
Portfolio management with semi-parametric bootstrapping

Beatriz Vaz de Melo Mendes*

holds a PhD in statistics from Rutgers University, New Jersey. Beatriz is an Associate Professor at the Institute of Mathematics/COPPEAD, Federal University at Rio de Janeiro, Brazil. She is an author and co-author of six books on extremes, finance, risk and copulas, and has had papers published in the *Annals of Statistics*, *Journal of Derivatives*, *Computational Statistics*, *Journal of Statistical Computation and Simulation* and the *Journal of Risk*, among many others. Her research areas include financial modelling and risk estimation, robust inference, extremes and copulas.

Ricardo Pereira Câmara Leal

is Professor of Finance at the Coppead Graduate School of Business in Brazil and has taught at Georgetown University and at the University of Nevada. Ricardo's research and consulting is on emerging securities markets and corporate governance. He is a former president of the Brazilian Finance Society and of the Business Association of Latin America Studies, and is the editor of the Brazilian Finance Review. Ricardo is a senior research fellow of the Brazilian Corporate Governance Institute.

*The COPPEAD Graduate School of Business at the Federal University of Rio de Janeiro, PO Box 68514, Rio de Janeiro, RJ 21941-972, Brazil
E-mail: beatriz@im.ufrj.br

Abstract Estimation risk is an important topic within the area of risk management. Uncertainties regarding parameter estimates carry on to the final statistical product, such as investment strategies, and need to be estimated and accounted for. Unless the exact expressions for the estimators' variances are known, the product's variability will be assessed through bootstrap techniques. The present paper addresses this issue, proposing a semi-parametric bootstrap method for reproducing the data, a method which parametrically takes care of all marginal characteristics of the returns data, and also takes care of the dependence structure existing in the data, in a very simple and clever non-parametric way. The technique is applied to the problem of assessing the variability of the Markowitz-efficient frontier. Simulation experiments are conducted to assess the out-of-sample forecasting usefulness of the semi-parametric bootstrap methodology.

Keywords: *resampling, efficient frontier, portfolio management, estimation risk, bootstrap*

INTRODUCTION

Markowitz¹ mean-variance (MV) optimisation has been the standard tool for efficient portfolio allocation and diversification for the last 50 years. It is implemented in probably all commercial portfolio optimisers for asset allocation and equity portfolio management. For a given level of expected return, the

Markowitz model obtains the portfolio composition with the lowest risk (standard deviation). It is a quadratic optimisation problem, where the inputs are just estimates of the population mean μ and covariance matrix Σ .

However, the inputs $(\hat{\mu}, \hat{\Sigma})$ are estimated with errors, and the optimisation routines are known as error

maximisation algorithms. The errors carry on to the weights with magnified variability, resulting in unstable portfolios, extreme weights and poor out-of-sample performance. The problem is how to assess this variability. Unless the exact expressions for the estimator variances are known, estimate variability is typically assessed through bootstrap techniques²—computer-intensive techniques for resampling the data.

The basic idea of resampling the efficient frontier was first introduced to the finance literature by Jorion.³ Since then, discussions on the topic have usually addressed the statistical validity of the resampled curves, the role of the average resampled portfolio (Michaud⁴ patented the use of the average resampled portfolios), and how to obtain the average portfolio, either based on ranks as per Michaud⁴ or based on same risk/return trade-off.

Little effort, however, has been invested in the appropriate selection of a resampling methodology, for which there are essentially two methods: a non-parametric and a parametric bootstrap. The non-parametric bootstrap is the simplest method, in which the bootstrapped observations are sampled from the original data with replacement. The parametric bootstrap usually assumes multivariate normal distribution for the data (more generally, an elliptical distribution), and then simulates from the assumed distribution. Under the parametric approach, either $(\hat{\mu}, \hat{\Sigma})$ are assumed to be the true parameter values and a new data set is simulated, or a new set of inputs is simulated from the distribution assumed for $(\hat{\mu}, \hat{\Sigma})$ (see Scherer and Martin⁵). To deal with serial dependences in the data, the moving

block bootstrap was introduced by Carlstein⁶ and Künch.⁷

To be faithful to the data at hand and to obtain a correct set of replications of the efficient frontier containing the same information brought in by the original one taking into account the nature provided and uncontrolled variability, one must be as close as possible to the true multivariate distribution generating the data. This problem is addressed in the present paper, which proposes a semi-parametric bootstrap for reproducing the data, a method which parametrically takes care of all marginal characteristics of the returns data, and also takes care of the dependence structure existing in the data in a very simple and clever non-parametric way. The semi-parametric bootstrap approach to risk estimation is an alternative to the confidence regions of Jobson⁸ and to the cited bootstrap methods.

