



An optimal tradeoff between explorations and repetitions in global sensitivity analysis for stochastic computer models

Gildas Mazo

► To cite this version:

Gildas Mazo. An optimal tradeoff between explorations and repetitions in global sensitivity analysis for stochastic computer models. 2019. hal-02113448v2

HAL Id: hal-02113448

<https://hal.science/hal-02113448v2>

Preprint submitted on 8 Jul 2019 (v2), last revised 7 Jun 2021 (v7)

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

An optimal tradeoff between explorations and repetitions in global sensitivity analysis for stochastic computer models

Gildas Mazo

MaIAGE, INRA, Université Paris-Saclay, 78350, Jouy-en-Josas,
France

Abstract

Global sensitivity analysis often accompanies computer modeling to understand what are the important factors of a model of interest. In particular, Sobol indices, naturally estimated by Monte-Carlo methods, permit to quantify the contribution of the inputs to the variability of the output. However, stochastic computer models raise difficulties. There is no unique definition of Sobol indices and their estimation is difficult because a good balance between repetitions of the computer code and explorations of the input space must be found. The problem of finding an optimal tradeoff between explorations and repetitions is addressed. Two Sobol indices are considered, their estimators constructed and their asymptotic properties established. To find an optimal tradeoff between repetitions and explorations, a tractable error criterion, which is small when the inputs of the model are ranked correctly, is built and minimized under a fixed computing budget. Then, Sobol estimates based on the balance found beforehand are produced. Convergence rates are given and it is shown that this method is asymptotically oracle. Numerical tests and a sensitivity analysis of a Susceptible-Infectious-Recovered (SIR) model are performed.

Keywords: Explorations, repetitions, Sobol, estimation, sensitivity, stochastic, random, model.

Contents

1	Introduction	2
2	Stochastic sensitivity analysis	3
2.1	Definition of the sensitivity indices	3
2.2	Estimation of the sensitivity indices	5
3	Balancing explorations and repetitions	7
3.1	Defining an optimal number of repetitions	8
3.2	Estimation	9

4	A two-stage procedure to estimate sensitivity indices	10
4.1	The procedure	11
4.2	Theoretical analysis	11
5	Numerical tests	13
5.1	Linear model	14
5.1.1	High noise context	14
5.1.2	Low noise context	15
5.2	The Ishigami model	18
6	Sensitivity analysis for a SIR model	18
7	Discussion	22
A	Proofs	25
B	A unified treatment of the asymptotics	31
C	Explicit moment calculations	34
D	Calculations for the linear model	38

1 Introduction

Sensitivity analysis often accompanies computer modeling to understand what are the important factors of a model of interest [17, 18]. In particular, Sobol indices [20, 21] permit to quantify the contribution of the inputs to the variability of the output. The Sobol index S_j is defined as

$$(1) \quad S_j = \frac{\text{Var } \mathbb{E}(f(X)|X_j)}{\text{Var } f(X)}, \quad j = 1, \dots, p,$$

where $Y = f(X)$ is the output of the computer model f evaluated at the input $X = (X_1, \dots, X_p)$. The larger S_j , the more X_j is important in the following sense: if $X_j = x_j$ were fixed, $\text{Var } Y$ would be reduced by at least $S_j \times 100\%$. As a consequence, the Sobol indices satisfy $S_1 + \dots + S_p \leq 1$, with equality in the absence of interaction effects, which we shall not account for in this paper.

The estimation of Sobol indices is naturally performed by Monte-Carlo methods [7, 15, 20, 21], which permit to build estimators with statistical guarantees [5, 12]. Sobol indices for multivariate, functional outputs [4, 13] or functional inputs [10] have been proposed as well.

Computer models employed to simulate physical systems/natural phenomena are increasingly stochastic. That is, two runs of the computer with the same input may lead to two different outputs. Examples can be found in epidemiology [1, 3, 16, 19] or ecology [22].

It is still unclear how sensitivity analysis should be performed when the models are stochastic. First, there is no unique definition of Sobol indices [6]. Second, it is unclear how to account for noise in the inference. Monte-Carlo sampling with repetitions is natural, but what is a good balance between the number of repetitions of the model and the number of explorations of the input space [22]? Having efficient estimators would permit to achieve the same level of

precision but with less computations, an important practical issue. An approach based on meta-models has been proposed [14], but it is difficult to control the induced bias and the construction of the meta-model itself can be challenging.

The problem of finding an optimal balance between explorations and repetitions to estimate Sobol indices in stochastic computer models is addressed. Two definitions of Sobol indices are given. Their estimators, based on Monte-Carlo sampling with repetitions, are built and their asymptotic properties are established. To find an optimal tradeoff between repetitions and explorations, a tractable error criterion, which is small when the inputs of the model are ranked correctly, is built and minimized under a fixed computing budget. A two-stage procedure is given for estimating the Sobol indices, the first stage serving to find the optimal number of repetitions that should be used in the Monte-Carlo experiment of the second stage.

This paper is organized as follows. The sensitivity indices and their estimators are defined in Section 2. The optimal number of repetitions is defined in Section 3, where an estimator is also given. In Section 4, a two-stage procedure is proposed that permit to efficiently estimate the sensitivity indices by exploiting the results of Section 2 and Section 3. Numerical experiments are provided in Section 5 to test the theory. In Section 6, a sensitivity analysis of a SIR model is performed to illustrate the proposed method. A discussion closes the paper. The proofs are given in Appendix A.

2 Stochastic sensitivity analysis

2.1 Definition of the sensitivity indices

In the case of a stochastic computer model, the output is

$$Y = f(X, Z),$$

where Z is an unobserved random element that represents the “noise” of the model. That is to say, even if $X = x$ is fixed, the output, distributed as $f(x, Z)$, exhibits a residual variability due to the randomness of Z . In all this paper, we assume that X and Z are independent. Note that f does *not* represent the computer model, as seen by the user. For the user, the computer model is not even a mapping, since to runs of the computer at the same input can lead to two different outputs. Here, f is any mapping that, together with X and Z , produces the output Y . The existence of Z and f such that $Y = f(X, Z)$ is assumed.

Two kinds of Sobol indices can be built.

Definition 1. *The Sobol indices of the first kind are defined as*

$$S'_j = \frac{\text{Var E}(f(X, Z)|X_j)}{\text{Var } f(X, Z)}$$

for $j = 1, \dots, p$.

Definition 2. *The Sobol indices of the second kind are defined as*

$$S''_j = \frac{\text{Var E}(E[f(X, Z)|X]|X_j)}{\text{Var E}[f(X, Z)|X]}$$

for $j = 1, \dots, p$.

The first index is built by substituting $f(X, Z)$ for $f(X)$ in (1). The second by substituting $E[f(X, Z)|X]$. Thus, the index of the first kind is a direct application of (1): one does *as if* Z were another input, even though it is not observable. The index of the second kind is also an application of (1) but to the function $x \mapsto E[f(x, Z)|X = x]$ with the noise smoothed out.

For estimation purposes, it is convenient to rewrite the indices as

$$(2) \quad S'_j = \frac{E E[f(X, Z)|X] E[f(\tilde{X}_{-j}, Z)|\tilde{X}_{-j}] - (E E[f(X, Z)|X])^2}{E E[f(X, Z)^2|X] - (E E[f(X, Z)|X])^2}$$

and

$$(3) \quad S''_j = \frac{E E[f(X, Z)|X] E[f(\tilde{X}_{-j}, Z)|\tilde{X}_{-j}] - (E E[f(X, Z)|X])^2}{E E[f(X, Z)|X]^2 - (E E[f(X, Z)|X])^2},$$

where $\tilde{X} = (\tilde{X}_1, \dots, \tilde{X}_p)$ is an independent copy of X and

$$\tilde{X}_{-j} = (\tilde{X}_1, \dots, \tilde{X}_{j-1}, X_j, \tilde{X}_{j+1}, \dots, \tilde{X}_p),$$

for $j = 1, \dots, p$. Note that S'_j and S''_j differ only by the lower left term. In particular, the upper left term is the same in both formula and is the only term that depends on j , and hence the only term that permits to discriminate between any two indices of the same kind. For this reason, it is called the discriminator and is denoted by D_j . Notice that $S'_j \leq S''_j$.

Example 1. Let $Y = aX_1 + cX_2h(Z)$, where X_1, X_2, Z are independent standard normal variables, a, c are real coefficients and h is a function such that $E h(Z) = 0$. Then

$$S'_1 = \frac{a^2}{a^2 + c^2 E h(Z)^2}, S'_2 = 0, S''_1 = 1 \text{ and } S''_2 = 0.$$

Example 1 illustrates that the two kinds of sensitivity indices measure different aspects of the sensitivity of a computer model. While the sensitivity indices of the first kind depend on the coefficients of the model, those of the second kind do not.

Example 2. Let $f(X_1, X_2, Z) = \sin(X_1) + a \sin(X_2)^2 + bZ^4 \sin(X_1)$, where X_1, X_2, Z are independent uniform random variables on $(-\pi, \pi)$ and a, b are real coefficients. The Sobol indices of the first kind are given by

$$S'_1 = \frac{b\pi^4/5 + b^2\pi^8/50 + 1/2}{a^2/8 + b\pi^4/5 + b^2\pi^8/18 + 1/2}$$

and

$$S'_2 = \frac{a^2/8}{a^2/8 + b\pi^4/5 + b^2\pi^8/18 + 1/2}.$$

The Sobol indices of the second kind are given by

$$S''_1 = \frac{b\pi^4/5 + b^2\pi^8/50 + 1/2}{(1 + b\pi^4/5)^2/2 + a^2/8}$$

and

$$S''_2 = \frac{a^2/8}{(1 + b\pi^4/5)^2/2 + a^2/8}.$$

The model in Example 2 with $a = 7$ and $b = 0.1$ was used to test sensitivity analysis methods based on meta-models [14]. When $Z = X_3$ is a controllable input, this model is a standard benchmark for deterministic models [7, 11]. Although, in the examples above, Z is a scalar, it is easy to build other examples in which Z is a random vector.

2.2 Estimation of the sensitivity indices

Estimation is based on Monte-Carlo sampling. The sampling scheme is given in Algorithm 1. The input space is explored n times and, for each exploration, the computer is run m times to smooth out the noise. Thus, the total number of calls to the computer is proportional to mn . Denote by $T = mn(p + 1)$ the total number of calls.

Algorithm 1 Generate a Monte-Carlo sample

```

for  $i = 1$  to  $n$  do
  draw two independent copies  $X^{(i)}, \tilde{X}^{(i)}$ 
  for  $j = 0, 1, \dots, p$  do
    for  $k = 1$  to  $m$  do
      run the computer model at  $\tilde{X}_{-j}^{(i)}$  to get an output  $Y_j^{(i,k)}$ 
    end for
  end for
end for

```

The data generated by the algorithm are

$$(Y_j^{(i,k)}, \tilde{X}_{-j}^{(i)}),$$

for $j = 0, 1, \dots, p$, $i = 1, \dots, n$ and $k = 1, \dots, m$, with the convention $\tilde{X}_{-0}^{(i)} = X^{(i)}$. By assumption, there are independent random elements $(Z_j^{(i,k)})$ such that

$$(4) \quad Y_j^{(i,k)} = f(\tilde{X}_{-j}^{(i)}, Z_j^{(i,k)}).$$

The estimators of the sensitivity indices are built by substituting empirical averages for expectations in (2) and (3), that is,

$$(5) \quad \hat{S}'_{j;n,m} = \frac{\frac{1}{n} \sum_{i=1}^n \frac{1}{m} \sum_{k=1}^m Y_0^{(i,k)} \frac{1}{m} \sum_{k'=1}^m Y_j^{(i,k')} - \left(\frac{1}{n} \sum_{i=1}^n \frac{1}{m} \sum_{k=1}^m Y_0^{(i,k)} \right)^2}{\frac{1}{n} \sum_{i=1}^n \frac{1}{m} \sum_{k=1}^m Y_0^{(i,k)2} - \left(\frac{1}{n} \sum_{i=1}^n \frac{1}{m} \sum_{k=1}^m Y_0^{(i,k)} \right)^2}$$

and

$$(6) \quad \hat{S}''_{j;n,m} = \frac{\frac{1}{n} \sum_{i=1}^n \frac{1}{m} \sum_{k=1}^m Y_0^{(i,k)} \frac{1}{m} \sum_{k'=1}^m Y_j^{(i,k')} - \left(\frac{1}{n} \sum_{i=1}^n \frac{1}{m} \sum_{k=1}^m Y_0^{(i,k)} \right)^2}{\frac{1}{n} \sum_{i=1}^n \left(\frac{1}{m} \sum_{k=1}^m Y_0^{(i,k)} \right)^2 - \left(\frac{1}{n} \sum_{i=1}^n \frac{1}{m} \sum_{k=1}^m Y_0^{(i,k)} \right)^2}.$$

To our knowledge (personal communication), when faced with stochastic computer models, practitioners tend to use softwares for deterministic sensitivity analysis in which an average over repetitions is given to the program as a

substitute for the value of the output. Thus, the second estimator is used in practice, albeit implicitly. The first estimator, to the best of our knowledge, was not formally defined. The second estimator appeared in [8, 9], where it was studied only in the case $m = n$ (to the best of our understanding).

In (5) and (6), take $m = 1$ and assume $Y_j^{(i,1)} = f(\tilde{X}_{-j}^{(i)})$ instead of (4). Then the estimators reduce to those of Sobol [20, 21] in the deterministic case. In the later case, these estimators are sometimes called pick-freeze estimators [5].

In the rest of this paper, we assume that, for all x and all z ,

$$(7) \quad f(x, z) \leq F(x)$$

for some F with $\mathbb{E} F(X)^8 < \infty$. The existence of the bound F to control the stochastic part of the model is needed to apply Lindeberg-Feller's central limit theorem. The condition $\mathbb{E} F(X)^8 < \infty$ is the weakest moment condition needed to be able to apply classical central limit theorems to all of the estimators in this paper.

These conditions are verified for the model in Example 2. Indeed we have

$$f(X_1, X_2, Z) \leq \sin(X_1)(1 + |b|\pi^4) + a \sin(X_2)^2.$$

For the model in Example 1, it suffices to take

$$h(Z) = \begin{cases} -t & \text{if } Z \leq -t, \\ Z & \text{if } -t \leq Z \leq t, \\ t & \text{otherwise} \end{cases}$$

for some threshold $t > 0$.

We now establish asymptotic properties of the sensitivity estimators as the number of explorations goes to infinity. The number of repetitions can be fixed or go to infinity as well. In the later, $m = m_n$ is assumed to be a sequence growing with n . Denote by $\mathbf{S}'$ (resp. $\mathbf{S}''$) the (column) vector with coordinates S'_j (resp. S''_j), $j = 1, \dots, p$, and denote by $\hat{\mathbf{S}}'_{n,m}$ (resp. $\hat{\mathbf{S}}''_{n,m}$) the vector with coordinates $\hat{S}'_{j;n,m}$ (resp. $\hat{S}''_{j;n,m}$).

Theorem 1. *Assume (7) holds and let $n \rightarrow \infty$. Then*

$$\sqrt{n} \left(\hat{\mathbf{S}}'_{n,m} - \mathbf{S}' \left[1 - \frac{\mathbb{E} \text{Var}[f(X, Z)|X]}{\mathbb{E} \text{Var}[f(X, Z)|X] + m \text{Var} \mathbb{E}[f(X, Z)|X]} \right] \right) \xrightarrow{d} N(0, \Xi_m),$$

for some nonnegative matrix Ξ_m of size $2p \times 2p$. If $m = m_n \rightarrow \infty$ as $n \rightarrow \infty$, then, elementwise, $\Xi_m \rightarrow \Xi$ for some Ξ . Moreover, the convergence in distribution above still holds with Ξ in place of Ξ_m .

