



SCHOOL OF ECONOMICS AND MANAGEMENT
FACULTY OF SOCIAL SCIENCES
AARHUS UNIVERSITY



CREATES
Center for Research in Econometric
Analysis of Time Series

CREATES Research Paper 2009-24

A Meta-Distribution for Non-Stationary Samples

Dominique Guégan

School of Economics and Management
Aarhus University
Bartholins Allé 10, Building 1322, DK-8000 Aarhus C
Denmark

A Meta-Distribution for Non-Stationary Samples

Dominique Guégan*

June 1, 2009

Abstract

In this paper, we focus on the building of an invariant distribution function associated to a non-stationary sample. After discussing some specific problems encountered by non-stationarity inside samples like the "spurious" long memory effect, we build a sequence of stationary processes permitting to define the concept of meta-distribution for a given non-stationary sample. We use this new approach to discuss some interesting econometric issues in a non-stationary setting, namely forecasting and risk management strategy.

Keywords: Non-Stationarity - Copula - Long-memory - Switching - Cumulants - Estimation. theory.

JEL classification: C32, C51, G12

1 Introduction

During decades time series modelling has focused on stationary linear models, then on stationary non-linear models. Recently the question of non-stationarity has arisen, and a sudden interest focuses on the modelling of non-stationary and/or non-linear time series.

*PSE, Centre d'Economie de la Sorbonne, University Paris1 Panthéon-Sorbonne, MSE, 106 bd de l'Hôpital, 75013 Paris, France. Email: dguegan@univ-paris1.fr, Tel: +33 1 40 07 82 98.

An interesting type of questions has emerged when it has been observed that stationary non-linear processes exhibited empirical autocorrelation function with a hyperbolic decreasing rate, although they are characterized by a theoretical short memory behavior (exponential decreasing of the autocorrelation function). This constat has highlighted the fact that the behavior of some statistical tools under non-stationarity must be questioned. Indeed, the statistical tools we are using are meaningful only under certain assumptions, the most crucial one being the stationarity. Hence, the question arises what the statistical tools are telling us when used on non-stationary data.

Therefore, it has appeared that modelling data sets need to adress almost two major questions. Is the underlying process stationary and do we use linear or non-linear modellings? The order in which these questions are adresssed is essential. Indeed, answering to the second question, need before establishing stationarity.

Thus, defining a correct framework in which we can analyse data sets is fundamental before chosing the class of models we will use. First, the nature of asset prices behavior is necessary to get robust forecasts. Second the knowledge of the probabilistic properties of these asset prices is fundamental to the formulation of the concept of risks: indeed the measurement of risks depends heavily on properties of the empirical distribution such as stationarity, long tailedness, finiteness of the second and higher order moments. Third, various tests for the empirical validity of financial models and the application of these models rely on the robustness of statistical tools which can be deficient in specific context. Fourth, several important pricing models for stock options and other similar financial instruments usually require explicit estimates of stock return variances. The usefulness of these models depends largely on the adequacy and the stationarity of almost the second order moments. Indeed, to model real data sets using classical stochastic processes imposes that the data sets verify almost the second order stationarity condition.

We propose here a new methodology considering that it is not always possible to remove the non-stationarity observed in a data set. We will also show how different types of non-stationary corrupt the behavior of statistical tools making confusion in the modelling analysis.

Notion of stationarity includes the notion of weak and strict stationarity. In most theoretical results (asymptotic theory), we need the process $(Y_t)_t$ strictly stationary which means that it is characterized by an invariant measure (the moments do not depend on time t). Thereby, the stochastic process

$(Y_t)_t$ whose a representation can be, $\forall t \ Y_t = f(Y_{t-1}, \dots) + \varepsilon_t$, has an invariant measure which is linked to the invariant measure of the noise $(\varepsilon_t)_t$. In case of a Gaussian noise and f linear then the distribution function of $(Y_t)_t$ is Gaussian; saying that it is invariant means that the mean and the variance are finite and constant (Gaussian distributions are characterized by only the two first moments).

Since the eighties', several non-linear models have been introduced and extensively investigated: the SETAR, EXPAR, Bilinear, related GARCH, Markov switching, breaks and jumps models, among others. All the theory which concerns these models (existence of a solution, inference, etc.) assumes that they are (strictly) stationary; the theory is based on stationary unconditional moments and then on the existence of an invariant measure. We can note that for nearly all these models this invariant measure is unknown. For instance for GARCH models it is the conditional law which is known as soon as the residuals distribution is known, and not the unconditional law.

