\Gamma_\ell \geq 0$ that depend on the $\mathbf{a}_i$. This expression can also be obtained through a base b Haar wavelet decomposition [30]. Our Γ_ℓ equals nG_ℓ from [8]. The variance of $\hat{\mu}_n$ under independent and identically distributed MC sampling is $(1/n) \sum_{\ell \in \mathbb{N}_0^d \setminus \{\mathbf{0}\}} \sigma_\ell^2(f)$, so $\Gamma_\ell < 1$ corresponds to integrating the term $\beta_\ell(\mathbf{x})$ with less variance than MC does.

If scrambling of [29] or [25] is applied to $\mathbf{a}_i$, then

$$(3.2) \quad \Gamma_\ell = \frac{1}{n} \sum_{i,i'=0}^{n-1} \prod_{j=1}^d \frac{b \mathbf{1} \left\{ \lfloor b^{\ell_j} a_{i,j} \rfloor = \lfloor b^{\ell_j} a_{i',j} \rfloor \right\} - 1 \left\{ \lfloor b^{\ell_j-1} a_{i,j} \rfloor = \lfloor b^{\ell_j-1} a_{i',j} \rfloor \right\}}{b-1}.$$

This holds for any $\mathbf{a}_i$, not just digital nets. When $\mathbf{a}_i$ are the first b^m points of a (t, d) -sequence in base b , then $\Gamma = \sup_\ell \Gamma_\ell < \infty$ (uniformly in m) so that $\text{Var}(\hat{\mu}) \leq \Gamma \sigma^2/n$. Similarly, for any $\ell \in \mathbb{N}_0^d$ we have $\Gamma_\ell \rightarrow 0$ as $n = b^m \rightarrow \infty$ in a (t, d) -sequence in base b from which we get $\text{Var}(\hat{\mu}_n) = o(1/n)$. For a (t, m, d) -net in base b , one can show that the gain coefficients $\Gamma_\ell = 0$ for all ℓ with $|\ell| \leq m - t$.

3.2. Walsh decomposition after preintegration.

PROPOSITION 3.1. For $f \in L^2[0, 1]^d$ and $j \in 1:d$, let g be f preintegrated over x_j . Then for $\mathbf{k} \in \mathbb{N}_0^d$,

$$(3.3) \quad \hat{g}(\mathbf{k}) = \begin{cases} \hat{f}(\mathbf{k}), & k_j = 0 \\ 0, & k_j > 0. \end{cases}$$

Proof. We write

$$\hat{g}(\mathbf{k}) = \int_{[0,1]^{d-1}} \int_0^1 g(\mathbf{x}) \overline{b \text{wal}_{\mathbf{k}}(\mathbf{x})} dx_j d\mathbf{x}_{-j} = \int_{[0,1]^{d-1}} g(\mathbf{x}) \int_0^1 \overline{b \text{wal}_{\mathbf{k}}(\mathbf{x})} dx_j d\mathbf{x}_{-j}$$

because $g(\mathbf{x})$ does not depend on x_j . If $k_j > 0$, then the inner integral vanishes, establishing the second clause in (3.3). If $k_j = 0$, then $b \text{wal}_{k_j}(x_j) = 1$ for all x_j and the inner integral equals $\prod_{\ell \neq j} b \text{wal}_{k_\ell}(x_j) = b \text{wal}_{\mathbf{k}}(\mathbf{x})$, establishing the first clause. $\square$

THEOREM 3.2. For $\mathbf{a}_0, \dots, \mathbf{a}_{n-1} \in [0, 1]^d$ let $\mathbf{x}_0, \dots, \mathbf{x}_{n-1}$ be a scrambled version of them using the algorithm from [29] or [25]. Let $f \in L^2[0, 1]^d$, and for $j \in 1:d$, let g be f preintegrated over x_j . Then

$$\text{Var} \left(\frac{1}{n} \sum_{i=0}^{n-1} g(\mathbf{x}_i) \right) \leq \text{Var} \left(\frac{1}{n} \sum_{i=0}^{n-1} f(\mathbf{x}_i) \right).$$

Proof. With either f or g we have the same gain coefficients Γ_ℓ for $\ell \in \mathbb{N}_0^d$. However

$$\sigma_\ell^2(g) = \sum_{\mathbf{k} \in C_\ell} |\hat{g}(\mathbf{k})|^2 = \sum_{\mathbf{k} \in C_\ell, k_j=0} |\hat{f}(\mathbf{k})|^2 \leq \sum_{\mathbf{k} \in C_\ell} |\hat{f}(\mathbf{k})|^2 = \sigma_\ell^2(f).$$

The result now follows from (3.1). $\square$

Theorem 3.2 shows that preintegration does not increase the variance under scrambling. This holds whether or not the underlying points are a digital net, though of course the main case of interest is for scrambling of digital nets and sequences. It holds because the RQMC error is a sum of uncorrelated errors (one per Walsh coefficient) and preintegration can decrease but not increase their variances. A similar phenomenon happens for RQMC by random shifting of lattice rules or by digital shifting of digital nets. For those RQMC methods also, preintegration cannot increase the variance. The proofs for those cases are in supplementary material, section SM3.

A variance reduction for conditional RQMC was previously found by [1]. Their Theorem 4.3 shows that after conditioning on a certain linear combination of variables, the variance of a randomly shifted lattice rule approach for some barrier options cannot be greater than the unconditional sampling variance. Their proof is specific to barrier options. They note that their argument would also apply to digital shifts.

Preintegration has another benefit that is not captured by Theorem 3.2. By reducing the input dimension from d to $d-1$ we will be using a $(t', m, d-1)$ -net in base b with $t' \leq t$ and possibly $t' < t$, since the lowest possible quality parameter commonly increases with the dimension. For scrambled net sampling, reducing the dimension reduces an upper bound on the variance. For any function $f \in L^2[0, 1]^d$, the variance using a scrambled (t, m, d) -net is at most b^{t+d} times the MC variance. Reducing the dimension reduces the bound to $b^{t'+d-1}$ times the MC variance.

As remarked above, preintegration of f over a variable that f uses will reduce the variance under scrambled net sampling. This reduction does not require f to be monotone in x_j , though such cases have the potential to bring a greater improvement [14, 15].

3.3. Choice of preintegration variable. In order to choose x_j to preintegrate over, we can look at the variance reduction we get. Preintegrating over x_j reduces the scrambled net variance by

$$(3.4) \quad \frac{1}{n} \sum_{\ell \in \mathbb{N}_0^d \setminus \{\mathbf{0}\}} \Gamma_{\ell} \sigma_{\ell}^2(f) \mathbf{1}_{\{\ell_j > 0\}} = \frac{1}{n} \sum_{\ell \in \mathbb{N}_0^d \setminus \{\mathbf{0}\}} \Gamma_{\ell} \sum_{\mathbf{k} \in C_{\ell}} |\hat{f}(\mathbf{k})|^2 \mathbf{1}_{\{k_j > 0\}}.$$

Evaluating this quantity for each $j \in 1:d$ might be more expensive than getting a good estimate of μ . However, we don't need to find the best j . Any j where f depends on x_j will bring some improvement. Below we develop a principled and computationally convenient strategy of choosing the j which is most important as measured by a Sobol' index [43] from global sensitivity analysis [39].

