

Monte Carlo Method Using Different Types of Quasi-Random Sequences for Estimating Low-dimensional Nonlinear Mixed Effects Model

Sukanta Das^{1,2,*}

¹ Department of Statistics, Begum Rokeya University, Rangpur, Bangladesh

² Institute of Statistical Research and Training, University of Dhaka, Dhaka, Bangladesh

Received: 9 Oct. 2020, Revised: 22 Dec. 2020, Accepted: 12 Jul. 2022

Published online: 1 Sep. 2022

Abstract: Nonlinear mixed effects models involve both fixed effects and random effects in which some of the fixed and random effects parameters enter nonlinearly to the model function. These models are very popular for analyzing clustered data or unbalanced repeated measures data. There are several methods for estimating the parameters of these models. Most of the available methods are based on the approximation technique. In this study, instead of approximation based methods, quasi-Monte Carlo method is used, which directly solves the intractable multidimensional integrations in the nonlinear mixed effects model. To apply this method, the Michaelis-Menten model with one random effects parameter is used as a nonlinear mixed effects model. For this simulation study, different types of quasi-random sequences (e.g. Halton, Sobol, and Faure) are used with different combination of subjects and intra-subject correlation coefficient. The results of the simulation studies show that the quasi-Monte Carlo method correctly estimates the parameters of the model, and also these estimates are slightly different for using different types of quasi-random sequences. Among the quasi-random sequences, Halton sequences give better estimates than Sobol and Faure sequences for both fixed and random effects parameters of the nonlinear mixed effects models.

Keywords: Nonlinear mixed effects model, estimation, quasi-random sequences, simulation, quasi-Monte Carlo method

1 Introduction

Statistical models that incorporate both fixed effects and random effects are called mixed-effects models. Nonlinear mixed effects (NLME) models involve both fixed effects and random effects in which some of the fixed and random effects parameters enter non-linearly to the model function. Nonlinear mixed effects models extend linear mixed models by allowing the regression function to depend non-linearly on fixed and random effects. Thus NLME models have properties of both the nonlinear model and mixed effects model. NLME models are a popular platform for analysis when interest focuses on individual-specific characteristics [1,2]. The popularity of nonlinear mixed effects model lies in its interpretability, parsimony and validity beyond the observed range of data [3,4,5]. These models have received a great deal of attention in the statistical literature in recent years because of the flexibility they offer in handling unbalanced repeated measures data that arise in different areas of investigation, such as pharmacokinetics, agriculture, biochemistry, environment, medicine, and economics [6,7,8,9]. NLME models are very important and extensively used in various studies that are costly to conduct and when the underlying scenario cannot be explained or analyzed by linear model. Nonlinear models are parsimonious so that it can capture the nonlinear variation with a minimum number of parameters. Due to their great importance, fitting of nonlinear mixed effects models are also of crucial matters.

Several different formulations of nonlinear mixed effects models are available in literature [10,11,12,13,14,15]. Most commonly used one is proposed by Lindstrom and Bates [12]. Also, there are different estimation methods for the parameters in the nonlinear mixed effects model [16,17,18,19,20] and there is an ongoing debate in the literature about the most adequate methods [21,22]. One of the reasons for this variety of estimation methods is related to the numerical complexity involved in the derivation of restricted maximum likelihood (REML) estimates in the nonlinear mixed effects

* Corresponding author e-mail: sdas1@isrt.ac.bd

model. This complexity is due to the fact that the likelihood function in the nonlinear mixed effects model does not usually have a closed form expression. Different approximations to the log-likelihood in nonlinear mixed effects model have been proposed try to circumvent this problem [12, 14, 18]. Pinheiro and Bates [9] provided primary reference for the theory and computational techniques of NLME models. They describe and compare several different integrated likelihood approximations and provide evidence that adaptive Gaussian quadrature is one of the best methods. Devidian and Gallant [18] also use Gaussian quadrature for nonlinear mixed models, they advocate smooth nonparametric density for the random effects. Traditional approaches to fitting nonlinear mixed models involve Taylor series expansions, expanding around either zero or the empirical best linear unbiased predictions of the random effects [16, 23, 24, 25]. The Laplacian approximation technique and its relationship to the Lindstrom-Bates method are discussed by several authors [23, 25, 26, 27].

Sometimes in statistical problems, the integrand may be a probability density function not easily expressible in a form suitable for computation, but at the same time it may be easy to sample from the distribution. In such cases Monte Carlo methods of Integration is a plausible candidate. Quasi Monte Carlo methods converge much faster than normal Monte Carlo methods. Quasi-Monte Carlo methods are based on the idea that random Monte Carlo techniques can often be improved by replacing the underlying source of random numbers with a more uniformly distributed deterministic quasi-random sequence. These methods are now widely used in scientific computation, especially in estimating integrals over multidimensional domains and in many different financial computations [28, 29]. The standard Monte Carlo method is frequently used when the quadrature methods are difficult or expensive to implement [25, 30]. Marokoff and Caflisch [30] studied the performance of Monte Carlo and quasi-Monte Carlo methods for integration. In the paper, Halton, Sobol, and Faure sequences for quasi-Monte Carlo are compared with the standard Monte Carlo method that uses pseudo-random sequences. They found that the Halton sequence performs best for dimensions up to around 6; the Sobol sequence performs best for higher dimensions; and the Faure sequence, while outperformed by the other two, still performs better than a pseudo-random sequence. They also remark that the advantage of quasi-Monte Carlo method is greater if the integrand is smooth, and the number of dimensions “s” of the integral is small. Hoque and Latif [31] proposed a new method for fitting nonlinear mixed effects models through quasi-Monte Carlo integration method using Sobol’s sequences. Without any approximations, they directly solved the intractable integration in the likelihood function of nonlinear mixed effects models using this quasi-Monte Carlo integration technique. They also showed that, the performance of Sobol’s sequence based method is reasonably well by comparing among other existing approximation based methods.

This study mainly addresses the estimation of nonlinear mixed effects model based on quasi-Monte Carlo (QMC) integration method using quasi-random sequences such as Halton, Sobol, and Faure sequences. Quasi-random sequences are also called low-discrepancy sequences; and discrepancy is a measure of deviation from the uniformity of a sequence of points. Discrepancy contributes to the error in quasi-Monte Carlo methods. Successful use of QMC methods in high dimensions depends on the construction of good sequences and the intelligent use or the sequences for path generation. Quasi-Monte Carlo integration method uses low-discrepancy sequences and most of the cases it offer faster rate of convergence and also delivers more accurate results in a shorter time than the usual Monte Carlo method that uses pseudo-random sequences [32]. In general, this study focuses on the behavior of quasi-random sequences (e.g. Halton, Sobol, Faure) for estimating NLME models. For this study the Michaelis-Menten model with one random effects parameter is used as a low-dimensional nonlinear mixed effects model, and then the quasi-Monte Carlo method is applied for estimating the parameters of the model. Finally, the simulation is performed with different combination of subjects and intra-subject correlation coefficient.

2 Nonlinear Mixed Effects Model

Nonlinear mixed effects (NLME) models are mixed effects models in which some, or all, of the fixed and random effects parameters enter nonlinearly to the model function. Nonlinear mixed effects models can be regarded either as an extension of linear mixed-effects models in which the conditional expectation of the response given the random effects is allowed to be a nonlinear function of the parameters, or as an extension of the nonlinear regression models [33, 34] to fit data from several individuals in which some of the regression parameters are considered as random. In this section a general formulation of nonlinear mixed effects models is presented with suitable notations.

