

Extrapolation and generative algorithms for three applications in finance

Philippe G. LeFloch*, Jean-Marc Mercier†, and Shohruh Miryusupov†

April 2024

Abstract

For three applications of central interest in finance, we demonstrate the relevance of numerical algorithms based on reproducing kernel Hilbert space (RKHS) techniques. Three use cases are investigated. First, we show that extrapolating from few pricer examples leads to sufficiently accurate and computationally efficient results so that our algorithm can serve as a pricing framework. The second use case concerns reverse stress testing, which is formulated as an inversion function problem and is treated here via an optimal transport technique in combination with the notions of kernel-based *encoders*, *decoders*, and *generators*. Third, we show that standard techniques for time series analysis can be enhanced by using the proposed *generative algorithms*. Namely, we use our algorithm in order to extend the validity of any given quantitative model. Our approach allows for conditional analysis as well as for escaping the ‘Gaussian world’. This latter property is illustrated here with a portfolio investment strategy.

1 Introduction

Motivation. We build on, and further expand, our earlier research study [5, 6], which led us to a novel class of kernel-based algorithms and allowed us to deal with applications of central interest in finance. Kernel-based methods are extremely efficient for financial analytics thanks to several fundamental advantages: they provide critical interpretability for audit and regulatory compliance, offer robustness in sparse data sce-

narios, and maintain computational efficiency which is critical for real-time analysis. Such methods have become increasingly valuable in finance, and become recognized for their ability to perform complex nonlinear transformations—in turn enhancing the predictive power of any given model without compromising algorithmic efficiency.

State of the art and main contribution. This Note introduces to several applications which employ (either now standard or recently developed) kernel-based algorithms. Among the traditional techniques, we focus first on the so-called predictive machines, formulated as methods for interpolation and extrapolation of data, showcasing their relevance in different financial contexts. We present a novel generative method, which combines ideas from optimal transport theory and kernel-based methods. Generative methods became popular under the name of Generative Adversarial Networks (GANs), which have emerged as a powerful class of models and allow one to generate very realistic images. The latter have not only captured the imagination of the public through, for instance, Midjourney¹, but have also spurred extensive research and development, eventually leading to the creation of over 500 GAN variants². Among the diverse families of GANs, conditional GANs stand out for their unsupervised learning capability and produce images based on conditional inputs, while TimeGan [11] allows the generation of realistic time series data. In parallel with the evolution of GANs, another significant strand of research in generative modelling takes its roots in traditional kernel methods, such as Sinkhorn autoencoders [10] and Schrödinger bridge generative models [4] for image and time series data. In agreement with optimal transport-based generative techniques, our methodology for constructing

*Laboratoire Jacques-Louis Lions, Sorbonne University and Centre National de la Recherche Scientifique, 4 Place Jussieu, 75258 Paris, France. Email: contact@philippelefloch.org

†MPG-Partners, 136 Boulevard Haussmann, 75008 Paris, France. Email: jean-marc.mercier@mpg-partners.com, shohruh.miryusupov@mpg-partners.com.

¹<https://www.midjourney.com/home>

²<https://github.com/hindupuravinash/the-gan-zoo>

suitable algorithms was inspired by the method of Nadaraya-Watson kernel regression [9], which guided us to design our sampling algorithm for conditioned density estimations. At the core of our method is a sampling algorithm that employs optimal transport and maps a white-noise latent space directly to the target distribution space. This process is uniquely implemented by using permutation indices of samples: the mappings arising in the computation can be traced back and interpreted in the applications, which is an essential feature in complex distribution analysis.

Applications in finance. We illustrate the performance of our algorithms in three practical applications within quantitative finance.

- **Online predictions of PnL and its sensitivities.** By adapting our prediction algorithm, we provide a framework for online forecasting of profit and loss (PnL) statements and their sensitivities, enhancing the robustness of financial decision-making processes.
- **Reverse-stress test (PnL function inversion).** A novel application of our permutation algorithms allows us to perform reverse-stress testing. By inverting the PnL function, we can better understand the conditions that lead to extreme financial outcomes, preparing for worst-case scenarios in risk management.
- **Quantitative models and generative methods.** Traditional financial models frequently rely on predefined stochastic processes, such as Brownian motions, which might not capture the full spectrum of market complexities and dynamics. Our methodology progresses in two significant directions: Initially, we ascertain that a majority of quantitative models can be conceptualized as mappings, transforming time series data into white noise. This insight paves the way for our novel contribution, where we employ optimal permutations and mappings to enhance the sophistication of quantitative models, exemplified through the refinement of the GARCH process. This innovation marks a significant advancement in our ability to model intricate market behaviors. Additionally, we employ a conditional probability estimator to devise a portfolio management strategy, anchored in precise

market indicators. This strategy not only deepens our understanding of market dynamics but also facilitates the formulation of superior investment strategies.