In summary, the semi-parametric bootstrap methodology finds the margins' distribution by fitting unconditional or conditional models, and finds the dependence structure linking the margins by using the data ranks. The first step may be as simple as fitting a normal distribution to a return series, and the second step just uses the empirical distribution. Alternative sophisticated univariate time series models for the mean and volatility may be used and usually provide good fits, especially for daily returns.

In what follows, the paper explains the semi-parametric bootstrap. It then goes on to provide a real data illustration of its usefulness when assessing the variability of the efficient frontier and optimal weights. Next, the paper conducts several simulation experiments designed to assess the forecasting ability of the semi-parametric bootstrap method. The

final section includes the concluding remarks.

REPLICATING THE EFFICIENT FRONTIER

Let $\mathbf{r}$ represent a d -dimensional vector of financial returns from a distribution F with mean $\boldsymbol{\mu}$ and (positive definite) covariance matrix $\boldsymbol{\Sigma}$, and let $\mathbf{w} = (w_1, w_2, \dots, w_d)$ be a vector of portfolio weights. The MV algorithm looks for a portfolio allocation that maximises the expected utility $\mathbf{w}'\boldsymbol{\mu} - \lambda \mathbf{w}'\boldsymbol{\Sigma}\mathbf{w}$ over $\mathbf{w} \in \mathbb{R}^d$, subject to $\mathbf{w}'\mathbf{1} = 1$ and $w_i \geq 0$, $i = 1, \dots, d$ (long only), and where λ represents the investor risk aversion, the rate between portfolio return and standard deviation. The smaller the λ , the higher the risk aversion. For example, an extremely risk-averse investor will choose the portfolio with the smallest possible standard deviation, the so-called 'minimum variance portfolio', which has a λ equal to zero.

Assuming that stock returns are jointly normally distributed implies that investors are using an expected utility given by the quadratic form. Even though normality seldom holds for financial returns, many theoretical results are obtained under this assumption. For instance, under normality, weights are unbiased and their variance may be obtained (see Britten-Jones⁹ and Fusai and Meucci¹⁰). Michaud and Michaud¹¹ review and summarise recent research and new developments in estimation error and MV portfolio optimisation.

Despite the method used to estimate $\boldsymbol{\mu}$ and $\boldsymbol{\Sigma}$, when it comes to replicating the efficient frontier, what really matters is how well the data are replicated. To this end, it is necessary to understand the data generating process (DGP). Figure 1 schematically represents the returns DGP.

Consider working with *daily* d -dimensional log-returns $\mathbf{r} = (r_1, r_2, \dots, r_d)$, and consider a period of length T . Every business day t , $t = 1, 2, \dots, T$, nature generates a d -dimensional vector of errors $\boldsymbol{\epsilon} = (\epsilon_1, \epsilon_2, \dots, \epsilon_d)$ according to some zero mean and unit variance marginal distributions $F_1, F_2, \dots, F_d$ not necessarily the same. Linking these marginal distributions is a dependence structure (the oval box in Figure 1), which creates all relationships between the marginal errors, including the interrelationships caused by economic, political and geographic factors and market macro-structure. At each day t , the errors vector passes through a black box holding information on market micro-structure, occasional (bad) news, seasonalities, and so on. At the end of the period T , the output is the observable return data, which may possess serial correlations, and other characteristics besides a mean vector $\boldsymbol{\mu}$ and a $d \times d$ covariance matrix $\boldsymbol{\Sigma}$.

