

EUROPEAN OPTION PRICING WITH MODEL CONSTRAINED GAUSSIAN PROCESS REGRESSIONS

Donatien Hainaut, Frédéric Vrins

LIDAM Discussion Paper LFIN
2024 / 05

LFIN

Voie du Roman Pays 34, L1.03.01 B-1348 Louvain-la-Neuve

Tel (32 10) 47 43 04

Email: lidam-library@uclouvain.be

<https://uclouvain.be/en/research-institutes/lidam/lfin/publications.html>

European option pricing with model constrained Gaussian process regressions

Donatien Hainaut*

UCLouvain, ISBA-LIDAM

Frédéric Vrins†

UCLouvain, LFIN-LIDAM

We propose a method for pricing European options based on Gaussian processes. We convert the problem of solving the Feynman-Kac (FK) partial differential equation (PDE) into a model-constrained regression. We form two training sets by sampling state variables from the PDEs inner domain and boundary. The regression function is then estimated to fit the option values on the boundary sample while satisfying the FK PDE on the inner sample. We adopt a Bayesian framework in which payoffs and the value of the FK PDE in the boundary and inner samples are noised. Assuming the regression function is a Gaussian process, we find closed-form expressions of approximated option prices and their Greeks. We demonstrate the performance of the procedure on call options in the Black-Scholes and Heston models.

KEYWORDS: Gaussian process regression, Option pricing, Feynman-Kac equation, partial differential equation, Heston model, machine learning

1 Introduction

We address the problem of computing the price of European-style derivative securities in an arbitrage-free market. In such markets, the prices of non-dividend paying assets (including derivative securities up to their expiry) are governed by a partial differential equation (PDE) subject to a terminal condition. In this context, the price of the asset is given by a function f evaluated at the prevailing state variables, $\mathbf{x}$. The PDE governing f is known as the Feynman-Kac (FK) equation (see e.g. Jeanblanc et al. 2009, section 2.2.3) and takes the form $\mathcal{L}_{\mathbf{x}}f(\mathbf{x}) = z(\mathbf{x})$ where $\mathcal{L}_{\mathbf{x}}$ is a linear differential operator. The boundary constraint reads as $f(\bar{\mathbf{x}}) = h(\bar{\mathbf{x}})$ where h is the value of derivative and $\bar{\mathbf{x}}$ is on the PDE boundary. The former captures the dynamics of the market model at hand while the latter uniquely characterizes the considered option. As an illustration, the time- t price of a European call with strike K and expiry $T \geq t$ written on a stock whose price process $S = (S_t)_{t \geq 0}$ follows a geometric Brownian motion with volatility σ and assuming a constant risk-free rate leads to $V_t = f(\mathbf{x})$ with $\mathbf{x} = (t, S_t)$ for $t < T$, $z(\mathbf{x}) = 0$,

$$\mathcal{L}_{\mathbf{x}}f(t, x) = \frac{\partial}{\partial t}f(t, x) + rx\frac{\partial}{\partial x}f(t, x) + \frac{\sigma^2 x^2}{2}\frac{\partial^2}{\partial x^2}f(t, x) - rf(t, x).$$

*Corresponding author. Postal address: Voie du Roman Pays 20, 1348 Louvain-la-Neuve (Belgium). E-mail to: donatien.hainaut(at)uclouvain.be

†E-mail to: frederic.vrins@uclouvain.be

The option values on the boundary are respectively null for $\bar{\mathbf{x}} = (t, 0)$ and equal to $(S_T - K)_+$ for $\bar{\mathbf{x}} = (T, S_T)$. Except in some special cases like the Black-Scholes model, it is not possible to solve this equation analytically. Computing the price of simple European derivative securities in a stochastic volatility model, for instance, requires numerical methods such as lattice, Laplace transform or Monte Carlo methods. Although efficient algorithms have been designed for this purpose, the most efficient ones are model-specific (e.g., they require the knowledge of the characteristic function). Moreover, they can suffer from convergence or stability issues, and often provide disappointing estimations of the sensitivities (Greeks). In this article, we propose an efficient method to tackle this problem in its general form by leveraging Constrained Gaussian Process Regression (CGPR).

In this paper we propose to solve the FK equation by considering a constrained regression model. To this end, we assume to have two samples of points, $\mathring{\mathbf{X}}$ and $\bar{\mathbf{X}}$, respectively in the inner domain of the PDE and on the boundary. We seek for a regression function $\hat{f}(\mathbf{x})$, approximating f subject to

$$\mathcal{L}_{\mathbf{x}} \hat{f}(\mathbf{x}) = z(\mathbf{x}) + \epsilon \quad \forall \mathbf{x} \in \mathring{\mathbf{X}} \quad \text{and} \quad \hat{f}(\bar{\mathbf{x}}) = h(\bar{\mathbf{x}}) + \bar{\epsilon} \quad \forall \bar{\mathbf{x}} \in \bar{\mathbf{X}}$$

for some independent Gaussian noises $\epsilon, \bar{\epsilon}$. In a constrained Gaussian process regression, we are looking for the random functions $\hat{f}$ taking the form of a Gaussian process. One can thus take the expectation of $\hat{f}$ as an approximation of the true price function, f .

The unconstrained Gaussian process regression (GPR) is a powerful tool for regression that has been successfully applied in many different fields. We refer the reader to Rasmussen and Williams (2006, Chapter 6) and to (2012, Chapter 15) for a general presentation. Gaussian processes (GPs) are also used to regress responses to explanatory variables constrained to satisfy a partial differential equation (PDE). Similarly to Physics-Inspired Neural Networks (PINNs), this PDE represents the physical law governing the experiment for which responses are measured. There is a vast literature on this topic, and we refer to Swiler et al. (2020) for a survey of CGPRs, including constraints provided by linear PDEs and boundary conditions. Among the useful references therein, we highlight Graepel (2003), who approximates the solution of PDEs with noisy responses using a CGPR. Calderhead et al. (2009) present an accelerated sampling procedure to solve nonlinear ordinary and delay differential equations with Gaussian processes. Cialenco et al. (2012) combine an Euler discretization scheme on the time axis with a CGPR. Dondelinger (2013), Barber and Wang (2014), and Cockayne et al. (2019) propose a CGPR that improves the gradient matching technique. Särkkä (2011) and Nguyen and Peraire (2015) develop procedures for handling responses of the PDE solution, given in the form of linear functionals of the state variables. Raissi et al. (2017) use an approach based on GPs to discover governing equations expressed by parametric linear operators. Pförtner et al. (2022) consider physics-informed Gaussian process regression for solving both weak and strong formulations of linear PDEs. Chen et al. (2021) propose a method to solve nonlinear PDEs based on the maximum a posteriori estimator of a Gaussian process.

When it comes to finance, GPR is mainly used for time-series forecasting. Wu et al. (2014) introduce a non-parametric model for time-changing variances based on Gaussian processes. Han et al. (2016) combine a Gaussian process framework with a stochastic volatility (SV) model for stock prices. Petelin et al. (2019) fit a GPR to daily data from U.S. commodity markets. Gonzalvez et al. (2019) present a method using GPR for fitting the U.S. yield curve and portfolio allocation. GPs are also used for the regression of risk measures or derivatives, computed with a model, on their parameters to speed up the valuation procedure. For instance, Crépey and Dixon (2019) use GPR for derivative portfolio modeling and computing credit valuation adjustments. Wilkens (2019) proposes an approach based on GPR for forecasting risk management indicators. Goudenège et al. (2021) fit a GP to regress prices of variable annuities on market risk factors. Closer to our work is De Spiegeleer et al. (2018), who use GPR to

predict option prices valued with trees or fast Fourier transforms from market parameters. Our approach is different, however, because we exploit the FK PDE. In particular, De Spiegeleer et al. (2018) build their regression on a sample of prices produced via other pricing techniques (such as FFT or Monte Carlo). By contrast, we only need to sample the state variables in the inner domain ($\hat{\mathbf{X}}$) and on the boundary ($\bar{\mathbf{X}}$), ensuring that the regression function complies with the payoff and the PDE at those points, respectively.

This work contributes to the financial literature in several ways. To our knowledge, this article is among the first to propose a method for pricing derivatives by solving the FK equation with a CGPR. Variants of the CGPR have been successfully applied in mechanics and engineering (see references above) but have not yet been applied to option valuation. Physics-Inspired Neural Networks (PINNs) also aim to solve PDE-constrained regression and have been considered for option pricing (see Sirignano and Spiliopoulos, 2018; Al-Arabi et al., 2022; Hainaut, 2024 a; and Hainaut and Casas, 2024). However, PINNs rely on neural networks instead of GPs. Second, by taking advantage of the FK equation, we do not rely on numerical differentiation, in contrast to PINN-based methods. Third, compared to classical finite difference approaches, GPR naturally justifies adding a Tikhonov regularization factor, which helps tune the robustness of numerical computations. Fourth, we improve upon the approach of Graepel (2003), which is designed for solving PDEs with a nullity constraint on the boundary. For other constraints, Graepel (2003) assumes that the solution is the sum of an a priori chosen function, which fulfills the terminal condition, and a GP that solves a PDE with a nullity constraint. Choosing the function that satisfies the boundary constraints a priori is the critical step of this method. In contrast, we sample state variables on the boundary as well and consider a CGPR model that fits the boundary condition within this sample. Finally, we provide fully analytical expressions for CGPR option prices and their Greeks using a Gaussian kernel.

The outline of the paper is as follows. Section 2 presents a general market model in a Brownian setting and a general version of the FK equation. The next section reviews the classic framework of GPR and introduces the PDE-constrained extension (CGPR). Section 4 provides detailed developments for solving this FK equation using a Bayesian approach. It also presents a general analytical expressions of approximated prices and Greeks when the kernel is Gaussian. Sections 5 and 6 test the efficiency of our method in two particular cases. First, we evaluate the capacity of the CGPR, for pricing a call option in a Black and Scholes market. Next, we price European call options in the Heston stochastic volatility model, using the Monte-Carlo method as benchmark.

2 The Feynman-Kac equation

We consider a financial market driven by $p - 1$ risk processes, stored in a vector $(\mathbf{y}_t)_{t \geq 0}$. These processes are defined on a complete probability space endowed with a filtration $\mathbb{F} = (\mathcal{F}_t)_{t \geq 0}$ and a risk neutral probability measure, $\mathbb{Q}$. The risk-free rate is function of state variables and denoted by $r_t = r(t, \mathbf{y}_t)$. Asset prices and risk factors (e.g. the stochastic variance) are ruled by the following stochastic differential equation (SDE)¹:

$$d\mathbf{y}_t = \boldsymbol{\mu}_{\mathbf{y}}(t, \mathbf{y}_t)dt + \Sigma_{\mathbf{y}}(t, \mathbf{y}_t)d\mathbf{B}_t, \quad (1)$$

where $\mathbf{B} = (\mathbf{B}_t)_{t \geq 0}$ is a $(p - 1)$ -vector of independent Brownian motions, $\boldsymbol{\mu}_{\mathbf{y}}(\cdot)$ is a vector of dimension and $\Sigma_{\mathbf{y}}(\cdot)$ is a $(p - 1) \times (p - 1)$ matrix such that

$$\begin{aligned} \mathbb{E} \left(\int_0^t \boldsymbol{\mu}_{\mathbf{y}}(s, \mathbf{y}_s) ds \right) &\leq \infty, \quad \forall t < \infty, \\ \mathbb{E} \left(\int_0^t \Sigma_{\mathbf{y}}(s, \mathbf{y}_s) \Sigma_{\mathbf{y}}(s, \mathbf{y}_s)^\top ds \right) &\leq \infty, \quad \forall t < \infty \end{aligned}$$

¹Note that if y is a tradeable asset, the corresponding drift takes the form $\mu_y(t, \mathbf{y}_t) = r_t y_t$.

where $\mathbb{E}$ is the expectation operator under $\mathbb{Q}$. We consider a European derivative security expiring at time T with payoff $\mathcal{P}(T, \mathbf{y}_T)$. We assume that its price at any time $t \leq T$ is a function of the current value of the state variables, i.e., $f(t, \mathbf{y}_t)$ where f is referred to as the *pricing function*, $f(t, \mathbf{y}_t)$ of the derivative. Under the assumption of arbitrage-free market, all traded securities paying no cash-flow are $(\mathbb{Q}, \mathbb{F})$ -martingales. Hence, the pricing function satisfies

$$\mathbb{E} \left(e^{-\int_s^t r(u, \mathbf{y}_u) du} f(t, \mathbf{y}_t) \mid \mathcal{F}_s \right) = f(s, \mathbf{y}_s) \quad \forall s \leq t \leq T \quad (2)$$

subject to the terminal boundary condition determined by the payoff

$$f(T, \mathbf{y}_T) = \mathcal{P}(T, \mathbf{y}_T) \quad \forall \mathbf{y}_T. \quad (3)$$

An immediate consequence of (2) is that the expected rate of return of the derivative under $\mathbb{Q}$ is the risk-free rate:

$$\mathbb{E}(df(t, \mathbf{y}_t) \mid \mathcal{F}_t) = r(t, \mathbf{y}_t) f(t, \mathbf{y}_t) dt. \quad (4)$$

Computing the left-hand side of (4) where $df(t, \mathbf{y}_t)$ is found using Itô's lemma yields the Feynman-Kac PDE:

$$\partial_t f + \boldsymbol{\mu}_{\mathbf{y}}(t, \mathbf{y}_t)^\top \nabla_{\mathbf{y}} f + \frac{1}{2} \text{tr} \left(\Sigma_{\mathbf{y}}(t, \mathbf{y}_t) \Sigma_{\mathbf{y}}(t, \mathbf{y}_t)^\top \mathcal{H}_{\mathbf{y}} f \right) = r(t, \mathbf{y}_t) f, \quad (5)$$

where $\nabla_{\mathbf{y}} f$ is the gradient of f with respect to $\mathbf{y}$, $\mathcal{H}_{\mathbf{y}} f$ is the Hessian and $\text{tr}(\cdot)$ is the trace operator. To lighten further developments, we note $\mathbf{x} = (t, \mathbf{y}_t)$, the p -vector of time and state variables at time t and define the linear² differential operator associated with the FK equation:

$$\mathcal{L}_{\mathbf{x}} f(\mathbf{x}) = \partial_t f - r(\mathbf{x}) f + \boldsymbol{\mu}_{\mathbf{y}}(\mathbf{x})^\top \nabla_{\mathbf{y}} f + \frac{1}{2} \text{tr} \left(\Sigma_{\mathbf{y}}(\mathbf{x}) \Sigma_{\mathbf{y}}(\mathbf{x})^\top \mathcal{H}_{\mathbf{y}} f \right). \quad (6)$$

We denote by $\mathcal{X} = [0, T] \times \mathbb{R}_0^{p-1}$ the domain on which the pricing function f is defined. The boundary is $\bar{\mathcal{X}} = \bar{\mathcal{X}}_b \cup \bar{\mathcal{X}}_T$ where

$$\bar{\mathcal{X}}_T = \{(T, s) \mid s \in \mathbb{R}^+\}$$

is the set of $\bar{\mathbf{x}}$ such the option is equal to the payoff $f(\bar{\mathbf{x}}) = \mathcal{P}(\bar{\mathbf{x}})$. Whereas $\bar{\mathcal{X}}_b$ is the set of $\bar{\mathbf{x}}$ such that $f(\bar{\mathbf{x}}) = b(\bar{\mathbf{x}})$ where b is a known function (e.g. $b(\bar{\mathbf{x}}) = 0$). $\mathring{\mathcal{X}}$, that we call the interior domain, is such that $\mathcal{X} = \mathring{\mathcal{X}} \cup \bar{\mathcal{X}}$. We infer that the pricing function f of any European derivative with a payoff function h , solves

$$\begin{cases} \mathcal{L}_{\mathbf{x}} f(\mathbf{x}) = 0 & \mathbf{x} \in \mathring{\mathcal{X}}, \\ f(\bar{\mathbf{x}}) = h(\bar{\mathbf{x}}) & \bar{\mathbf{x}} \in \bar{\mathcal{X}}, \end{cases} \quad (7)$$

where $h(\bar{\mathbf{x}})$ is the derivative value function on $\bar{\mathcal{X}}$:

$$h(\bar{\mathbf{x}}) = \begin{cases} b(\bar{\mathbf{x}}) & \mathbf{x} \in \bar{\mathcal{X}}_b \\ \mathcal{P}(\bar{\mathbf{x}}) & \bar{\mathbf{x}} \in \bar{\mathcal{X}}_T \end{cases}.$$

In later developments, $h(\bar{\mathbf{x}})$ is called the boundary function.

