

## Chapter 1

# RECENT ADVANCES IN RANDOMIZED QUASI-MONTE CARLO METHODS

Pierre L'Ecuyer

*Département d'Informatique et de Recherche Opérationnelle*

*Université de Montréal, C.P. 6128, Succ. Centre-Ville, Montréal, H3C 3J7, CANADA*

lecuyer@iro.umontreal.ca

Christiane Lemieux

*Department of Mathematics and Statistics*

*University of Calgary, 2500 University Drive N.W., Calgary, T2N 1N4, CANADA*

lemieux@math.ucalgary.ca

**Abstract** We survey some of the recent developments on quasi-Monte Carlo (QMC) methods, which, in their basic form, are a deterministic counterpart to the Monte Carlo (MC) method. Our main focus is the applicability of these methods to practical problems that involve the estimation of a high-dimensional integral. We review several QMC constructions and different randomizations that have been proposed to provide unbiased estimators and for error estimation. Randomizing QMC methods allows us to view them as variance reduction techniques. New and old results on this topic are used to explain how these methods can improve over the MC method in practice. We also discuss how this methodology can be coupled with clever transformations of the integrand in order to reduce the variance further. Additional topics included in this survey are the description of figures of merit used to measure the quality of the constructions underlying these methods, and other related techniques for multidimensional integration.

## 1. Introduction

To approximate the integral of a real-valued function  $f$  defined over the unit hypercube  $[0, 1]^s$ , given by

$$\mu = \int_{[0,1]^s} f(\mathbf{u}) d\mathbf{u}, \quad (1.1)$$

a frequently-used approach is to choose a point set  $P_n = \{\mathbf{u}_1, \dots, \mathbf{u}_n\} \subset [0, 1]^s$  and then take the average value of  $f$  over  $P_n$ ,

$$Q_n = \frac{1}{n} \sum_{i=1}^n f(\mathbf{u}_i), \quad (1.2)$$

as an approximation of  $\mu$ . Many problems can be formulated as in (1.1), e.g., when simulation is used to estimate an expectation and each simulation run requires  $s$  calls to a pseudorandom number generator that outputs numbers between 0 and 1; see Section 2 for more details.

If  $f$  is very smooth and  $s$  is small, the product of one-dimensional integration rules such as the rectangular or trapezoidal rules can be used to define  $P_n$  (Davis and Rabinowitz 1984). When these conditions are not met, the *Monte Carlo method* (MC) is usually more appropriate. It amounts to choosing  $P_n$  as a set of  $n$  i.i.d. uniformly distributed vectors  $\mathbf{u}_i$  over  $[0, 1]^s$ . With this method,  $Q_n$  is an unbiased estimator of  $\mu$  whose error can be approximated via the central limit theorem, and whose variance is given by

$$\frac{\sigma^2}{n} = \frac{1}{n} \left( \int_{[0,1]^s} f^2(\mathbf{u}) d\mathbf{u} - \mu^2 \right),$$

assuming that  $\sigma^2 < \infty$  (i.e.,  $f$  is square-integrable). This means that the error  $|Q_n - \mu|$  associated with the MC method has probabilistic order  $O_p(n^{-1/2})$ .

*Quasi-Monte Carlo* (QMC) methods can be seen as a deterministic counterpart to the MC method. They are based on the idea of using more regularly distributed point sets  $P_n$  to construct the approximation (1.2) than the random point set associated with MC. The fact that QMC methods are deterministic suggests that one has to make assumptions on the integrand  $f$  in order to guarantee a certain level of accuracy for  $Q_n$ . In other words, the improved regularity of  $P_n$  comes with worst-case functions for which the QMC approximation  $Q_n$  is bad. For this reason, the usual way to analyze QMC methods consists in choosing a set  $\mathcal{F}$  of functions and a definition of *discrepancy*  $D(P_n)$  to measure, in some

way, how far from the uniform distribution on  $[0, 1]^s$  is the empirical distribution induced by  $P_n$ . Once  $\mathcal{F}$  and  $D(P_n)$  are determined, one can usually derive upper bounds on the deterministic error, of the following form (Niederreiter 1992b; Hickernell 1998a):

$$|Q_n - \mu| \leq D(P_n)V(f), \text{ for any function } f \in \mathcal{F}, \quad (1.3)$$

where  $V(f)$  is a measure of the variability of  $f$  such that  $V(f) < \infty$  for all  $f \in \mathcal{F}$ . A well-known special case of (1.3) is the Koksma-Hlawka inequality (Hlawka 1961), in which  $\mathcal{F}$  is the set of functions having bounded variation in the sense of Hardy and Krause, and  $D(P_n)$  is the *rectangular-star discrepancy*. To compute this particular definition of  $D(P_n)$ , one considers all rectangular boxes in  $[0, 1]^s$  aligned with the axes and with a “corner” at the origin, and then take the supremum, over all these boxes, of the absolute difference between the volume of a box and the fraction of points of  $P_n$  that fall in it. The requirement that  $V(f) < \infty$  in this case implies in particular that  $f$  has continuous partial derivatives; see Niederreiter (1992b), page 19, for a precise definition.

It is clear from (1.3) that a small value of  $D(P_n)$  is desirable for the set  $P_n$ . This leads to the notion of *low-discrepancy sequences*, which are sequences  $P_\infty$  of points such that if  $P_n$  is constructed by taking the first  $n$  points of  $P_\infty$ , then  $D(P_n)$  is significantly smaller than the discrepancy of a typical set of  $n$  i.i.d. uniform points. The term *low-discrepancy point set* usually refers to a set obtained by taking the first  $n$  points of a low-discrepancy sequence, although it is sometimes used in a looser way to describe a point set that has a better uniformity than an independent and uniform point set.

In the case where  $D(P_n)$  is the rectangular-star discrepancy, it is common to say that  $P_\infty$  is a low-discrepancy sequence if  $D(P_n) = O(n^{-1} \log^s n)$ . Following this, the usual argument supporting the superiority of QMC methods over MC is to say that if  $Q_n$  is obtained using a low-discrepancy point set, then the error bound (1.3) is in  $O(n^{-1} \log^s n)$ , which for a fixed dimension  $s$  is a better asymptotic rate than the  $n^{-1/2}$  rate associated with MC. For this reason, one expects QMC methods to approximate  $\mu$  with a smaller error than MC if  $n$  is sufficiently large. However, the dimension  $s$  does not need to be very large in order to have  $n^{-1} \log^s n > n^{-1/2}$  for large values of  $n$ . For example, if  $s = 10$ , one must have  $n \geq 1.2 \times 10^{39}$  to ensure that  $n^{-1} \log^s n \leq n^{-1/2}$ , and thus the superiority of the convergence rate of the QMC bound over MC is meaningful only for values of  $n$  that are much too large for practical purposes.

Nevertheless, this does not mean that QMC methods cannot improve upon MC in practice, even for problems of large dimension. Arguments

supporting this are that first, the upper bound given in (1.3) is a worst-case bound for the whole set  $\mathcal{F}$ . It does not necessarily reflect the behavior of  $Q_n$  on a given function in this set. Second, it happens often in practice that even if the dimension  $s$  is large, the integrand  $f$  can be well approximated (in a sense to be specified in the next section) by a sum of low-dimensional functions. In that case, a good approximation  $Q_n$  for  $\mu$  can be obtained by simply making sure that the corresponding projections of  $P_n$  on these low-dimensional subspaces are well distributed. These observations have recently led many researchers to turn to other tools than the setup that goes with (1.3) for analyzing and improving the application of QMC methods to practical problems, where the dimension  $s$  is typically large, or even infinite (i.e., there is no *a priori* bound on  $s$ ). In connection with these new tools, the idea of randomizing QMC point sets has been an important contribution that has extended the practical use of these methods. The purpose of this chapter is to give a survey of these recent findings, with an emphasis on the theoretical results that appear most useful in practice. Along with explanations describing why these methods work, our goal is to provide, to the reader, tools for applying QMC methods to his own specific problems.

Our choice of topics is subjective. We do not cover all the recent developments regarding QMC methods. We refer the readers to, e.g., Fox (1999), Hellekalek and Larcher (1998), Niederreiter (1992b), and Spanier and Maize (1994) for other viewpoints. Also, the fact that we chose not to talk more about inequalities like (1.3) does not mean that they are useless. In fact, the concept of discrepancy turns out to be useful for defining selection criteria on which exhaustive or random searches to find “good” sets  $P_n$  can be based, as we will see later. Furthermore, we think it is important to be aware of the discouraging order of magnitude for  $n$  required for the rate  $n^{-1} \log^s n$  to be better than  $n^{-1/2}$ , and to understand that this problem is simply a consequence of the fact that placing points uniformly in  $[0, 1]^s$  is harder and harder as  $s$  increases because the space to fill becomes too large. This suggests that the success of QMC methods in practice is due to a clever choice of point sets exploiting the features of the functions that are likely to be encountered, rather than to an unexplainable way of breaking the “curse of dimensionality”.

Highly-uniform point sets can also be used for estimating the *minimum* of a function instead of its integral, sometimes in a context where function evaluations are noisy. This is discussed in Chapter 6 of Niederreiter (1992b), Chapter 13 of Fox (1999), and was also the subject of collaborative work between Sid Yakowitz and the first author (Yakowitz, L’Ecuyer, and Vázquez-Abad 2000).

This chapter is organized as follows. In Section 2, we give some insight on how point sets  $P_n$  should be constructed by using an ANOVA decomposition of the integrand over low-dimensional subspaces. Section 3 recalls the definition of different families of low-discrepancy point sets. In Section 4, we present measures of quality (or selection criteria) for low-discrepancy point sets that take into account the properties of the decomposition discussed in Section 2. Various randomizations that have been proposed for QMC methods are described in Section 5. Results on the error and variance of approximations based on (randomized) QMC methods are presented in Section 6. The purpose of Section 7 is to briefly review different classes of transformations that can be applied to the integrand  $f$  for reducing the variance further by exploiting, or not, the structure of the point set  $P_n$ . Integration methods that are somewhere between MC and QMC but that exploit specific properties of the integrand more directly are discussed in Section 8. Conclusions and ideas for further research are given in Section 9.

## 2. A Closer Look at Low-Dimensional Projections

We mentioned earlier that as the dimension  $s$  increases, it becomes difficult to cover the unit hypercube  $[0, 1]^s$  very well with a fixed number  $n$  of points. However, if instead our goal is to make sure that for some chosen subsets  $I \subset \{1, \dots, s\}$ , the projections  $P_n(I)$  over the subspace of  $[0, 1]^s$  indexed by the coordinates in  $I$  are evenly distributed, the task is easier. By doing this, one can get a small integration error if the chosen subsets  $I$  match the most important terms in the functional ANOVA decomposition of  $f$ , which we now explain.

The functional ANOVA decomposition (Efron and Stein 1981; Hoeffding 1948; Owen 1998a) writes a square-integrable function  $f$  as a sum of orthogonal functions as follows:

$$f(\mathbf{u}) = \sum_{I \subseteq \{1, \dots, s\}} f_I(\mathbf{u}),$$

where  $f_I(\mathbf{u}) = f_I(u_1, \dots, u_s)$  corresponds to the part of  $f$  that depends only on the variables in  $\{u_j, j \in I\}$ . Moreover, this decomposition is such that  $f_\emptyset(\mathbf{u}) = \mu$ ,

$$\int_{[0,1]^s} f_I(\mathbf{u}) d\mathbf{u} = 0$$

if  $I$  is non-empty, and

$$\int_{[0,1]^s} f_I(\mathbf{u}) f_J(\mathbf{u}) d\mathbf{u} = 0,$$

if  $I \neq J$ . Defining  $\sigma_I^2 = \text{Var}(f_I(\mathbf{u}))$ , we then have

$$\sigma^2 = \sum_{\emptyset \neq I \subseteq S} \sigma_I^2.$$

The best mean-square approximation of  $f(\cdot)$  by a sum of  $d$ -dimensional (or less) functions is  $\sum_{I: |I| \leq d} f_I(\cdot)$ . Also, the relative importance  $\sigma_I^2/\sigma^2$  of each component  $f_I$  indicates which variables or which subsets of variables are the most important (Hickernell 1998b).

A function  $f$  has *effective dimension*  $d$  (in the superposition sense) if  $\sum_{I: |I| \leq d} \sigma_I^2 \approx \sigma^2$  (Caflisch, Morokoff, and Owen 1997). Functions defined over many variables but having a low effective dimension often arise in practical applications (Caflisch, Morokoff, and Owen 1997; Lemieux and Owen 2001). The concept of effective dimension has actually been introduced (in a different form than above) by Paskov and Traub (1995) in the context of financial pricing to explain the much smaller error obtained with QMC methods compared with MC, for a problem defined over a 360-dimensional space.

A broad class of problems that are likely to have low effective dimension (relative to  $s$ ) are those arising from simulation applications. To see this, note that simulation is typically used to estimate the expectation of some measure of performance defined over a stochastic system, and proceeds by transforming in a more or less complicated way a sequence of numbers between 0 and 1 produced by a pseudorandom generator into an observation of the measure of interest. Hence it fits the framework of equation (1.1), with  $s$  equal to the number of uniforms required for each simulation, and  $f$  taken as the mapping that transforms a point in  $[0, 1]^s$  into an observation of the measure of performance. In that context, it is frequent that the uniform numbers that are generated close to each other in the simulation (i.e., corresponding to dimensions that are close together) are associated to random variables that interact more together. In other words, for these applications it is often the subsets  $I$  containing nearby indices and not too many of them that are the most important in the ANOVA decomposition. This suggests that to design point sets  $P_n$  that will work well for this type of problems, one should consider the quality of the projections  $P_n(I)$  corresponding to these “important” subsets. We present, in Section 4, measures of quality defined on this basis.

We conclude this section by recalling two important properties related to the projections of a point set  $P_n$  in  $[0, 1]^s$ . Firstly, we say that  $P_n$  is *fully projection-regular* (Sloan and Joe 1994; L’Ecuyer and Lemieux 2000) if each of its projections  $P_n(I)$  over a non-empty subset of dimensions  $I \subseteq \{1, \dots, s\}$  contains  $n$  distinct points. Such a prop-

erty is certainly desirable, for the lack of it means that some of the ANOVA components of  $f$  are integrated by less than  $n$  points even if  $n$  evaluations of  $f$  have been done. Secondly, we say that  $P_n$  is *dimension-stationary* (L'Ecuyer and Lemieux 2000) if  $P_n(I) = P_n(I + j)$  for any  $I = \{i_1, \dots, i_d\}$ ,  $1 \leq d, j < s$ , such that  $1 \leq i_1 < \dots < i_d < i_d + j \leq s$ ; that is, only the spacings between the indices in  $I$  are relevant in the definition of the projections  $P_n(I)$  of a dimension-stationary point set, and not their individual values. Hence not all  $2^s - 1$  non-empty projections of  $P_n$  need to be considered when measuring the quality of  $P_n$  since many are the same. Another advantage of dimension-stationary point sets is that because the quality of their projections does not deteriorate as the first index  $i_1$  increases, they can be used to integrate functions that have important ANOVA components associated with subsets  $I$  having a large value of  $i_1$ . Therefore, when working with those point sets it is not necessary to try rewriting  $f$  so that the important ANOVA components are associated with subsets  $I$  having a small first index  $i_1$  (as is often done; see, e.g., Fox 1999). We underline that not all types of QMC point sets have these properties.

### 3. Main Constructions

In this section, we present constructions for low-discrepancy point sets that are often used in practice. We first introduce *lattice rules* (Sloan and Joe 1994), and a special case of this construction called *Korobov rules* (Korobov 1959), which turn out to fit in another type of construction based on successive overlapping  $s$ -tuples produced by a recurrence defined over a finite ring. This type of construction is also used to define pseudorandom number generators (PRNGs) with huge period length when the underlying ring has a very large cardinality (e.g.,  $\geq 2^{100}$ ); low-discrepancy point sets are, in contrast, obtained by using a ring with a small cardinality (e.g., between  $10^2$  and  $10^6$ ). For this reason, we refer to this type of construction as *small PRNGs*, and discuss it after having introduced *digital nets* (Niederreiter 1992b), which form an important family of low-discrepancy point sets that also provides examples of small PRNGs. Various digital net constructions are described. We also recall the Halton sequence (Halton 1960), and discuss a method by which the number of points in a Korobov rule can be increased sequentially, thus offering an alternative to *digital sequences*. Additional references regarding the implementation of QMC methods are provided at the end of the section.

### 3.1 Lattice Rules

The general construction for lattice rules has been introduced by Sloan and his collaborators (see Sloan and Joe 1994 and the references therein) by building upon ideas developed by Korobov (1959, 1960), Bakhvalov (1959), and Hlawka (1962). The following definition is taken from the expository book of Sloan and Joe (1994).

**Definition:** Let  $\{\mathbf{v}_1, \dots, \mathbf{v}_s\}$  be a set of  $s$ -dimensional vectors linearly independent over  $\mathbb{R}^s$ , and with coordinates in  $[0, 1)$ . Define

$$L_s = \{\mathbf{v} = \sum_{j=1}^s z_j \mathbf{v}_j \text{ such that each } z_j \in \mathbb{Z}\}, \quad (1.4)$$

and assume that  $\mathbb{Z}^s \subseteq L_s$ . The approximation  $Q_n$  based on the set  $P_n = L_s \cap [0, 1)^s$  is a lattice rule. The number of points  $n$  in the rule is equal to the inverse of the absolute value of the determinant of the matrix  $\mathbf{V}$  whose rows are the vectors  $\mathbf{v}_j$ ,  $j = 1, \dots, s$ . This number  $n$  is called the order of the rule.

Note that the basis for  $L_s$  is not unique, but the determinant of the matrix  $\mathbf{V}$  remains constant for all choices of basis.

Figure 1.1 gives an example of a point set  $P_n$  that corresponds to a two-dimensional lattice rule, with  $n = 101$ . Here, the two vectors shown in the figure,  $\mathbf{v}_1 = (-8/101, 5/101)$  and  $\mathbf{v}_2 = (9/101, 7/101)$ , form a basis for the lattice  $L_2$ . Another basis for the same lattice is formed by  $\mathbf{v}_1 = (1/101, 12/101)$  and  $\mathbf{v}_2 = (0, 1)$ .

These 101 points cover the unit square quite uniformly. They are also placed very regularly on equidistant parallel lines, for several families of lines. For example, any of the vectors  $\mathbf{v}_1$  or  $\mathbf{v}_2$  given above determines one family of lines that are parallel to this vector. This regularity property stems from the lattice structure and it holds for *any* lattice rule—in more than two dimensions, the lines are simply replaced by equidistant parallel hyperplanes (Knuth 1998; Conway and Sloane 1999). For  $P_n$  to cover the unit hypercube quite well, the successive hyperplanes should never be too far apart (to avoid wide uncovered gaps), for any choice of family of parallel hyperplanes. Selection criteria for lattice rules, based on the idea of minimizing the distance between successive hyperplanes for the “worst-case” family, are discussed in Section 4.1.

From now on, we refer to point sets  $P_n$  giving rise to lattice rules as *lattice point sets*. Each lattice point set  $P_n$  has a *rank* associated with it, which can be defined as the smallest integer  $r$  such that  $P_n$  can be obtained by taking all integer combinations, modulo 1, of  $r$  vectors  $\mathbf{v}_1, \dots, \mathbf{v}_r$  independent over  $\mathbb{R}^s$ . Alternatively, the rank can be defined as the smallest number of cyclic groups whose direct sum yields  $P_n$  (Sloan

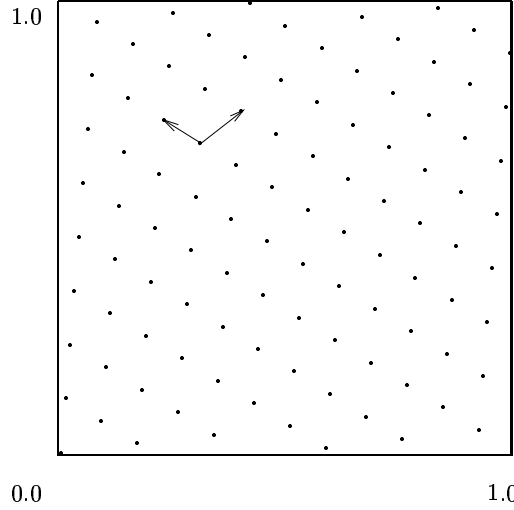


Figure 1.1. Korobov lattice point set with  $n = 101$  and  $a = 12$ ; the two vectors shown correspond to the basis  $\{(-8/101, 5/101), (9/101, 7/101)\}$ .

and Joe 1994). For example, if  $a_1, \dots, a_s$  are positive integers such that  $\gcd(a_j, n) = 1$  for at least one  $j$ , then the lattice point set

$$P_n = \left\{ \frac{i}{n}(a_1, \dots, a_s) \bmod 1, i = 0, \dots, n-1 \right\} \quad (1.5)$$

has rank 1 and contains  $n$  distinct points. It can also be obtained by taking  $P_n = L_s \cap [0, 1)^s$  with  $\mathbf{v}_1 = (a_1, \dots, a_s)/n$  and  $\mathbf{v}_j = \mathbf{e}_j$  for  $j = 2, \dots, s$  in (1.4), where  $\mathbf{e}_j$  is a vector of zeros with a one in the  $j^{\text{th}}$  position. The condition  $\gcd(a_j, n) = 1$  for all  $j$  is necessary and sufficient for a rank-1 lattice point set  $P_n$  to be fully projection-regular.