A convenient proxy replacement for (3.4) is

$$(3.5) \quad \frac{1}{n} \sum_{\mathbf{k} \in \mathbb{N}_0^d \setminus \{\mathbf{0}\}} |\hat{f}(\mathbf{k})|^2 \mathbf{1}_{\{k_j > 0\}} = \frac{1}{n} \sum_{u \subseteq 1:d} \sigma_u^2 \mathbf{1}_{\{j \in u\}},$$

where σ_u^2 is the ANOVA variance component of f for the set u . The equality above follows because the ANOVA can be defined by collecting up the terms involving x_j for $j \in u$ from the orthogonal decomposition as Sobol' did similarly in [42] using Haar functions. The right-hand side of (3.5) equals $\bar{\tau}_j^2/n$. It counts all the variance components in which variable j participates. Conditioning on x_j removes $\bar{\tau}_j^2/n$ from the plain MC variance [46]. From the orthogonality properties of ANOVA effects it follows that

$$\bar{\tau}_j^2 = \frac{1}{2} \int_{[0,1]^{d+1}} (f(\mathbf{z}_j : \mathbf{x}_{-j}) - f(\mathbf{x}))^2 d\mathbf{z}_j d\mathbf{x}.$$

The Jansen estimator [21] is an estimate of the above integral that can be done by a $d + 1$ dimensional MC or QMC or RQMC sampling algorithm. Our main interest in this Sobol' index estimator is that we use it as a point of comparison to the use of active subspaces in choosing a projection of a Gaussian vector along which to preintegrate.

4. Active subspace method. An expectation defined with respect to $\mathbf{x} \sim \mathcal{N}(0, \Sigma)$ for positive definite $\Sigma \in \mathbb{R}^{d \times d}$ can always be written as an expectation with respect to $\mathbf{x} \sim \mathcal{N}(0, I)$. For a unit vector $\theta \in \mathbb{R}^d$, we will preintegrate over $\mathbf{x}^\top \theta \sim \mathcal{N}(0, 1)$ and then the problem is to make a principled choice of θ . It would not be practical to seek an optimal choice.

Our proposal is to use active subspaces [5]. As mentioned in the introduction we let

$$C = \mathbb{E}(\nabla f(\mathbf{x}) \nabla f(\mathbf{x})^\top)$$

and then let $\Theta[1:r]$ comprise the r leading eigenvectors of C . The original purpose for using active subspaces is to approximate $f(\mathbf{x}) \approx \tilde{f}(\Theta[1:r]^\top \mathbf{x})$ for some function $\tilde{f}$ on $\mathbb{R}^r$. It is well known that one can construct functions where the active subspace will be a bad choice over which to approximate. For instance, with $f(\mathbf{x}) = \sin(10^6 x_1) + 100x_2$ the $r = 1$ active subspace provides a function of x_1 alone, while a function of x_2 alone can provide a better approximation. Active subspaces remain useful for approximation because the motivating problems are usually not so pathological and there is a human in the loop to catch such things. They also have an enormous practical advantage in that one set of evaluations of ∇f can be used in the search for Θ instead of having every candidate Θ require its own evaluations of ∇f . Using active subspaces for integration retains that advantage.

In our setting, we take $r = 1$ and preintegrate over $\theta^\top \mathbf{x}$, where θ is the leading eigenvector of C . That is, θ maximizes $\theta^\top \mathbb{E}(\nabla f(\mathbf{x}) \nabla f(\mathbf{x})^\top) \theta$ over all d -dimensional unit vectors. Now suppose that instead of using $f(\mathbf{x})$ we use $f_Q(\mathbf{x}) = f(Q\mathbf{x})$ for an orthogonal matrix $Q \in \mathbb{R}^{d \times d}$. Then $\mathbb{E}(\nabla f_Q(\mathbf{x}) \nabla f_Q(\mathbf{x})^\top) = Q^\top C Q$, which is similar to C . It has the same eigenvalues and the leading eigenvector is $\tilde{\theta} = Q^\top \theta$.

4.1. Connection to a Sobol' index. The discussion in section 3 motivates preintegration of $f(\mathbf{x})$ for $\mathbf{x} \sim \mathcal{N}(0, I)$ over a linear combination $\theta^\top \mathbf{x}$ having the largest Sobol' index over unit vectors θ . For $\theta_1 = \theta$, let $\Theta = (\theta_1, \theta_2, \dots, \theta_d) \in \mathbb{R}^{d \times d}$ be an orthogonal matrix and write

$$f_\Theta(\mathbf{x}) = f(\Theta \mathbf{x}) = f(x_1 \theta_1 + x_2 \theta_2 + \dots + x_d \theta_d).$$

Then we define $\bar{\tau}_\theta^2(f)$ to be $\bar{\tau}_1^2$ in the ANOVA decomposition of $f_\Theta(\mathbf{x})$. First we show that $\bar{\tau}_\theta^2$ does not depend on $\Theta_{-1} \equiv \Theta[2:d]$, the last $d - 1$ columns of Θ .

Let z , $\tilde{z}$, and $\mathbf{y}$ be independent with distributions $\mathcal{N}(0, 1)$, $\mathcal{N}(0, 1)$, and $\mathcal{N}(0, I_{d-1})$, respectively. Let $\mathbf{x} = z\theta_1 + \Theta_{-1}\mathbf{y}$ and $\tilde{\mathbf{x}} = \tilde{z}\theta_1 + \Theta_{-1}\mathbf{y}$. Using the Jansen formula for $\theta = \theta_1$,

$$(4.1) \quad \bar{\tau}_\theta^2 = \frac{1}{2} \int_{\mathbb{R}} \int_{\mathbb{R}} \int_{\mathbb{R}^{d-1}} (f(z\theta_1 + \Theta_{-1}\mathbf{y}) - f(\tilde{z}\theta_1 + \Theta_{-1}\mathbf{y}))^2 \varphi(\mathbf{y}) d\mathbf{y} \varphi(\tilde{z}) d\tilde{z} \varphi(z) dz.$$

Now for an orthogonal matrix $Q \in \mathbb{R}^{(d-1) \times (d-1)}$, let

$$\tilde{\Theta} = \Theta \begin{pmatrix} 1 & \mathbf{0}_{d-1}^\top \\ \mathbf{0}_{d-1} & Q \end{pmatrix} = (\tilde{\theta}_1, \tilde{\theta}_2, \dots, \tilde{\theta}_d)$$

where $\tilde{\theta}_1 = \theta_1$. In this parameterization we get

$$\bar{\tau}_\theta^2 = \frac{1}{2} \int_{\mathbb{R}} \int_{\mathbb{R}} \int_{\mathbb{R}^{d-1}} (f(z\theta_1 + \Theta_{-1}Q\mathbf{y}) - f(\tilde{z}\theta_1 + \Theta_{-1}Q\mathbf{y}))^2 \varphi(\mathbf{y}) d\mathbf{y} \varphi(\tilde{z}) d\tilde{z} \varphi(z) dz$$

which matches (4.1) after a change of variable. There is an even stronger invariance property in this setup. The random variable $\mathbb{E}(f_\Theta(\mathbf{x}) | \mathbf{x}_{-1})$ has a distribution that does not depend on $\theta_2, \dots, \theta_d$.

THEOREM 4.1. *Let f satisfy $\mathbb{E}(f(\mathbf{x})^2) < \infty$ for $\mathbf{x} \sim \mathcal{N}(0, I_d)$ and let $\Theta \in \mathbb{R}^{d \times d}$ be an orthogonal matrix with columns θ_j for $j = 1, \dots, d$. Then the distribution of $\mathbb{E}(f_\Theta(\mathbf{x}) | \mathbf{x}_{-1})$ does not depend on the last $d-1$ columns of Θ .*

The proof is in the supplementary material (SM2). Another consequence of Theorem 4.1 is that $\bar{\tau}_\theta^2(f) = \bar{\tau}_1^2(f_\Theta)$ is unaffected by $\theta_2, \dots, \theta_d$. Because the variance of f is unchanged by making an orthogonal matrix transformation of its inputs, the normalized Sobol' indices $\bar{\tau}_\theta^2/\sigma^2$ and $\bar{\tau}_\theta^2/\sigma^2$ are also invariant.