Under the authors' approach, to identify the models properly and obtain faithful replications of the data, the analyst needs to identify the square (black) box and the oval (black) box. They see the oval box as a d -dimensional



Figure 1: The returns data generating process

joint distribution F containing all marginal and joint information, in particular, the errors covariance matrix of the sequence of independent and identically distributed (i.i.d.) random vectors $(\epsilon_1, \epsilon_2, \dots, \epsilon_T)$. To model the information provided by the square black box, the analyst uses the sequence of T (no longer i.i.d.) d -dimensional returns. Univariate time series models are chosen and fitted, in order to capture the temporal dynamics in the mean and in the volatility, for example, some ARFIMA(p, q) and FIEGARCH(r, d, s) type models.

In practice, this is done in the reverse order. The identification of the square box is parametric. Using available computing facilities nowadays it is possible to obtain excellent univariate (conditional or unconditional) fits tailored for each series of log-returns. Note that, in particular, the ARFIMA(p, q)-FIEGARCH(r, d, s) models may have $p = q = r = d = s = 0$, signalling a return to *unconditional models*, with just a few parameters to estimate, such as the mean, standard deviation, asymmetry and kurtosis. In this case, a skew- t (Hansen¹²) distribution is a powerful and flexible option for univariate fitting.

From the d models fitted (say, $M_1, M_2, \dots, M_d$) one obtains d series of standardised (zero mean and unit variance) residuals and identifies their marginal distributions $(\hat{F}_1, \hat{F}_2, \dots, \hat{F}_d)$. The filtered data still contain the information about the oval box; to obtain this information *non-parametrically*, one substitutes the $T \times d$ matrix of standardised residuals by their *matrix of ranks* $\mathbf{R}$. That is, one substitutes each column of residuals by their ranks, numbers between 1 and T . For example,

suppose $T = 100$. Row j of $\mathbf{R}$ may be (9, 65, 34, 78). The closer the numbers, the stronger the dependence structure.

A reader familiarised with copulas may identify the rank matrix $\mathbf{R}$ with the support set for the empirical copula. Genest and Favre¹³ argue that, among all data functions which are invariant under monotone increasing transformations, the data ranks in $\mathbf{R}$ are the statistics retaining the greatest amount of information about the data dependence structure.¹⁴

In summary, to obtain the proposed semi-parametric replications of the data one should:

- (1) Fit conditional or unconditional models $(M_1, M_2, \dots, M_d)$ to the univariate series. In the case of conditional modelling, obtain the standardised residuals and their estimated distributions $\hat{F}_1, \hat{F}_2, \dots, \hat{F}_d$. Derive the rank matrix $\mathbf{R}$.
- (2) For $k = 1, \dots, B$, B large,
 - (a) Generate T i.i.d. random values from each $\hat{F}_i, i = 1, 2, \dots, d$, forming the $T \times d$ matrix of i.i.d. innovations, $\mathbf{Z}^{(k)}$.
 - (b) Order each column of $\mathbf{Z}^{(k)}$ according to the corresponding column in $\mathbf{R}$.
 - (c) Take each ordered column j as a new set of innovations and apply the corresponding model M_j , obtaining a new $T \times d$ data matrix $\mathbf{X}^{(k)}$ in the original scale.

Once the data replications are available, they may be used to compute any quantity of interest. Questions related to the size B are discussed in Efron and Tibshirani.² Each one out of the B new data sets $\mathbf{X}^{(k)}, k = 1, \dots, B$ contains the same marginal (dynamic or not) information found in the original data, as

well as their dependence structure. The novelties are using the rank matrix $\mathbf{R}$ in step 1 to reorder the innovations and reproducing the dependence structure in step 2.2. The idea of resampling using ranks has been inspired by the work of Mendes.¹⁵ The method is amazingly simple and when fitting unconditional models it becomes even simpler, *not* requiring any sophisticated computer software.

ILLUSTRATION

This illustration uses a six-dimensional data set to represent a moderate risk profile investor willing to invest in emerging and developed markets. The data are 1,629 contemporaneous daily log-returns from 2nd January, 2002, to 20th October, 2008, on: (V1) a Brazilian hedge fund index, the Arsenal Composite Index (ACI); (V2) a Brazilian index for inflation-indexed treasury bonds, the IMA-C; (V3) a Brazilian market value weighted stock index, the IBrX; (V4) an index of large world companies (WLDLg, by MSCI Barra); (V5) an index of small world companies (WLDSm, by MSCI Barra); (V6) the Total Return US T-Bonds Index (LBTBond, by Lehman Bros.).

