



FiNANCE FOR ENERGY MARKET RESEARCH CENTRE



A Non-Intrusive Stratified Resampler for Regression Monte Carlo: Application to Solving Non-Linear Equations

EMMANUEL GOBET, GANG LIU AND JORGE P. ZUBELLI

**Working Paper
RR-FiME-16-10**

September 2016

A NON-INTRUSIVE STRATIFIED RESAMPLER FOR REGRESSION MONTE CARLO: APPLICATION TO SOLVING NON-LINEAR EQUATIONS*

EMMANUEL GOBET[†], GANG LIU[†], AND JORGE P. ZUBELLI[‡]

Abstract. Our goal is to solve certain dynamic programming equations associated to a given Markov chain X , using a regression-based Monte Carlo algorithm. More specifically, we assume that the model for X is not known in full detail and only a root sample $X^1, \dots, X^M$ of such process is available. By a stratification of the space and a suitable choice of a probability measure ν , we design a new resampling scheme that allows to compute local regressions (on basis functions) in each stratum. The combination of the stratification and the resampling allows to compute the solution to the dynamic programming equation (possibly in large dimensions) using only a relatively small set of root paths. To assess the accuracy of the algorithm, we establish non-asymptotic error estimates in $L^2(\nu)$. Our numerical experiments illustrate the good performance, even with $M = 20 - 40$ root paths.

Key words. discrete Dynamic Programming Equations, empirical regression scheme, resampling methods, small-size sample

AMS subject classifications. 62G08, 62G09, 93Exx

1. Introduction. Stochastic dynamic programming equations are classic equations arising in the resolution of nonlinear evolution equations, like in stochastic control (see [18, 4]) or non-linear PDEs (see [6, 9]). In a discrete-time setting they take the form:

$$Y_N = g_N(X_N), \quad Y_i = \mathbb{E}[g_i(Y_{i+1}, \dots, Y_N, X_i, \dots, X_N) \mid X_i], \quad i = N-1, \dots, 0,$$

for some functions g_N and g_i which depend on the non-linear problem under consideration. Here $X = (X_0, \dots, X_N)$ is a Markov chain valued in $\mathbb{R}^d$, entering also in the definition of the problem. The aim is to compute the value function y_i such that $Y_i = y_i(X_i)$.

Among the popular methods to solve this kind of problem, we are concerned with Regression Monte Carlo (RMC) methods that take as input M simulated paths of X , say $(X^1, \dots, X^M) =: X^{1:M}$, and provide as output simulation-based approximations $y_i^{M,\mathcal{L}}$ using Ordinary Least Squares (OLS) within a vector space of functions $\mathcal{L}$:

$$y_i^{M,\mathcal{L}} = \arg \inf_{\varphi \in \mathcal{L}} \frac{1}{M} \sum_{m=1}^M \left| g_i(y_{i+1}^{M,\mathcal{L}}(X_{i+1}^m), \dots, y_N^{M,\mathcal{L}}(X_N^m), X_i^m, \dots, X_N^m) - \varphi(X_i^m) \right|^2.$$

This Regression Monte Carlo methodology has been investigated in [9] to solve Backward Stochastic Differential Equations associated to semi-linear partial differential equations (PDEs) [16], with some tight error estimates. Generally speaking, it is well known that the number of simulations M has to be much larger than the dimension of the vector space $\mathcal{L}$ and thus the number of coefficients we are seeking.

*This work is part of the Chair *Financial Risks* of the *Risk Foundation*, the *Finance for Energy Market Research Centre* and the ANR project *CAESARS* (ANR-15-CE05-0024).

[†]Centre de Mathématiques Appliquées (CMAP), Ecole Polytechnique and CNRS, Université Paris-Saclay, Route de Saclay, 91128 Palaiseau Cedex, France (Email: emmanuel.gobet@polytechnique.edu (Corresponding author), gang.liu1988@gmail.com).

[‡]IMPA, Est. D. Castorina 110, Rio de Janeiro, RJ 22460-320, Brazil (Email: zubelli@impa.br).

In contradistinction, throughout this work, we focus on the case where M is relatively small (a few hundreds) and the simulations are not sampled by the user but are directly taken from historical data ($X^{1:M}$ is called **root sample**), in the spirit of [17]. This is the most realistic situation when we collect data and when the model which fits the data is unknown.

Thus, as a main difference with the aforementioned references:

- We do not assume that we have full information about the model for X and we do not assume that we can generate as many simulations as needed to have convergent Regression Monte Carlo methods.
- The size M of the learning samples $X^1, \dots, X^M$ is relatively small, which discards the use of a direct RMC with large dimensional $\mathcal{L}$.

To overcome these major obstacles, we elaborate on two ingredients:

1. First, we partition $\mathbb{R}^d$ in strata $(\mathcal{H}_k)_k$, so that the regression functions can be computed locally on each stratum $\mathcal{H}_k$; for *small* stratum this allows to use only a small dimensional approximation space $\mathcal{L}_k$, and therefore it puts a lower constraint on M . In general, this stratification breaks the properties for having well-behaved error propagation and we provide a precise way to sample in order to be able to aggregate the error estimates in different strata. We use a probabilistic distribution ν that has good norm-stability properties with X (see Assumptions 3.2 and 4.2).
2. Second, by assuming a mild model condition on X , we are able to resample from the root sample of size M , a *training sample* of M simulations suitable for the stratum $\mathcal{H}_k$. This resampler is non intrusive in the sense that it only requires to know the form of the model but not its coefficients: for example, we can handle models with independent increments (discrete inhomogeneous Levy process) or Ornstein-Uhlenbeck processes. See Examples 2.1-2.2-2.3-2.4. We call this scheme NISR (Non Intrusive Stratified Resampler), it is described in Definition 2.1 and Proposition 2.1.

The resulting regression scheme is, to the best of our knowledge, completely new.

To sum up, the contributions of this work are the following:

- We design a non-intrusive stratified resample (NISR) scheme that allows to sample from M paths of the root sample restarting from any stratum $\mathcal{H}_k$. See Section 2.
- We combine this with regression Monte Carlo schemes, in order to solve one-step ahead dynamic programming equations (Section 3), discrete backward stochastic differential equations (BSDEs) and semi-linear PDEs (Section 4).
- In Theorems 3.4 and 4.1, we provide quadratic error estimates of the form

$$\begin{aligned} \text{quadratic error on } y_i \leq & \text{approximation error} + \text{statistical error} \\ & + \text{interdependency error} . \end{aligned}$$

The approximation error is related to the best approximation of y_i on each stratum $\mathcal{H}_k$, and averaged over all the strata. The statistical error is bounded by C/M with a constant C which does not depend on the number of strata: only relatively small M is necessary to get low statistical errors. This is in agreement with the motivation that the root sample has a relatively small size. The interdependency error is an unusual issue, it is related to the strong dependency between regression problems (because they all use the same root sample). The analysis as well as the framework are original. The error estimates take different forms according to the problem at hand (Section 3 or

Section 4).

- Finally we illustrate the performance of the methods on two types of examples: first, approximation of non-linear PDEs arising in reaction-diffusion biological models (Subsection 5.1) and optimal sequential decision (Subsection 5.2), where we illustrate that root samples of size $M = 20 - 40$ only can lead to remarkably accurate numerical solutions.

The paper is organized as follows. In Section 2 we present the model structure that leads to the non-intrusive stratified resampler for regression Monte Carlo (NISR), together with the stratification. Main notations will be also introduced. The algorithm is presented in a generic form of dynamic programming equations in Algorithm 1. In Section 3 we analyze the convergence of the algorithm in the case of one-step ahead dynamic programming equations (for instance optimal stopping problems). Section 4 is devoted to the convergence analysis for discrete BSDEs (probabilistic representation of semi-linear PDEs arising in stochastic control problems). Section 5 is devoted to numerical examples. Technical results are postponed to the Appendix.

2. Setting and the general algorithm.

2.1. General dynamic programming equation. Suppose we have N discrete dates, and we aim at solving numerically the following dynamic programming equation (DPE for short), written in general form:

$$Y_N = g_N(X_N), \quad Y_i = \mathbb{E}[g_i(Y_{i+1:N}, X_{i:N}) \mid X_i], \quad 0 \leq i < N.$$

Here, $(X_i)_{0 \leq i \leq N}$ is a Markov chain with state space $\mathbb{R}^d$, $(Y_i)_{0 \leq i \leq N}$ is a random process taking values in $\mathbb{R}$ and we use for convenience the generic short notation $z_{i:N} := (z_i, \dots, z_N)$. Note that the argument of the conditional expectation is path-dependent, thus allowing greater generality. Had we considered Y to be multidimensional, the subsequent algorithm and analysis would remain essentially the same.

Later (Sections 3 and 4), specific forms for g_i will be considered, depending on the model of DPE to solve at hand: it will have an impact on the error estimates that we can derive. However, the description of the algorithm can be the same for all the DPEs, as seen below, and this justifies our choice of unifying the presentation.

Thanks to the Markovian property of X , under mild assumptions we can easily prove by induction that there exists a measurable function y_i such that $Y_i = y_i(X_i)$, our aim is to compute an approximation of the value functions $y_i(\cdot)$ for all i . We assume that a bound on y_i is available.

ASSUMPTION 2.1 (A priori bound). *The solution y_i is bounded by a constant $|y_i|_\infty$.*

2.2. Model structure and root sample. We will represent $y_i(\cdot)$ through its coefficients on a vector space, and the coefficients will be computed thanks to learning samples of X .

ASSUMPTION 2.2 (Data). *We have the observation of M independent paths of X , which are denoted by $((X_i^m : 0 \leq i \leq N), 1 \leq m \leq M)$. We refer to this data as the root sample.*

For our needs, we adopt a representation of the flow of the Markov chain for different initial conditions, i.e., the Markov chain $X^{i,x}$ starting at different times $i \in \{0, \dots, N\}$ and points $x \in \mathbb{R}^d$. Namely, we write

$$(2.1) \quad X_j^{i,x} = \theta_{i,j}(x, U), \quad i \leq j \leq N,$$

where

- U is some random vector, called random source,
- $\theta_{i,j}$ are (deterministic) measurable functions.

We emphasize that, for the sake of convenience, U is the same for representing all $X_j^{i,x}, 0 \leq i \leq j \leq N, x \in \mathbb{R}^d$.

ASSUMPTION 2.3 (Noise extraction). *We assume that $\theta_{i,j}$ are known and we can retrieve the random sources $(U^1, \dots, U^M)$ associated to the root sample $X^{1:M} = (X^m : 1 \leq m \leq M)$, i.e.,*

$$X_j^m = X_j^{0,x_0^m,m} = \theta_{0,j}(x_0^m, U^m).$$

Observe that this assumption is much less stringent than identifying the distribution of the model. We exemplify this now.

EXAMPLE 2.1 (Arithmetic Brownian motion with time dependent parameters). *Let $(t_i : 0 \leq i \leq N)$ be N times and define the arithmetic Brownian motion by*

$$X_i = x_0 + \int_0^{t_i} \mu_s ds + \int_0^{t_i} \sigma_s dW_s$$

where $\mu_t \in \mathbb{R}^d, \sigma_t \in \mathbb{R}^{d \times q}, W_t \in \mathbb{R}^q$ and μ, σ are deterministic functions of time. In this case, the random source is given by

$$U := (X_{i+1} - X_i)_{0 \leq i \leq N-1}$$

and the functions by

$$\theta_{ij}(x, U) := x + \sum_{i \leq k < j} U_k.$$

This works since $U_i = \int_{t_i}^{t_{i+1}} \mu_s ds + \int_{t_i}^{t_{i+1}} \sigma_s dW_s$. The crucial point is that, in order to extract U from X , we do not assume that μ and σ are known.

