



Detecting performance persistence of hedge funds



Rania Hentati-Kaffel*, Philippe de Peretti

Centre d'Economie de la Sorbonne, Université Paris 1 Panthéon-Sorbonne, Maison des Sciences Economiques, 106-112 Boulevard de l'Hôpital, 75013 Paris, France

ARTICLE INFO

Article history:
Accepted 14 August 2014
Available online xxxx

Keywords:
Hedge funds
Runs tests
Persistence
Clustering

ABSTRACT

In this paper, we use nonparametric runs-based tests to analyze the randomness and the persistence of relative returns of hedge funds. Runs tests are implemented on a universe of hedge funds extracted from HFR database over the period spanning January 2000 to December 2012. Our findings suggest that i) slightly less than 80% of the studied universe has returns at random, ii) a similar figure is found out when focusing on relative returns, iii) hedge funds that do present clustering in their relative returns are mainly found within Event Driven and Relative Value strategies, iv) and for relative returns, results vary with the type of the benchmark nature (peer group average or traditional). This paper also emphasizes that runs tests may be a useful tool for investors in their fund's selection process.

© 2015 Elsevier B.V. All rights reserved.

1. Introduction

The hedge fund industry has long been, naively, seen as being able to generate “all weather” positive returns, no matter what the market conditions were. Nevertheless, the recent financial crisis has cast some doubts on this opinion, leading investors to question whether this industry was significantly able to over-perform the traditional management (Gupta et al., 2003). The question of over-performances, or equivalently of the persistence of relative returns, is of key importance for investors. Indeed, assessing persistence is a milestone in the decision making process. For instance, one of the main strategies used by investors, e.g. funds of hedge funds strategy, to pick-up top hedge funds, relies on realized relative returns (versus HFR representative strategy index or traditional indices) momentum. Thus, selecting a hedge fund for its ability to significantly over-perform the market during large periods may be a very useful tool.

Persistence has been studied by many authors using various methodologies¹ as the Cross-Product Ratio (CPR) (De Souza and Gokcan, 2004), Chi-square tests (Carpenter and Lynch, 1999), regression models (Agarwal and Naik, 2000a; Fama and MacBeth, 1973), or the test of Hurst (Amenc et al., 2003; De Souza and Gokcan, 2004; Edwards and Caglayan, 2001; Eling, 2008). Clearly, three conclusions are to be drawn: i) Results vary with both the database (TASS, HFR, Tremont, ...)

and the methods ii) Most studies agree to find a persistence for a one to a six-month horizon (short-term) (Barès et al., 2003; Boyson and Cooper, 2004; Brorsen and Harri, 2004; Herzberg and Mozes, 2003), but results are contradictory for longer periods, iii) There is no agreement whether the persistence is related to the nature of the strategy of the hedge fund.

The goal of this paper is to re-examine the questions of persistence, and randomness of returns for a given hedge fund relatively to a set of indices. For both analyses, we use the HFR data base, with a universe spanning more than 4000 hedge over the period spanning January 2000 to December 2012. Relative returns are computed using a blend of traditional and alternative indices: i) the median of the returns of funds having a common primary strategy, ii) an HFRI index computed for each primary strategy, iii) an overall index for the hedge fund market, and iv) the S&P500 index. Performances of hedge funds are thus analyzed with regard to peer groups, the whole hedge fund universe, and an external market.

To extract information about randomness and persistence, we use tests based on runs (Gibbons and Chakraborti, 1992; Wald and Wolfowitz, 1940). Runs tests are very versatile and powerful tools. Used as two-sided tests, they allow to check for randomness. Used as a one-sided test, they allow to test for randomness against a pre-specified alternative: Either clustering, i.e. persistence, implying the ability for a fund to significantly over (under)-perform a given market, or mixing, i.e. systematically alternating over and under performances.

Our main findings suggest that i) slightly less than 80% of the studied universe has returns at random, ii) a similar outcome is obtained when relative returns are used, iii) hedge fund strategies displaying the highest percentage of funds generating clusters are Event-Driven and Relative Value, emphasizing the link between the strategy and the

* Corresponding author.

E-mail addresses: rania.kaffel@univ-paris1.fr (R. Hentati-Kaffel), philippe.de-peretti@univ-paris1.fr (P. de Peretti).

¹ See also Edwards and Caglayan (2001), Harri and Brorsen (2004), Brown et al. (1999), Kat and Menexe (2003), Koh, Koh and Teo (2003), Baquero et al. (2004), Kouwenberg (2003), Jagannathan et al. (2006).

persistence, and iv) for the relative returns, results deeply vary with the benchmark.

This paper is organized as follows. In Section 2, we details how our series are computed, and introduce runs-based tests. An empirical application is also presented. In Section 3, we implement the tests on the HFR database, and present results in contingency tables crossing the results on runs tests with strategies. Results are also provided when we split our sample into two sub-samples, before and after the 2007 crisis, and re-run the tests. On Section 4, we use Monte Carlo simulation to show the robustness of our approach even if ARCH effects are present and/or there are structural breaks in the persistence. Finally Section 5 discusses our main results and concludes.

2. Runs-based tests

Let $\{r_{it}^j\}_{t=1}^T$, be an observed track record of T observations of returns for fund i having a main strategy $j, j \in (1, 4)$, where $j = 1$ corresponds to Equity Hedge, $j = 2$ to Event-Driven, $j = 3$ to Macro, and $j = 4$ to Relative Value.