A *Korobov rule* is obtained by choosing an integer  $a \in \{1, \dots, n-1\}$ , and taking  $a_j = a^{j-1} \bmod n$  in (1.5), for all  $j = 1, \dots, s$ . In this case, having  $n$  and  $a$  relatively prime is a necessary and sufficient condition for  $P_n$  to be both fully projection-regular and dimension-stationary (L'Ecuier and Lemieux 2000). The integer  $a$  is usually referred to as the *generator* of the lattice point set  $P_n$ . For instance, the point set given on Figure 1.1 has a generator  $a$  equal to 12. In Section 3.3, we describe an efficient way of constructing  $P_n$  for Korobov rules when  $n$  is prime and  $a$  is a primitive element modulo  $n$ .

In the definition of a lattice rule, we assumed  $\mathbb{Z}^s \subseteq L_s$ . This implies that  $L_s$  has a period of 1 in each dimension. A lattice  $L_s$  with this property is called an *integration lattice* (Sloan and Joe 1994). A necessary and sufficient condition for  $L_s$  to be an integration lattice is that the inverse  $\mathbf{V}^{-1}$  of the matrix  $\mathbf{V}$  has only integer entries. In this case, it can

be shown that if the determinant of  $\mathbf{V}$  is  $1/n$ , then there exists a basis for  $L_s$  with coordinates of the form  $a/n$ , where  $0 \leq a \leq n$ . We assume from now on that all lattices considered are integration lattices.

The columns of the inverse matrix  $\mathbf{V}^{-1}$  form a basis of the *dual lattice* of  $L_s$ , defined as

$$L_s^* = \{\mathbf{h} \in \mathbb{Z}^s : \mathbf{h} \cdot \mathbf{v} \in \mathbb{Z}, \text{ for all } \mathbf{v} \in L_s\}.$$

If the determinant of  $\mathbf{V}^{-1}$  is  $n$ , then  $L_s^*$  contains  $n$  times less points per unit of volume than  $\mathbb{Z}^s$ . Also,  $L_s^*$  is periodic with a period of  $n$  in each dimension. As we will see in Sections 4.1 and 6.1, this dual lattice plays an important role in the error and variance analysis for lattice rules, and in the definition of selection criteria.

### 3.2 Digital Nets

We first recall the general definition of a digital net in base  $b$ , a concept that was first introduced by Sobol' (1967) in base 2, and subsequently generalized by Faure (1982), Niederreiter (1987), and Tezuka (1995). The following definition is from Niederreiter (1992b), with the convention, as in Tezuka (1995), that the *generating matrices*  $\mathbf{C}^j$  contain an infinite number of rows (although often, only a finite number of these rows are nonzero).

**Definition 1** *Let  $s \geq 1$  and  $k \geq 1$  be integers. Choose*

1. *a commutative ring  $R$  with identity and with cardinality  $b$  (usually  $\mathbb{Z}_b$ );*
2. *bijections  $\psi_r : \mathbb{Z}_b \rightarrow R$  for  $0 \leq r \leq k-1$ ;*
3. *bijections  $\eta_{jl} : R \rightarrow \mathbb{Z}_b$  for  $1 \leq j \leq s$  and  $1 \leq l \leq k$ ;*
4. *Generating matrices  $\mathbf{C}^1, \dots, \mathbf{C}^s$  of dimension  $\infty \times k$  over  $R$ .*

*For  $i = 0, \dots, b^k - 1$  let*

$$i = \sum_{r=0}^{k-1} a_r b^r,$$

*with  $a_r \in \mathbb{Z}_b$ , be the digit expansion of  $i$  in base  $b$ . Consider the vector*

$$\mathbf{y} = (\psi_0(a_0), \dots, \psi_{k-1}(a_{k-1}))^T \in R^k$$

*and compute*

$$(b_{j,1}, b_{j,2}, \dots)^T = \mathbf{C}^j \cdot \mathbf{y},$$

where each element  $b_{j,l}$  is in  $R$ . For  $j = 1, \dots, s$ , let

$$u_{ij} = \frac{\eta_{j,1}(b_{j,1})}{b} + \frac{\eta_{j,2}(b_{j,2})}{b^2} + \dots \quad (1.6)$$

Then  $P_n = \{\mathbf{u}_i = (u_{i1}, \dots, u_{is}), i = 0, \dots, n-1\}$ , with  $n = b^k$ , is a digital net over  $R$  in base  $b$ .

This scheme has been used to construct point sets having a low-discrepancy property that can be described by introducing the notion of  $(q_1, \dots, q_s)$ -*equidistribution* (Lemieux and L'Ecuyer 2001). Let  $q = q_1 + \dots + q_s$ , where the  $q_j$  are non-negative integers, and consider the  $b^q$   $b$ -ary boxes obtained by partitioning  $[0, 1]^s$  into  $b^{q_j}$  equal intervals along the  $j$ th axis. If each of these  $b^q$  boxes contains exactly  $b^{k-q}$  points from a point set  $P_n$ , where  $n = b^k$ , then  $P_n$  is said to be  $(q_1, \dots, q_s)$ -*equidistributed*. If a digital net  $P_n$  is  $(q_1, \dots, q_s)$ -equidistributed whenever  $q \leq k - t$ , for some integer  $t \geq 0$ , it is called a  $(t, k, s)$ -*net* (Niederreiter 1992b). The smallest integer  $t$  having this property is a widely-used measure of uniformity for digital nets and we call it the  $t$ -value of  $P_n$ . Note that the  $t$ -value is meaningless if  $k \leq t$ , and that the smaller  $t$  is, the better is the quality of  $P_n$ . Criteria for measuring the equidistribution of digital nets are discussed in more details in Section 4.2.

Figure 1.2 shows an example of a two-dimensional point set with  $n = 3^4 = 81$  points in base 3, having the best possible equidistribution; that is, its  $t$ -value is 0 and thus, any partition of the unit square into ternary boxes of size  $3^{-4}$  is such that exactly one point is in each box. The figure shows two examples of such a partition, into rectangles of sizes  $3^{-2} \times 3^{-2}$  and  $3^{-3} \times 3^{-1}$ , respectively. The other partitions (not shown here) are into rectangles of sizes  $1 \times 3^{-4}$ ,  $3^{-1} \times 3^{-3}$ , and  $3^{-4} \times 1$ . For all of these partitions, each rectangle contains one point of  $P_n$ . This point set contains the first 81 points of the two-dimensional Faure sequence in base 3. In this case, in the definition above,  $R = \mathbb{Z}_3$ , all bijections are the identity function over  $\mathbb{Z}_3$ , the generating matrix  $\mathbf{C}^1$  for the first dimension is the identity matrix, and  $\mathbf{C}^2$  is given by

$$\begin{pmatrix} 1 & 0 & 0 & 0 \\ 1 & 1 & 0 & 0 \\ 1 & 2 & 1 & 0 \\ 1 & 3 & 3 & 1 \end{pmatrix}.$$

The general definition of this type of construction is given below.

From now on,  $b$  is assumed to be prime power, and in all the constructions described below,  $R$  is taken equal to  $\mathbb{F}_b$ . Also, we assume that a bijection from  $\mathbb{F}_b$  to  $\mathbb{Z}_b$  has been chosen to identify the elements

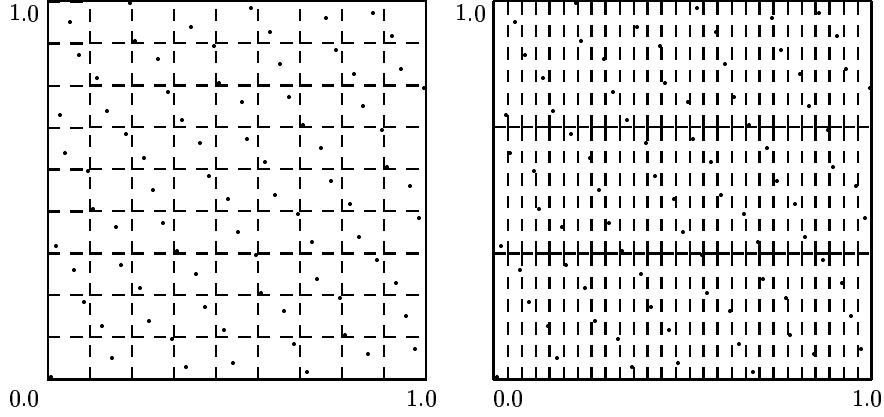


Figure 1.2. Digital net  $P_n$  in base 3 with  $n = 3^4$  points obtained from a Faure sequence. Left: Each of the 81 squares of size  $3^{-2} \times 3^{-2}$  contains one point from  $P_n$ ; Right: Each of the 81 rectangles of size  $3^{-3} \times 3^{-1}$  contains one point from  $P_n$ .

of  $\mathbb{F}_b$  with the “digits”  $\{0, \dots, b-1\}$ , and all bijections  $\psi_r$  and  $\eta_{jl}$  are defined according to this bijection. In particular, when  $b$  is prime, these bijections are equal to the identity function over  $\mathbb{Z}_b$ , and all operations are performed modulo  $b$ . The base  $b$  and the generating matrices  $\mathbf{C}^j$  therefore completely describe these constructions, because the  $b$ -ary expansion of a given coordinate  $u_{ij}$  of  $P_n$  is obtained by simply multiplying  $\mathbf{C}^j$  with the digit expansion of  $i$  in base  $b$ . The goal is then to choose the generating matrices so that the equidistribution property mentioned above holds for  $b$ -ary boxes that are as small as possible. In terms of these matrices, this roughly means that we want them to have a large number of rows that are linearly independent. For example, if for each  $j$  the first  $k$  rows of the matrix  $\mathbf{C}^j$  are linearly independent, then each  $b$ -ary box obtained by partitioning the  $j$ th axis into  $n$  equal intervals of length  $b^{-k}$  has one point from  $P_n$ . In particular, this implies that  $P_n$  is fully projection-regular.

*Digital sequences* in base  $b$  (see, e.g., Larcher 1998 and Niederreiter 1992b) are infinite sequences obtained in the same way as digital nets except that the generating matrices  $\mathbf{C}^j$  have an infinite number of columns; the first  $b^k$  points of the sequence thus form a digital net for each  $k \geq 1$ . For example, Sobol’, Generalized Faure, and Niederreiter sequences (to be discussed below) are all defined as digital sequences since the “recipe” to add extra columns in the generating matrices  $\mathbf{C}^j$  is given by the method.

We now describe specific well-known constructions of digital nets with a good equidistribution. We do not discuss recent constructions pro-

posed by Niederreiter and Xing (1997, 1998), as they require the introduction of many concepts from algebraic function fields that go well beyond the scope of this chapter. These sequences are built so as to optimize the asymptotic behavior of their  $t$ -value as a function of the dimension, for a fixed base  $b$ . See Pirsic (2001) for a definition of these sequences and a description of a software implementation.

**Sobol' Sequences.** Here the base is  $b = 2$  and the specification of each generating matrix  $\mathbf{C}^j$  requires a primitive polynomial  $f_j(z)$  over  $\mathbb{F}_2$ , and integers  $m_{j,q}$ , for  $1 \leq q \leq k = \deg(f_j(z))$ , to initialize a recurrence based on  $f_j(z)$  that generates the *direction numbers* defining  $\mathbf{C}^j$ . The method specifies that the polynomial  $f_j(z)$  should be the  $j$ th one in the list of primitive polynomials over  $\mathbb{F}_2$  sorted by increasing degree (within each degree, Sobol' specifies a certain order which is given in the code of Bratley and Fox (1988) for  $j \leq 40$ ). There remains the parameters  $m_{j,q}$  to control the quality of the point set.

Assume  $f_j(z) = z^k + a_{j,1}z^{k-1} + \dots + a_{j,k}$ , where  $a_{j,l} \in \mathbb{F}_2$  for each  $j, l$ . The direction numbers  $v_{j,1}, v_{j,2}, \dots$  are rationals of the form

$$v_{j,q} = \frac{m_{j,q}}{2^q} = \sum_{l=1}^q v_{j,q,l} 2^{-l},$$

where  $m_{j,q}$  is an odd integer smaller than  $2^q$ , for  $q \geq 1$ . The first  $k$  values  $v_{j,1}, \dots, v_{j,k}$  must be (carefully) chosen, and the following ones are obtained through the recurrence

$$v_{j,q} = a_{j,1}v_{j,q-1} \oplus \dots \oplus a_{j,k-1}v_{j,q-k+1} \oplus v_{j,q-k} \oplus (v_{j,q-k} \gg k),$$

where  $\oplus$  denotes a bit-by-bit exclusive-or operation, and  $v_{j,q-k} \gg k$  means that the binary expansion of  $v_{j,q-k}$  is shifted by  $k$  positions to the right (i.e.,  $v_{j,q-k}$  is divided by  $2^k$ ). These direction numbers are then used to define  $\mathbf{C}^j$ , whose entry in the  $l^{\text{th}}$  row and  $q^{\text{th}}$  column is given by  $v_{j,q,l}$ .

A good choice of the initial values  $v_{j,1}, \dots, v_{j,k}$  (or  $m_{j,1}, \dots, m_{j,k}$ ) in each dimension  $j$  is important for the success of this method, especially when  $n$  is small. The implementation of Bratley and Fox (1988) uses the initial values given by Sobol' and Levitan (1976) for  $m_{j,q}$ , with  $q \leq k = \deg(f_j(z))$  and  $j \leq 40$ . More details on how to choose these initial values to optimize certain quality criteria are given in Section 4.2.

**Generalized Faure Sequences.** These sequences were introduced by Faure (1982) and generalized by Tezuka (1995). For this type of sequence, the base  $b$  is the smallest prime power larger or equal to the

dimension  $s$  (which means that these sequences are practical only for small values of  $s$ ). An important feature of this construction is that their  $t$ -value has the best possible value ( $t = 0$ ), provided  $n$  is of the form  $n = \lambda b^d$ , for some integers  $\lambda, d \geq 1$ . Assuming  $n = b^k$ , the matrices  $\mathbf{C}^j$  take the form:

$$\mathbf{C}^j = \mathbf{A}_j(\mathbf{P}^T)^{j-1},$$

where  $\mathbf{P}$  is the  $\infty \times k$  Pascal's matrix (i.e., with entries  $\mathbf{P}_{i,j} = \binom{i-1}{j-1} \in \mathbb{F}_b$ ), and  $\mathbf{A}_j$  is an arbitrary  $\infty \times \infty$  non-singular lower-triangular matrix. The original Faure sequence is obtained by taking  $b$  prime and  $\mathbf{A}_j$  as the identity matrix for all  $j = 1, \dots, s$ . By allowing the matrices  $\mathbf{A}_j$  to be different from the identity matrix, point sets that avoid some of the defects observed on the original Faure sequence can be built (Morokoff and Caffisch 1994; Boyle, Broadie, and Glasserman 1997; Tan and Boyle 2000). For instance, the Generalized Faure sequence of Tezuka and Tokuyama (1994) amounts to take  $\mathbf{A}_j = \mathbf{P}^{j-1}$ . Recently, Faure suggested another form of Generalized Faure sequence that consists in taking  $\mathbf{A}_j$  equal to the lower-triangular matrix with all nonzero entries equal to 1, for all  $j$  (Faure 2001).

**Niederreiter Sequences.** This construction has been proposed by Niederreiter (1988). We first describe a special case (Niederreiter 1992b, Section 4.5), and then briefly explain how it can be generalized. The base  $b$  is assumed to be a prime power, and for this special case the generating matrices are defined by distinct monic irreducible polynomials  $p_1(z), \dots, p_s(z)$  in  $\mathbb{F}_b[z]$ . Let  $d_j = \deg(p_j(z))$  for  $j = 1, \dots, s$ . In what follows,  $\mathbb{F}_b((z^{-1}))$  represents the field of formal Laurent series over  $\mathbb{F}_b$ ; that is,

$$\mathbb{F}_b((z^{-1})) = \left\{ \sum_{l=w}^{\infty} d_l z^{-l} \text{ such that } w \in \mathbb{Z} \text{ and } d_l \in \mathbb{F}_b \text{ for each } l \right\}.$$

To find the element on the  $l$ th row and  $q$ th column of  $\mathbf{C}^j$ , for  $1 \leq j \leq s$ ,  $l \geq 1$ ,  $1 \leq q \leq k$ , consider the expansion

$$\frac{z^r}{p_j(z)^l} = \sum_{w=0}^{\infty} a_{l,r,w}^{(j)} z^{-w-1}, \quad (1.7)$$

where  $r = r(l, q)$  is the unique integer satisfying  $0 \leq r < d_j$  and  $l - 1 = \alpha d_j + r$ , and  $\alpha$  is also uniquely determined. The element on the  $l$ th row and  $q$ th column of  $\mathbf{C}^j$  is then given by  $a_{\alpha+1,r,q}^{(j)}$ . For these sequences, the  $t$ -value is given by  $\sum_{j=1}^s (d_j - 1)$ , which suggests that to minimize  $t$ , the

$p_j(z)$ 's should be taken equal to the  $s$  monic irreducible polynomials of smallest degree over  $\mathbb{F}_b$ . In the general definition of Niederreiter (1988), the numerator  $z^r$  in (1.7) can be multiplied by a different polynomial for each pair  $(j, l)$  of dimension and row index, and the  $p_j(z)$ 's just need to be pairwise relatively prime polynomials. Tezuka (1995), Section 6.1.2, generalizes this concept a step further by removing more restrictions on this numerator.

**Polynomial Lattice Rules.** This construction can be seen as a lattice rule in which the ring of integers is replaced by a ring of polynomials over a finite field. As we explain below, it generalizes the construction discussed by Niederreiter (1992a), Niederreiter (1992b), Section 4.4, and Larcher, Lauss, Niederreiter, and Schmid (1996), and is studied in more details by Lemieux and L'Ecuyer (2001).

**Definition:** Let  $\{\mathbf{v}_1(z), \dots, \mathbf{v}_s(z)\}$  be a set of  $s$ -dimensional vectors of formal Laurent series over  $\mathbb{F}_b$ , linearly independent over  $(\mathbb{F}_b((z^{-1})))^s$ , where  $b$  is a prime power. Define

$$\mathcal{L}_s = \left\{ \mathbf{v}(z) = \sum_{j=1}^s q_j(z) \mathbf{v}_j(z) \text{ such that each } q_j(z) \in \mathbb{F}_b[z] \right\}, \quad (1.8)$$

and assume that  $(\mathbb{F}_b[z])^s \subseteq \mathcal{L}_s$ . Let  $\varphi : (\mathbb{F}_b((z^{-1})))^s \rightarrow [0, 1]^s$  be the mapping that evaluates each component of a vector  $\mathbf{v}(z)$  in  $(\mathbb{F}_b((z^{-1})))^s$  at  $z = 2$ . The approximation  $Q_n$  based on the set  $P_n = \varphi(\mathcal{L}_s) \cap [0, 1]^s$  is a polynomial lattice rule. The number of points  $n$  in the rule is called the order of the rule and is equal to  $b^k$ , where  $k$  is the degree of the inverse of the determinant of the matrix  $\mathbf{V}$  whose rows are the vectors  $\mathbf{v}_j(z)$ .