The general formulation of NLME models for repeated measures data proposed by Lindstrom and Bates [12] can be thought of as a hierarchical model. At one level the j^{th} observation on the response y in the i^{th} group is modeled as

$$y_{ij} = f(\phi_{ij}, \mathbf{x}_{ij}) + \varepsilon_{ij}, \quad i = 1, \dots, m, \quad j = 1, \dots, n_i, \quad (1)$$

where m is the number of groups, n_i is the number of observations in the i^{th} group, f is a general, real-valued, differentiable function of a group specific parameter vector ϕ_{ij} and a covariate vector $\mathbf{x}_{ij}$, and ε_{ij} is a normally distributed within-group

error term. The function f is nonlinear in at least one component of the group-specific parameter vector ϕ_{ij} , which is modeled as

$$\phi_{ij} = \mathbf{A}_{ij}\boldsymbol{\beta} + \mathbf{B}_{ij}\mathbf{b}_i, \quad \mathbf{b}_i \sim N(\mathbf{0}, \mathbf{D}), \quad (2)$$

where $\boldsymbol{\beta}$ is a p -dimensional vector of fixed effects and $\mathbf{b}_i$ is a q -dimensional random effects vector associated with the i^{th} group (not varying with j) with variance-covariance matrix $\mathbf{D}$. The matrices $\mathbf{A}_{ij}$ and $\mathbf{B}_{ij}$ are of appropriate dimensions and depend on the group and possibly on the values of some covariates at the j^{th} observation. It is assumed that observations corresponding to different groups are independent and that the within-group errors ε_{ij} are independently distributed as $N(0, \sigma^2)$ and independent of the $\mathbf{b}_i$.

Because $f(\cdot)$ can be any nonlinear function of ϕ_{ij} , the representation of the group-specific coefficients ϕ_{ij} could be chosen so that $\mathbf{A}_{ij}$ and $\mathbf{B}_{ij}$ are always simple incidence matrices. However, it is desirable to encapsulate as much modeling of the ϕ_{ij} as possible in this second stage, as this simplifies the calculation of the derivatives of the model function with respect to $\boldsymbol{\beta}$ and $\mathbf{b}_i$, used in the optimization algorithm. The arguments *fixed* and *random* are used to specify the $\mathbf{A}_{ij}$ and $\mathbf{B}_{ij}$ matrices, respectively.

For i^{th} group we can write equation (1) and equation (2) in matrix form as

$$\begin{aligned} \mathbf{y}_i &= \mathbf{f}_i(\phi_i, \mathbf{x}_i) + \boldsymbol{\varepsilon}_i \\ \phi_i &= \mathbf{A}_i\boldsymbol{\beta} + \mathbf{B}_i\mathbf{b}_i, \end{aligned} \quad (3)$$

where for $i = 1, \dots, m$

$$\begin{aligned} \mathbf{y}_i &= \begin{bmatrix} y_{i1} \\ \vdots \\ y_{in_i} \end{bmatrix}, \quad \phi_i = \begin{bmatrix} \phi_{i1} \\ \vdots \\ \phi_{in_i} \end{bmatrix}, \quad \boldsymbol{\varepsilon}_i = \begin{bmatrix} \varepsilon_{i1} \\ \vdots \\ \varepsilon_{in_i} \end{bmatrix}, \quad \mathbf{f}_i(\phi_i, \mathbf{x}_i) = \begin{bmatrix} f(\phi_{i1}, \mathbf{x}_{i1}) \\ \vdots \\ f(\phi_{in_i}, \mathbf{x}_{in_i}) \end{bmatrix}, \\ \mathbf{x}_i &= \begin{bmatrix} \mathbf{x}_{i1} \\ \vdots \\ \mathbf{x}_{in_i} \end{bmatrix}, \quad \mathbf{A}_i = \begin{bmatrix} \mathbf{A}_{i1} \\ \vdots \\ \mathbf{A}_{in_i} \end{bmatrix}, \quad \mathbf{B}_i = \begin{bmatrix} \mathbf{B}_{i1} \\ \vdots \\ \mathbf{B}_{in_i} \end{bmatrix}. \end{aligned}$$

3 Estimation of Nonlinear Mixed Effects Model

Different methods have been proposed to estimate the parameters in the nonlinear mixed effects models. The methods may be either likelihood-based or Bayesian, parametric or nonparametric. Maximum likelihood parametric method proposed by Pinheiro and Bates [5] is the most commonly used method for estimating parameters of the NLME model. Descriptions and comparisons of other estimation methods proposed for NLME models can be found here [2, 20, 35, 36]. In the next subsection a brief description of maximum likelihood estimation technique of NLME models is given.

3.1 Maximum Likelihood Estimation of NLME Model

Because the random effects $\mathbf{b}_i$ in equation (3) are unobserved quantities, maximum likelihood estimation in mixed-effects models is based on the marginal density of the responses $\mathbf{y}$, which is calculated as

$$p(\mathbf{y}|\boldsymbol{\beta}, \sigma^2, \mathbf{D}) = \prod_{i=1}^m p(\mathbf{y}_i|\boldsymbol{\beta}, \sigma^2, \mathbf{D}) = \prod_{i=1}^m \int_{\mathbf{b}_i} p(\mathbf{y}_i|\mathbf{b}_i; \boldsymbol{\beta}, \sigma^2) p(\mathbf{b}_i|\mathbf{D}, \sigma^2) d\mathbf{b}_i, \quad (4)$$

where the conditional density of $\mathbf{y}_i$ given random effects $\mathbf{b}_i$ is multivariate normal and can be expressed as

$$p(\mathbf{y}_i|\mathbf{b}_i; \boldsymbol{\beta}, \sigma^2) = \frac{1}{(2\pi\sigma^2)^{n_i/2}} \exp \left[-\frac{1}{2\sigma^2} \|\mathbf{y}_i - \mathbf{f}_i(\boldsymbol{\beta}, \mathbf{b}_i)\|^2 \right] \quad (5)$$

where $\mathbf{f}_i(\boldsymbol{\beta}, \mathbf{b}_i) = \mathbf{f}_i(\phi_i, \mathbf{x}_i)$, and the marginal density of $\mathbf{b}_i$ is also multivariate normal, which is

$$p(\mathbf{b}_i|\boldsymbol{\Delta}, \sigma^2) = \frac{1}{(2\pi)^{q/2} \sqrt{|\mathbf{D}|}} \exp \left[-\frac{1}{2} \mathbf{b}_i' \mathbf{D}^{-1} \mathbf{b}_i \right] = \frac{1}{(2\pi\sigma^2)^{q/2} |\boldsymbol{\Delta}^{-1}|} \exp(-\|\boldsymbol{\Delta} \mathbf{b}_i\|^2 / 2\sigma^2), \quad (6)$$

where the expression $|\Delta|$ denotes the determinant of the random effects variance-covariance matrix in terms of the precision factor Δ , so that $\Delta^{-1} = \sigma^{-2} \Delta' \Delta$. Therefore, the marginal density of $\mathbf{y}$ is

$$p(\mathbf{y}_i | \boldsymbol{\beta}, \sigma^2, \Delta) = \prod_{i=1}^m \int_{\mathbf{b}_i} \frac{|\Delta|}{(2\pi\sigma^2)^{(n_i+q)/2}} \exp \left\{ \frac{\|\mathbf{y}_i - \mathbf{f}_i(\boldsymbol{\beta}; \mathbf{b}_i)\|^2 + \|\Delta \mathbf{b}_i\|^2}{-2\sigma^2} \right\} d\mathbf{b}_i. \quad (7)$$

To make the numerical optimization of this likelihood function a tractable problem arise because the model function $f(\cdot)$ can be nonlinear in the random effects, the integral in equation (7) generally does not have a closed-form expression. Different approximations of equation (7) have been proposed, some of these methods consist of taking a first-order Taylor expansion of the model function $f(\cdot)$ around the expected value of the random effects [14,23], or around the conditional (on Δ) modes of the random effects [12]. Gaussian quadrature rules have also been used [25]. These approximation methods have increasing degrees of accuracy, at the cost of increasing computational complexity.

3.2 Specific Example of Nonlinear Mixed Effects Model

There are many nonlinear mixed effects models that are widely used in different fields, especially in medical statistics and Biochemistry. In this section a brief description of Michaelis-Menten (MM) model for enzyme kinetics has given as a specific example of NLME models. Here, it will be shown that how MM model is a nonlinear model containing both fixed and random effects parameters.