The result of Theorem 1 suggests that, when the number of explorations is large enough,

$$\hat{\mathbf{S}}'_{n,m} \overset{d}{\approx} N(\mathbf{S}', \frac{1}{n} \Xi_{1m})$$

and

$$\hat{\mathbf{S}}''_{n,m} \overset{d}{\approx} N \left(\mathbf{S}'' \left[1 - \frac{\mathbb{E} \text{Var}[f(X, Z)|X]}{\mathbb{E} \text{Var}[f(X, Z)|X] + m \text{Var} \mathbb{E}[f(X, Z)|X]} \right], \frac{1}{n} \Xi_{2m} \right)$$

for some variance-covariance matrices Ξ_{1m}, Ξ_{2m} of size $p \times p$ and for any m . Since Ξ_m has a limit, these approximations hold for m large as well. This permits to draw inferences about the Sobol indices for any number of repetitions.

When the number of repetitions is much smaller than the number of explorations, the sensitivity estimators of the second kind underestimate the corresponding Sobol indices. Fortunately, the bias is explicit and can be estimated in practice. Note that the bias is proportional to $\mathbb{E} \text{Var}[f(X, Z)|X]$, which is zero whenever f actually does not depend on Z . This term is a noise term: it is expected to be large whenever the computer model is highly stochastic. When m grows, the bias diminishes but it could be that the sensitivity estimator is arbitrarily tightly concentrated around the wrong target. This phenomenon is avoided by choosing a number of repetitions much larger than the square root of the number of explorations, as stated in Corollary 1.

Corollary 1. *Let $\sqrt{n}/m \rightarrow 0$. Then, under the assumptions of Theorem 1,*

$$\sqrt{n} \begin{pmatrix} \hat{\mathbf{S}}'_{n,m} - \mathbf{S}' \\ \hat{\mathbf{S}}''_{n,m} - \mathbf{S}'' \end{pmatrix} \xrightarrow{d} N(0, \Xi).$$

3 Balancing explorations and repetitions

Clearly, to get the best estimators, one should do as many explorations and repetitions as possible. However, in practice, each call to the computer model is costly and the budget is limited. Denote the total available budget by T and remember from Section 2 that $T = mn(p+1)$. Which, among all couples (m, n) that satisfy $T = mn(p+1)$, yields the best performance? The criterion often used to optimize statistical methods, the mean squared error, is intractable for Sobol estimators.

We propose to consider a more convenient criterion, given by

$$\text{MRE} = \mathbb{E} \sum_{j=1}^p |\hat{R}_{j;n,m} - R_j|,$$

where R_j is the rank of D_j among $D_1, \dots, D_p$, that is, $R_j = \sum_{i=1}^p \mathbf{1}(D_i \leq D_j)$, and $\hat{R}_{j;n,m}$ is an estimator of R_j . Recall that $D_1, \dots, D_p$, the upper-left terms in (2) and (3), determine the ranks of the sensitivity indices and that the ranks of the sensitivity indices of the first and of the second kind are the same. Thus, the MRE permits to find a unique solution to the balance problem for both kinds of sensitivity indices. The MRE is an error which is small when one succeeds in ranking the inputs of the computer model from the most to the least important. That is, one seeks the input which, if fixed, would lead to the greatest reduction of the output variance; then the second, etc. This is called *Factors Prioritisation* in [18, p. 52].

The MRE is more tractable than the MSE because it depends only on the discriminators, but it is still difficult to minimize it directly. Therefore, it is a bound of the MRE that is used to find an optimal balance between explorations and repetitions.

3.1 Defining an optimal number of repetitions

In what follows, the vector $(X, \tilde{X})$ of size $2p$ is denoted by $\mathbf{X}$. Denote by $\hat{D}_{j;n,m}$ the upper-left term in (5) and (6), the estimator of D_j . Let Y_j be a shorthand for $Y_j^{(1,1)} = f(\tilde{X}_{-j}^{(1)}, Z_j^{(1,1)})$, $j = 0, 1, \dots, p$ and recall that $\tilde{X}_{-0}^{(1)} = X^{(1)}$.

Proposition 1. *Let $\hat{R}_{j;n,m} = \sum_{i=1}^p \mathbf{1}(\hat{D}_{i;n,m} \leq \hat{D}_{j;n,m})$. Then*

$$\begin{aligned} \mathbb{E} \sum_{j=1}^p |\hat{R}_{j;n,m} - R_j| &\leq \frac{2(p-1)(p+1)^2}{\min_{j < j'} (|D_j - D_{j'}|^2) T} \\ &\quad \times \sum_{j=1}^p (m \operatorname{Var} \mathbb{E}[Y_0 Y_j | \mathbf{X}] \\ &\quad + \mathbb{E}(\operatorname{Var}[Y_0 Y_j | \mathbf{X}] - \operatorname{Var}[Y_0 | \mathbf{X}] \operatorname{Var}[Y_j | \mathbf{X}]) \\ &\quad + \frac{1}{m} \mathbb{E} \operatorname{Var}[Y_0 | \mathbf{X}] \operatorname{Var}[Y_j | \mathbf{X}]). \end{aligned}$$

This bound provides us with some insights. It is small when the Sobol indices are well separated and increases as the cube of the number of inputs. Denote $\sum_j \operatorname{Var} \mathbb{E}[Y_0 Y_j | \mathbf{X}]$ by ζ_1 , $\sum_j \mathbb{E}(\operatorname{Var}[Y_0 Y_j | \mathbf{X}] - \operatorname{Var}[Y_0 | \mathbf{X}] \operatorname{Var}[Y_j | \mathbf{X}])$ by ζ_2 and $\sum_j \mathbb{E} \operatorname{Var}[Y_0 | \mathbf{X}] \operatorname{Var}[Y_j | \mathbf{X}]$ by ζ_3 . Then $\zeta_l \geq 0$ for $l = 1, 2, 3$ and the bound is proportional to $n^{-1}(\zeta_1 + m^{-1}\zeta_2 + m^{-2}\zeta_3)$, which decreases as n increases or m increases. The function $T^{-1}(m\zeta_1 + \zeta_2 + m^{-1}\zeta_3)$ is a convex function of m and, forgetting at the moment that m is an integer, it is minimized at

$$(8) \quad m^* := \sqrt{\frac{\sum_{j=1}^p \mathbb{E} \operatorname{Var}[Y_0 | \mathbf{X}] \operatorname{Var}[Y_j | \mathbf{X}]}{\sum_{j=1}^p \operatorname{Var} \mathbb{E}[Y_0 Y_j | \mathbf{X}]}.$$

The number m^* is called the *optimal continuous number of repetitions* and can be interpreted as a noise-to-signal ratio for ranking the Sobol indices. Indeed, recall that a variance can be decomposed into a variance of a conditional expectation and an expectation of a conditional variance, where the last represents a noise term and the former a signal term [18, p. 12]. The numerator in (8) is the noise term: if the computer model is deterministic, then the function $f(X, Z)$ does not in fact depend on Z , and the numerator is zero. The denominator is the signal term.

The *optimal number of repetitions* is naturally defined as the integer $m^\dagger$ that minimizes the upper bound in Proposition 1 over all compatible integers, that is, over all integers m such that $T = mn(p+1)$. For instance, for $T = 300$ and $p = 2$, the set of compatible integers is $\{1, 2, 4, 5, 10, 20, 25, 50\}$. By convexity of the bound, the optimal number of repetitions is either the greatest compatible integer less than or equal to m^* , denoted by $\lfloor m^* \rfloor$, or the smallest compatible integer greater than or equal to m^* , denoted by $\lceil m^* \rceil$.

Proposition 2. *The optimal number of repetitions $m^\dagger$ is given according to the following three cases.*

- (i) If $m^* \leq 1$ then $m^\dagger = 1$.
- (ii) If $m^* \geq T/(p+1)$ then $m^\dagger = T/(p+1)$.

(iii) If $1 < m^* < T/(p+1)$ then

$$m^\dagger = \begin{cases} \lceil m^\top \rceil^* & \text{if } \lfloor m^\top \rfloor^* \lceil m^\top \rceil^* \leq m^{*2} \\ \lfloor m^\top \rfloor^* & \text{if } \lfloor m^\top \rfloor^* \lceil m^\top \rceil^* \geq m^{*2}. \end{cases}$$

Moreover, if $\lfloor m^\top \rfloor^* = \lceil m^\top \rceil^* = m^*$, then $m^* = m^\dagger$.

Provided that the Monte-Carlo experiment is large enough so that $m^* < T/(p+1)$, the optimal continuous number of repetitions m^* and hence the optimal number of repetitions $m^\dagger$, do *not* depend on the size of the experiment. For instance, if it happens that $m^* < 50$, then the optimal number of repetitions will be the same whether the size of the experiment is $T = 300$ or $T = 3000000$.

In the sequel, by an abuse of language, both m^* and $m^\dagger$ shall be referred to as the optimal number of repetitions.

3.2 Estimation

The goal of this section is to build an estimator of the optimal number of repetitions (8) based on the same outputs as those produced in Algorithm 1. In view of Proposition 2, this will yield immediately an estimator for the optimal number of repetitions.

Denote $\text{Var } E[Y_0 Y_j | \mathbf{X}]$ by $\zeta_{1,j}$, $E(\text{Var}[Y_0 Y_j | \mathbf{X}] - \text{Var}[Y_0 | \mathbf{X}] \text{Var}[Y_j | \mathbf{X}])$ by $\zeta_{2,j}$ and $E \text{Var}[Y_0 | \mathbf{X}] \text{Var}[Y_j | \mathbf{X}]$ by $\zeta_{3,j}$, $j = 1, \dots, p$. Then $m^* = \sqrt{\sum_{j=1}^p \zeta_{3,j} / \sum_{j=1}^p \zeta_{1,j}}$ and natural estimators for $\zeta_{1,j}$ and $\zeta_{3,j}$ can be built using the same principles as those that were used to build the sensitivity estimators. Thus, let

$$\begin{aligned} \hat{\zeta}_{3,j;n,m} &= \\ (9) \quad & \frac{1}{n} \sum_{i=1}^n \frac{1}{m} \sum_{k_1=1}^m f(X^{(i)}, Z_0^{(i,k_1)})^2 \frac{1}{m} \sum_{k_2=1}^m f(\tilde{X}_{-j}^{(i)}, Z_j^{(i,k_2)})^2 \\ (10) \quad & + \frac{1}{n} \sum_{i=1}^n \left(\frac{1}{m} \sum_{k_1=1}^m f(X^{(i)}, Z_0^{(i,k_1)}) \right)^2 \left(\frac{1}{m} \sum_{k_2=1}^m f(\tilde{X}_{-j}^{(i)}, Z_j^{(i,k_2)}) \right)^2 \\ (11) \quad & - \frac{1}{n} \sum_{i=1}^n \left(\frac{1}{m} \sum_{k_1=1}^m f(X^{(i)}, Z_0^{(i,k_1)}) \right)^2 \frac{1}{m} \sum_{k_2=1}^m f(\tilde{X}_{-j}^{(i)}, Z_j^{(i,k_2)})^2 \\ (12) \quad & - \frac{1}{n} \sum_{i=1}^n \frac{1}{m} \sum_{k_1=1}^m f(X^{(i)}, Z_0^{(i,k_1)})^2 \left(\frac{1}{m} \sum_{k_2=1}^m f(\tilde{X}_{-j}^{(i)}, Z_j^{(i,k_2)}) \right)^2, \end{aligned}$$

and

$$\begin{aligned} \hat{\zeta}_{1,j;n,m} &= \\ (13) \quad & \frac{1}{n} \sum_{i=1}^n \left(\frac{1}{m} \sum_{k=1}^m f(X^{(i)}, Z_0^{(i,k)}) f(\tilde{X}_{-j}^{(i)}, Z_j^{(i,k)}) \right)^2 \\ (14) \quad & - \left(\frac{1}{n} \sum_{i=1}^n \frac{1}{m} \sum_{k=1}^m f(X^{(i)}, Z_0^{(i,k)}) f(\tilde{X}_{-j}^{(i)}, Z_j^{(i,k)}) \right)^2 \end{aligned}$$

so that m^* is estimated by

$$(15) \quad \hat{m}_{n,m}^* := \sqrt{\sum_{j=1}^p \hat{\zeta}_{3,j;n,m} / \sum_{j=1}^p \hat{\zeta}_{1,j;n,m}}.$$

It is easy to check that $\hat{\zeta}_{1,j;n,m} \geq 0$ and $\hat{\zeta}_{3,j;n,m} \geq 0$ so that $\hat{m}_{n,m}^* \geq 0$. Also, notice that for $m = 1$, we have $\hat{\zeta}_{3,j;n,1} = 0$, so that $\hat{m}_{n,1}^* = 0$. Asymptotic properties are given in Theorem 2.

Theorem 2. *Assume (7) holds and let $n \rightarrow \infty$ and $m \rightarrow \infty$. Then*

$$\sqrt{n} \left(\hat{m}_{n,m}^* - \left[m^* + \frac{C + o(1)}{m} \right] \right) \rightarrow N(0, \sigma^2),$$

for some constant C and variance σ^2 .

The term $o(1)$ is a sequence of constants indexed by m that goes to zero as $m \rightarrow \infty$. The bias is subjected to the same phenomenon observed for the sensitivity estimators of the second kind: as shown in Corollary 2, it can be annihilated when the number of repetitions grow faster than the square root of the number of explorations.

Corollary 2. *Let $n \rightarrow \infty$ and $m \rightarrow \infty$ such that $\sqrt{n}/m \rightarrow 0$. Then, under the assumptions of Theorem 2, $\sqrt{n}(\hat{m}_{n,m}^* - m^*) \rightarrow N(0, \sigma^2)$.*

Theorem 2 is useful to study the statistical performance of the procedure proposed in Section 4.

4 A two-stage procedure to estimate sensitivity indices

The goal is to estimate the sensitivity indices of Section 2 by using the results of Section 3. It is assumed that the total budget available is T . To estimate the sensitivity indices, the following procedure is natural.

1. Generate a Monte-Carlo sample to get an estimate of the optimal number of repetitions.
2. Use that estimate to generate another Monte-Carlo sample with which the sensitivity estimators are built.

Although simple, the above procedure raises some questions. The sum of the sizes of the two Monte-Carlo experiments must be T . Thus, the size of the second Monte-Carlo sample is smaller than it could have been if no share of the budget was spent to estimate the optimal number of repetitions. To what extent is the performance affected? How to calibrate the procedure? How should the budget be split? These questions are addressed.

Algorithm 2 Estimate the sensitivity indices in a two-stage procedure

Stage 1. Generate a Monte-Carlo sample with K runs of the computer model, n_0 explorations and m_0 repetitions to get an estimate $\hat{m}_{K,m_0}^\dagger$ of $m^\dagger$. If $K = 0$, simply return m_0 .

Stage 2. Generate a Monte-Carlo sample with $T - K$ runs and $\hat{m}_{K,m_0}^\dagger$ repetitions to estimate the sensitivity indices.

4.1 The procedure

Assume that only $T = mn(p + 1)$ runs of the computer model are allowed. Choose integers K, m_0, n_0 such that $m_0 n_0 (p + 1) = K < T$. Algorithm 2 gives the details of the procedure.

The integer K is the share of the total budget T dedicated to the estimation of the optimal number of repetitions, which is performed in the first stage. The estimator is (15) with $m = m_0$ and $n = n_0$. In the second stage, the sensitivity indices are estimated with the remaining budget $T - K$. The estimators are (5) and (6) with $m = \hat{m}_{K,m_0}^\dagger$ and n the integer satisfying $T - K = \hat{m}_{K,m_0}^\dagger n(p + 1)$. (If this equation has no solution, take m near $\hat{m}_{K,m_0}^\dagger$ such that the equation has a solution.) We allow for the case $K = 0$. That is, the sensitivity indices are estimated directly with $m = m_0$ based on the whole budget T . If, moreover, it happens accidentally that $m_0 = m^\dagger$, then the best performance is to be expected.