2 Performing stress tests via an extrapolation algorithm

Aim. We present our extrapolation algorithm, particularly designed for asset pricing applications when dealing with potentially large multi-asset portfolios. Our motivation for proposing this test stems from the fact that pricing engines in the finance industry are often computationally intensive and therefore struggle with real-time deployment. Our aim here is to demonstrate that learning a pricing function offline and then using extrapolation provide an approach that is sufficiently accurate for dynamical use as a real-time risk/pricing framework for stress testing or for Profit and Loss (PnL) analysis.

Notation: market data and time series. We consider time series denoted by $t \mapsto X(t) \in \mathbb{R}^D$ which is observed on a time grid $t^{-T_x} < \dots < t^0$. In this context, t^0 denotes the pricing date and our notation is as follows:

$$X = \left(x_d^{n,k} \right)_{d=1 \dots D}^{k=1, \dots, T_x} \in \mathbb{R}^{D, T_x}. \quad (2.1)$$

In our example, the observed data comprise 253 closing values denoted $x^{-252}, \dots, x^0$, for the S&P500 market index during the period of time from $t^{-252} = \text{June 1, 2021}$ to $t^0 = \text{June 1, 2022}$. (These data were retrieved from Yahoo Finance.) In the notation (2.1), this dataset is represented by a matrix with dimensions $D = 3, T_x = 253$ and the corresponding charts are displayed in Figure 1.

The integer T_x is the number of historical observations of the time series X , D represents the number of components of the observed process, and, for the sake of simplicity in the presentation, our example is based on three assets.

Notation: portfolio. We focus on a function $P(t, x) \in \mathbb{R}^{D_P}$, where $x \in \mathbb{R}^D$ and which represents an external pricing engine, evaluating a portfolio of D_P instruments, based on assets valued at x at time t . This setup allows us to calculate the portfolio's value at any given time, especially at its maturity T , where the payoff

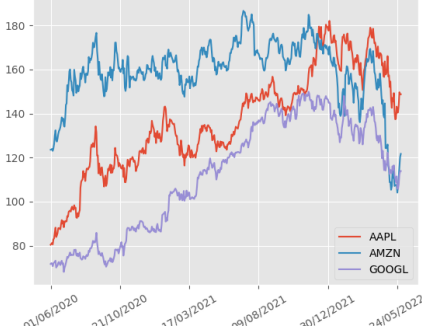


Figure 1: Charts for Apple, Amazon, and Google

is specifically defined as $P(T, x)$. To simplify our presentation, we assume a single instrument, $D_P = 1$, focusing on a basket option with underlying assets. The payoff for this basket option is computed as $P(T, x) = \max(\langle \omega, x \rangle - K, 0)$, where $\langle \omega, x \rangle$ represents the weighted sum of the basket's asset values with ω being the equal weights. Here, K represents the strike price of the option and T denotes its maturity. Figure 2 displays, on the left-hand side, the payoff as a function of the basket values and, on the right-hand side, the pricing engine values as a function of time. For demonstration purposes, we use a simplified Black-Scholes model, denoted as $P(t, x) := \text{BS}(S, T - t, \sigma, K)$ with $S = K = \langle \omega, X^0 \rangle$, implying the initial basket value equals the strike price. Of course, our framework can accommodate many different types of pricing functions.

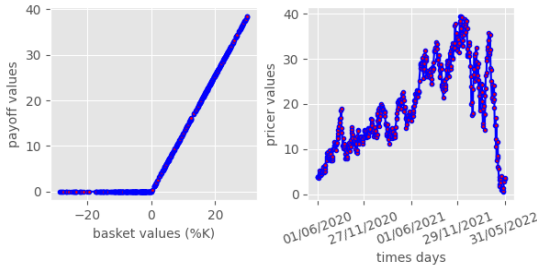


Figure 2: Payoff (left), and pricer (right) values

A kernel-based extrapolation algorithm. Using historical data, we generate synthetic data for a future date $t^1 = t^0 + H$ days through time series forecasting, as described earlier, with $H = 10$ days in this example, in order to simulate stress tests scenarios. These data are denoted by $X := (x_d^n)_{d=1 \dots D}^{n=1 \dots N_x} \in \mathbb{R}^{N_x, D}$. Similarly, we

produce another set of scenarios $Z \in \mathbb{R}^{N_z, D}$ using the same methodology. Consequently, we simulate learning the pricing engine values $P(X) \equiv P(t^1, X)$ from the market data X , referred to as the *training set*, and then extrapolate this function onto the set Z , referred to as the *test set*, to compare against the actual pricing engine values $P(Z)$.