3 The constrained Gaussian process regression (CGPR)

Solving the PDE (7) can be viewed as a regression problem constrained by a PDE, with its structure depending on the stochastic model chosen for the market. In this work, we propose an approach based on Gaussian process regression. We will briefly review the main features of this type of regression before introducing the PDE-constrained extension.

²By linear operator, we mean that $\mathcal{L}_{\mathbf{x}}(g^{(1)}(\mathbf{x}) + g^{(2)}(\mathbf{x})) = \mathcal{L}_{\mathbf{x}}g^{(1)}(\mathbf{x}) + \mathcal{L}_{\mathbf{x}}g^{(2)}(\mathbf{x})$.

3.1 The Gaussian process regression (GPR)

A Gaussian process is a collection noted $\{g(\mathbf{x}), \mathbf{x} \in \mathcal{X}\}$ such that for any $d \in \mathbb{N}$ and $\mathbf{x}_1, \dots, \mathbf{x}_d \in \mathcal{X}$, the random vector $(g(\mathbf{x}_1), \dots, g(\mathbf{x}_d))$ has a joint multivariate Gaussian distribution, characterized by a mean and a covariance function. Without loss of generality, the mean is set to zero and the covariance function is defined by a kernel $k(\mathbf{x}, \mathbf{x}')$:

$$\mathbb{C}(g(\mathbf{x}), g(\mathbf{x}')) = k(\mathbf{x}, \mathbf{x}') \quad \forall (\mathbf{x}, \mathbf{x}') \in \mathcal{X} \times \mathcal{X}.$$

The value of the kernel $k(\mathbf{x}, \mathbf{x}')$ controls the interaction between the states $(\mathbf{x}, \mathbf{x}')$. A necessary and sufficient condition for the function $k : \mathcal{X} \times \mathcal{X} \rightarrow \mathbb{R}$ to be a valid kernel is that the $d \times d$ matrix of $(k(\mathbf{x}_i, \mathbf{x}_j))_{i,j=1,\dots,d}$ is positive semidefinite for all possible samples, i.e., $\sum_{i=1}^d k(\mathbf{x}_i, \mathbf{x}_j) c_i c_j \geq 0$ holds for any $d \in \mathbb{N}$, $(\mathbf{x}_i)_{i=1,\dots,d} \in \mathcal{X}$ and $c_1, \dots, c_d \in \mathbb{R}$.

The general principle of GPR is to provide an approximation of a function $h(\mathbf{x}) : \mathcal{X} \rightarrow \mathbb{R}$ whose functional form is unknown but whose values are observed at certain points. The training set contains covariates, $\mathbf{X} = \{\mathbf{x}_i\}_{i=1,\dots,d}$, randomly sampled in $\mathcal{X}$, and noised responses $\mathbf{h} = (h_i)_{i=1,\dots,d}$ with $h_i = h(\mathbf{x}_i) + \epsilon$, $\epsilon \sim \mathcal{N}(0, \sigma_*^2)$. From a Bayesian perspective, we encode our belief that instances of $h(\mathbf{x})$ are drawn from a Gaussian process g with zero mean and covariance function $k(\mathbf{x}, \mathbf{x}')$, prior to taking into account observations. The covariance matrix of g on the training set is given by the kernel matrix

$$k(\mathbf{X}, \mathbf{X}) = (k(\mathbf{x}_i, \mathbf{x}_j))_{i,j=1,\dots,d}.$$

Our regression problem thus consists in finding a Gaussian process g such that

$$g(\mathbf{x}_i) = h_i \quad \forall \mathbf{x}_i \in \mathbf{X}. \quad (8)$$

In the framework of the GPR, $\mathbf{h}$ is the realization of d -dimensional normal denoted by $\mathbf{H}$, which is the Gaussian process at points $(\mathbf{x}_i)_{i=1,\dots,d}$, noised by $\epsilon \sim \mathcal{N}(0, \sigma_*^2)$. Therefore, the a priori joint-distribution of $g(\mathbf{x})$ for $\mathbf{x} \in \mathcal{X} \setminus \mathbf{X}$ and $\mathbf{H}$ is :

$$\begin{pmatrix} g(\mathbf{x}) \\ \mathbf{H} \end{pmatrix} \sim \mathcal{N} \left(\mathbf{0}; \begin{pmatrix} k(\mathbf{x}, \mathbf{x}) & k(\mathbf{X}, \mathbf{x}) \\ k(\mathbf{X}, \mathbf{x})^\top & \sigma_*^2 I_d + k(\mathbf{X}, \mathbf{X}) \end{pmatrix} \right), \quad (9)$$

where $k(\mathbf{X}, \mathbf{x}) = (k(\mathbf{x}_i, \mathbf{x}))_{i=1,\dots,d}$. Equation (9) should be regarded as the prior of $(g(\mathbf{x}), \mathbf{H})$. We consider as estimator $\hat{h}(\mathbf{x})$ of $h(\mathbf{x})$ the conditional expectation of g given $\mathbf{H} = \mathbf{h}$, which is found thanks to the posterior density of $g(\mathbf{x}) \mid \mathbf{h}$:

$$\begin{aligned} \hat{h}(\mathbf{x}) &= \mathbb{E}(g(\mathbf{x}) \mid \mathbf{h}) \\ &= k(\mathbf{x}, \mathbf{X})^\top (\sigma_*^2 I_d + k(\mathbf{X}, \mathbf{X}))^{-1} \mathbf{h}. \end{aligned}$$

Note that assuming noisy observations of $h(\mathbf{x}_i)$ for $\mathbf{x}_i \in \mathbf{X}$ provides regularization via a linear shrinkage of the covariance matrix $k(\mathbf{X}, \mathbf{X})$, analogous to Ridge estimators. We also infer from (9) that the conditional variance of the estimator does not depend on $\mathbf{h}$, and is equal to:

$$\mathbb{V}(g(\mathbf{x}) \mid \mathbf{h}) = k(\mathbf{x}, \mathbf{x}) - k(\mathbf{x}, \mathbf{X})^\top (\sigma_*^2 I_d + k(\mathbf{X}, \mathbf{X}))^{-1} k(\mathbf{X}, \mathbf{x}).$$

3.2 GPR constrained by the FK PDE

Applied to derivatives valuation, $g(\mathbf{x})$ can be estimated by performing a GPR on a sample of state variables $\mathbf{X} = \{\mathbf{x}_i\}_{i=1,\dots,d}$ (e.g., current time, spot price and volatility) using the corresponding prices $\{h_i\}_{i=1,\dots,d}$ as targets. This approach first requires generating the sample of prices used for training, obtained through other pricing techniques (such as trees, lattice,

Laplace transform or Monte Carlo) to estimate the target values h_i . This is the route taken by De Spiegeleer et al. (2018). Under model (1), this appears to be sub-optimal for two reasons. First, the price function is completely determined by the FK PDE combined with the terminal condition, meaning there is no need to rely on a training set of prices produced by alternative methods. Second, such a regression exercise fails to leverage (2), which is a crucial functional constraint for driving arbitrage-free prices. For these reasons, we adopt a different approach that consists of solving a GPR constrained by the FK PDE (the CGPR), without relying on a set of exogenous prices.

We consider two training sets for the CGPR. The first set contains d combinations of risk factors on the boundary, denoted as $\bar{\mathbf{X}} = \{\bar{\mathbf{x}}_i\}_{i=1,\dots,d}$, along with the corresponding noised values, $\mathbf{h} = (h_i)_{i=1,\dots,d}$. The $\bar{\mathbf{x}}_i$'s are randomly sampled from $\bar{\mathcal{X}}$ while $h_i = h(\bar{\mathbf{x}}_i) + \bar{\epsilon}$ are observations of the boundary function (e.g. payoffs), noised by $\bar{\epsilon} \sim \mathcal{N}(0, \sigma_*^2)$. The second training set consists of m combinations of risk factors before expiry, denoted as $\mathring{\mathbf{X}} = \{\mathring{\mathbf{x}}_i\}_{i=1,\dots,m}$ with $\mathring{\mathbf{x}}_i \in \mathring{\mathcal{X}}$. The noised values of the FK PDE are stored in a m -vector, $\mathbf{z} = (z_i)_{i=1,\dots,m}$ where $z_i \sim \mathcal{N}(0, \sigma_*^2)$. We seek a Gaussian process g approximating the price function f , and satisfying the following equalities on the two training sets:

$$\begin{cases} \mathcal{L}_{\mathbf{x}} g(\mathring{\mathbf{x}}_i) &= z_i \quad \forall \mathring{\mathbf{x}}_i \in \mathring{\mathbf{X}}, \quad i = 1, 2, \dots, m, \\ g(\bar{\mathbf{x}}_i) &= h_i \quad \forall \bar{\mathbf{x}}_i \in \bar{\mathbf{X}}, \quad i = 1, 2, \dots, d. \end{cases} \quad (10)$$

Let us make two important points here. Firstly, the price function at expiry coincides with the (noisy) payoff, meaning that $h_i = f(\bar{\mathbf{x}}_i) + \bar{\epsilon}$ for $\bar{\mathbf{x}}_i \in \bar{\mathcal{X}}_T$. If we ignore the first PDE constraint in Equation (10), we recognize the GPR framework and Equation (8), in which we approximate the price function on the boundary by a Gaussian process. Secondly, the training sets do not contain any market or computed prices. Indeed, $\mathbf{h}$ is a vector of noisy boundary values and $\mathbf{z}$ is a vector of Gaussian noises. It is important to note that the noise serves to provide regularization.

In the sequel, we denote by $(\mathbf{z}, \mathbf{h})$ a realization of $(\mathbf{Z}, \mathbf{H})$, two random d - and m -vectors with a Gaussian prior. We show in the next section that the joint distribution of $(g(\mathbf{x}), \mathbf{Z}, \mathbf{H})$ for $\mathbf{x} \notin \mathring{\mathbf{X}}$ follows a multivariate normal, similar to $(g(\mathbf{x}), \mathbf{H})$, as presented in Equation (9).

4 Estimation of the CGPR

We encode a-priori our belief that instances of $g(\mathbf{x})$, satisfying Equation (10), are drawn from a Gaussian process with zero mean and covariance or kernel function $k(\mathbf{x}, \mathbf{x}')$, prior to taking into account observations $\mathbf{h}$ and $\mathbf{z}$ of the boundary and FK PDE. From a Bayesian perspective, the estimator of $f(\mathbf{x})$ is the *a posteriori* expectation of $g(\mathbf{x})$ conditional upon $(\mathbf{Z}, \mathbf{H}) = (\mathbf{z}, \mathbf{h})$, which we note $\hat{f}(\mathbf{x}) = \mathbb{E}(g(\mathbf{x}) | \mathbf{z}, \mathbf{h})$. From a mathematical point of view, the function $g(\mathbf{x})$ resides in a reproducing kernel Hilbert space spanned by eigenfunctions of the covariance function (see e.g. Chapter 6 of Rasmussen and Williams, 2006).

4.1 Estimator and properties

Without loss of generality, we consider kernels defined by a finite number of eigenfunctions. In this case, there exist n basis functions $\boldsymbol{\varphi}(\mathbf{x}) = (\varphi_j(\mathbf{x}))_{j=1,\dots,n}$ such that

$$k(\mathbf{x}, \mathbf{x}') = \boldsymbol{\varphi}(\mathbf{x})^\top \boldsymbol{\varphi}(\mathbf{x}'), \quad \forall \mathbf{x}, \mathbf{x}' \in \mathcal{X}.$$

This assumption is not restrictive as any kernel can be approximated by a finite sum for a given accuracy. Because we consider a Gaussian prior with null mean and covariance matrix

$\mathbb{C}(g(\mathbf{x}), g(\mathbf{x}')) = k(\mathbf{x}, \mathbf{x}')$ for $g(\mathbf{x})$, we rewrite $g(\mathbf{x})$ as a weighted sum of functions $\varphi_j(\cdot)$:

$$g(\mathbf{x}) = \sum_{j=1}^n w_j \varphi_j(\mathbf{x}) = \mathbf{w}^\top \boldsymbol{\varphi}(\mathbf{x}) , \quad (11)$$

where the weight vector is $\mathbf{w} \sim \mathcal{N}(0, I_n)$. We define the function $\boldsymbol{\varphi}_{\mathcal{L}}(\mathbf{x}) = (\mathcal{L}_{\mathbf{x}} \varphi_j(\mathbf{x}))_{j=1, \dots, n}$ from $\mathcal{X} \rightarrow \mathbb{R}^n$. Using the linearity of the differential operator (6), $\mathcal{L}_{\mathbf{x}} g(\mathbf{x})$ is the scalar product:

$$\mathcal{L}_{\mathbf{x}} g(\mathbf{x}) = \sum_{j=1}^n w_j (\mathcal{L}_{\mathbf{x}} \varphi_j(\mathbf{x})) = \mathbf{w}^\top \boldsymbol{\varphi}_{\mathcal{L}}(\mathbf{x}) .$$

Next, we denote by $\dot{\boldsymbol{\varphi}}_{\mathcal{L}}$ the $m \times n$ matrix of $(\dot{\boldsymbol{\varphi}}_{\mathcal{L}})_{i,j} = \mathcal{L}_{\dot{\mathbf{x}}_i} \varphi_j(\dot{\mathbf{x}}_i)$ where $\dot{\mathbf{x}}_i \in \dot{\mathbf{X}}$ and by $\bar{\boldsymbol{\varphi}}$, the $d \times n$ matrix of $(\bar{\boldsymbol{\varphi}})_{i,j} = \varphi_j(\bar{\mathbf{x}}_i)$ with $\bar{\mathbf{x}}_i \in \bar{\mathbf{X}}$. Conditionally to $\mathbf{w}$, the joint distribution of $(\mathbf{Z}, \mathbf{H})$, the random vectors of right-hand terms in Equation (10) is a multivariate normal:

$$\begin{pmatrix} \mathbf{Z} \\ \mathbf{H} \end{pmatrix} \Big| \mathbf{w} \sim \mathcal{N} \left(\begin{pmatrix} \dot{\boldsymbol{\varphi}}_{\mathcal{L}} \mathbf{w} \\ \bar{\boldsymbol{\varphi}} \mathbf{w} \end{pmatrix}, \begin{pmatrix} \sigma_*^2 I_m & \mathbf{0} \\ \mathbf{0} & \sigma_*^2 I_d \end{pmatrix} \right) . \quad (12)$$

The next proposition provides the posterior distribution of weights.