Most definitions and results that we mentioned for lattice rules, which from now on are referred to as *standard* lattice rules, have their counterpart for polynomial lattice rules, as we now explain. First, we refer to point sets  $P_n$  that define polynomial lattice rules as *polynomial lattice point sets*, whose rank  $r$  is equal to the smallest number of basis vectors  $\mathbf{v}_1(z), \dots, \mathbf{v}_r(z)$  required to write  $P_n$  as

$$P_n = \left\{ \varphi(\mathbf{v}(z)) : \mathbf{v}(z) = \sum_{j=1}^r q_j(z) \mathbf{v}_j(z) \bmod \mathbb{F}_b[z], q_j(z) \in \mathbb{F}_b[z] \right\}.$$

In the expression above, “ $\bmod \mathbb{F}_b[z]$ ” represents the operation by which all non-negative powers in  $\mathbf{v}(z)$  are dropped. The construction introduced by Niederreiter (1992a), also discussed in Niederreiter (1992b), Section 4.4, and Larcher, Lauss, Niederreiter, and Schmid (1996), and

sometimes referred to as the “method of optimal polynomials”, is a polynomial lattice point set of rank 1, since it can be obtained as

$$P_n = \left\{ \varphi(\mathbf{v}(z)) : \mathbf{v}(z) = \frac{q(z)}{P(z)}(g_1(z), \dots, g_s(z)) \bmod \mathbb{F}_b[z], \right. \\ \left. q(z) \in \mathbb{F}_b[z]/(P) \right\}, \quad (1.9)$$

where  $n = b^k$ ,  $P(z) = z^k + a_1 z^{k-1} + \dots + a_k$  is a polynomial of degree  $k$  over  $\mathbb{F}_b$ , and the  $g_j(z)$  are polynomials in  $\mathbb{F}_b[z]/(P)$ , the ring of polynomials over  $\mathbb{F}_b$  modulo  $P(z)$ . A *Korobov polynomial lattice point set* is obtained by taking  $g_j(z) = g(z)^{j-1} \bmod P(z)$  in (1.9) for some  $g(z) \in \mathbb{F}_b[z]/(P)$ . As in the standard case, the condition  $\gcd(g(z), P(z))$  is necessary and sufficient to guarantee that a Korobov polynomial lattice point set is dimension-stationary and fully projection-regular (Lemieux and L’Ecuyer 2001). An efficient way of generating this type of point set when  $P(z)$  is a primitive polynomial is described in the next subsection.

The condition that  $(\mathbb{F}_b[z])^s \subseteq \mathcal{L}_s$  implies that  $\varphi(\mathcal{L}_s)$  has a period of 1 in each dimension, and we call  $\mathcal{L}_s$  a *polynomial integration lattice* in this case. If  $\mathbf{V}$  represents a matrix whose rows are formed by basis vectors of  $\mathcal{L}_s$ , then  $\mathcal{L}_s$  is a polynomial integration lattice if and only if all entries in  $\mathbf{V}^{-1}$  are in  $\mathbb{F}_b[z]$ . In that case, a basis for  $\mathcal{L}_s$  with vectors having coordinates of the form  $q(z)/P(z)$  can be found, where  $P(z)$  is the determinant of  $\mathbf{V}^{-1}$  and  $q(z)$  is in  $\mathbb{F}_b[z]/(P)$  or equal to  $P(z)$ . Finally, the columns of  $\mathbf{V}^{-1}$  form a basis of the *dual lattice* of  $\mathcal{L}_s$ , which is defined by

$$\mathcal{L}_s^* = \{\mathbf{h}(z) \in (\mathbb{F}_b[z])^s : \mathbf{h}(z) \cdot \mathbf{v}(z) \in \mathbb{F}_b[z], \text{ for all } \mathbf{v}(z) \in \mathcal{L}_s\},$$

where the scalar product  $\cdot$  is defined as  $\mathbf{h}(z) \cdot \mathbf{v}(z) = \sum_{j=1}^s h_j(z) v_j(z)$ .

This construction is a special case of digital nets in which the generating matrices are defined by

$$\mathbf{C}_{l,q}^j = v_{ij}(z)|_{l+q-1},$$

where  $v_{ij}(z)|_m$  denotes the coefficient of  $z^{-m}$  in the formal series  $v_{ij}(z)$ , the  $j$ th coordinate of  $\mathbf{v}_i(z)$ , and the index  $i$  is determined by  $q$  and the structure of  $\mathcal{L}_s$  as follows:  $i$  has the property that for some non-zero polynomials  $h_{11}(z), \dots, h_{ss}(z)$  coming from a triangular basis of  $\mathcal{L}_s^*$  and such that  $\prod_{m=1}^s h_{m,m}(z) = P(z)$  (see Lemieux and L’Ecuyer 2001, Lemma A.2, for more details), we have

$$\deg(h_{11}(z) + \dots + h_{i-1,i-1}(z)) < q \leq \deg(h_{11}(z) + \dots + h_{i,i}(z)).$$

In other words, the  $\deg(h_{11}(z))$  first columns of each matrix  $\mathbf{C}^j$  contain the coefficients associated with the first basis vector  $\mathbf{v}_1(z)$ ; the  $\deg(h_{22}(z))$  next columns are associated with the second basis vector  $\mathbf{v}_2(z)$ , and so on.

For polynomial lattice point sets of rank 1, the corresponding generating matrices can be described a bit more easily (Niederreiter 1992b, Section 4.4), as we now explain. For each dimension  $j = 1, \dots, s$ , consider the formal Laurent series expansion

$$\frac{g_j(z)}{P(z)} = \sum_{l=1}^{\infty} w_{j,l} z^{-l}.$$

The first  $k$  coefficients  $w_{j,1}, \dots, w_{j,k}$  satisfy

$$\mathbf{A} \begin{pmatrix} w_{j,1} \\ \vdots \\ w_{j,k} \end{pmatrix} = \begin{pmatrix} g_{j,1} \\ \vdots \\ g_{j,k} \end{pmatrix},$$

where  $\mathbf{A}$  is defined as

$$\mathbf{A} = \begin{pmatrix} 1 & 0 & \dots & 0 \\ a_1 & 1 & \dots & 0 \\ \vdots & \ddots & \ddots & \vdots \\ a_{k-1} & \dots & a_1 & 1 \end{pmatrix}$$

(see, e.g., Lemieux and L'Ecuyer 2001), and the coefficients  $g_{j,l}$  are such that  $g_j(z) = \sum_{l=1}^k g_{j,l} z^{k-l}$ . The coefficients  $w_{j,l}$  for  $l > k$  satisfy the recurrence determined by  $P(z)$ , i.e.,

$$w_{j,l} = -a_1 w_{j,l-1} - \dots - a_k w_{j,l-k},$$

where the minus sign represents the subtraction in  $\mathbb{F}_b$ . The entries of the generating matrices  $\mathbf{C}^1, \dots, \mathbf{C}^s$  are defined by

$$\mathbf{C}_{l,q}^j = w_{j,l+q-1}, \text{ for } l \geq 1, 1 \leq q \leq k.$$

Note that in the definition of Niederreiter (1992b) and Pirsic and Schmid (2001), the matrices  $\mathbf{C}^j$  are restricted to  $k$  rows.

### 3.3 Constructions Based on Small PRNGs

Consider a PRNG based on a recurrence over a finite ring  $R$  and a specific output function from  $R$  to  $[0, 1)$ . The idea proposed here is to define  $P_n$  as the set of all overlapping  $s$ -tuples in  $[0, 1)^s$  than can be

obtained by running the PRNG from all possible initial states in  $R$ . More precisely, let  $\phi : R \rightarrow R$  be a bijection over  $R$  called the *transition function* of the PRNG, and define

$$r_i = \phi(r_{i-1}),$$

for  $i \geq 0$ . The sequence  $r_0, r_1, \dots$  obtained from any seed  $r_0 \in R$  has a period of at most  $|R|$ . Now let  $g : R \rightarrow [0, 1)$  be an *output function*, and define

$$u_i = g(r_i)$$

for  $i \geq 0$ . We call the point set

$$P_n = \{(u_0, \dots, u_{s-1}), r_0 \in R\}, \quad (1.10)$$

where  $n = |R|$ , a *recurrence-based* point set. The requirement that  $\phi$  be a bijection guarantees that  $P_n$  is dimension-stationary, and therefore fully-projection regular (L'Ecuyer and Lemieux 2000). In concrete realizations, it is often the case that the recurrence  $r_0, r_1, \dots$  has two main cycles, one of period length  $|R| - 1$ , and the other of period length 1 and which contains 0, the zero of  $R$ . In this case,  $P_n$  can be generated very efficiently (provided the function  $\phi$  is easy to compute) by running the PRNG for  $n + s - 3$  steps to obtain the point set

$$\{(u_i, \dots, u_{i+s-1}), i = 0, \dots, n - 2\}, \quad (1.11)$$

and adding the vector  $(g(0), \dots, g(0))$  to this set.

An important advantage of this construction is that it can be used easily on problems for which  $f$  depends on a random an unbounded number of variables. For this type of function, one needs a point set whose defining parameters are independent of  $s$ , and for which points of arbitrary size can be generated. Recurrence-based point sets satisfy these requirements since they are determined by the functions  $\phi, g$  and the set  $R$ , and those are independent of the dimension. Also, points of any size can be obtained by running the PRNG as long as required. Note that when  $s > n$ , the coordinates of each point in  $P_n$  have a periodic structure but by randomizing  $P_n$  as in Section 5, this periodicity disappears and the coordinates of each point become mutually independent, so one can have  $s \gg n$  without any problem.

The link between PRNG constructions and QMC point sets was pointed out by Niederreiter (1986). Not only similar constructions have been proposed in both fields but in addition, some of the selection criteria used in the two settings have interesting connections. Criteria measuring the uniformity of the  $s$ -dimensional point set  $P_n$  given in (1.10) are also

required in the PRNG context because in this case, this point set can be seen as the sampling space for the PRNG when it is used on a problem requiring  $s$  uniform numbers per run. See L'Ecuyer (1994) for details on this aspect of PRNGs and more.

We now describe two particular cases of this type of construction that provide an alternative way of generating Korobov-type lattice point sets (either standard or polynomial).

**Example 1** Let  $R = \mathbb{Z}_n$  for some prime  $n$ , define

$$r_i = \phi(r_{i-1}) \equiv ar_{i-1} \pmod{n}, \quad (1.12)$$

for some nonzero  $a \in \mathbb{Z}_n$ , and let  $g(r_i) = r_i/n$ . This type of PRNG is a *linear congruential generator* (LCG) (Knuth 1998; L'Ecuyer 1994), and the recurrence-based point set  $P_n$  is a Korobov lattice point set. When  $a$  is a *primitive element modulo  $n$*  (see Lidl and Niederreiter 1994 for the definition), the recurrence (1.12) has the maximal period of  $n - 1$  for any nonzero seed, and  $P_n$  can thus be generated efficiently using (1.12) and (1.11).

**Example 2** Let  $P(z) = z^k + a_1 z^{k-1} + \dots + a_k$  be a *primitive polynomial* over  $\mathbb{F}_2$  (see Lidl and Niederreiter 1994 for the definition),  $\nu$  be a positive integer,  $R = \mathbb{F}_2[z]/(P)$ , and

$$r_i(z) = \phi(r_{i-1}(z)) \equiv z^\nu r_{i-1}(z) \pmod{P(z)}. \quad (1.13)$$

Let  $g : \mathbb{F}_2[z]/(P) \rightarrow [0, 1)$  be given by the composition

$$g(r_i(z)) = \psi(r_i(z)/P(z)),$$

where  $\psi : \mathbb{F}_2((z^{-1})) \rightarrow [0, 1)$  is defined as the evaluation of a formal Laurent series over  $\mathbb{F}_2$  at  $z = 2$ ; that is,

$$\psi \left( \sum_{l=w}^{\infty} d_l z^{-l} \right) = \sum_{l=0}^{\infty} d_l 2^{-l}.$$

This type of PRNG is called a *polynomial LCG* (L'Ecuyer 1994; Tezuka 1995), and the recurrence-based point set  $P_n$  is a Korobov polynomial lattice point set. When  $\nu$  and  $2^k - 1$  are relatively prime, the recurrence (1.13) has maximal period length  $2^k - 1$ . Polynomial LCGs can also be obtained via Tausworthe-type linear feedback shift-register sequences (Tausworthe 1965): the idea is to use the recurrence

$$x_i = (a_1 x_{i-1} + \dots + a_k x_{i-k}) \pmod{2}$$

over  $\mathbb{F}_2$ , and to define the output at step  $i$  as

$$u_i = \sum_{l=1}^{\infty} x_{i\nu+l-1} 2^{-l},$$

which in practice is typically truncated to the word-length of the computer. Tezuka and L'Ecuyer give an efficient algorithm to generate the output  $u_i$  under some specific conditions (Tezuka and L'Ecuyer 1991; L'Ecuyer 1996).

Combining a few Tausworthe generators to define the output  $u_i$  can greatly help improving the quality of the associated set  $P_n$ , as explained by L'Ecuyer (1996), Tezuka and L'Ecuyer (1991) and Wang and Compagner (1993). Another way of enhancing the quality of point sets based on linear recurrences modulo 2 is to use *tempering transformations* (Matsumoto and Kurita 1994; Matsumoto and Nishimura 1998; L'Ecuyer and Panneton 2000). Note that these transformations generally destroy the lattice structure of  $P_n$  (Lemieux and L'Ecuyer 2001). However the point set obtained is still a digital net and therefore, it can be studied under this more general setting. Conditions that preserve the dimension-stationarity of  $P_n$  under these transformations are given by Lemieux and L'Ecuyer (2001). The idea of combining different constructions to build sets  $P_n$  with better equidistribution properties is also discussed in Niederreiter and Pirsic (2001) in the more general setting of digital nets.

### 3.4 Halton sequence

This sequence was introduced by Halton (1960) for constructing point sets of arbitrary length, and is a generalization of the one-dimensional van der Corput sequence (Niederreiter 1992b). Although it is not a digital sequence, it uses similar ideas and can thus be thought of as an ancestor of those sequences. For  $i \geq 0$ , the  $i^{th}$  point in the sequence is given by

$$\mathbf{u}_i = (\phi_{b_1}(i), \dots, \phi_{b_s}(i)),$$

where the integers  $b_1, \dots, b_s$  are typically chosen as the  $s$  first prime numbers sorted in increasing order, and  $\phi_b(i)$  is the *radical-inverse function* in base  $b$ , defined by

$$\phi_b(i) = \sum_{r=0}^{\infty} a_r b^{-r-1},$$

where the integers  $a_r \in \mathbb{Z}_b$  are the coefficients in the  $b$ -ary expansion of  $i$ , i.e.,  $i = \sum_{r=0}^{\infty} a_r b^r$ , as in the digital net definition.

### 3.5 Sequences of Korobov rules

With (infinite) digital sequences, one can always add new points to  $P_n$  until the estimation error is deemed small enough. The lattice point sets that we have discussed so far, on the other hand, contain only a fixed number of points. We now consider a method discussed by (Sobol' 1967; Maize 1981; Hickernell and Hong 1997; Hickernell, Hong, L'Ecuyer, and Lemieux 2001) for generating an infinite sequence of point sets  $\{P_{2^k}, k \geq 1\}$  in  $[0, 1]^s$  such that for  $k \geq 1$ ,  $P_{2^k}$  is a Korobov (polynomial) lattice point set and  $P_{2^k} \subseteq P_{2^{k+1}}$ .

Sequences based on nested Korobov lattice point sets can be constructed by choosing a nonzero odd integer  $a$  and defining

$$P_{2^k} = \left\{ \frac{i}{n} (1, a_k, \dots, a_k^{s-1}) \bmod 1, i = 0, \dots, n-1 \right\},$$

where  $a_k = a \bmod 2^k$ . Hickernell, Hong, L'Ecuyer, and Lemieux (2001) give tables of generators  $a$ , to build these sequences, that were chosen with respect to different selection criteria to be explained in Section 4.1.

One way of constructing a sequence based on Korobov polynomial lattice point sets is to choose a polynomial  $p_1(z)$  of degree 1 in  $\mathbb{F}_2[z]$  (i.e.,  $p_1(z) = z$  or  $p_1(z) = z + 1$ ), and a generating polynomial  $g(z) \in \mathbb{F}_2[z]$  such that  $\gcd(g(z), p_1(z)) = 1$ . Then define

$$P_{2^k} = \left\{ \varphi(\mathbf{v}(z)) : \mathbf{v}(z) = \frac{q(z)}{P_k(z)} (1, g_k(z), \dots, g_k^{s-1}(z)) \bmod \mathbb{F}_2[z], \right. \\ \left. q(z) \in \mathbb{F}_2[z]/(P_k) \right\},$$

where  $\varphi$  is defined as on page 15,  $P_k(z) = p_1^k(z)$ , and  $g_k(z) = g(z) \bmod P_k(z)$ . In this case, the sequence  $\{P_{2^k}, k = 1, 2, \dots\}$  turns out to be a special case of a digital sequence; see also Larcher (1998), page 187, for the particular case where  $p_1(z) = z$ , and for a more general setting based on irrational formal Laurent series.

### 3.6 Implementations

The points of a digital net  $P_n$  in base  $b$  can be generated efficiently using a Gray code. This idea was first suggested by Antonov and Saleev (1979) for the Sobol' sequence, and for other constructions by, e.g., Hong and Hickernell (2001), Pirsic and Schmid (2001), Tezuka (1995), and Bratley, Fox, and Niederreiter (1992). Assuming  $n = b^k$ , the idea is to modify the order in which the points are generated by replacing the digits  $(a_0, \dots, a_{k-1})$  from the  $b$ -ary expansion of  $i$  by the Gray code

$(g_0, \dots, g_{k-1})$  for  $i$ , which satisfies the following property: the Gray code for  $i$  and  $i + 1$  only differ in one position; if  $c$  is the smallest index such that  $a_c \neq b - 1$ , then  $g_c$  is the digit whose value changes, and it becomes  $g_c + 1$  in the Gray code for  $i + 1$  (Tezuka 1995, Theorem 6.6). This reduces the generation time because only a dot product of two vectors has to be performed in each dimension to generate a point.

It is sometimes suggested in the literature (Bratley, Fox, and Niederreiter 1992; Matoušek 1998; Acworth, Broadie, and Glasserman 1997; Fox 1986; Morokoff and Caflisch 1994) that for most low-discrepancy sequences, an initial portion of the sequence should be discarded because of the so-called “leading-zeros phenomenon”. For sequences such as Sobol’s that require initial parameters, this problem can be avoided (at least partially) by choosing these parameters carefully. Using a sufficiently large number of points and randomizing the point set can also help alleviate this problem. We refer the reader to the papers mentioned above for more details.

The FORTRAN code of Bratley and Fox (1988) can be used to generate Sobol’ sequence for  $s \leq 40$ , and is available from the Collective Algorithms of the ACM at [www.acm.org/calgo](http://www.acm.org/calgo), where the code of Fox (1986) for generating the Faure sequence can also be found, as well as a code for Niederreiter’s sequence (Bratley, Fox, and Niederreiter 1994). A code to generate Generalized Faure sequences (provided the matrices  $\mathbf{C}^j$  have been precomputed) is given by Tezuka (1995). Owen has written code to generate scrambled nets; it is available at [www-stat.stanford.edu/~owen/scramblednets](http://www-stat.stanford.edu/~owen/scramblednets), and is free for noncommercial uses. A C++ library called *libseq* was recently developed by Friedel and Keller (2001), in which they use efficient algorithms to generate *scrambled* digital sequences, Halton sequences, and other techniques such as Latin Hypercube and Supercube Sampling (Mckay, Beckman, and Conover 1979; Owen 1998a). This library can be found at [www.multires.caltech.edu/software/libseq/index.html](http://www.multires.caltech.edu/software/libseq/index.html). There are also a few commercial software packages to generate different QMC point sets; e.g., *QR Streams* at [www.mathdirect.com/products/qrn/](http://www.mathdirect.com/products/qrn/), and the *FinDer* software of Paskov and Traub (1995) at [www.cs.columbia.edu/~ap/html/finder.html](http://www.cs.columbia.edu/~ap/html/finder.html).

## 4. Measures of Quality

In this section, we present a number of criteria that have been proposed in the literature for measuring the uniformity (or non-uniformity) of a point set  $P_n$  in the unit hypercube  $[0, 1)^s$ , i.e., for measuring the

*discrepancy* between the distribution of the points of  $P_n$  and the uniform distribution, in the context of QMC integration.

In one dimension (i.e.,  $s = 1$ ), several such measures are widely used in statistics for testing the goodness of fit of a data set with the uniform distribution; e.g., the Kolmogorov-Smirnov (KS), Anderson-Darling, and chi-square test statistics. Chi-square tests also apply in more than one dimension, but their efficiency vanishes quickly as the dimension increases. The rectangular-star discrepancy discussed earlier turns out to be one possible multivariate generalization of the KS test statistic. Other variants of this measure are discussed by Niederreiter (1992b), and connections between discrepancy measures and goodness-of-fit tests used in statistics are studied in Hickernell (1999).

The asymptotic behavior of quality measures like the rectangular-star discrepancy is known for many constructions. For example, the Halton sequence has a rectangular-star discrepancy in  $O(n^{-1} \log^s n)$ , but with a hidden constant that grows superexponentially with  $s$ . This is also true for Sobol' sequence, but with a smaller hidden constant. By contrast, Faure and Niederreiter-Xing sequences are built so that this hidden constant goes to zero exponentially fast as the dimension  $s$  goes to infinity.

In practice however, as soon as the number of dimensions exceeds a few units, these general measures of discrepancy are unsatisfactory because they are too hard to compute. A more practical and less general approach is to define measures of uniformity that exploit the structure of a given class of highly structured point sets. Here we concentrate our discussion on these types of criteria, starting with criteria for standard lattice rules, then covering those for digital nets.