Now, define $\{d_{it}^j\}_{t=1}^T, j \in (1, 4)$ as follows:

$$d_{it}^j = \begin{cases} 1 & \text{if } r_{it}^j \geq b_{it}^j, \\ 0 & \text{otherwise.} \end{cases} \quad (1)$$

where b_{it}^j is either defined as:

$$b_{it}^j = b_i^j = \text{median}(r_{it}^j, r_{it}^j, \dots, r_{it}^j), j \in (1, 4). \quad (2)$$

or:

$$b_{it}^j = b_t^j = \text{median}(r_{it}^j, r_{it}^j, \dots, r_{it}^j), j \in (1, 4). \quad (3)$$

$$b_{it}^j = b_t^j = \text{HFRI}_t^j, j \in (1, 4). \quad (4)$$

$$b_{it}^j = b_t = \text{HFRGI}_t. \quad (5)$$

$$b_{it}^j = b_t = \text{SP500}_t. \quad (6)$$

where: $r_{it}^j, r_{it}^j, \dots, r_{it}^j$ are the returns of funds having a common main strategy j , HFRI_t^j is a return computed using a performance index corresponding to the primary strategy j , HFRGI_t is a return computed using the HFRI global performance index at time, SP500_t is a return computed using the S&P500 index at time t .

Remark 1. Definition (2) allows us to analyze the randomness of the series, whereas definitions of b_{it}^j given by (3) to (6) return an information about the relative performance of the fund, i.e. the possible persistence of the returns with regard to a benchmark, indicated by large clusters of 1s or 0s. Using (3) to compute b_{it}^j returns a straightforward information about the location of the return of the fund i in the distribution of the returns of a main strategy, i.e. if the returns are located in the right (left) tail of the distribution during large periods of times, or is randomly distributed on the right or left tail.

Remark 2. In our opinion, the definition of a skilled manager should emphasize its ability to outperform its peers (representative HFRI hedge index), as well as the overall sample (HFRI Global Hedge index). Comparing performance to the overall sample attenuates the selection bias effects (Databases have their own classification criteria which could differ from one provider to another). Thus, the second and third benchmarks, (3) and (4) are used to study how a fund performs compared to its peers (funds classified in the same class), whereas the fourth, (5), is used to study the relative performance of the fund with regard to whole hedge fund sample. The last benchmark,

(6), is used as an external reference, to see if funds are able to outperform traditional market (equity in our case).

We next use runs-based tests to analyze the information returned by the d_{it}^j s. Define a run of one kind of element, say of 1's, as a successions of 1's immediately preceded or followed by at least one 0, or nothing. Let T_1 be the number of 1's and T_0 be the 0's with $T_1 + T_0 = T$, and let r_{1j} be the number of runs of 1's of length j and r_{0j} be the number of runs of 0's of length j . Let $r_1 = \sum_j j r_{1j}$ be the total number of runs of 1's, and $r_0 = \sum_j j r_{0j}$ the total number of runs of 0's. At last let $r = r_1 + r_0$ be the total number of runs of both kinds.

Testing for randomness amounts to testing if we have either too few runs or too many runs by using a two-sided test, whereas testing for the null of randomness against the alternative of clustering i.e. persistence, amounts to using a one-sided test (focusing on the left tail of the distribution), testing for too low values of r_1 (or r).

Following Gibbons and Chakraborti (1992), exact and approximate distributions can be used to test for the null. Concerning the former, using combinatorial, the marginal (exact) distribution function of r_1 is given by:

$$P(r_1) = \frac{\binom{T_1-1}{r_1-1} \binom{T_0+1}{r_0}}{\binom{T}{T_1}} \quad (7)$$

where $\binom{T_1-1}{r_1-1}$ is a binomial coefficient.

Similarly, the (exact) distribution function of r is given by:

$$P(r) = \begin{cases} 2 \frac{\binom{T_1-1}{\frac{1}{2}r-1} \binom{T_0-1}{\frac{1}{2}r-1}}{\binom{T}{T_1}} & \text{if } r \text{ is even,} \\ \frac{\binom{T_1-1}{\frac{1}{2}r-1} \binom{T_0-1}{\frac{1}{2}r-2}}{\binom{T}{T_1}} + \frac{\binom{T_1-1}{\frac{1}{2}r-2} \binom{T_0-1}{\frac{1}{2}r-1}}{\binom{T}{T_1}} & \text{if } r \text{ is odd.} \end{cases} \quad (8)$$

Among many other, Gibbons and Chakraborti (1992) provide tabulations for small values of T_0 and T_1 , i.e. for $T_0 \leq T_1 \leq 12$, such that (7) and (8) can be used to build one or two-sided tests.

For 'large' values of T_1 and T_0 , i.e. for $T_0 > 12$ and $T_1 > 12$ a normal approximation can be used. Define the first two moments of r_1 and r as:

$$E[r_1] = \frac{T_1(T_0+1)}{T} \quad (9)$$

$$V[r_1] = \frac{(T_0+1)^2(T_1)^2}{T(T)^2} \quad (10)$$

where $x^{[a]} = x(x-1)(x-2) \dots (x-a+1)$,

$$E[r] = E[r_1] + E[r_0] = \frac{2T_1T_0}{T} + 1 \quad (11)$$

$$V[r] = V[r_1] + V[r_0] + 2\text{cov}[r_1, r_0] = \frac{2T_1T_0(2T_1T_0-T)}{T^2(T-1)} \quad (12)$$

Then, using a continuity correction, the corresponding Z -stats are defined as:

$$Z_{r_1} = \frac{r_1 + 0.5 - T_1(T_0+1)T^{-1}}{\sqrt{\frac{(T_0+1)^2(T_1)^2}{T(T)^2}}} \quad (13)$$

Download English Version:

<https://daneshyari.com/en/article/5053911>

Download Persian Version:

<https://daneshyari.com/article/5053911>

[Daneshyari.com](https://daneshyari.com)