Proposition 1. *The conditional distribution of regression weights in Equation (11) is:*

$$\mathbf{W} | \mathbf{z}, \mathbf{h} \sim \mathcal{N}(\boldsymbol{\mu}_{\mathbf{z}, \mathbf{h}}, \Sigma_{\mathbf{z}, \mathbf{h}})$$

where the mean $\boldsymbol{\mu}_{\mathbf{z}, \mathbf{h}}$ and covariance $\Sigma_{\mathbf{z}, \mathbf{h}}$ are equal to

$$\begin{cases} \Sigma_{\mathbf{z}, \mathbf{h}} &= \left(I_n + \sigma_*^{-2} \dot{\boldsymbol{\varphi}}_{\mathcal{L}}^\top \dot{\boldsymbol{\varphi}}_{\mathcal{L}} + \sigma_*^{-2} \bar{\boldsymbol{\varphi}}^\top \bar{\boldsymbol{\varphi}} \right)^{-1} , \\ \boldsymbol{\mu}_{\mathbf{z}, \mathbf{h}} &= \sigma_*^{-2} \Sigma_{\mathbf{z}, \mathbf{h}} \left(\dot{\boldsymbol{\varphi}}_{\mathcal{L}}^\top \mathbf{z} + \bar{\boldsymbol{\varphi}}^\top \mathbf{h} \right) . \end{cases} \quad (13)$$

Proof. We infer the posterior density of $\mathbf{W}$ from the Bayes theorem:

$$f_{\mathbf{W} | \mathbf{z}, \mathbf{h}}(\mathbf{w}) = \frac{f_{\mathbf{Z}, \mathbf{H} | \mathbf{w}}(\mathbf{z}, \mathbf{h}) f_{\mathbf{W}}(\mathbf{w})}{f_{\mathbf{Z}, \mathbf{H}}(\mathbf{z}, \mathbf{h})} ,$$

where $f_{\mathbf{W}}(\cdot)$ is the standard normal density of dimension n ,

$$f_{\mathbf{W}}(\mathbf{w}) = \frac{1}{(2\pi)^{n/2}} \exp \left(-\frac{1}{2} \mathbf{w}^\top \mathbf{w} \right) ,$$

and $f_{\mathbf{Z}, \mathbf{H} | \mathbf{w}}(\cdot)$ is the multivariate normal density (12):

$$\begin{aligned} f_{\mathbf{Z}, \mathbf{H} | \mathbf{w}}(\mathbf{z}, \mathbf{h}) &= \frac{1}{(2\pi)^{(m+d)/2} \sigma_*^{m+d}} \\ &\times \exp \left(-\frac{1}{2\sigma_*^2} (\mathbf{z} - \dot{\boldsymbol{\varphi}}_{\mathcal{L}} \mathbf{w})^\top (\mathbf{z} - \dot{\boldsymbol{\varphi}}_{\mathcal{L}} \mathbf{w}) \right) \\ &\times \exp \left(-\frac{1}{2\sigma_*^2} (\mathbf{h} - \bar{\boldsymbol{\varphi}} \mathbf{w})^\top (\mathbf{h} - \bar{\boldsymbol{\varphi}} \mathbf{w}) \right) . \end{aligned}$$

Ignoring terms not depending upon $\mathbf{w}$, we infer that $f_{\mathbf{W} | \mathbf{z}, \mathbf{h}}(\mathbf{w})$ is proportional to :

$$\begin{aligned} f_{\mathbf{W} | \mathbf{z}, \mathbf{h}}(\mathbf{w}) &\propto \exp \left(-\frac{1}{2} \mathbf{w}^\top \mathbf{w} - \frac{1}{2\sigma_*^2} (\mathbf{z} - \dot{\boldsymbol{\varphi}}_{\mathcal{L}} \mathbf{w})^\top (\mathbf{z} - \dot{\boldsymbol{\varphi}}_{\mathcal{L}} \mathbf{w}) \right) \\ &\times \exp \left(-\frac{1}{2\sigma_*^2} (\mathbf{h} - \bar{\boldsymbol{\varphi}} \mathbf{w})^\top (\mathbf{h} - \bar{\boldsymbol{\varphi}} \mathbf{w}) \right) . \end{aligned}$$

To prove that this expression complies with the claim, we show that the quadratic and linear terms in the above exponential (scaled by $-1/2$) agree with those of

$$(\mathbf{w} - \boldsymbol{\mu}_{\mathbf{z}, \mathbf{h}})^\top \Sigma_{\mathbf{z}, \mathbf{h}}^{-1} (\mathbf{w} - \boldsymbol{\mu}_{\mathbf{z}, \mathbf{h}}) = \mathbf{w}^\top \Sigma_{\mathbf{z}, \mathbf{h}}^{-1} \mathbf{w} - \mathbf{w}^\top \Sigma_{\mathbf{z}, \mathbf{h}}^{-1} \boldsymbol{\mu}_{\mathbf{z}, \mathbf{h}} - \boldsymbol{\mu}_{\mathbf{z}, \mathbf{h}}^\top \Sigma_{\mathbf{z}, \mathbf{h}}^{-1} \mathbf{w} + \boldsymbol{\mu}_{\mathbf{z}, \mathbf{h}}^\top \Sigma_{\mathbf{z}, \mathbf{h}}^{-1} \boldsymbol{\mu}_{\mathbf{z}, \mathbf{h}}.$$

This is indeed the case because the quadratic term in the above equation is equal to

$$\begin{aligned} & \mathbf{w}^\top \left(I_n + \frac{1}{\sigma_*^2} \left(k_{\mathcal{L}^2}(\mathring{\mathbf{X}}, \mathring{\mathbf{X}}) + k(\bar{\mathbf{X}}, \bar{\mathbf{X}}) \right) \right) \mathbf{w} \\ &= \mathbf{w}^\top \mathbf{w} + \frac{1}{\sigma_*^2} \left(\mathbf{w}^\top (\dot{\varphi}_{\mathcal{L}}^\top \dot{\varphi}_{\mathcal{L}}) \mathbf{w} + \mathbf{w}^\top (\bar{\varphi}^\top \bar{\varphi}) \mathbf{w} \right) \\ &= \mathbf{w}^\top \mathbf{w} + \frac{1}{\sigma_*^2} \left((\dot{\varphi}_{\mathcal{L}} \mathbf{w})^\top (\dot{\varphi}_{\mathcal{L}} \mathbf{w}) + (\bar{\varphi} \mathbf{w})^\top (\bar{\varphi} \mathbf{w}) \right), \end{aligned}$$

while the linear term is the following sum:

$$-\mathbf{w}^\top \Sigma_{\mathbf{z}, \mathbf{h}}^{-1} \boldsymbol{\mu}_{\mathbf{z}, \mathbf{h}} - \Sigma_{\mathbf{z}, \mathbf{h}}^{-1} \boldsymbol{\mu}_{\mathbf{z}, \mathbf{h}}^\top \mathbf{w} = -\frac{1}{\sigma_*^2} \left((\dot{\varphi}_{\mathcal{L}} \mathbf{w})^\top \mathbf{z} + \mathbf{z}^\top (\dot{\varphi}_{\mathcal{L}} \mathbf{w}) + (\bar{\varphi} \mathbf{w})^\top \mathbf{h} + \mathbf{h}^\top (\bar{\varphi} \mathbf{w}) \right).$$

□

From (11), $\mathcal{L}_{\mathbf{x}} g(\mathbf{x})$ depends on the weights. Therefore, we infer from Proposition (1) that, given $(\mathbf{z}, \mathbf{h})$, $\mathcal{L}_{\mathbf{x}} g(\mathbf{x})$ is a random variable distributed as:

$$\mathcal{L}_{\mathbf{x}} g(\mathbf{x}) | (\mathbf{z}, \mathbf{h}) \sim \mathcal{N} \left(\varphi_{\mathcal{L}}^\top(\mathbf{x}) \boldsymbol{\mu}_{\mathbf{z}, \mathbf{h}}, \varphi_{\mathcal{L}}^\top(\mathbf{x}) \Sigma_{\mathbf{z}, \mathbf{h}} \varphi_{\mathcal{L}}(\mathbf{x}) \right). \quad (14)$$

The conditional expectation of the differential operator given $(\mathbf{z}, \mathbf{h})$ is obtained by integrating $g(\mathbf{x})$ with respect to the conditional density of $\mathbf{W}$.

Proposition 2. *Let us define*

$$\left\{ \begin{array}{lll} k_{\mathcal{L}^2}(\mathring{\mathbf{X}}, \mathbf{x}) &= \dot{\varphi}_{\mathcal{L}} \varphi_{\mathcal{L}}(\mathbf{x}) &= (\mathcal{L}_{\hat{\mathbf{x}}_i} \mathcal{L}_{\mathbf{x}} k(\hat{\mathbf{x}}_i, \mathbf{x}))_{i=1, \dots, m} \\ k_{\mathcal{L}}(\bar{\mathbf{X}}, \mathbf{x}) &= \bar{\varphi} \varphi_{\mathcal{L}}(\mathbf{x}) &= (\mathcal{L}_{\mathbf{x}} k(\bar{\mathbf{x}}_i, \mathbf{x}))_{i=1, \dots, d} \\ k_{\mathcal{L}^2}(\mathring{\mathbf{X}}, \mathring{\mathbf{X}}) &= \dot{\varphi}_{\mathcal{L}} \dot{\varphi}_{\mathcal{L}}^\top &= (\mathcal{L}_{\hat{\mathbf{x}}_i} \mathcal{L}_{\hat{\mathbf{x}}_j} k(\hat{\mathbf{x}}_i, \hat{\mathbf{x}}_j))_{i, j=1, \dots, m} \\ k(\bar{\mathbf{X}}, \bar{\mathbf{X}}) &= \bar{\varphi} \bar{\varphi}^\top &= (k(\bar{\mathbf{x}}_i, \bar{\mathbf{x}}_j))_{i, j=1, \dots, d} \\ k_{\mathcal{L}}(\mathring{\mathbf{X}}, \bar{\mathbf{X}}) &= \dot{\varphi}_{\mathcal{L}} \bar{\varphi}^\top &= (\mathcal{L}_{\hat{\mathbf{x}}_i} k(\bar{\mathbf{x}}_j, \hat{\mathbf{x}}_i))_{i=1, \dots, m, j=1, \dots, d} \end{array} \right.$$

and the regularized Gram matrix $C(\mathring{\mathbf{X}}, \bar{\mathbf{X}})$,

$$C(\mathring{\mathbf{X}}, \bar{\mathbf{X}}) = \sigma_*^2 I_{m+d} + \begin{pmatrix} k_{\mathcal{L}^2}(\mathring{\mathbf{X}}, \mathring{\mathbf{X}}) & k_{\mathcal{L}}(\mathring{\mathbf{X}}, \bar{\mathbf{X}}) \\ k_{\mathcal{L}}(\mathring{\mathbf{X}}, \bar{\mathbf{X}})^\top & k(\bar{\mathbf{X}}, \bar{\mathbf{X}}) \end{pmatrix}. \quad (15)$$

The expectation of $\mathcal{L}_{\mathbf{x}} g(\mathbf{x})$ given $(\mathbf{z}, \mathbf{h})$ is equal to

$$\begin{aligned} \mathbb{E}(\mathcal{L}_{\mathbf{x}} g(\mathbf{x}) | \mathbf{z}, \mathbf{h}) &= \\ & \begin{pmatrix} k_{\mathcal{L}^2}(\mathring{\mathbf{X}}, \mathbf{x}) \\ k_{\mathcal{L}}(\bar{\mathbf{X}}, \mathbf{x}) \end{pmatrix}^\top C(\mathring{\mathbf{X}}, \bar{\mathbf{X}})^{-1} \begin{pmatrix} \mathbf{z} \\ \mathbf{h} \end{pmatrix}. \end{aligned} \quad (16)$$

Proof. From equations (14) and (13), using the properties of conditional expectations of the multivariate normal law, we find that

$$\mathbb{E}(\mathcal{L}_{\mathbf{x}} g(\mathbf{x}) | \mathbf{z}, \mathbf{h}) = \varphi_{\mathcal{L}}^\top(\mathbf{x}) \left(\sigma_*^2 I_n + \dot{\varphi}_{\mathcal{L}}^\top \dot{\varphi}_{\mathcal{L}} + \bar{\varphi}^\top \bar{\varphi} \right)^{-1} \left(\dot{\varphi}_{\mathcal{L}}^\top \mathbf{z} + \bar{\varphi}^\top \mathbf{h} \right).$$

The matrix inverted in this equation can be rewritten as follows:

$$\sigma_*^2 I_n + \dot{\varphi}_{\mathcal{L}}^\top \dot{\varphi}_{\mathcal{L}} + \bar{\varphi}^\top \bar{\varphi} = \sigma_*^2 \left(I_n + \begin{pmatrix} \sigma_*^{-1} \dot{\varphi}_{\mathcal{L}} \\ \sigma_*^{-1} \bar{\varphi} \end{pmatrix}^\top \begin{pmatrix} \sigma_*^{-1} \dot{\varphi}_{\mathcal{L}} \\ \sigma_*^{-1} \bar{\varphi} \end{pmatrix} \right). \quad (17)$$

We next use the Woodbury formula which states that for a $n \times (m+d)$ matrix U and a $(m+d) \times n$ matrix V ,

$$(I_n + UV)^{-1} = I_n - U(I_{m+d} + VU)^{-1}V. \quad (18)$$

This equality allows us to rewrite (17) as:

$$\begin{aligned} & \left(I_n + \begin{pmatrix} \sigma_*^{-1} \dot{\varphi}_{\mathcal{L}} \\ \sigma_*^{-1} \bar{\varphi} \end{pmatrix}^\top \begin{pmatrix} \sigma_*^{-1} \dot{\varphi}_{\mathcal{L}} \\ \sigma_*^{-1} \bar{\varphi} \end{pmatrix} \right)^{-1} = \\ & I_n - \begin{pmatrix} \sigma_*^{-1} \dot{\varphi}_{\mathcal{L}} \\ \sigma_*^{-1} \bar{\varphi} \end{pmatrix}^\top \left(I_{m+d} + \sigma_*^{-2} \begin{pmatrix} \dot{\varphi}_{\mathcal{L}} \dot{\varphi}_{\mathcal{L}}^\top & \dot{\varphi}_{\mathcal{L}} \bar{\varphi}^\top \\ \bar{\varphi} \dot{\varphi}_{\mathcal{L}}^\top & \bar{\varphi} \bar{\varphi}^\top \end{pmatrix} \right)^{-1} \begin{pmatrix} \sigma_*^{-1} \dot{\varphi}_{\mathcal{L}} \\ \sigma_*^{-1} \bar{\varphi} \end{pmatrix}. \end{aligned} \quad (19)$$

Injecting this expression into the expectation, one gets

$$\begin{aligned} \mathbb{E}(\mathcal{L}_x g(\mathbf{x}) | \mathbf{z}, \mathbf{h}) &= \sigma_*^{-2} \varphi_{\mathcal{L}}^\top(\mathbf{x}) \begin{pmatrix} \dot{\varphi}_{\mathcal{L}} \\ \bar{\varphi} \end{pmatrix}^\top \begin{pmatrix} \mathbf{z} \\ \mathbf{h} \end{pmatrix} - \varphi_{\mathcal{L}}^\top(\mathbf{x}) \begin{pmatrix} \dot{\varphi}_{\mathcal{L}} \\ \bar{\varphi} \end{pmatrix}^\top \times \\ & \left(I_{m+d} + \sigma_*^{-2} \begin{pmatrix} \dot{\varphi}_{\mathcal{L}} \dot{\varphi}_{\mathcal{L}}^\top & \dot{\varphi}_{\mathcal{L}} \bar{\varphi}^\top \\ \bar{\varphi} \dot{\varphi}_{\mathcal{L}}^\top & \bar{\varphi} \bar{\varphi}^\top \end{pmatrix} \right)^{-1} \begin{pmatrix} \dot{\varphi}_{\mathcal{L}} \dot{\varphi}_{\mathcal{L}}^\top & \dot{\varphi}_{\mathcal{L}} \bar{\varphi}^\top \\ \bar{\varphi} \dot{\varphi}_{\mathcal{L}}^\top & \bar{\varphi} \bar{\varphi}^\top \end{pmatrix} \begin{pmatrix} \mathbf{z} \\ \mathbf{h} \end{pmatrix}. \end{aligned}$$

We obtain (16) with the following relation:

$$\begin{pmatrix} \dot{\varphi}_{\mathcal{L}} \dot{\varphi}_{\mathcal{L}}^\top & \dot{\varphi}_{\mathcal{L}} \bar{\varphi}^\top \\ \bar{\varphi} \dot{\varphi}_{\mathcal{L}}^\top & \bar{\varphi} \bar{\varphi}^\top \end{pmatrix} = \sigma_*^2 \left[I_{m+d} + \sigma_*^{-2} \begin{pmatrix} \dot{\varphi}_{\mathcal{L}} \dot{\varphi}_{\mathcal{L}}^\top & \dot{\varphi}_{\mathcal{L}} \bar{\varphi}^\top \\ \bar{\varphi} \dot{\varphi}_{\mathcal{L}}^\top & \bar{\varphi} \bar{\varphi}^\top \end{pmatrix} - I_{m+d} \right].$$

The claim follows from the definition of basis functions. $\square$

Notice that Equations (15) and (16) do not directly depend on basis functions $\varphi_j(\cdot)$, but only on the kernel. Therefore, we do not need to know the analytical expressions of the basis functions to compute $C(\dot{\mathbf{X}}, \bar{\mathbf{X}})$ and $\mathbb{E}(\mathcal{L}_x g(\mathbf{x}) | \mathbf{z}, \mathbf{h})$.