#### 4.1 Criteria for standard lattice rules

The criteria discussed here all relate to the dual lattice  $L_s^*$ , defined in Section 3. The first criterion we introduce has a nice geometrical interpretation, and is often used to measure the quality of LCGs through the so-called *spectral test* (Fishman and Moore III 1986; Fishman 1990; Knuth 1998; L'Ecuyer 1999). It was introduced by Coveyou and MacPherson (1967) and was used by Entacher, Hellekalek, and L'Ecuyer (2000) to choose lattice point sets. It amounts to computing the euclidean length  $l_s$  of the shortest nonzero vector in  $L_s^*$ , i.e.,

$$l_s = \min_{\mathbf{0} \neq \mathbf{h} \in L_s^*} \|\mathbf{h}\|_2.$$

The quantity  $l_s$  turns out to be equal to the inverse of the distance  $d_s$  between the successive hyperplanes in the family of most distant parallel

hyperplanes on which the points of  $L_s$  lie. This distance  $d_s$  should be as small as possible so there are no wide gaps in  $[0, 1]^s$  without any point from  $P_n$ . Equivalently,  $l_s$  should be as large as possible. Algorithms for computing  $l_s$  are discussed in, e.g., Dieter (1975), Fincke and Pohst (1985), Knuth (1998), and L'Ecuyer and Couture (1997). For instance, the dual of the basis shown in Figure 1.1 contains the shortest vector in  $L_s^*$ , given by  $\mathbf{h} = (5, 8)$ , so  $d_2 = 1/\sqrt{5^2 + 8^2} = 1/\sqrt{89}$  in this case.

This test can be applied not only to  $P_n$ , but also to any projection  $P_n(I)$  of  $P_n$ . More precisely, assume  $|I| = t$  and let  $L_I$  be the  $t$ -dimensional integration lattice such that  $P_n(I) = L_I \cap [0, 1]^t$ . The idea is to compute

$$l_I = \min_{\mathbf{0} \neq \mathbf{h} \in L_I^*} \|\mathbf{h}\|_2,$$

where  $L_I^* = \{\mathbf{h} \in \mathbb{Z}^t : \mathbf{h} \cdot \mathbf{v} \in \mathbb{Z} \text{ for all } \mathbf{v} \in L_I\}$ . To define a criterion in which  $l_I$  is computed for several subsets  $I$ , it is convenient to normalize  $l_I$  so that the same scale is used to compare the different projections. This can be achieved by using upper bounds  $l_t^*$  derived from the “best” possible lattice in  $t$  dimensions (not necessarily an integration lattice). Such bounds can be found in, e.g., Conway and Sloane (1999) and L'Ecuyer (1999). Criteria using these ideas have been used for measuring LCGs (Fishman and Moore III 1986; Fishman 1990; L'Ecuyer 1999), usually for subsets  $I$  containing successive indices, i.e., of the form  $I = \{1, 2, \dots, t\}$ . The following criterion is more general and has been used to construct tables of “good” Korobov lattice point sets by L'Ecuyer and Lemieux (2000):

$$M_{t_1, \dots, t_d} = \min \left( \min_{I=\{1, \dots, t\}, t \leq t_1} l_I/l_t^*, \min_{2 \leq j \leq d} \min_{I \in S(j, t_j)} l_I/l_j^* \right), \quad (1.14)$$

where  $S(j, t_j) = \{\{1, i_2, \dots, i_j\}, 1 < i_2 < \dots < i_j \leq t_j\}$ . Let the *range* of a subset  $I = \{i_1, \dots, i_t\}$  be defined as  $i_t - i_1 + 1$ . The criterion  $M_{t_1, \dots, t_d}$  considers projections over successive indices whose range is at most  $t_1$ ; over pairs of indices whose range is at most  $t_2$ ; over triplets of indices whose range is at most  $t_3$ , etc. Note that for dimension-stationary point sets, it is sufficient to do as in the definition of  $M_{t_1, \dots, t_d}$  and to only consider subsets  $I$  having 1 as their first index.

The next criterion is called the *Babenko-Zaremba index*, and is similar to  $l_s$  except that a different norm is used for measuring the vectors  $\mathbf{h}$  in  $L_s^*$ . It is defined as follows (Niederreiter 1992b):

$$\rho_s = \min_{\mathbf{0} \neq \mathbf{h} \in L_s^*} \|\mathbf{h}\|,$$

where  $\|\mathbf{h}\| = \prod_{j=1}^s \max(1, |h_j|)$ . It has been used by Maisonneuve (1972) to provide tables of “good” Korobov rules, but its computation is typ-

ically much more time-consuming than computing  $l_s$ . Both  $l_s$  and  $\rho_s$  can be seen as special cases of more general criteria such as general discrepancy measures defined, e.g., by Hickernell (1998b), Theorem 3.8, or the *generalized spectral test* of Hellekalek (1998), Definition 5.2. These criteria use a general norm to measure  $\mathbf{h}$  and apply to point sets that do not necessarily come from a lattice.

The following criterion, called *p-alpha*, uses the same norm as the Babenko-Zaremba index, but it sums a fixed power of the length of *all* vectors in the dual lattice  $L_s^*$  instead of only considering the smallest one. For an arbitrary integer  $\alpha > 1$ , it is defined as (Sloan and Joe 1994; Hickernell 1998b)

$$\mathcal{P}_\alpha = \sum_{\mathbf{0} \neq \mathbf{h} \in L_s^*} \|\mathbf{h}\|^{-\alpha}.$$

When  $\alpha$  is even, it simplifies to (Sloan and Joe 1994; Hickernell 1998b)

$$\mathcal{P}_\alpha = -1 + \frac{1}{n} \sum_{\mathbf{u} \in P_n} \prod_{j=1}^s \left[ 1 - \frac{(-1)^{\alpha/2} (2\pi)^\alpha}{\alpha!} B_\alpha(u_j) \right],$$

where  $B_\alpha(\cdot)$  is the Bernoulli polynomial of degree  $\alpha$  (Sloan and Joe 1994), and it can then be computed in  $O(ns)$  time. This criterion has been used in, e.g., Haber (1983), Sloan and Walsh (1990), and Sloan and Joe (1994) to choose lattice point sets.

The definition of  $\mathcal{P}_\alpha$  has been generalized by Hickernell (1998b) by the introduction of weights  $\beta_I$  that take into account the relative importance of each subset  $I$  (e.g., with respect to the ANOVA decomposition of  $f$ ). The generalization is defined as

$$\tilde{\mathcal{P}}_\alpha = \sum_{\mathbf{0} \neq \mathbf{h} \in L_s^*} \beta_{I(\mathbf{h})}^2 \|\mathbf{h}\|^{-\alpha},$$

where  $I(\mathbf{h}) = \{j : h_j \neq 0\}$  is the set of nonzero indices of  $\mathbf{h}$ . Assuming that  $\alpha$  is even and the  $\beta_I$  are product-type weights, i.e., that

$$\beta_I = \beta_0 \prod_{j \in I} \beta_j^\alpha, \quad (1.15)$$

this “weighted”  $\mathcal{P}_\alpha$  becomes

$$\tilde{\mathcal{P}}_\alpha = \beta_0^2 \left\{ -1 + \frac{1}{n} \sum_{\mathbf{u} \in P_n} \prod_{j=1}^s \left[ 1 - \frac{(-1)^{\alpha/2} (2\pi \beta_j)^\alpha}{\alpha!} B_\alpha(u_j) \right] \right\},$$

which can still be computed in  $O(ns)$  time. By letting  $\beta_j = 1$  for  $j = 0, \dots, s$ , we recover the criterion  $\mathcal{P}_\alpha$ . A normalized version of this

criterion and the quality measure  $M_{t_1}$  described above have been used for selecting generators  $a$  defining sequences of nested Korobov lattice point sets by Hickernell, Hong, L'Ecuyer, and Lemieux (2001). Note that  $\tilde{\mathcal{P}}_\alpha$  can be considered as a special case of the *weighted spectral test* of Hellekalek (1998), Definition 6.1. Various other measures of quality for lattice rules can be found in Niederreiter (1992b), Sloan and Joe (1994), Hickernell (1998b), and the references therein.

## 4.2 Criteria for digital nets

As we mentioned in Section 3.2, the  $t$ -value of a digital net is often used to assess its quality. To compute this value, one has to find  $q^* = k - t$ , which is the largest integer  $q$  such that for any integers  $q_j \geq 0$  satisfying  $q_1 + \dots + q_s \leq q$ , the vectors  $\{\mathbf{c}_r^j, r = 1, \dots, q_j, j = 1, \dots, s\}$  are linearly independent, where  $\mathbf{c}_r^j$  denote the  $r^{\text{th}}$  row of the generating matrix  $\mathbf{C}^j$  of  $P_n$ , for  $j = 1, \dots, s$ . Hence, computing  $t$  is typically quite time-consuming since for a given integer  $q$ , there are  $\binom{q+s-1}{q}$  different vectors  $(q_1, \dots, q_s)$  satisfying  $q_j \geq 0$  and  $q_1 + \dots + q_s = q$  (Pirsic and Schmid 2001).

If we define  $t_I$  as the value of the  $t$ -value for the projection  $P_n(I)$ , an equivalent definition of  $t$  is

$$t = \max_{\emptyset \neq I \subseteq \{1, \dots, s\}} t_I.$$

That is,  $t$  measures the regularity of all projections  $P_n(I)$  and returns the worst case. Inside this definition of  $t$ , we can also normalize the value of  $t_I$  so that projections over subspaces of different dimensions are judged more equitably, in the same way as the value  $l_t^*$  is used to normalize  $l_I$  in the criterion  $M_{t_1, \dots, t_d}$ . To do so, we can use the lower bound for  $t$  in  $s$  dimensions, given by Niederreiter and Xing (1998)

$$t_s^* = \max \left( k, \frac{s-1}{b} - \log_b \frac{(b-1)(s-1) + b + 1}{2} \right),$$

and define,

$$\tilde{t} = \max_{\emptyset \neq I \subseteq \{1, \dots, s\}} t_{|I|}^* / t_I, \quad (1.16)$$

for example. The idea behind this is that a large value of  $t_I$  for a low-dimensional subset  $I$  is usually worse than when  $I$  is high-dimensional and therefore it should be more penalized.

The definition of the  $t$ -value that we used so far is of a geometrical nature, similarly to the interpretation of  $l_s$  as being the inverse of the distance  $d_s$  between the hyperplanes of a lattice. Interestingly, just like

$l_s$  the  $t$ -value can also be related to the length of a shortest vector in a certain dual space (Tezuka 1995; Larcher, Niederreiter, and Schmid 1996; Niederreiter and Pirsic 2001; Couture and L'Ecuyer 2000) as we now explain. Our presentation follows Niederreiter and Pirsic (2001).

Let  $\mathbf{C}^1, \dots, \mathbf{C}^s$  be the  $\infty \times k$  generating matrices associated with a digital net in base  $b$  with  $n = b^k$  points. Let  $\mathbf{C}$  be the  $k \times sk$  matrix obtained by concatenating the transposed of the first  $k$  rows of each  $\mathbf{C}^j$ ; that is, if  $\tilde{\mathbf{C}}^j$  denotes the matrix containing the first  $k$  rows of  $\mathbf{C}^j$ , then

$$\mathbf{C} = ((\tilde{\mathbf{C}}^1)^T | \dots | (\tilde{\mathbf{C}}^s)^T).$$

The analysis of Niederreiter and Pirsic (2001) assumes that the generating matrices are  $k \times k$ , but we will explain shortly why it remains valid even if we start with  $\infty \times k$  matrices and truncate them.

Let  $\mathcal{C}_s^*$  be the null space of the row space of  $\mathbf{C}$ , i.e.,

$$\mathcal{C}_s^* = \{\mathbf{h} \in \mathbb{F}_b^{sk} : \mathbf{C}\mathbf{h} = \mathbf{0}\}. \quad (1.17)$$

We refer to  $\mathcal{C}_s^*$  as the *dual space* of the digital net  $P_n$  from now on. Define the following norm on  $\mathbb{F}_b^k$ : for any nonzero  $h_j = (h_{j,1}, \dots, h_{j,k}) \in \mathbb{F}_b^k$ , let  $v(h_j) = \max\{1 \leq l \leq k : h_{j,l} \neq 0\}$ , and  $v(0) = 0$ . Define the norm of a vector  $\mathbf{h} = (h_1, \dots, h_s) \in \mathbb{F}_b^{sk}$  by

$$V(\mathbf{h}) = \sum_{j=1}^s v(h_j). \quad (1.18)$$

The following result about the  $t$ -value is proved by Niederreiter and Pirsic (2001):

$$t = k + 1 - \min_{\mathbf{0} \neq \mathbf{h} \in \mathcal{C}_s^*} V(\mathbf{h}). \quad (1.19)$$

We now explain why this result is valid even if the matrices  $\mathbf{C}^j$  have been truncated to their first  $k$  rows. Let  $\mathcal{C}_{s,\text{tot}}^*$  denote the dual space that would be obtained without the truncation. Observe that by definition,

$$\min_{\mathbf{0} \neq \mathbf{h} \in \mathcal{C}_{s,\text{tot}}^* \setminus \mathcal{C}_s^*} V(\mathbf{h}) \geq k + 1.$$

Also, using Proposition 1 of Niederreiter and Pirsic (2001) and the fact that the dimension of the row space of  $\mathbf{C}$  is not larger than  $k$ , we have that

$$\min_{\mathbf{0} \neq \mathbf{h} \in \mathcal{C}_s^*} V(\mathbf{h}) \leq k + 1.$$

Therefore,

$$\min_{\mathbf{0} \neq \mathbf{h} \in \mathcal{C}_{s,\text{tot}}^*} V(\mathbf{h}) = \min \left( \min_{\mathbf{0} \neq \mathbf{h} \in \mathcal{C}_{s,\text{tot}}^* \setminus \mathcal{C}_s^*} V(\mathbf{h}), \min_{\mathbf{0} \neq \mathbf{h} \in \mathcal{C}_s^*} V(\mathbf{h}) \right) = \min_{\mathbf{0} \neq \mathbf{h} \in \mathcal{C}_s^*} V(\mathbf{h}),$$

which means that (1.19) is also true if we replace  $\mathcal{C}_s^*$  by  $\mathcal{C}_{s,\text{tot}}^*$ . From now on, we assume that  $\mathcal{C}_s^*$  actually represents the dual space obtained without truncating the generating matrices  $\mathbf{C}^j$ . Also, we view  $\mathcal{C}_s^*$  as a subspace of  $(\mathbb{F}_b[z])^s$ , which means that each element in  $\mathcal{C}_s^*$  is represented by a vector of polynomials  $\mathbf{h} = (h_1(z), \dots, h_s(z))$  over  $\mathbb{F}_b[z]$  and the norm  $v(\cdot)$  is defined by  $v(h_j(z)) = \deg(h_j(z)) + 1$ , with the convention that  $\deg(0) = -1$ .

In the special case where  $P_n$  is a polynomial lattice point set, the dual space  $\mathcal{C}_s^*$  corresponds to the dual lattice  $\mathcal{L}_s^*$ . If we define the norm

$$\|\mathbf{h}\| = \max_{1 \leq j \leq s} v(h_j(z)), \quad (1.20)$$

then

$$\ell_s = -1 + \min_{\mathbf{0} \neq \mathbf{h} \in \mathcal{C}_s^*} \|\mathbf{h}\|,$$

where  $\ell_s$  is the *resolution* of the polynomial lattice point set. This result is discussed by Couture, L'Ecuyer, and Tezuka (1993), Couture and L'Ecuyer (2000), Lemieux and L'Ecuyer (2001), and Tezuka (1995). The resolution is often used for measuring the quality of PRNGs based on linear recurrences modulo 2 such as Tausworthe generators (Tootill, Robinson, and Eagle 1973; L'Ecuyer 1996). From a geometrical point of view, the resolution is the largest integer  $l$  for which  $P_n$  is  $(l, \dots, l)$ -equidistributed. Obviously,  $\ell_s \leq \lfloor k/s \rfloor$  if  $n = b^k$ .

This concept can be extended from polynomial lattice point sets to general digital nets by replacing  $\mathcal{L}_s^*$  by  $\mathcal{C}_s^*$  above. More precisely, we have:

**Proposition 1** *Let  $P_n$  be a digital net in base  $b$ , and let  $\mathcal{C}_s^*$  be the dual space of  $P_n$ . The resolution  $\ell_s$  of  $P_n$  satisfies:*

$$\ell_s = -1 + \min_{\mathbf{0} \neq \mathbf{h} \in \mathcal{C}_s^*} \|\mathbf{h}\|,$$

where  $\|\mathbf{h}\|$  is defined as in (1.20).

*Proof:* The proof of this proposition requires results given in the forthcoming sections, and it can be found in the appendix.

The resolution can be computed for any projection of  $P_n$ : for  $I = \{i_1, \dots, i_t\}$ , let

$$\ell_I = \min_{\mathbf{0} \neq \mathbf{h} \in \mathcal{C}_I^*} \|\mathbf{h}\|,$$

where  $\mathcal{C}_I^*$  is the dual of the row space of the  $k \times tk$  matrix  $(\tilde{\mathbf{C}}_{i_1}^T | \dots | \tilde{\mathbf{C}}_{i_t}^T)$ . The following criterion has been proposed to select polynomial lattice

point sets (Lemieux and L'Ecuyer 2001), and it could also be used for any digital net:

$$\Delta_{t_1, \dots, t_d} = \max \left( \max_{I=\{1, \dots, t\}, t \leq t_1} (\ell_t^* - \ell_I), \max_{2 \leq j \leq d} \max_{I \in S(j, t_j)} (\ell_j^* - \ell_I) \right), \quad (1.21)$$

where the set  $S(j, t_j)$  has the same meaning as in the definition of the criterion  $M_{t_1, \dots, t_d}$  in (1.14).

Another criterion is the digital version of the quality measure  $\tilde{\mathcal{P}}_\alpha$ . It is closely related to the *dyadic dyaphony* (Hellekalek and Leeb 1997) and the weighted spectral test (Hellekalek 1998), and was introduced by Lemieux and L'Ecuyer (2001) for polynomial lattice point sets in base 2. It uses a norm  $W(\mathbf{h})$  defined as

$$W(\mathbf{h}) = \sum_{\substack{1 \leq j \leq s \\ h_j(z) \neq 0}} \deg(h_j(z)),$$

and is defined for any integer  $\alpha > 1$  and weights  $\beta_I$ , as

$$\tilde{\mathcal{D}}_\alpha = \sum_{\mathbf{0} \neq \mathbf{h} \in \mathcal{C}_s^*} \beta_{I(\mathbf{h})}^2 W(\mathbf{h}),$$

where  $I(\mathbf{h}) = \{j : h_j(z) \neq 0\}$ . In the special case where  $b = 2$ ,  $\alpha$  is even, and the weights  $\beta_I$  are product-type weights as in (1.15), we have that

$$\tilde{\mathcal{D}}_\alpha = \frac{\beta_0^2}{n} \sum_{i=0}^{n-1} \tilde{\xi}(\mathbf{u}_i),$$

where

$$\tilde{\xi}(\mathbf{u}) = -1 + \prod_{j=1}^t \left[ 1 + 2\beta_j^2 \left( 1 - 3 \cdot 2^{\lfloor \log_2 u_j \rfloor} \right) \right],$$

and it is assumed that  $2^{\lfloor \log_2 u_j \rfloor} = 0$  when  $u_j = 0$ .

We conclude this subsection by giving references where numerical values of the previous criteria are given for specific point sets. The value of the  $t$ -value for Sobol' sequence can be found in Niederreiter and Xing (1998) for dimensions  $s \leq 20$ ; these values are compared in that paper with those obtained from the improved base-2 sequences proposed by Niederreiter and Xing (1997, 1998). The "Salzburg Tables" given in (Pirsic and Schmid 2001) list optimal polynomial pairs  $(g(z), P(z))$  and their associated  $t$ -value, to build Korobov polynomial rules. Generalized Faure sequences are built so that  $t = 0$ , but with the drawback that the base  $b$  must be at least as large as the dimension. Hence only small

powers  $k$  of the bases are typically used in practice, and the quality of  $P_n$  is measured only for  $k$ -dimensional (or less) projections, where  $k$  is small. This illustrates the fact that for comparisons to be fair between different constructions, the value of the  $t$ -value should be considered in conjunction with the base  $b$ . Algorithms to compute  $t$  are discussed by Pirsic and Schmid (2001).

The resolution has been used by Sobol' for finding optimal values to initialize recurrences defining the direction numbers in his construction (Sobol' 1967; Sobol' 1976). More precisely, his *Property A* means that the first  $2^s$  points of the sequence have the maximal resolution of 1, and his *Property A'* means that the first  $2^{2s}$  points have the maximal resolution of 2.