4.2 Theoretical analysis

As in Section 3, the statistical performance of the two-stage procedure is assessed with respect to the MRE. Remember that $\mathbf{X}^{(1)} = (X^{(1)}, \tilde{X}^{(1,1)})$, $Y_0^{(1,1)} = f(X^{(1)}, Z_0^{(1,1)})$ and $Y_j^{(1,1)} = f(\tilde{X}_{-j}^{(1)}, Z_j^{(1,1)})$, $j = 1, \dots, p$, and the superscripts are dropped for convenience. Put

$$\begin{aligned} v(m) &= \frac{2(p-1)(p+1)^2}{\min_{j < j'} (|D_j - D_{j'}|^2)} \\ &\quad \times \sum_{j=1}^p (m \operatorname{Var} \mathbb{E}[Y_0 Y_j | \mathbf{X}] \\ &\quad + \mathbb{E}(\operatorname{Var}[Y_0 Y_j | \mathbf{X}] - \operatorname{Var}[Y_0 | \mathbf{X}] \operatorname{Var}[Y_j | \mathbf{X}]) \\ &\quad + \frac{1}{m} \mathbb{E} \operatorname{Var}[Y_0 | \mathbf{X}] \operatorname{Var}[Y_j | \mathbf{X}]) \end{aligned}$$

so that the MRE in Proposition 1 is $v(m)/T$. The function v is convex and, over all compatible integers m , attains a minimum at $m = m^\dagger$. Over the nonnegative reals, it attains minimum at $m = m^*$ and $v''(m^*)$, the second derivative, is positive.

Conditionally on Stage 1, the MRE obeys

$$\text{MRE} \leq \frac{1}{T - K} v(\hat{m}_{K,m_0}^\dagger).$$

This bound shall be compared to the least possible bound, given by $v(m^\dagger)/T$, which would be obtained by a hypothetical “oracle” who would know the true value of the optimal number of repetitions.

For the sake of simplicity, we shall work with the noise-to-signal ratio. This permits to avoid issues related to the fact that $m^\dagger$ must be a compatible integer, while still providing useful insights.

Define the excess-of-variance by

$$\hat{E}_{K,m_0} = \frac{\frac{1}{T-K}v(\hat{m}_{K,m_0}^*) - \frac{1}{T}v(m^*)}{\frac{1}{T}v(m^*)},$$

where $\hat{m}_{K,m_0}^*$ is the optimal number of repetitions produced by Stage 1. The excess-of-variance, which is greater or equal to zero, is literally the excess of variance incurred by our ignorance about m^* . Thus, if $K = 0$ and $m_0 = m^*$, then $\hat{E}_{0,m^*} = 0$. The budget K is subject to a compromise: on the one hand it must be large enough to get a good estimate of the optimal number of repetitions, but on the other hand it must be small enough to keep the remaining budget $T-K$ large enough. The excess-of-variance has an asymptotic expansion given in Lemma 1.

Lemma 1. *Assume that $K/T \rightarrow c/(c+1)$ for some $c \geq 0$ as $m_0 \rightarrow \infty$ and $n_0 \rightarrow \infty$. Take $\sigma > 0$, C and the $o(1)$ term from Theorem 2 and put $V_K = \sqrt{n_0}(\hat{m}_{K,m_0}^* - [m^* + \{C + o(1)\}/m_0])/\sigma$, so that $V_K \xrightarrow{d} N(0,1)$. Then*

$$\begin{aligned} \hat{E}_{K,m_0} = & \frac{(c+1)v''(m^*)}{2v(m^*)} \left(\frac{C}{m_0} + \frac{\sigma V_K}{\sqrt{n_0}} \right)^2 \\ & + \frac{K}{T-K} + o_P\left(\frac{1}{m_0^2}\right) + o_P\left(\frac{1}{m_0\sqrt{n_0}}\right) + o_P\left(\frac{1}{n_0}\right). \end{aligned}$$

At the first order of approximation, that is, when $n_0 \rightarrow \infty$ and $m_0 \rightarrow \infty$, the behavior of the excess-of-variance is controlled by the growth rate of K/T . Indeed, all terms in the asymptotic expansion vanish except possibly $K/(T-K) = (K/T)(1-K/T)^{-1}$. Since $K/(T-K) \rightarrow c$, it holds that $\hat{E}_{K,m_0} \xrightarrow{P} c$. The case $c = 0$, which demands that $K/T \rightarrow 0$, is called the oracle property and is formally stated in Corollary 3.

Corollary 3. *Assume that $K/T \rightarrow 0$. Then, under the assumptions of Lemma 1, $\hat{E}_{K,m_0} \xrightarrow{P} 0$.*

At the second order of approximation, that is, removing all the $o_P(\cdot)$ terms, we see that the error increases with the curvature of v at m^* . A high curvature means that even a small deviation of the optimal number of repetitions estimate from its target can lead to a large error. It is also seen that the error depends on the asymptotic bias and variance of the optimal number of repetitions estimator, the former being controlled by m_0 and the later by $\sqrt{n_0}$. Since m_0 and $\sqrt{n_0}$ are related by the equation $K = m_0 n_0 (p+1)$, this means that there is a tradeoff to find. It turns out that it is exactly when $\sqrt{n_0}$ and m_0 are proportional that the best rate of convergence can be achieved for the excess-of-variance.

To see this, take the term $(Cm_0^{-1} + \sigma V_K n_0^{-1/2})^2$ in the asymptotic expansion of Lemma 1. Assume $m_0 = n_0^{1/2+\beta}$, $-1/2 < \beta < \infty$, so that $\sqrt{n_0}/m_0$ goes to

∞ , 0 or 1 whether $\beta < 0$, $\beta > 0$ or $\beta = 0$. Since $m_0 n_0(p+1) = K$, the taken term is written, up to a multiplicative constant,

$$\begin{cases} K^{-2/(3+2\beta)} (CK^{-2\beta/(3+2\beta)} + \sigma V_K)^2 & \text{if } \beta \geq 0, \\ K^{-2(1+2\beta)/(3+2\beta)} (C + \sigma V_K K^{2\beta/(3+2\beta)})^2 & \text{if } \beta \leq 0. \end{cases}$$

In both cases, the fastest rate of decrease toward zero is $K^{-2/3}$, attained for $\beta = 0$, meaning that m_0 is proportional to $\sqrt{n_0}$.

The above argument carries over to get the optimal rate for K . Let $K = T^\alpha$, $0 < \alpha \leq 1$. The rate $K^{-2/3}$ is then $T^{-2\alpha/3}$ and must be balanced with the term $K/(T-K) \sim T^{\alpha-1}$ in the asymptotic expansion. The sum of those two terms is

$$\begin{cases} T^{-2\alpha/3}(1 + T^{5\alpha/3-1}) & \text{if } 5\alpha/3 \leq 1, \\ T^{\alpha-1}(T^{-5\alpha/3+1} + 1) & \text{if } 5\alpha/3 \geq 1. \end{cases}$$

In both cases, the fastest rate of decrease is $T^{-2/5}$, attained for $\alpha = 3/5$.

These results are stated in Theorem 3, for which a formal proof is given in Appendix A.

Theorem 3. *Take $K = T^{3/5}$, $m_0 = T^{1/5}(p+1)^{-1/3}$ and $n_0 = T^{2/5}(p+1)^{-2/3}$. Then, under the assumptions of Lemma 1,*

$$T^{2/5} \hat{E}_{K,m_0} \xrightarrow{d} \frac{v''(m^*)(p+1)^{2/3}(C + \sigma W)^2}{2v(m^*)},$$

where $W \sim N(0,1)$. Moreover, if $|C| > 0$, then the rate $T^{2/5}$ is optimal: there exist no $\delta > 0$, $0 < \alpha \leq 1$, $-1/2 < \beta < \infty$ such that

$$T^{2/5+\delta} \hat{E}_{K,m_0} = O_P(1)$$

with K proportional to T^α and m_0 proportional to $n_0^{1/2+\beta}$.

Theorem 3 give the answers to the questions raised at the beginning of this section, albeit in an asymptotic framework. Asymptotically, to get the best performance, the share of the budget to be spent in the first stage of Algorithm 2 should be of order $T^{3/5}$ and the number of repetitions of order $T^{1/5}$. This way, the excess-of-variance is guaranteed to vanish at the rate $T^{2/5}$, which is the optimal one.

Remark 1. *The coefficients α and β are not equally important. For instance, if $\beta \geq 0$ and $\alpha \geq 3/5$, the logarithm of the rate is of order $\log(T)(-2\alpha/(3+2\beta) + \alpha - 1)$. The gradient with respect to (α, β) at $(3/5, 0)$ is about $(1.67, -0.27)$. This means that a change in α is likely to have a greater effect on the performance than a change in β . In practice, this means that the choice of m_0 may not be important.*

5 Numerical tests

The effect of the number of repetitions on the sensitivity indices estimators and the effect of the calibration in the two-stage procedure are examined in two kinds of experiments: the “direct” experiments and the “calibration” experiments.

In the direct experiments, the sensitivity indices are estimated directly with the given number of repetitions. Increasing numbers of repetitions m are tested. (Since the budget is fixed, this goes with decreasing numbers of explorations.) For each m , the mean squared errors (MSEs), given by $E(\hat{S}'_{1;n,m} - S'_1)^2 + (\hat{S}'_{2;n,m} - S'_2)^2$ and $E(\hat{S}''_{1;n,m} - S''_1)^2 + (\hat{S}''_{2;n,m} - S''_2)^2$, are estimated with replications. They are also split into the sum of the squared biases and the sum of the variances to get further insight about the behavior of the estimators. The MREs are estimated as well. A normalized version is considered: it is the MRE divided by the number of variables. For models with two inputs, the normalized MRE is interpreted directly as the probability that the two inputs are ranked incorrectly.

In the calibration experiments, the sensitivity indices are estimated with the two-stage procedure, the results of which depend on the calibration parameters K and m_0 . Various calibration parameters are tested to see their effect on the MRE. The budgets for the direct experiments and the calibration experiments are the same so that the numbers can be compared. In particular, the direct experiments correspond to the case $K = 0$ in the calibration experiments.

Two models have been considered: the linear model and a randomized Ishigami model. These models have been chosen because the sensitivity indices are explicit and hence the performance of the estimators can be evaluated easily. For the linear model, the optimal continuous number of repetitions is explicit, too. The formula is given in given in Appendix D.

5.1 Linear model

The model is of the form $Y = X_1 + \beta X_2 + \sigma Z$, where X_1, X_2, Z , are standard normal random variables and β, σ are real coefficients.

5.1.1 High noise context

The coefficients are $\beta = 1.2$ and $\sigma = 4$. The sensitivity indices are $S'_1 = 0.05$, $S'_2 = 0.08$, $S''_1 = 0.41$ and $S''_2 = 0.59$. The optimal continuous number of repetitions is $m^* = 5.8$. The total budget is $T/(p+1) = 500$ and hence the compatible numbers of repetitions are 1, 2, 4, 5, 10, 20, 25, 50, 100, 125, 250, 500. Thus, the optimal number of repetitions is either 5 or 10. The test given in Proposition 2 reveals that it is 5. Since the budget is kept fixed, the numbers of explorations are, respectively, 500, 250, 125, 100, 50, 25, 20, 10, 5, 4, 2, 1. The number of replications is 1500.

The results of the direct experiment are given in Figure 1 for $m = 1, 2, 4, 5, 10, 20, 25$. The MSE of first kind does not vary with the number of repetitions and is much lower than the MSE of second kind, see (c). The estimators of the second kind are highly biased for small numbers of repetitions (a) and they have a higher variance for larger numbers of repetitions (b). The fact that the bias is high for small numbers of repetitions agrees with the theory, according to which the bias should vanish as m goes to infinity. Overall, the sensitivity indices of the second kind seem to be much harder to estimate than the indices of the first kind, the estimators of which have a negligible bias and a very small variance whatever the number of repetitions.

According to Figure 1(c), the normalized MRE curve has a banana shape with a minimum of about slightly less than 30% reached around $m \in \{5, 10\}$ and endpoints with a value of about 35%. A value of 30% means that the probability

$K/3$	m_0				n_0			
	2	5	10	20	20	10	5	2
400	0.43	0.42	0.42	-	-	0.42	0.39	0.40
200	0.38	0.39	0.37	-	-	0.35	0.35	0.34
100	0.36	0.37	-	-	-	-	0.32	0.30
50	0.39	0.33	-	-	-	-	0.33	0.31

Table 1: Normalized MRE in the linear model with high noise for various calibrations: $K/(p+1) = 50, 100, 200, 400$ and $m_0 = 2, 5, 10, 20, \dots$. For instance, for $K/(p+1) = 200 = m_0 n_0$, the normalized MRE is available for $m_0 = 2, 5, 10, 20, 40, 100$.

of ranking the inputs correctly is about 70%. The region of observed optimal performance $m \in \{5, 10\}$ agrees with the optimal number of repetitions $m = 5$ predicted by the theory, although the last is, in fact, the minimizer of only a bound of the MRE.

The results of the calibration experiment is given in Table 1 for the normalized MRE. The lowest MREs are reached at the bottom right of the table, with values corresponding to $2 \leq m \leq 10$ in Figure 1 (c). Optimal performance is reached with very few explorations in the first stage of the two-stage procedure. This corresponds to an estimator of the optimal number of repetitions that has a small bias but a high variance. Such an estimator seems to perform better than an estimator with a small variance but a large bias. This might be explained by the low curvature of the MRE curve.

5.1.2 Low noise context

The coefficients are $\beta = 1.2$ and $\sigma = 0.9$. The sensitivity indices are $S'_1 = 0.31$, $S'_2 = 0.44$, $S''_1 = 0.41$ and $S''_2 = 0.59$. The optimal continuous number of repetitions is $m^* = 0.30$ and hence the optimal number of repetitions is $m^\dagger = 1$. As expected, the optimal number of repetition is smaller than the one found in the high noise context. The total budget is $T/(p+1) = 500$. The number of replications is 500.

The results for the direct experiment are given in Figure 2. The MSE of first kind increases with the number of repetitions, see (c): this is due to the increase of the variance (b), while the bias is negligible (a). As in the high noise context, the estimators of the second kind have a decreasing bias and an increasing variance, although the decrease of the bias is of much less magnitude. This agrees with the theory, where we have seen that, for the sensitivity indices of the second kind, the biases of the estimators are small when the noise of the model is low.

In Figure 2 (c), the normalized MRE varies a lot. It increases from about 2% at $m = 1$ to 30% at $m = 25$. Thus, unlike in the high noise setting, choosing a good number of repetitions is important. The best performance is achieved at $m = 1$, which corresponds to the predicted optimal number of repetitions.

The results of the calibration experiment for the normalized MRE is given in Table 2. The best performance is reached at the bottom left of the table with numbers that correspond to the optimal performance in Figure 2 (c). Moreover, notice that a large spectrum of calibration parameters (K, m_0) yield low errors.