Proposition 3. *The conditional variance of $\mathcal{L}_x g(\mathbf{x})$ given $\mathbf{z}$ and $\mathbf{h}$ does not depend on $(\mathbf{z}, \mathbf{h})$, and is equal to*

$$\mathbb{V}(\mathcal{L}_x g(\mathbf{x}) | \mathbf{z}, \mathbf{h}) = k_{\mathcal{L}^2}(\dot{\mathbf{X}}, \dot{\mathbf{X}}) - \begin{pmatrix} k_{\mathcal{L}^2}(\dot{\mathbf{X}}, \mathbf{x}) \\ k_{\mathcal{L}}(\bar{\mathbf{X}}, \mathbf{x}) \end{pmatrix}^\top C(\dot{\mathbf{X}}, \bar{\mathbf{X}})^{-1} \begin{pmatrix} k_{\mathcal{L}^2}(\dot{\mathbf{X}}, \mathbf{x}) \\ k_{\mathcal{L}}(\bar{\mathbf{X}}, \mathbf{x}) \end{pmatrix}. \quad (20)$$

Proof. This expression results from the definition of $\Sigma_{\mathbf{x}, \mathbf{z}}$ and from the application of the Woodbury formula, Equation (19). $\square$

The next proposition presents the closed-form expression for the estimator of the pricing function using the CGPR method.

Proposition 4. *Let us define $k_{\mathcal{L}}(\dot{\mathbf{X}}, \mathbf{x}) = (\mathcal{L}_{\dot{\mathbf{x}}_j} k(\mathbf{x}, \dot{\mathbf{x}}_j))_{j=1, \dots, m}$ and the $(m+d)$ -vector:*

$$\beta(\mathbf{x}) = \begin{pmatrix} k_{\mathcal{L}}(\dot{\mathbf{X}}, \mathbf{x}) \\ k(\bar{\mathbf{X}}, \mathbf{x}) \end{pmatrix}. \quad (21)$$

The CGPR estimator $\hat{f}(\mathbf{x}) = \mathbb{E}(g(\mathbf{x}) | \mathbf{z}, \mathbf{h})$ of option prices is equal to

$$\hat{f}(\mathbf{x}) = \beta(\mathbf{x})^\top C(\dot{\mathbf{X}}, \bar{\mathbf{X}})^{-1} \begin{pmatrix} \mathbf{z} \\ \mathbf{h} \end{pmatrix}. \quad (22)$$

Proof. From Eq. (16), the conditional expectation of $\mathcal{L}_{\mathbf{x}}g(\mathbf{x})$ is a linear combination of $\mathcal{L}_{\mathbf{x}}k(\mathbf{x}, \bar{\mathbf{x}}_j)$ and $\mathcal{L}_{\mathbf{x}}\mathcal{L}_{\mathbf{x}_j}k(\mathbf{x}, \mathbf{x}_j)$:

$$\begin{aligned}\mathbb{E}(\mathcal{L}_{\mathbf{x}}g(\mathbf{x}) | \mathbf{z}, \mathbf{h}) &= \begin{pmatrix} k_{\mathcal{L}^2}(\dot{\mathbf{X}}, \mathbf{x}) & k_{\mathcal{L}}(\bar{\mathbf{X}}, \mathbf{x}) \end{pmatrix} \hat{\boldsymbol{\alpha}} \\ &= \sum_{j=1}^m \hat{\alpha}_j \mathcal{L}_{\dot{\mathbf{x}}_j} \mathcal{L}_{\mathbf{x}} k(\dot{\mathbf{x}}_j, \mathbf{x}) + \sum_{j=1}^d \hat{\alpha}_{m+j} \mathcal{L}_{\mathbf{x}} k(\bar{\mathbf{x}}_j, \mathbf{x}),\end{aligned}$$

where the m -vector $\hat{\boldsymbol{\alpha}} = (\hat{\alpha}_j)_{j=1, \dots, m}$ is defined by

$$\hat{\boldsymbol{\alpha}} = \left(\sigma_*^2 I_{m+d} + \begin{pmatrix} k_{\mathcal{L}^2}(\dot{\mathbf{X}}, \dot{\mathbf{X}}) & k_{\mathcal{L}}(\dot{\mathbf{X}}, \bar{\mathbf{X}}) \\ k_{\mathcal{L}}(\dot{\mathbf{X}}, \bar{\mathbf{X}})^\top & k(\bar{\mathbf{X}}, \bar{\mathbf{X}}) \end{pmatrix} \right)^{-1} \begin{pmatrix} \mathbf{z} \\ \mathbf{h} \end{pmatrix}. \quad (23)$$

Using the linearity of the operator, the estimator $\hat{f}(\mathbf{x})$ is given by

$$\begin{aligned}\hat{f}(\mathbf{x}) &= \mathcal{L}_{\mathbf{x}}^{-1} \mathbb{E}(\mathcal{L}_{\mathbf{x}}g(\mathbf{x}) | \mathbf{z}, \mathbf{h}) \\ &= \sum_{j=1}^m \hat{\alpha}_j \mathcal{L}_{\dot{\mathbf{x}}_j} k(\mathbf{x}, \dot{\mathbf{x}}_j) + \sum_{j=1}^d \hat{\alpha}_{m+j} k(\bar{\mathbf{x}}_j, \mathbf{x}).\end{aligned} \quad (24)$$

Eq. (22) is the matrix reformulation of this last equation. $\square$

As previously mentioned, adding Gaussian noise to h_i 's and z_i 's introduces a diagonal matrix $\sigma_*^2 I_{m+d}$ in the definition (22) of $\hat{f}(\mathbf{x})$. This diagonal matrix is a Tikhonov regularization factor that stabilizes a possibly ill-conditioned inverse matrix. The next proposition provides the a posteriori conditional variance of the estimator.

Proposition 5. *Let $\boldsymbol{\beta}(\mathbf{x})$ be the $(m+d)$ vector defined in Eq. (21). Then the conditional variance of the CGPR estimator (22) does not depend on $(\mathbf{z}, \mathbf{h})$ and is equal to*

$$\mathbb{V}(g(\mathbf{x}) | \mathbf{z}, \mathbf{h}) = k(\mathbf{x}, \mathbf{x}) - \boldsymbol{\beta}(\mathbf{x})^\top C(\dot{\mathbf{X}}, \bar{\mathbf{X}})^{-1} \boldsymbol{\beta}(\mathbf{x}). \quad (25)$$

Proof. By definition of $g(\mathbf{x})$, the $\mathbf{z}, \mathbf{h}$ -conditional variance of $\mathcal{L}_{\mathbf{x}}g(\mathbf{x})$ is equal to:

$$\begin{aligned}\mathbb{V}(\mathcal{L}_{\mathbf{x}}g(\mathbf{x}) | \mathbf{z}, \mathbf{h}) &= \mathbb{V}\left(\sum_{k=1}^n W_k \mathcal{L}_{\mathbf{x}} \varphi_k(\mathbf{x}) \middle| \mathbf{z}, \mathbf{h}\right) \\ &= \boldsymbol{\varphi}_{\mathcal{L}}(\mathbf{x})^\top \Sigma_{\mathbf{z}, \mathbf{h}} \boldsymbol{\varphi}_{\mathcal{L}}(\mathbf{x}),\end{aligned}$$

where $\Sigma_{\mathbf{z}, \mathbf{h}}$ is the a posteriori covariance of weights defined in Equation (13). By comparison with Eq. (20), we deduce an alternative formulation³ of $\Sigma_{\mathbf{z}, \mathbf{h}}$:

$$\Sigma_{\mathbf{z}, \mathbf{h}} = I_n - \begin{pmatrix} \dot{\boldsymbol{\varphi}}_{\mathcal{L}} \\ \bar{\boldsymbol{\varphi}} \end{pmatrix}^\top C(\dot{\mathbf{X}}, \bar{\mathbf{X}})^{-1} \begin{pmatrix} \dot{\boldsymbol{\varphi}}_{\mathcal{L}} \\ \bar{\boldsymbol{\varphi}} \end{pmatrix}. \quad (26)$$

On the other hand, the variance of $g(\mathbf{x})$ is equal to

$$\begin{aligned}\mathbb{V}(g(\mathbf{x}) | \mathbf{z}, \mathbf{h}) &= \mathbb{V}\left(\sum_{i=1}^n W_i \varphi_j(\mathbf{x}) \middle| \mathbf{z}, \mathbf{h}\right) \\ &= \boldsymbol{\varphi}(\mathbf{x})^\top \Sigma_{\mathbf{z}, \mathbf{h}} \boldsymbol{\varphi}(\mathbf{x}).\end{aligned}$$

Combining this last expression with Eq. (26) allows us to infer that:

$$\mathbb{V}(g(\mathbf{x}) | \mathbf{z}, \mathbf{h}) = \boldsymbol{\varphi}(\mathbf{x})^\top \boldsymbol{\varphi}(\mathbf{x}) - \begin{pmatrix} \dot{\boldsymbol{\varphi}}_{\mathcal{L}} \boldsymbol{\varphi}(\mathbf{x}) \\ \bar{\boldsymbol{\varphi}} \boldsymbol{\varphi}(\mathbf{x}) \end{pmatrix}^\top C(\dot{\mathbf{X}}, \bar{\mathbf{X}})^{-1} \begin{pmatrix} \dot{\boldsymbol{\varphi}}_{\mathcal{L}} \boldsymbol{\varphi}(\mathbf{x}) \\ \bar{\boldsymbol{\varphi}} \boldsymbol{\varphi}(\mathbf{x}) \end{pmatrix}.$$

As $\boldsymbol{\varphi}(\mathbf{x})^\top \boldsymbol{\varphi}(\mathbf{x}) = k(\mathbf{x}, \mathbf{x})$, $(\dot{\boldsymbol{\varphi}}_{\mathcal{L}} \boldsymbol{\varphi}(\mathbf{x}))_i = \mathcal{L}_{\dot{\mathbf{x}}_i} k(\mathbf{x}, \dot{\mathbf{x}}_i)$ and $(\bar{\boldsymbol{\varphi}} \boldsymbol{\varphi}(\mathbf{x}))_i = k(\mathbf{x}, \bar{\mathbf{x}}_i)$, we obtain Equation (25). $\square$

³The same expression can be obtained by applying the Woodbury formula.

4.2 Algorithm and the Gaussian kernel

The frame 1 summarizes the algorithm for computing derivative prices by CGPR. In practice, we do not add noises ϵ and $\bar{\epsilon}$ to $(z_i)_{i=1,\dots,m}$ and $(h_i)_{i=1,\dots,d}$. This has no impact on calculations and slightly improves the accuracy of the algorithm.

Algorithm 1 CGPR approximation of the price $\widehat{f}(\mathbf{x})$ of a derivative with payoff $h(\mathbf{x})$.

1. Sample m points $\mathring{\mathbf{x}}_i = (t_i, y_i) \in \mathcal{X}$ and set $(z_i = 0)_{i=1,\dots,m}$.
2. Sample d points $\bar{\mathbf{x}}_i = (t, y_i) \in \bar{\mathcal{X}}$ and set $(h_i = h(\bar{\mathbf{x}}_i))_{i=1\dots d}$.
3. Compute the regularized Gram matrix

$$C(\mathring{\mathbf{X}}, \bar{\mathbf{X}}) = \sigma_*^2 I_{m+d} + \begin{pmatrix} k_{\mathcal{L}^2}(\mathring{\mathbf{X}}, \mathring{\mathbf{X}}) & k_{\mathcal{L}}(\mathring{\mathbf{X}}, \bar{\mathbf{X}}) \\ k_{\mathcal{L}}(\mathring{\mathbf{X}}, \bar{\mathbf{X}})^\top & k(\bar{\mathbf{X}}, \bar{\mathbf{X}}) \end{pmatrix}$$

where

$$\begin{cases} \left(k_{\mathcal{L}^2}(\mathring{\mathbf{X}}, \mathring{\mathbf{X}}) \right)_{i,j} &= \mathcal{L}_{\mathring{\mathbf{x}}_i} \mathcal{L}_{\mathring{\mathbf{x}}_j} k(\mathring{\mathbf{x}}_i, \mathring{\mathbf{x}}_j) & i, j = 1, \dots, m, \\ \left(k_{\mathcal{L}}(\mathring{\mathbf{X}}, \bar{\mathbf{X}}) \right)_{i,j} &= \mathcal{L}_{\mathring{\mathbf{x}}_i} k(\mathring{\mathbf{x}}_i, \bar{\mathbf{x}}_j) & i = 1, \dots, m, j = 1, \dots, d, \\ \left(k(\bar{\mathbf{X}}, \bar{\mathbf{X}}) \right)_{i,j} &= k(\bar{\mathbf{x}}_i, \bar{\mathbf{x}}_j) & i, j = 1, \dots, d. \end{cases} \quad (27)$$

4. Invert the Gram matrix $C(\mathring{\mathbf{X}}, \bar{\mathbf{X}})$ and compute for an out of sample $\mathbf{x} \in \mathcal{X}$:

$$\boldsymbol{\beta}(\mathbf{x}) = \begin{pmatrix} (\mathcal{L}_{\mathring{\mathbf{x}}_i} k(\mathbf{x}, \mathring{\mathbf{x}}_i))_{i=1,\dots,m} \\ (k(\bar{\mathbf{x}}_i, \mathbf{x}))_{i=1,\dots,d} \end{pmatrix}.$$

5. The pricing function is $\widehat{f}(\mathbf{x}) = \boldsymbol{\beta}(\mathbf{x})^\top C(\mathring{\mathbf{X}}, \bar{\mathbf{X}})^{-1} \begin{pmatrix} \mathbf{z} \\ \mathbf{h} \end{pmatrix}$.
-

The computation of estimates requires applying the differential operator $\mathcal{L}_{\mathbf{x}}$ twice to the kernel function. Many alternatives exist such as the rational quadratic (RQ), Matern 3/2 (M32), Matern 5/2 (M52) kernels, see e.g. Beckers (2021). We opt for a Gaussian kernel (also called radial basis function) in our numerical illustrations. This kernel is easy to differentiate and allows us to infer fully analytical expressions for the vector $\boldsymbol{\beta}(\mathbf{x})$ and sub-blocks of the Gram matrix. In a general setting, the Gaussian kernel has the form:

$$k(\mathbf{x}, \tilde{\mathbf{x}}) = \exp \left(-\frac{(t - \tilde{t})^2}{2\eta_t^2} - \sum_{k=1}^{p-1} \frac{(y_k - \tilde{y}_k)^2}{2\eta_k^2} \right) \quad \forall \mathbf{x}, \tilde{\mathbf{x}} \in \mathcal{X}. \quad (28)$$