Following ideas from Morokoff and Caflisch (1994) and Cheng and Druzdzel (2000), a criterion related to  $\Delta_{t_1, \dots, t_d}$  is used to find initial direction numbers for Sobol' sequence in dimensions  $s > 40$  in the forthcoming **RandQMC** library (Lemieux, Cieslak, and Luttmer 2001); the maximum in (1.21) is taken over all two-dimensional subsets  $I$  of the form  $I = \{i_1, j\}$ , where  $j > 40$  is the dimension for which we want to find initial direction numbers, and  $j - 8 \leq i_1 \leq j - 1$ . Examples of parameters for polynomial lattice point sets chosen with respect to  $\Delta_{t_1, \dots, t_4}$  for different values of  $t_1, \dots, t_4$  are given by Lemieux and L'Ecuyer (2001).

## 5. Randomizations

Once a construction is chosen for  $P_n$  and the approximation  $Q_n$  given by (1.2) is computed, it is usually important to have an estimate of the error  $|Q_n - \mu|$ . For that purpose, upper bounds of the form (1.3) are not very useful since they are usually much too conservative, in addition to being hard to compute and restricted to a possibly small set of functions. Instead, one can randomize the set  $P_n$  so that: 1) each point in the randomized point set  $\tilde{P}_n$  has a uniform distribution over  $[0, 1]^s$ ; 2) the regularity (or low-discrepancy property) of  $P_n$ , as measured by a specific quality criterion, is preserved under the randomization. The first property guarantees that the approximation

$$\hat{\mu}_{\text{RQMC}} = \frac{1}{n} \sum_{\mathbf{u}_i \in \tilde{P}_n} f(\mathbf{u}_i)$$

is an unbiased estimator of  $\mu$ . When the second property holds, the variance of the estimator  $\hat{\mu}_{\text{RQMC}}$  can usually be expressed in a way that establishes a relation between the optimization of the criterion whose value is preserved under the randomization, and the minimization of the estimator's variance. Hence these two properties help viewing random-

ized QMC methods as variance reduction techniques that preserve the unbiasedness of the MC estimator. In practice, the variance of  $\hat{\mu}_{\text{RQMC}}$  can be estimated by generating i.i.d. copies of  $\hat{\mu}_{\text{RQMC}}$  through independent replications of the randomization. This estimator can then be compared with the estimated variance of the MC estimator to assess the effectiveness of QMC for any particular problem.

We now describe some randomizations having these two properties and that have been proposed in the literature for the constructions presented in the preceding section.

### 5.1 Random shift modulo 1

The following randomization has originally been proposed by Cranley and Patterson (1976) for standard lattice rules. Some authors suggested that it could also be used for other low-discrepancy point sets (Tuffin 1996; Morohosi and Fushimi 2000).

Let  $P_n = \{\mathbf{u}_i, i = 0, \dots, n-1\}$  be a given point set in  $[0, 1)^s$ , and  $\Delta$  an  $s$ -dimensional random vector uniformly distributed over  $[0, 1)^s$ . The *randomly shifted estimator* based on  $P_n$  is defined as

$$\hat{\mu}_{\text{sh}} = \frac{1}{n} \sum_{i=0}^{n-1} f((\mathbf{u}_i + \Delta) \bmod 1). \quad (1.22)$$

When  $P_n$  is a lattice point set, the length  $l_I$  of the shortest vector associated with any projection  $P_n(I)$  is preserved under this randomization. An explicit expression for the variance of  $\hat{\mu}_{\text{sh}}$  in that case will be given in Section 6.1.

With this randomization, each shifted point  $\tilde{\mathbf{u}}_i = ((\mathbf{u}_i + \Delta) \bmod 1)$  is uniformly distributed over  $[0, 1)^s$ . Therefore, even if the dimension  $s$  is much larger than the number of points  $n$  and if many coordinates are equal within a given point  $\mathbf{u}_i$  (for instance, when  $P_n$  comes from a LCG with a small period, as in Section 3.3), these coordinates become mutually independent *after* the randomization. Hence each point has the same distribution as in the MC method; the difference with MC is that the  $n$  points of the shifted lattice are not independent. See L'Ecuyer and Lemieux (2000), Section 10.3, for a concrete numerical example with  $n = 1021$  and  $s = 10000$ . These properties also hold for the other randomizations described below.

### 5.2 Digital $b$ -ary shift

When  $P_n$  is a digital net in base  $b$ , the counterpart of the previous method is to consider the  $b$ -ary expansion of the random vector  $\Delta$ , and to add it to each point of  $P_n$  using operations over  $\mathbb{F}_b$ . More precisely,

if  $\Delta = (\Delta_1, \dots, \Delta_s)$  and

$$\Delta_j = \sum_{l=1}^{\infty} d_{j,l} b^{-l}, \quad u_{ij} = \sum_{l=1}^{\infty} u_{ijl} b^{-l},$$

we compute

$$\mathbf{u}_i \oplus \Delta = (\tilde{u}_{i1}, \dots, \tilde{u}_{is}), \quad (1.23)$$

where

$$\tilde{u}_{ij} = \sum_{l=1}^{\infty} (u_{ij,l} + d_{j,l}) b^{-l}, \quad (1.24)$$

and let  $\tilde{P}_n = \{\tilde{\mathbf{u}}_i, i = 0, \dots, n-1\}$ .

This randomization was suggested to us by Raymond Couture for point sets based on linear recurrences modulo 2 (see also Lemieux and L'Ecuyer 2001). It is also used in an arbitrary base (along with other more time-consuming randomizations) in Hong and Hickernell (2001) and Matousěk (1998) as we will see in Section 5.4. It is best suited for digital nets in base  $b$ , and its application preserves the resolution and the  $t$ -value of any projection.

### 5.3 Scrambling

This randomization has been proposed by Owen (1995), and it also preserves the  $t$ -value of a digital net and its resolution, for any projection. It works as follows: in each dimension  $j = 1, \dots, s$ , partition the interval  $[0, 1)$  in  $b$  equal parts and permute them uniformly and randomly; then, partition each of these sub-intervals into  $b$  equal parts and permute them uniformly and randomly; etc. More precisely, to scramble  $L$  digits one needs to randomly and uniformly generate several independent permutations  $\pi_{i_1, \dots, i_t}^j$  of the integers  $[0 \dots b-1]$  (assuming a specific bijection has been chosen to identify the elements in  $\mathbb{F}_b$  with those in  $\mathbb{Z}_b$ , if  $b$  is not prime), where  $i_1 = 0$ ,  $i_t \in \{0, \dots, b-1\}$ ,  $1 \leq t \leq L$ ,  $1 \leq j \leq s$ , and compute

$$\tilde{u}_{i,j} = \sum_{l=1}^{\infty} \tilde{u}_{i,j,l} b^{-l},$$

where

$$\tilde{u}_{i,j,l} = \begin{cases} \pi_0^j(u_{i,j,1}) & \text{for } l = 1, \\ \pi_{0, u_{i,j,1}, \dots, u_{i,j,l-1}}^j(u_{i,j,l}) & \text{for } 1 < l \leq L, \\ u_{i,j,l} & \text{for } l > L. \end{cases}$$

In practice,  $L$  may be chosen equal to the word-length of the computer, and the digits  $u_{i,j,l}$  for  $l > L$  are then dropped. However, as Matousěk

(1998) points out, if  $n = b^k$  and no two points have the same first  $k$  digits in each dimension (i.e., for each  $j$ , the unidimensional projection  $P_n(\{j\})$  has a maximal resolution of  $k$ ), then the permutations after level  $k$  are independent for each point and therefore, the random digits  $\tilde{u}_{i,j,l}$  for  $l > k$  can be generated uniformly and independently over  $\mathbb{Z}_b$ . Hence in this case we do not need to generate any permutation after level  $k$ . Owen (1995), Section 3.3, has suggested a similar implementation.

When  $b$  and  $s$  are large, the amount of memory required for storing all the permutations becomes very large, and only a *partial scrambling* might then be feasible; that is, scramble  $m < k$  digits, and generate the remaining ones randomly and uniformly (Fox 1999), or use the permutations for the  $m$ th digit (Tan and Boyle 2000). A clever way of avoiding storage problems is discussed by Matoušek (1998), and a related idea is used in Morohosi's code (which can be found at [www.misojiro.t.u-tokyo.ac.jp/~morohosi](http://www.misojiro.t.u-tokyo.ac.jp/~morohosi)) for scrambling Faure sequences. The idea is to avoid storing all the permutations by reinitializing appropriately the underlying PRNG so that the permutations can be regenerated as they are needed. This is especially useful when the base  $b$  is large, which happens when Faure sequences are used in large or even moderate dimensions.

## 5.4 Random Linear Scrambling

Matoušek (1998) proposes an alternative scrambling method that does not require as much randomness and storage. It borrows ideas from the scrambling technique and transformations proposed by Faure and Tezuka (2001) and Tezuka (1995). This method is also discussed by Hong and Hickernell (2001), where it is called "Owen-Scrambling"; our presentation follows theirs but we prefer the name used by Matoušek to avoid any confusion with the actual scrambling proposed by Owen. The idea is to generate  $s$  nonsingular lower-triangular  $k \times \infty$  matrices  $\mathbf{L}_1, \dots, \mathbf{L}_s$  with elements chosen randomly, independently, and uniformly over  $\mathbb{F}_b$  (the elements on the main diagonal of each  $\mathbf{L}_j$  are chosen over  $\mathbb{F}_b \setminus \{0\}$ ), where  $k$  is such that  $n = b^k$ , and  $s$   $\infty$ -dimensional vectors  $\mathbf{d}_1, \dots, \mathbf{d}_s$  with entries independently and uniformly distributed over  $\mathbb{F}_b$ . The digits  $\tilde{u}_{i,j,1}, \tilde{u}_{i,j,2}, \dots$  of a randomized coordinate  $\tilde{u}_{i,j}$  are then obtained as

$$\begin{pmatrix} \tilde{u}_{i,j,1} \\ \tilde{u}_{i,j,2} \\ \vdots \end{pmatrix} = \mathbf{L}_j \begin{pmatrix} u_{i,j,1} \\ u_{i,j,2} \\ \vdots \end{pmatrix} + \mathbf{d}_j,$$

where all operations are performed in  $\mathbb{F}_b$ .

## 5.5 Others

We briefly mention some other ideas that can be used to randomize QMC point sets. In addition to the linear scrambling, Matousěk (1998) proposes randomization techniques for digital sequences that are easier to generate than the scrambling method, while retaining enough randomness for the purpose of some specific theoretical analyses. Hong and Hickernell suggest another form of linear scrambling that incorporates transformations proposed by Faure and Tezuka (2001). Randomizations that use permutations in each dimension to reorder the Halton sequence are discussed by Braaten and Weller (1979) and Morokoff and Caflisch (1994). Wang and Hickernell (2000) propose to randomize this sequence by randomly generating its starting point.

Some authors (Ökten 1996; Spanier 1995) suggest partitioning the set of dimensions  $\{1, \dots, s\}$  into two subsets (typically, of successive indices, i.e.,  $\{1, \dots, t\}$  and  $\{t+1, \dots, s\}$ ), and then to use a QMC method (randomized or not) on one subset and MC on the other one. One of the justifications for this approach is that some digital nets (e.g., Sobol' sequence) are known to have projections  $P_n(I)$  with better properties when  $I$  contains small indices; this suggests using QMC on the first few dimensions and MC on the remaining ones. However, this argument becomes irrelevant if a dimension-stationary point set is used. More importantly, if the QMC point set is randomized and can be shown (or presumed) to do no worse than MC in terms of its variance, then there is no advantage or “safety net” gained by using MC on one part of the problem. Estimators obtained by “padding” randomized QMC point sets with MC are analyzed by Owen (1998a) and Fox (1999). Owen (1998a) discusses some other padding techniques, as well as a method called *Latin Supercube Sampling* to handle very high-dimensional problems.

## 6. Error and Variance Analysis

In this section, we study the error for the approximations  $Q_n$  based on low-discrepancy point sets, and the variance of estimators obtained by randomizing these sets. All the results mentioned here are obtained by using a particular basis for  $\mathcal{L}^2([0, 1]^s)$ , the set of square-integrable functions over  $[0, 1]^s$ , to expand the integrand  $f$ . In each case, the basis is carefully chosen, depending on the structure of  $P_n$ , so that the behavior of the approximation  $Q_n$  on the basis functions is easy to analyze.

The results presented here describe, for a given square-integrable function, the variance of different randomized QMC estimators. This is one possible approach for analyzing the performance of these estimators;

different aspects of randomized QMC methods have been studied in the literature. In particular, Hickernell and Woźniakowski (2001) have recently studied the behavior of different types of error on Hilbert spaces  $\mathcal{F}$  of integrands in a general setting that includes MC, scrambled nets, and randomly shifted lattices. They consider the “worst-case error” (error on the worst integrand for a given QMC approximation), the “random-case error” (expected error, e.g., in the root mean square sense, of a randomized QMC estimator on a given integrand), and the “average-case error” (average error over  $\mathcal{F}$ ), and they give explicit formulas for these errors. This type of analysis provides useful connections between variance expressions such as those studied here (which corresponds to the random-case setting), and the discrepancy measures discussed in the previous section (which corresponds to the worst-case setting). Specific results for scrambled-nets estimators are given by Heinrich, Hickernell, and Yue (2001b).

### 6.1 Standard Lattices and Fourier Expansion

For standard lattice rules, the following result suggests that expanding  $f$  in a Fourier series is appropriate for error and variance analysis. Recall that the Fourier basis for  $\mathcal{L}^2([0, 1]^s)$  is orthonormal and given by  $\{e^{2\pi i \mathbf{h} \cdot \mathbf{u}}, \mathbf{h} = (h_1, \dots, h_s) \in \mathbb{Z}^s\}$ , where  $\iota = \sqrt{-1}$ , and  $\mathbf{h} \cdot \mathbf{u} = \sum_{j=1}^s h_j u_j$ .

**Lemma 1** (Sloan and Joe 1994, Lemma 2.7) *If  $P_n = \{\mathbf{u}_i, i = 0, \dots, n-1\}$  is a lattice point set, then for any  $\mathbf{h} \in \mathbb{Z}^s$ ,*

$$\frac{1}{n} \sum_{i=0}^{n-1} e^{2\pi i \mathbf{h} \cdot \mathbf{u}_i} = \begin{cases} 1 & \text{if } \mathbf{h} \in L_s^*, \\ 0 & \text{otherwise.} \end{cases}$$

Hence, the lattice rule integrates a basis function  $e^{2\pi i \mathbf{h} \cdot \mathbf{u}}$  with no error when  $\mathbf{h} \notin L_s^*$  or  $\mathbf{h} = \mathbf{0}$ , and with error 1 otherwise. Using this, we get the following result:

**Proposition 2** *Let  $P_n$  be a lattice point set,  $\hat{\mu}_{\text{sh}}$  be defined as in (1.22), and*

$$\hat{f}(\mathbf{h}) = \int_{[0,1]^s} f(\mathbf{u}) e^{-2\pi i \mathbf{h} \cdot \mathbf{u}} d\mathbf{u}$$

*be the Fourier coefficient of  $f$  evaluated in  $\mathbf{h}$ . (From Sloan and Joe 1994) If  $f$  has an absolutely convergent Fourier series, then*

$$Q_n - \mu = \sum_{\mathbf{0} \neq \mathbf{h} \in L_s^*} |\hat{f}(\mathbf{h})|.$$

(From L'Ecuyer and Lemieux 2000) If  $f \in \mathcal{L}^2([0, 1]^s)$ , then

$$\text{Var}(\hat{\mu}_{\text{sh}}) = \sum_{\mathbf{0} \neq \mathbf{h} \in L_s^*} |\hat{f}(\mathbf{h})|^2, \quad (1.25)$$

whereas for the MC estimator  $\hat{\mu}_{\text{MC}}$  based on  $n$  points,

$$\text{Var}(\hat{\mu}_{\text{MC}}) = \frac{\sigma^2}{n} = \frac{1}{n} \sum_{\mathbf{0} \neq \mathbf{h} \in \mathbb{Z}^s} |\hat{f}(\mathbf{h})|^2.$$

The result (1.25) was proved independently by Tuffin (1998), but under the stronger assumption that  $f$  has an absolutely convergent Fourier series. Notice that by contrast with the MC estimator, there is no factor of  $1/n$  that multiplies the sum of squared Fourier coefficients for the randomly shifted lattice rule estimator  $\hat{\mu}_{\text{sh}}$ . Hence in the worst case, the variance of  $\hat{\mu}_{\text{sh}}$  could be  $n$  times as large as the MC estimator's variance. This worst case corresponds to an extremely unlucky pairing of function and point set for which  $\hat{f}(\mathbf{h}) = 0$  for all  $\mathbf{h} \notin L_s^*$ . However, in the expression for the variance of  $\hat{\mu}_{\text{sh}}$  the coefficients are summed only over the dual lattice  $L_s^*$ , which contains  $n$  times less points than the set  $\mathbb{Z}^s$  over which the sum is taken in the MC case. Therefore, if the dual lattice  $L_s^*$  is such that the squared Fourier coefficients are smaller “on average” over  $L_s^*$  than over  $\mathbb{Z}^s$ , then the variance of  $\hat{\mu}_{\text{sh}}$  will be smaller than the variance of  $\hat{\mu}_{\text{MC}}$ .

From the results given in the previous proposition, different bounds on the error and variance can be obtained by making additional assumptions on the integrand  $f$  (Sloan and Joe 1994; Hickernell 1998b; Hickernell 2000). Most of these bounds involve the quality measures  $\mathcal{P}_\alpha$  or  $\tilde{\mathcal{P}}_\alpha$ . Hence a point set  $P_n$  that minimizes one of these two criteria minimizes a bound on the error or variance for the class of functions for which those bounds hold. Such analyses often provide arguments in favor of these criteria. A different type of analysis, based on the belief that the largest squared Fourier coefficients tend to be associated with “short vectors”  $\mathbf{h}$ , corresponding to the low frequency terms of  $f$ , suggests that the lattice point set should be chosen so that  $L_s^*$  does not contain those “short” vectors. From this point of view, a criterion like  $M_{t_1, \dots, t_d}$  seems appropriate since it makes sure that  $L_s^*$  does not contain vectors with a small euclidean length. This criterion also has the advantage of being usually much faster to compute than  $\mathcal{P}_\alpha$  or  $\rho_s$  (Entacher, Hellekalek, and L'Ecuyer 2000; Hickernell, Hong, L'Ecuyer, and Lemieux 2001).

## 6.2 Digital Nets and Haar or Walsh Expansions

Recall that digital nets are usually built so as to satisfy different equidistribution properties with respect to partitions of the unit hypercube  $[0, 1]^s$  into  $b$ -ary boxes. For this reason, it is convenient to use a basis consisting of step functions that are constant over  $b$ -ary boxes for studying their associated error and variance. Both Walsh and Haar basis functions have this property. In addition, the Walsh functions form an orthonormal basis of  $\mathcal{L}^2([0, 1]^s)$ .