3.2.1 The Michaelis-Menten Model

In biochemistry, Michaelis-Menten kinetics is one of the simplest and best-known model of enzyme kinetics [37]. It is named after German biochemist Leonor Michaelis and Canadian pathologist Maud Menten. Michaelis-Menten mechanism takes the form of an equation relating reaction velocity to substrate concentration for a system where a substrate S binds reversibly to an enzyme E to form an enzyme-substrate complex ES , which then reacts irreversibly to generate a product P and to regenerate the free enzyme E . This system can be represented schematically as follows:



where k_1 , k_{-1} and k_2 denote the rate constants of the three reactions. The model takes the form of an equation describing the rate of enzymatic reactions, by relating reaction rate v to $[S]$, the concentration of a substrate S . Its formula is given by

$$v = \frac{V_{max}[S]}{K_m + [S]}, \quad (9)$$

where V_{max} represents the maximum rate achieved by the system, at maximum (saturating) substrate concentrations. The Michaelis constant K_m is the substrate concentration at which the reaction rate is half of V_{max} . Biochemical reactions involving a single substrate are often assumed to follow Michaelis-Menten kinetics, without regard to the model's underlying assumptions. Figure (1) shows the graphical representation of Michaelis-Menten model in equation (9).

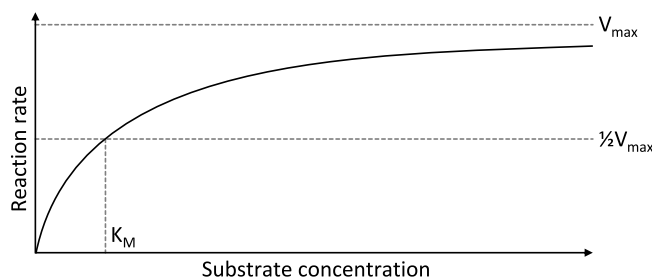


Fig. 1: Graphical representation of Michaelis-Menten model

Here, the function shows a steep rise that gradually flattens into a plateau. Initially, small increases in substrate concentration cause a substantial increase in reaction rate, but eventually, further increases in concentration cause progressively smaller increases in reaction rate, and eventually reaction rate plateaus and does not respond to further increases in concentration.

3.2.2 Michaelis-Menten Model with Random V_{max} Parameter: A NLME Model Perspective

For repeated measure data, suppose we have m subjects each having n repeated observations. For the j -th observation of the i -th subject the Michaelis-Menten model (9) can be written as follows

$$\begin{aligned} y_{ij} &= \frac{V_i x_{ij}}{K + x_{ij}} + \varepsilon_{ij} \\ &= \frac{e^{\theta_i} x_{ij}}{e^{\beta} + x_{ij}} + \varepsilon_{ij}, \quad i = 1, \dots, m, \quad j = 1, \dots, n, \end{aligned} \quad (10)$$

where y_{ij} is the reaction rate and x_{ij} is the concentration for the j -th measurements on the i -th individual. Also, $\theta_i = \log(V_i) = \log(V) + b_i$, $\beta = \log(K)$, $\varepsilon_{i,j} \sim N(0, \sigma^2)$ and b_i is random effects for subject i with $b_i \sim N(0, \sigma_b^2)$.

Here, V_i represents the maximum rate for the i -th subject achieved by the system, which varies from subject to subject due to random effect b_i with $V_{max} = V$ for notational convenience. And $K_m = K$, the Michaelis constant, is the substrate concentration x at which the reaction rate is half of the maximum rate. The fixed effects parameters to be estimated are V_{max} and K_m or, θ and β , and the random effects parameter to be estimated is σ_b . For the i -th subject the equation (10) can be written as

$$\mathbf{y}_i = \mathbf{f}_i(\beta, \theta, b_i) + \boldsymbol{\varepsilon}_i, \quad i = 1, \dots, m, \quad (11)$$

where $\mathbf{y}_i = [y_{i1}, \dots, y_{in}]^T$, $\mathbf{f}_i(\beta, \theta, b_i) = [f_{i1}(\beta, \theta, b_i), \dots, f_{in}(\beta, \theta, b_i)]^T$, $f_{ij}(\beta, \theta, b_i) = \frac{e^{\theta_i} x_{ij}}{e^{\beta} + x_{ij}}$, $\boldsymbol{\varepsilon}_i = [\varepsilon_{i1}, \dots, \varepsilon_{in}]^T$ with $\boldsymbol{\varepsilon}_i \sim N(\mathbf{0}, I\sigma^2)$.

Here, the expression of MM model in equation (11) is a nonlinear mixed effects model and the parameters to be estimated are the fixed effects θ and β , within subject variance σ^2 and between subject variance or variance of the random effects σ_b^2 . As the likelihood function of NLME models contain intractable integrations that have no closed form of expression. This intractable integration can directly be solved by using well known quasi-Monte Carlo integration technique for different types of quasi random sequences that is the main goal of this paper.

4 Quasi-Random Sequences

The use of quasi-random, rather than random, numbers in Monte Carlo methods, is called quasi-Monte Carlo methods, which converge much faster than normal Monte Carlo method. Quasi-Monte Carlo methods are now widely used in scientific computation, especially in estimating integrals over multidimensional domains and in many different financial computations. The quasi random sequences are also called low-discrepancy sequences, due to their common use as a replacement of uniformly distributed random numbers. The discrepancy is a measure of deviation from uniformity of a sequence of points in $U = ([0, 1]^s)$. The low-discrepancy sequences cover the unit cube as ‘uniformly’ as possible by reducing gaps and clustering of points.

A sequence of n -tuples that fills n -space more uniformly than uncorrelated random points, sometimes also called a low-discrepancy sequence. In another statement, a low-discrepancy sequence is a set of s -dimensional points, filling the sample area “efficiently” and that has a lower discrepancy than the straight pseudo-random number set. The quasi-random numbers have the low-discrepancy (LD) property that is a measure of uniformity for the distribution of the point mainly for the multidimensional case. The main advantage of the quasi-random sequence in comparison to pseudo-random sequence is it distributes evenly hence there is no larger gaps and no cluster formation, this leads to spread the number over the entire region. The concept of LD is associated with the property that the successive numbers are added is a position as away as possible from the other numbers that is, avoiding clustering (grouping of numbers close to each other). The sequence is constructed based on some pattern such that each point is separated from the others, this leads to maximal separation between the points. This process takes care of evenly distribution random numbers in the entire search space $[0, 1]$.

The low-discrepancy sequences are generated in a highly correlated manner, i.e. the next point knows where the previous points are. There are different types of low-discrepancy sequences such as Halton sequences, Faure sequences, Sobol sequences, etc. The most fundamental LD sequence for one dimension is generated by Van der Corput method,

further to continue random sequence in higher dimension Faure, Sobol and Halton method are used. This paper is designed as that the nonlinear mixed effects models are to be fitted by solving multidimensional integration directly using the various types of quasi-random sequences based on quasi-Monte Carlo integration method.

4.1 Importance of Quasi-random Sequences

Although the ordinary uniform random numbers and quasi-random sequences both produce uniformly distributed sequences, there is a big difference between these two. A uniform random generator on $[0,1)$ will produce outputs so that each trial has the same probability of generating a point on equal subintervals, for example $[0, 1/2)$ and $[1/2, 1)$. Therefore, it is possible for n trials to coincidentally all lie in the first half of the interval, while the $(n+1)^{st}$ point still falls within the other of the two halves with probability $1/2$. This is not the case with the quasi-random sequences, in which the outputs are constrained by a low-discrepancy requirement that has a net effect of points being generated in a highly correlated manner, i.e. the next point “knows” where the previous points are.