The interest for non stationarity is now a long tradition inside the econometrician community. It has been developed around the theory of I(d) processes through extended Dickey Fuller tests. Recently interesting works have been published assuming under the alternatives non-linear models. A limitation of this approach is due to the fact that the specification under the alternative generally concerns one non-linear feature, due to the complexity of the testing method. More, to obtain theoretical power for these tests is difficult, and often done under specific assumptions, Guégan and Pham (1992), Terasvirta (2003). In another hand, this approach privileges existence of polynomial trends, which can be interesting to study some economic data sets, but which is not the more important feature observed in financial data sets, Bec, Ben Salem and Carrasco (2004), and references therein. Our approach does not follow the same methodology (testing the non stationarity) and could be considered as a alternative as hers.

When we observe, some financial data sets given by a sample $(Y_1, \dots, Y_T)$, a certain number of features characterize these data like the existence of clusters, explosions, pseudo-seasonalities, existence of several states. These features are the origin of the non-invariant property of the distribution function associated to the underlying process, and can provoke misspecification. Indeed, most of these characteristics make the first four moments of the process $(Y_t)_t$ depending on time. Thus, we are interested to analyse the behavior of these empirical moments in presence of non stationarity. Then, we will build sequences of intervals on which these sample moments do not evolve

with time, in order to define a measure, with invariant properties, permitting to characterize the whole sample.

The problem of misspecification is crucial in the modelling of time series because it can be responsible of a lot of disasters if the predictions are false or if the risks associated with this modelling are underestimated for instance. An interesting problem has been raised with the famous example of "spurious" long memory behavior detected in some data sets. We comment this fact now.

A fundamental tool to analyse the time series is the autocovariance function. In many cases this tool has shown its limitation. For instance, we can use it to determine the orders of the autoregressive and moving average parts of linear models, Brockwell and Davis (1988), but this identification procedure fails as soon as we extend the class to the non-linear models: indeed the autocovariance function of autoregressive models can be similar to the one of bilinear or Markov switching processes, Granger and Andersen (1978), Guégan (1987), and Poskitt (1996), and then misspecification arises.

Another interesting point is the following. Many models are known to be short memory under stationarity conditions: they are the ARMA, GARCH, Markov switching, Bilinear, Stop break models, among others. These models, under stationary conditions have a theoretical autocovariance function $\gamma_Y(t, h) = cov(Y_t, Y_{t+h})$ which decreases towards zero with an exponential rate. In another hand, there exists long memory processes including FARMA, GARMA, FIEGARCH and GIGARCH models whose autocorrelation function decreases with an hyperbolic rate towards zero. Thus, to discriminate between these two classes of models a natural tool is the autocovariance function, and observing a data set $(Y_1, \dots, Y_T)$, we compute the sample autocovariance function

$$\hat{\gamma}_Y(h) = \frac{1}{T} \sum_{t=1}^{T-h} (Y_t - \bar{Y})(Y_{t+h} - \bar{Y}),$$

where T is the sample size, and $\bar{Y}$ is the sample mean. We know that the theory concerning stochastic stationary process lies on $\gamma_Y(t, h)$, assuming that this statistic does not depend on time t . Now, if the underlying process is non-stationary, then $\gamma_Y(t, h)$ and $\hat{\gamma}_Y(h)$ are different concepts, and confusion arises.

So, recent empirical studies have shown that stochastic processes - which exhibit an exponential decreasing for $\gamma_Y(t, h)$ - exhibit a sample autocorrela-

tion function $\tilde{\gamma}_Y(h)$ which has not this property, and the notion of "spurious" long memory appears. In the previous paper the author (Guégan (2005), and references therein) investigates this fact in details, exhibiting a lot of examples, discussing different kinds of long memory behaviors and raising many questions on this purpose. The objective of this paper is to answer few questions.

In the theory all the results are obtained assuming the existence of an invariant measure, and it appears that this property fails for specific samples characterized by particular features. We will analyse in more details this phenomenon, extending the discussion introduced in two interesting papers, Mikosch and Starica (2004) and Starica and Granger (2005). We will see that non-stationarity inside samples create distortions for the sample autocovariance function, and we will propose another way to avoid the misspecification coming from this distortion, considering a local processing for the data. Thus, another working framework is proposed based on a local methodology.