Finding the optimal θ would ordinarily require an expensive search because every estimate of $\bar{\tau}_\theta^2$ for a given θ would require its own collection of evaluations of f . Using a Poincaré inequality from [44] we can bound that Sobol' index by

$$\bar{\tau}_\theta^2(f) \leq \mathbb{E}((\theta^\top \nabla f(\mathbf{x}))^2) = \theta^\top C \theta.$$

The active subspace direction thus maximizes an upper bound on the Sobol' index for a projection. Next we develop a deeper correspondence between these two measures.

For a unit vector $\theta \in \mathbb{R}^d$, we can write $f(\mathbf{x}) = f(\theta\theta^\top \mathbf{x} + (I - \theta\theta^\top)\mathbf{x})$. If $\mathbf{x}, \mathbf{z}$ are independent $\mathcal{N}(0, I)$ vectors, then we can change the component of $\mathbf{x}$ parallel to θ by changing the argument of f to be $\theta\theta^\top \mathbf{z} + (I - \theta\theta^\top)\mathbf{x}$. This leaves the resulting point unchanged in the $d-1$ dimensional space orthogonal to θ . Let $\tilde{x} = \theta^\top \mathbf{x}$ and $\tilde{z} = \theta^\top \mathbf{z}$. Then $\tilde{x}, \tilde{z} \sim \mathcal{N}(0, 1)$ and $(I - \theta\theta^\top)\mathbf{x} \sim \mathcal{N}(0, I - \theta\theta^\top)$ are all independent. If f is differentiable, then by the mean value theorem

$$f(\theta\theta^\top \mathbf{z} + (I - \theta\theta^\top)\mathbf{x}) - f(\theta\theta^\top \mathbf{x} + (I - \theta\theta^\top)\mathbf{x}) = \theta^\top \nabla f(\theta\tilde{y} + (I - \theta\theta^\top)\mathbf{x})(\tilde{z} - \tilde{x})$$

for a real number $\tilde{y}$ between $\tilde{x}$ and $\tilde{z}$. Using the Jansen formula, the Sobol' index for this projection is

$$(4.2) \quad \frac{1}{2} \theta^\top \mathbb{E} \left((\tilde{z} - \tilde{x})^2 \nabla f(\theta\tilde{y} + (I - \theta\theta^\top)\mathbf{x}) \nabla f(\theta\tilde{y} + (I - \theta\theta^\top)\mathbf{x})^\top \right) \theta$$

which matches $\theta^\top \mathbb{E}(\nabla f(\mathbf{x}) \nabla f(\mathbf{x})^\top) \theta$ over a $d-1$ dimensional subspace but differs from it as follows. First, it includes a weight factor $(\tilde{z} - \tilde{x})^2$ that puts more emphasis on pairs of inputs where $\theta^\top \mathbf{x}$ and $\theta^\top \mathbf{z}$ are far from each other. Second, the evaluation point projected onto θ equals $\tilde{y}$, which lies between two independent $\mathcal{N}(0, 1)$ variables instead of having the $\mathcal{N}(0, 1)$ distribution, and its exact location depends on details of f and there could be more than one such $\tilde{y}$ for some f . The formula simplifies in an illustrative way for quadratic functions f .

PROPOSITION 4.2. *If $f : \mathbb{R}^d \rightarrow \mathbb{R}$ is a quadratic function and $\theta \in \mathbb{R}^d$ is a unit vector, then the Sobol' index $\bar{\tau}_\theta^2$ is*

$$(4.3) \quad \theta^\top \mathbb{E} \left(\nabla f \left(\frac{\theta\theta^\top \mathbf{x}}{\sqrt{2}} + (I - \theta\theta^\top)\mathbf{x} \right) \nabla f \left(\frac{\theta\theta^\top \mathbf{x}}{\sqrt{2}} + (I - \theta\theta^\top)\mathbf{x} \right)^\top \right) \theta.$$

Proof. If f is quadratic, then $\tilde{y} = (\tilde{z} + \tilde{x})/2 \sim \mathcal{N}(0, 1/2)$ and $\tilde{z} - \tilde{x} \sim \mathcal{N}(0, 2)$ and $(I - \theta\theta^\top)\mathbf{x}$ are all independent. Then $\mathbb{E}((\tilde{z} - \tilde{x})^2) = 2$ and $\theta\tilde{y}$ has the same distribution as $\theta\theta^\top\mathbf{x}/\sqrt{2}$ which is also independent of $(\tilde{z} - \tilde{x})$ and $(I - \theta\theta^\top)\mathbf{x}$. Making those substitutions in (4.2) yields (4.3). $\square$

The Sobol' index in (4.3) matches the quantity optimized by the first active subspace apart from the divisor $\sqrt{2}$ affecting one of the d dimensions. We can also show directly that for $\mathbf{x} \sim \mathcal{N}(0, I)$ and $f(\mathbf{x}) = (1/2)\mathbf{x}^\top A\mathbf{x} + b^\top\mathbf{x}$ for a symmetric matrix A , the Sobol' criterion (4.1) reduces to $\theta^\top A^2\theta + (\theta^\top b)^2 - (1/2)(\theta^\top A\theta)^2$ compared to an active subspace criterion of $\theta^\top A^2\theta + (\theta^\top b)^2$.

4.2. Active subspace for preintegration. Because $C = \mathbb{E}(\nabla f(\mathbf{x})\nabla f(\mathbf{x})^\top)$ is positive semidefinite, it has the eigen-decomposition $C = \Theta D \Theta^\top$, where $\Theta = (\theta_1, \dots, \theta_d) \in \mathbb{R}^{d \times d}$ is an orthogonal matrix consisting of eigenvectors of C , and $D = \text{diag}(\lambda_1, \dots, \lambda_d)$ with $\lambda_1 \geq \dots \geq \lambda_d \geq 0$ being the eigenvalues. Constantine, Dow, and Wang [6] prove that there exists a constant c such that

$$(4.4) \quad \mathbb{E}((f(\mathbf{x}) - \mathbb{E}(f(\mathbf{x})|\Theta[1:r]^\top\mathbf{x}))^2) \leq c(\lambda_{r+1} + \dots + \lambda_d)$$

for all f with a square integrable gradient. In general, the Poincaré constant c depends on the support of the function and the probability measure. But for the multivariate standard Gaussian distribution, the Poincaré constant is always 1 [37].

In our problem, we take $r = 1$ and compute $\mathbb{E}(f(\mathbf{x})|\theta^\top\mathbf{x})$, where θ is the first column of Θ . In practice, it is convenient to preintegrate the last variable of f_Θ , not the first. For instance, one would use the first $d - 1$ components in a Sobol' sequence not components 2 through d . Taking θ to be the first column of Θ , we compute

$$g(\mathbf{x}_{-d}) = \int_{-\infty}^{\infty} f(\theta x_d + \Psi\mathbf{x}_{-d})\varphi(x_d) dx_d,$$

where Ψ is orthogonal to θ (i.e., $\theta^\top\Psi = \mathbf{0}$), using a quadrature rule of negligible error or, if possible, a closed form. We then integrate this g over $d - 1$ variables by RQMC. We can use $\Psi = \Theta[2:d]$. Or if we want to avoid the cost of computing the full eigendecomposition of C we can find θ by a power iteration and then use a Householder transformation $\tilde{\Theta} = I - 2\mathbf{w}\mathbf{w}^\top$, where $\mathbf{w} = (\theta - e_1)/\|\theta - e_1\|$ and $e_1 = (1, 0, 0, \dots, 0)^\top$. This $\tilde{\Theta}$ is an orthogonal matrix whose first column is θ . In our numerical work, we have used $\Psi = \Theta[2:d]$ instead of the Householder transformation because of the effective dimension motivation for those eigenvectors given by [47].