To describe the extrapolation procedure, we introduce basic notions relative to kernel operators, and refer to [1] for a comprehensive overview of reproducing kernel Hilbert space RKHS theory. Let $\mathcal{X} \subset \mathbb{R}^D$ a convex set, we call a function $k : \mathcal{X} \times \mathcal{X} \mapsto \mathbb{R}$ a kernel if it is a symmetric and positive definite (see [1] for a definition). If $(X, Y) := (x_d^n)_{d=1 \dots D}^{n=1 \dots N_x}, (y_d^n)_{d=1 \dots D}^{n=1 \dots N_y}$ are two set of distinct points, we define $k(X, Y) := k(x^i, y^j)_{i=1 \dots N_x}^{j=1 \dots N_y}$ the Gram matrix. Let P be any function taking values on $\mathcal{X}$, and denote $P(X) := P(x^1), \dots, P(x^{N_x})$ its discrete values. Then, for any $z \in \mathbb{R}^D$, we define the *projection* as

$$P_k(X, Y)(z) := k(Y, z)k(X, Y)^{-1}P(X). \quad (2.2)$$

In this equation, the inverse of the Gram matrix $k(X, Y)^{-1}$ is computed by applying a standard least-square method. Extrapolation is defined as $z \mapsto P_k(X, X)(z)$, denoted $P_k^X(z)$ for short, and is a *reproducible* operation, in the sense that it satisfies $P_k(X) = P(X)$; namely, it is exact on the training set. In a similar way, we can define other types of derivative operators such as, for instance, the gradient operator

$$\nabla P_k(X, Y)(z) := (\nabla k(Y, z))k(X, Y)^{-1}P(X), \quad (2.3)$$

and we denote $\nabla P_k^X(z)$ for short in the extrapolation mode. To benchmark the proposed approach, we compute the prices on the test set Z using three methods and display our results in Figure 3.

- Analytical prices : it is computed as $P(t^1, Z)$, and is the reference values for our tests.
- Predicted prices or PnL: the price function $P(t^1, Z)$ is approximated using the formula (2.2), as $P_k(Z)$.
- Δ - Γ approximation: the price function $P(t^1, Z)$ is computed using a second order Taylor formula approximation around $P(t^0, x^0)$, hence involving the price, four

derivatives (three for each asset, one for time), and sixteen second order derivatives.

3

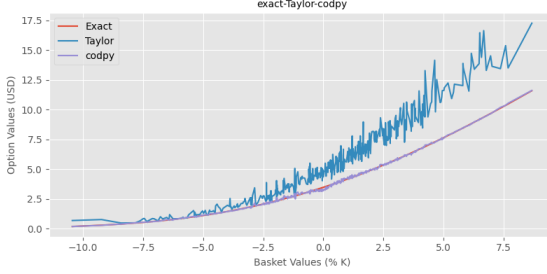


Figure 3: A benchmark of PnL extrapolation methods

In the same vein, we can also compute greeks using (2.3), resulting in four plot, three deltas for each asset, one theta for time, in Figure 4.

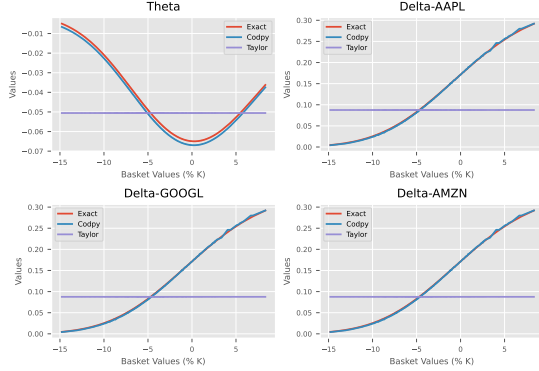


Figure 4: A benchmark of PnL greeks from an extrapolation methods

In conclusion. The above experiments show that extrapolation methods achieve basis point accurate method⁴, even with a limited number of pricing examples, demonstrating their viability for real-time pricing applications. We point out that our approach offers multiple ways for further improvement and optimization. For instance, the selection of stress test scenarios, denoted by X , can be refined through clustering-based strategies, for instance by relying on *sharp discrepancy* sequences; see [5].