**Scrambled-type estimators.** We first define the Haar basis functions in base  $b$ , following the presentation of Owen (1997) and (Heinrich, Hickernell, and Yue 2001a). Let  $I \subseteq \{1, \dots, s\}$ ,  $\kappa_I = (\kappa_j)_{j \in I}$  be a vector of positive integers,  $\tau = (\tau_j)_{j \in I}$  be a vector such that  $0 \leq \tau_j < b^{\kappa_j}$ , and  $\gamma = (\gamma_j)_{j \in I}$  be such that  $0 \leq \gamma_j < b$ . A *multivariate Haar wavelet basis function* is defined as

$$\psi_{\kappa_I, \tau, \gamma}(\mathbf{u}) = b^{(|\kappa_I| - |I|)/2} \prod_{j \in I} \left( b \mathbf{1}_{[b^{\kappa_j+1} u_j] = b\tau_j + \gamma_j} - \mathbf{1}_{[b^{\kappa_j} u_j] = \tau_j} \right),$$

where  $|\kappa_I| = \sum_{j \in I} \kappa_j$ . Now consider the part of the Haar expansion of  $f$  that depends on the basis functions associated with a given vector  $\kappa_I$ , i.e., let

$$\nu_{\kappa_I}(\mathbf{u}) = \sum_{\tau} \sum_{\gamma} \langle \psi_{\kappa_I, \tau, \gamma}(\mathbf{u}), f \rangle \psi_{\kappa_I, \tau, \gamma}(\mathbf{u}),$$

where

$$\langle \psi_{\kappa_I, \tau, \gamma}(\mathbf{u}), f \rangle = \int_{[0, 1]^s} f(\mathbf{u}) \psi_{\kappa_I, \tau, \gamma}(\mathbf{u}) d\mathbf{u}.$$

The function  $\nu_{\kappa_I}(\mathbf{u})$  is also a step function, which is constant within the  $b$ -ary boxes obtained by partitioning  $[0, 1]^s$  into  $b^{\kappa_j}$  equal intervals along the  $j$ th axis, for each  $j \in I$ . Owen (1997) shows that the variance of the estimator  $\hat{\mu}_{\text{scr}}$  based on a scrambled digital net with  $n$  points is given by

$$\text{Var}(\hat{\mu}_{\text{scr}}) = \frac{1}{n} \sum_I \sum_{\kappa_I} \Gamma_{\kappa_I} \sigma_{\kappa_I}^2,$$

where

$$\sigma_{\kappa_I}^2 = \text{Var}(\nu_{\kappa_I}(\mathbf{u})), \quad (1.26)$$

and  $\Gamma_{\kappa_I}$  depends on the equidistribution properties of the digital net and the definition of the scrambling. Assuming that the  $t$ -value of the net is

$t$  and  $n = b^k$ , we have (Owen 1998b):

$$\Gamma_{\kappa_I} \begin{cases} = 0 & \text{if } |\kappa_I| + |I| \leq k - t, \\ \leq b^t \left( \frac{b^{|I|} + (b-2)^{|I|}}{2(b-1)^{|I|}} \right) & \text{if } |\kappa_I| > k - t, \\ \leq b^t \left( \frac{b+1}{b-1} \right)^{|I|} & \text{otherwise.} \end{cases} \quad (1.27)$$

Using this, Owen obtains the following bound on the variance of the scrambled-net estimator:

**Proposition 3** (*Owen 1998b, Theorem 1*) *Let  $\hat{\mu}_{\text{scr}}$  be the estimator constructed from a scrambled digital net with  $n = b^k$  points. For any square-integrable function,*

$$\text{Var}(\hat{\mu}_{\text{scr}}) \leq \frac{b^t}{n} \left( \frac{b+1}{b-1} \right)^s \sigma^2. \quad (1.28)$$

Hence the variance of the scrambled-net estimator cannot be larger than the MC estimator's variance, up to a certain constant (independent of  $n$  but growing exponentially with  $s$  for a fixed base  $b$ ). For a  $(0, k, s)$ -net, which implies that  $b \geq s$ , this constant can be bounded by  $e \approx 2.718$  (Owen 1997). In the case where  $f$  satisfies certain smoothness properties (its mixed partial derivatives satisfy a Lipschitz condition), Owen shows that

$$\sigma_{\kappa_I}^2 = O(b^{-2|\kappa_I|}).$$

Under this assumption, he obtains a bound in  $O(n^{-3}(\log n)^s)$  for the variance of the scrambled-net estimator. Other results on the asymptotic properties of the scrambled-net estimator, that use Haar wavelets, can be found in Heinrich, Hickernell, and Yue (2001a), Heinrich, Hickernell, and Yue (2001b), and the references cited there. Haar series are also considered in the context of QMC integration by, e.g., Entacher (1997) and Sobol' (1969).

An important point to mention is that the scrambling is not the only randomization for which (1.27) holds; the result is valid for any randomization satisfying the following properties (Hong and Hickernell 2001; Matousěk 1998):

- 1 each point  $\tilde{\mathbf{u}}_i$ ,  $i = 0, \dots, n-1$  in the randomized point set  $\tilde{P}_n$  is uniformly distributed over  $[0, 1)^s$ ;
- 2 for  $0 \leq i, r < n$  and  $1 \leq j \leq s$ , if  $u_{i,j,l} = u_{r,j,l}$  for  $l = 1, \dots, k$  but  $u_{i,j,k+1} \neq u_{r,j,k+1}$ , then

- (a)  $\tilde{u}_{i,j,l} = \tilde{u}_{r,j,l}$  for  $l = 1, \dots, k$ ;
- (b)  $(\tilde{u}_{i,j,k+1}, \tilde{u}_{r,j,k+1})$  is uniformly distributed over  $\{(a_1, a_2) \in \mathbb{F}_b^2 : a_1 \neq a_2\}$ ;
- (c)  $(\tilde{u}_{i,j,p}, \tilde{u}_{i,j,q})$  are uncorrelated, for any  $p, q > k + 1$ .

Hong and Hickernell (2001) have shown that the random linear scrambling mentioned in Section 5.4 satisfies these properties, so the bound (1.28) given in Proposition 3 holds for linearly scrambled estimators as well. This is interesting since this method has a faster implementation than the scrambling of Section 5.3. Note that the digital  $b$ -ary shift does not satisfy 2(b) since the digits  $\tilde{u}_{i,j,k+1}, \tilde{u}_{r,j,k+1}$  are such that  $\tilde{u}_{r,j,k+1} - \tilde{u}_{i,j,k+1} = u_{r,j,k+1} - u_{i,j,k+1} \in \mathbb{F}_b$ .

**Digitally  $b$ -ary shifted estimators.** To study the variance of a  $b$ -ary shifted digital net we use a Walsh expansion for  $f$ . Walsh series have also been used to analyze the error produced by (non-randomized) digital nets by Larcher, Niederreiter, and Schmid (1996), Larcher, Lauss, Niederreiter, and Schmid (1996), and Larcher and Pirsic (1999). In the presentation of the forthcoming results, the vector  $\mathbf{h}$  will be used both to represent elements of  $\mathbb{N}_0^s$ , where  $\mathbb{N}_0 = \mathbb{N} \cup \{0\}$ , and elements of the ring  $(\mathbb{F}_b[z])^s$ . When required, we will use the bijection  $\zeta : \mathbb{N}_0^s \rightarrow (\mathbb{F}_b[z])^s$  defined by  $\zeta(h_1, \dots, h_s) = (\sum_{l=1}^{\infty} h_{1,l} z^{l-1}, \dots, \sum_{l=1}^{\infty} h_{s,l} z^{l-1})$ , where  $h_j = \sum_{l=1}^{\infty} h_{j,l} b^{l-1}$  for each  $j$ , to go back and forth between these two spaces.

For any  $\mathbf{h} \in \mathbb{N}_0^s$ , the Walsh basis function in  $\mathbf{h}$  is defined as

$$\phi_{\mathbf{h}}(\mathbf{u}) = e^{2\pi i \mathbf{h} \cdot \mathbf{u} / b},$$

where  $\mathbf{h} \cdot \mathbf{u} = \sum_{j=1}^s h_j \cdot u_j = \sum_{j=1}^s \sum_{l=1}^{\infty} h_{j,l} u_{j,l}$ , and the coefficients  $h_{j,l}$  and  $u_{j,l}$  are such that  $h_j = \sum_{l=1}^{\infty} h_{j,l} b^{l-1}$  and  $u_j = \sum_{l=1}^{\infty} u_{j,l} b^{-l}$ , and all operations are performed in  $\mathbb{F}_b$ . For  $\mathbf{h}_1, \mathbf{h}_2 \in \mathbb{N}_0^s$ , we have that

$$\phi_{\mathbf{h}_1 \oplus \mathbf{h}_2}(\mathbf{u}) = \phi_{\mathbf{h}_1}(\mathbf{u}) \phi_{\mathbf{h}_2}(\mathbf{u}),$$

where  $\mathbf{h}_1 \oplus \mathbf{h}_2$  corresponds to a digit-by-digit addition over  $\mathbb{F}_b$  (as if we were adding the corresponding elements in  $(\mathbb{F}_b[z])^s$ ). See Larcher, Niederreiter, and Schmid (1996) and Larcher and Pirsic (1999) for more information on generalized definitions of Walsh series in the context of QMC integration.

Let  $\tilde{f}(\mathbf{h})$  denote the Walsh coefficient of  $f$  in  $\mathbf{h}$ , that is

$$\tilde{f}(\mathbf{h}) = \int_{[0,1]^s} f(\mathbf{u}) e^{-2\pi i \mathbf{h} \cdot \mathbf{u} / b} d\mathbf{u}. \quad (1.29)$$

The following result may be interpreted as the digital counterpart of the result stated in Lemma 1. Recall that  $\mathcal{C}_s^*$  denotes the dual space of  $P_n$ .

**Lemma 2** *Let  $P_n$  be a digital net in base  $b$  with  $n = b^k$ . For any  $\mathbf{h} \in \mathbb{N}_0^s$ , we have*

$$\frac{1}{n} \sum_{i=0}^{n-1} e^{2\pi i \mathbf{h} \cdot \mathbf{u}_i / b} = \begin{cases} 1 & \text{if } \zeta(\mathbf{h}) \in \mathcal{C}_s^*, \\ 0 & \text{otherwise.} \end{cases}$$

Larcher, Niederreiter, and Schmid (1996) have shown that the above sum is 0 when  $\mathbf{h}$  satisfies  $V(\mathbf{h}) \leq k - t$ .

*Proof:* If  $\mathbf{h} \in \mathcal{C}_s^*$ , then  $\mathbf{h} \cdot \mathbf{u}_i = 0$  since  $(u_{i,1,1}, u_{i,1,2}, \dots, u_{i,s,1}, u_{i,s,2}, \dots)$  is in the row space of  $\mathbf{C}$  for all  $i = 0, \dots, n-1$ , and the result follows easily. If  $\mathbf{h} \notin \mathcal{C}_s^*$ , then  $\mathbf{C}\mathbf{h} = \mathbf{y} \neq \mathbf{0}$ , where  $\mathbf{y} \in \mathbb{F}_b^k$ . We are interested in the scalar product  $\mathbf{h} \cdot \mathbf{u}_i$  for  $i = 0, \dots, n-1$ . Notice that  $\{\mathbf{h} \cdot \mathbf{u}_i, i = 0, \dots, n-1\} = \{\mathbf{y}^T \mathbf{x}, \mathbf{x} \in \mathbb{F}_b^k\}$ , which is the image of  $\mathbb{F}_b^k$  under the application of a mapping that corresponds to the multiplication by  $\mathbf{y}^T$ . Since  $\mathbf{y} \neq \mathbf{0}$ , the dimension of this image is 1 and the dimension of the kernel of this mapping is thus  $k-1$ . Hence each element in  $\mathbb{F}_b$  has  $b^{k-1}$  pre-images in  $\mathbb{F}_b^k$  under this mapping, and therefore as a multiset  $\{\mathbf{h} \cdot \mathbf{u}_i, i = 0, \dots, n-1\}$  contains  $b^{k-1}$  copies of each element of  $\mathbb{F}_b$ . Using this and the fact that

$$\sum_{j=0}^{b-1} e^{2\pi i j / b} = 0,$$

the result immediately follows.  $\square$

Using this lemma, we get the following result, which is analogous to that presented in Proposition 2. It is proved by Lemieux and L'Ecuyer (2001) for the case where  $b = 2$  and  $P_n$  is a polynomial lattice point set.

**Proposition 4** *Let  $P_n$  be a digital net in base  $b$ . For any function  $f$  having an absolutely convergent Walsh series expansion, we have*

$$Q_n - \mu = \sum_{\mathbf{0} \neq \mathbf{h}: \zeta(\mathbf{h}) \in \mathcal{C}_s^*} \tilde{f}(\mathbf{h}).$$

Let  $\tilde{P}_n$  be a  $b$ -ary shifted digital net in base  $b$ , and  $\hat{\mu}_{\text{bsh}}$  be the associated estimator. For any square-integrable function  $f$ , we have

$$\text{Var}(\hat{\mu}_{\text{bsh}}) = \sum_{\mathbf{0} \neq \mathbf{h}: \zeta(\mathbf{h}) \in \mathcal{C}_s^*} |\tilde{f}(\mathbf{h})|^2, \quad (1.30)$$

and the variance of the MC estimator  $\hat{\mu}_{\text{MC}}$  based on  $n$  points is given by

$$\text{Var}(\hat{\mu}_{\text{MC}}) = \frac{\sigma^2}{n} = \frac{1}{n} \sum_{\mathbf{0} \neq \mathbf{h} \in \mathbb{N}_0^s} |\tilde{f}(\mathbf{h})|^2.$$

*Proof:* Assume  $\sum_{\mathbf{h} \in \mathbb{N}_0^s} |\tilde{f}(\mathbf{h})| < \infty$ . Then we can write

$$\begin{aligned}
Q_n - \mu &= \frac{1}{n} \sum_{i=0}^{n-1} f(\mathbf{u}_i) - \tilde{f}(\mathbf{0}) \\
&= \frac{1}{n} \sum_{i=0}^{n-1} \sum_{\mathbf{h} \in \mathbb{N}_0^s} \tilde{f}(\mathbf{h}) e^{2\pi i \mathbf{h} \cdot \mathbf{u}_i / b} - \tilde{f}(\mathbf{0}) \\
&= \frac{1}{n} \sum_{\mathbf{h} \in \mathbb{N}_0^s} \tilde{f}(\mathbf{h}) \sum_{i=0}^{n-1} e^{2\pi i \mathbf{h} \cdot \mathbf{u}_i / b} - \tilde{f}(\mathbf{0}) \\
&= \sum_{\mathbf{h}: \zeta(\mathbf{h}) \in \mathcal{C}_s^*} \tilde{f}(\mathbf{h}) - \tilde{f}(\mathbf{0}) = \sum_{\mathbf{0} \neq \mathbf{h}: \zeta(\mathbf{h}) \in \mathcal{C}_s^*} \tilde{f}(\mathbf{h}).
\end{aligned}$$

If  $f$  is square-integrable, then the function  $g : [0, 1)^s \rightarrow \mathbb{R}$  defined by

$$g(\mathbf{u}) = \frac{1}{n} \sum_{i=0}^{n-1} f(\mathbf{u}_i \oplus \mathbf{u}),$$

is also square-integrable, where  $\oplus$  corresponds to a digit-by-digit addition in  $\mathbb{F}_b$ . In addition,  $\text{Var}(\hat{\mu}_{\text{bsh}}) = \text{Var}(g(\mathbf{u})) = \sum_{\mathbf{0} \neq \mathbf{h} \in \mathbb{N}_0^s} |\tilde{g}(\mathbf{h})|^2$ , because Parseval's equality holds for the Walsh series expansion (see Golubov, Efimov, and Skvortsov 1991, for example). Now, for any  $\tilde{\mathbf{h}} \in \mathbb{N}_0^s$ , we have

$$\begin{aligned}
\tilde{g}(\tilde{\mathbf{h}}) &= \int_{[0,1)^s} g(\mathbf{u}) e^{-2\pi i \tilde{\mathbf{h}} \cdot \mathbf{u} / b} d\mathbf{u} \\
&= \int_{[0,1)^s} \frac{1}{n} \sum_{i=0}^{n-1} f(\mathbf{u}_i \oplus \mathbf{u}) e^{-2\pi i \tilde{\mathbf{h}} \cdot \mathbf{u} / b} d\mathbf{u} \\
&= \frac{1}{n} \sum_{i=0}^{n-1} \int_{[0,1)^s} f(\mathbf{v}_i) e^{-2\pi i \tilde{\mathbf{h}} \cdot (\mathbf{v}_i \ominus \mathbf{u}_i) / b} d\mathbf{v}_i \\
&= \frac{1}{n} \sum_{i=0}^{n-1} e^{2\pi i \tilde{\mathbf{h}} \cdot \mathbf{u}_i / b} \int_{[0,1)^s} f(\mathbf{v}_i) e^{-2\pi i \tilde{\mathbf{h}} \cdot \mathbf{v}_i / b} d\mathbf{v}_i \\
&= \frac{1}{n} \sum_{i=0}^{n-1} e^{2\pi i \tilde{\mathbf{h}} \cdot \mathbf{u}_i / b} \tilde{f}(\tilde{\mathbf{h}}) \\
&= \begin{cases} \tilde{f}(\tilde{\mathbf{h}}) & \text{if } \zeta(\tilde{\mathbf{h}}) \in \mathcal{C}_s^*, \\ 0 & \text{otherwise.} \end{cases} \tag{1.31}
\end{aligned}$$

In the above display, the third line is obtained by letting  $\mathbf{v}_i = \mathbf{u}_i \oplus \mathbf{u}$ , and thus  $\mathbf{u} = \mathbf{v}_i \ominus \mathbf{u}_i$ , where  $\ominus$  denotes a digit-by-digit subtraction in

$\mathbb{F}_b$ . From (1.31), the result follows. The variance of the MC estimator is obtained by applying Parseval's equality.  $\square$

To see the connection with scrambled-type estimators, we use the following result (proved for  $b = 2$  by Lemieux 2000), whose proof is given in the appendix. The result also makes a connection between variance expressions in terms of Haar and Walsh expansions, because  $\sigma_{\kappa_I}^2$  is defined in terms of Haar coefficients.

**Lemma 3** *Let  $\kappa_I$  be a vector of positive integers and  $\sigma_{\kappa_I}^2$  be defined as in (1.26). If  $f$  is square-integrable, then*

$$\sigma_{\kappa_I}^2 = \sum_{\mathbf{h} \in G(\kappa_I)} |\tilde{f}(\mathbf{h})|^2,$$

where  $G(\kappa_I) = \{\mathbf{h} \in \mathbb{N}_0^s : 2^{\kappa_j-1} \leq h_j < 2^{\kappa_j} \text{ if } j \in I, h_j = 0 \text{ otherwise}\}.$

Hence in the digital shift case, in comparison with MC, the contribution of a basis function  $\phi_{\mathbf{h}}(\cdot)$  to the variance expression is either multiplied by  $n$  (if  $\mathbf{h}$  is in the dual space) or by 0, whereas in the scrambled case, this contribution is multiplied by 0 for “small vectors”, and by a factor that can be upper-bounded by a quantity *independent* of  $n$  otherwise. This factor being sometimes  $n$  in the digital shift case prevents us from bounding  $\text{Var}(\hat{\mu}_{\text{bsh}})$  by a constant times  $\sigma^2$ . Similarly, the case of smooth functions yields a variance bound in  $O(n^{-2}(\log n)^s)$  for digitally-shifted estimators, which is larger by a factor  $n$  than the order of the bound obtained for scrambled-type estimators. On the other hand, the  $b$ -ary shift is a very simple randomization easy to implement; the estimator  $\hat{\mu}_{\text{bsh}}$  can typically be constructed in the same (or less) time as the MC estimator based on the same number of points.

Based on the expression (1.30) for the variance of the digitally shifted estimator, the same type of heuristic arguments as those given for randomly shifted lattice rules can be used to justify selection criteria such as  $\Delta_{t_1, \dots, t_d}$  to choose digital nets. That is, if we assume that the largest Walsh coefficients are those associated with “small” vectors  $\mathbf{h}$ , then it is reasonable to choose  $P_n$  so that the dual space  $\mathcal{C}_s^*$  does not contain those small vectors. If we use the norm  $\|\mathbf{h}\|$  defined in (1.20) to measure  $\mathbf{h}$ , this suggests using a criterion based on the resolution such as  $\Delta_{t_1, \dots, t_d}$ ; if instead we use the norm  $V(\mathbf{h})$  defined in (1.18), then the  $t$ -value or the variant  $\tilde{t}$  defined in (1.16) should be employed. We refer the reader to Hellekalek and Leeb (1997) and Tezuka (1987) for additional connections between Walsh expansions and nonuniformity measures (e.g., the so-called ‘Walsh-spectral test’ of Tezuka). Note that criteria based on the resolution are faster to compute than those based on the  $t$ -value,

because the latter verifies the  $(q_1, \dots, q_s)$ -equidistribution of  $P_n$  for a much larger number of vectors  $(q_1, \dots, q_s)$ .

## 7. Transformations of the Integrand

So far, our description of how to use QMC methods can be summarized as follows: Choose a construction and a randomization; choose a selection criterion; find a good point set with respect to this criterion (or use a precomputed table of “good” point sets); randomize the point set, and compute  $\hat{\mu}_{\text{RQMC}}$  as an estimator for  $\mu$ . If the selection criterion mimics the variance of  $\hat{\mu}_{\text{RQMC}}$  well enough, one should obtain a low variance estimator with this approach. Most of the selection criteria presented in Section 4 are defined so that they should imitate more or less the variance of  $\hat{\mu}_{\text{RQMC}}$  for a large class of functions, i.e., they provide “general-purpose” low-discrepancy point sets. However, once the problem at hand is known, the variance can sometimes be reduced further by making use of information on  $f$  in a clever way. In particular, techniques used to reduce the MC variance can also be used in combination with QMC methods. Examples of such techniques are antithetic variates, control variables, importance sampling, and conditional Monte Carlo. These methods can all be seen as transformations applied to  $f$  in order to reduce its variability; that is, one replaces  $f$  by a function  $g$  such that  $\int_{[0,1]^s} g(\mathbf{u}) d\mathbf{u} = \mu$  and (hopefully)  $\int_{[0,1]^s} g^2(\mathbf{u}) d\mathbf{u} < \int_{[0,1]^s} f^2(\mathbf{u}) d\mathbf{u}$ . If the function  $g$  requires more computation time for its evaluation, one should make sure that the variance reduction gained is worth the extra effort, i.e., that the *efficiency* is improved.