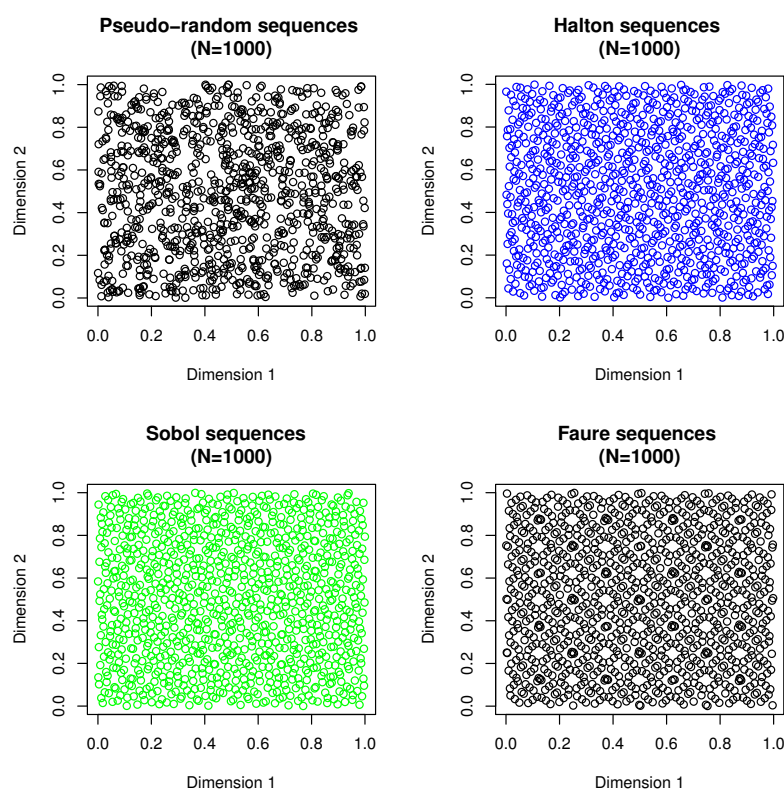


Fig. 2: Pattern of different low-discrepancy sequences with pseudo-random sequences. For $N=1000$ points and Dimensions =2

Quasi-random sequences or low-discrepancy sequences tend to sample space “*more uniformly*” than the uniform or pseudo-random numbers. Such a sequence is extremely useful in computational problems where numbers are computed on a grid, but it is not known in advance how fine the grid must be to obtain accurate results. Using a quasi-random sequence allows stopping at any point where convergence is observed, whereas the usual approach of halving the interval between subsequent computations requires a huge number of computations between stopping points.

Figure (2) shows the first 1000 points of different quasi-random sequences and pseudo-random sequences for dimension 2. From this figure we see that, quasi-random sequences (e.g. Halton, Sobol, Faure) are more uniformly and evenly distributed than the pseudo-random sequences. The quasi-random sequences are highly correlated than the pseudo-random sequences in the sense that quasi-random sequences fill-up the gap more deterministically. Here we also see that there are some small differences among different types of quasi-random sequences. The main purposes in this

study was to explore the general behavior of such quasi-random sequences to solve multidimensional integration numerically by direct quasi-Monte Carlo integration method.

4.2 Quasi-Monte Carlo Integration (QMCI)

In numerical analysis, quasi-Monte Carlo Integration (QMCI) method is a method for the computation of an integral that uses quasi-random sequences (i.e., also called low-discrepancy sequences) which have a more uniform behavior to compute the integral. Quasi-random numbers are generated algorithmically by computer, and are similar to pseudo-random numbers while having the additional important property of being deterministically chosen based on equally distributed sequences in order to minimize errors [38].

This is in contrast to a regular Monte Carlo method, which is based on sequences of pseudo-random numbers. The QMCI methods use points that are evenly distributed; the points are spread over the domain in such a way that there are no clusters. The classical QMCI method replaces the independent random points used in MCI by a deterministic set of distinct points that cover the region of integration more uniformly. The use of quasi-random sequences in place of the usual pseudo-random numbers often improves the convergence of the numerical integration and it is also possible to compute an absolute bound for the error [39].

Monte Carlo and quasi-Monte Carlo methods are stated in a similar way. The problem is to approximate the integral of a function $f(\cdot)$ as the average of the function evaluated at a set of points $x_1, \dots, x_N$.

$$\int_{[0,1]^s} f(u) du \approx \frac{1}{N} \sum_{i=1}^N f(x_i). \quad (12)$$

Since we are integrating over the s -dimensional unit cube, each x_i is a vector of s elements. The difference between quasi-Monte Carlo and Monte Carlo is the way that x_i are chosen. Quasi-Monte Carlo method uses a quasi-random sequences such as the Halton sequence, the Sobol sequence, or the Faure sequence, whereas simple Monte Carlo uses a pseudo-random sequence. The advantage of using quasi-random sequences is a faster rate of convergence. Quasi-Monte Carlo has a rate of convergence close to $O(1/N)$, whereas the rate for the simple Monte Carlo method is $O(N^{-1/2})$ [32].

4.3 Estimating Parameters of NLME Model using Quasi-random Sequences

The quasi-random sequences are very useful in solving Monte Carlo Integration and hence to construct likelihood. Monte Carlo Integration method that uses quasi random sequences is called Quasi Monte Carlo Integration method. By this method we directly solve the multidimensional integration (7) of NLME models without considering any approximation technique. The likelihood construction of Michaelis-Menten model (10) as a NLME models with one random effects parameter (V_{max} as random) is described below.

The Michaelis-Menten (10) can be written as,

$$y_{ij}|b_i = \frac{e^{\theta_i} x_{ij}}{e^{\beta} + x_{ij}} + \varepsilon_{ij} = f_{ij} + \varepsilon_{ij}, \quad i = 1, \dots, m, \quad j = 1, \dots, n,$$

where $f_{ij} = \frac{e^{\theta_i} x_{ij}}{e^{\beta} + x_{ij}}$, $\theta_i = \theta + b_i$, $b_i \sim N(0, \sigma_b^2)$ and $\varepsilon_{ij} \sim N(0, \sigma^2)$. From here parameters to be estimated are the fixed effects β and θ , within subject variance σ^2 and between subject variance or variance of random effects σ_b^2 .

From equation (4) the marginal density of response y_i or likelihood for i th subject can be written as

$$p(y_i|\beta, \theta, \sigma^2, \sigma_b^2) = \int_{b_i} p(y_i|b_i; \beta, \theta, \sigma^2) p(b_i|\sigma_b^2) db_i.$$

Here, $y_{ij}|b_i \sim N(f_{ij}, \sigma^2)$, for $j = 1, \dots, n$, $y_{ij}|b_i$'s are independent of one another. So, $p(y_i|b_i; \beta, \theta, \sigma^2) \sim N_n(\mathbf{f}_i, \mathbf{I}\sigma^2)$, where $\mathbf{f}_i = (f_{i1}, \dots, f_{in})'$ and $\mathbf{I}$ is an $(n \times n)$ identity matrix.

The conditional density $p(y_i|b_i; \beta, \theta, \sigma^2)$ can be expressed as

$$p(y_i|b_i; \beta, \theta, \sigma^2) = \frac{1}{(2\pi\sigma^2)^{n/2}} \exp \left[-\frac{1}{2\sigma^2} \sum_{j=1}^n (y_{ij} - f_{ij})^2 \right]$$

The marginal density of b_i is also multivariate normal can be expressed as

$$p(b_i|\sigma_b^2) = \frac{1}{(2\pi\sigma_b^2)^{1/2}} \exp\left[-\frac{1}{2\sigma_b^2}\right]$$

So, the likelihood function for i th subject is written as

$$p(\mathbf{y}_i|\beta, \theta, \sigma^2, \sigma_b^2) = \int_{b_i} \frac{1}{(2\pi\sigma^2)^{n/2}(2\pi\sigma_b^2)^{1/2}} \exp\left[-\frac{1}{2\sigma^2} \sum_{j=1}^n (y_{ij} - f_{ij})^2 - \frac{1}{2\sigma_b^2} b_i^2\right] db_i \quad (13)$$

Then the complete likelihood is obtained as

$$L_C = \prod_{i=1}^m p(\mathbf{y}_i|\beta, \theta, \sigma^2, \sigma_b^2)$$

and the log-likelihood function is

$$\ell_C = \sum_{i=1}^m \log p(\mathbf{y}_i|\beta, \theta, \sigma^2, \sigma_b^2).$$

In equation (13) the multidimensional integration is evaluated using the Monte Carlo integration rule with the help of quasi-random sequences. Here, the value of b_i are generated from normal quasi-random sequences (e.g. Halton, Sobol and Faure etc.), with specific mean 0 and variance σ^2 . This log-likelihood function is to be optimized for estimating all fixed effects and random effects parameters considered here as this is a function of these parameters. The next section will describe the simulation study for estimating parameters of Michaelis-Menten model in equation (10).