We note that active subspaces use an uncentered PCA analysis of the matrix of sample gradients. One could use instead a centered analysis of $\mathbb{E}((\nabla f(\mathbf{x}) - \eta)(\nabla f(\mathbf{x}) - \eta)^\top)$, where $\eta = \mathbb{E}(\nabla f(\mathbf{x}))$. The potential advantage of this is that $\nabla f(\mathbf{x}) - \eta$ is the gradient of $f(\mathbf{x}) - \eta^\top\mathbf{x}$ which subtracts a linear approximation from f before searching for θ . The rationale for this alternative is that RQMC might already do well integrating a linear function and we would then want to choose a direction θ that performs well for the nonlinear part of f . In our examples, we found very little difference between the two methods and so we proposed the simpler uncentered active subspace preintegration.

In practice, we must estimate C . In the above description, we replace C by

$$(4.5) \quad \hat{C} = \frac{1}{M} \sum_{i=0}^{M-1} \nabla f(\mathbf{x}_i) \nabla f(\mathbf{x}_i)^\top,$$

Algorithm 4.1 Preintegration with active subspace

Input: Integrand f , number of samples M to compute $\hat{C}$, number of samples n to compute $\hat{\mu}$
 $\triangleright$ Find active subspaces
 Take $\mathbf{x}_0, \dots, \mathbf{x}_{M-1} \sim \mathcal{N}(0, I_d)$ by RQMC.
 Compute $\hat{C} = \frac{1}{M} \sum_{i=0}^{M-1} \nabla f(\mathbf{x}_i) \nabla f(\mathbf{x}_i)^\top$.
 Compute the eigen-decomposition $\hat{C} = \hat{\Theta} \hat{D} \hat{\Theta}^\top$.
 $\triangleright$ Preintegration
 Let $\theta = \hat{\Theta}[1]$ and $\Psi = \hat{\Theta}[2:d]$
 Define $g: \mathbb{R}^{d-1} \rightarrow \mathbb{R}$ by $g(\mathbf{x}) = \int_{-\infty}^{\infty} f(\theta x_d + \Psi \mathbf{x}_{-d}) \varphi(x_d) dx_d$
 $\triangleright$ RQMC integration
 Take $\mathbf{x}_0, \dots, \mathbf{x}_{n-1} \sim \mathcal{N}(0, I_{d-1})$ by RQMC and compute $\Psi x_0, \dots, \Psi x_{n-1}$.
return $\hat{\mu} = \frac{1}{n} \sum_{i=0}^{n-1} g(\mathbf{x}_i)$.

for an RQMC generated sample with $\mathbf{x}_i \sim \mathcal{N}(0, I_d)$ and then define θ and Θ_{-1} using $\hat{C}$ in place of C . We summarize the procedure in Algorithm 4.1.

Using our prior notation we can now describe the approach of [46] more precisely. They first preintegrate one variable in closed form, producing a $d-1$ dimensional integrand. They then apply GPCA to the preintegrated function to find a good $d-1$ dimensional rotation matrix [6]. That is, they first find $h(\mathbf{x}_{2:d}) := \mathbb{E}(f(\mathbf{x}) | \mathbf{x}_{2:d})$, then compute

$$(4.6) \quad \hat{\hat{C}} = \frac{1}{M} \sum_{i=0}^{M-1} \nabla h(\mathbf{x}_i) \nabla h(\mathbf{x}_i)^\top \in \mathbb{R}^{(d-1) \times (d-1)}, \quad \mathbf{x}_i \sim \mathcal{N}(0, I_{d-1}),$$

using RQMC points $\mathbf{x}_i$. Then they find the eigen-decomposition $\hat{\hat{C}} = \hat{\hat{V}} \hat{\hat{\Lambda}} \hat{\hat{V}}^\top$. Finally, they use RQMC to integrate the function $h(\hat{\hat{V}} \mathbf{x})$, where $\mathbf{x} \sim \mathcal{N}(0, I_{d-1})$. The main difference is that they apply preintegration to the original integrand $f(\mathbf{x})$, while we apply preintegration to the rotated integrand $f_\Theta(\mathbf{x}) = f(\Theta \mathbf{x})$. They conduct GPCA in the end as an approach to reduce effective dimension, while we conduct a similar GPCA (active subspace method) at the beginning to find the important subspace.

5. Application to option pricing. Here we study some Gaussian integrals arising from financial valuation. We assume that an asset price S_t , such as a stock, follows a geometric Brownian motion satisfying the stochastic differential equation (SDE)

$$dS_t = rS_t dt + \sigma S_t dB_t,$$

where B_t is a Brownian motion under the risk-neutral measure. Here, r is the interest rate and $\sigma > 0$ is the constant volatility for the asset. For an initial price S_0 , the SDE above has a unique solution

$$S_t = S_0 \exp\left(\left(r - \frac{\sigma^2}{2}\right)t + \sigma B_t\right).$$

Suppose that the maturity time of the option is T . In practice, we simulate a discrete Brownian motion. We call B a d -dimensional discrete Brownian motion if B follows a multivariate Gaussian distribution with mean zero and covariance Σ with $\Sigma_{ij} =$

$\Delta t \min(i, j)$, where $\Delta t = T/d$ is the length of each time interval and $1 \leq i, j \leq d$. To sample a discrete Brownian motion, we can first find a $d \times d$ matrix R such that $RR^T = \Sigma$, then generate a standard Gaussian variable $\mathbf{z} \sim \mathcal{N}(0, I_d)$, and let $B = R\mathbf{z}$. Taking R to be the lower triangular matrix in the Cholesky decomposition of Σ yields the *standard construction*. Using the usual eigen-decomposition $\Sigma = U\Lambda U^T$, we can take $R = U\Lambda^{1/2}$. This is called the *principal component analysis (PCA)* construction. For explicit forms of both these choices of R , see [12].

5.1. Option with one asset. When we use the matrix R , we can approximate $S_{j\Delta t}$ by

$$S_j = S_0 \exp\left(\left(r - \frac{\sigma^2}{2}\right)j\Delta t + \sigma B_j\right), \quad 1 \leq j \leq d$$

where $B = R\mathbf{z}$ is the discrete Brownian motion and B_j is the j th coordinate of B . The arithmetic average of the stock price is given by

$$\bar{S}(R, \mathbf{z}) = \frac{S_0}{d} \sum_{j=1}^d \exp\left(\left(r - \frac{\sigma^2}{2}\right)j\Delta t + \sigma \sum_{k=1}^d R_{jk} z_k\right).$$

Then the expected payoff of the arithmetic average Asian call option with strike price K is $\mathbb{E}((\bar{S}(R, \mathbf{z}) - K)_+)$, where the expectation is taken over $\mathbf{z} \sim \mathcal{N}(0, I_d)$.

Suppose that we want to preintegrate with respect to z_1 before computing the expectation $\mathbb{E}((\bar{S}(R, \mathbf{z}) - K)_+)$. If $R_{j1} > 0$ for all $1 \leq j \leq d$, then $\bar{S}(R, \mathbf{z})$ is increasing in z_1 for any value of $\mathbf{z}_{2:d}$. If we can find $\gamma = \gamma(\mathbf{z}_{2:d})$ such that

$$(5.1) \quad \bar{S}(R, (\gamma, \mathbf{z}_{2:d})) = K,$$

then the preintegration step becomes

$$\begin{aligned} & \mathbb{E}((\bar{S}(R, \mathbf{z}) - K)_+ | \mathbf{z}_{2:d}) \\ &= \int_{z_1 \geq \gamma(\mathbf{z}_{2:d})} (\bar{S}(R, (z_1, \mathbf{z}_{2:d})) - K) \varphi(z_1) dz_1 \\ (5.2) \quad &= \frac{S_0}{d} \sum_{j=1}^d \exp\left(\left(r - \frac{\sigma^2}{2}\right)j\Delta t + \sigma \sum_{k=2}^d R_{jk} z_k + \frac{\sigma^2 R_{j1}^2}{2}\right) \bar{\Phi}(\gamma - \sigma R_{j1}) - K \bar{\Phi}(\gamma), \end{aligned}$$

where $\bar{\Phi}(x) = 1 - \Phi(x)$. In practice, (5.1) can be solved by a root finding algorithm [15]. For example, Newton iteration usually converges in only a few steps. The monotonicity in z_1 is important for the preintegration step to be carried out easily. Without monotonicity, there might exist multiple roots for (5.1). For instance, if there are two roots γ_1, γ_2 , we need to integrate over $[\gamma_1, \gamma_2]$ or $(-\infty, \gamma_1) \cup (\gamma_2, +\infty)$. In more general cases, we might have to use a high-precision quadrature rule for the preintegration step.