A second class of methods that can reduce the variance of QMC estimators for certain functions are *dimension reduction* methods. Among this class are the *Brownian bridge technique* of Caflisch and Moskowitz (1995), approaches based on principal components (Acworth, Broadie, and Glasserman 1997; Åkesson and Lehoczy 2000), and various methods discussed by Fox (1999) for generating Poisson and other stochastic processes.

Typically, these methods are used when  $f$  is defined in terms of a stochastic process for which a sample path is generated using the uniform numbers  $u_{i1}, \dots, u_{is}$  provided by a point  $\mathbf{u}_i$  in  $\tilde{P}_n$ . The goal is then to generate the sample path in a way that will decrease the effective dimension of  $f$ . This is usually achieved by using a method that gives a lot of importance to a few uniform numbers. As an illustration, we describe in the example below the Brownian bridge technique, which can be used to generate the sample paths of a Brownian motion.

**Example 3** As this often happens in financial simulations (see e.g., Boyle, Broadie, and Glasserman 1997; Caffisch, Morokoff, and Owen 1997), suppose we want to generate the sample path of a Brownian motion at  $s$  different times  $B(t_1), \dots, B(t_s)$ , using  $s$  uniform numbers  $u_1, \dots, u_s$ . For instance, this Brownian motion might be driving the price process of an asset on which an *option* has been written (Duffie 1996). Instead of generating these observations sequentially (that is, by using  $u_j$  to generate the Gaussian random variable  $B(t_j)/B(t_{j-1})$  given  $B(t_{j-1})$ ),  $u_1$  is used to generate  $B(t_s)$ ,  $u_2$  is used to generate  $B(t_{\lfloor s/2 \rfloor})$ ,  $u_3$  for  $B(t_{\lfloor s/4 \rfloor})$ ,  $u_4$  for  $B(t_{\lfloor 3s/4 \rfloor})$ , etc. This can be done easily since for  $u < v < w$ , the distribution of  $B(v)$  given  $B(u)$  and  $B(w)$  is Gaussian with parameters depending only on  $u, v, w$ . By generating the Brownian motion path this way, more importance is given to the first few uniform numbers since they determine important aspects of the path such as its value at the end, middle, first quarter, etc.

Another type of transformation that can sometimes reduce the actual dimension of the problem (and not only the effective one) is the conditional Monte Carlo method; see L'Ecuyer and Lemieux (2000), Section 10.1, for example, where this method is used with randomly shifted lattice rules. This method can also be viewed as a “smoothing technique”, i.e., it replaces the integrand by a smoother function (in this case, a conditional expectation) that, e.g., satisfies conditions required to obtain variance rates that are  $O(n^{-3} \log^s n)$ , as in Section 6.2. We refer the reader to Spanier and Maize (1994) and Fox (1999), Chapters 7 and 8, for more on smoothing techniques that can be used in combination with QMC methods.

## 8. Related Methods

We now discuss integration methods that are closely related to QMC methods, but that do not exactly fit the framework presented so far.

First, a natural extension for the estimator  $Q_n$  or  $\hat{\mu}_{\text{RQMC}}$  would be to assign weights to the different evaluation points; that is, for a point set  $P_n = \{\mathbf{u}_0, \dots, \mathbf{u}_{n-1}\}$ , define

$$\hat{\mu}_w = \sum_{i=0}^{n-1} w_i f(\mathbf{u}_i),$$

with  $\sum_{i=1}^n w_i = 1$ . Hickernell (1998b) proved that when  $P_n$  is a lattice point set, then a certain measure of discrepancy (called the  $\mathcal{L}^2$ -star discrepancy), defined so that these weights  $w_i$  are accounted for, is minimized by setting the weights to be all equal to  $1/n$ . In other words, using weights is useless in this case.

However, if  $P_n$  is not restricted to have a particular form, then it can be shown that in some cases allowing different weights can bring a significant improvement (Yakowitz, Krimmel, and Szidarovszky 1978). For example, one can use the MC method with weights  $w_i$  defined by the *Voronoi tessellation* induced by the uniform and random points  $\mathbf{u}_0, \dots, \mathbf{u}_{n-1}$ ; more precisely, define

$$w_i = \lambda(\{\mathbf{u} \in [0, 1]^s : \|\mathbf{u} - \mathbf{u}_i\|_2 \leq \|\mathbf{u} - \mathbf{u}_j\|_2 \text{ for all } j \neq i\}),$$

where  $\lambda$  denotes the Lebesgue measure on  $[0, 1]^s$ . This approach yields an estimator  $\hat{\mu}_w$  with variance in  $O(n^{-2})$  when  $s = 2$ . Weighted approximations also based on Voronoi tessellations are discussed by Pagès (1997).

A closely related idea is used in *stratified sampling* (Cochran 1977). In this method, the unit hypercube is partitioned into  $n$  cells  $W_0, \dots, W_{n-1}$ , and  $\mathbf{u}_i$  is uniformly distributed over  $W_i$ , for  $i = 0, \dots, n-1$ . The stratified sampling estimator is then

$$\hat{\mu}_{\text{str}} = \sum_{i=0}^{n-1} w_i f(\mathbf{u}_i),$$

where  $w_i = \lambda(W_i)$ . It can be shown (Cochran 1977) that for any square-integrable function,

$$\text{Var}(\hat{\mu}_{\text{str}}) \leq \text{Var}(\hat{\mu}_{\text{MC}}).$$

The amount of variance reduction depends on the definition of the cells  $W_0, \dots, W_{n-1}$  and their interaction with the integrand  $f$ .

An integration method that is guaranteed to yield an estimator with a variance not larger than the MC estimator for monotone functions is the *Latin Hypercube Sampling* (LHS) (Mckay, Beckman, and Conover 1979). It uses a point set whose unidimensional projections are evenly distributed (i.e., one point per interval  $[j/n, (j+1)/n)$ , for  $j = 0, \dots, n-1$ ). To construct this point set, one needs to generate  $s$  random, uniform, and independent permutations  $\pi_1, \dots, \pi_s$  of the integers from 0 to  $n-1$ , and  $n$  independent shifts  $\Delta_0, \dots, \Delta_{n-1}$  uniformly distributed over  $[0, 1/n)^s$ . Then define the point set

$$\tilde{P}_n = \left\{ \left( \frac{\pi_1(i)}{n} + \Delta_{i,1}, \dots, \frac{\pi_s(i)}{n} + \Delta_{i,s} \right), i = 0, \dots, n-1 \right\}.$$

Additional results on this method can be found in, e.g., Avramidis and Wilson (1996), Owen (1998a), and the references therein. In particular, see Owen (1992a) and Loh (1996b) for results showing that the LHS estimator obeys a central-limit theorem. A related method that extends

the uniformity property of LHS and has close connections with digital nets is a *randomized orthogonal array design*; we refer the reader to Owen (1992b, 1994, 1995, 1997) and Loh (1996a) for more on this method.

## 9. Conclusions and Discussion

We have described various QMC constructions that can be used for multidimensional numerical integration. Measures of quality that can help selecting parameters for a given construction have been presented. We also discussed different randomizations, and provided results on the variance of estimators obtained by applying these randomizations. In particular, we gave a new result that expresses the variance of an estimator based on a digitally shifted net  $P_n$  as a sum of squared Walsh coefficients over the dual space of  $P_n$ .

In the near future, we plan to compare empirically various constructions and randomizations on practical problems, to study selection criteria and compare their effectiveness, and to investigate in more details the effect of transformations, such as those discussed in Section 7, on the variance of the randomized QMC estimators.

## Appendix: Proofs

*Proof of Proposition 1:* The result is obtained by first generalizing Proposition 5.2 of Lemieux and L'Ecuyer (2001) to arbitrary digital nets in base  $b$ . This can be done by using Lemma 2 from Section 6.2. More precisely, we show that  $P_n$  is  $(q_1, \dots, q_s)$ -equidistributed if and only if  $\mathcal{C}_s^* \cap \mathcal{H}(q_1, \dots, q_s) = \{\mathbf{0}\}$ , where

$$\mathcal{H}(q_1, \dots, q_s) = \{\mathbf{h} \in (\mathbb{F}_b[z])^s \text{ such that } v(h_j(z)) \leq q_j \text{ for each } j\}.$$

Consider the class  $\mathcal{F}(q_1, \dots, q_s)$  of all real-valued functions that are constant on each of the  $b^q$   $b$ -ary boxes in the definition of  $(q_1, \dots, q_s)$ -equidistribution. Clearly,  $P_n$  is  $(q_1, \dots, q_s)$ -equidistributed if and only if the corresponding point set integrates every function  $f \in \mathcal{F}$  with zero error. But due to its periodic structure, each function  $f \in \mathcal{F}(q_1, \dots, q_s)$  has a Walsh expansion of the form

$$f(\mathbf{u}) = \sum_{\mathbf{h} \in H(q_1, \dots, q_s)} \tilde{f}(\mathbf{h}) e^{2\pi i \mathbf{h} \cdot \mathbf{u} / b}, \quad (1.A.1)$$

where  $H(q_1, \dots, q_s) = \{\mathbf{h} \in \mathbb{N}_0^s : \zeta(\mathbf{h}) \in \mathcal{H}(q_1, \dots, q_s)\}$ , and  $\zeta$  is the bijection between  $\mathbb{N}_0^s$  and  $(\mathbb{F}_b[z])^s$  mentioned on page 39. To see this, note that any  $f \in \mathcal{F}$  can be written as

$$f(\mathbf{u}) = \sum_{v_1=0}^{b^{q_1}-1} \dots \sum_{v_s=0}^{b^{q_s}-1} c_{v_1, \dots, v_s} \prod_{j=1}^s \mathbf{1}_{v_j b^{-q_j} \leq u_j < (v_j+1)b^{-q_j}}, \quad (1.A.2)$$

where the  $c_{v_1, \dots, v_s}$  are real numbers. When  $\mathbf{h} \notin H(q_1, \dots, q_s)$ , there exists  $j \in \{1, \dots, s\}$  and an integer  $w_j \geq q_j$  such that  $h_j > 2^{w_j}$ . Let  $d = w_j + 1 - q_j$ , and recall that  $h_j \cdot u = \sum_{l=1}^{w_j+1} h_{j,l} u_l$ , where the coefficients  $h_{j,l}$  and  $u_l$  come from the

representation  $h_j = \sum_{l=1}^{w_j+1} h_{j,l} b^{l-1}$  and  $u = \sum_{l=1}^{\infty} u_l b^{-l}$ , respectively. When  $u_j$  goes from  $v_j b^{-q_j}$  to  $(v_j + 1) b^{-q_j}$ ,  $l$  goes from 0 to  $b^d - 1$  in  $h_j \cdot (v_j b^{-q_j} + l b^{-w_j-1})$ , and this dot product is then equal to each number between 0 and  $b - 1$  exactly  $b^{d-1}$  times. Hence, if we first integrate  $f(\mathbf{u})$  with respect to  $u_j$  when computing  $\tilde{f}(\mathbf{h})$  via (1.29), any term from the sum (1.A.2) will be 0 because

$$\begin{aligned} & \int_0^1 c_{v_1, \dots, v_s} \mathbf{1}_{v_j b^{-q_j} \leq u_j < (v_j + 1) b^{-q_j}} e^{-2\pi i h_j \cdot u_j / b} du_j \\ &= c_{v_1, \dots, v_s} \sum_{l=0}^{b^d-1} e^{-2\pi i h_j \cdot (v_j b^{-q_j} + l 2^{-w_j-1}) / b} = 0. \end{aligned}$$

Therefore,  $\tilde{f}(\mathbf{h}) = 0$  if  $\mathbf{h} \notin H(q_1, \dots, q_s)$ , and (1.A.1) follows. Now, for any nonzero  $\tilde{\mathbf{h}} \in H(q_1, \dots, q_s)$ ,  $g(\mathbf{u}) \equiv e^{-2\pi i \tilde{\mathbf{h}} \cdot \mathbf{u} / b}$  is in  $\mathcal{F}(q_1, \dots, q_s)$ . Hence by Proposition 4, the error obtained by using  $P_n$  to integrate  $g$  is  $\sum_{\mathbf{0} \neq \mathbf{h}: \zeta(\mathbf{h}) \in \mathcal{C}_s^*} \tilde{g}(\mathbf{h}) = \mathbf{1}_{\zeta(\tilde{\mathbf{h}}) \in \mathcal{C}_s^*}$  since the only nonzero Walsh coefficient of  $g$  is the one evaluated in  $\tilde{\mathbf{h}}$  (and it is equal to 1). From this, we see that if  $P_n$  is  $(q_1, \dots, q_s)$ -equidistributed, then  $\mathcal{C}_s^* \cap \mathcal{H}(q_1, \dots, q_s) = \{\mathbf{0}\}$ . Hence if  $P_n$  has a resolution of  $\ell_s$ , then it is  $(\ell_s, \dots, \ell_s)$ -equidistributed and therefore  $\min_{\mathbf{0} \neq \mathbf{h} \in \mathcal{C}_s^*} \|\mathbf{h}\| > \ell_s$ .

We now show that  $\min_{\mathbf{0} \neq \mathbf{h} \in \mathcal{C}_s^*} \|\mathbf{h}\| \leq \ell_s + 1$ , which will prove the result. Since the resolution is  $\ell_s$ , it means that  $P_n$  is not  $(\ell_s + 1, \dots, \ell_s + 1)$ -equidistributed. Therefore, the  $k \times s(\ell_s + 1)$  matrix  $\mathbf{L}$  formed by concatenating the transposed of the first  $\ell_s + 1$  rows of each generating matrix  $\mathbf{C}^j$  has a row space whose dimension is strictly smaller than  $s(\ell_s + 1)$ . Hence there exists a nonzero vector  $\mathbf{x}$  in  $\mathbb{F}_b^{s(\ell_s + 1)}$  such that  $\mathbf{L}\mathbf{x} = \mathbf{0}$ . Furthermore, we can assume  $\|\mathbf{x}\| = \ell_s + 1$  since  $\|\mathbf{x}\| \leq \ell_s$  would contradict our assumption that  $P_n$  has a resolution of  $\ell_s$ . Define  $\mathbf{h} = (h_1(z), \dots, h_s(z)) \in (\mathbb{F}_b[z])^s$  by

$$h_j(z) = \sum_{l=1}^{\ell_s+1} x_{j,l} z^{l-1}.$$

Since  $\mathbf{L}$  is just a truncated version of  $\mathbf{C}$  and the coefficients of  $h_j(z)$  for powers of  $z$  larger than  $\ell_s + 1$  are zero for all  $j$ , we have that  $\mathbf{C}\mathbf{h} = \mathbf{0}$ , and therefore  $\mathbf{h} \in \mathcal{C}_s^*$  with  $\|\mathbf{h}\| = \|\mathbf{x}\| = \ell_s + 1$ , which proves the result.  $\square$

*Proof of Lemma 3:* Recall that

$$\sigma_{\kappa_I}^2 = \text{Var}(\nu_{\kappa_I}(\mathbf{u})),$$

where  $\nu_{\kappa_I}(\mathbf{u})$  is a step function constant on the  $b$ -ary boxes obtained by partitioning the  $j$ th axis of  $[0, 1]^s$  in  $b^{\kappa_j}$  equal intervals, for each  $j \in I$ . Using the same notation as for the preceding proof, we have that  $\nu_{\kappa_I} \in \mathcal{F}(q_1, \dots, q_s)$ , where

$$q_j = \begin{cases} \kappa_j & \text{if } j \in I, \\ 0 & \text{otherwise.} \end{cases}$$

Hence from the proof of Proposition 1, we know that  $\tilde{\nu}_{\kappa_I}(\mathbf{h}) = 0$  if  $\mathbf{h} \notin H(q_1, \dots, q_s)$ .

Assume  $\mathbf{h} \in H(q_1, \dots, q_s)$  and that there exists one  $j \in I$  such that  $h_j < 2^{\kappa_j-1}$ . We need to verify that  $\tilde{\nu}_{\kappa_I}(\mathbf{h}) = 0$ . Now,

$$\tilde{\nu}_{\kappa_I}(\mathbf{h}) = \int_{[0,1]^s} \sum_{\tau} \sum_{\gamma} \langle \psi_{\kappa_I, \tau, \gamma}, f \rangle \psi_{\kappa_I, \tau, \gamma}(\mathbf{u}) e^{-2\pi i \mathbf{h} \cdot \mathbf{u} / b} d\mathbf{u}$$

$$\begin{aligned}
&= \sum_{\tau} \sum_{\gamma} \langle \psi_{\kappa_I, \tau, \gamma}, f \rangle \int_{[0,1]^s} \psi_{\kappa_I, \tau, \gamma}(\mathbf{u}) e^{-2\pi i \mathbf{h} \cdot \mathbf{u}/b} d\mathbf{u} \\
&= \sum_{\tau} \sum_{\gamma} \langle \psi_{\kappa_I, \tau, \gamma}, f \rangle b^{(|\kappa_I| - |I|)/2} \cdot \\
&\quad \left( \prod_{k \in I, k \neq j} \int_0^1 (b \mathbf{1}_{\lfloor b^{\kappa_k} + 1 \rfloor u_j = b\tau_k + \gamma_k} - \mathbf{1}_{\lfloor b^{\kappa_k} u_k \rfloor = \tau_k}) e^{-2\pi i h_k \cdot u_k / b} du_k \right) \\
&\quad \int_0^1 (b \mathbf{1}_{\lfloor b^{\kappa_j} + 1 \rfloor u_j = b\tau_j + \gamma_j} - \mathbf{1}_{\lfloor b^{\kappa_j} u_j \rfloor = \tau_j}) e^{-2\pi i h_j \cdot u_j / b} du_j. \tag{1.A.3}
\end{aligned}$$

Now, since  $h_j < 2^{\kappa_j - 1}$ , the function  $e^{-2\pi i h_j \cdot u_j / b}$  is constant over any interval of the form  $[db^{-\kappa_j + 1}, (d+1)b^{-\kappa_j + 1})$ ,  $0 \leq d < b^{\kappa_j - 1}$ . Hence, if

$$\int_{db^{-\kappa_j + 1}}^{(d+1)b^{-\kappa_j + 1}} (b \mathbf{1}_{\lfloor b^{\kappa_j} + 1 \rfloor u_j = b\tau_j + \gamma_j} - \mathbf{1}_{\lfloor b^{\kappa_j} u_j \rfloor = \tau_j}) du_j = 0, \tag{1.A.4}$$

for any  $d \in [0, \dots, b^{\kappa_j - 1} - 1]$ , then (1.A.3) is equal to zero and the result is proved. To show (1.A.4), it suffices to observe that

$$\int_{db^{-\kappa_j + 1}}^{(d+1)b^{-\kappa_j + 1}} (b \mathbf{1}_{\lfloor b^{\kappa_j} + 1 \rfloor u_j = b\tau_j + \gamma_j} - \mathbf{1}_{\lfloor b^{\kappa_j} u_j \rfloor = \tau_j}) du_j = b \cdot b^{-\kappa_j - 1} - b^{-\kappa_j} = 0,$$

for any  $d$ , which proves the result.  $\square$

## Acknowledgments

This work was supported by NSERC-Canada individual grants to the two authors and by an FCAR-Québec grant to the first author. We thank Bennett L. Fox, Fred J. Hickernell, Harald Niederreiter, and Art B. Owen for helpful comments and suggestions.

## References

- Acworth, P., M. Broadie, and P. Glasserman. 1997. A comparison of some Monte Carlo and quasi-Monte Carlo techniques for option pricing. In *Monte Carlo and Quasi-Monte Carlo Methods in Scientific Computing*, ed. P. Hellekalek and H. Niederreiter, Number 127 in Lecture Notes in Statistics, 1–18. Springer-Verlag.
- Åkesson, F., and J. P. Lehoczy. 2000. Path generation for quasi-Monte Carlo simulation of mortgage-backed securities. *Management Science* 46:1171–1187.
- Antonov, I. A., and V. M. Saleev. 1979. An economic method of computing  $LP_{\tau}$ -sequences. *Zh. Vychisl. Mat. i. Mat. Fiz.* 19:243–245. In Russian.
- Avramidis, A. N., and J. R. Wilson. 1996. Integrated variance reduction strategies for simulation. *Operations Research* 44:327–346.