5 Simulation Study for Estimating Parameters of NLME Model

The mixed effects Michaelis-Menten model that we have described in section 3.2.2 at equation (10) is considered here for this simulation study. This NLME model considers m subject and each having n repeated observations. For j th observations of i th subject the model (10) is rewritten as follows:

$$\begin{aligned} y_{ij} &= \frac{V_i x_{ij}}{K + x_{ij}} + \varepsilon_{ij} \\ &= \frac{e^{\theta_i} x_{ij}}{e^{\beta} + x_{ij}} + \varepsilon_{ij}, \quad i = 1, \dots, m, \quad j = 1, \dots, n, \end{aligned}$$

where $\theta_i = \log V_i = \theta + b_i$, $\beta = \log K$, $\varepsilon_{ij} \sim N(0, \sigma^2)$ and b_i is the random effect for subject i with $b_i \sim N(0, \sigma_b^2)$, $V_{\max} = V$ and $K_m = K$. Here, x_{ij} denotes substrate concentration and y_{ij} denotes reaction rate. V_i represents the maximum rate of reaction for i th subject or individual, which varies from individual to individual due to the random effect b_i . The Michaelis constant K is the value of substrate concentration at which reaction rate is half of the maximum rate.

The considered Michaelis-Menten model (10) has two fixed effects parameters ($V_{\max}, K_m$) and two random effects parameters (σ_b^2, σ^2) (i.e., actually the variance components). The true values of the fixed effects parameters that we used for this simulation study are $V_{\max} = 5$ and $K_m = 2$, i.e. $\theta = \log(5)$ and $\beta = \log(2)$. So, from this simulation study the fixed effects parameters to be estimated are $V_{\max}$ and K_m or θ and β and the random effects σ_b , i.e. the variance component with which random effect is associated. Simulation results are examined for different values of the variance components depending on the values of the intra-class correlation efficient ($ICC = \frac{\sigma_b^2}{\sigma_b^2 + \sigma_e^2}$) and also for different sets of concentration levels, i.e. different designs. The design that we consider here contains 21 levels of substrate concentration ranging from 0.5 to 10, at increment 0.5 and 50, i.e. substrate concentration $x = (0.5, 1.0, 1.5, \dots, 9.5, 10, 50)$ is replicated for each subjects.

5.1 Simulation Settings

In this section, we will present our details simulation scenarios for this particular study. In order to investigate the general behavior of low-discrepancy sequences for solving intractable integration in the nonlinear mixed effects model by direct

quasi-Monte Carlo method, we considered some cases in this simulation study. The cases are different because different number of subjects are considered in each case. Within each of the cases, two situations with different set of variance components or random effects parameters were considered. That is, under each cases, the simulation scenarios were considered here: (i) with very high intra-subject correlation coefficient $\rho = .96$, and (ii) with low intra-subject correlation coefficient $\rho = .50$. The simulation were run under the settings summarized in Table (1). In all simulations, fixed effects parameters remain fixed and the true value considered for this parameters are $V_{max} = 5$ and $K_m = 2$, i.e. $\theta = \log(5)$ and $\beta = \log(2)$ under the model (10). For each scenarios, 1000 simulations were conducted.

Table 1: Different cases or scenarios are considered for the simulation study

CASES	Number of Subjects	ρ (ICC)	σ_b	σ
Case I	10	$\rho = .96$	$\sigma_b = .50$	$\sigma = .10$
		$\rho = .50$	$\sigma_b = .50$	$\sigma = .50$
Case II	20	$\rho = .96$	$\sigma_b = .50$	$\sigma = .10$
		$\rho = .50$	$\sigma_b = .50$	$\sigma = .50$
Case III	5	$\rho = .96$	$\sigma_b = .50$	$\sigma = .10$
		$\rho = .50$	$\sigma_b = .50$	$\sigma = .50$

The quasi-Monte Carlo method that we have considered for estimating the nonlinear mixed effects model, uses various types of quasi-random sequences such as Halton, Sobol and Faure sequences. The performance of these sequences for estimating the parameters of NLME model are of the main interest. In general, the effects and natural behavior of such sequences in the quasi-Monte Carlo method are investigated through this simulation study. This study also investigated the effect of the number of subjects considered, the random effects parameters (i.e., the variance components of the model) for different intra-class correlation coefficients on the parameter estimation of the model (10) by estimating the empirical bias and mean square error (MSE) of the point estimates. Bias in the estimates of the parameters was calculated as the differences between the true and estimated values (simulated average over 1000 simulations) for each parameters. The MSE was calculated as the expected value of the squared differences between the true and estimated values of the parameters. Here MSE is a risk function, corresponding to the expected value of the squared error loss or the quadratic loss. The mathematical formula of bias and MSE are given below,

$$\text{Bias}(V_{max}) = V_{max} - \hat{V}_{max} \quad (14)$$

$$\text{MSE}(V_{max}) = E[V_{max} - \hat{V}_{max}]^2, \quad (15)$$

where V_{max} is the true value of the parameter, $\hat{V}_{max}$ is the estimated value of the parameter V_{max} which is calculated by averaging the estimated values of 1000 simulations. Similarly, the bias and mean square error (MSE) for K_m and σ_b were calculated as stated above in equations (14) and (15). This study also used scatter plots of fixed and random effects parameter estimates for visually representing the differences.

5.2 Simulation with Case I (Number of Subjects, $n = 10$)

In this subsection, initially, the simulation study was started through Case I by selecting the number of subjects, $n = 10$. In Case I, two situations were also considered depending on the value of intra-subject or intra-class correlation coefficients ρ . Firstly, the simulation was conducted with very high intra-class correlation coefficient ($ICC = 0.96$) and then for low intra-class correlation coefficients ($ICC = 0.50$). The details simulation results of these two scenarios are shown in the next two paragraphs given below.

For Case I with 10 subjects and with very high intra-class correlation coefficient the simulation are conducted by considering true values of variance components $\sigma_b = 0.50$ and $\sigma = 0.1$. For this combination of σ_b and σ intra-subject correlation coefficient ρ is very high i.e. 0.96. The results of quasi-Monte Carlo method for this scenario are reported in Table (2), where mean, bias, mean square error (MSE) for each parameter are reported separately for each quasi-random sequences (e.g. Halton, Sobol, Faure).