The following results depends on the following p -vector:

$$\mathbf{d}_{\tilde{\mathbf{y}}}(\mathbf{x}, \tilde{\mathbf{x}}) = \begin{pmatrix} (y_1 - \tilde{y}_1) / \eta_1^2 \\ \vdots \\ (y_{p-1} - \tilde{y}_{p-1}) / \eta_{p-1}^2 \end{pmatrix},$$

and on the a $(p-1) \times (p-1)$ matrix:

$$\mathbf{H}_{\tilde{\mathbf{y}}}(\mathbf{x}, \tilde{\mathbf{x}}) = \begin{pmatrix} \frac{(y_1 - \tilde{y}_1)^2 - \eta_1^2}{\eta_1^4} & \frac{(y_1 - \tilde{y}_1)(y_2 - \tilde{y}_2)}{\eta_1^2 \eta_2^2} & \cdots & \frac{(y_1 - \tilde{y}_1)(y_{p-1} - \tilde{y}_{p-1})}{\eta_1^2 \eta_{p-1}^2} \\ \frac{(y_1 - \tilde{y}_1)(y_2 - \tilde{y}_2)}{\eta_1^2 \eta_2^2} & \ddots & \ddots & \frac{(y_2 - \tilde{y}_2)(y_{p-1} - \tilde{y}_{p-1})}{\eta_2^2 \eta_{p-1}^2} \\ \vdots & \ddots & \ddots & \vdots \\ \frac{(y_1 - \tilde{y}_1)(y_{p-1} - \tilde{y}_{p-1})}{\eta_1^2 \eta_{p-1}^2} & \frac{(y_2 - \tilde{y}_2)(y_{p-1} - \tilde{y}_{p-1})}{\eta_2^2 \eta_{p-1}^2} & \cdots & \frac{(y_{p-1} - \tilde{y}_{p-1})^2 - \eta_{p-1}^2}{\eta_{p-1}^4} \end{pmatrix}.$$

With a Gaussian kernel, the vector $\beta(\mathbf{x})$ admits a closed-form expression as stated in the next proposition. We assume in later developments that the risk-free rate $r(\mathbf{x}) = r \in \mathbb{R}$ is constant, without loss of generality.

Proposition 6. *For all $\mathbf{x} = (t, \mathbf{y}) \in \mathcal{X}$, the approximated derivative price is $\hat{f}(\mathbf{x}) = \beta(\mathbf{x})^\top \times C(\dot{\bar{\mathbf{X}}}, \bar{\mathbf{X}})^{-1}(\mathbf{z}, \mathbf{h})^\top$. With a Gaussian kernel (28), the $(m+d)$ -vector $\beta(\mathbf{x}) = (\beta_1(\mathbf{x}), \beta_2(\mathbf{x}))^\top$ has the following components:*

$$\begin{cases} \beta_1(\mathbf{x}) = \left(((t - \bar{t}_i) / \eta_t^2 - r + \boldsymbol{\mu}_{\mathbf{y}}(\hat{\mathbf{x}}_i)^\top \mathbf{d}_{\tilde{\mathbf{y}}_i}(\mathbf{x}, \hat{\mathbf{x}}_i) \right. \\ \quad \left. + \frac{1}{2} \text{tr}(\Sigma_{\mathbf{y}}(\hat{\mathbf{x}}_i) \Sigma_{\mathbf{y}}(\hat{\mathbf{x}}_i)^\top \mathbf{H}_{\tilde{\mathbf{y}}_i}(\mathbf{x}, \hat{\mathbf{x}}_i)) \right) k(\mathbf{x}, \hat{\mathbf{x}}_i)_{i=1, \dots, m}, \\ \beta_2(\mathbf{x}) = \left(e^{-(\bar{t}_i - t)^2 / (2\eta_t^2) - \sum_{k=1}^{p-1} (\bar{y}_{i,k} - t)^2 / (2\eta_k^2)} \right)_{i=1, \dots, d}. \end{cases}$$

Proof. The last d elements of $\beta(\mathbf{x})$ are equal to $(k(\bar{\mathbf{x}}_i, \mathbf{x}))_{i=1, \dots, d}$, directly computable with Equation (28). By direct differentiation, we check that the gradient and Hessian of the kernel are given by

$$\begin{cases} \nabla_{\tilde{\mathbf{y}}} k(\mathbf{x}, \tilde{\mathbf{x}}) &= \mathbf{d}_{\tilde{\mathbf{y}}}(\mathbf{x}, \tilde{\mathbf{x}}) k(\mathbf{x}, \tilde{\mathbf{x}}), \\ \mathcal{H}_{\tilde{\mathbf{y}}} k(\mathbf{x}, \tilde{\mathbf{x}}) &= \mathbf{H}_{\tilde{\mathbf{y}}}(\mathbf{x}, \tilde{\mathbf{x}}) k(\mathbf{x}, \tilde{\mathbf{x}}). \end{cases}$$

We infer from this last equation, that $\mathcal{L}_{\tilde{\mathbf{x}}} k(\mathbf{x}, \tilde{\mathbf{x}})$ depends on $\mathbf{d}_{\tilde{\mathbf{y}}}(\mathbf{x}, \tilde{\mathbf{x}})$ and $\mathbf{H}_{\tilde{\mathbf{y}}}(\mathbf{x}, \tilde{\mathbf{x}})$ as follows:

$$\begin{aligned} \mathcal{L}_{\tilde{\mathbf{x}}} k(\mathbf{x}, \tilde{\mathbf{x}}) &= \left((t - \tilde{t}) / \eta_t^2 - r + \boldsymbol{\mu}_{\mathbf{y}}(\tilde{\mathbf{x}})^\top \mathbf{d}_{\tilde{\mathbf{y}}}(\mathbf{x}, \tilde{\mathbf{x}}) \right. \\ &\quad \left. + \frac{1}{2} \text{tr}(\Sigma_{\mathbf{y}}(\tilde{\mathbf{x}}) \Sigma_{\mathbf{y}}(\tilde{\mathbf{x}})^\top \mathbf{H}_{\tilde{\mathbf{y}}}(\mathbf{x}, \tilde{\mathbf{x}})) \right) k(\mathbf{x}, \tilde{\mathbf{x}}). \end{aligned} \quad (29)$$

Equations (29) allows us to retrieve the expression of $\beta_1(\mathbf{x})$. □

The construction of the Gram matrix, $C(\dot{\bar{\mathbf{X}}}, \bar{\mathbf{X}})$, requires to compute the sub-blocks: $k(\bar{\mathbf{X}}, \bar{\mathbf{X}})$, $k_{\mathcal{L}}(\dot{\bar{\mathbf{X}}}, \bar{\mathbf{X}})$ and $k_{\mathcal{L}^2}(\dot{\bar{\mathbf{X}}}, \dot{\bar{\mathbf{X}}})$. When the kernel is Gaussian, they admit general closed-form expressions.

Proposition 7. *The sub-blocks $k(\bar{\mathbf{X}}, \bar{\mathbf{X}})$ and $k_{\mathcal{L}}(\dot{\bar{\mathbf{X}}}, \bar{\mathbf{X}})$ are computed with the following analytical formulas*

$$\begin{cases} k(\bar{\mathbf{X}}, \bar{\mathbf{X}}) = \left(e^{-(\bar{t}_i - \bar{t}_j)^2 / (2\eta_t^2) - \sum_{k=1}^{p-1} (\bar{y}_{i,k} - \bar{y}_{j,k})^2 / (2\eta_k^2)} \right)_{i,j=1, \dots, d}, \\ k_{\mathcal{L}}(\dot{\bar{\mathbf{X}}}, \bar{\mathbf{X}}) = \left(((\bar{t}_j - \bar{t}_i) / \eta_t^2 - r + \boldsymbol{\mu}_{\mathbf{y}}(\hat{\mathbf{x}}_i)^\top \mathbf{d}_{\tilde{\mathbf{y}}_i}(\bar{\mathbf{x}}_j, \hat{\mathbf{x}}_i) \right. \\ \quad \left. + \frac{1}{2} \text{tr}(\Sigma_{\mathbf{y}}(\hat{\mathbf{x}}_i) \Sigma_{\mathbf{y}}(\hat{\mathbf{x}}_i)^\top \mathbf{H}_{\tilde{\mathbf{y}}_i}(\bar{\mathbf{x}}_j, \hat{\mathbf{x}}_i)) \right) k(\bar{\mathbf{x}}_j, \hat{\mathbf{x}}_i)_{i=1, \dots, m, j=1, \dots, d}. \end{cases}$$

The expression of $k(\bar{\mathbf{X}}, \bar{\mathbf{X}})$ is a direct consequence of the kernel choice. The second equation comes from the symmetry of the kernel and Equation (29). Let us denote by $\mathbf{e}_k$'s the canonical vectors of $\mathbb{R}^{p-1}$. We define the following matrix:

$$\mathbf{H}_{\mathbf{y}}(\mathbf{x}, \tilde{\mathbf{x}}) = \mathbf{H}_{\tilde{\mathbf{y}}}(\mathbf{x}, \tilde{\mathbf{x}}) + \begin{pmatrix} 2/\eta_1^2 & 0 & 0 \\ 0 & \cdots & 0 \\ 0 & 0 & 2/\eta_{p-1}^2 \end{pmatrix}.$$

Proposition 8. The sub-block $k_{\mathcal{L}^2}(\overset{\circ}{\mathbf{X}}, \overset{\circ}{\mathbf{X}})$ of $C(\overset{\circ}{\mathbf{X}}, \bar{\mathbf{X}})$ is the matrix of $\mathcal{L}_{\hat{\mathbf{x}}_i} \mathcal{L}_{\hat{\mathbf{x}}_j} k(\hat{\mathbf{x}}_i, \hat{\mathbf{x}}_j)$ for $i, j = 1, \dots, m$ where

$$\begin{aligned} \mathcal{L}_{\mathbf{x}} \mathcal{L}_{\tilde{\mathbf{x}}} k(\mathbf{x}, \tilde{\mathbf{x}}) &= \partial_t \mathcal{L}_{\tilde{\mathbf{x}}} k(\mathbf{x}, \tilde{\mathbf{x}}) + \boldsymbol{\mu}_{\mathbf{y}}(\mathbf{x})^\top \nabla_{\mathbf{y}} \mathcal{L}_{\tilde{\mathbf{x}}} k(\mathbf{x}, \tilde{\mathbf{x}}) \\ &\quad - r \mathcal{L}_{\tilde{\mathbf{x}}} k(\mathbf{x}, \tilde{\mathbf{x}}) + \frac{1}{2} \text{tr} \left(\Sigma_{\mathbf{y}}(\mathbf{x}) \Sigma_{\mathbf{y}}(\mathbf{x})^\top \mathcal{H}_{\mathbf{y}} \mathcal{L}_{\tilde{\mathbf{x}}} k(\mathbf{x}, \tilde{\mathbf{x}}) \right). \end{aligned} \quad (30)$$

Using Equation (29), the partial derivative in (30) is given by

$$\partial_t \mathcal{L}_{\tilde{\mathbf{x}}} k(\mathbf{x}, \tilde{\mathbf{x}}) = \frac{1}{\eta_t^2} k(\mathbf{x}, \tilde{\mathbf{x}}) - (\mathcal{L}_{\tilde{\mathbf{x}}} k(\mathbf{x}, \tilde{\mathbf{x}})) \frac{(t - \tilde{t})}{\eta_t^2}, \quad (31)$$

while the gradient $\nabla_{\mathbf{y}} \mathcal{L}_{\tilde{\mathbf{x}}} k(\mathbf{x}, \tilde{\mathbf{x}}) = (\partial_{y_k} \mathcal{L}_{\tilde{\mathbf{x}}} k(\mathbf{x}, \tilde{\mathbf{x}}))_{k=1, \dots, p-1}$ is a $(p-1)$ vector filled by :

$$\begin{aligned} \partial_{y_k} \mathcal{L}_{\tilde{\mathbf{x}}} k(\mathbf{x}, \tilde{\mathbf{x}}) &= \left[\mathbf{e}_k^\top \Sigma_{\mathbf{y}}(\tilde{\mathbf{x}}) \Sigma_{\mathbf{y}}(\tilde{\mathbf{x}})^\top \mathbf{d}_{\tilde{\mathbf{y}}}(\mathbf{x}, \tilde{\mathbf{x}}) \right. \\ &\quad \left. + \mathbf{e}_k^\top \boldsymbol{\mu}_{\mathbf{y}}(\tilde{\mathbf{x}}) \right] \frac{k(\mathbf{x}, \tilde{\mathbf{x}})}{\eta_k^2} - \mathcal{L}_{\tilde{\mathbf{x}}} k(\mathbf{x}, \tilde{\mathbf{x}}) \mathbf{e}_k^\top \mathbf{d}_{\tilde{\mathbf{y}}}(\mathbf{x}, \tilde{\mathbf{x}}), \end{aligned} \quad (32)$$

Finally The Hessian $\mathcal{H}_{\mathbf{y}} \mathcal{L}_{\tilde{\mathbf{x}}} k(\mathbf{x}, \tilde{\mathbf{x}}) = (\partial_{y_l} \partial_{y_k} \mathcal{L}_{\tilde{\mathbf{x}}} k(\mathbf{x}, \tilde{\mathbf{x}}))_{k,l=1, \dots, p-1}$ in Equation (30) is a $(p-1) \times (p-1)$ matrix with elements:

$$\begin{aligned} \partial_{y_l} \partial_{y_k} \mathcal{L}_{\tilde{\mathbf{x}}} k(\mathbf{x}, \tilde{\mathbf{x}}) &= \left(\mathbf{e}_k^\top \Sigma_{\mathbf{y}}(\tilde{\mathbf{x}}) \Sigma_{\mathbf{y}}(\tilde{\mathbf{x}})^\top \mathbf{e}_l \right) \frac{k(\mathbf{x}, \tilde{\mathbf{x}})}{\eta_k^2 \eta_l^2} \\ &\quad - (\partial_{y_k} \mathcal{L}_{\tilde{\mathbf{x}}} k(\mathbf{x}, \tilde{\mathbf{x}})) \mathbf{e}_l^\top \mathbf{d}_{\tilde{\mathbf{y}}}(\mathbf{x}, \tilde{\mathbf{x}}) - (\partial_{y_l} \mathcal{L}_{\tilde{\mathbf{x}}} k(\mathbf{x}, \tilde{\mathbf{x}})) \mathbf{e}_k^\top \mathbf{d}_{\tilde{\mathbf{y}}}(\mathbf{x}, \tilde{\mathbf{x}}) \\ &\quad - \mathcal{L}_{\tilde{\mathbf{x}}} k(\mathbf{x}, \tilde{\mathbf{x}}) \left(\mathbf{e}_l^\top (\mathbf{H}_{\mathbf{y}}(\mathbf{x}, \tilde{\mathbf{x}})) \mathbf{e}_k \right). \end{aligned} \quad (33)$$

In Section 6 and 5, we will present the specific forms of $\mathcal{L}_{\tilde{\mathbf{x}}} k(\mathbf{x}, \tilde{\mathbf{x}})$ and $\mathcal{L}_{\mathbf{x}} \mathcal{L}_{\tilde{\mathbf{x}}} k(\mathbf{x}, \tilde{\mathbf{x}})$ in the Heston and Geometric diffusive models.