The local approach that we propose consists in building "homogeneity" intervals (we use the terminology introduced by Starica and Granger (2005)), in which the data are stationary up to the first four moments. This means that inside these intervals, the data are characterized by a mean, a variance, a skewness and a kurtosis which do not depend on time t . In practice, these first four moments are important and permit to characterize most of the features of the financial data sets. Then, we will use this sequence of "homogeneity" intervals, on which processes are characterized by approximately invariant measure (up to order four), to build a new distribution function for the whole sample, that we called meta-distribution, based on the copula's concept, Nielsen (1999). Thank to this new flexible approach, we discuss new methodologies to do forecast, and risks management strategy.

Our plan is the following. In section two we discuss the main features observed inside financial data sets which lead to different kinds of non-stationarity in samples. In Section three, we specify the behavior of the sample autocovariance function in presence of non-stationarity in mean or in volatility. We exhibit examples. Section four is devoted to the construction of "homogeneity" intervals; it permits to provide the new concept of meta-distribution and develop its interest for several topics. Section five concludes.

2 Non-stationary stylized facts

In this Section, we analyse some stylized facts observed in financial data sets which can provoke non-stationarity. We focus on structural behaviors like volatility, jumps, explosions and seasonality, and also on specific transformations like concatenation, aggregation or distortion which are also at the origin of non-stationarity.

Indeed, the stationary conditions are essential to guarantee the asymptotic properties of the sample mean, variance and covariances and are different for each of these estimators. In a non-linear setting, ergodicity is necessary, it requires that observations sufficiently far apart should be almost uncorrelated. Under all these conditions, the process is globally stationary: this means that this stationary property remains true on the whole sample. In all cases this means that inference can be done for such processes and that the asymptotic theory works. In particular forecasts are available and confidence intervals can be provided. Thus, it appears necessary that, in practice these conditions are verified in order to apply the theory developed in that context.

Many features imply that the property of global stationarity fails. Indeed, existence of volatility imposes that the variance depends on time. In presence of seasonality the covariance depends on time. The existence of states induces changes in mean or in variance all along the trajectory. Concatenated or distorted models cannot have the same probability distribution function on the whole period. Aggregation is a source of specific features. For some of the previous situations the higher order moments do not exist implying that the autocovariance function is not defined. Thus, we need - in terms of modelling - to specify the framework in which we are going to work. More precisely:

- Existence of volatility characterized by clusters provokes important changes around the mean price that we model. A famous way to model it is to use conditional properties through the ARCH model (Engle, 1982), and related extensions like for instance the APARCH model introduced by Ding and Granger (1996), using any power of the conditional variance permitting to introduce another kind of non-linearity inside the modelling. Using related GARCH modellings means that we work with a conditional approach, but often the observations exhibit non-stationarity in the non-conditional variance which affects the invariance of the global distribution function.
- A strong cyclical component inside financial data (monthly for in-

stance or hourly for high frequency data) produces evidence of non-stationarity, which is not always removed by transformations. Presence of seasonals indicate that there exist significant correlations between random variables at present t , and in the past or in the future. This correlation creates dependence in each subsequence, which implies non-stationarity on the whole sample. To take into account this feature can be difficult mainly when the seasonals are not fixed. An approach is proposed using the k -factor Gegenbauer processes, Gray, Zhang and Woodward (1989), Guégan (2001). In that case it is the long memory behavior which is privileged more than the seasonals, and attention need to be done to the possibly remaining seasonals.

- Specific shocks can produce jumps, and then create different regimes inside data sets, in means or in variance. These changes affect the invariance property of the distribution function of the underlying process. These behaviors can be modelled in part through Markov switching models, SETAR or STAR processes, Stop-Break or sign processes. The stationarity of these models is questionable, and generally we do not pay attention of the unconditional distribution of these models.
- Distorsion effects are characterized by explosions that cannot be removed from any transformation. Explosions imply that some higher order moments of the distribution function do not exist. Models with coefficients close to the non stationary domain can also create this kind of effect. In that case the global distribution function of the process loses its invariant property. The juxtaposition of different stationary linear or non linear processes creates non-stationarity by building. Finally, to aggregate independent or weakly dependent random variables can create specific dependence and also non-stationarities, Robinson (1978) and Granger (1980).

In order to understand the impact of these specific features in a sample, we focus now on the study of the sample autocovariance function when non-stationarity is detected in mean or in variance.