Table (2) shows that the estimates of the fixed effects parameters (V_{max} and K_m) both have negligible bias for all sequences under the estimation by quasi-Monte Carlo method. The Halton sequences has less bias and MSE compare to the other sequences. In the estimate of random effects parameter σ_b , all the sequences and methods produce negatively biased estimate but the bias is also negligible in the sense that the bias is not more than 10 times greater than the simulated

Table 2: Simulation results of parameter estimation in the Michaelis-Menten model, when $n = 10$, and $\rho = 0.96$

Q.R.S.	V_{max}			K_m			σ_b		
	Mean	Bias	MSE	Mean	Bias	MSE	Mean	Bias	MSE
Halton	5.0025	0.0025	0.0204	2.0001	0.0001	0.0010	0.4699	-0.0301	0.0140
Sobol	4.9958	-0.0042	0.0263	2.0001	0.0001	0.0010	0.4599	-0.0401	0.0140
Faure	4.9951	-0.0049	0.0274	2.0002	0.0002	0.0010	0.4591	-0.0409	0.0141

1000 simulations are used with true parameter values $V_{max} = 5$, $K_m = 2$, $\sigma_b = 0.5$, MSE: mean square error, Q.R.S: quasi-random sequences.

standard deviation of the corresponding estimates (the results is not shown here). Here, the estimate of fixed effects parameters are more precise than the random effects parameter σ_b , because the precision of the estimate of fixed effects (as well as σ^2) is determined by the total number of observations, while the precision of the estimate of σ_b is determined by the number of cluster of subjects [9]. This interesting property of random effects estimates for other two cases are shown in the next subsection (5.3) and (5.4). Figure (3) represents the scatter plots of the fixed (V_{max} and K_m) and random effects (σ_b) estimates for all combination of Halton, Sobol and Faure sequences in the quasi-Monte Carlo method. Figure (3) shows the reflection of results from Table (2), Halton and Sobol sequences produce less mean square error (MSE) compare to the Faure sequences for estimating V_{max} parameter, and also the Faure sequences produce more MSE than the Halton and Sobol sequences for estimating K_m and random effects σ_b .

Table 3: Simulation results of parameter estimation in the Michaelis-Menten model, when $n = 10$, and $\rho = 0.50$

Q.R.S.	V_{max}			K_m			σ_b		
	Mean	Bias	MSE	Mean	Bias	MSE	Mean	Bias	MSE
Halton	5.0025	0.0025	0.0406	1.9958	-0.0042	0.0253	0.4521	-0.0479	0.0156
Sobol	5.0046	0.0046	0.0413	1.9958	-0.0042	0.0253	0.4519	-0.0481	0.0157
Faure	4.9939	-0.0061	0.0420	1.9957	-0.0043	0.0254	0.4513	-0.0487	0.0163

1000 simulations are used with true parameter values $V_{max} = 5$, $K_m = 2$, $\sigma_b = 0.5$, MSE: mean square error, Q.R.S: quasi-random sequences.

The simulations by two scenarios under Case I were mainly based on the same number of subjects but with different values of intra-subject correlation coefficients. Here, in this scenario of simulation, a low values of intra-subject correlation coefficient (i.e. $\rho = .50$) was considered for the combination of $\sigma_b = 0.50$ and $\sigma = 0.50$. The results of these simulations study are summarized in Table (3). The parameter estimates are almost same as the previous results. The fixed effects estimates of K_m parameter are very similar for all sequences but in the V_{max} and random effects σ_b estimate, the Halton and Sobol sequences produce less mean square error (MSE) than the Faure sequences. The scatter plots (not shown here) are also exhibit same results as the tabulated results.

5.3 Simulation with Case II (Number of Subjects, $n = 20$)

Now in this subsection, we want to show the effects of the number of subjects in parameter estimation of Michelis Menten model (10) by quasi-Monte Carlo method considering different quasi-random sequences. This study considered the Case II and Case III, where we have 20 and 5 number of subjects, which are illustrated in this subsection and the next subsection (5.4). First, we increase the number of subjects from 10 to 20 and check the results of the simulation study for two different scenarios with respect to the high and low intra-subject correlation coefficients (ρ). These results are given in two separate Tables (4), and (5) by varying different intra-class correlation coefficients (i.e., $\rho = 0.96$ and $\rho = 0.50$).

Table 4: Simulation results of parameter estimation in the Michaelis-Menten model, when $n = 20$, and $\rho = 0.96$

Q.R.S.	V_{max}			K_m			σ_b		
	Mean	Bias	MSE	Mean	Bias	MSE	Mean	Bias	MSE
Halton	5.0006	0.0006	0.0176	1.9998	-0.0002	0.0006	0.4875	-0.0125	0.0074
Sobol	4.9991	-0.0009	0.0179	1.9997	-0.0003	0.0006	0.4870	-0.0130	0.0075
Faure	5.0013	0.0013	0.0256	1.9997	-0.0003	0.0006	0.4853	-0.0147	0.0077

1000 simulations are used with true parameter values $V_{max} = 5$, $K_m = 2$, $\sigma_b = 0.5$, MSE: mean square error, Q.R.S: quasi-random sequences.

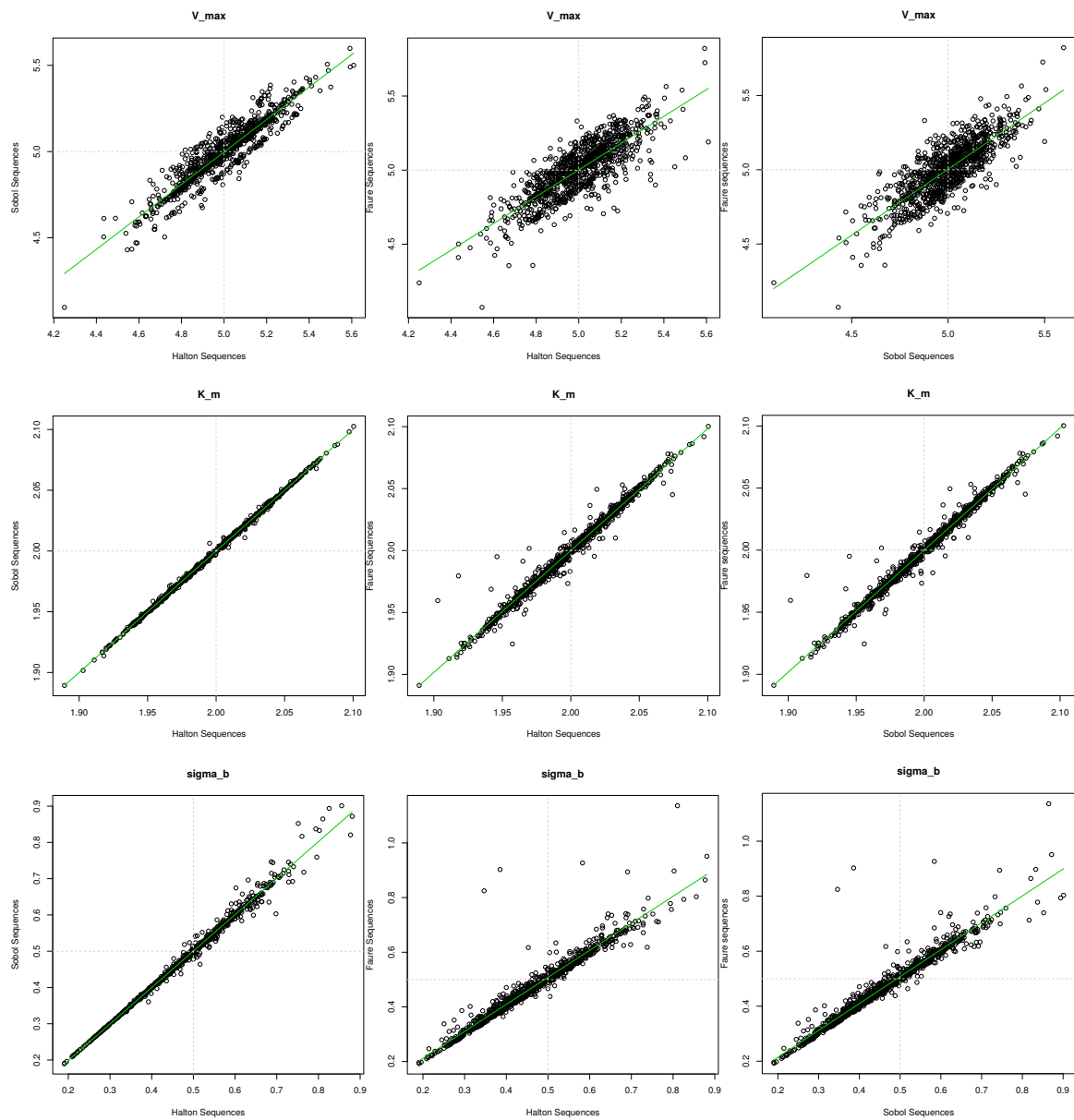


Fig. 3: Scatter plots of fixed and random effects estimates for all combination of the Halton, Sobol and Faure sequences in the quasi-Monte Carlo method of estimating the model (10). Here, the simulation is conducted by considering $n = 10$ subjects and intra-subject correlation coefficient $\rho = .96$. The dashed lines indicate the true values of the parameters.