EXAMPLE 2.2 (Levy process). *More generally, we can set $X_i = \mathbf{X}_{t_i}$ with a time-inhomogeneous Levy process $\mathbf{X}$. Then take*

$$U := (X_{i+1} - X_i)_{0 \leq i \leq N-1}, \quad \theta_{ij}(x, U) := x + \sum_{i \leq k < j} U_k.$$

EXAMPLE 2.3 (Geometric Brownian motion with time dependent parameters). *With the same kind of parameters as for Example 2.1, define the geometric Brownian motion (component by component)*

$$X_i = X_0 \exp \left(\int_0^{t_i} \mu_s ds + \int_0^{t_i} \sigma_s dW_s \right).$$

Then, we have that

$$U := \left(\log \left(\frac{X_{i+1}}{X_i} \right) \right)_{0 \leq i \leq N-1}, \quad \theta_{ij}(x, U) := x \prod_{i \leq k < j} \exp(U_k).$$

EXAMPLE 2.4 (Ornstein-Uhlenbeck process with time dependent parameters). *Given N times $(t_i : 0 \leq i \leq N)$, set $X_i = \mathbf{X}_{t_i}$ where $\mathbf{X}$ has the following dynamics:*

$$\mathbf{X}_t = \mathbf{x}_0 - \int_0^t A(\mathbf{X}_s - \bar{\mathbf{X}}_s) ds + \int_0^t \Sigma_s dW_s$$

where A is $d \times d$ -matrix, $\mathbf{X}_t$ and $\bar{\mathbf{X}}_t$ are in $\mathbb{R}^d$, Σ_t is a $d \times q$ -matrix, $W_t \in \mathbb{R}^q$. $\bar{\mathbf{X}}_t$ and Σ_t are both deterministic functions of time. The explicit solution is

$$\mathbf{X}_t = e^{-A(t-s)}\mathbf{X}_s + e^{-At} \int_s^t e^{Ar} (A\bar{\mathbf{X}}_r dr + \Sigma_r dW_r).$$

Assume that we know A : in this case, an observation of $X_{0:N}$ enables to retrieve the random source

$$U := \left(X_j - e^{-A(t_j - t_i)} X_i \right)_{0 \leq i \leq j \leq N}$$

and then

$$\theta_{ij}(x, U) := e^{-A(t_j - t_i)} x + U_{i,j}.$$

142 The noise extraction works since $U_{i,j} = e^{-At_j} \int_{t_i}^{t_j} e^{Ar} (A\bar{\mathbf{X}}_r dr + \Sigma_r dW_r)$.

143 As illustrated above, through Assumption 2.2, all we need to know is the general
144 structure of the Markov chain model but we do not need to estimate all the model
145 parameters, and sometimes none of them (Examples 2.1, 2.2, 2.3). Our approach is
146 non intrusive in this sense.

147 **2.3. Stratification and resampling algorithm.** On the one hand, we can rely
148 on a root sample of size M only (possibly with a relatively small M , constrained by
149 the available data), which is very little to perform accurate Regression Monte-Carlo
150 methods (usually M has to be much larger than the dimension of approximation
151 spaces, as reminded in introduction).

152 On the other hand, we are able to access the random sources so that resampling
153 the M paths is possible. The degree of freedom comes from the flexibility of initial
154 conditions (i, x) , thanks to the flow representation (2.1). We now explain how we take
155 advantage of this property.

156 The idea is to resample the model paths for different starting points in different
157 parts of the space $\mathbb{R}^d$ and on each part, we will perform a regression Monte Carlo
158 using M paths and a low-dimensional approximation space. These ingredients give
159 the ground reasons for getting accurate results.

160 Let us proceed to the details of the algorithm. We design a stratification approach:
161 suppose there exist K strata $(\mathcal{H}_k)_{1 \leq k \leq K}$ such that

$$162 \quad \mathcal{H}_k \cap \mathcal{H}_l = \emptyset \quad \text{for } k \neq l, \quad \bigcup_{k=1}^K \mathcal{H}_k = \mathbb{R}^d.$$

163 An example for $\mathcal{H}_k$ is a hypercube of the form $\mathcal{H}_k = \prod_{l=1}^d [x_{k,l}^-, x_{k,l}^+)$. Then, we are
164 given a probability measure ν on $\mathbb{R}^d$ and denote its restriction on $\mathcal{H}_k$ by

$$165 \quad \nu_k(dx) := \frac{1}{\nu(\mathcal{H}_k)} 1_{\mathcal{H}_k}(x) \nu(dx).$$

166 The measure ν will serve as a reference to control the errors. See Paragraph 3.1.2 and
167 Section 5 for choices of ν .

168 **DEFINITION 2.1** (Non-intrusive stratified resampler, NISR for short). *We define the*
169 *M -sample used for regression at time i and in the k -th stratum $\mathcal{H}_k$:*

- 170 • let $(X_i^{i,k,m})_{1 \leq m \leq M}$ be an i.i.d. sample according to the law ν_k ;

- for $j = i + 1, \dots, N$, set

$$X_j^{i,k,m} = \theta_{i,j}(X_i^{i,k,m}, U^m),$$

where $U^{1:M}$ are the random sources from Assumption 2.3.

In view of Assumptions 2.2 and 2.3, the random sources $U^1, \dots, U^M$ are independent, therefore we easily prove the following.

PROPOSITION 2.1. *The M paths $(X_{i:N}^{i,k,m}, 1 \leq m \leq M)$ are independent and identically distributed as $X_{i:N}$ with $X_i \stackrel{d}{\sim} \nu_k$.*

2.4. Approximation spaces and regression Monte Carlo schemes. On each stratum, we approximate the value functions y_i using basis functions. We can take different kinds of basis functions:

- **LP₀** (partitioning estimate): $\mathcal{L}_k = \text{span}(1_{\mathcal{H}_k})$,
- **LP₁** (piecewise linear): $\mathcal{L}_k = \text{span}(1_{\mathcal{H}_k}, x_1 1_{\mathcal{H}_k}, \dots, x_d 1_{\mathcal{H}_k})$,
- **LP_n** (piecewise polynomial): $\mathcal{L}_k = \text{span}(\text{all the polynomials of degree less than or equal to } n \text{ on } \mathcal{H}_k)$.

To simplify the presentation, we assume hereafter that the dimension of $\mathcal{L}_k$ does not depend on k , we write

$$\dim(\mathcal{L}_k) =: \dim(\mathcal{L}).$$

To compute the approximation of y_i on each stratum $\mathcal{H}_k$, we will use the M samples of Definition 2.1. Our NISR-regression Monte Carlo algorithm takes the form:

Algorithm 1 General NISR-regression Monte Carlo algorithm

```

1: set  $y_N^{(M)}(\cdot) = g_N(\cdot)$ 
2: for  $i = N - 1$  until 0 do
3:   for  $k = 1$  until  $K$  do
4:     sample  $(X_{i:N}^{i,k,m})_{1 \leq m \leq M}$  using the NISR (Definition 2.1)
5:     set  $S^{(M)}(x_{i:N}) = g_i(y_{i+1}^{(M)}(x_{i+1}), \dots, y_N^{(M)}(x_N), x_{i:N})$ 
6:     compute  $\psi_i^{(M),k} = \text{OLS}(S^{(M)}, \mathcal{L}_k, X_{i:N}^{i,k,1:M})$ 
7:     set  $y_i^{(M),k} = \mathcal{T}_{|y_i|_\infty}(\psi_i^{(M),k})$  where  $\mathcal{T}_L$  is the truncation operator,
8:     defined by  $\mathcal{T}_L(x) = -L \vee x \wedge L$ 
9:   end for
10:  set  $y_i^{(M)} = \sum_{k=1}^K y_i^{(M),k} 1_{\mathcal{H}_k}$ 
11: end for
```

In the above, the Ordinary Least Squares approximation of the response function $\tilde{S} : (\mathbb{R}^d)^{N-i+1} \mapsto \mathbb{R}$ in the function space $\mathcal{L}_k$ using the M sample $X_{i:N}^{i,k,1:M}$ is defined and denoted by

$$\text{OLS}(\tilde{S}, \mathcal{L}_k, X_{i:N}^{i,k,1:M}) = \arg \inf_{\varphi \in \mathcal{L}_k} \frac{1}{M} \sum_{m=1}^M |\tilde{S}(X_{i:N}^{i,k,m}) - \varphi(X_i^{i,k,m})|^2.$$

The main difference with the usual regression Monte-Carlo schemes (see [8] for instance) is that here we use the common random numbers $U^{1:M}$ for all the regression problems. This is the effect of resampling. The convergence analysis becomes more delicate because we lose nice independence properties. Figure 1 describes a key part in the algorithm, namely the process of using the root paths to generate new paths.

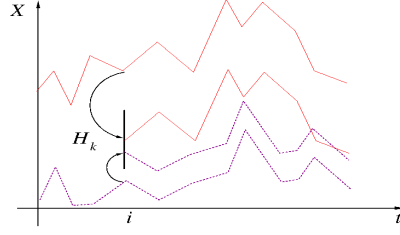


FIG. 1. Description of the use of the root paths to produce new paths in an arbitrary hypercube.

3. Convergence analysis in the case of the one-step ahead dynamic programming equation. We consider here the case

$$Y_N = g_N(X_N), \quad Y_i = \mathbb{E}[g_i(Y_{i+1}, X_i, \dots, X_N) \mid X_i], \quad 0 \leq i < N,$$

where we need the value of Y_{i+1} (one step ahead) to compute the value Y_i (at the current date) through a conditional expectation. To compare with Algorithm 1, we take $g_i(Y_{i+1:N}, X_{i:N}) = g_i(Y_{i+1}, X_{i:N})$.

Equations of this form are quite natural when solving optimal stopping problems in the Markovian case. Indeed, if V_i is the related value function at time i , i.e., the essential supremum over stopping times $\tau \in \{i, \dots, N\}$ of a reward process $f_\tau(X_\tau)$, then $V_i = \max(Y_i, f_i(X_i))$ where Y_i is the continuation value defined by

$$Y_i = \mathbb{E}[\max(Y_{i+1}, f_{i+1}(X_{i+1})) \mid X_i],$$

see [18] for instance. This corresponds to our setting with

$$g_i(y_{i+1}, x_{i:N}) = \max(y_{i+1}, f_{i+1}(x_{i+1})).$$

Similar dynamic programming equations appear in stochastic control problems. See [4].

3.1. Standing assumptions. The following assumptions enable us to provide error estimates (Theorem 3.4 and Corollary 3.1) for the convergence of Algorithm 1.

3.1.1. Assumptions on g_i .

ASSUMPTION 3.1 (Functions g_i). Each function g_i is Lipschitz w.r.t. the variable y_{i+1} , with Lipschitz constant L_{g_i} and $C_{g_i} := \sup_{x_{i:N}} |g_i(0, x_{i:N})| < +\infty$.

It is then easy to justify that y_i (such that $y_i(X_i) = Y_i$) is bounded (Assumption 2.1).

3.1.2. Assumptions on the distribution ν . We assume a condition on the probability measure ν and the Markov chain X , which ensures a suitable stability in the propagation of errors.

ASSUMPTION 3.2 (norm-stability). There exists a constant $\underline{C}_{(3.1)} \geq 1$ such that for any $\varphi : \mathbb{R}^d \mapsto \mathbb{R} \in L^2(\nu)$ and any $0 \leq i \leq N-1$, we have

$$(3.1) \quad \int_{\mathbb{R}^d} \mathbb{E}[\varphi^2(X_{i+1}^{i,x})] \nu(dx) \leq \underline{C}_{(3.1)} \int_{\mathbb{R}^d} \varphi^2(x) \nu(dx).$$

We now provide some examples of distribution ν where the above assumption holds, in connection with Examples 2.1, 2.2 and 2.4.

PROPOSITION 3.1. Let $\alpha = (\alpha^1, \dots, \alpha^d) \in]0, +\infty[^d$ and assume that $X_{i+1}^{i,x} = x + U_i$ (as in Examples 2.1 and 2.2) with $\mathbb{E} \left[\prod_{j=1}^d e^{\alpha^j |U_i^j|} \right] < +\infty$. Then, the tensor-product Laplace distribution $\nu(dx) := \prod_{j=1}^d \frac{\alpha^j}{2} e^{-\alpha^j |x^j|} dx$ satisfies Assumption 3.2.