The condition that $R_{j1} \geq 0$ for $1 \leq j \leq d$ is satisfied when we use the standard construction or PCA construction of Brownian motion. With the active subspace method, we are using $\tilde{R} = R\Theta$ in the place of R where Θ consists of the eigenvectors of $C = \mathbb{E}(\nabla f(\mathbf{z}) \nabla f(\mathbf{z})^T)$, and here $f(\mathbf{z}) = (\bar{S}(R, \mathbf{z}) - K)_+$. In this example, f is not differentiable at those $\mathbf{z}$ where $\bar{S}(R, \mathbf{z}) = K$. But those nondifferentiable points have zero probability, so with probability one $\nabla f(\mathbf{z})$ exists. We use the analytic form of the gradients when computing $\hat{C}$ in (4.5). Now we show that $\tilde{R}_{j1} \geq 0$ for all j . We start from the standard construction, meaning that $R_{ij} = \sqrt{\Delta t} \mathbf{1}\{i \geq j\}$; then $f(\mathbf{z})$ is

increasing in all coordinates of $\mathbf{z}$. So $\nabla f(\mathbf{z})$, whenever it exists, is always nonnegative in all coordinates, which means that $\widehat{C}_{ij} \geq 0$ for all $1 \leq i, j \leq d$. Therefore, the first eigenvector of $\widehat{C}$ has the same sign in all of its components. Without loss of generality, we can take $\Theta_{j1} \geq 0$. So the first column of $\tilde{R} = R\Theta$ is also nonnegative. This proves that the preintegration step for the active subspace integrand can also be easily conducted by a Newton search.

The arXiv preprint of this paper [24] includes numerical results for some of the Greeks for the Asian option. Those Greeks have integrands that are not only non-differentiable but actually discontinuous at $\mathbf{z}$, where $\bar{S}(R, \mathbf{z}) = K$, so the motivation for active subspaces does not apply to them. Our method in that report estimated C by the expected value of $\nabla f(\mathbf{z})\nabla f(\mathbf{z})^\top$ taken over $\mathbf{z}$ where $\nabla f(\mathbf{z})$ exists. This leads to a direction θ influenced by the gradient information while ignoring the jump discontinuities. Preintegrating over the resulting active subspace direction cannot raise variance and it produced very competitive results for those examples but this may be because the options ended up in the money (i.e., with $\bar{S}(R, \mathbf{z}) > K$) about half of the time. Choosing a direction θ to preintegrate over for discontinuous integrands is out of the scope of the present paper.

We compare Algorithm 4.1 with other methods in the option pricing example considered in [46] and [17]. We take the parameters $d = 50$, $T = 1$, $\sigma = 0.4$, $r = 0.1$, $S_0 = K = 100$ the same as in [17]. We consider five methods:

- **AS+preint**: our proposed active subspace preintegration method (Algorithm 4.1), which applies the active subspace method to find the direction to preintegrate;
- **preint+DimRed**: the method proposed in [46], which first preintegrates z_1 and applies GPCA to conduct dimension reduction for the other $d - 1$ variables;
- **preint**: preintegrating z_1 with no dimension reduction;
- **RQMC**: usual RQMC; and
- **MC**: plain Monte Carlo.

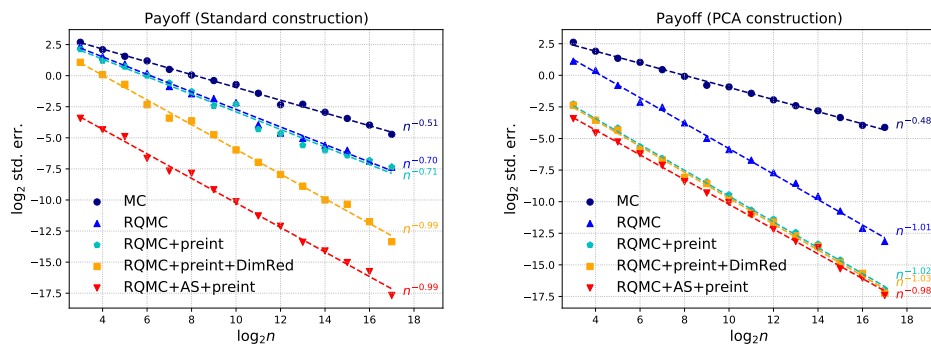
The first three methods, **AS+preint**, **preint+DimRed**, and **preint**, all use RQMC sampling unless otherwise specified.

For each sample size n , we repeat the simulation 50 times and compute the standard error across the 50 replicates. The standard error (std. err.) is plotted versus the sample size (n) on the log-log scale. For the methods **AS+preint** and **preint+DimRed**, we use $M = 128$ samples to estimate C as in (4.5). We approximate the gradients of the preintegrated integrand g by the finite difference

$$\nabla g(\mathbf{x}) \approx \left(\frac{g(\mathbf{x} + \varepsilon \mathbf{e}_1) - g(\mathbf{x})}{\varepsilon}, \dots, \frac{g(\mathbf{x} + \varepsilon \mathbf{e}_{d-1}) - g(\mathbf{x})}{\varepsilon} \right)^\top, \quad \varepsilon = 10^{-6},$$

matching the choice in [46]. We chose a small value of M to keep the costs comparable to plain RQMC. Also because θ is a local maximizer of $\theta^\top C \theta$, we have $\hat{\theta}^\top C \hat{\theta} = \theta^\top C \theta + O(\|\hat{\theta} - \theta\|^2)$, so there are diminishing benefits to accurate estimation of θ . Finally, any θ where f varies along $\theta^\top \mathbf{x}$ brings a variance reduction.

In Figure 1, we show the results under the standard construction (left) and the PCA construction (right) of the Brownian motion. We observe that **AS+preint** makes a huge improvement under the standard construction and is about the same as **preint+DimRed** under the PCA construction. The performance of active subspace preintegration is the same under either the standard or the PCA construction by invariance. For the Asian call option, it is already well known that the PCA construction is especially effective. Active subspace preintegration finds something almost as good without special knowledge.

FIG 1. *Single-asset call option.*

Running time is another measure of efficiency. In our examples, preintegration increases the running time by roughly 12-fold compared to methods that do not preintegrate. But that is not enough to offset the gains from preintegrating the first principal component or the first active subspace component for large n . In Figure 1, the improvement of **AS+preint** over plain RQMC is about 1000-fold under standard construction and 30-fold under PCA construction for large n . We also observe that the proposed **AS+preint** method is faster than **preint+DimRed** in terms of computing $\hat{C}$ and $\hat{\tilde{C}}$. This is because we only need to evaluate the original integrand to compute $\hat{C}$ but we need to evaluate the preintegrated integrand to compute $\hat{\tilde{C}}$. Evaluating the preintegrated integrand is more expensive because we need to apply Newton iterations for root-finding. The cost of computing estimates of C is negligible because it is done by inspecting only $M \ll n$ points in the domain of f . We put more details of the timings in the supplement (SM4).