Table 5: Simulation results of parameter estimation in the Michaelis-Menten model, when $n = 20$, and $\rho = 0.50$

Q.R.S.	V_{max}			K_m			σ_b		
	Mean	Bias	MSE	Mean	Bias	MSE	Mean	Bias	MSE
Halton	5.0011	0.0011	0.0192	1.9989	-0.0011	0.0136	0.4814	-0.0186	0.0076
Sobol	4.9927	-0.0073	0.0189	1.9988	-0.0012	0.0136	0.4774	-0.0226	0.0076
Faure	4.9909	-0.0091	0.0185	2.0016	0.0016	0.0136	0.4775	-0.0225	0.0083

1000 simulations are used with true parameter values $V_{max} = 5$, $K_m = 2$, $\sigma_b = 0.5$, MSE: mean square error, Q.R.S: quasi-random sequences.

The results are quite similar with the results obtained for 10 subjects. The fixed effects have negligible bias for all scenarios and in case of comparing among the quasi-random sequences, Halton sequences produce less bias and mean square error (MSE) (especially for V_{max}) than Sobol and Faure sequences. The only fluctuation occurred due to the estimation of the random effects parameter. For 20 subjects, all quasi-random sequences have comparatively smaller bias and MSE compare to the 10 subjects Cases in section (5.2). Also, all the quasi-random sequences produce almost similar results. For this Case II with two different intra-subject correlation coefficient ($\rho = 0.50$ & $\rho = 0.96$), the scatter plots of the parameter estimates for different quasi-random sequences in the quasi-Monte Carlo method of the Michaelis-Menten model (10) are also give similar result as in the Figure (3).

5.4 Simulation with Case III (Number of Subjects, $n = 5$)

In this subsection we decrease the number of subjects from 10 to 5 and again perform the similar simulation study by considering low and high intra-subject correlation coefficient ρ for checking the results and finding the effects of number of subjects in parameter estimation. Two scenarios are considered here for simulation study under the Case III (for same 5 subjects) depending on high ($\rho = .96$), and low ($\rho = .50$) intra-subjects correlation coefficient. These results are given below in the two separate Tables (6), and (7).

Table 6: Simulation results of parameter estimation in the Michaelis-Menten model, when $n = 5$, and $\rho = 0.96$

Q.R.S.	V_{max}			K_m			σ_b		
	Mean	Bias	MSE	Mean	Bias	MSE	Mean	Bias	MSE
Halton	5.0026	0.0026	0.0523	2.0010	0.0010	0.0021	0.4302	-0.0698	0.0300
Sobol	5.0033	0.0033	0.0525	2.0010	0.0010	0.0021	0.4213	-0.0787	0.0299
Faure	5.0036	0.0036	0.0527	2.0013	0.0013	0.0020	0.4214	-0.0786	0.0300

1000 simulations are used with true parameter values $V_{max} = 5$, $K_m = 2$, $\sigma_b = 0.5$, MSE: mean square error, Q.R.S: quasi-random sequences.

Table 7: Simulation results of parameter estimation in the Michaelis-Menten mode, when $n = 5$, and $\rho = 0.50$

Q.R.S.	V_{max}			K_m			σ_b		
	Mean	Bias	MSE	Mean	Bias	MSE	Mean	Bias	MSE
Halton	4.9900	-0.0100	0.0777	2.0020	0.0020	0.0565	0.4005	-0.0995	0.0404
Sobol	4.9895	-0.0105	0.0781	2.0020	0.0020	0.0565	0.4004	-0.0995	0.0403
Faure	4.9888	-0.0112	0.0778	2.0034	0.0034	0.0567	0.4015	-0.0985	0.0408

1000 simulations are used with true parameter values $V_{max} = 5$, $K_m = 2$, $\sigma_b = 0.5$, MSE: mean square error, Q.R.S: quasi random sequences.

From the above two tables, we see that the results are again quit similar with the results obtained previously. The fixed effects estimates are unbiased for all quasi-random sequences in the quasi-Monte Carlo method. In some situations the estimates of Halton and Sobol sequences have less bias (especially for V_{max}) and mean square error (MSE) than the Faure sequences. Again the only fluctuation occurred due to the estimation of variance component or random effects parameter. For 5 subjects, all quasi-random sequences in the quasi-Monte Carlo method have comparatively larger bias and MSE in comparison with the 10 and 20 subjects cases.

6 Conclusions

Nonlinear mixed effects models are very popular platform and a powerful tool for modeling and analysis of clustered or repeated measure data that arise in different areas of scientific investigation, such as, first order compartment model in Pharmacokinetics, Michaelis-Menten model in Biochemistry and Pharmacokinetics, logistic growth model in Forestry etc. The regression function of these models nonlinearly depends on fixed and random effects parameters. The popularity of nonlinear mixed effects model lies in its interpretability, parsimony and validity beyond the observed range of data. So, fitting of nonlinear mixed effects models take a great deal of interest.

In this study, the Michaelis-Menten model (10) was used with one random parameter as a specific example of nonlinear mixed effects model. Fitting of this model is an important issue because of the parameters involved within it, and the way of their estimation procedure. There are two types of parameters associated with such model. These are classified as regression parameters and variance-covariance parameters. Random effects are associated with variance-covariance parameters. Another important insight is that the score functions are nonlinear in parameters; so, no close form solution is available. There are many estimation procedures available for fitting nonlinear mixed effects models, but most of the techniques use approximation in different stages while fitting.

Here, in this paper, no approximation technique was used; instead of approximation based methods, this study considered Monte Carlo integration technique using different types of quasi-random sequences, which directly solves the intractable integrations in the likelihood function of nonlinear mixed effects models. For this method, Monte Carlo integration technique is performed using the points from the quasi-random sequences. Simulation studies were conducted for Michaelis-Menten model by considering one random effects parameter, that is the maximum rate achieved by the system. This simulation study considers some different cases with different combinations of number of subjects and intra-subject correlation coefficients. It was also considered and compared the performance of different types of quasi-random sequences (e.g. Halton, Sobol, Faure) in the parameter estimation of NLME model by quasi-Monte Carlo method.

The details plan and results of this simulation study for the Michaelis-Menten model (10) are given in Section 5. The results of these simulation studies were summarized with some interesting findings. In all cases and scenarios, the fixed effects estimates of parameters are almost similar and also unbiased for all types of quasi-random sequences in the estimation by quasi-Monte Carlo method. But, with deep insight and in comparison among the quasi-random sequences, Halton sequences provide less bias and mean square error (MSE) estimates in some situation (especially for V_{max} estimates) than the Sobol and Faure sequences. In some scenarios, Faure sequences provide the worst estimate in terms of bias and MSE because of the Box-Muller transformation used for generating Faure sequences. These results are very similar to the results in the paper of Morokoff and Caflisch [30]. They show that Halton sequences performs best for dimensions up to around 6; and Sobol sequences performs best for higher dimensions; and the Faure sequences is worst because the possible reason may be the Box-Muller transformation within it. But, in this study, the considered model (10) for parameter estimation was a low-dimensional non-linear mixed effects model with only one random effects parameter; so, the dimension of quasi-random sequences was one and there was no interest to show the effects of dimensional variety.

The main and significant difference occurs in the random effects parameter estimation. For any number of subjects, the fixed effects estimates are almost same, but the random effects parameter estimates are different. There is a significant effect of subjects in estimation of random effects parameter. As the number of subjects increases, the random effects parameter estimates tend to more close to the true value, that is, the biases and MSE are reduced for increasing the number of subjects. In this simulation study, the bias was least for 20 subjects cases and highest for 5 subjects cases for all types of quasi-random sequences in the quasi-Monte Carlo method for estimating parameters of nonlinear mixed effects models.