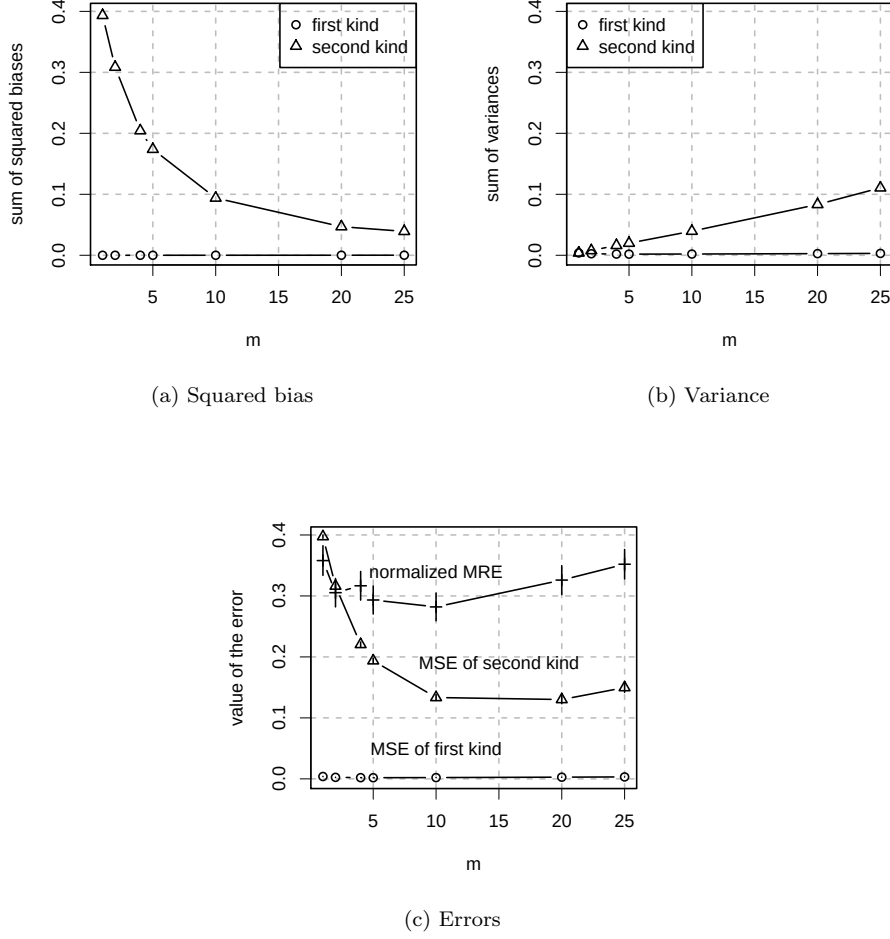
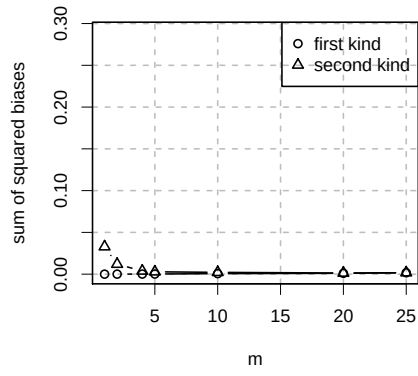


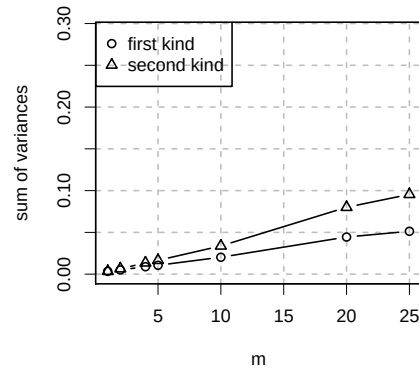
Figure 1: Sum of squared biases (a), sum of variances (b) and errors (c) of the sensitivity indices estimators for the linear model in the high noise setting. Confidence intervals of level 95% are added in (c).

	m_0				n_0			
$K/3$	2	5	10	20	20	10	5	2
400	0.18	0.15	0.17	-	-	0.16	0.18	0.20
200	0.05	0.04	0.04	-	-	0.06	0.05	0.07
100	0.02	0.04	-	-	-	-	0.04	0.04
50	0.03	0.02	-	-	-	-	0.02	0.04

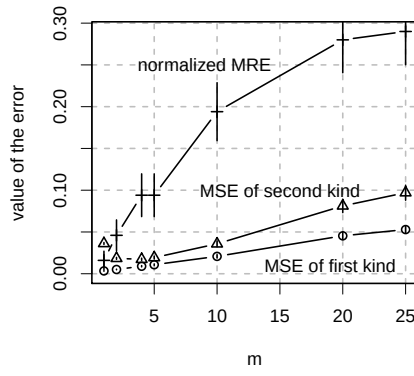
Table 2: Normalized MRE in the linear model with low noise for various calibrations: $K/(p+1) = 50, 100, 200, 400$ and $m_0 = 2, 5, 10, 20, \dots$. For instance, for $K/(p+1) = 200 = m_0 n_0$, the normalized MRE is available for $m_0 = 2, 5, 10, 20, 40, 100$.



(a) Squared bias



(b) Variance



(c) Error

Figure 2: Sum of squared biases (a), sum of variances (b) and errors (c) of the sensitivity indices estimators for the linear model in the low noise context. Confidence intervals of level 95% are added in (c).

$K/3$	m_0				n_0			
	2	5	10	20	20	10	5	2
400	0.22	0.27	0.22	-	-	0.21	0.27	0.24
200	0.11	0.12	0.10	-	-	0.10	0.08	0.10
100	0.08	0.08	-	-	-	-	0.08	0.11
50	0.10	0.08	-	-	-	-	0.07	0.06

Table 3: Normalized MRE in the Ishigami model for various calibrations: $K/(p+1) = 50, 100, 200, 400$ and $m_0 = 2, 5, 10, 20, \dots$. For instance, for $K/(p+1) = 200 = m_0 n_0$, the normalized MRE is available for $m_0 = 2, 5, 10, 20, 40, 100$.

5.2 The Ishigami model

The Ishigami model, a benchmark model in the sensitivity analysis literature [7, 14], is $Y = \sin(X_1) + \beta \sin(X_2)^2 + \sigma Z^4 \sin(X_1)$, where $\beta = 7$, $\sigma = 0.1$ and $X_1, X_2, Z \sim \mathcal{U}(-\pi, \pi)$. The sensitivity indices are $S'_1 = 0.31$, $S'_2 = 0.44$, $S''_1 = 0.42$ and $S''_2 = 0.58$. The total budget is $T/(p+1) = 500$. The number of replications is 500.

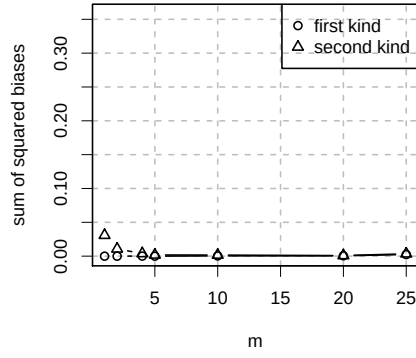
The results of the direct experiment are given in Figure 3. The plots are quite similar to the linear model with low noise.

The results of the calibration experiment for the normalized MRE are given in Table 3. The lowest values are reached at the bottom of the table and correspond to optimal performance in Figure 3 (c).

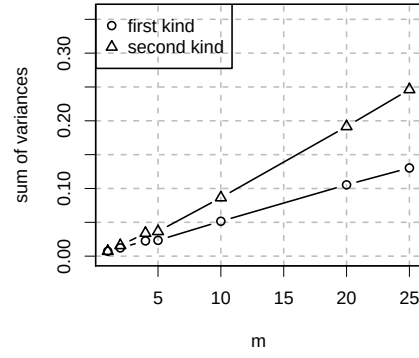
6 Sensitivity analysis for a SIR model

The Susceptible-Infectious-Recovered (aka SIR) model underpins many of epidemics models. A closed population of size N is followed at each time event, where an infection or a recovery occurs. At each time step $i = 0, 1, \dots$, it holds $N = S^i + I^i + R^i$, where S^i , I^i and R^i are the number of susceptible, infectious and recovered individuals, respectively. In case of an infection, the number of infectious is increased by one unit and the number of susceptible is decreased by the same amount. In case of a recovery, the number of infectious is decreased by one unit and the number of recovered is incremented. Recovered individuals cannot be sick again. The time between two consecutive events T^{i-1} and T^i is an exponential random variable depending on parameters R_0 , τ , I^{i-1} , S^{i-1} and N . The probability according to which an infection or a recovery happens also depends on those parameters. The parameters S^0 , I^0 and N are assumed to be known and fixed, so that the SIR model can be seen as a model with two parameters: $R_0 = X_1$ and $\tau = X_2$. For the sake of illustration, we shall be interested in the output $S^0 - S^{i^*} = Y$, where i^* is the smallest value of i such that $S^i = 0$ or $I^i = 0$, which represents the end of the epidemic. The SIR model is stochastic: even if the values of R_0 and τ are fixed, the output is still a random variable. The larger the population, the more the SIR model is deterministic. This is because the SIR model converges to a deterministic model when $N \rightarrow \infty$. For more details about SIR models, see, e.g., [3]. For the sake of completeness, the SIR model used in this section is given in Algorithm 3.

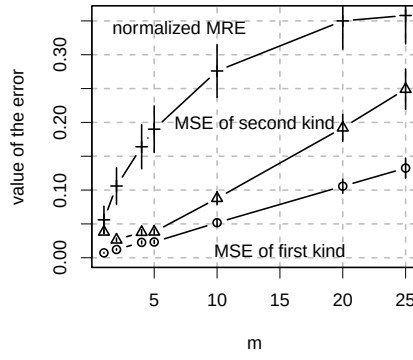
To perform the sensitivity analysis, the parameters are drawn independently as $R_0 \sim \text{Unif}(2 - \epsilon_{R_0}, 2 + \epsilon_{R_0})$ and $\tau \sim \text{Unif}(4 - \epsilon_\tau, 4 + \epsilon_\tau)$. By default, $\epsilon_{R_0} = 0.9$



(a) Squared bias



(b) Variance



(c) Error

Figure 3: Sum of squared biases (a), sum of variances (b) and errors (c) of the sensitivity indices estimators for the randomized Ishigami model. Confidence intervals of level 95% are added in (c).

Algorithm 3 A SIR model

Require: R_0, τ, N, S^0, I^0 $i = 0$ **while** $S^i > 0$ and $I^i > 0$ **do** $i = i + 1$ draw $T^i \sim \text{Exp}(\text{mean} = \tau/[R_0 S^{i-1} I^{i-1}/N + I^{i-1}])$ draw $u \sim \text{Unif}(0, 1)$ **if** $u \leq [R_0 S^{i-1} I^{i-1}/N]/[R_0 S^{i-1} I^{i-1}/N + I^{i-1}]$ **then** $I^i = I^{i-1} + 1$ $S^i = S^{i-1} - 1$ **else** $I^i = I^{i-1} - 1$ $R^i = R^{i-1} + 1$ **end if****end while****return** $S^0 - S^i$

and $\epsilon_\tau = 2$. The total budget is set to $T/3 = 1000$, the share dedicated to the estimation of the optimal number of repetition is set to $K/3 = 100$ and $m_0 = 10$. To get insight about the variability of the estimators, the sensitivity analyses are replicated. The number of replications is 100. The proportion of infectious at start is 10%. The population size is 100, unless stated otherwise.

Since no closed form expression of the sensitivity indices exist for the SIR model, our goal is to recover characteristics of sensitivity indices and/or SIR models. Since the model becomes more and more deterministic as the population size increases, we expect the optimal number of repetition to decrease along the way. The estimated sensitivity indices of the first kind should increase because the variance of the noise Z should vanish. (This is not true for the indices of the second kind.) Finally, when ϵ_{R_0} or ϵ_τ increase, the corresponding sensitivity indices are expected to increase as well. (In fact, this is true only under certain conditions, see [2])

The results are shown in Figure 4 and Figure 5. In Figure 4 (a), we see that the optimal number of repetitions (in fact, its continuous version $\hat{m}^*$) decreases from one to zero as the population size decreases from 10 to 500. Thus, it seems that, even for a very small population size, repeating the SIR model more than once is suboptimal. Remember that when there is no repetition ($m = 1$), the sensitivity estimators of the first and the second kind coincide. This is why, in Figure 4 (b), there is only two curves. The one which stays at zero are the sensitivity indices corresponding to τ , which does not seem to influence the output. The curve that goes up represents the sensitivity indices for R_0 . When the population is small, there are about zero, indicating that the variability of the output essentially stems from the noise of the model. As the population gets larger, R_0 gets more and more influential (and hence the noise is less and less influential).

In Figure 5, the sensitivity indices are displayed with respect to a change in ϵ_{R_0} or ϵ_τ . The sensitivity indices for R_0 increase with ϵ_{R_0} but this is not so for τ . Thus, it appears that the parameter τ hardly influences the output our SIR model, even when it is drawn on a wider interval.

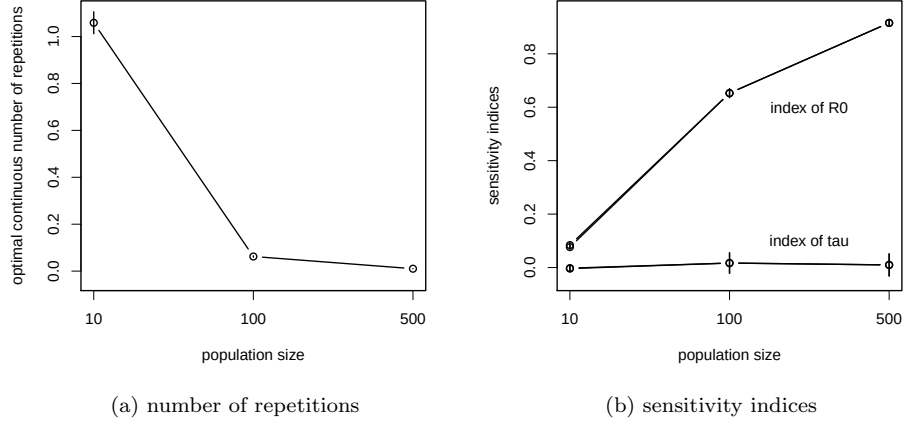


Figure 4: Optimal (continuous) number of repetitions estimates and sensitivity indices estimates in terms of population size. The vertical bars, of length four standard deviations divided by the square root of the number of replications, represent the variability of the estimates.

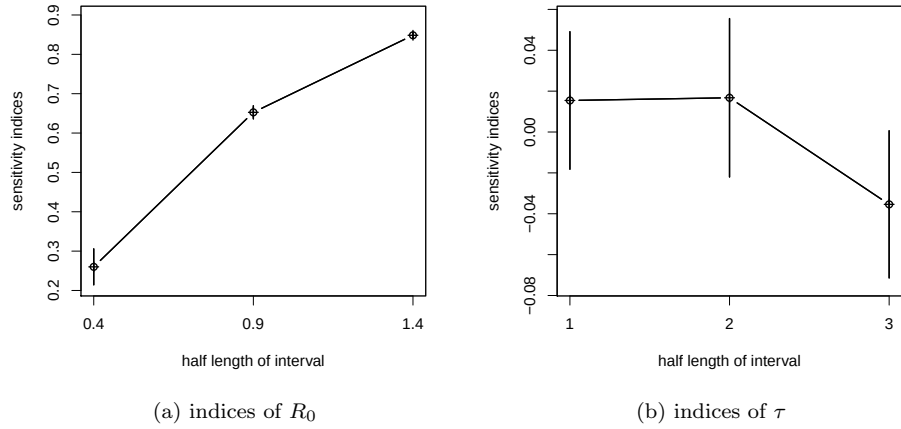


Figure 5: Sensitivity indices estimates in terms of the length of the interval on which the SIR parameters are drawn. The vertical bars, of length four standard deviations divided by the square root of the number of replications, represent the variability of the estimates.