The long only constraint classical efficient frontier is obtained by computing the classical sample estimates. (Note that any estimation method for the inputs could have been used.) Applying the proposed semi-parametric bootstrap and for each replicated data set $\mathbf{X}^{(k)}$ one computes the classical inputs $(\hat{\mu}^{(k)}, \hat{\Sigma}^{(k)})$ obtaining the replications of the efficient frontier. For the sake of comparisons, the study also computes replications based on the parametric bootstrap assuming multivariate normality (as in Scherer and Martin⁵).

When replicating the efficient frontier, what really matters is the weights stability across replications. To investigate this issue, Figure 2 examines the box-plots of weight distribution for portfolios ranked 2, 8 and 15 (out of 20). The original weights for each variable are signed with arrows. The box-plots corresponding to the parametric bootstrap based on normality are marked with ‘-N’. It is impressive how the weight distribution based on the semi-parametric approach is much more concentrated. This is crucial to reduce portfolio rebalancing costs.

One can observe that, in many situations, the parametric bootstrap weight distribution is shifted with respect to the original weight. For example, for portfolio 8 and variables 1 and 2, the original weight values are approximately located at the 0.75 and 0.25 quantiles of the weight distribution. For the same portfolio and variables, the arrows point to the centre of the distribution provided by the new method. This means that, in the d -dimensional space, the distance between the original portfolio's weights and the centre of the distribution of weights from the replications is smaller for the new method. It is possible to measure this.

For any fixed portfolio, one can measure the distance between the centre of its weight distribution and the original weights using the squared Euclidean distance. For example, consider the portfolio ranked position 8: let $\mathbf{w}^*$ represent the original weights and let $\mathbf{w} \in \mathcal{R}^d$ represent the distribution centre, $\mathbf{w} = 1/B \sum_{k=1}^B \mathbf{w}_k$, where $\mathbf{w}_k$ is the optimal set of weights for portfolio 8 for the k -th simulation. The Euclidean distance is $D = (\bar{\mathbf{w}} - \mathbf{w}^*)'(\bar{\mathbf{w}} - \mathbf{w}^*)$. For portfolio 8 and for the new semi-parametric bootstrap, $D = 0.016460$ is