Acknowledgement

The author acknowledges and expresses his sincere gratitude to Md. Rashedul Hoque, Assistant Professor, Institute of Statistical Research and Training (ISRT), University of Dhaka for his scholastic assistance and valuable contribution during the simulation study and development of this article.

The author is grateful to the anonymous referee for a careful checking of the details and for helpful comments that improved this paper.

References

- [1] Davidian, M., & Giltinan, D. M. (2003). Nonlinear models for repeated measurement data: an overview and update. *Journal of agricultural, biological, and environmental statistics*, 8(4), 387.
- [2] Davidian, M., & Giltinan, D. M. (2017). Nonlinear models for repeated measurement data. *Routledge*.
- [3] Yang, Y., & Huang, S. (2014). Suitability of five cross validation methods for performance evaluation of nonlinear mixed-effects forest models—a case study. *Forestry: An International Journal of Forest Research*, 87(5), 654-662.
- [4] Olofsen, E., Dinges, D. F., & Van Dongen, H. (2004). Nonlinear mixed-effects modeling: individualization and prediction. *Aviation, space, and environmental medicine*, 75(3), A134-A140.
- [5] Pinheiro, J., & Bates, D. (2006). Mixed-effects models in S and S-PLUS. *Springer science & business media*.

- [6] Comets, E., Brendel, K., & Mentré, F. (2010). Model evaluation in nonlinear mixed effect models, with applications to pharmacokinetics. *Journal de la Société Française de Statistique*, 151(1), 106-127.
- [7] Hickman, D. A., Ignatowich, M. J., Caracotsios, M., Sheehan, J. D., & D'Ottaviano, F. (2019). Nonlinear mixed-effects models for kinetic parameter estimation with batch reactor data. *Chemical Engineering Journal*, 377, 119817.
- [8] Oddi, F. J., Miguez, F. E., Ghermandi, L., Bianchi, L. O., & Garibaldi, L. A. (2019). A nonlinear mixed-effects modeling approach for ecological data: Using temporal dynamics of vegetation moisture as an example. *Ecology and evolution*, 9(18), 10225-10240.
- [9] Pinheiro, J. C., & Bates, D. M. (1995). Approximations to the log-likelihood function in the nonlinear mixed-effects model. *Journal of computational and Graphical Statistics*, 4(1), 12-35.
- [10] Ke, C., & Wang, Y. (2001). Semiparametric nonlinear mixed-effects models and their applications. *Journal of the American Statistical Association*, 96(456), 1272-1298.
- [11] Hall, D. B., & Clutter, M. (2004). Multivariate multilevel nonlinear mixed effects models for timber yield predictions. *Biometrics*, 60(1), 16-24.
- [12] Lindstrom, M. J., & Bates, D. M. (1990). Nonlinear mixed effects models for repeated measures data. *Biometrics*, 673-687.
- [13] Pillai, G. C., Mentré, F., & Steimer, J. L. (2005). Non-linear mixed effects modeling—from methodology and software development to driving implementation in drug development science. *Journal of pharmacokinetics and pharmacodynamics*, 32(2), 161-183.
- [14] Vonesh, E. F., & Carter, R. L. (1992). Mixed-effects nonlinear regression for unbalanced repeated measures. *Biometrics*, 1-17.
- [15] Bonate, P. L. (2011). Nonlinear mixed effects models: theory. In *Pharmacokinetic-pharmacodynamic modeling and simulation* (pp. 233-301). Springer, Boston, MA.
- [16] Meza, C., Osorio, F., & De la Cruz, R. (2012). Estimation in nonlinear mixed-effects models using heavy-tailed distributions. *Statistics and Computing*, 22(1), 121-139.
- [17] Fu Liyong, C. A. F., Zhang Huiru, C. A. F., & Li Chunming, C. A. F. (2013). Analysis of nonlinear mixed effects model parameter estimation methods. *Scientia Silvae Sinicae*.
- [18] Davidian, M., & Gallant, A. R. (1993). The nonlinear mixed effects model with a smooth random effects density. *Biometrika*, 80(3), 475-488.
- [19] Comets, E., Lavenu, A., & Lavielle, M. (2017). Parameter estimation in nonlinear mixed effect models using saemix, an R implementation of the SAEM algorithm. *Journal of Statistical Software*, 80, 1-41.
- [20] Wang, J. (2007). EM algorithms for nonlinear mixed effects models. *Computational statistics & data analysis*, 51(6), 3244-3256.
- [21] Plan, E. L., Maloney, A., Mentré, F., Karlsson, M. O., & Bertrand, J. (2012). Performance comparison of various maximum likelihood nonlinear mixed-effects estimation methods for dose-response models. *The AAPS journal*, 14(3), 420-432.
- [22] Girard, P., & Mentré, F. (2005, June). A comparison of estimation methods in nonlinear mixed effects models using a blind analysis. *In Annual Meeting of the Population Approach Group in Europe*, Pamplona, Spain.
- [23] Wolfinger, R. D., & Lin, X. (1997). Two Taylor-series approximation methods for nonlinear mixed models. *Computational Statistics & Data Analysis*, 25(4), 465-490.
- [24] Schumacher, F. L., Dey, D. K., & Lachos, V. H. (2021). Approximate inferences for nonlinear mixed effects models with scale mixtures of skew-normal distributions. *Journal of Statistical Theory and Practice*, 15(3), 1-26.
- [25] Smith, N., & Blozis, S. (2015). Options in estimating nonlinear mixed models: quadrature points and approximation methods. *Western Users of SAS Software*.
- [26] Nie, L. (2002). Laplace approximation in nonlinear mixed-effect models. *University of Illinois at Chicago*.
- [27] Piersol, L. J. (2000). Fitting nonlinear mixed effect models by Laplace approximation. University of California, Los Angeles.
- [28] Creal, D. (2012). A survey of sequential Monte Carlo methods for economics and finance. *Econometric reviews*, 31(3), 245-296.
- [29] Dimov, I. T. (2008). Monte Carlo methods for applied scientists. *World Scientific*.
- [30] Morokoff, W. J., & Caflisch, R. E. (1995). Quasi-monte carlo integration. *Journal of computational physics*, 122(2), 218-230.
- [31] Hoque, M. R., & Latif, A. M. (2013). Sobol's Sequence Based Method for Fitting Nonlinear Mixed Effects Model: A Comparative View. *Baboulin, Marc*.
- [32] Asmussen, S., & Glynn, P. W. (2007). Stochastic simulation: algorithms and analysis (Vol. 57). *Springer Science & Business Media*.
- [33] Smyth, G. K. (2002). Nonlinear regression. *Encyclopedia of environmetrics*, 3, 1405-1411.
- [34] Ritz, C., & Streibig, J. C. (Eds.). (2008). Nonlinear regression with R. *New York, NY: Springer New York*.
- [35] Krumpolc, T., Trahan, D. W., Hickman, D. A., & Biegler, L. T. (2022). Kinetic parameter estimation with nonlinear mixed-effects models. *Chemical Engineering Journal*, 136319.
- [36] Lai, T. L., & Shih, M. C. (2003). Nonparametric estimation in nonlinear mixed effects models. *Biometrika*, 90(1), 1-13.
- [37] Menten, L., & Michaelis, M. I. (1913). Die kinetik der invertinwirkung. *Biochem Z*, 49, 333-369.
- [38] Ueberhuber, C. W. (2012). Numerical computation 1: methods, software, and analysis. *Springer Science & Business Media*.
- [39] Legrand, X., Xémard, A., Fleury, G., Auriol, P., & Nucci, C. A. (2008). A quasi-Monte Carlo integration method applied to the computation of the Pollaczek integral. *IEEE Transactions on Power Delivery*, 23(3), 1527-1534.