5.2. Basket option. A basket option depends on a weighted average of several assets. Suppose that under the risk-neutral measure the L assets $S^{(1)}, \dots, S^{(L)}$ follow the SDE

$$dS_t^{(\ell)} = rS_t^{(\ell)} dt + \sigma_\ell S_t^{(\ell)} dB_t^{(\ell)},$$

where $\{B^{(\ell)}\}_{1 \leq \ell \leq L}$ are standard Brownian motions with correlation $\text{Corr}(B_t^{(\ell)}, B_t^{(k)}) = \rho_{\ell k}$ for all $t > 0$. For some nonnegative weights $w_1 + \dots + w_L = 1$, the payoff function of the Asian basket call option is given by

$$\left(\sum_{\ell=1}^L w_\ell \bar{S}^{(\ell)} - K \right)_+$$

where $\bar{S}^{(\ell)}$ is the arithmetic average of $S_t^{(\ell)}$ in the time interval $[0, T]$. Here, we only consider $L = 2$ assets. To generate $B^{(1)}, B^{(2)}$ with correlation ρ , we can generate two independent standard Brownian motions $W^{(1)}, W^{(2)}$ and let $B^{(1)} = W^{(1)}$, $B^{(2)} = \rho W^{(1)} + \sqrt{1 - \rho^2} W^{(2)}$. Following the same discretization as before, we can generate $(z^\top, \tilde{z}^\top) \sim \mathcal{N}(0, I_{2d})$. Then for time steps $j = 1, \dots, d$, let

$$S_j^{(1)} = S_0^{(1)} \exp \left(\left(r - \frac{\sigma_1^2}{2} \right) j \Delta t + \sigma_1 \sum_{k=1}^d R_{jk} z_k \right), \quad \text{and}$$

$$S_j^{(2)} = S_0^{(2)} \exp \left(\left(r - \frac{\sigma_2^2}{2} \right) j \Delta t + \sigma_2 \left(\rho \sum_{k=1}^d R_{jk} z_k + \sqrt{1 - \rho^2} \sum_{k=1}^d R_{jk} \tilde{z}_k \right) \right).$$

Again, the matrix R can be constructed by the standard construction or the PCA construction. We call these methods the ordinary standard construction and ordinary PCA construction.

A sharper principal components analysis would merge the two Brownian motions into a single $2d$ -dimensional process and use the principal components from their joint covariance matrix. The processes $B^{(1),\sigma_1} = \sigma_1 B^{(1)}$ and $B^{(2),\sigma_2} = \sigma_2 B^{(2)}$ have joint distribution

$$\begin{pmatrix} B^{(1),\sigma_1} \\ B^{(2),\sigma_2} \end{pmatrix} \sim \mathcal{N}\left(0, \begin{pmatrix} \sigma_1^2 \Sigma & \rho \sigma_1 \sigma_2 \Sigma \\ \rho \sigma_1 \sigma_2 \Sigma & \sigma_2^2 \Sigma \end{pmatrix}\right).$$

Let $\tilde{\Sigma}$ be the joint covariance matrix above. We can pick $\tilde{R}$ with $\tilde{R}\tilde{R}^\top = \tilde{\Sigma}$ and let

$$\begin{pmatrix} B^{(1),\sigma_1} \\ B^{(2),\sigma_2} \end{pmatrix} = \tilde{R} \begin{pmatrix} z \\ \tilde{z} \end{pmatrix}, \quad \text{where} \quad \begin{pmatrix} z \\ \tilde{z} \end{pmatrix} \sim \mathcal{N}(0, I_{2d}).$$

The matrix $\tilde{R}$ can be found by either a Cholesky decomposition or eigendecomposition of $\tilde{\Sigma}$. We call this method joint standard construction or joint PCA construction. With $B^{(1),\sigma_1}$ and $B^{(2),\sigma_1}$ generated, we can compute

$$S_j^{(\ell)} = S_0^{(\ell)} \exp\left(\left(r - \frac{\sigma_\ell^2}{2}\right)j\Delta t + \sum_{k=1}^d B_k^{(1),\sigma_\ell}\right).$$

In the preintegration step, we choose to integrate out z_1 . This can be easily carried out similarly as in (5.1) and (5.2) provided that the first column of $\tilde{R}$ is nonnegative. This is true if $\tilde{R}$ is found by the active subspaces method, following a similar argument in the previous example. We take $d = 50$, $T = 1$, $\rho = 0.5$, $S_0^{(1)} = S_0^{(2)} = 100$, $r = 0.1$, $\sigma_1 = 0.1$, $\sigma_2 = 0.2$, $w_1 = 0.8$, $w_2 = 0.2$. The standard errors under ordinary standard construction and PCA construction are plotted in Figures 2 and 3. The results for joint standard and PCA constructions are in Figure 4. The three columns correspond to $K = 80$, 100, and 120, respectively. A few observations are in order:

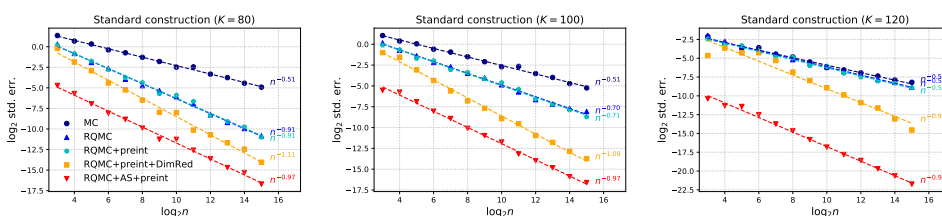


FIG 2. Basket options under ordinary standard construction.

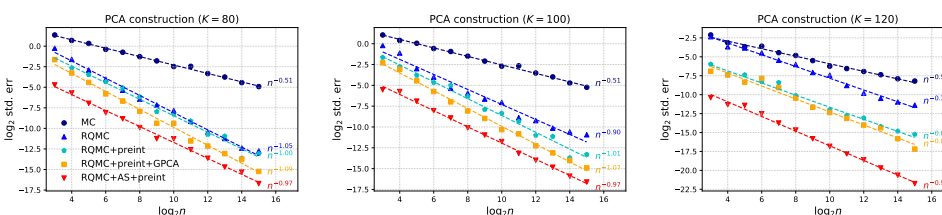


FIG 3. Basket options under ordinary PCA construction.

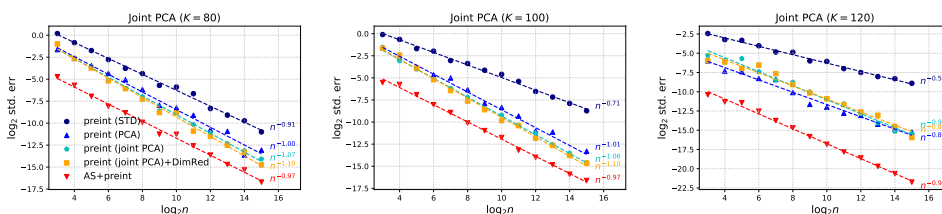


FIG 4. Basket options under joint standard or joint PCA construction with preintegration.

1. For the standard construction, preintegrating over z_1 brings little improvement over plain RQMC. But for the PCA construction, preintegrating over z_1 brings a big variance reduction for $K = 100$ and 120 .
2. The dimension reduction technique from [46] largely improves the error rate compared to preintegration without dimension reduction under the ordinary standard construction and ordinary PCA construction. This improvement is particularly significant for the ordinary standard construction and for smaller K .
3. The proposed method **AS+preint** has the best performance in all settings. It is even better than preintegrating out the first joint principal component of the Brownian motion with dimension reduction. In this example, the active subspace method is able to find a better direction than the first principal component over which to preintegrate.
4. For the $K = 120$ case, the probability that the option is exercised is about 2%. When we compute $\hat{C}$ with 128 samples, we only get five gradients that are not zero vector. But it is enough to find a better preintegration direction.