Proof. The L.H.S. of (3.1) writes

$$\begin{aligned} \mathbb{E} \left[\int_{\mathbb{R}^d} \varphi^2(x + U_i) \nu(dx) \right] &= \mathbb{E} \left[\int_{\mathbb{R}^d} \varphi^2(x) \prod_{j=1}^d \frac{\alpha^j}{2} e^{-\alpha^j |x^j - U_i^j|} dx \right] \\ &\leq \mathbb{E} \left[\int_{\mathbb{R}^d} \varphi^2(x) \prod_{j=1}^d \frac{\alpha^j}{2} e^{-\alpha^j |x^j| + \alpha^j |U_i^j|} dx \right] \end{aligned}$$

which leads to the announced inequality (3.1) with $\underline{C}_{(3.1)} := \mathbb{E} \left[\prod_{j=1}^d e^{\alpha^j |U_i^j|} \right]$. $\square$

PROPOSITION 3.2. Let $k > 0$ and assume that $X_{i+1}^{i,x} = Dx + U_i$ for a diagonal invertible matrix $D := \text{diag}(D^1, \dots, D^d)$ (a form similar to Example 2.4) with $\mathbb{E} \left[(1 + |U_i|)^{d(k+1)} \right] < +\infty$. Then, the tensor-product Pareto-type distribution $\nu(dx) := \prod_{j=1}^d \frac{k}{2} (1 + |x^j|)^{-k-1} dx$ satisfies Assumption 3.2.

Proof. The L.H.S. of (3.1) equals

$$\begin{aligned} &\mathbb{E} \left[\int_{\mathbb{R}^d} \varphi^2(Dx + U_i) \nu(dx) \right] \\ &= \mathbb{E} \left[\int_{\mathbb{R}^d} \varphi^2(x) \det(D^{-1}) \prod_{j=1}^d \frac{k}{2} (1 + |(x^j - U_i^j)/D^j|)^{-k-1} dx \right] \\ (3.2) \quad &\leq \int_{\mathbb{R}^d} \varphi^2(x) \det(D^{-1}) \prod_{j=1}^d \frac{k}{2} \left(\mathbb{E} \left[(1 + |(x^j - U_i^j)/D^j|)^{-d(k+1)} \right] \right)^{1/d} dx. \end{aligned}$$

On the set $\{|U_i^j| \leq |x^j|/2\}$ we have $(1 + |(x^j - U_i^j)/D^j|) \geq (1 + (|x^j| - |U_i^j|)/D^j) \geq (1 + |x^j|/(2D^j))$. On the complementary set $\{|U_i^j| > |x^j|/2\}$, the random variable inside the j -th expectation in (3.2) is bounded by 1 and furthermore

$$\mathbb{P} \left(|U_i^j| > |x^j|/2 \right) \leq \frac{\mathbb{E} \left[(1 + 2|U_i^j|)^{d(k+1)} \right]}{(1 + |x^j|)^{d(k+1)}}.$$

By gathering the two cases, we observe that we have shown that the j -th expectation in (3.2) is bounded by $\text{Cst}(1 + |x^j|)^{-d(k+1)}$, for any x^j , whence the advertised result. $\square$

Remarks.

- Since we will apply the inequality (3.1) only to functions in a finite dimensional space, the norm equivalence property of finite dimensional space may also give the existence of a constant $\underline{C}_{(3.1)}$. But the constant built in this way could depend on the finite dimensional space (and may blow up when its dimension increases) while here the constant is valid for any φ .
- The previous examples on ν are related to distributions with independent components: this is especially convenient when one has to sample ν restricted to hypercubes $\mathcal{H}_k$, since we are reduced to independent one-dimensional simulations.

- In Proposition 3.2, had the matrix D been symmetric instead of diagonal, we would have applied an appropriate rotation to the density ν .

3.1.3. Covering number of an approximation space. To analyze how the M -samples $(X_{i:N}^{i,k,m}, 1 \leq m \leq M)$ from NISR approximates the exact distribution of $X_{i:N}$ with $X_i \stackrel{d}{\sim} \nu_k$ over test functions in the space $\mathcal{L}_k$, we will invoke concentration of measure inequalities (uniform in $\mathcal{L}_k$). This is possible thanks to complexity estimates related to $\mathcal{L}_k$, expressed in terms of covering numbers. Note that the concept of covering numbers is mainly used to introduce Assumption 3.3 and it intervenes in the main theorems only through the proof of Proposition 3.5.

We briefly recall the definition of a covering number of a dictionary of functions $\mathcal{G}$, see [10, Chapter 9] for more details. For a dictionary $\mathcal{G}$ of functions from $\mathbb{R}^d$ to $\mathbb{R}$ and for M points $x^{1:M} := \{x^{(1)}, \dots, x^{(M)}\}$ in $\mathbb{R}^d$, an ε -cover ($\varepsilon > 0$) of $\mathcal{G}$ w.r.t. the L^1 -empirical norm $\|g\|_1 := \frac{1}{M} \sum_{m=1}^M |g(x^{(m)})|$ is a finite collection of functions $g_1, \dots, g_n$ such that for any $g \in \mathcal{G}$, we can find a $j \in \{1, \dots, n\}$ such that $\|g - g_j\|_1 \leq \varepsilon$. The smallest possible integer n is called the ε -covering number and is denoted by $\mathcal{N}_1(\varepsilon, \mathcal{G}, x^{1:M})$.

ASSUMPTION 3.3 (Covering the approximation space). *There exist three constants*

$$\alpha_{(3.3)} \geq \frac{1}{4}, \quad \beta_{(3.3)} > 0, \quad \gamma_{(3.3)} \geq 1$$

such that for any $B > 0, \varepsilon \in (0, \frac{4}{15}B]$ and stratum index $1 \leq k \leq K$, the minimal size of an ε -covering number of $\mathcal{T}_B \mathcal{L}_k := \{\mathcal{T}_B \varphi : \varphi \in \mathcal{L}_k\}$ is bounded as follows:

$$(3.3) \quad \mathcal{N}_1(\varepsilon, \mathcal{T}_B \mathcal{L}_k, x^{1:M}) \leq \alpha_{(3.3)} \left(\frac{\beta_{(3.3)} B}{\varepsilon} \right)^{\gamma_{(3.3)}}$$

independently of the points sample $x^{1:M}$.

We assume that the above constants do not depend on k , mainly for the sake of simplicity. In the error analysis (see also Proposition A.1), the constants $\alpha_{(3.3)}$ and $\beta_{(3.3)}$ appear in log and thus, they have a small impact on error bounds. On the contrary, $\gamma_{(3.3)}$ appears as a multiplicative factor and we seek to have the smallest estimate.

PROPOSITION 3.3. *In the case of approximation spaces $\mathcal{L}_k$ like $\mathbf{LP}_0$, $\mathbf{LP}_1$ or $\mathbf{LP}_n$, Assumption 3.3 is satisfied with the following parameters: for any given $\eta > 0$, we have*

	$\alpha_{(3.3)}$	$\beta_{(3.3)}$	$\gamma_{(3.3)}$
$\mathbf{LP}_0$	1	7/5	1
$\mathbf{LP}_1$	3	$[4c_\eta 6^\eta]^{1/(1+\eta)} e$	$(d+2)(1+\eta)$
$\mathbf{LP}_n$	3	$[4c_\eta 6^\eta]^{1/(1+\eta)} e$	$((d+1)^n + 1)(1+\eta)$

where $c_\eta = \sup_{x \geq \frac{45e}{2}} x^{-\eta} \log(x)$.

The proof is postponed to the Appendix.

3.2. Main result: Error estimate. We are now in the position to state a convergence result, expressed in terms of the quadratic error of the best approximation of y_i on the stratum $\mathcal{H}_k$:

$$T_{i,k} := \inf_{\varphi \in \mathcal{L}_k} |y_i - \varphi|_{\nu_k}^2 \quad \text{where} \quad |\varphi|_{\nu_k}^2 := \int_{\mathbb{R}^d} |\varphi|^2(x) \nu_k(dx).$$

Our goal is to find an upper bound for the error $\mathbb{E} \left[|y_i^{(M)} - y_i|_\nu^2 \right]$ where

$$|\varphi|_\nu^2 := \int_{\mathbb{R}^d} |\varphi|^2(x) \nu(dx).$$

Note that the above expectation is taken over all the random variables, including the random sources $U^{1:M}$, i.e., we estimate the quadratic error averaged on the root sample.

THEOREM 3.4. *Assume Assumptions 2.2-2.3-3.2-3.3 and define $y_i^{(M)}$ as in Algorithm 1. Then, for any $\varepsilon > 0$, we have*

$$\begin{aligned} \mathbb{E} \left[|y_i^{(M)} - y_i|_\nu^2 \right] &\leq 4(1 + \varepsilon) L_{g_i}^2 \underline{C}_{(3.1)} \mathbb{E} \left[|y_{i+1}^{(M)} - y_{i+1}|_\nu^2 \right] + 2 \sum_{k=1}^K \nu(H_k) T_{i,k} + 4c_{(3.8)}(M) \frac{|y_i|_\infty^2}{M} \\ &\quad + 2(1 + \frac{1}{\varepsilon}) \frac{\dim(\mathcal{L})}{M} (C_{g_i} + L_{g_i} |y_{i+1}|_\infty)^2 + 8(1 + \varepsilon) L_{g_i}^2 c_{(3.7)}(M) \frac{|y_{i+1}|_\infty^2}{M}. \end{aligned}$$

We emphasize that whenever useful, the constant $4(1 + \varepsilon) L_{g_i}^2 \underline{C}_{(3.1)}$ could be reduced to $(1 + \delta)(1 + \varepsilon) L_{g_i}^2 \underline{C}_{(3.1)}$ (for any given $\delta > 0$) by slightly adapting the proof: namely, the term $4 = 2^2$ comes from two applications of deviation inequalities stated in Proposition A.1. These inequalities are valid with $(1 + \delta)^{\frac{1}{2}}$ instead of 2, up to modifying the constants $c_{(A.2)}(M)$ and $c_{(A.3)}(M)$.

As a very significant difference with usual Regression Monte-Carlo methods (see [9, Theorem 4.11]), in our algorithm there is no competition between the bias term (approximation error) and the variance term (statistical error), while in usual algorithms as the dimension of the approximation space $K \dim(\mathcal{L})$ goes to infinity, the statistical term (of size $\frac{K \dim(\mathcal{L})}{M}$) blows up. This significant improvement comes from the stratification which gives rise to decoupled and low-dimensional regression problems.

Since $y_N^{(M)} = y_N$, we easily derive global error bounds.

COROLLARY 3.1. *Under the assumptions and notations of Theorem 3.4, there exists a constant $C_{(3.4)}(N)$ (depending only on N , $\sup_{0 \leq i < N} L_{g_i}$, $\underline{C}_{(3.1)}$), such that for any $j \in \{0, \dots, N - 1\}$,*

$$\begin{aligned} (3.4) \quad \mathbb{E} \left[|y_j^{(M)} - y_j|_\nu^2 \right] &\leq C_{(3.4)}(N) \sum_{i=j}^{N-1} \left[\sum_{k=1}^K \nu(H_k) T_{i,k} \right. \\ &\quad \left. + \frac{1}{M} \left(c_{(3.8)}(M) |y_i|_\infty^2 + \dim(\mathcal{L}) (C_{g_i} + L_{g_i} |y_{i+1}|_\infty)^2 + L_{g_i}^2 c_{(3.7)}(M) |y_{i+1}|_\infty^2 \right) \right]. \end{aligned}$$

It is easy to see that if $4(1 + \varepsilon) L_{g_i}^2 \underline{C}_{(3.1)} \leq 1$, then interestingly $C_{(3.4)}(N)$ can be taken uniformly in N . This case corresponds to a small Lipschitz constant of g_i . In the case $4(1 + \varepsilon) L_{g_i}^2 \underline{C}_{(3.1)} \gg 1$, the above error estimates deteriorate quickly as N increases. We shall discuss that in Section 4 which deals with BSDEs and where we propose a different scheme that allows both large Lipschitz constant and large N .

3.3. Proof of Theorem 3.4. Let us start by setting up some useful notations:

$$S(x_{i:N}) := g_i(y_{i+1}(x_{i+1}), x_{i:N}), \quad \psi_i^k := \text{OLS}(S, \mathcal{L}_k, X_{i:N}^{i,k,1:M}),$$

$$|f|_{i,k,M}^2 := \frac{1}{M} \sum_{m=1}^M f^2(X_{i:N}^{i,k,m})$$

308

309 (or $|f|_{i,k,M}^2 := \frac{1}{M} \sum_{m=1}^M f^2(X_i^{i,k,m})$ when f depends only on one argument).