7 Discussion

We have considered two sensitivity indices for stochastic models. Asymptotic normality of the estimators, which depend both on the number of explorations and the number of repetitions, has been established, and it was noticed that the second kind, that which arises from smoothing out the computer model, suffers from a bias term which vanishes only when the number of repetitions goes to infinity.

Assuming a fixed computing budget, a definition of the optimal number of repetitions has been proposed and an estimator has been built. The optimal number of repetitions minimizes an upper-bound of the missranking error (MRE), an error which is small when the inputs of the model are ranked correctly. This approach has been exploited in a two-stage procedure to efficiently estimate the sensitivity indices. The procedure is simple: in the first stage the optimal number of repetitions is estimated and plugged into the sensitivity estimators in the second stage.

Although a share of the total budget must be sacrificed in the first stage, we have shown that the proposed procedure is asymptotically oracle, that is, performs as well as the hypothetical procedure in which the optimal number of repetitions would be known. The optimal convergence rates have been given.

To test the procedure, simulation experiments were conducted, where the bias of the sensitivity estimator of the second kind was confirmed. Optimal compromises between repetitions and explorations have been identified to which the two-stage procedure was compared for different values of the tuning parameters (among which the share of the budget spent in the first stage).

Sensitivity analysis of a SIR model has been carried out and it has been found that the τ parameters almost has no influence on the considered output. Most of the variability is due to the parameter R_0 for large populations. For small populations, it is the noise that contributes the most. The optimal number of repetitions was found be one. That is, doing repetitions in the considered SIR model appears suboptimal for estimating the sensitivity indices: it is best to explore as much as possible the input space.

Research questions remain open. First, the sensitivity estimators of the two stages could be aggregated to build estimators with a lower variance. Second, higher-order Sobol indices could be studied in the light of the theory introduced in this paper. Third, it would be interesting to look for other optimality criteria or Monte-Carlo sampling schemes. Finally, the obtained results could be extended to the multivariate case, that in which the output is a multivariate vector.

Acknowledgment

The author thanks Hervé Monod and Elisabeta Vergu for stimulating discussions from which this work arose and Bertrand Iooss for the references [6, 8].

References

- [1] A. Courcoul, H. Monod, M. Nielen, D. Klinkenberg, L. Hogerwerf, F. Beaudou, and E. Vergu. Modelling the effect of heterogeneity of shed-

- ding on the within herd coxiella burnetii spread and identification of key parameters by sensitivity analysis. *Journal of Theoretical Biology*, 284:130–141, 2011.
- [2] A. Cousin, A. Janon, V. M. Deschamps, and I. Niang. On the consistency of sobol indices with respect to stochastic ordering of model parameters. *ESAIM: Probability and Statistics*, 23:387–408, 2019.
 - [3] D. J. Daley and J. Gani. *Epidemic Modelling*. Cambridge, 1999.
 - [4] F. Gamboa, A. Janon, T. Klein, and A. Lagnoux. Sensitivity analysis for multidimensional and functional outputs. *Electron. J. Stat.*, 8(1):575–603, 2014.
 - [5] F. Gamboa, A. Janon, T. Klein, A. Lagnoux, and C. Prieur. Statistical inference for Sobol pick-freeze Monte Carlo method. *Statistics*, 50(4):881–902, 2016.
 - [6] J. L. Hart, A. Alexanderian, and P. A. Gremaud. Efficient computation of Sobol’ indices for stochastic models. *SIAM Journal on Scientific Computing*, 39(4):A1514–A1530, 2017.
 - [7] T. Homma and A. Saltelli. Importance measures in global sensitivity analysis of nonlinear models. *Reliability Engineering & System Safety*, 52(1):1–17, 1996.
 - [8] B. Iooss, L. L. Gratiet, A. Lagnoux, and T. Klein. Sensitivity analysis for stochastic computer codes: Theory and estimation methods. Technical report, EDF R&D, 2014.
 - [9] B. Iooss, T. Klein, and A. Lagnoux. Sobol’ sensitivity analysis for stochastic numerical codes. In *8th International Conference Sensitivity Analysis of Model Output*, pages 48–49, 2016.
 - [10] B. Iooss and M. Ribatet. Global sensitivity analysis of computer models with functional inputs. *Reliability Engineering & System Safety*, 94(7):1194 – 1204, 2009.
 - [11] T. Ishigami and T. Homma. An importance quantification technique in uncertainty analysis for computer models. Technical report, JAERI-M-89-111, 1989.
 - [12] A. Janon, T. Klein, A. Lagnoux, M. Nodet, and C. Prieur. Asymptotic normality and efficiency of two Sobol index estimators. *ESAIM: Probability and Statistics*, 18:342–364, 2014.
 - [13] M. Lamboni, H. Monod, and D. Makowski. Multivariate sensitivity analysis to measure global contribution of input factors in dynamic models. *Reliability Engineering and System Safety*, 96:450–459, 2010.
 - [14] A. Marrel, B. Iooss, S. Da Veiga, and M. Ribatet. Global sensitivity analysis of stochastic computer models with joint metamodels. *Statistics and Computing*, 22(3):833–847, 2012.

- [15] H. Monod, C. Naud, and D. Makowski. Working with dynamic crop models: Evaluation, analysis, parameterization, and applications. In *Uncertainty and sensitivity analysis for crop models*, pages 55–100. Elsevier, 2006.
- [16] L. Rimbaud, C. Bruchou, S. Dallot, D. R. J. Pleydell, E. Jacquot, S. Soubeyrand, and G. Thébaud. Using sensitivity analysis to identify key factors for the propagation of a plant epidemic. *Open Science*, 5(1), 2018.
- [17] A. Saltelli, S. Tarantola, and F. Campolongo. Sensitivity analysis as an ingredient of modeling. *Statistical Science*, 15(4):377–395, 2000.
- [18] A. Saltelli, S. Tarantola, F. Campolongo, and M. Ratto. *Sensitivity analysis in practice*. Wiley, 2004.
- [19] J. F. Savall, C. Bidot, M. Leblanc-Maridor, C. Belloc, and S. Touzeau. Modelling salmonella transmission among pigs from farm to slaughterhouse: interplay between management variability and epidemiological uncertainty. *International Journal of Food Microbiology*, 229:33–43, 2016.
- [20] I. M. Sobol. Sensitivity estimates for nonlinear mathematical models. *Mathematical modelling and computational experiments*, 1(4):407–414, 1993.
- [21] I. M. Sobol. Global sensitivity indices for nonlinear mathematical models and their Monte Carlo estimates. *Mathematics and computers in simulation*, 55(1-3):271–280, 2001.
- [22] M. Spence. *Statistical issues in ecological simulation models*. PhD thesis, University of Sheffield, 2015. <http://etheses.whiterose.ac.uk/10517>.
- [23] A. W. van der Vaart. *Asymptotic Statistics*. Cambridge University Press, 1998.

A Proofs

Proof of Theorem 1

The proof is based on the results in Appendix B. The Sobol estimators in (5) and (6) are made of empirical averages of the form (20) and (21) with $L = 2$ and coefficients

$$\xi_{1;m,i}^{\text{UL}} \simeq \begin{pmatrix} 1 & 0 & \cdots & 0 \\ 0 & 1 & \cdots & 0 \end{pmatrix}, \dots, \xi_{p;m,i}^{\text{UL}} \simeq \begin{pmatrix} 1 & 0 & \cdots & 0 \\ 0 & 0 & \cdots & 1 \end{pmatrix}$$

for the upper left (UL) terms,

$$\xi_{m,i}^{\text{UR}} \simeq \begin{pmatrix} 1 & 0 & \cdots & 0 \\ 0 & 0 & \cdots & 0 \end{pmatrix}$$

for the upper right (UR) term,

$$\xi_{m,i}^{\text{LL}} \simeq \begin{pmatrix} 2 & 0 & \cdots & 0 \\ 0 & 0 & \cdots & 0 \end{pmatrix}$$

for the lower left (LL) term of $\widehat{S}_{j;n,m}'$ and

$$\xi_{m,i}^{\text{LL}} \simeq \begin{pmatrix} 1 & 0 & \cdots & 0 \\ 1 & 0 & \cdots & 0 \end{pmatrix}$$

for the lower left (LL) term of $\widehat{S}_{j;n,m}''$ (see Appendix B). Therefore, denoting by

$$\boldsymbol{\xi}_{m,i} := (\xi_{1;m,i}^{\text{UL}}, \dots, \xi_{p;m,i}^{\text{UL}}, \xi_{m,i}^{\text{UR}}, \xi_{m,i}^{\text{LL}}, \xi_{m,i}^{\text{LL}})$$

the corresponding summands, Lemma 3 in Appendix B yields

$$\sqrt{n} \left(\frac{1}{n} \sum_{i=1}^n \boldsymbol{\xi}_{m,i} - \mathbb{E} \boldsymbol{\xi}_{m,1} \right) \xrightarrow{d} N(0, \Sigma_m),$$

for some nonnegative matrix Σ_m of size $p+3$, where the expectation is given by

$$\mathbb{E} \boldsymbol{\xi}_{m,1} = \begin{pmatrix} \mathbb{E} \mathbb{E}[f(X, Z)|X] \mathbb{E}[f(\widetilde{X}_{-1}, Z)|\widetilde{X}_{-1}] \\ \vdots \\ \mathbb{E} \mathbb{E}[f(X, Z)|X] \mathbb{E}[f(\widetilde{X}_{-p}, Z)|\widetilde{X}_{-p}] \\ \mathbb{E} f(X, Z) \\ \mathbb{E} f(X, Z)^2 \\ \mathbb{E} \mathbb{E}[f(X, Z)|X]^2 + \frac{\mathbb{E} \text{Var}[f(X, Z)|X]}{m} \end{pmatrix}.$$

(Some calculations are carried out in Appendix C.) Define the function

$$\begin{aligned} & s(x_1, \dots, x_p, x_{p+1}, x_{p+2}, x_{p+3}) \\ &= \left(\frac{x_1 - x_{p+1}^2}{x_{p+2} - x_{p+1}^2}, \dots, \frac{x_p - x_{p+1}^2}{x_{p+2} - x_{p+1}^2}, \frac{x_1 - x_{p+1}^2}{x_{p+3} - x_{p+1}^2}, \dots, \frac{x_p - x_{p+1}^2}{x_{p+3} - x_{p+1}^2} \right). \end{aligned}$$

Then clearly

$$s(n^{-1} \sum_{i=1}^n \boldsymbol{\xi}_{m,i}) = \begin{pmatrix} \widehat{\mathbf{S}}'_{m,n} \\ \widehat{\mathbf{S}}_{m,n} \end{pmatrix}$$

and

$$s(\mathbb{E} \boldsymbol{\xi}_{m,1}) = \begin{pmatrix} \mathbf{S}' \\ \mathbf{S}'' \left[1 - \frac{\mathbb{E} \text{Var}[f(X,Z)|X]}{\mathbb{E} \text{Var}[f(X,Z)|X] + m \text{Var} \mathbb{E}[f(X,Z)|X]} \right] \end{pmatrix}.$$

The result follows by the delta-method.

If $m = m_n \rightarrow \infty$, Lemma 3 still holds with the variance-covariance matrix replaced by its limit.

Proof of Proposition 1

Assume without loss of generality that $D_1 < \dots < D_p$. We first prove the following Lemma. For convenience, the subscripts n and m are left out.

Lemma 2. *Let $i < j$. Then*

$$P(\widehat{D}_i - \widehat{D}_j \geq 0) \leq \frac{\text{Var} \widehat{D}_i + \text{Var} \widehat{D}_j}{\frac{1}{2}|D_i - D_j|^2}.$$

Proof. We have

$$\begin{aligned} P(\widehat{D}_i - \widehat{D}_j \geq 0) &\leq P(|\widehat{D}_i - D_i| + |\widehat{D}_j - D_j| \geq D_j - D_i) \\ &\leq P(|\widehat{D}_i - D_i|^2 + |\widehat{D}_j - D_j|^2 \geq \frac{1}{2}|D_j - D_i|^2) \end{aligned}$$

and the claim follows from Markov's inequality.

We now prove Proposition 1. Recall that $D_1 < \dots < D_p$. We have

$$\begin{aligned} \sum_{i=1}^p \mathbb{E} |\widehat{R}_i - R_i| &\leq \sum_{i=1}^p \sum_{j=1}^p \mathbb{E} |\mathbf{1}(\widehat{D}_j \leq \widehat{D}_i) - \mathbf{1}(D_j \leq D_i)| \\ &\leq \sum_{i=1}^p \sum_{j \neq i} \frac{\text{Var} \widehat{D}_i + \text{Var} \widehat{D}_j}{\frac{1}{2}|D_i - D_j|^2} \\ &\leq \frac{2(p-1)(p+1)}{\min_{j < j'} |D_j - D_{j'}|^2} \sum_{i=1}^p \text{Var} \widehat{D}_i, \end{aligned}$$

where the second inequality holds by Lemma 2 and because

$$\mathbb{E} |\mathbf{1}(\widehat{D}_j \leq \widehat{D}_i) - \mathbf{1}(D_j \leq D_i)| = \begin{cases} \mathbb{E} |\mathbf{1}(\widehat{D}_j > \widehat{D}_i)| & \text{if } j < i, \\ 0 & \text{if } j = i, \\ \mathbb{E} |\mathbf{1}(\widehat{D}_j \leq \widehat{D}_i)| & \text{if } j > i. \end{cases}$$

It remains to calculate the variances. But this is done in Lemma 6 in Appendix C, where it is found that

$$\begin{aligned} \text{Var} \widehat{D}_j &= \frac{1}{n} \{ \text{Var} \mathbb{E}[Y_0 Y_j | \mathbf{X}] + \frac{1}{m} (\mathbb{E} \text{Var}[Y_0 Y_j | \mathbf{X}] - \text{Var}[Y_0 | \mathbf{X}] \text{Var}[Y_j | \mathbf{X}]) \\ &\quad + \frac{1}{m^2} \mathbb{E} \text{Var}[Y_0 | \mathbf{X}] \text{Var}[Y_j | \mathbf{X}] \}. \end{aligned}$$

Proof of Proposition 2

Denote by $f(m)$ the upper-bound in Proposition 1. By convexity of the function f , if $m^* \leq 1$, then $f(1) \leq f(m)$ for all (compatible) integers m and hence $m^\dagger = 1$. Likewise, if $m^* \geq T/(p+1)$, then $f(T/(p+1)) \leq f(m)$ for all m and hence $m^\dagger = T/(p+1)$. For (iii), note that $m^\dagger$ must be either $\lfloor m \rfloor^*$ or $\lceil m \rceil^*$. Since $\lfloor m \rfloor^* \leq \lceil m \rceil^*$, $f(\lfloor m \rfloor^*) \geq f(\lceil m \rceil^*)$ is equivalent to $\lfloor m \rfloor^* \lceil m \rceil^* \leq m^{*2}$.