³We benchmarked against a Taylor approximation, as this method is currently used by some banks to estimate their PnL on a real time basis.

⁴See [7] for further testing.

3 Reverse stress tests as encoders

Aim. Reverse stress testing stands in contrast to traditional stress testing by starting with a specific outcome, such as portfolio values or PnL losses denoted as $p \in \mathbb{R}^{D_p}$, and backtracking to uncover the market scenarios that could lead to such outcomes. This approach is invaluable for identifying vulnerabilities within a portfolio and enhancing risk management strategies. We illustrate this concept by considering portfolio or PnL values $p \in \mathbb{R}^{D_p}$, aiming to reverse-engineer the market data scenarios $x = P^{-1}(t^1, p)$ using a pricer function P , known from discrete values $P := P^1, \dots, P^{N_x} \in \mathbb{R}^{N_x, D_p}$, where $P^n = P(t^1, X^n)$ is evaluated on market data $X := X^1, \dots, X^{N_x}$. The challenge arises when this mapping, $x \mapsto P(x)$, from $\mathbb{R}^D$ to $\mathbb{R}^{D_p}$, lacks an obvious inverse due to its non-invertibility. We address this challenge by introducing the concept of encoders and decoders, borrowed from machine learning, but using kernel-RKHS methods.

Encoders, decoders and generators. Encoders in this particular numerical experiment are conceptualized as smooth, invertible maps, $x \mapsto P(x)$ from $\mathbb{R}^D$ to $\mathbb{R}^{D_p}$ that bridge the gap between the market data X and the portfolio values P , and we denote this maps as $x \mapsto \mathcal{L}(X, P)(x)$, while its inverse is formally denoted $p \mapsto \mathcal{L}(P, X)(p)$. The extrapolation operator (2.2), denoted by the equation below, initially suggests a direct inversion approach:

$$k(P, p)k(P, P)^{-1}X. \quad (3.1)$$

Yet, this direct approach may falter due to the inherent non-invertibility of P as a mapping from $\mathbb{R}^{D_p}$ to $\mathbb{R}^D$. To enhance stability, we propose a refined method, relying on a permutation $\sigma : [1, \dots, N_x] \mapsto [1, \dots, N_x]$ of the original data set, written as

$$\mathcal{L}(P, X)(p) := k(P, p)k(P, P)^{-1}X^\sigma \quad (3.2)$$

In the particular case where spaces have matching dimensions, that is $D_x = D$ in our notation (cf. also [2]), considering a distance function $d(x, y)$, optimal transport theory proposes to determine this permutation as

$$\bar{\sigma} = \arg \inf_{\sigma \in \Sigma} \sum_{n=1}^{N_x} d(X^{\sigma(n)}, P^n) \quad (3.3)$$

where Σ , the set of all permutations. For kernels methods, a natural distance is given by $d_k(x, y) = k(x, x) + k(y, y) - 2k(x, y)$, called the kernel discrepancy or maximum mean discrepancy, see [3].

However, the existing literature seems less profuse when the function maps unrelated spaces, that is $D_x \neq D$ in our notation. Hence we introduced the following Ansatz in this case, based on the gradient formula (2.3)

$$\bar{\sigma} = \arg \inf_{\sigma \in \Sigma} \|(\nabla k(P, P))k(P, P)^{-1} X^\sigma\|_2^2. \quad (3.4)$$

This approach is reminiscent of a generalized *traveling salesman problem*⁵ and is not equivalent to the one in 3.3, however it allows for the determination of a smooth mapping between the original and target spaces, as desired.

Encoders and decoders allows to define the notion of **generators**, needed later on. Let $\mathbb{X}, \mathbb{Y}$ two continuous distributions taking values in $\mathbb{R}^{D_x}, \mathbb{R}^{D_y}$, and consider $X \in \mathbb{R}^{N, D_x}, Y \in \mathbb{R}^{N, D_y}$ two variates of equals length. Then (3.1) provides a method to *generate* a sample y , from a sample x , according to the formula $y = k(X, x)k(X, X)^{-1}Y^\sigma$. In particular, if one consider $\mathbb{X}$ as a known distribution, as for instance a uniform one, this provides a modeling of any distribution $\mathbb{Y}$ by a continuous one, that is statistically similar to the observed data Y .

In conclusion. Using (3.4), we produced the following picture, corresponding to a reverse stress test, that we comment thereafter.