5.3. Rainbow option. A rainbow option is also a multiasset option whose payoff depends on the maximum or the minimum price across the assets. Here we consider the Asian *call on max* and *call on min* options, whose payoffs are defined as

$$\begin{aligned} \text{call on max: } & \mathbb{E} \left[\left(\max(\bar{S}^{(1)}, \bar{S}^{(2)}) - K \right)_+ \right] \quad \text{and} \\ \text{call on min: } & \mathbb{E} \left[\left(\min(\bar{S}^{(1)}, \bar{S}^{(2)}) - K \right)_+ \right], \end{aligned}$$

respectively. Here, $\bar{S}^{(1)}$ and $\bar{S}^{(2)}$ are the same as in the basket option.

For the above two integrands, we still have the monotonicity in z_1 if $R_{\cdot 1} \geq 0$; thus we can find the threshold $\gamma = \gamma(z_{-1})$ such that $\max(\bar{S}^{(1)}, \bar{S}^{(2)}) = K$ or $\min(\bar{S}^{(1)}, \bar{S}^{(2)}) = K$ by Newton iterations. After finding γ , the expectation over x_1 can be written as

$$\int_{\gamma}^{\infty} \max(\bar{S}^{(1)}, \bar{S}^{(2)}) \varphi(x_1) dx_1 - K \bar{\Phi}(\gamma).$$

For the call on min option, “max” is replaced by “min” in the last display. However, the above integral does not have a closed form. Therefore, we use a Gaussian quadrature to compute the integral over the interval $[\gamma, 10]$ using the quadrature function from SciPy [41].

We take $r = 0.1$, $\sigma_1 = 0.2$, $\sigma_2 = 0.1$, $S_0^{(1)} = 120$, $S_0^{(2)} = 80$. For the call on max option, we take $K \in \{100, 120\}$. For the call on min option, we take $K \in \{80, 100\}$. The results are shown in Figure 5. The call on min option with $K = 100$ is extremely out-of-the-money so we increase M to 2048 when computing $\hat{C}$ for the active subspace

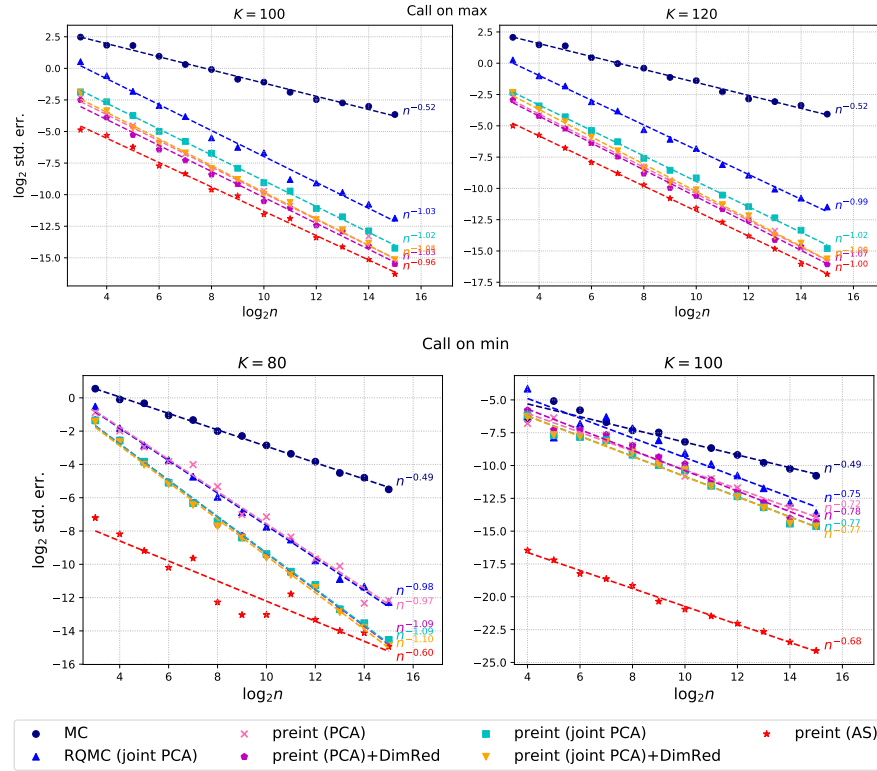


FIG 5. Rainbow options. The top panel is for the call on max option with $K \in \{100, 120\}$ and the bottom panel is for the call on min option with $K \in \{80, 100\}$.

method. We get three nonzero gradients. We also increase the number of replicates to 100 for this integrand. In the plot, n starts from 2^4 because for $n = 2^3$ some methods only get 0 in all 100 replicates.

Several observations are in order:

1. For the call on max option, preintegration with ordinary PCA is slightly better than preintegration with joint PCA. This is because for the call on max option, the maximum is most likely achieved by $\bar{S}^{(1)}$. So the integrand behaves like the single asset payoff in section 5.1. When we preintegrate z_1 under the ordinary PCA construction, we are preintegrating the first principal component of the Brownian motion for $\bar{S}^{(1)}$. And this is known to be quite effective.
2. For the call on min option, preintegration with joint PCA is better than preintegration with ordinary PCA. This is because the minimum is mostly achieved by $\bar{S}^{(2)}$. So to make preintegration under ordinary PCA more effective, one might consider exchanging the order of $\bar{S}^{(1)}$ and $\bar{S}^{(2)}$. But it is still not as effective as the active subspace method. See Figure SM4 in the supplement.
3. The proposed method, preintegration with the active subspace integrand, outperforms all other methods without using any specific knowledge about the integrands.

6. Discussion. In this paper we have studied a kind of conditional RQMC known as preintegration. We found that, just like conditional MC, the procedure can reduce variance but cannot increase it. We proposed to preintegrate over the first component in the active subspace. We showed a close relationship between this choice of preintegration variable and what one would get using a computationally infeasible but well-motivated choice by maximizing the Sobol' index of a linear combination of variables. The proposed method outperforms previous methods in some option pricing problems, especially when there is no strong incumbent construction of the Brownian motion.

Acknowledgments. We thank Dirk Nuyens and two anonymous reviewers for comments that have helped us improve this paper.