We conclude this section by discussing the fit of kernel bandwidths. Equations (22) and (25) correspond to the conditional expectation and variance of $g(\mathbf{x}) | \mathbf{z}, \mathbf{h}$ with a multivariate normal distribution:

$$\begin{bmatrix} g(\mathbf{x}) \\ \mathbf{Z} \\ \mathbf{H} \end{bmatrix} \sim \mathcal{N} \left(\begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix}, \begin{bmatrix} k(\mathbf{x}, \mathbf{x}) & \boldsymbol{\beta}(\mathbf{x})^\top \\ \boldsymbol{\beta}(\mathbf{x}) & C(\overset{\circ}{\mathbf{X}}, \bar{\mathbf{X}}) \end{bmatrix} \right).$$

Based on this relation, we can tune the bandwidths of Gaussian kernels by maximizing the marginal log-likelihood of $(\mathbf{Z}, \mathbf{H})$. The kernel parameters, $\boldsymbol{\eta} = \{\eta_t, \eta_1, \dots, \eta_{p-1}\}$ are optimized by maximization of the following expression:

$$\log l(\mathbf{z}, \mathbf{h}) = -\frac{1}{2} \log |C(\overset{\circ}{\mathbf{X}}, \bar{\mathbf{X}})| - \frac{1}{2} \begin{pmatrix} \mathbf{z} \\ \mathbf{h} \end{pmatrix}^\top C(\overset{\circ}{\mathbf{X}}, \bar{\mathbf{X}})^{-1} \begin{pmatrix} \mathbf{z} \\ \mathbf{h} \end{pmatrix} + \frac{(m+d)}{2} \log 2\pi. \quad (34)$$

In practice, this maximization is subject to numerical instabilities because the log-likelihood is a non-convex function of hyper-parameters. We propose an efficient and robust alternative in which the bandwidths minimize the errors of approximation on the inner and boundary domains. For this purpose, we generate two new samples, exclusively used for tuning $\boldsymbol{\eta}$. The first sample set contains m points in the inner domain and is denoted by $\overset{\circ}{\mathbf{X}}' = \{\hat{\mathbf{x}}'_i\}_{i=1, \dots, m}$ with $\hat{\mathbf{x}}'_i \in \overset{\circ}{\mathcal{X}}$.

We draw a second sample set of d points on the boundary. This set is $\bar{\mathbf{X}}' = \{\tilde{\mathbf{x}}'_i\}_{i=1, \dots, d}$ with

$\bar{\mathbf{x}}'_i \in \bar{\mathcal{X}}$. We define the mean square errors on the inner and terminal domains as follows:

$$\begin{aligned} \mathring{Err}(\boldsymbol{\eta}) &= \frac{1}{m} \sum_{\hat{\mathbf{x}}'_i \in \hat{\mathcal{X}}'} \left(\mathcal{L}_{\hat{\mathbf{x}}'_i} \hat{f}(\hat{\mathbf{x}}'_i) \right)^2, \\ \bar{Err}(\boldsymbol{\eta}) &= \frac{1}{d} \sum_{\bar{\mathbf{x}}'_i \in \bar{\mathcal{X}}'} \left(\hat{f}(\bar{\mathbf{x}}'_i) - h(\bar{\mathbf{x}}'_i) \right)^2. \end{aligned}$$

The optimal hyper-parameters $\boldsymbol{\eta}^*$, are obtained by minimization of the total error, e.g. with a Nelder-Mead algorithm:

$$\boldsymbol{\eta}^* = \arg \min_{\boldsymbol{\eta}} \left(\mathring{Err}(\boldsymbol{\eta}) + \bar{Err}(\boldsymbol{\eta}) \right). \quad (35)$$

We will see that this method yields good results for the two optimization problems analyzed in the following section.

4.3 The Greeks

Since the pricing function $\hat{f}(\mathbf{x})$ is linear in $\boldsymbol{\beta}(\mathbf{x})$, we can efficiently approximate the Greeks with analytical expressions. The Greeks are the partial derivatives of the pricing function, representing the sensitivity of the option's value to changes in one or more underlying risk factors. They are essential tools for risk management, but their computation via finite differences is often prone to numerical instabilities. The sensitivity of the option price to time, called the “theta”, is in our case given by

$$\partial_t \hat{f}(\mathbf{x}) = (\partial_t \boldsymbol{\beta}(\mathbf{x}))^\top C(\dot{\mathbf{X}}, \bar{\mathbf{X}})^{-1} \begin{pmatrix} \mathbf{z} \\ \mathbf{h} \end{pmatrix}. \quad (36)$$

The sensitivity of $\hat{f}(\mathbf{x})$ to the i^{th} state variable, the “delta”, is measured by the following derivative:

$$\partial_{y_k} \hat{f}(\mathbf{x}) = (\partial_{y_k} \boldsymbol{\beta}(\mathbf{x}))^\top C(\dot{\mathbf{X}}, \bar{\mathbf{X}})^{-1} \begin{pmatrix} \mathbf{z} \\ \mathbf{h} \end{pmatrix}. \quad (37)$$

Whereas the second order sensitivity, called “gamma” or “cross-gamma” is the second order derivative of the approximated prices with respect to state processes:

$$\partial_{y_l} \partial_{y_k} \hat{f}(\mathbf{x}) = (\partial_{y_l} \partial_{y_k} \boldsymbol{\beta}(\mathbf{x}))^\top C(\dot{\mathbf{X}}, \bar{\mathbf{X}})^{-1} \begin{pmatrix} \mathbf{z} \\ \mathbf{h} \end{pmatrix}. \quad (38)$$

When the kernel is Gaussian, partial derivatives of $\boldsymbol{\beta}(\mathbf{x})$ admits closed-form expressions. The partial derivative of $\boldsymbol{\beta}(\mathbf{x})$ w.r.t. time is in this case equal to:

$$\partial_t \boldsymbol{\beta}(\mathbf{x}) = \begin{pmatrix} (\partial_t \mathcal{L}_{\hat{\mathbf{x}}_i} k(\mathbf{x}, \hat{\mathbf{x}}_i))_{i=1, \dots, m} \\ (\partial_t k(\bar{\mathbf{x}}_j, \mathbf{x}))_{j=1, \dots, d} \end{pmatrix}. \quad (39)$$

Using Equation (29), the terms of $\partial_t \boldsymbol{\beta}(\mathbf{x})$ are:

$$\begin{cases} \partial_t \mathcal{L}_{\hat{\mathbf{x}}_i} k(\mathbf{x}, \hat{\mathbf{x}}_i) &= \frac{k(\mathbf{x}, \hat{\mathbf{x}}_i)}{\eta_t^2} - \frac{(t - \bar{t}_i)}{\eta_t^2} \mathcal{L}_{\hat{\mathbf{x}}_i} k(\mathbf{x}, \hat{\mathbf{x}}_i) \\ \partial_t k(\bar{\mathbf{x}}_j, \mathbf{x}) &= \frac{(\bar{t}_j - t)}{\eta_t^2} k(\bar{\mathbf{x}}_j, \mathbf{x}). \end{cases}$$

For the “delta's”, the derivative of $\boldsymbol{\beta}(\mathbf{x})$ w.r.t. y_k in Equation (37), is a $(m + d)$ vector given by

$$\partial_{y_k} \boldsymbol{\beta}(\mathbf{x}) = \begin{pmatrix} (\partial_{y_k} \mathcal{L}_{\hat{\mathbf{x}}_i} k(\mathbf{x}, \hat{\mathbf{x}}_i))_{i=1, \dots, m} \\ (\partial_{y_k} k(\bar{\mathbf{x}}_j, \mathbf{x}))_{j=1, \dots, d} \end{pmatrix}, \quad (40)$$

with elements equal to

$$\begin{cases} \partial_{y_k} \mathcal{L}_{\hat{\mathbf{x}}_i} k(\mathbf{x}, \hat{\mathbf{x}}_i) &= [\mathbf{e}_k^\top \Sigma_{\mathbf{y}}(\hat{\mathbf{x}}_i) \Sigma_{\mathbf{y}}(\hat{\mathbf{x}}_i)^\top \mathbf{d}_{\hat{\mathbf{y}}_i}(\mathbf{x}, \hat{\mathbf{x}}_i) \\ &+ \boldsymbol{\mu}_{\mathbf{y}}(\hat{\mathbf{x}}_i)^\top \mathbf{e}_k] \frac{k(\mathbf{x}, \hat{\mathbf{x}}_i)}{\eta_k^2} \\ &- \mathcal{L}_{\hat{\mathbf{x}}_i} k(\mathbf{x}, \hat{\mathbf{x}}_i) \mathbf{e}_k^\top \mathbf{d}_{\hat{\mathbf{y}}_i}(\mathbf{x}, \hat{\mathbf{x}}_i), \\ \partial_{y_k} k(\bar{\mathbf{x}}_j, \mathbf{x}) &= \frac{(\bar{y}_{j,k} - y_k)}{\eta_k^2} k(\bar{\mathbf{x}}_j, \mathbf{x}). \end{cases}$$

The “cross-Gamma’s” depend on the $(m + d)$ vectors:

$$\partial_{y_l} \partial_{y_k} \boldsymbol{\beta}(\mathbf{x}) = \begin{pmatrix} (\partial_{y_l} \partial_{y_k} \mathcal{L}_{\hat{\mathbf{x}}_i} k(\mathbf{x}, \hat{\mathbf{x}}_i))_{i=1, \dots, m} \\ (\partial_{y_l} \partial_{y_k} k(\bar{\mathbf{x}}_j, \mathbf{x}))_{j=1, \dots, d} \end{pmatrix}. \quad (41)$$

Using Equation (29) and (32), the elements of this vector have the following closed-form expressions:

$$\begin{cases} \partial_{y_l} \partial_{y_k} \mathcal{L}_{\hat{\mathbf{x}}_i} k(\mathbf{x}, \hat{\mathbf{x}}_i) &= (\mathbf{e}_k^\top \Sigma_{\mathbf{y}}(\hat{\mathbf{x}}_i) \Sigma_{\mathbf{y}}(\hat{\mathbf{x}}_i)^\top \mathbf{e}_l) \frac{k(\mathbf{x}, \hat{\mathbf{x}}_i)}{\eta_k^2 \eta_l^2} \\ &- (\partial_{y_k} \mathcal{L}_{\hat{\mathbf{x}}_i} k(\mathbf{x}, \hat{\mathbf{x}}_i)) \mathbf{e}_l^\top \mathbf{d}_{\hat{\mathbf{y}}_i}(\mathbf{x}, \hat{\mathbf{x}}_i) \\ &- (\partial_{y_l} \mathcal{L}_{\hat{\mathbf{x}}_i} k(\mathbf{x}, \hat{\mathbf{x}}_i)) \mathbf{e}_k^\top \mathbf{d}_{\hat{\mathbf{y}}_i}(\mathbf{x}, \hat{\mathbf{x}}_i) \\ &- \mathcal{L}_{\hat{\mathbf{x}}_i} k(\mathbf{x}, \hat{\mathbf{x}}_i) (\mathbf{e}_l^\top \mathbf{H}_{\mathbf{y}}(\mathbf{x}, \hat{\mathbf{x}}_i) \mathbf{e}_k) \\ \partial_{y_l} \partial_{y_k} k(\bar{\mathbf{x}}_j, \mathbf{x}) &= \left(\frac{(\bar{y}_{j,k} - y_k)(\bar{y}_{j,l} - y_l)}{\eta_k^2 \eta_l^2} - \mathbf{1}_{\{k=l\}} \frac{1}{\eta_k^2} \right) k(\bar{\mathbf{x}}_j, \mathbf{x}). \end{cases}$$

The calculation of Greeks is challenging in absence of analytical formulas due to numerical instabilities when approximating partial derivatives. Having (approximated) closed-form expressions is a real added value for managing risks.

5 European options in a geometric diffusive market

In this first example, we benchmark our approach with the Black & Scholes formula. We consider a financial market with a single risky assets, noted $(S_t)_{t \geq 0}$ ruled by a Brownian motion $(B_t)_{t \geq 0}$. The state variable is $\mathbf{y}_t = S_t$. This is ruled by the stochastic differential equation (1) where $\boldsymbol{\mu}_{\mathbf{y}}(\cdot)$ and $\Sigma_{\mathbf{y}}(\cdot)$ are scalar functions:

$$\boldsymbol{\mu}_{\mathbf{y}}(t, \mathbf{y}_t) = r S_t \quad \text{and} \quad \Sigma_{\mathbf{y}}(t, \mathbf{y}_t) = \sigma S_t,$$

where r is the risk-free rate and $\sigma \in \mathbb{R}^+$. $\mathbf{x} = (t, s)$ is a bivariate vector ($p = 2$) combining time t and a realization s of $\mathbf{y}_t = S_t$. The domain of $\mathbf{x}$ is $\mathcal{X} = \mathbb{R}^{2+}$. We evaluate a call option with a strike price, K and a payoff $\mathcal{P}(\bar{\mathbf{x}}) = (\bar{s} - K)_+$. We solve the FK equation on the inner domain $\bar{\mathcal{X}} = [0, T) \times (0, s_{max})$ where s_{max} is large. The boundary set is the union of three subsets, $\bar{\mathcal{X}} = \bar{\mathcal{X}}_0 \cup \bar{\mathcal{X}}_{max} \cup \bar{\mathcal{X}}_T$, where

$$\begin{aligned} \bar{\mathcal{X}}_0 &= \{(t, 0) \mid t \in [0, T)\}, \\ \bar{\mathcal{X}}_{max} &= \{(t, s_{max}) \mid t \in [0, T)\}, \\ \bar{\mathcal{X}}_T &= \{(T, s) \mid s \in \mathbb{R}^+\}. \end{aligned} \quad (42)$$

The boundary function for $\bar{\mathbf{x}} = (\bar{t}, \bar{s})$, is in this particular case defined as follows

$$h(\bar{\mathbf{x}}) = \begin{cases} 0 & \bar{\mathbf{x}} \in \bar{\mathcal{X}}_0, \\ e^{-r(T-\bar{t})}(\bar{s} - K) & \bar{\mathbf{x}} \in \bar{\mathcal{X}}_{max}, \\ (\bar{s} - K)_+ & \bar{\mathbf{x}} \in \bar{\mathcal{X}}_T. \end{cases} \quad (43)$$

The value of the option in state $\mathbf{x}$, $f(\mathbf{x})$, solves the FK equation (5) where the gradient and the Hessian of f with respect to $\mathbf{y}$ by $\nabla_{\mathbf{y}}f$ and $\mathcal{H}_{\mathbf{y}}(f)$ are this time given by:

$$\nabla_{\mathbf{y}}f = \partial_s f \quad \text{and} \quad \mathcal{H}_{\mathbf{y}}(f) = \partial_{ss}f.$$

We consider the Gaussian kernel, $k(\mathbf{x}, \tilde{\mathbf{x}})$, defined by a positive vector of bandwidths, $\boldsymbol{\eta} = (\eta_t, \eta_1) \in \mathbb{R}^2$:

$$k(\mathbf{x}, \tilde{\mathbf{x}}) = \exp \left(-\frac{(t - \tilde{t})^2}{2\eta_t^2} - \frac{(s - \tilde{s})^2}{2\eta_1^2} \right).$$

We compare CGPR prices, obtained with Proposition 6, to those calculated using the Black-Scholes-Merton (BSM) formula (see Black and Scholes, 1976). This special case is interesting because exact prices are available. Table 1 presents the stock parameters, the features of the priced call options, σ_* and the kernel bandwidths. We optimize bandwidths by minimization of errors, see Equation (35). The points in training sets $\hat{\mathbf{X}}$ and $\bar{\mathbf{X}}$ are randomly drawn according to a uniform distribution over the intervals reported in Table 1.