We first aim at deriving a bound on $\mathbb{E} \left[|y_i^{(M)} - y_i|_{i,k,M}^2 \right]$. First of all, note that

$$|y_i^{(M)} - y_i|_{i,k,M}^2 = \left| \mathcal{T}_{|y_i|_\infty}(\psi_i^{(M),k}) - \mathcal{T}_{|y_i|_\infty}(y_i) \right|_{i,k,M}^2 \leq |\psi_i^{(M),k} - y_i|_{i,k,M}^2$$

310 since the truncation operator is 1-Lipschitz. Now we define

$$311 \quad (3.5) \quad \mathbb{E} \left[S(X_{i:N}^{i,k,m}) | X_i^{i,k,1:M} \right] = \mathbb{E} \left[S(X_{i:N}^{i,k,m}) | X_i^{i,k,m} \right] = y_i(X_i^{i,k,m})$$

312 where the first equality is due to the independence of the paths $(X_{i:N}^{i,k,m}, 1 \leq m \leq M)$
 313 (Proposition 2.1) and where the last equality stems from the definition of y_i .

According to [9, Proposition 4.12] which allows to interchange conditional expectation and OLS, we have

$$\mathbb{E} \left[\psi_i^k(\cdot) | X_i^{i,k,1:M} \right] = \text{OLS}(y_i, \mathcal{L}_k, X_{i:N}^{i,k,1:M}).$$

314 Since the expected values $\left(\mathbb{E} \left[\psi_i^k(X_i^{i,k,m}) | X_i^{i,k,1:M} \right] \right)_{1 \leq m \leq M}$ can be seen as the projec-
 315 tions of $(y_i(X_i^{i,k,m}))_{1 \leq m \leq M}$ on the subspace of $\mathbb{R}^M$ spanned by $\{(\varphi(X_i^{i,k,m}))_{1 \leq m \leq M}, \varphi \in$
 316 $\mathcal{L}_k\}$ and $(\psi_i^{(M),k}(X_i^{i,k,m}))_{1 \leq m \leq M}$ is an element in this subspace, Pythagoras theorem
 317 yields

$$\begin{aligned} 318 \quad |\psi_i^{(M),k} - y_i|_{i,k,M}^2 &= \left| \psi_i^{(M),k} - \mathbb{E} \left[\psi_i^k(\cdot) | X_i^{i,k,1:M} \right] \right|_{i,k,M}^2 + \left| \mathbb{E} \left[\psi_i^k(\cdot) | X_i^{i,k,1:M} \right] - y_i \right|_{i,k,M}^2 \\ 319 \quad &= \left| \psi_i^{(M),k} - \mathbb{E} \left[\psi_i^k(\cdot) | X_i^{i,k,1:M} \right] \right|_{i,k,M}^2 + \inf_{\varphi \in \mathcal{L}_k} |\varphi - y_i|_{i,k,M}^2. \end{aligned}$$

321 For any given $\phi \in \mathcal{L}_k$, we have

$$\begin{aligned} 322 \quad \mathbb{E} \left[\inf_{\varphi \in \mathcal{L}_k} |\varphi - y_i|_{i,k,M}^2 \right] &\leq \mathbb{E} \left[|\phi - y_i|_{i,k,M}^2 \right] = \mathbb{E} \left[\frac{1}{M} \sum_{m=1}^M |\phi(X_i^{i,k,m}) - y_i(X_i^{i,k,m})|^2 \right] \\ 323 \quad &= \int_{\mathbb{R}^d} |\phi(x) - y_i(x)|^2 \nu_k(dx). \end{aligned}$$

Taking the infimum over all functions ϕ on the R.H.S. gives

$$\mathbb{E} \left[\inf_{\varphi \in \mathcal{L}_k} |\varphi - y_i|_{i,k,M}^2 \right] \leq T_{i,k}.$$

325 So, for any $\varepsilon > 0$, we have

$$\begin{aligned} 326 \quad \mathbb{E} \left[|\psi_i^{(M),k} - y_i|_{i,k,M}^2 \right] &\leq T_{i,k} + (1 + \varepsilon) \mathbb{E} \left[|\psi_i^{(M),k} - \psi_i^k|_{i,k,M}^2 \right] \\ 327 \quad &+ (1 + \frac{1}{\varepsilon}) \mathbb{E} \left[\left| \psi_i^k - \mathbb{E} \left[\psi_i^k(\cdot) | X_i^{i,k,1:M} \right] \right|_{i,k,M}^2 \right]. \end{aligned}$$

329 By [9, Proposition 4.12], the last term is bounded by $\frac{\dim(\mathcal{L})}{M} (C_{g_i} + L_{g_i} |y_{i+1}|_\infty)^2$ where
 330 $(C_{g_i} + L_{g_i} |y_{i+1}|_\infty)^2$ clearly bounds the conditional variance of $S(X_{i:N}^{i,k})$. This is the

statistical error contribution. Here, we have used the independence of $(X_{i:N}^{i,k,m}, 1 \leq m \leq M)$ (Proposition 2.1).

The control of the term $\mathbb{E} \left[|\psi_i^{(M),k} - \psi_{i,k,M}^k|^2 \right]$ is possible due to the linearity and stability of OLS [9, Proposition 4.12]:

$$|\psi_i^{(M),k} - \psi_{i,k,M}^k|^2 \leq |S^{(M)} - S|_{i,k,M}^2 \leq L_{g_i}^2 \frac{1}{M} \sum_{m=1}^M (y_{i+1}^{(M)} - y_{i+1})^2 (X_{i+1}^{i,k,m}),$$

where we have taken advantage of the Lipschitz property of g_i w.r.t. the component y_{i+1} . So far we have shown

$$\begin{aligned} \mathbb{E} \left[|y_i^{(M)} - y_i|_{i,k,M}^2 \right] &\leq T_{i,k} + (1 + \varepsilon) L_{g_i}^2 \mathbb{E} \left[\frac{1}{M} \sum_{m=1}^M (y_{i+1}^{(M)} - y_{i+1})^2 (X_{i+1}^{i,k,m}) \right] \\ &\quad + (1 + \frac{1}{\varepsilon}) \frac{\dim(\mathcal{L})}{M} (C_{g_i} + L_{g_i} |y_{i+1}|_\infty)^2. \end{aligned}$$

This shows a relation between the errors at time i and time $i + 1$, but measured in different norms. In order to retrieve the same $L^2(\nu)$ -norm and continue the analysis, we will use the norm-stability property (Assumption 3.2) and the following result about concentration of measures. The proof is a particular case of Proposition A.1 in the Appendix, with $\psi(x) = (-2|y_{i+1}|_\infty \vee x \wedge 2|y_{i+1}|_\infty)^2$, $B = |y_{i+1}|_\infty$, $\mathcal{K} = \mathcal{L}_k$, $\eta = y_{i+1}$.

PROPOSITION 3.5. Define $(c_{(3.7)}(M), c_{(3.8)}(M))$ by considering $(c_{(A.2)}(M), c_{(A.3)}(M))$ from Proposition A.1 with the values $(\alpha_{(3.3)}, \beta_{(3.3)}, \gamma_{(3.3)})$ instead of (α, β, γ) . Then we have

$$\begin{aligned} \mathbb{E} \left[\frac{1}{M} \sum_{m=1}^M (y_{i+1}^{(M)} - y_{i+1})^2 (X_{i+1}^{i,k,m}) \right] &\leq 2\mathbb{E} \left[|y_{i+1}^{(M)}(X_{i+1}^{i,\nu_k}) - y_{i+1}(X_{i+1}^{i,\nu_k})|^2 \right] \\ &\quad + 4c_{(3.7)}(M) \frac{|y_{i+1}|_\infty^2}{M}, \end{aligned}$$

$$\mathbb{E} \left[|y_i^{(M)} - y_i|_{\nu_k}^2 \right] \leq 2\mathbb{E} \left[|y_i^{(M)} - y_i|_{i,k,M}^2 \right] + 4c_{(3.8)}(M) \frac{|y_i|_\infty^2}{M}.$$

Multiply both sides of Equation (3.7) by $\nu(H_k)$, sum over k , and use the norm-stability property (Assumption 3.2): it readily follows that

$$\begin{aligned} \sum_{k=1}^K \nu(H_k) \mathbb{E} \left[\frac{1}{M} \sum_{m=1}^M (y_{i+1}^{(M)} - y_{i+1})^2 (X_{i+1}^{i,k,m}) \right] \\ \leq 2\mathbb{E} \left[|y_{i+1}^{(M)}(X_{i+1}^{i,\nu}) - y_{i+1}(X_{i+1}^{i,\nu})|^2 \right] + 4c_{(3.7)}(M) \frac{|y_{i+1}|_\infty^2}{M} \\ \leq 2C_{(3.1)} \mathbb{E} \left[|y_{i+1}^{(M)} - y_{i+1}|_\nu^2 \right] + 4c_{(3.7)}(M) \frac{|y_{i+1}|_\infty^2}{M}. \end{aligned}$$

Similarly, we can get from Equation (3.8) that

$$\mathbb{E} \left[|y_i^{(M)} - y_i|_\nu^2 \right] \leq 2 \sum_{k=1}^K \nu(H_k) \mathbb{E} \left[|y_i^{(M)} - y_i|_{i,k,M}^2 \right] + 4c_{(3.8)}(M) \frac{|y_i|_\infty^2}{M}.$$

364 Finally, by combining the above estimates with (3.6), we get

$$\begin{aligned}
 365 \quad \mathbb{E} \left[|y_i^{(M)} - y_i|_\nu^2 \right] &\leq 4c_{(3.8)}(M) \frac{|y_i|_\infty^2}{M} + 2 \sum_{k=1}^K \nu(H_k) T_{i,k} + 2(1 + \frac{1}{\varepsilon}) \frac{\dim(\mathcal{L})}{M} (C_{g_i} + L_{g_i} |y_{i+1}|_\infty)^2 \\
 366 \quad &+ 2(1 + \varepsilon) L_{g_i}^2 \left(2C_{(3.1)} \mathbb{E} \left[|y_{i+1}^{(M)} - y_{i+1}|_\nu^2 \right] + 4c_{(3.7)}(M) \frac{|y_{i+1}|_\infty^2}{M} \right). \\
 367
 \end{aligned}$$

368 This links $\mathbb{E} \left[|y_i^{(M)} - y_i|_\nu^2 \right]$ with $\mathbb{E} \left[|y_{i+1}^{(M)} - y_{i+1}|_\nu^2 \right]$ as announced.

4. Convergence analysis for the solution of BSDEs with the MDP representation. Let us consider the semi-linear final value problem for a parabolic PDE of the form

$$\begin{cases} \partial_t u(t, x) + \frac{1}{2} \Delta u(t, x) + f(t, u(t, x), x) = 0, & t < 1, x \in \mathbb{R}^d, \\ u(1, x) = g(x). \end{cases}$$

This is a simple form of the Hamilton-Jacobi-Bellman equation of stochastic control problems [13]. Under fairly mild assumptions (see [16] for instance), the solution to the above PDE is related to a Backward Stochastic Differential Equation $(\mathcal{Y}, \mathcal{Z})$ driven by a d -dimensional Brownian motion W . Namely,

$$\mathcal{Y}_t = g(W_1) + \int_t^1 f(s, \mathcal{Y}_s, W_s) ds - \int_t^1 \mathcal{Z}_s dW_s$$

and $\mathcal{Y}_t = u(t, W_t)$, $\mathcal{Z}_t = \nabla u(t, W_t)$. Needless to say, the Laplacian Δ and the process W could be replaced by a more general second order operator and its related diffusion process, and that f could depend on the gradient Z as well. We stick to the above setting which is consistent with this work. Taking conditional expectation reduces to

$$\mathcal{Y}_t = \mathbb{E} \left[g(W_1) + \int_t^1 f(s, \mathcal{Y}_s, W_s) ds \mid W_t \right].$$

369 There are several time discretization schemes of $\mathcal{Y}$ (explicit or implicit Euler
 370 schemes, high order schemes [6]) but here we follow the Multi-Step Forward Dynamic
 371 Programming (MDP for short) Equation of [9], which allows a better error propagation
 372 compared to the One-Step Dynamic Programming Equation:

$$\begin{aligned}
 373 \quad Y_i &= \mathbb{E} \left[g_N(X_N) + \frac{1}{N} \sum_{j=i+1}^N f_j(Y_j, X_j, \dots, X_N) \mid X_i \right] = y_i(X_i), \quad 0 \leq i < N. \\
 374
 \end{aligned}$$

Here, we consider a more general path-dependency on f_j , actually this does not affect the error analysis. In comparison with Algorithm 1, we take

$$g_i(y_{i+1:N}, x_{i:N}) = g_N(x_N) + \frac{1}{N} \sum_{j=i+1}^N f_j(y_j, x_{j:N}).$$

375 In [2] similar discrete BSDEs appear but with an external noise. That corresponds to
 376 time-discretization of Backward Doubly SDEs, which in turn are related to stochastic
 377 semi-linear PDEs.