REFERENCES

- [1] N. ACHTSIS, R. COOLS, AND D. NUYENS, *Conditional sampling for barrier option pricing under the LT method*, SIAM J. Financial Math., 4 (2013), pp. 327–352.
- [2] P. ACWORTH, M. BROADIE, AND P. GLASSERMAN, *A comparison of some Monte Carlo techniques for option pricing*, in Monte Carlo and Quasi-Monte Carlo Methods '96, Springer, New York, 1997, pp. 1–18.
- [3] D. BIGONI, Y. MARZOUK, C. PRIEUR, AND O. ZAHM, *Nonlinear Dimension Reduction for Surrogate Modeling Using Gradient Information*, Tech. report, <https://arxiv.org/abs/2102.10351>, 2021.
- [4] W. CHEN, A. SRIVASTAV, AND G. TRAVAGLINI, eds., *A Panorama of Discrepancy Theory*, Springer, Cham, Switzerland, 2014.
- [5] P. G. CONSTANTINE, *Active Subspaces: Emerging Ideas for Dimension Reduction in Parameter Studies*, SIAM, Philadelphia, 2015.
- [6] P. G. CONSTANTINE, E. DOW, AND Q. WANG, *Active subspace methods in theory and practice: Applications to Kriging surfaces*, SIAM J. Sci. Comput., 36 (2014), pp. A1500–A1524.
- [7] P. J. DAVIS AND P. RABINOWITZ, *Methods of Numerical Integration*, 2nd ed., Academic Press, San Diego, 1984.
- [8] J. DICK AND F. PILlichshammer, *Digital Sequences, Discrepancy and Quasi-Monte Carlo Integration*, Cambridge University Press, Cambridge, UK, 2010.
- [9] B. EFRON AND C. STEIN, *The jackknife estimate of variance*, Ann. Statist., 9 (1981), pp. 586–596.
- [10] C. J. GEYER, *Conditioning in Markov chain Monte Carlo*, J. Comput. Graph. Statist., 4 (1995), pp. 148–154.
- [11] A. D. GILBERT, F. Y. KUO, AND I. H. SLOAN, *Preintegration Is Not Smoothing when Monotonicity Fails*, Tech. report, <https://arxiv.org/abs/2112.11621>, 2021.
- [12] P. GLASSERMAN, *Monte Carlo Methods in Financial Engineering*, Stoch. Model. Appl. Probab. 53, Springer, New York, 2004.
- [13] P. GLASSERMAN, P. HEIDELBERGER, AND P. SHAHABUDDIN, *Asymptotically optimal importance sampling and stratification for pricing path-dependent options*, Math. Finance, 9 (1999), pp. 117–152.
- [14] M. GRIEBEL, F. Y. KUO, AND I. H. SLOAN, *The smoothing effect of the ANOVA decomposition*, J. Complexity, 26 (2010), pp. 523–551.
- [15] A. GRIEWANK, F. Y. KUO, H. LEÖVEY, AND I. H. SLOAN, *High dimensional integration of kinks and jumps—smoothing by preintegration*, J. Comput. Appl. Math., 344 (2018), pp. 259–274.
- [16] J. M. HAMMERSLEY, *Conditional Monte Carlo*, J. ACM, 3 (1956), pp. 73–76.
- [17] Z. HE, *On the error rate of conditional quasi-Monte Carlo for discontinuous functions*, SIAM J. Numer. Anal., 57 (2019), pp. 854–874.
- [18] W. HOEFFDING, *A class of statistics with asymptotically normal distribution*, Ann. Math. Stat., 19 (1948), pp. 293–325.
- [19] C. R. HOYT AND A. B. OWEN, *Mean dimension of ridge functions*, SIAM J. Numer. Anal., 58 (2020), pp. 1195–1216.
- [20] J. IMAI AND K. S. TAN, *Minimizing effective dimension using linear transformation*, in Monte Carlo and Quasi-Monte Carlo Methods 2002, Springer, New York, 2004, pp. 275–292.
- [21] M. J. W. JANSSEN, *Analysis of variance designs for model output*, Comput. Phys. Commun., 117 (1999), pp. 35–43.

- [22] I. T. JOLLIFFE, *A note on the use of principal components in regression*, J. R. Stat. Soc. Ser. C. Appl. Stat., 31 (1982), pp. 300–303.
- [23] P. L'ECUYER AND C. LEMIEUX, *Variance reduction via lattice rules*, Management Sci., 46 (2000), pp. 1214–1235.
- [24] S. LIU AND A. B. OWEN, *Pre-integration via Active Subspaces*, Tech. report, <https://arxiv.org/abs/2202.02682>, 2022.
- [25] J. MATOUŠEK, *On the L^2 -discrepancy for anchored boxes*, J. Complexity, 14 (1998), pp. 527–556.
- [26] B. MOSKOWITZ AND R. E. CAFLISCH, *Smoothness and dimension reduction in quasi-Monte Carlo methods*, Math. Comput. Model., 23 (1996), pp. 37–54.
- [27] H. NIEDERREITER, *Random Number Generation and Quasi-Monte Carlo Methods*, SIAM, Philadelphia, 1992.
- [28] G. ÖKTEN AND A. GÖNCÜ, *Generating low-discrepancy sequences from the normal distribution: Box-Muller or inverse transform?*, Math. Comput. Model., 53 (2011), pp. 1268–1281.
- [29] A. B. OWEN, *Randomly permuted (t, m, s) -nets and (t, s) -sequences*, in Monte Carlo and Quasi-Monte Carlo Methods in Scientific Computing, Springer-Verlag, New York, 1995, pp. 299–317.
- [30] A. B. OWEN, *Monte Carlo variance of scrambled net quadrature*, SIAM J. Numer. Anal., 34 (1997), pp. 1884–1910.
- [31] A. B. OWEN, *Scrambled net variance for integrals of smooth functions*, Ann. Statist., 25 (1997), pp. 1541–1562.
- [32] A. B. OWEN, *Scrambling Sobol' and Niederreiter-Xing points*, J. Complexity, 14 (1998), pp. 466–489.
- [33] A. B. OWEN, *Local antithetic sampling with scrambled nets*, Ann. Statist., 36 (2008), pp. 2319–2343.
- [34] A. B. OWEN, *Monte Carlo Theory, Methods and Examples*, statweb.stanford.edu/~owen/mc, 2013.
- [35] A. B. OWEN AND D. RUDOLF, *A strong law of large numbers for scrambled net integration*, SIAM Rev., 63 (2021), pp. 360–372.
- [36] A. PAPAGEORGIOU, *The Brownian bridge does not offer a consistent advantage in quasi-Monte Carlo integration*, J. Complexity, 18 (2002), pp. 171–186.
- [37] M. T. PARENTE, J. WALLIN, AND B. WOHLMUTH, *Generalized bounds for active subspaces*, Electron. J. Stat., 14 (2020), pp. 917–943.
- [38] G. PIRSIC, *Schnell konvergierende Walshreihen über gruppen*, Master's thesis, University of Salzburg, Institute for Mathematics, 1995.
- [39] S. RAZAVI ET AL., *The future of sensitivity analysis: An essential discipline for systems modeling and policy support*, Environ. Modell. Soft., 137 (2021), 104954.
- [40] C. P. ROBERT AND G. O. ROBERTS, *Rao-Blackwellization in the MCMC Era*, University of Warwick, Tech. report, <https://arxiv.org/abs/2101.01011>, 2021.
- [41] *SciPy 1.0: Fundamental algorithms for scientific computing in Python*, Nat. Met., 17 (2020), pp. 261–272, <https://doi.org/10.1038/s41592-019-0686-2>.
- [42] I. M. SOBOLOV, *Multidimensional Quadrature Formulas and Haar Functions*, Nauka, Moscow, 1969 (in Russian).
- [43] I. M. SOBOLOV, *Sensitivity estimates for nonlinear mathematical models*, Math. Model. Comput. Experiment, 1 (1993), pp. 407–414.
- [44] I. M. SOBOLOV AND S. KUCHERENKO, *Derivative based global sensitivity measures and their link with global sensitivity indices*, Math. Comput. Simulation, 79 (2009), pp. 3009–3017.
- [45] H. F. TROTTER AND J. W. TUKEY, *Conditional Monte Carlo for normal samples*, in Symposium on Monte Carlo Methods, Wiley, New York, 1956, pp. 64–79.
- [46] Y. XIAO AND X. WANG, *Conditional quasi-Monte Carlo methods and dimension reduction for option pricing and hedging with discontinuous functions*, J. Comput. Appl. Math., 343 (2018), pp. 289–308.
- [47] Y. XIAO AND X. WANG, *Enhancing quasi-Monte Carlo simulation by minimizing effective dimension for derivative pricing*, Comput. Econ., 54 (2019), pp. 343–366.
- [48] R.-X. YUE AND S.-S. MAO, *On the variance of quadrature over scrambled nets and sequences*, Statist. Probab. Lett., 44 (1999), pp. 267–280.
- [49] O. ZAHM, P. CONSTANTINE, C. PRIEUR, AND Y. MARZOUK, *Gradient-based dimension reduction of multivariate vector-valued functions*, SIAM J. Sci. Comput., 42 (2020), pp. A534–A558.