- Bakhvalov, N. S. 1959. On approximate calculation of multiple integrals. *Vestnik Moskovskogo Universiteta, Seriya Matematiki, Mekhaniki, Astronomi, Fiziki, Himii* 4:3–18. In Russian.
- Boyle, P., M. Broadie, and P. Glasserman. 1997. Monte Carlo methods for security pricing. *Journal of Economic Dynamics & Control* 21 (8-9): 1267–1321. Computational financial modelling.
- Braaten, E., and G. Weller. 1979. An improved low-discrepancy sequence for multidimensional quasi-Monte Carlo integration. *Journal of Computational Physics* 33:249–258.
- Bratley, P., and B. L. Fox. 1988. Algorithm 659: Implementing Sobol's quasirandom sequence generator. *ACM Transactions on Mathematical Software* 14 (1): 88–100.
- Bratley, P., B. L. Fox, and H. Niederreiter. 1992. Implementation and tests of low-discrepancy sequences. *ACM Transactions on Modeling and Computer Simulation* 2:195–213.
- Bratley, P., B. L. Fox, and H. Niederreiter. 1994. Algorithm 738: Programs to generate Niederreiter's low-discrepancy sequences. *ACM Transactions on Mathematical Software* 20:494–495.
- Caffisch, R. E., W. Morokoff, and A. Owen. 1997. Valuation of mortgage-backed securities using Brownian bridges to reduce effective dimension. *The Journal of Computational Finance* 1 (1): 27–46.
- Caffisch, R. E., and B. Moskowitz. 1995. Modified Monte Carlo methods using quasi-random sequences. In *Monte Carlo and Quasi-Monte Carlo Methods in Scientific Computing*, ed. H. Niederreiter and P. J.-S. Shiue, Number 106 in Lecture Notes in Statistics, 1–16. New York: Springer-Verlag.
- Cheng, J., and M. J. Druzdzel. 2000. Computational investigation of low-discrepancy sequences in simulation algorithms for bayesian networks. In *Uncertainty in Artificial Intelligence Proceedings 2000*, 72–81.
- Cochran, W. G. 1977. *Sampling techniques*. Second ed. New York: John Wiley and Sons.
- Conway, J. H., and N. J. A. Sloane. 1999. *Sphere packings, lattices and groups*. 3rd ed. Grundlehren der Mathematischen Wissenschaften 290. New York: Springer-Verlag.
- Couture, R., and P. L'Ecuyer. 2000. Lattice computations for random numbers. *Mathematics of Computation* 69 (230): 757–765.
- Couture, R., P. L'Ecuyer, and S. Tezuka. 1993. On the distribution of  $k$ -dimensional vectors for simple and combined Tausworthe sequences. *Mathematics of Computation* 60 (202): 749–761, S11–S16.
- Coveyou, R. R., and R. D. MacPherson. 1967. Fourier analysis of uniform random number generators. *Journal of the ACM* 14:100–119.

- Cranley, R., and T. N. L. Patterson. 1976. Randomization of number theoretic methods for multiple integration. *SIAM Journal on Numerical Analysis* 13 (6): 904–914.
- Davis, P., and P. Rabinowitz. 1984. *Methods of numerical integration*. second ed. New York: Academic Press.
- Dieter, U. 1975. How to calculate shortest vectors in a lattice. *Mathematics of Computation* 29 (131): 827–833.
- Duffie, D. 1996. *Dynamic asset pricing theory*. second ed. Princeton University Press.
- Efron, B., and C. Stein. 1981. The jackknife estimator of variance. *Annals of Statistics* 9:586–596.
- Entacher, K. 1997. Quasi-Monte Carlo methods for numerical integration of multivariate Haar series. *BIT* 37:846–861.
- Entacher, K., P. Hellekalek, and P. L’Ecuyer. 2000. Quasi-Monte Carlo node sets from linear congruential generators. In *Monte Carlo and Quasi-Monte Carlo Methods 1998*, ed. H. Niederreiter and J. Spanier, 188–198. Berlin: Springer.
- Faure, H. 1982. Discrépance des suites associées à un système de numération. *Acta Arithmetica* 61:337–351.
- Faure, H. 2001. Variations on  $(0, s)$ -sequences. *Journal of Complexity*. To appear.
- Faure, H., and S. Tezuka. 2001. A new generation of  $(0, s)$ -sequences. To appear.
- Fincke, U., and M. Pohst. 1985. Improved methods for calculating vectors of short length in a lattice, including a complexity analysis. *Mathematics of Computation* 44:463–471.
- Fishman, G. S. 1990, Jan. Multiplicative congruential random number generators with modulus  $2^\beta$ : An exhaustive analysis for  $\beta = 32$  and a partial analysis for  $\beta = 48$ . *Mathematics of Computation* 54 (189): 331–344.
- Fishman, G. S., and L. S. Moore III. 1986. An exhaustive analysis of multiplicative congruential random number generators with modulus  $2^{31} - 1$ . *SIAM Journal on Scientific and Statistical Computing* 7 (1): 24–45.
- Fox, B. L. 1986. Implementation and relative efficiency of quasirandom sequence generators. *ACM Transactions on Mathematical Software* 12:362–376.
- Fox, B. L. 1999. *Strategies for quasi-Monte Carlo*. Boston, MA: Kluwer Academic.
- Friedel, I., and A. Keller. 2001. Fast generation of randomized low-discrepancy point sets. In *Monte Carlo and Quasi-Monte Carlo Methods 2000*, ed. K.-T. Fang, F. J. Hickernell, and H. Niederreiter: Springer. To appear.

- Golubov, B., A. Efimov, and V. Skvortsov. 1991. *Walsh series and transforms: Theory and applications*, Volume 64 of *Mathematics and Applications: Soviet Series*. Boston: Kluwer Academic Publishers.
- Haber, S. 1983. Parameters for integrating periodic functions of several variables. *Mathematics of Computation* 41:115–129.
- Halton, J. H. 1960. On the efficiency of certain quasi-random sequences of points in evaluating multi-dimensional integrals. *Numerische Mathematik* 2:84–90.
- Heinrich, S., F. J. Hickernell, and R.-X. Yue. 2001a. Integration of multivariate Haar wavelet series. Submitted.
- Heinrich, S., F. J. Hickernell, and R. X. Yue. 2001b. Optimal quadrature for Haar wavelet spaces. submitted.
- Hellekalek, P. 1998. On the assessment of random and quasi-random point sets. In *Random and Quasi-Random Point Sets*, ed. P. Hellekalek and G. Larcher, Volume 138 of *Lecture Notes in Statistics*, 49–108. New York: Springer.
- Hellekalek, P., and G. Larcher. (Eds.) 1998. *Random and quasi-random point sets*, Volume 138 of *Lecture Notes in Statistics*. New York: Springer.
- Hellekalek, P., and H. Leeb. 1997. Dyadic diaphony. *Acta Arithmetica* 80:187–196.
- Hickernell, F. J. 1998a. A generalized discrepancy and quadrature error bound. *Mathematics of Computation* 67:299–322.
- Hickernell, F. J. 1998b. Lattice rules: How well do they measure up? In *Random and Quasi-Random Point Sets*, ed. P. Hellekalek and G. Larcher, Volume 138 of *Lecture Notes in Statistics*, 109–166. New York: Springer.
- Hickernell, F. J. 1999. Goodness-of-fit statistics, discrepancies and robust designs. *Statistical and Probability Letters* 44:73–78.
- Hickernell, F. J. 2000. What affects accuracy of quasi-Monte Carlo quadrature? In *Monte Carlo and Quasi-Monte Carlo Methods 1998*, ed. H. Niederreiter and J. Spanier, 16–55. Berlin: Springer.
- Hickernell, F. J., and H. S. Hong. 1997. Computing multivariate normal probabilities using rank-1 lattice sequences. In *Proceedings of the Workshop on Scientific Computing (Hong Kong)*, ed. G. H. Golub, S. H. Lui, F. T. Luk, and R. J. Plemmons, 209–215. Singapore: Springer-Verlag.
- Hickernell, F. J., H. S. Hong, P. L’Ecuyer, and C. Lemieux. 2001. Extensible lattice sequences for quasi-Monte Carlo quadrature. *SIAM Journal on Scientific Computing* 22 (3): 1117–1138.
- Hickernell, F. J., and H. Woźniakowski. 2001. The price of pessimism for multidimensional quadrature. *Journal of Complexity* 17. To appear.

- Hlawka, E. 1961. Funktionen von beschränkter variation in der theorie der gleichverteilung. *Ann. Mat. Pura. Appl.* 54:325–333.
- Hlawka, E. 1962. Zur angenäherten berechnung mehrfacher integrale. *Monatshefte für Mathematik* 66:140–151.
- Hoeffding, W. 1948. A class of statistics with asymptotically normal distributions. *Annals of Mathematical Statistics* 19:293–325.
- Hong, H. S., and F. H. Hickernell. 2001. Implementing scrambled digital sequences. Submitted for publication.
- Knuth, D. E. 1998. *The art of computer programming, volume 2: Seminumerical algorithms*. Third ed. Reading, Mass.: Addison-Wesley.
- Korobov, N. M. 1959. The approximate computation of multiple integrals. *Dokl. Akad. Nauk SSSR* 124:1207–1210. in Russian.
- Korobov, N. M. 1960. Properties and calculation of optimal coefficients. *Dokl. Akad. Nauk SSSR* 132:1009–1012. in Russian.
- Larcher, G. 1998. Digital point sets: Analysis and applications. In *Random and Quasi-Random Point Sets*, ed. P. Hellekalek and G. Larcher, Volume 138 of *Lecture Notes in Statistics*, 167–222. New York: Springer.
- Larcher, G., A. Lauss, H. Niederreiter, and W. C. Schmid. 1996. Optimal polynomials for  $(t, m, s)$ -nets and numerical integration of multivariate Walsh series. *SIAM Journal on Numerical Analysis* 33 (6): 2239–2253.
- Larcher, G., H. Niederreiter, and W. C. Schmid. 1996. Digital nets and sequences constructed over finite rings and their application to quasi-Monte Carlo integration. *Monatshefte für Mathematik* 121 (3): 231–253.
- Larcher, G., and G. Pirsic. 1999. Base change problems for generalized Walsh series and multivariate numerical integration. *Pacific Journal of Mathematics* 189:75–105.
- L’Ecuyer, P. 1994. Uniform random number generation. *Annals of Operations Research* 53:77–120.
- L’Ecuyer, P. 1996. Maximally equidistributed combined Tausworthe generators. *Mathematics of Computation* 65 (213): 203–213.
- L’Ecuyer, P. 1999. Tables of linear congruential generators of different sizes and good lattice structure. *Mathematics of Computation* 68 (225): 249–260.
- L’Ecuyer, P., and R. Couture. 1997. An implementation of the lattice and spectral tests for multiple recursive linear random number generators. *INFORMS Journal on Computing* 9 (2): 206–217.
- L’Ecuyer, P., and C. Lemieux. 2000. Variance reduction via lattice rules. *Management Science* 46 (9): 1214–1235.

- L'Ecuyer, P., and F. Panneton. 2000. A new class of linear feedback shift register generators. In *Proceedings of the 2000 Winter Simulation Conference*, ed. J. A. Joines, R. R. Barton, K. Kang, and P. A. Fishwick, 690–696. Piscataway, NJ: IEEE Press.
- Lemieux, C. 2000, May. *L'utilisation de règles de réseau en simulation comme technique de réduction de la variance*. Ph. D. thesis, Université de Montréal.
- Lemieux, C., M. Cieslak, and K. Luttmer. 2001. RandQMC user's guide. In preparation.
- Lemieux, C., and P. L'Ecuyer. 2001. Selection criteria for lattice rules and other low-discrepancy point sets. *Mathematics and Computers in Simulation* 55 (1–3): 139–148.
- Lemieux, C., and A. B. Owen. 2001. Quasi-regression and the relative importance of the ANOVA components of a function. In *Monte Carlo and Quasi-Monte Carlo Methods 2000*, ed. K.-T. Fang, F. J. Hickernell, and H. Niederreiter: Springer. To appear.
- Lidl, R., and H. Niederreiter. 1994. *Introduction to finite fields and their applications*. Revised ed. Cambridge: Cambridge University Press.
- Loh, W.-L. 1996a. A combinatorial central limit theorem for randomized orthogonal array sampling designs. *Annals of Statistics* 24:1209–1224.
- Loh, W.-L. 1996b. On Latin hypercube sampling. *The Annals of Statistics* 24:2058–2080.
- Maisonneuve, D. 1972. Recherche et utilisation des “bons treillis”, programmation et résultats numériques. In *Applications of Number Theory to Numerical Analysis*, ed. S. K. Zaremba, 121–201. New York: Academic Press.
- Maize, E. 1981. *Contributions to the theory of error reduction in quasi-Monte Carlo methods*. Ph. D. thesis, Claremont Graduate School, Claremont, CA.
- Matoušek, J. 1998. On the L2-discrepancy for anchored boxes. *Journal of Complexity* 14:527–556.
- Matsumoto, M., and Y. Kurita. 1994. Twisted GFSR generators II. *ACM Transactions on Modeling and Computer Simulation* 4 (3): 254–266.
- Matsumoto, M., and T. Nishimura. 1998. Mersenne twister: A 623-dimensionally equidistributed uniform pseudo-random number generator. *ACM Transactions on Modeling and Computer Simulation* 8 (1): 3–30.
- Mckay, M. D., R. J. Beckman, and W. J. Conover. 1979. A comparison of three methods for selecting values of input variables in the analysis of output from a computer code. *Technometrics* 21:239–245.

- Morohosi, H., and M. Fushimi. 2000. A practical approach to the error estimation of quasi-Monte Carlo integration. In *Monte Carlo and Quasi-Monte Carlo Methods 1998*, ed. H. Niederreiter and J. Spanier, 377–390. Berlin: Springer.
- Morokoff, W. J., and R. E. Caflisch. 1994. Quasi-random sequences and their discrepancies. *SIAM Journal on Scientific Computing* 15:1251–1279.
- Niederreiter, H. 1986. Multidimensional numerical integration using pseudorandom numbers. *Mathematical Programming Study* 27:17–38.
- Niederreiter, H. 1987. Point sets and sequences with small discrepancy. *Monatshefte für Mathematik* 104:273–337.
- Niederreiter, H. 1988. Low-discrepancy and low-dispersion sequences. *Journal of Number Theory* 30:51–70.
- Niederreiter, H. 1992a. Low-discrepancy point sets obtained by digital constructions over finite fields. *Czechoslovak Math. Journal* 42:143–166.
- Niederreiter, H. 1992b. *Random number generation and quasi-Monte Carlo methods*, Volume 63 of *SIAM CBMS-NSF Regional Conference Series in Applied Mathematics*. Philadelphia: SIAM.
- Niederreiter, H., and G. Pirsic. 2001. Duality for digital nets and its applications. *Acta Arithmetica* 97:173–182.
- Niederreiter, H., and C. Xing. 1997. The algebraic-geometry approach to low-discrepancy sequences. In *Monte Carlo and Quasi-Monte Carlo Methods in Scientific Computing*, ed. P. Hellekalek, G. Larcher, H. Niederreiter, and P. Zinterhof, Volume 127 of *Lecture Notes in Statistics*, 139–160. New York: Springer-Verlag.
- Niederreiter, H., and C. Xing. 1998. Nets,  $(t, s)$ -sequences, and algebraic geometry. In *Random and Quasi-Random Point Sets*, ed. P. Hellekalek and G. Larcher, Volume 138 of *Lecture Notes in Statistics*, 267–302. New York: Springer.
- Ökten, G. 1996. A probabilistic result on the discrepancy of a hybrid-Monte Carlo sequence and applications. *Monte Carlo methods and Applications* 2:255–270.
- Owen, A. B. 1992a. A central limit theorem for Latin hypercube sampling. *Journal of the Royal Statistical Society B* 54 (2): 541–551.
- Owen, A. B. 1992b. Orthogonal arrays for computer experiments, integration and visualization. *Statistica Sinica* 2:439–452.
- Owen, A. B. 1994. Lattice sampling revisited: Monte Carlo variance of means over randomized orthogonal arrays. *Annals of Statistics* 22:930–945.
- Owen, A. B. 1995. Randomly permuted  $(t, m, s)$ -nets and  $(t, s)$ -sequences. In *Monte Carlo and Quasi-Monte Carlo Methods in Sci-*

- entific Computing*, ed. H. Niederreiter and P. J.-S. Shiue, Number 106 in *Lecture Notes in Statistics*, 299–317. Springer-Verlag.
- Owen, A. B. 1997. Monte Carlo variance of scrambled equidistribution quadrature. *SIAM Journal on Numerical Analysis* 34 (5): 1884–1910.
- Owen, A. B. 1998a. Latin supercube sampling for very high-dimensional simulations. *ACM Transactions of Modeling and Computer Simulation* 8 (1): 71–102.
- Owen, A. B. 1998b. Scrambling Sobol and Niederreiter-Xing points. *Journal of Complexity* 14:466–489.
- Pages, G. 1997. A space quantization method for numerical integration. *Journal of Computational and Applied Mathematics* 89:1–38.
- Paskov, S., and J. Traub. 1995. Faster valuation of financial derivatives. *Journal of Portfolio Management* 22:113–120.
- Pirsic, G. 2001. A software implementation of Niederreiter-Xing sequences. In *Monte Carlo and Quasi-Monte Carlo Methods 2000*, ed. K.-T. Fang, F. J. Hickernell, and H. Niederreiter: Springer. To appear.
- Pirsic, G., and W. C. Schmid. 2001. Calculation of the quality parameter of digital nets and application to their construction. *Journal of Complexity*. To appear.
- Sloan, I. H., and S. Joe. 1994. *Lattice methods for multiple integration*. Oxford: Clarendon Press.
- Sloan, I. H., and L. Walsh. 1990. A computer search of rank-2 lattice rules for multidimensional quadrature. *Mathematics of Computation* 54:281–302.
- Sobol', I. M. 1967. The distribution of points in a cube and the approximate evaluation of integrals. *U.S.S.R. Comput. Math. and Math. Phys.* 7:86–112.
- Sobol', I. M. 1969. *Multidimensional quadrature formulas and Haar functions*. Moscow: Nauka. In Russian.
- Sobol', I. M. 1976. Uniformly distributed sequences with an additional uniform property. *USSR Comput. Math. Math. Phys. Academy of Sciences* 16:236–242.
- Sobol', I. M., and Y. L. Levitan. 1976. The production of points uniformly distributed in a multidimensional. Technical Report Preprint 40, Institute of Applied Mathematics, USSR Academy of Sciences. In Russian.
- Spanier, J. 1995. Quasi-Monte Carlo methods for particle transport problems. In *Monte Carlo and Quasi-Monte Carlo Methods in Scientific Computing*, ed. H. Niederreiter and P. J.-S. Shiue, Volume 106 of *Lecture Notes in Statistics*, 121–148. New York: Springer-Verlag.

- Spanier, J., and E. H. Maize. 1994. Quasi-random methods for estimating integrals using relatively small samples. *SIAM Review* 36:18–44.
- Tan, K. S., and P. P. Boyle. 2000. Applications of randomized low discrepancy sequences to the valuation of complex securities. *Journal of Economic Dynamics and Control* 24:1747–1782.
- Tausworthe, R. C. 1965. Random numbers generated by linear recurrence modulo two. *Mathematics of Computation* 19:201–209.
- Tezuka, S. 1987. Walsh-spectral test for GFSR pseudorandom numbers. *Communications of the ACM* 30 (8): 731–735.
- Tezuka, S. 1995. *Uniform random numbers: Theory and practice*. Norwell, Mass.: Kluwer Academic Publishers.
- Tezuka, S., and P. L'Ecuyer. 1991. Efficient and portable combined Tausworthe random number generators. *ACM Transactions on Modeling and Computer Simulation* 1 (2): 99–112.
- Tezuka, S., and T. Tokuyama. 1994. A note on polynomial arithmetic analogue of Halton sequences. *ACM Transactions on Modeling and Computer Simulation* 4:279–284.
- Tootill, J. P. R., W. D. Robinson, and D. J. Eagle. 1973. An asymptotically random Tausworthe sequence. *Journal of the ACM* 20:469–481.
- Tuffin, B. 1996. On the use of low-discrepancy sequences in Monte Carlo methods. Technical Report No. 1060, I.R.I.S.A., Rennes, France.
- Tuffin, B. 1998. Variance reduction order using good lattice points in Monte Carlo methods. *Computing* 61:371–378.
- Wang, D., and A. Compagner. 1993. On the use of reducible polynomials as random number generators. *Mathematics of Computation* 60:363–374.
- Wang, X., and F. J. Hickernell. 2000. Randomized Halton sequences. *Math. Comput. Modelling* 32:887–899.
- Yakowitz, S., J. E. Krimmel, and F. Szidarovszky. 1978. Weighted Monte Carlo integration. *SIAM Journal on Numerical Analysis* 15:1289–1300.
- Yakowitz, S., P. L'Ecuyer, and F. Vázquez-Abad. 2000. Global stochastic optimization with low-discrepancy point sets. *Operations Research* 48 (6): 939–950.