4.1. Standing assumptions. We shall now describe the main assumptions that are needed in the methodology proposed in this paper.

4.1.1. Assumptions on f_i and g_N .

ASSUMPTION 4.1 (Functions f_i and g_N). *Each f_i is Lipschitz w.r.t. y_i , with Lipschitz constant L_{f_i} and $C_{f_i} = \sup_{x_{i:N}} |f_i(0, x_{i:N})| < +\infty$. Moreover g_N is bounded.*

The reader can easily check that y_i is bounded.

4.1.2. Assumptions on the distribution ν .

ASSUMPTION 4.2 (norm-stability). *There exists a constant $\underline{C}_{(4.1)} \geq 1$ such that for any $\varphi \in L^2(\nu)$ and any $0 \leq i < j \leq N$, we have*

$$(4.1) \quad \int_{\mathbb{R}^d} \mathbb{E} \left[\varphi^2(X_j^{i,x}) \right] \nu(dx) \leq \underline{C}_{(4.1)} \int_{\mathbb{R}^d} |\varphi(x)|^2 \nu(dx).$$

It is straightforward to extend Propositions 3.1 and 3.2 to fulfill the above assumption.

4.2. Main result: error estimate. We express the error in terms of the best local approximation error and the averaged one:

$$T_{i,k} := \inf_{\varphi \in \mathcal{L}_k} |y_i - \varphi|_{\nu_k}^2, \quad \nu(T_{i,\cdot}) := \sum_{k=1}^K \nu(H_k) T_{i,k}.$$

In this discrete time BSDE context, Theorem 3.4 becomes the following.

THEOREM 4.1. *Assume Assumptions 2.2-2.3-3.3-4.2 and define $y_i^{(M)}$ as in Algorithm 1. Set*

$$\bar{\mathcal{E}}(Y, M, i) := \mathbb{E} \left[|y_i^{(M)} - y_i|_{\nu}^2 \right] = \sum_{k=1}^K \nu(\mathcal{H}_k) \mathbb{E} \left[|y_i^{(M)} - y_i|_{\nu_k}^2 \right].$$

Define

$$\begin{aligned} \delta_i &= 4c_{(3.8)}(M) \frac{|y_i|_{\infty}^2}{M} + 2\nu(T_{i,\cdot}) + 16 \frac{1}{N} \sum_{j=i+1}^{N-1} L_{f_j}^2 c_{(3.7)}(M) \frac{|y_j|_{\infty}^2}{M} + \\ &+ 4 \frac{\dim(\mathcal{L})}{M} \left(|y_N|_{\infty} + \frac{1}{N} \sum_{j=i+1}^N (C_{f_j} + L_{f_j} |y_j|_{\infty}) \right)^2. \end{aligned}$$

Then, letting $L_f := \sup_j L_{f_j}$, we have

$$\bar{\mathcal{E}}(Y, M, i) \leq \delta_i + 8\underline{C}_{(4.1)} L_f^2 \exp \left(8\underline{C}_{(4.1)} L_f^2 \right) \frac{1}{N} \sum_{j=i+1}^{N-1} \delta_j.$$

The above general error estimates become simpler when the parameters are uniform in i .

COROLLARY 4.1. *Under the assumptions of Theorem 4.1 and assuming that C_{f_i}, L_{f_i} and $|y_i|_{\infty}$ are bounded uniformly in i and N , there exists a constant $C_{(4.2)}$ (independent of N and of approximation spaces $\mathcal{L}_k$) such that*

$$(4.2) \quad \bar{\mathcal{E}}(Y, M, i) \leq C_{(4.2)} \left(\frac{c_{(3.8)}(M) + c_{(3.7)}(M) + \dim(\mathcal{L})}{M} + \nu(T_{i,\cdot}) + \frac{1}{N} \sum_{j=i+1}^{N-1} \nu(T_{j,\cdot}) \right).$$

We observe that this upper bound is expressed in a quite convenient form to let $N \rightarrow +\infty$ and $K \rightarrow +\infty$. As a major difference with the usual Regression Monte Carlo schemes, the impact of the statistical error (through the parameter M) is not affected by the number K of strata.

4.3. Proof of Theorem 4.1. We follow the arguments of the proof of Theorem 3.4 with the following notation:

$$S(x_{i:N}) := g_N(x_N) + \frac{1}{N} \sum_{j=i+1}^N f_j(y_j(x_j), x_{j:N}),$$

$$S^{(M)}(x_{i:N}) := g_N(x_N) + \frac{1}{N} \sum_{j=i+1}^N f_j(y_j^{(M)}(x_j), x_{j:N}),$$

$$\psi_i^k := \text{OLS}(S, \mathcal{L}_k, X_{i:N}^{i,k,1:M}), \quad \psi_i^{(M),k} := \text{OLS}(S^{(M)}, \mathcal{L}_k, X_{i:N}^{i,k,1:M}).$$

The beginning of the proof is similar and we obtain (here, there is no need to optimize ε and we take $\varepsilon = 1$)

$$\begin{aligned} \mathbb{E} \left[|y_i^{(M)} - y_{i,k,M}|^2 \right] &\leq \mathbb{E} \left[|\psi_i^{(M),k} - y_{i,k,M}|^2 \right] \\ &\leq T_{i,k} + 2\mathbb{E} \left[|\psi_i^{(M),k} - \psi_i^k|_{i,k,M}^2 \right] + 2\mathbb{E} \left[\left| \psi_i^k - \mathbb{E} \left[\psi_i^k(\cdot) | X_i^{i,k,1:M} \right] \right|_{i,k,M}^2 \right]. \end{aligned}$$

The last term is a statistical error term, which can be controlled as follows:

$$\mathbb{E} \left[\left| \psi_i^k - \mathbb{E} \left[\psi_i^k(\cdot) | X_i^{i,k,1:M} \right] \right|_{i,k,M}^2 \right] \leq \frac{\dim(\mathcal{L})}{M} \left(|y_N|_\infty + \frac{1}{N} \sum_{j=i+1}^N (C_{f_j} + L_{f_j} |y_j|_\infty) \right)^2$$

where $(\dots)^2$ is a rough bound of the conditional variance of $S(X_{i:N}^{i,k})$.

We handle the control of the term $\mathbb{E} \left[|\psi_i^{(M),k} - \psi_i^k|_{i,k,M}^2 \right]$ as in Theorem 3.4 but the results are different because the dynamic programming equation differs:

$$\begin{aligned} \mathbb{E} \left[|\psi_i^{(M),k} - \psi_i^k|_{i,k,M}^2 \right] &\leq \mathbb{E} \left[|S^{(M)} - S|_{i,k,M}^2 \right] \\ &\leq \mathbb{E} \left[\frac{1}{M} \sum_{m=1}^M \left(\frac{1}{N} \sum_{j=i+1}^{N-1} L_{f_j} |y_j^{(M)} - y_j|(X_j^{i,k,m}) \right)^2 \right] \\ &\leq \mathbb{E} \left[\frac{1}{M} \sum_{m=1}^M \frac{1}{N} \sum_{j=i+1}^{N-1} L_{f_j}^2 |y_j^{(M)} - y_j|^2(X_j^{i,k,m}) \right]. \end{aligned}$$

We multiply the above by $\nu(H_k)$, sum over k , apply the *extended* Proposition 3.5 valid also for the problem at hand, and the Assumption 4.2. Then, it follows that

$$\sum_{k=1}^K \nu(H_k) \mathbb{E} \left[|\psi_i^{(M),k} - \psi_i^k|_{i,k,M}^2 \right] \leq 2 \frac{1}{N} \sum_{j=i+1}^{N-1} L_{f_j}^2 \left(\underline{C}_{(4.1)} \mathbb{E} \left[|y_j^{(M)} - y_j|^2 \right] + 2c_{(3.7)}(M) \frac{|y_j|_\infty^2}{M} \right).$$

On the other hand, from Equation (3.8) we have

$$\mathbb{E} \left[|y_i^{(M)} - y_{i,\nu}|^2 \right] \leq 2 \sum_{k=1}^K \nu(H_k) \mathbb{E} \left[|y_i^{(M)} - y_{i,k,M}|^2 \right] + 4c_{(3.8)}(M) \frac{|y_i|_\infty^2}{M}.$$

Now collect the different estimates: it writes

$$\begin{aligned}
\bar{\mathcal{E}}(Y, M, i) &:= \mathbb{E} \left[|y_i^{(M)} - y_i|_\nu^2 \right] \leq 4c_{(3.8)}(M) \frac{|y_i|_\infty^2}{M} + 2 \sum_{k=1}^K \nu(H_k) T_{i,k} \\
&+ 8 \frac{1}{N} \sum_{j=i+1}^{N-1} L_{f_j}^2 \left(\underline{C}_{(4.1)} \mathbb{E} \left[|y_j^{(M)} - y_j|_\nu^2 \right] + 2c_{(3.7)}(M) \frac{|y_j|_\infty^2}{M} \right) + \\
&+ 4 \frac{\dim(\mathcal{L})}{M} \left(|y_N|_\infty + \frac{1}{N} \sum_{j=i+1}^N (C_{f_j} + L_{f_j} |y_j|_\infty) \right)^2 \\
&:= \delta_i + 8\underline{C}_{(4.1)} \frac{1}{N} \sum_{j=i+1}^{N-1} L_{f_j}^2 \bar{\mathcal{E}}(Y, M, j).
\end{aligned}$$

It takes the form of a discrete Gronwall lemma, which easily allows to derive the following upper bound (see [2, Appendix A.3]):

$$\begin{aligned}
\bar{\mathcal{E}}(Y, M, i) &\leq \delta_i + 8\underline{C}_{(4.1)} \frac{1}{N} \sum_{j=i+1}^{N-1} \Gamma_{i,j} L_{f_j}^2 \delta_j, \\
\text{where } \Gamma_{i,j} &:= \begin{cases} \prod_{i < k < j} (1 + 8\underline{C}_{(4.1)} \frac{1}{N} L_{f_k}^2), & \text{for } i+1 < j, \\ 1, & \text{otherwise.} \end{cases}
\end{aligned}$$

Using now $L_f = \sup_j L_{f_j}$, we get $\Gamma_{i,j} \leq \exp \left(\sum_{i < k < j} 8\underline{C}_{(4.1)} \frac{1}{N} L_{f_k}^2 \right) \leq \exp(8\underline{C}_{(4.1)} L_f^2)$. This completes the proof.

5. Numerical tests. We shall now illustrate the methodology in two numerical examples coming from practical problems. The first one concerns a reaction-diffusion PDE connected to spatially distributed populations, whereas the second one deals with a stochastic control problem.

5.1. An Application to Reaction-Diffusion Models in Spatially Distributed Populations. In this section we consider a biologically motivated example to illustrate the strength of the stratified resampling regression methodology presented in the previous sections. We selected an application to spatially distributed populations that evolve under reaction diffusion equations. Besides the theoretical challenges behind the models, it has recently attracted attention due to its impact in the spread of infectious diseases [14, 15] and even to the modeling of Wolbachia infected mosquitoes in the fight of disease spreading *Aedes aegypti* [3, 11].