Proof of Theorem 2

We want to show

$$\sqrt{n} \left(\sqrt{\frac{\sum_{j=1}^p \hat{\zeta}_{3,j}}{\sum_{j=1}^p \hat{\zeta}_{1,j}}} - \left[m^* + \frac{C}{m} + o\left(\frac{1}{m}\right) \right] \right) \rightarrow N(0, \sigma^2),$$

for some C and $\sigma \geq 0$, as $n \rightarrow \infty$, $m = m_n \rightarrow \infty$. In view of (9)–(14), the estimators $\hat{\zeta}_{3,j}$ and $\hat{\zeta}_{1,j}$ can be expressed in terms of estimators of the form (20) and (21), namely,

$$\begin{aligned} \hat{\zeta}_{3,j} &= \frac{1}{n} \sum_{i=1}^n \xi_{j;m,i}^{(9)} + \xi_{j;m,i}^{(10)} - \xi_{j;m,i}^{(11)} - \xi_{j;m,i}^{(12)}, \quad \text{and,} \\ \hat{\zeta}_{1,j} &= \frac{1}{n} \sum_{i=1}^n \xi_{j;m,i}^{(13)} - \left(\frac{1}{n} \sum_{i=1}^n \xi_{j;m,i}^{(14)} \right)^2, \end{aligned}$$

where $\xi_{j;m,i}^{(e)}$, $j = 1, \dots, p$, $e = 9, \dots, 14$, are all of the form (21), that is, they are all elements of the set $\mathcal{E}_{m,i}(4)$, introduced in Appendix B. Their coefficients are given by

$$\hat{\xi}_{j;m,i}^{(9)} \simeq \begin{pmatrix} 2 & 0 \\ 0 & 2 \\ 0 & 0 \\ 0 & 0 \end{pmatrix}, \hat{\xi}_{j;m,i}^{(10)} \simeq \begin{pmatrix} 1 & 0 \\ 1 & 0 \\ 0 & 1 \\ 0 & 1 \end{pmatrix}, \hat{\xi}_{j;m,i}^{(11)} \simeq \begin{pmatrix} 1 & 0 \\ 1 & 0 \\ 0 & 2 \\ 0 & 0 \end{pmatrix}, \hat{\xi}_{j;m,i}^{(12)} \simeq \begin{pmatrix} 2 & 0 \\ 0 & 1 \\ 0 & 1 \\ 0 & 0 \end{pmatrix},$$

and

$$\hat{\xi}_{j;m,i}^{(13)} \simeq \begin{pmatrix} 1 & 1 \\ 1 & 1 \\ 0 & 0 \\ 0 & 0 \end{pmatrix}, \hat{\xi}_{j;m,i}^{(14)} \simeq \begin{pmatrix} 1 & 1 \\ 0 & 0 \\ 0 & 0 \\ 0 & 0 \end{pmatrix}.$$

Stack all the estimators in a vector $\boldsymbol{\xi}_{m,i}$ of size $6p$, that is, $\boldsymbol{\xi}_{m,i} = (\boldsymbol{\xi}_{1;m,i}^\top, \dots, \boldsymbol{\xi}_{p;m,i}^\top)^\top$, where $\boldsymbol{\xi}_{j;m,i} = (\xi_{j;m,i}^{(9)}, \dots, \xi_{j;m,i}^{(14)})^\top$ for $j = 1, \dots, p$. Let $\bar{\boldsymbol{\xi}} = n^{-1} \sum_{i=1}^n \boldsymbol{\xi}_{m,i}$. Since, for each element of $\mathcal{E}_{m,i}(4)$, the sum of the coefficients is less than four, the condition in Lemma 3 in Appendix B is fulfilled because we assumed $f(X, Z) \leq F(X)$ with $\mathbb{E} F(X)^8 < \infty$. Therefore,

$$\sqrt{n}(\bar{\boldsymbol{\xi}} - \mathbb{E} \boldsymbol{\xi}_{m,1}) \rightarrow N(0, \Sigma),$$

for some variance-covariance matrix Σ of size $6p \times 6p$. A delta-method will yield the result. Let $\mathbf{x}_j = (x_j^{(9)}, \dots, x_j^{(14)})$, $j = 1, \dots, p$, and

$$s(\mathbf{x}_1, \dots, \mathbf{x}_p) = \sqrt{\frac{\sum_{j=1}^p x_j^{(9)} + x_j^{(10)} - x_j^{(11)} - x_j^{(12)}}{\sum_{j=1}^p x_j^{(13)} - x_j^{(14)2}}},$$

so that $s(\bar{\xi}) = \hat{m}^*$. Denote $E\xi_{m,1}$ by θ_m for simplicity. A Taylor expansion of s around $E\xi_{m,1}$ yields

$$\sqrt{n}(\hat{m}^* - s(\theta_m)) = \sqrt{n}(\bar{\xi} - \theta_m)^\top \dot{s}(\theta_m) + \frac{1}{2}(\bar{\xi} - \theta_m)^\top \ddot{s}_{n,m} \sqrt{n}(\bar{\xi} - \theta_m),$$

where $\ddot{s}_{n,m}$ is the Hessian matrix of s at a point between $\bar{\xi}$ and θ_m . The second order term is $o_P(1)$ by Cauchy-Schwartz's inequality, the continuity of $\ddot{s}$, and the fact that $\sqrt{n}(\bar{\xi} - \theta_m) = O_P(1)$. By Proposition 3 in Appendix B, there exists θ such that $\theta_m \rightarrow \theta$ and hence the first order term goes to $N(0, \dot{s}^\top \Sigma \dot{s})$, where $\dot{s}$ is the gradient of s at θ .

The image of θ_m by s is not m^* but the discrepancy can be controlled when $m \rightarrow \infty$. Elementary calculations show that

$$E\xi_{j;m,1} = E \underbrace{\begin{pmatrix} Y_0^{(1,1)2} Y_j^{(1,1)2} \\ Y_0^{(1,1)} Y_0^{(1,2)} Y_j^{(1,1)} Y_j^{(1,2)} \\ Y_0^{(1,1)} Y_0^{(1,2)} Y_j^{(1,1)2} \\ Y_j^{(1,1)} Y_j^{(1,2)} Y_0^{(1,1)2} \\ Y_0^{(1,1)} Y_0^{(1,2)} Y_j^{(1,1)} Y_j^{(1,2)} \\ Y_j^{(1,1)} Y_0^{(1,1)} \end{pmatrix}}_{\theta_j} + \frac{\mathbf{C}_j}{m} + o\left(\frac{1}{m}\right)$$

for some constant $\mathbf{C}_j \in \mathbf{R}^p$. It can be checked that $s(\theta_1, \dots, \theta_p) = m^*$. Thus, for some constant C ,

$$\begin{aligned} \sqrt{n}(\hat{m}^* - s(\theta_m)) &= \sqrt{n}(\hat{m}^* - s(\theta)) - (\theta_m - \theta)^\top \dot{s} + o(\|\theta_m - \theta\|) \\ &= \sqrt{n}(\hat{m}^* - [m^* + C/m + o(1/m)]), \end{aligned}$$

and hence the proof is complete with $\sigma^2 = \dot{s}^\top \Sigma \dot{s}$.

Proof of Corollary 2

With the same notations as those of the proof of Theorem 2, if $\sqrt{n}/m \rightarrow 0$, then $\sqrt{n}(\hat{m}^* - m^*) = \sqrt{n}(\hat{m}^* - s(\theta_m)) + O(\sqrt{n}/m) \rightarrow N(0, \sigma^2)$. The proof is complete.

Proof of Lemma 1

Denote $m^* + (C + o(1))/m_0$ by m_K^* . We have the following Taylor expansions:

$$(16) \quad v(m_K^*) = v(m^*) + \frac{1}{2}(m_K^* - m^*)^2 v''(\bar{m}_K^*),$$

$$(17) \quad v(\hat{m}_K^*) = v(m_K^*) + (\hat{m}_K^* - m_K^*) v'(m_K^*) + \frac{1}{2}(\hat{m}_K^* - m_K^*)^2 v''(\tilde{m}_K^*),$$

where $\tilde{m}_K^*$ is between $\hat{m}_K^*$ and m_K^* , $\bar{m}_K^*$ is between m_K^* and m^* . By using (17), we get

$$(18) \quad \hat{E}_K = E_K + \frac{(\hat{m}_K^* - m_K^*)v'(m_K^*)T}{(T - K)v(m^*)} + \frac{(\hat{m}_K^* - m_K^*)^2 v''(\tilde{m}_K^*)T}{2(T - K)v(m^*)},$$

where

$$E_K = \frac{\frac{1}{T-K}v(m_K^*) - \frac{1}{T}v(m^*)}{\frac{1}{T}v(m^*)}.$$

Let us find asymptotic expansions for the three terms in the right-hand side of (18). Since v'' is continuous, $\tilde{m}_K^*$ goes to m^* in probability and $T/(T - K) \rightarrow c + 1$ as $n_0 \rightarrow \infty$ and $m_0 \rightarrow \infty$, the third term satisfies

$$\frac{(\hat{m}_K^* - m_K^*)^2 v''(\tilde{m}_K^*)T}{2(T - K)v(m^*)} = \frac{(c + 1)\sigma^2 V_K^2 v''(m^*)}{2v(m^*)n_0} + o_P\left(\frac{1}{n_0}\right).$$

For the second term, use Taylor's expansion of v' at m^* to get

$$\frac{(\hat{m}_K^* - m_K^*)v'(m_K^*)T}{(T - K)v(m^*)} = \frac{(c + 1)Cv''(m^*)\sigma V_K}{v(m^*)m_0\sqrt{n_0}} + o_P\left(\frac{1}{m_0\sqrt{n_0}}\right).$$

Finally, using (16), we get

$$\begin{aligned} E_K &= \frac{T[v(m_K^*) - v(m^*)] + Kv(m^*)}{(T - K)v(m^*)} \\ &= \frac{Tv''(\bar{m}_K^*)(m_K^* - m^*)^2}{2(T - K)v(m^*)} + \frac{K}{T - K} \\ &= \frac{(c + 1)v''(m^*)C^2}{2v(m^*)m_0^2} + \frac{K}{T - K} + o_P\left(\frac{1}{m_0^2}\right). \end{aligned}$$

The proof is complete.

Proof of Theorem 3

With the given choice for K , m_0 and n_0 , the rates m_0^{-2} , $m_0^{-1}n_0^{-1/2}$ and n_0^{-1} are all equal to $T^{-2/5}(p + 1)^{2/3}$. The formula in Lemma 1 simplifies to

$$\begin{aligned} T^{2/5}\hat{E}_{K,m_0} &= \frac{v''(m^*)(p + 1)^{2/3}(C + \sigma V_K)^2}{2v(m^*)} + o_P(1) \\ &\xrightarrow{d} \frac{v''(m^*)(p + 1)^{2/3}(C + \sigma W)^2}{2v(m^*)}, \end{aligned}$$

where $W \sim N(0, 1)$. The proof of the first statement is complete.

Let us prove by contradiction that the rate $T^{2/5}$ is optimal. Suppose there are $\delta > 0$, $0 < \alpha \leq 1$ and $-1/2 < \beta < \infty$ such that

$$T^{2/5+\delta}\hat{E}_{K,m_0} = O_P(1)$$

with $K \propto T^\alpha$ and $m_0 \propto n_0^{1/2+\beta}$, where $\propto$ denotes proportionality. There are three cases to consider: $\sqrt{n_0}/m_0 \rightarrow c' > 0$, $\sqrt{n_0}/m_0 \rightarrow 0$ or $\sqrt{n_0}/m_0 \rightarrow \infty$.

Note that K/T must converge, and recall that the limit is denoted by $c/(c+1) \geq 0$. Note also that

$$m_0^2 \propto T^{\frac{2\alpha(1+2\beta)}{3+2\beta}}, m_0\sqrt{n_0} \propto T^{2\alpha(1+\beta)/(3+2\beta)} \quad \text{and} \quad n_0 \propto T^{2\alpha/(3+2\beta)}.$$

Put $A = (c+1)v''(m^*)/v(m^*)$.

If $\sqrt{n_0}/m_0 \rightarrow c' > 0$ then $\beta = 0$ and, by Lemma 1, we have

$$(19) \quad T^{2/5+\delta} \widehat{E}_{K,m_0} = \frac{T^{2/5+\delta}}{m_0\sqrt{n_0}} U_K^2 + \frac{T^{2/5+\delta} K}{T-K},$$

where (U_K^2) is sequence of random variables such that

$$U_K^2 \xrightarrow{d} \frac{A}{2} \left(\sqrt{c'} C + \frac{1}{\sqrt{c'}} \sigma W \right)^2,$$

for some $W \sim N(0,1)$. Since both terms in the right hand side of (19) are positive and U_K^2 goes to a positive random variable, we must have

$$\frac{2}{5} + \delta - \frac{2\alpha(1+\beta)}{3+2\beta} \leq 0 \quad \text{and} \quad \frac{2}{5} + \delta + \alpha - 1 \leq 0.$$

But this leads to a contradiction because the first inequality implies that the second is false.

If $\sqrt{n_0}/m_0 \rightarrow 0$, then $\beta > 0$ and

$$T^{2/5+\delta} \widehat{E}_{K,m_0} = \frac{T^{2/5+\delta}}{n_0} U_K^2 + T^{2/5+\delta} \frac{K}{T-K}$$

with

$$U_K^2 \xrightarrow{d} \frac{A\sigma^2 W^2}{2}.$$

We proceed as before. We must have

$$\frac{2}{5} + \delta - \frac{2\alpha}{3+2\beta} \leq 0 \quad \text{and} \quad \frac{2}{5} + \delta + \alpha - 1 \leq 0.$$

But the first inequality implies that the second is false.

If $\sqrt{n_0}/m_0 \rightarrow \infty$, then $-1/2 < \beta < 0$ and

$$T^{2/5+\delta} \widehat{E}_{K,m_0} = \frac{T^{2/5+\delta}}{m_0^2} U_K^2 + T^{2/5+\delta} \frac{K}{T-K}$$

with

$$U_K^2 \xrightarrow{d} \frac{AC^2}{2}.$$

Again, we must have

$$\frac{2}{5} + \delta - \frac{2\alpha(1+2\beta)}{(3+2\beta)} \leq 0 \quad \text{and} \quad \frac{2}{5} + \delta + \alpha - 1 \leq 0,$$

where the first inequality implies the second is false. The proof of the second statement is complete.

B A unified treatment of the asymptotics

All estimators in this paper have a common form, given by

$$(20) \quad \frac{1}{n} \sum_{i=1}^n \xi_{m,i},$$

with

$$(21) \quad \xi_{m,i} = \prod_{l=1}^L \frac{1}{m} \sum_{k=1}^m \prod_{j=0}^p Y_j^{(i,k) b_{j;l}},$$

where $Y_0^{(i,k)} = Y^{(i,k)} = f(X^{(i)}, Z_0^{(i,k)})$, $Y_j^{(i,k)} = f(\tilde{X}_{-j}^{(i)}, Z_j^{(i,k)})$ for $j = 1, \dots, p$, and $b_{j;l}$, $j = 0, \dots, p$, $l = 1, \dots, L$, are nonnegative coefficients. The coefficients are arranged in a matrix $(b_{j;l})$ with L rows and $p+1$ columns, where $b_{j;l}$ is the element in the l th row and $(j+1)$ th column. This way, all estimators of the form (20) and (21), or, equivalently, all summands (21), can be identified with a matrix. For instance, we have

$$\frac{1}{n} \sum_{i=1}^n \frac{1}{m} \sum_{k=1}^m Y_0^{(i,k)} \frac{1}{m} \sum_{k'=1}^m Y_j^{(i,k')} \simeq \begin{pmatrix} 1 & 0 & \cdots & 0 & 0 & 0 & \cdots & 0 \\ 0 & 0 & \cdots & 0 & 1 & 0 & \cdots & 0 \end{pmatrix},$$

where the non-null columns are the first and the $(j+1)$ th ones. This identity is reduced to

$$\frac{1}{n} \sum_{i=1}^n \frac{1}{m} \sum_{k=1}^m Y_0^{(i,k)} \frac{1}{m} \sum_{k'=1}^m Y_j^{(i,k')} \simeq \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix},$$

where it is implicitly understood that we have taken the first and $(j+1)$ th columns only, the remaining columns being zero. Note that the identification of an estimator to a matrix is done up to permutations between the rows. Thus, we also have

$$\frac{1}{n} \sum_{i=1}^n \frac{1}{m} \sum_{k=1}^m Y_0^{(i,k)} \frac{1}{m} \sum_{k'=1}^m Y_j^{(i,k')} \simeq \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \simeq \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}.$$

Proceeding the same way, we get

$$\frac{1}{n} \sum_{i=1}^n \frac{1}{m} \sum_{k=1}^m Y_0^{(i,k)} \simeq \begin{pmatrix} 1 & 0 \\ 0 & 0 \end{pmatrix}, \quad \frac{1}{n} \sum_{i=1}^n \frac{1}{m} \sum_{k=1}^m Y_0^{(i,k)^2} \simeq \begin{pmatrix} 2 & 0 \\ 0 & 0 \end{pmatrix}$$

and

$$\frac{1}{n} \sum_{i=1}^n \left(\frac{1}{m} \sum_{k=1}^m Y_0^{(i,k)} \right)^2 \simeq \begin{pmatrix} 1 & 0 \\ 1 & 0 \end{pmatrix}.$$

The above matrices can always be completed with rows full of zeros whenever the dimensions of the matrices need to match (mainly for comparison purposes).