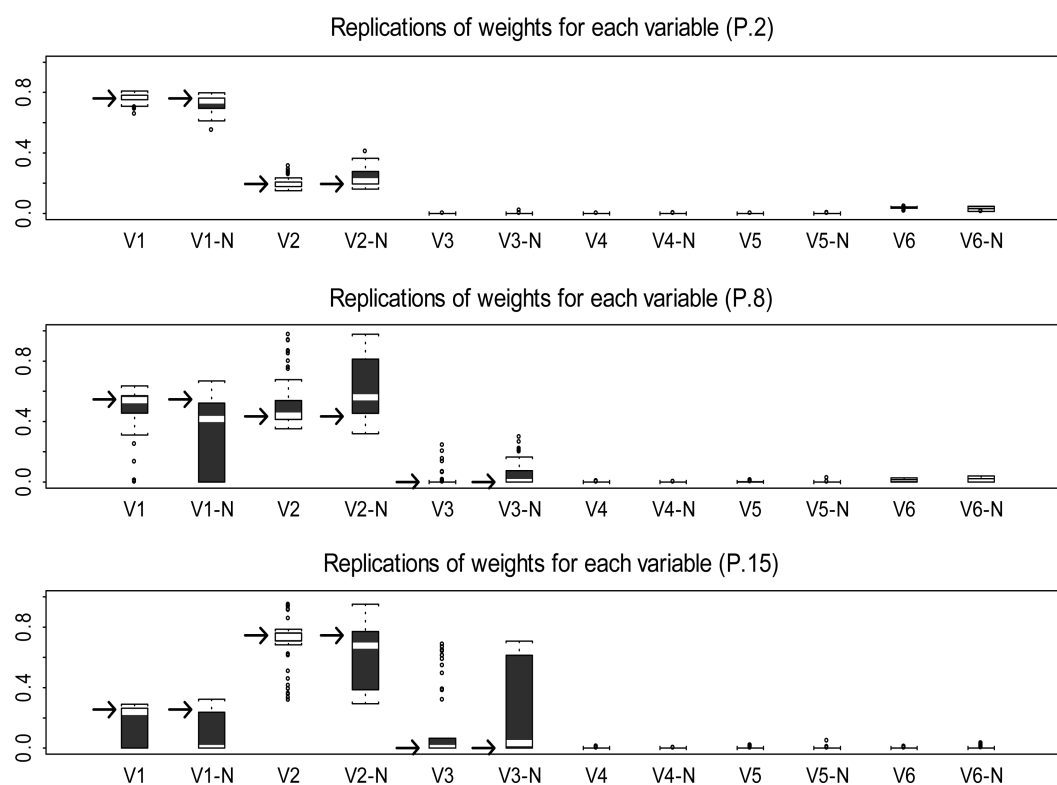


Figure 2: Weight distributions for portfolios ranked 2 (P.2), 8 (P.8), and 15 (P.15), from the semi-parametric method and the parametric (normal) method

computed, whereas for the multivariate normal assumption, one obtains $D = 0.1009245$, a distance 6.13 times greater.

The statistical equivalence between any resampled portfolios may be tested using the Mahalanobis distance, a statistical distance based on the weights covariance matrix. The resulting statistic is chi-squared distributed only under normality of the weight distributions. For any fixed (rank) portfolio, one can visually inspect its convex hull associated with the 90 per cent statistically equivalent portfolios, using only the replications for which all d -variables weights were simultaneously inside a 90 per cent confidence interval.

As pointed out by a referee, ‘showing that simulations from the model fit have less variability is not the same as showing that the method is more efficient’.

To investigate this issue, the paper continues simulation experiments in the next section, comparing the performance of the semi-parametric method with the parametric and non-parametric alternatives.

OUT-OF-SAMPLE VALIDATIONS

This part of the study simulates a real-life situation where a trader manages three clients’ optimal portfolios and wants to forecast the portfolios’ expected utilities for the next period based on the current allocations. The portfolios are defined through three λ values reflecting the clients’ different risk preferences, namely $\lambda = 0.5, 1, 2$. The manager takes the available data along with the current allocations, and replicates the data (500 replications) according to their

resampling bootstrap methodology: non-parametric (NPB), parametric (PB) or semi-parametric (SPB). For each replication, they compute the respective expected utilities, finally taking the average values over the 500 replications as the forecast for the next day. The aim here is to determine which resampling methodology should be used by the manager, the methodology resulting in an expected utility closest to the true one.

In a simulation experiment, the true DGP is known, and thus the ‘true’ expected utility is known, namely:

$$EU_{\lambda,0} = w_0' \mu_0 - \lambda w_0' \Sigma_0 w_0,$$

where (μ_0, Σ_0) are the true parameters (inputs) which provide the true weights w_0 . To measure distance between the expected utilities derived from the resampling methods — $EU_{\lambda,NPB}$, $EU_{\lambda,PB}$ and $EU_{\lambda,SPB}$ — and the true $EU_{\lambda,0}$, the absolute differences are computed.

As per a real data example provided in Chapter 2 of Michaud,⁴ and used later by other authors,^{16,17} the study assumes that the data could represent a collection of monthly percentage returns from eight assets (six equity indexes and two bond indexes), over a period of 216 months. For simplicity, the parameter estimates used in these references are taken as the true parameter values, although any other value could indifferently be used.

The experiments run as follows:

- (1) Assume some multivariate distribution F and the true parameter values (μ_0, Σ_0) .
- (2) For $i = 1, \dots, 100$

- (a) Generate a 216×8 data set by independently drawing from the 8-variate distribution assumed in 1. These are the manager’s data at period i , used to define the current allocations (three λ -based portfolios).
- (b) For each bootstrap method (NPB, PB, SPB)
 - (i) Obtain 500 replications of the 216×8 data.
 - (ii) For each sample, compute the sample estimates, which combined with the current λ -based portfolios’ weights, give rise to the corresponding expected utilities.
 - (iii) For fixed λ , return the average expected utility over the 500 simulations, denoted by $EU_{i,\lambda,NPB}$, $EU_{i,\lambda,PB}$ and $EU_{i,\lambda,SPB}$. These are used to compute the absolute distances between the average expected utilities and the true one, that is, the absolute values of $EU_{i,\lambda,NPB} - EU_{\lambda,0}$, $EU_{i,\lambda,PB} - EU_{\lambda,0}$ and $EU_{i,\lambda,SPB} - EU_{\lambda,0}$.