Parameter	Value	Parameter	Value
r	0.05	K	100
σ_1	0.25	T	0.5 , 1 , 1.5 , 2 , 2.5 , 3
$s_i^{(1)}$	$\in [0, 200]$, uniform law	t_i	$\in [0, T]$, uniform law
σ_*	0.01		

Table 1: Parameters of the BSM model, features of the call options and boundaries for sampling

$$\hat{\mathbf{x}}_i = (t_i, s_i^{(1)}) \quad \text{and} \quad \bar{\mathbf{x}}_i = (T, s_i^{(1)}).$$

We consider a validation set of $6 \times 50 = 300$ call options with 6 maturities T , from 6 months to 3 years and 50 initial values $S_0^{(1)}$ from 50 to 150, in increments of 2. We use training sets $\hat{\mathbf{X}}$ of size m varying between 200 and 1000 points. We respectively sample $m/5$, $m/4$ and $m/3$ points in $\bar{\mathcal{X}}_0$, $\bar{\mathcal{X}}_{max}$ and $\bar{\mathcal{X}}_T$. The size d of $\bar{\mathbf{X}}$ is then equal to $\frac{3}{5}m$, $\frac{3}{4}m$ and m . Table 2 reports the mean absolute errors (MAEs) between CGPR and BSM prices. Regardless the combination of d and m , the MAEs are negligible but the lowest one is achieved with $d = m$. Table 2 also shows the mean relative errors (MREs) for in-the-Money options. The MRE ranges from 11 to 35 basis points, which is acceptable for industrial applications.

m	Spread CGPR - BSM prices					
	Absolute			Relative (ITM)		
	$d = 3/5m$	$d = 3/4m$	$d = m$	$d = 3/5m$	$d = 3/4m$	$d = m$
200	0.0546	0.0643	0.0655	0.0011	0.0017	0.0018
400	0.0380	0.0388	0.0400	0.0016	0.0018	0.0018
600	0.0791	0.0619	0.0282	0.0035	0.0027	0.0012

Table 2: Absolute: MAEs between CGPR and BSM prices. Relative (ITM) : MREs for in-the-Money (ITM) options ($S_0^{(1)} > K$).

Table 3 reports the computation times along with the MAEs and MREs of CGRP with $m = d$. As a numerical method, CGPR is naturally slower than the analytical valuation using the Black-Scholes formula. Nevertheless, the time required to value 300 options does not exceed 13 seconds. Table 3 also presents the average standard deviation of CGPR prices computed using Equation (25). These values are comparable regardless of the sample sizes.

m and $d = m$	Spread CGPR - BSM prices		Std	Times (seconds)		
	Absolute	Relative (ITM)		Opt.	CGPR	BSM
200	0.0655	0.0018	0.0018	5.70	1.32	0.03
400	0.0400	0.0018	0.0023	6.65	1.59	0.03
600	0.0282	0.0012	0.0018	10.58	1.80	0.04

Table 3: Absolute: MAEs between CGPR and BSM prices. Relative : MREs for in-the-Money (ITM) options ($S_0^{(1)} > K$). Std: standard deviation of CGPR estimates. Times CGPR and BSM: computation times with CGPR and BSM formula.

T	Spread CGPR - BSM prices		St. dev.	Hyper-parameters	
	Absolute	Relative (ITM)		η_t	η_1
0.5	0.0217	0.0017	0.0014	0.75	11.96
1.0	0.0248	0.0013	0.0014	0.76	13.48
1.5	0.0246	0.0010	0.0016	0.77	13.96
2.0	0.0272	0.0010	0.0018	0.76	15.19
2.5	0.0299	0.0010	0.0022	0.76	15.03
3.0	0.0409	0.0013	0.0025	0.76	15.94

Table 4: MAEs and MREs for call options and in-the-Money (ITM) options ($S_0^{(1)} > K$) per maturity. CGPR prices are computed with $m = 600$ and $d = 600$. St.dev.: standard deviations of CGPR estimates.

Table 4 presents the MAEs and MREs per maturity for options in the validation sets, and the estimated hyper-parameters. For each maturity, the spread is based on 50 option prices with $S_0^{(1)}$ ranging from 50 to 150, in increments of 2. In relative terms, the largest errors are observed for 0.5 year options (0.17%). The left plot of Figure 1 compares CGPR and BSM prices of 1.5 year options. The right plot displays a heatmap of absolute spreads between BSM and CGPR call prices with respect to T and $S_0^{(1)}$. These graphs confirm the overall accuracy of the CGPR method and reveal that the highest errors occurs for deep in the money long-term options (near the boundary on which the FK equation is solved).

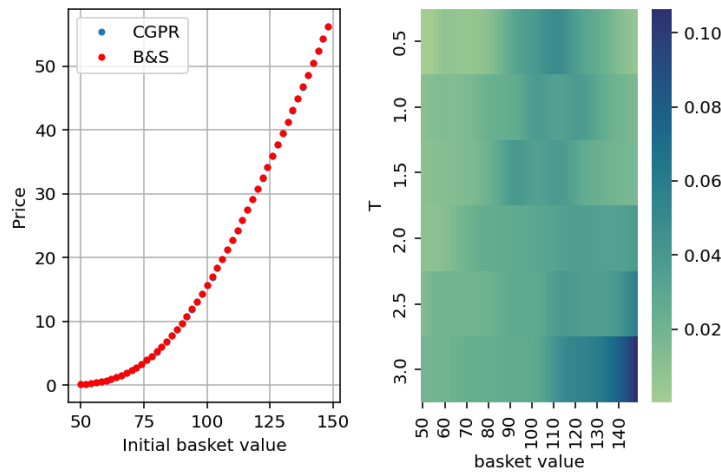


Figure 1: Left: BSM (B&S) and CGPR (Kernel) 1.5-year call prices computed with $m = 400$ and $d = 240$ for $S_0^{(1)}$ ranging from 50 to 150, by step of 2. Right: heatmap of absolute spreads between BSM and CGPR call prices with respect to T and $S_0^{(1)}$.

6 Options valuation in the Heston model

We test the CGPR Algorithm 1 for pricing European options in the Heston model. We consider a market with two assets. The cash account earns a constant risk-free rate $r \in \mathbb{R}^+$. The stock price process $S = (S_t)_{t \geq 0}$, is ruled by a geometric Brownian diffusion with a stochastic variance process, $V = (V_t)_{t \geq 0}$. The price and variance processes are defined on a complete filtered probability space $(\Omega, \mathcal{F}, \mathbb{F}, \mathbb{Q})$ generated by two independent Brownian motions $\mathbf{B}_t = (B_t^{(1)}, B_t^{(2)})_{t \geq 0}^\top$. The vector of state variables $\mathbf{y}_t = (S_t, V_t)$ is driven by the SDE (1) with

$$\boldsymbol{\mu}_{\mathbf{y}}(t, \mathbf{y}_t) = \begin{pmatrix} r S_t \\ \kappa(\gamma - V_t) \end{pmatrix} \quad \text{and} \quad \Sigma_{\mathbf{y}}(t, \mathbf{y}_t) = \begin{pmatrix} S_t \sqrt{V_t} \Sigma_S^\top \\ \sigma \sqrt{V_t} \Sigma_V^\top \end{pmatrix}.$$

Parameters $\kappa, \gamma \in \mathbb{R}^+$ are respectively the speed of mean reversion and the long term average of the stock variance while $\sigma \in \mathbb{R}^+$ is the volatility of the stock variance. In this expression, Σ_S and Σ_V are 1×2 vectors such that the matrix

$$\Sigma = \begin{pmatrix} \Sigma_S^\top \\ \Sigma_V^\top \end{pmatrix} = \begin{pmatrix} \sqrt{1 - \rho_{SV}^2} & \rho_{SV} \\ 0 & 1 \end{pmatrix} \quad \text{satisfies} \quad \Sigma \Sigma^\top = \begin{pmatrix} 1 & \rho_{SV} \\ \rho_{SV} & 1 \end{pmatrix},$$

i.e., ρ_{SV} is the instantaneous coefficient of correlation between S_t and V_t and Σ is the upper Choleski decomposition of the correlation matrix.

As previously, we note $\mathbf{x} = (t, \mathbf{y})$, where t is the time and $\mathbf{y} = (v, s)$ are the state variables at the corresponding time. The domain of $\mathbf{x}$ is $\mathcal{X} = [0, T] \times \mathbb{R}^{+,2}$. In the numerical illustration, we consider a European call option on S with expiry T and strike price K . The boundary set and boundary function are defined by Equations (42) and (43). The Heston model is a popular framework for option valuation, with many extensions. Prices are computed by a Fast Fourier Transform (FFT), Monte-Carlo simulations or finite difference approximation. For instance, Düring and Miles (2017) propose an alternating direction implicate finite difference scheme for pricing options. Hainaut (2023 b) evaluate spread and exchange options in a rough-Heston model. Souto Arias (2025) use the COS method for valuing an extension of the Heston model with a Q-Hawkes jump process. Cuomo et al. (2020) assume that the option price in the Heston model, is a weighted sum of radial basis function on a grid. This is very close to the idea underlying Physics Inspired Neural Network (PINN's). In this article, we opt for the Monte-Carlo (MC) method because of its accuracy.

The price of the derivative is equal to the risk-neutral expectation of the discounted payoff. It solves the FK equation (5) where the gradient and the Hessian of f with respect to $\mathbf{y}$ by $\nabla_{\mathbf{y}} f$ and $\mathcal{H}_{\mathbf{y}}(f)$ are here given by:

$$\nabla_{\mathbf{y}} f = \begin{pmatrix} \partial_s f \\ \partial_v f \end{pmatrix} \quad \text{and} \quad \mathcal{H}_{\mathbf{y}}(f) = \begin{pmatrix} \partial_{ss} f & \partial_{sv} f \\ \partial_{sv} f & \partial_{vv} f \end{pmatrix}.$$

We consider a Gaussian kernel with bandwidth parameters $\boldsymbol{\eta} = (\eta_t, \eta_s, \eta_v) \in \mathbb{R}^{+,3}$:

$$k(\mathbf{x}, \tilde{\mathbf{x}}) = \exp \left(-\frac{(t - \tilde{t})^2}{2\eta_t^2} - \frac{(s - \tilde{s})^2}{2\eta_s^2} - \frac{(v - \tilde{v})^2}{2\eta_v^2} \right) \quad \forall \mathbf{x}, \tilde{\mathbf{x}} \in \mathcal{X}. \quad (44)$$

The estimator of the price function, $\hat{f}(\mathbf{x}) = \mathbb{E}(g(\mathbf{x} | \mathbf{z}, \mathbf{h}))$ is provided in Proposition 4.

Let us now analyze the performance of our CGPR estimator. To this end, we will benchmark our results against MC option prices. The CGPR algorithm is implemented in Python

and tested on a laptop with an Intel Core Ultra 7, 1.40 Ghz with 32Gb of RAM. Table 5 presents the parameters of the Heston model. We evaluate standard call options with maturities ranging from 6 months to 3 years. The samples in $\mathring{\mathcal{X}}$ and $\bar{\mathcal{X}}$ are obtained by randomly drawing from independent uniform distributions over intervals reported in Table 5.

Parameter	Value	Parameter	Value
r	0.05	K	100
σ	0.30	T	0.5 , 1 , 1.5 , 2 , 2.5 , 3
γ	0.15^2	s_i	$\in [0, 300]$, uniform law
κ	0.95	t_i	$\in [0, T]$, uniform law
ρ_{SV}	-0.30	v_i	$\in [0, 8\gamma]$, uniform law

Table 5: Parameters of the Heston model, of the call options, and of the sampling model used for building the training sets $\mathring{\mathcal{X}}$ and $\bar{\mathcal{X}}$. We draw m samples $\mathring{\mathbf{x}}_i = (t_i, s_i, v_i)$ and d samples $\bar{\mathbf{x}}_i = (T, s_i, v_i)$. The samples s_i, v_i (and t_i for $\mathring{\mathbf{x}}_i$) are drawn independently of each other.

Parameter	Value
T	0.5 , 1 , 1.5 , 2 , 2.5 , 3
$\bar{S}_0$	20 equispaced values from 50 to 150
$\bar{V}_0$	20 equispaced values from $\frac{\gamma}{5}$ to 5γ .

Table 6: Features of 2400 options in the validation set.

We compute CGPR and MC prices at time $t = 0$ for $6 \times 20 \times 20 = 2400$ call options with features specified in Tables 5 and 6. The features of the 2400 options are selected to uniformly cover the hypercube $[0, 3] \times [50, 150] \times [\gamma/5, 5\gamma]$, representing expiry, stock and variance values, respectively. We consider various sample sizes m for $\mathring{\mathcal{X}}$, ranging from 400 to 800 points. Specifically, we respectively sample $m/5$, $m/4$, $m/3$ points in $\mathring{\mathcal{X}}_0$, $\mathring{\mathcal{X}}_{max}$ and $\mathring{\mathcal{X}}_T$. The boundary sample $\bar{\mathcal{X}}$ has a size d equal to $\frac{3}{5}m$, $\frac{3m}{4}$ and m . Table 7 presents MAEs and MREs between CGPR and MC prices, obtained using 5000 simulations. The MC stock and volatility paths are generated using Euler discretization with 100 steps per year. In absolute terms, the MAEs between CGPR and MC prices are consistently small across all combinations of d and m . The lowest MAEs for deep in-the-money options are achieved when $m = d$ and $m = 600$ or 800. For these configurations, the MRE drops to 0.74%, which is acceptable given that the MC benchmark also incurs estimation error due to the Euler discretization.

m	Spread CGPR - MC prices					
	Absolute			Relative (ITM)		
	$d = 3/5m$	$d = 3/4m$	$d = m$	$d = 3/5m$	$d = 3/4m$	$d = m$
400	0.3296	0.3927	0.4555	0.0132	0.0122	0.0117
600	0.3660	0.5516	0.3944	0.0125	0.0228	0.0074
800	0.3781	0.3906	0.3548	0.0140	0.0172	0.0074

Table 7: Absolute: MAEs between CGPR and MC prices. Relative (ITM) : MREs for in-the-Money (ITM) options ($S_0^{(1)} > K$).

Table 8 reports for $d = m$, computation times of CGPR and MC, errors, and standard deviations of CGPR prices, computed with formula (25). We observe that the a posteriori standard deviation falls to 0.46% when $m = 800$. In term of computational time, the CGPR is more efficient than MC simulations. The left plot of Figure 2 compares MC and CGPR values of

0.5-year call options and $\sqrt{V_0} = 24.76\%$. This visually confirms the accuracy of the CGPR in comparison to a standard MC method. The right plot of the same figure, shows the heatmap of MAEs between MC and CGPR call prices with respect to S_0 and $\sqrt{V_0}$. Figure 3 presents the heatmaps⁴ for combinations of T with S_0 and of T with $\sqrt{V_0}$. Regardless of the combination of T , S_0 and $\sqrt{V_0}$, the spreads between MC and CGPR prices remain at acceptable levels. The highest spreads are obtained for combinations with either small or large values for S_0 and $\sqrt{V_0}$. We attribute this to the truncation of the domain $\mathcal{X}$ by a closed rectangular parallelepiped. The accuracy of the CGPR method decreases for combinations of state variables that are close to the boundaries delimiting this box.

$m = d$	Spread CGPR - MC prices		St.dev.	Times (seconds)		
	Absolute	Relative (ITM)		Opt.	CGPR	MC
400	0.4555	0.0117	0.0056	14.92	4.02	174.82
600	0.3944	0.0074	0.0046	33.29	7.46	176.15
800	0.3548	0.0074	0.0046	46.74	8.76	173.21

Table 8: Absolute: MAEs between CGPR and MC call prices. Relative (ITM) : MREs for in-the-Money options ($S_0 > K$). St.dev.: standard deviations of CGPR estimates. Opt. : optimization times for computing η_t , η_s and η_v . CGPR and MC: pricing times of 2400 options with CGPR and MC methods.