The use of reaction diffusion models to describe the population dynamics of a single species or genetic trait expanding into new territory dominated by another one goes back to the work of R. A. Fisher [7] and A. Kolmogorov et al. [12]. The mathematical model behind it is known as the (celebrated) Fisher-Kolmogorov-Petrovski-Piscounov (FKPP) equation.

In a conveniently chosen scale it takes the form, in dimension 1,

$$(5.1) \quad \partial_t u + \partial_x^2 u + au(1-u) = 0, \quad u(T, x) = h(x), \quad x \in \mathbb{R}, t \leq T,$$

where $u = u(t, x)$ refers to the proportion of members of an invading species in a spatially distributed population on a straight line. The equation is chosen with time

running backwards and as a final value problem to allow direct connection with the standard probabilistic interpretation.

It is well known [1] that for any arbitrary positive C , if we define

$$(5.2) \quad h(x) := \left(1 + C \exp\left(\pm \frac{\sqrt{6a}}{6}x\right)\right)^{-2}$$

then

$$u(t, x) = \left(1 + C \exp\left(\frac{5a}{6}(t - T) \pm \frac{\sqrt{6a}}{6}x\right)\right)^{-2}$$

is a traveling wave solution to Equation (5.1). The behavior of $h(x)$ as $x \rightarrow \pm\infty$ is either one or zero according to the sign chosen inside the exponential. Thus describing full dominance of the invading species or its absence.

The probabilistic formulation goes as follows: Introduce the system, as in Section 4,

$$\begin{aligned} dP_s &= \sqrt{2}dW_s, \\ dY_s &= -f(Y_s)ds + Z_s dW_s, \text{ where } f(x) = ax(1-x), \text{ and } Z_s = \sqrt{2}\partial_x u(s, P_s) \\ Y_T &= u(T, P_T) = h(P_T). \end{aligned}$$

Then, the process $Y_t = \mathbb{E}\left[Y_T + \int_t^T f(Y_s)ds \mid P_t\right]$ satisfies $Y_t = u(t, P_t)$.

To test the algorithms presented herein, we shall start with the following more general parabolic PDE

$$(5.3) \quad \partial_t W + \sum_{1 \leq i, j \leq d} A_{ij} \partial_{y_i} \partial_{y_j} W + aW(1 - W) = 0, \quad t \leq T, \text{ and } y \in \mathbb{R}^d.$$

Here, the matrix A is chosen as an arbitrary *positive-definite* constant $d \times d$ matrix. Furthermore, we choose, for convenience, the final condition

$$(5.4) \quad W(T, y) = h(y' \Sigma^{-1} \theta),$$

where $\Sigma = \Sigma' = \sqrt{A}$ and θ is arbitrary unit vector. We stress that this special choice of the final condition has the sole purpose of bypassing the need of solving Equation (5.4) by other numerical methods for comparison with the present methodology. Indeed, the fact that we are able to exhibit an explicit solution to Equation (5.3) with final condition (5.4) allows an easy checking of the accuracy of the method. We also stress that the method developed in this work does not require an explicit knowledge of the diffusion coefficient matrix A of Equation (5.3) since we shall make use of the observed paths. Yet the knowledge of the function $W \mapsto aW(1 - W)$ is crucial.

It is easy to see that if $u = u(t, x)$ satisfies Equation (5.1) with final condition given by Equation (5.2) then

$$(5.5) \quad W(t, y) := u(t, y' \Sigma^{-1} \theta)$$

satisfies Equation (5.3) with final condition (5.4).

An interpretation of the methodology proposed here is the following: If we were able to observe the trajectories performed by a small number of free individuals according to the diffusion process associated to Equation (5.3), even if we did not know the explicit form of the diffusion (i.e., we did not have a good calibration of the covariance matrix) we could use such trajectories to produce a reliable solution to the final value problem (5.3) and (5.4).

We firstly present some numerical results in dimension 1 (Tables 1-2-3). We have tested both the one-step (Section 3) and multi-step schemes (Section 4). The final time T is fixed to 1 and we use time discretization $t_i = \frac{i}{N}T, 0 \leq i \leq N$ with $N = 10$ or 20. We divide the real line $\mathbb{R}$ into K subintervals $(I_i)_{1 \leq i \leq K}$ by fixing $A = 25$ and dividing $[-A, A]$ into $K - 2$ equal length intervals and then adding $(-\infty, -A)$ and $(A, +\infty)$. We implement our method by using piecewise constant estimation on each interval. Then finally we get a piecewise constant estimation of $u(0, y)$, noted as $\hat{u}(0, y)$. Then we approximate the squared $L^2(\nu)$ error of our estimation by

$$\sum_{1 \leq k \leq K} |u(0, y_k) - \hat{u}(0, y_k)|^2 \nu(I_k)$$

where y_k is chosen as the middle point of the rectangle if I_k is finite and the boundary point if I_k is infinite. We take $\nu(dx) = \frac{1}{2}e^{-|x|}dx$ and we use the restriction of ν on I_k to sample initial points. The squared $L^2(\nu)$ norm of $u(0, \cdot)$ is around 0.25. Finally remark that the error of our method includes three parts: time discretization error, approximation error due to the use of piecewise constant estimation on hypercubes and statistical error due to the randomness of trajectories. In the following tables, M is the number of trajectories that we use (i.e., the root sample).

We observe in Tables 1 and 3 that the approximation error (visible for small K) contributes much more to the global error for the one-step scheme, compared to the multi-step one. This observation is in agreement with those of [5, 9].

When N gets larger with fixed K and M (Tables 1 and 2), we may observe an increase of the global error for the one-step scheme, this is coherent with the estimates of Corollary 3.1.

	$K = 10$	$K = 20$	$K = 50$	$K = 100$	$K = 200$	$K = 400$
$M = 20$	0.0993	0.0253	0.0038	0.0014	0.0014	0.0019
$M = 40$	0.0997	0.0252	0.0034	9.01e-04	5.16e-04	6.17e-04
$M = 80$	0.0993	0.0249	0.0029	6.15e-04	3.92e-04	3.91e-04
$M = 160$	0.0990	0.0248	0.0029	3.15e-04	1.57e-04	1.71e-04
$M = 320$	0.0990	0.0248	0.0028	2.47e-04	1.02e-04	1.19e-04
$M = 640$	0.0990	0.0246	0.0028	2.26e-04	5.46e-05	4.94e-05

TABLE 1

Average squared L^2 errors with 50 macro runs, $N = 10$, one-step scheme.

	$K = 10$	$K = 20$	$K = 50$	$K = 100$	$K = 200$	$K = 400$
$M = 20$	0.1031	0.0299	0.0073	0.0018	0.0011	0.0012
$M = 40$	0.1031	0.0294	0.0066	0.0014	7.86e-04	7.28e-04
$M = 80$	0.1027	0.0293	0.0065	0.0010	3.18e-04	3.86e-04
$M = 160$	0.1027	0.0294	0.0064	8.91e-04	2.46e-04	1.04e-04
$M = 320$	0.1026	0.0293	0.0064	8.39e-04	1.42e-04	7.03e-05
$M = 640$	0.1027	0.0292	0.0063	8.04e-04	8.16e-05	5.60e-05

TABLE 2

Average squared L^2 errors with 50 macro runs, $N = 20$, one-step scheme.

Table 4 below describes numerical results in dimension 2. The final time T is fixed to 1 and we use the time discretization $t_i = \frac{i}{N}T, 0 \leq i \leq N$ with $N = 10$. We divide the real line $\mathbb{R}$ into K subintervals $(I_i)_{1 \leq i \leq K}$ by fixing $A = 25$ and dividing

	$K = 10$	$K = 20$	$K = 50$	$K = 100$	$K = 200$	$K = 400$
$M = 20$	0.0484	0.0066	0.0017	0.0015	0.0011	0.0013
$M = 40$	0.0488	0.0058	8.45e-04	5.81e-04	6.35e-04	5.68e-04
$M = 80$	0.0478	0.0053	4.33e-04	2.96e-04	3.45e-04	4.06e-04
$M = 160$	0.0481	0.0051	2.98e-04	2.23e-04	1.71e-04	1.08e-04
$M = 320$	0.0479	0.0051	1.79e-04	6.48e-05	8.38e-05	1.04e-04
$M = 640$	0.0478	0.0050	1.50e-04	6.49e-05	6.66e-05	5.70e-05

TABLE 3

Average squared L^2 errors with 50 macro runs, $N = 10$, multi-step scheme.

$[-A, A]$ into $K - 2$ equal length intervals and then adding $(-\infty, -A)$ and $(A, +\infty)$. We take $\Sigma = [1, \beta; \beta, 1]$ with $\beta = 0.25$ and $\theta = \frac{[1;1]}{\sqrt{2}}$. We implement our method by using piecewise constant estimation on each finite (or infinite) rectangle $I_i \times I_j$. Then finally we get a piecewise constant estimation of $W(0, y)$, noted as $\hat{W}(0, y)$. Then we approximate the squared $L^2(\nu \otimes \nu)$ error of our estimation by

$$\sum_{1 \leq k_1 \leq K, 1 \leq k_2 \leq K} |W(0, y_{k_1}, y_{k_2}) - \hat{W}(0, y_{k_1}, y_{k_2})|^2 \nu \otimes \nu(I_{k_1} \times I_{k_2})$$

where (y_{k_1}, y_{k_2}) is chosen as the middle point of the rectangle if $I_{k_1} \times I_{k_2}$ is finite and the boundary point if one or both of I_{k_1} and I_{k_2} are infinite. We take $\nu(dx) = \frac{1}{2}e^{-|x|}dx$ and we use the restriction of $\nu \otimes \nu$ on $I_i \times I_j$ to sample initial points. The squared $L^2(\nu \otimes \nu)$ norm of $W(0, \cdot, \cdot)$ is around 0.25.

	$K = 10$	$K = 20$	$K = 50$	$K = 100$	$K = 200$
$M = 20$	0.0592	0.0167	0.0027	0.0018	0.0010
$M = 40$	0.0588	0.0163	0.0022	5.34e-04	5.00e-04
$M = 80$	0.0588	0.0160	0.0019	3.74e-04	2.98e-04
$M = 160$	0.0586	0.0160	0.0018	3.08e-04	9.16e-05
$M = 320$	0.0586	0.0159	0.0017	1.1e-04	9.24e-05

TABLE 4

Average squared L^2 errors with 50 macro runs, $N = 10$, one-step scheme.

As for the previous case in dimension 1, we observe that when K is small, it is useless to increase M . This is because in such case the approximation error is dominant. But when K is large enough, the performance of our method improves when M becomes larger, since this time it is the statistical error which becomes dominant and larger M means smaller statistical error.

In the perspective of a given root sample (M fixed), it is recommended to take K large: indeed, in agreement with Theorems 3.4 and 4.1, we observe from the numerical results that the global error decreases up to the statistical error term (depending on M but not K). In this way, for $M = 20$ (resp. $M = 40$) the relative squared L^2 error is about 0.4% (resp. 0.22%).

5.2. Travel agency problem: when to offer travels, according to currency and weather forecast.... In this section we illustrate the stratified resampler methodology in the solution of an optimal investment problem. The underlying model will have two sources of stochasticity, one related to the weather and the other one to the exchange rate. The corresponding stochastic processes shall be denoted by X_t^1 and X_t^2 .

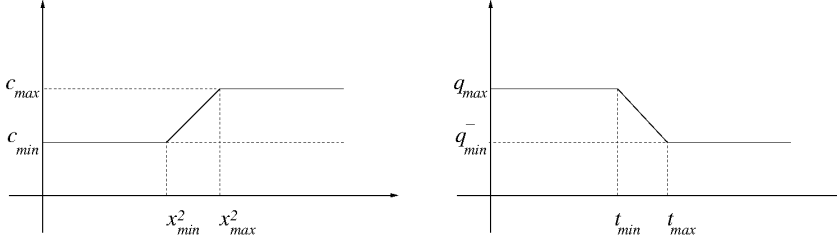


FIG. 2. Pictorial description of the cost function c (left) and of the campaign effectiveness q (right).