The following multivariate distributions were considered in step 1: normal; t-student with 5 degrees of freedom (df); multivariate distribution with normal univariate margins linked by a t-copula with 5 df; and multivariate distribution with t-student with 5 df univariate margins linked by a Gaussian-copula.

For each fixed λ , Table 1 reports the proportion of winnings for each bootstrap method over the 100 periods. The study also reports the utility functions’ grand mean and standard deviation, and notes that the true

Table 1: Out-of-sample one-step-ahead expected utilities, under different probability distributions and for three risk aversion profiles, $\lambda = 0.5, 1.0$ and 2.0

	$\lambda = 0.5$			$\lambda = 1$			$\lambda = 2$		
	NPB	PB	SPB	NPB	PB	SPB	NPB	PB	SPB
Multivariate normal	36%	17%	47%	39%	14%	47%	42%	12%	46%
EU Gr.Mean	0.01563	0.01404	0.01158	0.01750	0.01733	0.01682	0.01557	0.01399	0.01151
EU Std. Dev.	0.00367	0.00342	0.00293	0.00396	0.00399	0.00404	0.00370	0.00344	0.00293
Multivariate t-student(5)	41%	12%	47%	41%	12%	47%	45%	10%	45%
EU Gr.Mean	0.01629	0.01463	0.01198	0.01814	0.01800	0.01754	0.01615	0.01448	0.01184
EU Std. Dev.	0.00378	0.00351	0.00289	0.00405	0.00412	0.00422	0.00375	0.00350	0.00291
t-copula & Normal margins	40%	11%	49%	42%	6%	52%	48%	2%	50%
EU Gr.Mean	0.01624	0.01502	0.01300	0.01752	0.01746	0.01717	0.01620	0.01498	0.01297
EU Std. Dev.	0.00313	0.00299	0.00265	0.00329	0.00331	0.00340	0.00317	0.00304	0.00269
Gaussian-copula & t-student (5) margins	33%	19%	48%	32%	20%	48%	38%	21%	41%
EU Gr.Mean	0.01568	0.01338	0.01025	0.01857	0.01820	0.01717	0.01549	0.01323	0.01012
EU Std. Dev.	0.00429	0.00386	0.00309	0.00486	0.00487	0.00483	0.00439	0.00392	0.00311
True EU	0.01286	0.01286	0.01286	0.01173	0.01173	0.01173	0.01016	0.01016	0.01016

Numbers in table are the proportions of times the average expected utility under each method (NPB, PB, SPB) was the closest to the true one.
 NPB = non-parametric bootstrap; PB = parametric bootstrap; SPB = semi-parametric bootstrap; EU = expected utility

expected utilities values are $EU_{0.5,0} = 0.0128581$, $EU_{1,0} = 0.0117284$ and $EU_{2,0} = 0.0101549$, which are also given in the last row of Table 1. The figures in Table 1 indicate that the semi-parametric bootstrap resampling methodology was able to provide better forecasts under all distributional assumptions.

In summary, the experiments indicate the accuracy of the long run forecast of an investor that consistently assumes some resampling methodology and forecasts his expected utility according to it. Note that most financial institutions assume the multivariate normal distribution for parametrically resampling the data, and to this end, the paper recommends looking to the data marginal and joint structures.

CONCLUDING REMARKS

This paper proposed a semi-parametric bootstrap method for replicating multivariate data, a method which parametrically takes care of all marginal characteristics of the returns series, including their temporal dependencies, and also captures the correlations linking the data in a very simple non-parametric way. The method's flexibility allows one to go beyond the assumption of any particular fixed margins multivariate distribution. This is very important in finance, as it is well known that returns on financial assets and indexes possess specific characteristics following different conditional and unconditional probability distributions. Therefore, no multivariate distribution with fixed margins would successfully model such data sets.

The replication scheme requires the same computational effort required by a parametric bootstrap. It can be implemented in any accounting oriented

software where an unconditional distribution is fitted to the margins, or in any statistical software should time series models be chosen for the univariate series. The two-step semi-parametric bootstrap may be quickly applied to any (high) dimensional data, not suffering from the well-known 'curse of dimensionality'.