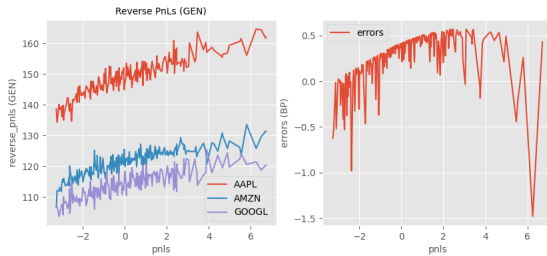


Figure 5: Reverse prices (left) and benchmark (right)

Consider the original set of market data $X = X^1, \dots, X^{N_x}$ and corresponding PnL $P = P^1, \dots, P^{N_x}$. Using a generator, as described

⁵see wikipedia page https://en.wikipedia.org/wiki/Travelling_salesman_problem

in the previous section, we sampled new examples $\bar{P} = \bar{P}^1, \dots, \bar{P}^{N_x}$. These simulated distribution of prices is used with the reverse stress test method to compute the corresponding scenarios at left, notation $\bar{X} = \bar{X}^1, \dots, \bar{X}^{N_x}$. To benchmark this result, we computed the error distribution $P(\bar{X}^1) - \bar{P}^1, \dots, P(\bar{X}^{N_x}) - \bar{P}^{N_x}$. This distribution is plot as the red line at right, expressed in basis points.

4 Modeling time series via a generative algorithm

Mapping time series. Our approach consists in a general model framework where a time series X is transformed into a latent variable ε , interpreted as white noise, through a continuous and invertible map F . With this approach, the transformation $F(X) = \varepsilon$ isolate the inherent noise within the data, facilitating a deeper analysis and manipulation of the underlying stochastic processes. Let us give a first simple example to fix ideas. Consider any time series X of the form (2.1), then the following simple random walk provides such a modeling:

$$\varepsilon^k = X^{k+1} - X^k, F^{-1}(\varepsilon)^k = X^0 + \sum_{n \leq k} \varepsilon^n.$$

A time series model can then be interpreted as an invertible mapping from $X \in \mathbb{R}^{D_x, T_x}$ into $\varepsilon \in \mathbb{R}^{D_\varepsilon, T_\varepsilon}$. As demonstrated with the GARCH model in the following, most quantitative models can be interpreted through such mappings F , and we provide numerous other examples in [7]. The purpose of such an approach is to provide a framework to deal with time series enjoying the following properties.

- The time series X is *reproducible*, meaning that $X = F^{-1}(\varepsilon)$. This property ensures that X , the historical dataset, is in the range of the model.
- The variable $\varepsilon \in \mathbb{R}^{D_\varepsilon, T_\varepsilon}$ can be reproduced from any other random variable $\eta \in \mathbb{R}^{D_\eta, T_\eta}$, for instance a uniform sampling, using a *generator*, as described earlier. Note that the distribution that generates the variate ε can be rather arbitrary, and we are not restricted to assuming Gaussian distributions for time series, as is usually the case in classical modeling of time series.
- It provides a simple Monte-Carlo framework,

as sampling new $\bar{\varepsilon}$ generates a new trajectory defined as $\bar{X} = F^{-1}(\bar{\varepsilon})$.

From an economical point of view, these models simulates new trajectories that are calibrated to the observed dynamics of the financial time series, as they rely on reproducing historically observed white noise. This way of modeling ensures that short-term dynamics of time series are easily and better captured. It provides also a powerful tool for forecasting, risk assessment, and strategy benchmarking, as we can apply this framework in the following areas.

- Benchmarking strategies, where we resample the original signal X on the same time-lattice to draw several simulated trajectories and compare them to the original one using various performance indicators.
- Monte-Carlo forecast simulations for future time points, allowing for the exploration of potential future scenarios based on historical data.
- Forward Calibration, where we frame it as a minimization problem with constraints, optimizing the generation of new samples that meet specific financial criteria.
- PDE pricers, enabling the computation of forward prices or sensitivities by solving backward Kolmogorov equations in a multi-dimensional tree structure, see [5].

The GARCH(p, q) model as an example. We show an example of such an model extension considering the GARCH model, in order to illustrate its flexibility and adaptability to any quantitative model, but also its potential to enhance the modeling and simulation of financial time series. This approach is not exclusive to GARCH models; it is equally applicable to the broader ARIMA family, as well as a wide range of both discrete and continuous stochastic processes, as local or rough volatility ones.