We envision the following situation: A travel agency wants to launch a campaign for the promotion of vacations in a warm region abroad during the Fall-Winter season. Such travel agency would receive a fixed value $\underline{c}$ in local currency from the customers and on the other hand would have to pay the costs $c = c(\exp(X_{\tau+1/12}^2))$ in a future time $\tau + 1/12$, where τ is the launching time of the campaign and $X_{\tau+1/12}^2$ is the prevailing logarithm of exchange rate one month after the launching, with the time unit set to be one year. The initial time $t = 0$ is by convention October 1st. In other words, the costs are fixed to the traveler and variable for the agency. A pictorial description of the cost function is presented in Figure 2.

The effectiveness of the campaign will depend on the local temperature $(t - 0.25)^2 \times 240 + X_t^1$ (in Celsius) and will be denoted by $q((t - 0.25)^2 \times 240 + X_t^1) \exp(-|t - 1/6|)$, where $(t - 0.25)^2 \times 240$ represents the seasonal component and X_t^1 represents the random part. Its purpose is to capture the idea that if the local temperature is very low, then people would be more interested in spending some days in a warm region, whereas if the weather is mild then people would just stay at home. A pictorial description of the function q is presented in Figure 2. The second part of this function $\exp(-|t - 1/6|)$ is created to represent the fact that there are likely more registrations at beginning of December for the period of new year holidays.

Thus, our problem consists of finding the function v defined by

$$\begin{aligned} v(X_0^1, X_0^2) &= \operatorname{ess\,sup}_{\tau \in \mathcal{T}} \mathbb{E} \left[q((\tau - 0.25)^2 \times 240 + X_\tau^1) e^{-|\tau - 1/6|} \left(\underline{c} - c(e^{X_{\tau+1/12}^2}) \right) \mid X_0^1, X_0^2 \right] \\ &= \operatorname{ess\,sup}_{\tau \in \mathcal{T}} \mathbb{E} \left[q((\tau - 0.25)^2 \times 240 + X_\tau^1) e^{-|\tau - 1/6|} \left(\underline{c} - \mathbb{E} \left[c(e^{X_{\tau+1/12}^2}) \mid X_\tau^2 \right] \right) \mid X_0^1, X_0^2 \right], \end{aligned}$$

where $\mathcal{T}$ denotes the set of stopping times valued in the weeks of the Fall-Winter seasons $\{\frac{k}{48}, k = 0, 1, \dots, 24\}$, which corresponds to possible weekly choices for the travel agency to launch the campaign. The above function v models the optimal expected benefit for the travel agency and the optimal τ gives the best launching time. We shall assume, for simplicity, that the processes X^1 and X^2 are uncorrelated since we do not expect much influence of the weather on the exchange rate or vice-versa.

The problem is tackled by formulating it as a dynamic programming one related to optimal stopping problems (as exposed in Section 3) using a mean-reversion process for the underlying process X^1 and a drifted Brownian motion for X^2 . Their dynamics are given as follows:

$$dX_t^1 = -aX_t^1 dt + \sigma_1 dW_t, X_0^1 = 0, \quad X_t^2 = -\frac{\sigma_2^2}{2}t + \sigma_2 B_t.$$

	$K = 10$	$K = 20$	$K = 50$	$K = 100$
$M = 20$	0.1827	0.0512	0.0349	0.0269
$M = 40$	0.1982	0.0361	0.0249	0.0114
$M = 80$	0.2063	0.0325	0.0051	0.0047
$M = 160$	0.1928	0.0264	0.0058	0.0067

TABLE 5

Average squared L^2 errors with 20 macro runs. Simple regression.

	$K = 10$	$K = 20$	$K = 50$	$K = 100$
$M = 20$	0.1711	0.0458	0.0436	0.0252
$M = 40$	0.1648	0.0361	0.0130	0.0169
$M = 80$	0.1534	0.0273	0.0109	0.0085
$M = 160$	0.1510	0.0296	0.0048	0.0058

TABLE 6

Average squared L^2 errors with 20 macro runs. Nested regression.

The cost function c is chosen piecewise linear so that we can get $\mathbb{E}(c(e^{X_\tau^2+1/12})|X_\tau^2)$ explicitly as a function of X_τ^2 using the Black-Scholes formula in mathematical finance. Thus we can run our method in two different ways: either using this explicit expression and apply directly the regression scheme of Section 3; or first estimating $\mathbb{E}(c(e^{X_\tau^2+1/12})|X_\tau^2)$ by stratified regression then plugging the estimate in our method again to get a final estimation. We refer to these two different ways as simple regression and nested regression. The latter case corresponds to a coupled two-component regression problem (that could be mathematically analyzed very similarly to Section 3).

The parameter's values are given as: $a = 2, \sigma_1 = 10, \sigma_2 = 0.2, \underline{c} = 3, x_{min}^2 = e^{-0.5}, x_{max}^2 = e^{0.5}, c_{min} = 1, c_{max} = c_{min} + x_{max}^2 - x_{min}^2, t_{min} = 0, t_{max} = 15, q_{min} = 1, q_{max} = 4$. We use the restriction of $\mu(dx) = \frac{k}{2}(1+|x|)^{-k-1}dx$ with $k = 6$ to sample point for X^1 and the restriction of $\nu(dx) = \frac{1}{2}e^{-|x|}dx$ to sample points for X^2 . Note that $k = 6$ means that, in the error estimation, more weight is distributed to the region around $X_0^1 = 0$, which is the real interesting information for the travel agency.

We will firstly run our method with $M = 320$ and $K = 300$ to get a reference value for v then our estimators will be compared to this reference value in a similar way as in the previous example. The squared $L^2(\mu \otimes \nu)$ norm of our reference estimation is 32.0844. The results are displayed in the Tables 5 and 6.

As in Subsection 5.1 and in agreement with Theorem 3.4, we observe an improved accuracy as K and M increases, independently of each other. The relative error is rather small even for small M .

Interestingly, the nested regression algorithm (which is the most realistic scheme to use in practice when the model is unknown) is as accurate as the scheme using the explicit form of the internal conditional expectation $\mathbb{E}[c(e^{X_\tau^2+1/12}) | X_\tau^2]$. Surprisingly, the simple regression scheme takes much more time than the nested regression one because of the numerous evaluations of the Gaussian CDF in the Black-Scholes formula.

Appendix A. Appendix.

A.1. Proof of Proposition 3.3. Consider first the case of the partitioning estimate $(\mathbf{LP}_0)$ and let $\varepsilon \in (0, \frac{4}{15}B]$. We use an ε -cover in the L^∞ -norm, which

simply reduces to cover $[-B, B]$ with intervals of size 2ϵ . A solution is to take the interval center defined by $h_j = -B + \epsilon + 2\epsilon j$, $0 \leq j \leq n$, where n is the smallest integer such that $h_n \geq B$ (i.e., $n = \lceil \frac{B}{\epsilon} - \frac{1}{2} \rceil$). Thus, we obtain

$$\mathcal{N}_1(\epsilon, \mathcal{T}_B \mathcal{L}_k, x^{1:M}) \leq n + 1 \leq \frac{B}{\epsilon} + \frac{3}{2} \leq \frac{7}{5} \frac{B}{\epsilon}$$

604 where we use the constraint on ϵ .

In the case of general vector space of dimension K , from [10, Lemma 9.2, Theorem 9.4 and Theorem 9.5], we obtain

$$\mathcal{N}_1(\epsilon, \mathcal{T}_B \mathcal{K}, x^{1:M}) \leq 3 \left(\frac{4eB}{\epsilon} \log \left(\frac{6eB}{\epsilon} \right) \right)^{K+1}$$

whenever $\epsilon < B/2$. For ϵ as in the statement of Assumption 3.3, we have $\frac{6eB}{\epsilon} \geq \frac{45e}{2}$. Let $\eta > 0$, since $\log(x) \leq c_\eta x^\eta$ for any $x \geq \frac{45e}{2}$ with $c_\eta = \sup_{x \geq \frac{45e}{2}} \frac{\log(x)}{x^\eta}$, we get

$$\mathcal{N}_1(\epsilon, \mathcal{T}_B \mathcal{K}, x^{1:M}) \leq 3 \left([4c_\eta 6^\eta]^{1/(1+\eta)} \frac{eB}{\epsilon} \right)^{(K+1)(1+\eta)}.$$

605 For $\mathbf{LP}_1$ and $\mathbf{LP}_n$, we have respectively $K = d + 1$ and $K = (d + 1)^n$, therefore the
606 announced result. Whenever useful, the choice $\eta = 1$ gives $\beta_{(3.3)} \leq 3.5$.

607 For the partitioning estimate (case $\mathbf{LP}_0$), we could also use this estimate with
608 $K = 1$ but with the first arguments, we get better parameters (especially for γ).

609 A.2. Probability of uniform deviation.

610 LEMMA A.1 ([9, Lemma B.2]). Let $\mathcal{G}$ be a countable set of functions $g : \mathbb{R}^d \mapsto$
611 $[0, B]$ with $B > 0$. Let $\mathcal{X}, \mathcal{X}^{(1)}, \dots, \mathcal{X}^{(M)}$ ($M \geq 1$) be i.i.d. $\mathbb{R}^d$ valued random
612 variables. For any $\alpha > 0$ and $\epsilon \in (0, 1)$ one has

$$\begin{aligned} 613 \quad & \mathbb{P} \left(\sup_{g \in \mathcal{G}} \frac{\frac{1}{M} \sum_{m=1}^M g(\mathcal{X}^{(m)}) - \mathbb{E}[g(\mathcal{X})]}{\alpha + \frac{1}{M} \sum_{m=1}^M g(\mathcal{X}^{(m)}) + \mathbb{E}[g(\mathcal{X})]} > \epsilon \right) \\ 614 \quad & \leq 4\mathbb{E} \left[\mathcal{N}_1 \left(\frac{\alpha\epsilon}{5}, \mathcal{G}, \mathcal{X}^{1:M} \right) \right] \exp \left(- \frac{3\epsilon^2 \alpha M}{40B} \right), \\ 615 \quad & \mathbb{P} \left(\sup_{g \in \mathcal{G}} \frac{\mathbb{E}[g(\mathcal{X})] - \frac{1}{M} \sum_{m=1}^M g(\mathcal{X}^{(m)})}{\alpha + \frac{1}{M} \sum_{m=1}^M g(\mathcal{X}^{(m)}) + \mathbb{E}[g(\mathcal{X})]} > \epsilon \right) \\ 616 \quad & \leq 4\mathbb{E} \left[\mathcal{N}_1 \left(\frac{\alpha\epsilon}{8}, \mathcal{G}, \mathcal{X}^{1:M} \right) \right] \exp \left(- \frac{6\epsilon^2 \alpha M}{169B} \right). \end{aligned}$$

618 A.3. Expected uniform deviation.

619 PROPOSITION A.1. For finite $B > 0$, let $\mathcal{G} := \{\psi(\mathcal{T}_B \phi(\cdot) - \eta(\cdot)) : \phi \in \mathcal{K}\}$, where
620 $\psi : \mathbb{R} \rightarrow [0, \infty)$ is Lipschitz continuous with $\psi(0) = 0$ and Lipschitz constant L_ψ ,
621 $\eta : \mathbb{R}^d \rightarrow [-B, B]$, and $\mathcal{K}$ is a finite K -dimensional vector space of functions with

$$622 \quad (\text{A.1}) \quad \mathcal{N}_1(\epsilon, \mathcal{T}_B \mathcal{K}, \mathcal{X}^{1:M}) \leq \alpha \left(\frac{\beta B}{\epsilon} \right)^\gamma \quad \text{for } \epsilon \in (0, \frac{4}{15} B]$$

623 for some positive constants α, β, γ with $\alpha \geq 1/4$ and $\gamma \geq 1$. Then, for $\mathcal{X}^{(1)}, \dots, \mathcal{X}^{(M)}$
624 i.i.d. copies of $\mathcal{X}$, we have

$$625 \quad \mathbb{E} \left[\sup_{g \in \mathcal{G}} \left(\frac{1}{M} \sum_{m=1}^M g(\mathcal{X}^{(m)}) - 2 \int_{\mathbb{R}^d} g(x) \mathbb{P} \circ \mathcal{X}^{-1}(\mathrm{d}x) \right)_+ \right] \leq c_{(\text{A.2})}(M) \frac{BL_\Psi}{M}$$

(A.2)

$$\text{with } c_{(\text{A.2})}(M) := 120 \left(1 + \log(4\alpha) + \gamma \log \left(\left(1 + \frac{\beta}{16} \right) M \right) \right),$$

$$\mathbb{E} \left[\sup_{g \in \mathcal{G}} \left(\int_{\mathbb{R}^d} g(x) \mathbb{P} \circ \mathcal{X}^{-1}(\mathrm{d}x) - \frac{2}{M} \sum_{m=1}^M g(\mathcal{X}^{(m)}) \right)_+ \right] \leq c_{(\text{A.3})}(M) \frac{BL_\Psi}{M}$$

(A.3)

$$\text{with } c_{(\text{A.3})}(M) := \frac{507}{2} \left(1 + \log(4\alpha) + \gamma \log \left(\left(1 + \frac{8\beta}{169} \right) M \right) \right).$$

Proof. The idea is to adapt the arguments of [9, Proposition 4.9].