The main reason for obtaining replications of a data set is to assess estimation risk. This means obtaining the variability of selected statistics, for example, portfolios' optimal weights or the value-at-risk. When constructing efficient frontiers, the resulting portfolios' weight distributions may be used to identify similar portfolios. Statistical tests may be conducted to assess how far one needs to be away from the original set of weights to obtain a statistically different portfolio. This is measured with the concept of distance applied to the resampled weights. The illustration provided showed that the new method has the potential to provide more stable weight distributions for any given portfolio.

Simulations were carried out to assess the out-of-sample performance of the semi-parametric bootstrap method when forecasting three portfolios' expected utilities. The portfolios reflect the risk preferences of three clients. The experiments indicated that the semi-parametric resampling methodology provided better expected utility forecasts in the long run, as measured by their distances to the true one, under all distributional assumptions made.

References

- 1 Markowitz, H. M. (1959) 'Portfolio Selection: Efficient Diversification of Investments', J. Wiley, New York, NY.

- 2 Efron, B. and Tibshirani, R. J. (1998) 'An Introduction to the Bootstrap', Chapman and Hall, CRC, Boca Raton, FL.
- 3 Jorion, P. (1992) 'Portfolio optimization in practice', *Financial Analysts Journal*, Vol. 48, No. 1, pp. 68–74.
- 4 Michaud, R. (1998) 'Efficient Asset Management: A Practical Guide to Stock Portfolio Optimization and Asset Allocation', Harvard Business School Press, Boston, MA.
- 5 Scherer, B. and Martin, R. D. (2005) 'Introduction to Modern Optimization with Nlopt and SPlus', Springer, New York.
- 6 Carlstein, E. (1986) 'The use of subseries values for estimating the variance of a general statistics from a stationary sequence', *Annals of Statistics*, Vol. 14, No. 3, pp. 1171–1195.
- 7 Künch, H. R. (1989) 'The jackknife and the bootstrap for general stationary observations', *Annals of Statistics*, Vol. 17, No. 3, pp. 1217–1241.
- 8 Jobson, J. D. (1991) 'Confidence regions for the mean-variance efficient set: An alternative approach to estimation risk', *Review of Quantitative Finance and Accounting*, Vol. 1, No. 3, pp. 235–257.
- 9 Britten-Jones, M. (1999) 'The sampling error in estimates of mean-variance efficient portfolio weights', *The Journal of Finance*, Vol. 2, No. 2, pp. 655–655.
- 10 Fusai, G. and Meucci, A. (2003) 'Assessing views', *Risk*, March, pp. 18–21.
- 11 Michaud, R. and Michaud, R. (2008) 'Estimation error and portfolio optimization: A resampling solution', *Journal of Investment Management*, Vol. 6, No. 1, pp. 8–28.
- 12 Hansen, B. (1994) 'Autoregressive conditional density estimation', *International Economic Review*, Vol. 35, No. 3, pp. 705–730.
- 13 Genest, C. and Favre, C. (2007) 'Everything you always wanted to know about copula modeling but were afraid to ask', *Journal of Hydrologic Engineering*, Vol. 12, No. 4, pp. 347–368.
- 14 Oakes, D. (1982) 'A model for association in bivariate survival data', *Journal of the Royal Statistical Society Series B (Statistical Methodology)*, Vol. 44, No. 3, pp. 414–422.
- 15 Mendes, B. V. M. (2008) 'A conditional approach for risk estimation', *Journal of Risk*, Vol. 11, No. 1, pp. 1–23.
- 16 Markowitz, H. M. and Usman, N. (2003) 'Resampled frontiers versus diffuse Bayes: An experiment', *Journal of Investment Management*, Vol. 1, No. 4, pp. 9–25.
- 17 Harvey, C. R., Liechty, J. and Liechty, M. W. (2008) 'Bayes vs resampling: A rematch', *Journal of Investment Management*, Vol. 6, No. 1, pp. 1–17, available at: <http://ssrn.com/abstract=898241>.

Copyright of Journal of Risk Management in Financial Institutions is the property of Henry Stewart Publications LLP and its content may not be copied or emailed to multiple sites or posted to a listserv without the copyright holder's express written permission. However, users may print, download, or email articles for individual use.