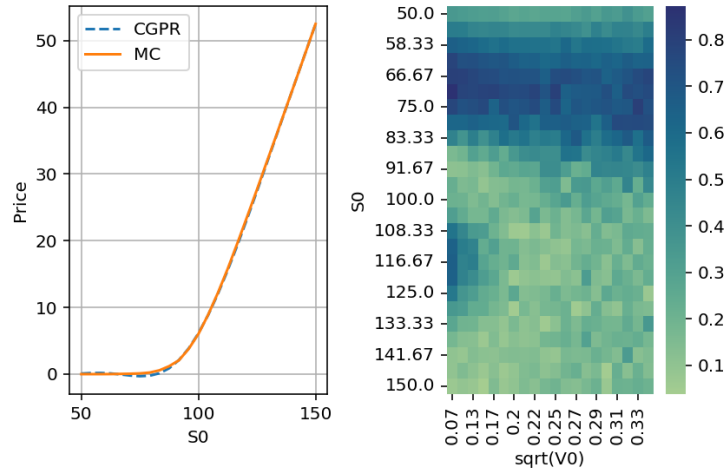


Figure 2: Left: MC and CGPR (kernel) 0.5-year call prices computed with $m = 800$, $d = 800$ and $\sqrt{V_0} = 24.76\%$. Right: heatmap of average absolute spreads between MC and CGPR call prices with respect to S_0 and $\sqrt{V_0}$ (all maturities).

We conclude this section by discussing the impact of σ_* on accuracy and stability of the CGPR prices. In practice, σ_* serves as a Tikhonov regularization parameter that prevents numerical instability when inverting the matrix $C(\dot{\mathbf{X}}, \bar{\mathbf{X}})$, defined in Equation (15). If σ_* is too small, the CGPR algorithm may fail to invert $C(\dot{\mathbf{X}}, \bar{\mathbf{X}})$. Conversely, increasing σ_* introduces a bias in the CGPR prices. To better understand the impact of σ_* , we evaluate 3-years call options with S_0 and V_0 from Table 6 under various noise levels. Table 9 reports MAEs and MREs (ITM) between MC and CGPR prices for σ_* values ranging from 0.001 to 0.50. The table also presents

⁴The heatmaps present an average of errors with respect to 2 state variables. These errors are averaged over the other simulated state variables.

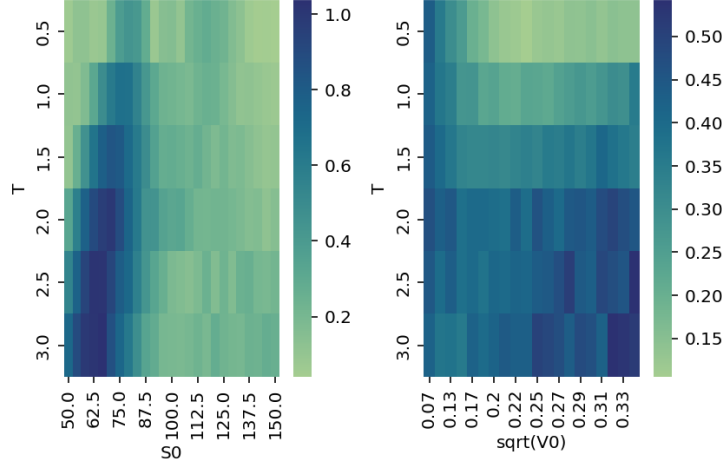


Figure 3: Right: heatmap of average absolute spreads between MC and CGPR ($m=800$, $d=800$) call prices with respect to T and $\sqrt{V_0}$. Left: same heatmap but with respect to T and S_0 .

the optimal hyper-parameters. We observe that the highest absolute error is obtained with the highest value of σ_* . When $\sigma_* = 0.001$, the Tikhonov regularization factor is not sufficient to prevent some numerical instability that generates higher errors mainly for out-of-the-money options. The minimal absolute error is obtained with $\sigma_* = 0.01$ and increases for larger σ_* . Figure 4 visualizes the gap between MC and CGPR prices as σ_* increases, highlighting that a high σ_* introduces a valuation bias.

σ_*	Spread CGPR - MC prices			Hyper-parameters		
	Absolute	Relative (ITM)	St. dev.	η_t	η_s	η_v
0.001	0.8354	0.0144	0.0037	1.49	45.41	0.12
0.01	0.4825	0.0060	0.0080	1.51	51.56	0.13
0.1	0.5213	0.0069	0.0416	1.55	64.19	0.14
0.25	1.0763	0.0212	0.0839	1.58	80.89	0.15
0.5	1.5301	0.0361	0.1176	2.84	73.95	0.99

Table 9: Absolute values and relative spreads for call options and in-the-Money (ITM) options ($S_0 > K$) per maturity. $m=400$ and σ_* ranges from 0.001 to 0.5. St.dev.: standard deviations of CGPR estimates.

7 Conclusions

This article presents a novel method for pricing European options using constrained Gaussian process regression (CGPR), leveraging the Feynman-Kac PDE that governs arbitrage-free derivative prices. By reframing the task of solving the Feynman-Kac partial differential equation as a constrained regression problem, this approach provides a robust and efficient alternative to traditional numerical methods. The method’s reliance on sampling state variables from both the inner domain of the PDE and the terminal boundary, along with a Bayesian framework, ensures high accuracy and stability in option pricing.

One of the key advantages of this method is its avoidance of numerical differentiation, a common challenge in Physics-Inspired Neural Networks (PINNs). Instead, CGPR harnesses the inherent

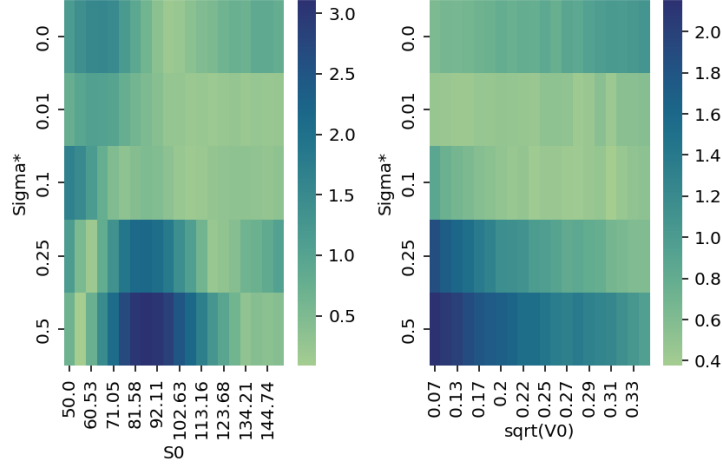


Figure 4: Right: heatmap of average absolute spreads between MC and CGPR ($m=800$, $d = 800$) 3 years call prices with respect to σ_* and $\sqrt{V_0}$. Left: same heatmap but with respect to σ_* and S_0 .

properties of Gaussian processes to deliver a closed-form approximation of option prices. This approach not only enhances numerical robustness but also supports the inclusion of a Tikhonov regularization factor, further improving the numerical stability of the method.

The numerical illustrations, including the evaluation of European call options in the Black and Scholes and Heston models demonstrate the practical applicability and effectiveness of the proposed framework. The results indicate that CGPR achieves accuracy comparable to established techniques such as the Monte Carlo (MC) methods. This positions CGPR as a robust alternative for pricing derivatives, showcasing its potential in both theoretical and practical contexts.

Nevertheless, the CGPR presents a few drawbacks, such as the time required to optimize the kernel bandwidths. Similar to PINNs, CGPR necessitates large training datasets when addressing higher-dimensional problems, complicating the inversion of the Gram matrix $C(\hat{\mathbf{X}}, \bar{\mathbf{X}})$. Nevertheless, there exists numerical solutions for managing large datasets, such as discussed in Chapter 8 of Rasmussen and Williams (2006).

However, the CGPR framework offers a novel and efficient approach to derivative pricing. By addressing the limitations of existing methods and providing a robust solution to the Feynman-Kac PDE, this framework has the potential to inspire further research and development in this area. In particular, the CGPR can be extended for pricing American options with an iterative penalized method, as studied in Hainaut (2024 b). Another extension of CGPR is presented in Hainaut and Dupret (2025) who combine in this article, the CGPR with policy improvements for solving stochastic control problems.

This work paves the way for further research. For instance, we can consider other kernels or other dynamics for state variables. It is probably possible to apply our framework to jump or switching processes. In this case, the differential operator is non-linear but we can adapt the framework of Hainaut (2024 b) for developing an iterative method.

Acknowledgement

Donatien Hainaut thanks the FNRS (Fonds de la recherche scientifique) for the financial support through the EOS project Asterisk (research grant 40007517). The work of Frédéric Vrins is supported by the F.S.R.-FNRS (Fonds de la Recherche Scientifique) through the CDR grant J.0037.18 and by the ARC project NEMO (grant 18-23/089) of the Belgian Federal Science Policy Office.

Disclosure statement

No potential conflict of interest was reported by authors.

References

- [1] Al-Aradi A., Correia A., Jardim G., de Freitas Naiff D., Saporito Y. 2022. Extensions of the deep Galerkin method. *Applied Mathematics and Computation*.
- [2] Barber D., Wang Y. 2014. Gaussian processes for Bayesian estimation in ordinary differential equations. In *Proceedings of the 31st International Conference on Machine Learning (ICML-14)*, pages 1485-1493. *JMLR Workshop and Conference Proceedings*.
- [3] Beckers T. 2021. An introduction to Gaussian process models. *arXiv:2102.05497v1*.
- [4] Black F., Scholes M. 1976. The Pricing of options and corporate liabilities. *Journal of political economy*, 81, 637-654.
- [5] Calderhead B. , Girolami M., Lawrence N.D. 2009. Accelerating Bayesian inference over nonlinear differential equations with Gaussian processes, in *Advances in Neural Information Processing Systems*, pp. 217-224.
- [6] Cialenco I., Fasshauer G.E., Ye Q. 2012. Approximation of Stochastic Partial Differential Equations by a Kernel-based Collocation Method. *International Journal of Computer Mathematics*. Vol 89 (18), p. 2543-2561.
- [7] Chen Y., Hosseini B., Owhadi H., Stuart A.M. 2021. Solving and Learning Nonlinear PDEs with Gaussian Processes. *Journal of Computational Physics*, 447, 110668.
- [8] Cockayne J., Oates C., Ipsen I., Girolami M. 2019. A Bayesian conjugate gradient method (with discussion). *Bayesian Analysis*, 14(3), p. 937–1012.
- [9] Crépey S. Dixon M. 2019. Gaussian Process Regression for Derivative Portfolio Modeling and Application to CVA Computations. *arXiv:1901.11081*
- [10] Cuomo S., Di Somma V., Sica F. 2020. A note on the numerical resolution of Heston PDEs. *Ricerche di Matematica* (2020) 69:501–508.
- [11] De Spiegeleer J., Madan D.B., Reyners S., Schoutens W. 2018. Machine learning for quantitative finance: fast derivative pricing, hedging and fitting. *Quantitative finance*, 18 (10), 1635–1643.
- [12] Düring B., Miles J. 2017. High-order ADI scheme for option pricing in stochastic volatility models. *J. of Comp. and App. Math.* 316, 109-121
- [13] Dondelinger F. , Rogers S., Husmeier D. 2013. ODE parameter inference using adaptive gradient matching with Gaussian processes. In *Sixteenth International Conference on Artificial Intelligence and Statistics; AISTATS*.

- [14] Goudenège, L., Molent, A., Zanette, A., 2021. Gaussian process regression for pricing variable annuities with stochastic volatility and interest rate. *Decisions in Economics and Finance*, 44, 57–72.
- [15] Gonzalvez J., Lezmi E., Roncalli T., Xu J. 2019. Financial Applications of Gaussian Processes and Bayesian Optimization. arXiv:1903.04841.
- [16] Graepel T. 2003. Solving Noisy Linear Operator Equations by Gaussian Processes: Application to Ordinary and Partial Differential Equations. ICML'03: Proceedings of the Twentieth International Conference on International Conference on Machine Learning Pages 234 - 241
- [17] Han J., Zhang X.P., Wang F., 2016. Gaussian process regression stochastic volatility model for financial time series, in *IEEE Journal of Selected Topics in Signal Processing*, 10 (6), p. 1015-1028.
- [18] Hainaut D. 2023 b. Pricing of spread and exchange options in a rough jump-diffusion market. *J. of Comp. and App. Math.* 419, 114752.
- [19] Hainaut D., Casas A., 2024. Option pricing in the Heston model with physics inspired neural networks. *Annals of finance*. <https://doi.org/10.1007/s10436-024-00452-7>
- [20] Hainaut D. 2024 a. Valuation of guaranteed minimum accumulation benefits (GMABs) with physics-inspired neural networks. *Annals of actuarial sciences*. <https://doi.org/10.1017/S1748499524000095>
- [21] Hainaut D. 2024 b. American option pricing with model constrained Gaussian process regressions. UCLouvain LIDAM working paper. <https://dial.uclouvain.be/pr/boreal/object/boreal:292665>
- [22] Hainaut D., Dupret J.L. 2025. Optimal control by policy improvements and constrained Gaussian process regressions. UCLouvain LIDAM working paper. https://drive.google.com/file/d/1iHK5nEF2TI5cdnAaxeopslS6GtH_wJud/view
- [23] Jeanblanc M., Yor M., Chesney M. 2009. *Mathematical methods for financial markets*. Springer-Verlag London.
- [24] Murphy K.P. 2012. *Machine learning: a probabilistic perspective*, MIT press.
- [25] Nguyen N.C., Peraire J. 2015. Gaussian functional regression for linear partial differential equations. *Computational methods in applied mechanics and engineering*, 287, p. 69-89.
- [26] Petelin D., Sindelar J., Prikryl J., Kocijan J., 2011. Financial modeling using Gaussian process models. *Proceedings of the 6th IEEE International Conference on Intelligent Data Acquisition and Advanced Computing Systems*, Prague, Czech Republic, 2011, pp. 672-677.
- [27] Rasmussen C.E., Williams C.K.I. 2006. *Gaussian processes for machine learning*, MIT press.
- [28] Raissi M., Perdikaris P., Em Karniadakis E. 2017. Machine learning of linear differential equations using Gaussian processes. *Journal of Computational Physics* 348, p. 683-693.
- [29] Sarkka S. 2011. Linear operators and stochastic partial differential equations in Gaussian process regression. In *Artificial Neural Networks and Machine Learning - ICANN 2011: 21st International Conference on Artificial Neural Networks*, Espoo, Finland, June 14-17, 2011, *Proceedings, Part II*, pages 151-158. Springer, Berlin, Heidelberg.
- [30] Sirignano J., Spiliopoulos K., 2018. DGM: a deep learning algorithm for solving partial differential equations, *Journal of Computational Physics*, 375, 1339–1364.

- [31] Souto Arias L.A. , Cirillo P., Oosterlee C.W. 2025 The Heston–Queue–Hawkes process: A new self-exciting jump–diffusion model for options pricing, and an extension of the COS method for discrete distributions, *J. of Comp. and App. Math.* 454, 116177.
- [32] Swiler L., Gulian M., Frankel A., Safta C., Jakeman J. 2020. A survey of constrained Gaussian process regression: approaches and implementation challenges. *Journal of Machine Learning for Modeling and Computing* 1(2) p. 119-156.
- [33] Pförtner, M., Steinwart I., Hennig P., Wenger J. 2022. Physics-Informed Gaussian Process Regression Generalizes Linear PDE Solvers. *arXiv:2212.12474*
- [34] Wilkens S., 2019. Machine learning in risk measurement: Gaussian process regression for value-at-risk and expected shortfall. *Journal of risk management in financial institutions.* 12 (4), p. 374-383.
- [35] Wu Y., Hernández-Lobato J.M., Ghahramani Z. 2014. Gaussian process volatility model. *Advances in Neural Information Processing Systems* 27, proceedings NIPS 2014.