The generalized autoregressive conditional heteroskedasticity (GARCH) model, particularly in its (p, q) order formulation, captures financial markets' volatility clustering—a phenomenon where high-volatility events tend to cluster together in time. The defining equations of a GARCH(p, q) model are as follows:

$$\begin{cases} X^k = \mu + \sigma^k Z^k, \\ (\sigma^k)^2 = \alpha_0 + \sum_{i=1}^p \alpha_i (X^{k-i})^2 + \sum_{i=1}^q \beta_i (\sigma^{k-i})^2, \end{cases}$$

where μ is the mean, σ^k denotes the time-varying

volatility, and Z^k symbolizes a white noise process. The coefficients α_i and β_i govern the model's responsiveness to changes in volatility and the inertia of past volatility, respectively.

The variance equation of GARCH can be compactly expressed using the backshift operator B , leading to

$$(1 - \beta(B))(\sigma^k)^2 = \alpha_0 + \alpha(B)(X^k)^2,$$

where $\alpha(B) = \sum_{i=1}^p \alpha_i B^i$ and $\beta(B) = \sum_{i=1}^q \beta_i B^i$ and backshift $B^i X^k = X^{k-i}$ allows for the compact expression of lagged effects on the stochastic process. By defining $\varphi(B) = \alpha_0 + \sum_{i=1}^p \alpha_i B^i$, $\theta(B) = 1 - \sum_{i=1}^q \beta_i B^i$, the stochastic variance σ^k can be derived as

$$\sigma^k = \sqrt{\varphi^{-1}(B)\theta(B)(X^k)^2} = \sqrt{\pi(B)(X^k)^2}.$$

with $\pi(B) = \varphi^{-1}(B)\theta(B)$ encapsulating the variance transformation.

Our methodology integrates the GARCH model within a broader generative framework, identifying the 'GARCH map', $G : X^k \mapsto Z^k$, which translates the observable time series X^k into a latent white noise process Z^k :

$$Z^k = G(X^k) = \sqrt{\pi^{-1}(B)(X^k)^2}(X^k - \mu).$$

This GARCH map forms the basis of our three-stage generative process for financial time series, specifically demonstrated through the example of stock prices in Figure 1.

- By applying the GARCH map to the historical stock prices X , we extract the latent noise $\varepsilon = G(X)$. This latent noise encapsulates the fundamental stochastic processes driving the volatility modeled by the GARCH framework.
- Utilizing the permutation algorithm (3.3) as a generator, we produce new instances of the latent variable $\tilde{\varepsilon}$, simulating alternative realizations of the market's stochastic behavior.
- Employing the inverse GARCH transformation, symbolically represented as G^{-1} , we map these new latent samples back to the domain of stock prices, thereby generating new trajectories $\tilde{X} = G^{-1}(\tilde{\varepsilon})$, consistent with the GARCH dynamics.

Figure 6 displays ten simulated trajectories of one of the three stock, Amazon's ones, using the

GARCH(1,1) model, the others two producing quite similar patterns. This visualization exemplifies how the generative process can give us a diverse set of plausible future paths for the stock price, rooted in the historical volatility patterns captured by the GARCH model.

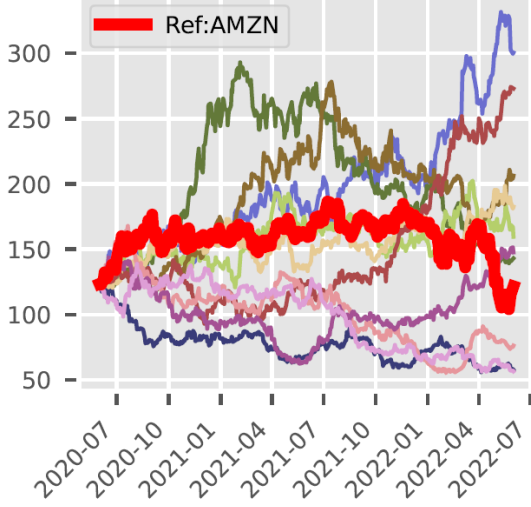


Figure 6: Ten examples of generated paths with the GARCH(1,1) model

5 Extending quantitative models through conditioning

Further applications. We finally demonstrate that the generative algorithms presented so far can be safely used for conditional analysis; this is done on a numerical experiment in the context of portfolio management. For the sake of illustration, we consider a time series given by closing of a basket of 106 crypto currencies during the period from 19/07/2021 to 28/04/2023, which was observed on a daily basis and corresponds to a time series X with $D, T_x = 106, 649$. We rely here on our notation (2.1), and we present the plot in Figure 7 after suitable normalization at the initial time.