For each n , the random variables $\xi_{m,1}, \dots, \xi_{m,n}$ are independent and identically distributed. Denote by $\mathcal{E}_{m,i}(L)$ the set of all summands (21). In other words, $\mathcal{E}_{m,i}(L)$ is the set of all nonnegative matrices of size $L \times (p+1)$. This set has useful properties, gathered in Proposition 3 for subsequent use.

Proposition 3. Let $L \geq 0$ and take $(b_{j;l}) \simeq \xi \in \mathcal{E}_{m,i}(L)$. The following statements are true.

(i) If $(b'_{j;l}) \simeq \xi' \in \mathcal{E}_{m,i}(L)$, then

$$\xi \xi' \simeq \begin{pmatrix} b_{0;1} & \cdots & b_{p;1} \\ \vdots & & \vdots \\ b_{0;L} & \cdots & b_{p;L} \\ b'_{0;1} & \cdots & b'_{p;1} \\ \vdots & & \vdots \\ b'_{0;L} & \cdots & b'_{p;L} \end{pmatrix}$$

belongs to $\mathcal{E}_{m,i}(2L)$.

(ii) There exists a real θ such that $E \xi \rightarrow \theta$ whenever $m = m_n \rightarrow \infty$ as $n \rightarrow \infty$.

(iii) Under the assumption $f(X, Z) \leq F(X)$ for all X and Z , it holds that

$$|\xi| \leq \left(\bigvee_{j=0}^p F_j(\mathbf{X}^{(i)}) \right)^{\sum_{j=0}^p \sum_{l=1}^L b_{j;l}},$$

where $F_j(\mathbf{X}^{(i)})$ is $F(X^{(i)}) \vee 1$ if $j = 0$ and $F(\tilde{X}_{-j}^{(i)}) \vee 1$ if $j \geq 1$.

Proof. The proof of (i) is trivial. Let us prove (ii). Put

$$a(k_1, \dots, k_L) = E \prod_{l=1}^L \prod_{j=0}^p Y_j^{(1, k_l) b_{j;l}}.$$

Then

$$(22) \quad E \xi = \frac{1}{m^L} \sum_{(k_1, \dots, k_L) \in \{1, \dots, m\}^L} a(k_1, \dots, k_L).$$

This sum can be decomposed into a finite number of sub-sums such that each sub-sum has a polynomial number of equal terms. The first sub-sum is built as follows. Take $(k_1, \dots, k_L)$ all distinct and $m \geq L$. Let $\mathbf{X}^{(1)} = (X^{(1)}, \tilde{X}^{(1)})$ and $\mathbf{Z}^{(1,k)} = (Z_0^{(1,k)}, Z_1^{(1,k)}, \dots, Z_p^{(1,k)})$ for $k = 1, \dots, m$. Since $\mathbf{X}^{(1)}$ and $\{\mathbf{Z}^{(1,k)}, k = 1, \dots, m\}$ are independent, and since the law of $(\mathbf{Z}^{(1,k_1)}, \dots, \mathbf{Z}^{(1,k_L)})$ is invariant through any permutation of $(k_1, \dots, k_L)$, the corresponding coefficients $a(k_1, \dots, k_L)$ are all equal. Thus, by summing over all permutations of $(k_1, \dots, k_L)$ for distinct $k_1, \dots, k_L$, we get that the sub-sum is

$$m(m-1) \cdots (m-L+1) a(1, \dots, L).$$

For the second sub-sum, take $(k_1, \dots, k_L)$ such that exactly two of the indices are equal. The same arguments can be used to get a polynomial in m with degree at most L . Go on until all the sub-sums are calculated. An example is given in the proof of Lemma 5 in Appendix C. The result is a polynomial in m with degree at most L , and hence, divided by m^L , the expectation in (22) must converge to some constant (possibly zero).

To prove (iii), simply remember that, by assumption, $|Y_j^{(1,k)}| \leq F(X^{(1)})$ and $|Y_j^{(1,k)}| \leq F(\tilde{X}_{-j}^{(1)})$ for all k and all j . $\square$

Lemma 3. Take $\xi_{m,i}^{(I)} \simeq (b_{j;l}^{(I)})$, $I = 1, \dots, N$ in $\mathcal{E}_{m,i}(L)$ for $i = 1, 2, \dots$ and $m = 1, 2, \dots$. Assume

$$\mathbb{E} F(X^{(1)})^2 \sum_{j=0}^p \sum_{l=1}^L b_{j;l}^{(I)} < \infty$$

for all $I = 1, \dots, N$. Let $n \rightarrow \infty$. Then

$$\sqrt{n} \left[\frac{1}{n} \sum_{i=1}^n \xi_{m,i}^{(1)} - \mathbb{E} \xi_{m,1}^{(1)}, \dots, \frac{1}{n} \sum_{i=1}^n \xi_{m,i}^{(N)} - \mathbb{E} \xi_{m,1}^{(N)} \right]^\top \xrightarrow{d} N(0, \Sigma_m),$$

where Σ_m is the variance-covariance matrix of $\boldsymbol{\xi}_{m,i} = (\xi_{m,i}^{(1)}, \dots, \xi_{m,i}^{(N)})$. If $m = m_n \rightarrow \infty$, then $\Sigma_m \rightarrow \Sigma$ elementwise for some variance-covariance matrix Σ and the convergence in distribution still holds with Σ in place of Σ_m .

Proof. For any $I = 1, \dots, N$, by Proposition 3 (i), $\xi_{m,i}^{(I)2}$ belongs to $\mathcal{E}_{m,i}(2L)$ and has coefficients

$$\xi_{m,i}^{(I)2} \simeq \begin{pmatrix} b_{0;1} & \cdots & b_{p;1} \\ \vdots & & \vdots \\ b_{0;L} & \cdots & b_{p;L} \\ b_{0;1} & \cdots & b_{p;1} \\ \vdots & & \vdots \\ b_{0;L} & \cdots & b_{p;L} \end{pmatrix}.$$

Thus, denoting $\sum_{j=0}^p \sum_{l=1}^L b_{j;l}$ by β , Proposition 3 (iii) yields

$$\mathbb{E} \xi_{m,1}^{(I)2} = \mathbb{E} \bigvee_{j=0}^p F_j(\mathbf{X}^{(1)})^{2\beta} \leq (p+1) \mathbb{E} \left(1 \vee F(X^{(1)}) \right)^{2\beta},$$

which is finite by assumption. Therefore, Σ_m exists and we can apply the central limit theorem, which concludes the proof for m fixed.

Let $m = m_n \rightarrow \infty$. According to Lindeberg-Feller's central limit theorem (see e.g. [23]), it suffices to show

(i) for all $\epsilon > 0$,

$$\sum_{i=1}^n \mathbb{E} \left\| \frac{1}{\sqrt{n}} \boldsymbol{\xi}_{m_n,i} \right\|^2 \mathbf{1} \left\{ \left\| \frac{1}{\sqrt{n}} \boldsymbol{\xi}_{m_n,i} \right\| > \epsilon \right\} \rightarrow 0,$$

and

(ii) there exists Σ such that

$$\sum_{i=1}^n \text{Cov} \left(\frac{1}{\sqrt{n}} \boldsymbol{\xi}_{m_n,i} \right) \rightarrow \Sigma.$$

Let us show (i). We have

$$\begin{aligned} \sum_{i=1}^n \mathbb{E} \left\| \frac{\boldsymbol{\xi}_{m_n,i}}{\sqrt{n}} \right\|^2 \mathbf{1} \{ \|\boldsymbol{\xi}_{m_n,i}\| > \sqrt{n}\epsilon \} &= \mathbb{E} \|\boldsymbol{\xi}_{m_n,1}\|^2 \mathbf{1} \{ \|\boldsymbol{\xi}_{m_n,1}\| > \sqrt{n}\epsilon \} \\ &= \mathbb{E} \sum_{I=1}^N \xi_{m_n,1}^{(I)2} \mathbf{1} \{ \|\boldsymbol{\xi}_{m_n,1}\| > \sqrt{n}\epsilon \}. \end{aligned}$$

Write $\mathbf{X} = (X^{(1)}, \tilde{X}^{(1)})$. By Proposition 3 (iii), there exists an integrable function $G(\mathbf{X})$ such that

$$\mathbb{E} \left(\xi_{m_n,1}^{(I)2} \mathbf{1}_{\{\|\xi_{m_n,1}\| > \sqrt{n}\epsilon\}} | \mathbf{X} \right) \leq G(\mathbf{X}) P(\|\xi_{m_n,1}\| > \sqrt{n}\epsilon | \mathbf{X}).$$

By Proposition 3 (ii) and Chebyshev's inequality, the upper bound goes to zero and is dominated by $G(\mathbf{X})$, which is integrable. Thus, we can apply the dominated convergence theorem to complete the proof.

Let us show that (ii) holds. We have $\sum_{i=1}^n \text{Cov}(\xi_{m_n,i}/\sqrt{n}) = \text{Cov}(\xi_{m_n,1})$. The element (I, J) in this matrix is given by $\mathbb{E} \xi_{m_n,1}^{(I)} \xi_{m_n,1}^{(J)} - \mathbb{E} \xi_{m_n,1}^{(I)} \mathbb{E} \xi_{m_n,1}^{(J)}$. But since $\xi_{m_n,1}^{(I)} \xi_{m_n,1}^{(J)} \in \mathcal{E}$, the limit of $\text{Cov} \xi_{m_n,1}$ exists and is finite. The proof is complete. $\square$

C Explicit moment calculations

Explicit moment calculations are given for the summands in the proof of Theorem 1. In this section, $\mathbb{E} f(X, Z)$ and $\mathbb{E} \mathbb{E}[f(X, Z) | X]^2$ are denoted by μ and D , respectively. Recall that the upper-left term in (2) and (3) is denoted by D_j . The moments are given in Lemma 4 and Lemma 5. The variances and covariances are given in Lemma 6. Let $\mathbf{X} = (X^{(1)}, \tilde{X}^{(1)})$. Whenever there is a superscript $\mathbf{X}$ added to the expectation symbol $\mathbb{E}$ or the variance symbol Var , this means that these operators are to be understood conditionally on $\mathbf{X}$. An integral with respect to $P(d\mathbf{x})$ means that we integrate with respect to the law of $\mathbf{X}$.

Lemma 4 (Moments of order 1). *The moments of order 1 are given by*

$$\begin{aligned} \mathbb{E} \xi_{m1}^{UL} &= D_j, \\ \mathbb{E} \xi_{m1}^{UR} &= \mu, \\ \mathbb{E} \xi_{m1}^{\prime LL} &= \frac{1}{m} \mathbb{E} \text{Var}^X f(X^{(1)}, Z^{(1,1)}) + D. \end{aligned}$$

Proof. One has

$$\begin{aligned} \mathbb{E} \xi_{m1}^{UL} &= \frac{1}{m^2} \sum_{k,k'} \mathbb{E} f(X^{(1)}, Z^{(1,k)}) f(\tilde{X}_{-j}^{(1)}, Z_j^{(1,k')}) \\ &= \frac{1}{m^2} \sum_{k,k'} \int \mathbb{E} f(x, Z^{(1,k)}) f(\tilde{x}_{-j}, Z_j^{(1,k')}) P(d\mathbf{x}) \\ &= \mathbb{E} f(X^{(1)}, Z^{(1,1)}) f(\tilde{X}_{-j}^{(1)}, Z_j^{(1,1)}) \\ &= D_j, \end{aligned}$$

where the integral is taken with respect to the law of $\mathbf{x} = (x, \tilde{x})$, and,

$$\begin{aligned} \mathbb{E} \xi_{m1}^{\prime LL} &= \frac{1}{m^2} \sum_{k,k'} \mathbb{E} f(X^{(1)}, Z^{(1,k)}) f(X^{(1)}, Z^{(1,k')}) \\ &= \frac{1}{m} \mathbb{E} \text{Var}^X f(X, Z) + \mathbb{E} (\mathbb{E}^X f(X, Z))^2 \\ &= \frac{1}{m} \mathbb{E} \text{Var}^X f(X, Z) + D. \end{aligned}$$

The proof for ξ_{m1}^{UR} is similar. $\square$

Lemma 5 (Moments of order 2). *The moments of order 2 are given by*

$$\begin{aligned}
\mathbb{E} \xi_{m1}^{(UL)2} &= \text{Var} \mathbb{E}^{\mathbf{X}} f(X^{(1)}, Z^{(1,1)}) f(\tilde{X}_{-j}^{(1)}, Z_j^{(1,1)}) + D_j^2 \\
&\quad + \frac{1}{m} [\mathbb{E} \text{Var}^{\mathbf{X}} f(X^{(1)}, Z^{(1,1)}) f(\tilde{X}_{-j}^{(1)}, Z_j^{(1,1)}) \\
&\quad - \text{Var}^{\mathbf{X}} f(X^{(1)}, Z^{(1,1)}) \text{Var}^{\mathbf{X}} f(\tilde{X}_{-j}^{(1)}, Z_j^{(1,1)})] \\
&\quad + \frac{1}{m^2} \mathbb{E} \text{Var}^{\mathbf{X}} f(X^{(1)}, Z^{(1,1)}) \text{Var}^{\mathbf{X}} f(\tilde{X}_{-j}^{(1)}, Z_j^{(1,1)}), \\
\mathbb{E} \xi_{m1}^{(UR)2} &= \frac{1}{m} \mathbb{E} \text{Var}^{\mathbf{X}} f(X^{(1)}, Z^{(1,1)}) + \mathbb{E} (\mathbb{E}^{\mathbf{X}} f(X^{(1)}, Z^{(1,1)}))^2, \\
\mathbb{E} \xi_{m1}^{(LL)2} &= \frac{m(m-1)(m-2)(m-3)}{m^4} \\
&\quad \mathbb{E} f(X^{(1)}, Z^{(1,1)}) f(X^{(1)}, Z^{(1,2)}) f(X^{(1)}, Z^{(1,3)}) f(X^{(1)}, Z^{(1,4)}) \\
&\quad + \frac{\binom{4}{2} m(m-1)(m-2)}{m^4} \mathbb{E} f(X^{(1)}, Z^{(1,1)})^2 f(X^{(1)}, Z^{(1,2)}) f(X^{(1)}, Z^{(1,3)}) \\
&\quad + \frac{\binom{4}{3} m(m-1)}{m^4} \mathbb{E} f(X^{(1)}, Z^{(1,1)})^3 f(X^{(1)}, Z^{(1,2)}) \\
&\quad + \frac{m}{m^4} \mathbb{E} f(X^{(1)}, Z^{(1,1)})^4 \\
&\quad + \frac{\binom{4}{2} m(m-1)/2}{m^4} \mathbb{E} f(X^{(1)}, Z^{(1,1)})^2 f(X^{(1)}, Z^{(1,2)})^2
\end{aligned}$$