▷ We first show (A.2). Set $\mathcal{Z} := \sup_{g \in \mathcal{G}} \left(\frac{1}{M} \sum_{m=1}^M g(\mathcal{X}^{(m)}) - 2 \int_{\mathbb{R}^d} g(x) \mathbb{P} \circ \mathcal{X}^{-1}(\mathrm{d}x) \right)_+$. Let us find an upper bound for $\mathbb{P}(\mathcal{Z} > \varepsilon)$ in order to bound $\mathbb{E}[\mathcal{Z}] = \int_0^\infty \mathbb{P}(\mathcal{Z} > \varepsilon) \mathrm{d}\varepsilon$. Using the equality

$$\mathbb{P}(\mathcal{Z} > \varepsilon) = \mathbb{P} \left(\exists g \in \mathcal{G} : \frac{\frac{1}{M} \sum_{m=1}^M g(\mathcal{X}^{(m)}) - \int_{\mathbb{R}^d} g(x) \mathbb{P} \circ \mathcal{X}^{-1}(\mathrm{d}x)}{2\varepsilon + \int_{\mathbb{R}^d} g(x) \mathbb{P} \circ \mathcal{X}^{-1}(\mathrm{d}x) + \frac{1}{M} \sum_{m=1}^M g(\mathcal{X}^{(m)})} > \frac{1}{3} \right),$$

and that the elements of $\mathcal{G}$ take values in $[0, 2BL_\psi]$, it follows from Lemma A.1 that

$$\mathbb{P}(\mathcal{Z} > \varepsilon) \leq 4\mathbb{E} \left[\mathcal{N}_1 \left(\frac{2\varepsilon}{15}, \mathcal{G}, \mathcal{X}^{1:M} \right) \right] \exp \left(- \frac{\varepsilon M}{120BL_\psi} \right).$$

Define $\mathcal{T}_B \mathcal{K}$ as in Proposition A.1. Since $|\psi(\phi_1(x) - \eta(x)) - \psi(\phi_2(x) - \eta(x))| \leq L_\psi |\phi_1(x) - \phi_2(x)|$ for all $x \in \mathbb{R}^d$ and all (ϕ_1, ϕ_2) , it follows that

$$\mathcal{N}_1 \left(\frac{2\varepsilon}{15}, \mathcal{G}, \mathcal{X}^{1:M} \right) \leq \mathcal{N}_1 \left(\frac{2\varepsilon}{15L_\psi}, \mathcal{T}_B \mathcal{K}, \mathcal{X}^{1:M} \right).$$

Due to Equation (A.1), we deduce that

$$(A.4) \quad \mathbb{P}(\mathcal{Z} > \varepsilon) \leq 4\alpha \left(\frac{15\beta BL_\psi}{2\varepsilon} \right)^\gamma \exp \left(- \frac{\varepsilon M}{120BL_\psi} \right)$$

whenever $\frac{2\varepsilon}{15L_\psi} \leq \frac{4}{15}B$, i.e., $\varepsilon \leq 2BL_\Psi$. On the other hand, $\mathbb{P}(\mathcal{Z} > \varepsilon) = 0$ for all $\varepsilon > 2BL_\Psi$. Setting $a = \frac{15\beta BL_\psi}{2}$, $b = \frac{1}{120BL_\psi}$, it follows from (A.4) that

$$\mathbb{P}(\mathcal{Z} > \varepsilon) \leq 4\alpha \left(\frac{a}{\varepsilon} \right)^\gamma \exp(-bM\varepsilon), \quad \forall \varepsilon > 0.$$

Fix ε_0 to be some finite value (to be determined later) such that

$$(A.5) \quad \varepsilon_0 \geq \frac{a}{M(1+ab)}.$$

It readily follows that

$$\begin{aligned} \mathbb{E}[\mathcal{Z}] &= \int_0^\infty \mathbb{P}(\mathcal{Z} > \varepsilon) \mathrm{d}\varepsilon \leq \varepsilon_0 + \int_{\varepsilon_0}^\infty 4\alpha \left(\frac{a}{\varepsilon} \right)^\gamma \exp(-bM\varepsilon) \mathrm{d}\varepsilon \\ &\leq \varepsilon_0 + \frac{4\alpha}{bM} (M(1+ab))^\gamma \exp(-bM\varepsilon_0). \end{aligned}$$

We choose $\varepsilon_0 = \frac{1}{bM} \log \left(4\alpha((1+ab)M)^\gamma \right)$: It satisfies (A.5) since

$$\frac{1}{bM} \log \left(4\alpha((1+ab)M)^\gamma \right) \geq \frac{a}{M} \frac{\log(1+ab)}{ab} \geq \frac{a}{M} \frac{1}{1+ab}$$

(use $\alpha \geq 1/4$, $\gamma \geq 1$, $M \geq 1$ and $\log(1+x) \geq x/(1+x)$ for all $x \geq 0$). Moreover, this choice of ε_0 implies that

$$\begin{aligned} \mathbb{E}[Z] &\leq \frac{1}{bM} \left(1 + \log(4\alpha) + \gamma \log((1+ab)M) \right) \\ &= \frac{120BL_\psi}{M} \left(1 + \log(4\alpha) + \gamma \log \left(\left(1 + \frac{\beta}{16} \right) M \right) \right). \end{aligned}$$

The inequality (A.2) is proved.

▷ We now justify (A.3) by similar arguments. Set $\mathcal{Z} := \sup_{g \in \mathcal{G}} \left(\int_{\mathbb{R}^d} g(x) \mathbb{P} \circ \mathcal{X}^{-1}(dx) - \frac{2}{M} \sum_{m=1}^M g(\mathcal{X}^{(m)}) \right)_+$. From Lemma A.1, we get

$$\mathbb{P}(\mathcal{Z} > \varepsilon) \leq 4\mathbb{E} \left[\mathcal{N}_1 \left(\frac{\varepsilon}{12}, \mathcal{G}, \mathcal{X}^{1:M} \right) \right] \exp \left(- \frac{2\varepsilon M}{507BL_\psi} \right).$$

Since $\mathcal{N}_1 \left(\frac{\varepsilon}{12}, \mathcal{G}, \mathcal{X}^{1:M} \right) \leq \mathcal{N}_1 \left(\frac{\varepsilon}{12L_\psi}, \mathcal{T}_B \mathcal{K}, \mathcal{X}^{1:M} \right)$ and thanks to (A.1), we derive

$$\mathbb{P}(\mathcal{Z} > \varepsilon) \leq 4\alpha \left(\frac{12\beta BL_\psi}{\varepsilon} \right)^\gamma \exp \left(- \frac{2\varepsilon M}{507BL_\psi} \right)$$

whenever $\frac{\varepsilon}{12L_\psi} \leq \frac{4}{15}B$. For other values of ε the above probability is zero, therefore (A.7) holds for any $\varepsilon > 0$. The end of the computations is now very similar to the previous case: we finally get the inequality (A.6) for the new $\mathcal{Z}$ with adjusted values $a = 12\beta BL_\psi$, $b = \frac{2}{507BL_\psi}$. Thus inequality (A.3) is thus proved. $\square$

REFERENCES

- [1] M. ABLOWITZ AND A. ZEPPELELLA, *Explicit solutions of Fisher's equation for a special wave speed*, Bull. Math. Biol., 41 (1979), pp. 835–840.
- [2] A. BACHOUCH, E. GOBET, AND A. MATOUSSI, *Empirical regression method for backward doubly stochastic differential equations*, To appear in SIAM ASA Journal on Uncertainty Quantification, (2015).
- [3] N. H. BARTON AND M. TURELLI, *Spatial waves of advance with bistable dynamics: Cytoplasmic and genetic analogues of allee effects*, The American Naturalist, 178 (2011), pp. E48–E75.
- [4] D. BELOMESTNY, A. KOLODKO, AND J. SCHOENMAKERS, *Regression methods for stochastic control problems and their convergence analysis*, SIAM Journal on Control and Optimization, 48 (2010), pp. 3562–3588.
- [5] C. BENDER AND R. DENK, *A forward scheme for backward SDEs*, Stochastic Processes and their Applications, 117 (2007), pp. 1793–1823.
- [6] J. CHASSAGNEUX AND D. CRISAN, *Runge-Kutta schemes for backward stochastic differential equations*, Ann. Appl. Probab., 24 (2014), pp. 679–720.
- [7] R. A. FISHER, *The wave of advance of advantageous genes*, Annals of Eugenics, 7 (1937), pp. 353–369.
- [8] E. GOBET, J. LOPEZ-SALAS, P. TURKEDJIEV, AND C. VASQUEZ, *Stratified regression Monte-Carlo scheme for semilinear PDEs and BSDEs with large scale parallelization on GPUs*, In revision for SIAM Journal of Scientific Computing, Hal preprint hal-01186000, (2015).
- [9] E. GOBET AND P. TURKEDJIEV, *Linear regression MDP scheme for discrete backward stochastic differential equations under general conditions*, Math. Comp., 85 (2016), pp. 1359–1391.

- [10] L. GYORFI, M. KOHLER, A. KRZYZAK, AND H. WALK, *A distribution-free theory of nonparametric regression*, Springer Series in Statistics, 2002.
- [11] A. HOFFMANN, B. MONTGOMERY, J. POPOVICI, I. ITURBE-ORMAETXE, P. JOHNSON, F. MUZZI, M. GREENFIELD, M. DURKAN, Y. LEONG, Y. DONG, ET AL., *Successful establishment of wolbachia in aedes populations to suppress dengue transmission*, Nature, 476 (2011), pp. 454–457.
- [12] A. KOLMOGOROFF, I. PRETROVSKY, AND N. PISCOUNOFF, *Étude de l'équation de la diffusion avec croissance de la quantité de matière et son application à un problème biologique*. Bull. Univ. État Moscou, Sér. Int., Sect. A: Math. et Mécan. 1, Fasc. 6, 1-25, 1937.
- [13] J. MA AND J. YONG, *Forward-Backward Stochastic Differential Equations*, Lecture Notes in Mathematics, 1702, Springer-Verlag, 1999. A course on stochastic processes.
- [14] J. D. MURRAY, *Mathematical biology. I*, vol. 17 of Interdisciplinary Applied Mathematics, Springer-Verlag, New York, third ed., 2002. An introduction.
- [15] J. D. MURRAY, *Mathematical biology. II*, vol. 18 of Interdisciplinary Applied Mathematics, Springer-Verlag, New York, third ed., 2003. Spatial models and biomedical applications.
- [16] E. PARDOUX AND A. RASCANU, *Stochastic Differential Equations, Backward SDEs, Partial Differential Equations*, vol. 69 of Stochastic Modelling and Applied Probability, Springer-Verlag, 2014.
- [17] M. POTTERS, J.-P. BOUCHAUD, AND D. SESTOVIC, *Hedged Monte Carlo: low variance derivative pricing with objective probabilities*, Physica A. Statistical Mechanics and its Applications, 289 (2001), pp. 517–525.
- [18] J. TSITSIKLIS AND B. V. ROY, *Regression methods for pricing complex american-style options*, IEEE Transactions on Neural Networks, 12 (2001), pp. 694–703.

Finance for Energy Market Research Centre

Institut de Finance de Dauphine, Université Paris-Dauphine

1 place du Maréchal de Lattre de Tassigny

75775 PARIS Cedex 16

www.fime-lab.org