To this aim, consider a integer W defining a sliding window $t^{n+k}, k = 0, \dots, W, n = 0, \dots, T_x - W$. At each time labelled n , we define several portfolios determined by their components $\omega_{n,d}$, and their wealth $P^n := \langle \omega_n, X^{n+W} \rangle = \sum_d \omega_{n,d} x_d^{n+W}$. Each strategy corresponds to an investment at the time t^{n+W} , and portfolio

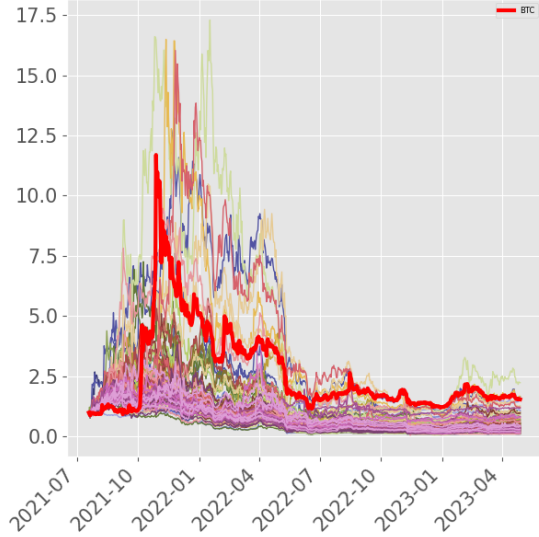


Figure 7: A time series of 106 crypto assets

performances are compared together with, and without, conditional analysis. In the next paragraph, we describe the investment strategies that we used, and how conditional analysis can be used to enhance them.

Portfolio management usually build upon a simple quantitative model, determined by *returns* of the assets, that is, following the notation introduced in the previous section, returns are determined as the discrete distributions $\epsilon_n = \left(\frac{X^{n+k+1}}{X^{n+k}} - 1 \right)_{k=1, \dots, W}$, and the inverse map is $F_n^{-1}(\epsilon) = \left(X^n \prod_{k=1}^W (1 + \epsilon^{n+k}) \right)_{k=0, \dots, W}$.

Efficient portfolio. Let us recall the Markowitz mean-variance optimization, providing a classical investment strategy. It consists in finding a set of portfolio weights $\omega := (\omega_d)_{d=1, \dots, D}$ as the solution of the quadratic programming problem

$$\bar{\omega}(\epsilon) = \arg \inf_{\omega} \frac{1}{2} \omega^T Q \omega - \epsilon \omega^T \bar{\epsilon} + \beta (|\omega - \omega^0|). \quad (5.1)$$

The terms β represent the transaction costs, which are proportional to the portfolio change in weights, symbolized here in the term $|\omega - \omega^0|$. This comes usually with constraints over the weights, and we used $\sum_d \omega_d = 0$ (long / short strategy) together with $|\omega_i| \leq 1$.

As our strategies are determined via a random variable modeling the asset returns ϵ ⁶, we con-

⁶as point out earlier, this conditional method can be used with any other quantitative modeling of the assets

sider two quite comparable situations, where this random variable is conditioned by two different observed random variable η , that is we considered the distribution $\epsilon|\eta = \eta^n$ for each time t^n , $n = W, \dots T_x$. Now, we are going to describe how the generative methods of the previous section can be applied to conditioning random variables.

Conditioned random variables and latent spaces. Let $\mathbb{X} \in \mathbb{R}^{D_x}$, $\mathbb{Y} \in \mathbb{R}^{D_y}$ two dependent random variables and consider $\mathbb{Z} = (\mathbb{X}, \mathbb{Y}) \in \mathbb{R}^{D_x+D_y}$ the joint random variable. Let $Z = (X, Y) \in \mathbb{R}^{N, D_x+D_y}$ be a variate of $\mathbb{Z}$. Relying on (3.1), we consider another variate ϵ drawn from any known distribution, say $\epsilon = (\epsilon^n)_{n=1\dots N}$, which is called the *latent* and is decomposed as $\epsilon := (\epsilon_x, \epsilon_y)$. Elaborating on the two encoders map in (3.1), namely $\mathcal{L}(X, \epsilon_x)$ and $\mathcal{L}(\epsilon, Z)$, the following construction provides a generator of the conditioned law $\mathbb{Y}|\mathbb{X} = x$:

$$\eta_y \mapsto \mathcal{L}(\epsilon, Y)(\eta_x, \eta_y), \quad \eta_x = \mathcal{L}(X, \epsilon_x)(x).$$

There is a lot of freedom in choosing the distributions ϵ_x, ϵ_y . We can pick up, for instance, uniform distributions, or the trivial map $\epsilon_x = X$, which can be handy in certain applications.