Proof. Let us first deal with ξ_{m1}^{UL} . We have

$$\mathbb{E} \xi_{m1}^{(\text{UL})2} = \frac{1}{m^4} \sum_{k_1, k_2, k_3, k_4} \mathbb{E} f(X^{(1)}, Z^{(1, k_1)}) f(X^{(1)}, Z^{(1, k_2)}) f(\tilde{X}_{-j}^{(1)}, Z_j^{(1, k_3)}) f(\tilde{X}_{-j}^{(1)}, Z_j^{(1, k_4)})$$

where, in the sum, the indices run over $1, \dots, m$. We split the sum into four parts. The first contains the $m^2(m-1)^2$ terms that satisfy $k_1 \neq k_3$ and $k_2 \neq k_4$. In this part, all the terms are equal to

$$(\text{term 1}) \quad \mathbb{E} \left(\mathbb{E}^{\mathbf{X}} f(X^{(1)}, Z^{(1,1)}) f(\tilde{X}_{-j}^{(1)}, Z_j^{(1,1)}) \right)^2.$$

The second part contains the $m^2(m-1)$ terms that satisfy $k_1 \neq k_3$ and $k_2 = k_4$ and that are equal to

$$(\text{term 2}) \quad \mathbb{E} f(X^{(1)}, Z^{(1,1)}) f(X^{(1)}, Z^{(1,2)}) f(\tilde{X}_{-j}^{(1)}, Z_j^{(1,1)})^2.$$

The third part contains the $m^2(m-1)$ terms that satisfy $k_1 = k_3$ and $k_2 \neq k_4$ and that are equal to

$$(\text{term 3}) \quad \mathbb{E} f(X^{(1)}, Z^{(1,1)})^2 f(\tilde{X}_{-j}^{(1)}, Z_j^{(1,1)}) f(\tilde{X}_{-j}^{(1)}, Z_j^{(1,2)}).$$

Finally, the fourth part contains the m^2 terms that satisfy $k_1 = k_4$ and $k_2 = k_4$ and are equal to

$$(\text{term 4}) \quad \mathbb{E} f(X^{(1)}, Z^{(1,1)})^2 f(\tilde{X}_{-j}^{(1)}, Z_j^{(1,1)})^2.$$

(One can see that the number of terms is m^4 .) Thus,

$$\begin{aligned} \mathbb{E} \xi_{m1}^{(\text{UL})2} = & (\text{term 1}) \\ & + \frac{1}{m} [(\text{term 2}) + (\text{term 3}) - 2(\text{term 1})] \\ & + \frac{1}{m^2} [(\text{term 1}) - (\text{term 2}) - (\text{term 3}) + (\text{term 4})]. \end{aligned}$$

Furthermore, $[(\text{term 1}) - (\text{term 2}) - (\text{term 3}) + (\text{term 4})]$ is equal to

$$\begin{aligned} & \int \left(\mathbb{E}^{\mathbf{X}} f(x, Z) f(\tilde{x}_{-j}, Z_j) \right)^2 \\ & - \mathbb{E}^{\mathbf{X}} f(x, Z^{(1,1)}) f(x, Z^{(1,2)}) f(\tilde{x}_{-j}, Z_j^{(1,1)})^2 \\ & - \mathbb{E}^{\mathbf{X}} f(x, Z^{(1,1)})^2 f(\tilde{x}_{-j}, Z_j^{(1,1)}) f(\tilde{x}_{-j}, Z_j^{(1,2)}) \\ & + \mathbb{E}^{\mathbf{X}} f(x, Z^{(1,1)})^2 f(\tilde{x}_{-j}, Z_j^{(1,1)})^2 dP(\mathbf{x}) \\ & = \int \left(\mathbb{E}^{\mathbf{X}} f(x, Z) \right)^2 \left(\mathbb{E}^{\mathbf{X}} f(\tilde{x}_{-j}, Z_j) \right)^2 \\ & - \left(\mathbb{E}^{\mathbf{X}} f(x, Z) \right)^2 \mathbb{E}^{\mathbf{X}} f(\tilde{x}_{-j}, Z_j)^2 \\ & - \mathbb{E}^{\mathbf{X}} f(x, Z)^2 \left(\mathbb{E}^{\mathbf{X}} f(\tilde{x}_{-j}, Z_j) \right)^2 \\ & + \mathbb{E}^{\mathbf{X}} f(x, Z)^2 \mathbb{E}^{\mathbf{X}} f(\tilde{x}_{-j}, Z_j)^2 dP(\mathbf{x}) \\ & = \int \text{Var}^{\mathbf{X}} f(X, Z) \text{Var}^{\mathbf{X}} f(\tilde{X}_{-j}, Z_j) dP(\mathbf{x}). \end{aligned}$$

Likewise, we find that $[(\text{term 2}) + (\text{term 3}) - 2(\text{term 1})]$ is equal to

$$\mathbb{E} \text{Var}^{\mathbf{X}} f(X, Z) f(\tilde{X}_{-j}, Z_j) - \text{Var}^{\mathbf{X}} f(X, Z) \text{Var}^{\mathbf{X}} f(\tilde{X}_{-j}, Z_j),$$

and term 1 is $\text{Var} \mathbb{E}^{\mathbf{X}} f(X, Z) f(\tilde{X}_{-j}, \tilde{Z}) + D_j^2$.

We now deal with $\xi_{m1}^{\prime\prime\text{LL}}$. We have

$$\mathbb{E} \xi_{m1}^{\prime\prime\text{LL}2} = \frac{1}{m^4} \sum_{k_1, k_2, k_3, k_4} \mathbb{E} f(X^{(1)}, Z^{(1, k_1)}) f(X^{(1)}, Z^{(1, k_2)}) f(X^{(1)}, Z^{(1, k_3)}) f(X^{(1)}, Z^{(1, k_4)}).$$

The sum is split into five parts. The first part consists of the $m(m-1)(m-2)(m-3)$ terms with different indices; those terms are equal to

$$\mathbb{E} f(X^{(1)}, Z^{(1,1)}) f(X^{(1)}, Z^{(1,2)}) f(X^{(1)}, Z^{(1,3)}) f(X^{(1)}, Z^{(1,4)}).$$

The second part consists of the $\binom{4}{2} m(m-1)(m-2)$ terms with exactly two equal indices; those terms are equal to

$$\mathbb{E} f(X^{(1)}, Z^{(1,1)})^2 f(X^{(1)}, Z^{(1,2)}) f(X^{(1)}, Z^{(1,3)}).$$

The third part consists of the $\binom{4}{3} m(m-1)$ terms with exactly three equal indices; those terms are equal to

$$\mathbb{E} f(X^{(1)}, Z^{(1,1)})^3 f(X^{(1)}, Z^{(1,2)}).$$

The fourth part consists of the m terms with exactly four equal indices; those terms are equal to

$$\mathbb{E} f(X^{(1)}, Z^{(1,1)})^4.$$

The fifth and last part consists of the $\binom{4}{2}m(m-1)/2$ terms with exactly two pairs of equal indices; those terms are equal to

$$\mathbb{E} f(X^{(1)}, Z^{(1,1)})^2 f(X^{(1)}, Z^{(1,2)})^2.$$

(One can check that the total number of terms is m^4 .) $\square$

Lemma 6 (Variances and covariances).

(i)

$$\begin{aligned} \text{Var } \xi_{m1}^{UL} &= \text{Var } \mathbb{E}^{\mathbf{X}} f(X^{(1)}, Z^{(1,1)}) f(\tilde{X}_{-j}^{(1)}, Z_j^{(1,1)}) \\ &\quad + \frac{1}{m} [\mathbb{E} \text{Var}^{\mathbf{X}} f(X^{(1)}, Z^{(1,1)}) f(\tilde{X}_{-j}^{(1)}, Z_j^{(1,1)}) - \text{Var}^{\mathbf{X}} f(X^{(1)}, Z^{(1,1)}) \text{Var}^{\mathbf{X}} f(\tilde{X}_{-j}^{(1)}, Z_j^{(1,1)})] \\ &\quad + \frac{1}{m^2} \mathbb{E} \text{Var}^{\mathbf{X}} f(X^{(1)}, Z^{(1,1)}) \text{Var}^{\mathbf{X}} f(\tilde{X}_{-j}^{(1)}, Z_j^{(1,1)}), \end{aligned}$$

(ii)

$$\begin{aligned} \text{Cov}(\xi_{m1}^{UL}, \xi_{m1}^{UR}) &= \frac{m-1}{m} \mathbb{E} f(X^{(1)}, Z^{(1,1)}) f(X^{(1)}, Z^{(1,2)}) f(\tilde{X}_{-j}^{(1)}, Z_j^{(1,1)}) \\ &\quad + \frac{1}{m} \mathbb{E} f(X^{(1)}, Z^{(1,1)})^2 f(\tilde{X}_{-j}^{(1)}, Z_j^{(1,1)}) - D_j \mu \end{aligned}$$

(iii)

$$\text{Cov}(\xi_{m1}^{UL}, f(X, Z)^2) = \frac{1}{m} \mathbb{E} f(X^{(1)}, Z^{(1,1)})^3 f(\tilde{X}_{-j}^{(1)}, Z_j^{(1,1)})$$

(iii)

$$+ \frac{m-1}{m} \mathbb{E} f(X^{(1)}, Z^{(1,1)})^2 f(X^{(1)}, Z^{(1,2)}) f(\tilde{X}_{-j}^{(1)}, Z_j^{(1,1)}) - D_j \kappa$$

(iv)

$$\text{Var } \xi_{m1}^{UR} = \frac{1}{m} \text{Var } f(X, Z)$$

(v)

$$\begin{aligned} \text{Cov}(\xi_{m1}^{UR}, f(X, Z)^2) &= \frac{1}{m} f(X, Z)^3 \\ &\quad + \frac{m-1}{m} \mathbb{E} f(X^{(1)}, Z^{(1,1)})^2 f(X^{(1)}, Z^{(1,2)}) - \mu \kappa \end{aligned}$$

(vi)

$$\begin{aligned} \text{Cov}(\xi_{m1}^{UL}, \xi_{m1}^{ULL}) &= \frac{m}{m^3} \mathbb{E} f(X^{(1)}, Z^{(1,1)})^3 f(\tilde{X}_{-j}^{(1)}, Z_j^{(1,1)}) \\ &\quad + \frac{3m(m-1)}{m^3} \mathbb{E} f(X^{(1)}, Z^{(1,1)})^2 f(X^{(1)}, Z^{(1,2)}) f(\tilde{X}_{-j}^{(1)}, Z_j^{(1,1)}) \\ &\quad + \frac{m(m-1)(m-2)}{m^3} \mathbb{E} f(X^{(1)}, Z^{(1,1)}) f(X^{(1)}, Z^{(1,2)}) \\ &\quad \quad f(X^{(1)}, Z^{(1,3)}) f(\tilde{X}_{-j}^{(1)}, Z_j^{(1,1)}) \\ &\quad - \mathbb{E} f(X^{(1)}, Z^{(1,1)}) f(\tilde{X}_{-j}^{(1)}, Z_j^{(1,1)}) \\ &\quad \left\{ \frac{1}{m} \mathbb{E} f(X^{(1)}, Z^{(1,1)})^2 + \frac{m-1}{m} \mathbb{E} f(X^{(1)}, Z^{(1,1)}) f(X^{(1)}, Z^{(1,2)}) \right\} \end{aligned}$$

Proof. The proof follows from direct calculations. $\square$

D Calculations for the linear model

Lemma 7. Suppose that $f(X, Z) = \beta_0 + \beta_{p+1}Z + \sum_{j=1}^p \beta_j X_j$ where $X = (X_1, \dots, X_p), Z_k, \tilde{Z}_{ik}$ are independent, $\mathbb{E} X_j = \mathbb{E} Z = 0$, $\mathbb{E} X_j^2 = \mathbb{E} Z^2 = 1$, $\mathbb{E} X_j^3 = 0$, $\mathbb{E} X_j^4 = 3$. Then the squared optimal number of repetitions is given by

$$(m_i^*)^2 = \frac{\beta_{p+1}^4}{(\beta_0 + \beta_i)^2 - 2\beta_0^4 + (\sum_{j=0}^p \beta_j^2)^2}$$

and the discriminator (the upper-left term in (2) and (3)) is

$$\beta_0^2 + \beta_i^2.$$

Proof. We have

$$m_i^* = \frac{A_i + B_i + C_i + D_i}{E_i},$$

with

$$\begin{aligned} A_i &= \mathbb{E} f(X, Z_1)^2 f(\tilde{X}_{-i}, \tilde{Z}_{i1})^2 \\ B_i &= \mathbb{E} f(X, Z_1) f(\tilde{X}_{-i}, \tilde{Z}_{i1}) f(X, Z_2) f(\tilde{X}_{-i}, \tilde{Z}_{i2}) \\ C_i &= -\mathbb{E} f(X, Z_1)^2 f(\tilde{X}_{-i}, \tilde{Z}_{i1}) f(\tilde{X}_{-i}, \tilde{Z}_{i2}) \\ D_i &= -\mathbb{E} f(\tilde{X}_{-i}, \tilde{Z}_{i1})^2 f(X, Z_1) f(X, Z_2) \\ E_i &= B - [\mathbb{E} f(X, Z_1) f(\tilde{X}_{-i}, \tilde{Z}_{i1})]^2 \end{aligned}$$

where $X = (X_1, \dots, X_p), Z_k, \tilde{Z}_{ik}$ are independent, $\mathbb{E} X_j = \mathbb{E} Z = 0$, $\mathbb{E} X_j^2 = \mathbb{E} Z^2 = 1$, $\mathbb{E} X_j^3 = 0$, $\mathbb{E} X_j^4 = 3$. We deal with the case

$$f(X, Z) = \beta_0 + \beta_{p+1}Z + \sum_{j=1}^p \beta_j X_j.$$

We calculate the terms one by one as follows. We have

$$\begin{aligned} (A_{j1}) \quad A_j &= \mathbb{E} \left(\beta_0 + \sum_{j=1}^p \beta_j X_j \right)^2 \left(\beta_0 + \beta_i X_i + \sum_{j:1 \leq j \neq i} \beta_j \tilde{X}_j \right)^2 \\ (A_{j2}) \quad &+ \left(\beta_0 + \sum_{j=1}^p \beta_j X_j \right)^2 \beta_{p+1}^2 \tilde{Z}_{i1}^2 + \beta_{p+1}^4 Z_1^2 \tilde{Z}_{i1}^2 \\ (A_{j3}) \quad &+ \beta_{p+1}^2 Z_1^2 \left(\beta_0 + \beta_i X_i + \sum_{j:1 \leq j \neq i} \beta_j \tilde{X}_j \right)^2, \end{aligned}$$

where $\mathbb{E} (A2) = \beta_{p+1}^4 + \beta_{p+1}^2 \sum_{j=0}^p \beta_j^2$, $\mathbb{E} (A3) = \beta_{p+1}^2 \sum_{j=0}^p \beta_j^2$. Elementary but somewhat tedious calculations yield

$$\mathbb{E} (A1) = \beta_0^4 + 3\beta_i^4 + 6\beta_0^2 \beta_i^2 + 2(\beta_0^2 + \beta_i^2) \sum_{j:1 \leq j \neq i} \beta_j^2 + \left(\sum_{j:1 \leq j \neq i} \beta_j^2 \right)^2.$$

Similar calculations show that $B_j = A_{j1}$, $C_j = -A_{j1} - A_{j3}$, $D_j = -A_{j1} - A_{j3}$, $E_j = A_{j1} - (\beta_0^2 + \beta_i^2)^2$. Thus,

$$(m_i^*)^2 = \frac{\beta_{p+1}^4}{(\beta_0 + \beta_i)^2 - 2\beta_0^4 + (\sum_{j=0}^p \beta_j^2)^2}.$$

□