Benchmarks results. Next, we provide an illustration of allocation strategies. We compared four strategies in figure 8, listed as follows.

- The first one is given by an equiweighted portfolio, which we call ‘index’ and serves as a reference.
- The second one (LS) is a Long Short strategy, determined as approximating the mean variance problem (5.1).
- The third one is also a Long Short strategy, but where the return distribution at time t^n , ϵ^n , is conditioned to the capital asset pricing model, that is, $b^n = r_f + \beta^k(\bar{\omega}^k - r_f)$, in which β^k is the regression coefficients of the weights ω^k .
- The third one is very similar to the previous one, except that the distribution ϵ^n at time t^n is conditioned with a distribution b^n contains for each assets liquidity values, time averages on different windows as well as differences between them.

We reported the daily performance of each portfolio $P_j, j = 1, \dots, 4$, as $P_j^n := \Pi_{k \leq n} \frac{\langle \omega_j^k, X^{k+1} \rangle}{\langle \omega_j^k, X^k \rangle}$.

As we can see, strategies based on conditional returns outperformed non conditioned ones on

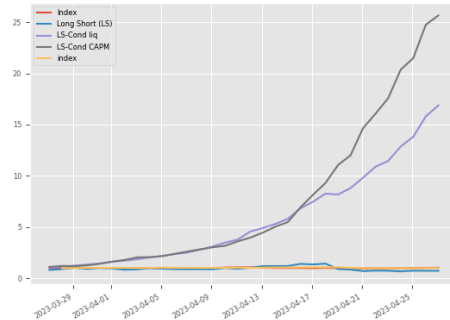


Figure 8: Benchmark of different portfolio strategies

these examples. Of course, this experiment has to be considered as a simple validation of our algorithms, as we conditioned our returns values ϵ^n with distributions that are known at time t^n . In an operational context, the available information for conditioning comes from the past, that is the time t^{n-1} . In particular, given a conditioning distribution, this test does not answer the quite interesting question whether there exists a time frequency of observations below which such an approach could be profitable or not.

References

- [1] A. BERLINET AND C. THOMAS-AGNAN, *Reproducing kernel Hilbert spaces in probability and statistics*, Springer US, Kluwer Academic Publishers, 2004.
- [2] H. BREZIS, Remarques sur le problème de Monge–Kantorovich dans le cas discret, *Comptes Rendus Math.* 356 (2018), 207–213.
- [3] A. GRETTON, K.M. BORGWARDT, M. RASCH, B. SCHÖLKOPF, AND A.J. SMOLA, A kernel method for the two sample problems, *Proc. 19th Int. Conf. on Neural Information Processing Systems*, 2006, pp. 513–520.
- [4] M. HAMDOUCHE, P. HENRY-LABORDERE, AND H. PHAM, Generative modeling for time series via Schrödinger bridge. Available as arXiv:2304.05093.
- [5] P.G. LEFLOCH AND J.-M. MERCIER, A class of mesh-free algorithms for some problems arising in finance and machine learning, *J. Scientific Comput.* 95 (2023), 75.

- [6] P.G. LEFLOCH AND J.-M. MERCIER, The transport-based mesh-free method: a short review, *Wilmott journal*, 2020, pp. 52–57.
- [7] P.G. LEFLOCH, J.-M. MERCIER, AND S. MIRYUSPOV, *CodPy: a Python library for numerics, machine learning, and statistics*, Monograph, 2024. Available as ArXiv:2402.07084.
- [8] H.D. LIU, Y. GUO, N. LEI, Z. SHU, S.-T. YAU, D. SAMARAS, AND X. GU, Latent space optimal transport for generative models, 2018. Available as arXiv:1809.05964.
- [9] E. A. NADARAYA, On estimating regression, *Theory Proba. and Appl.* 9 (1964)), 141.
- [10] G. PATRINI, R. VAN DEN BERG, P. FORRÉ, M. CARIONI, S. BHARGAV, M. WELLING, T. GENEWEIN, AND F. NIELSEN, Sinkhorn auto-encoders, 2018. Available as arXiv:1810.01118.
- [11] J. YOON, D. JARRETT, AND M. VAN DER SCHAAR, Time-series generative adversarial networks, *Neural Information Processing Systems (NeurIPS)*, 2019.
- [12] O. ZHANG, R.-S. LIN, AND Y. GOU, Optimal transport-based generative autoencoders, 2019. Available as arXiv:1910.07636.