3 Sample autocovariance behavior in presence of non-stationarity

In this section we highlight that $\gamma_Y(t, h)$ and $\tilde{\gamma}_Y(h)$ are different concepts in presence of non-stationary. We assume that we observe a non-stationary

sample $(Y_1, \dots, Y_T)$, corresponding to an underlying process $(Y_t)_t$. We consider the sample ACF in situations when structural breaks occur for instance. The sample autocovariance function is equal to:

$$\tilde{\gamma}_Y(h) = \frac{1}{T} \sum_{t=1}^{T-h} (Y_t - \bar{Y})(Y_{t+h} - \bar{Y}),$$

where $\bar{Y}$ is the sample size. We divide the previous sample in r subsamples consisting each of distinct stationary ergodic processes with finite second order moments. We denote $p_j \in R^+$, $j = 1, \dots, r$ such that $p_1 + p_2 + \dots + p_r = 1$. Here p_j is the proportion of observations from the j th subsample in the full sample. If we define now $q_j = p_1 + p_2 + \dots + p_j$, $j = 1, \dots, r$, then the whole sample is written as: $Y = ((Y_1^{(1)}, \dots, Y_{T_{q_1}}^{(1)}), (Y_{T_{q_1}+1}^{(2)}, \dots, Y_{T_{q_2}}^{(2)}), \dots, (Y_{T_{q_{r-1}}+1}^{(r)}, \dots, Y_{T_{q_r}}^{(r)}))$, $T_{q_r} = T$ and in the following we denote $Y^{(i)} = (Y_1^{(i)}, \dots, Y_{T_{q_i}}^{(i)})$, $i = 1, \dots, r$. Thus, the i subsamples come from distinct stationary ergodic models with finite second moment, and then the resulting sample is not stationary. Now, to get $\tilde{\gamma}_Y(h)$, we compute the sample autocovariance on each subsample and then sum it:

Proposition 3.1 *Let be r subsamples $Y_1^{(i)}, \dots, Y_{T_{q_r}}^{(i)}$, $i = 1, \dots, r$, coming from the sample Y , each subsample corresponding to a distinct stationary ergodic process with finite second order moments, whose sample covariance is equal to $\tilde{\gamma}_{Y^{(i)}}(h)$, then*

$$\tilde{\gamma}_Y(h) \rightarrow \sum_{i=1}^r p_i \tilde{\gamma}_{Y^{(i)}}(h) + \sum_{1 \leq i \leq j \leq r} p_i p_j (EY^{(j)} - EY^{(i)})^2, \quad h \rightarrow \infty. \quad (1)$$

If the expectations of the subsequences $(Y^{(i)})_i$ differ, and because the autocovariances $\tilde{\gamma}_{Y^{(i)}}(h)$ decay to zero exponentially as $h \rightarrow \infty$ (due to the ergodic property of the subsequences), the sample ACF $\tilde{\gamma}_Y(h)$ for sufficiently large h is close to a strictly positive constant given by the second term of the equation (1). Indeed, the term $\sum_{1 \leq i \leq j \leq r} p_i p_j (EY^{(j)} - EY^{(i)})^2$ dominates and determines the behavior of $\tilde{\gamma}_Y(h)$.

Since a long memory process is exclusively determined by the decay pattern of its autocovariance, then we say that the process $(Y_t)_t$ has a long memory behavior. It is this fact ("spurious" long memory) which has been observed for different simple models with structural breaks or switches in the mean. In that case, the sample autocovariance of these models $\tilde{\gamma}_Y(h)$ does not converge towards zero while the theoretical autocovariance function $\gamma_Y(h)$ does.

The latter one is computed under strong stationarity conditions (existence of an invariant measure) and the former one from a sample which is affected by non-stationarity (the empirical distribution function is not invariant).

It is exactly this behavior which is observed when we simulate stationary models like the Markov switching, the stopbreak model, or the SETAR models, for a complete review Guégan (2005). In order to illustrate the proposition 3.1, we simulate a particular case of model with switches, introduced by Granger and Terasvirta (1999). It has the following form :

$$Y_t = \mu_1 I(Y_{t-1} > 0) + \mu_2 I(Y_{t-1} \leq 0) + \varepsilon_t, \quad (2)$$

where $I(\cdot)$ is the indicator function. This model permits to shift from the mean μ_1 to the mean μ_2 with respect to the value taken by Y_{t-1} . SETAR processes are known to be short memory, but it is also possible to exhibit sample ACFs which present slow decay, and this slow decay can also be explained by the second term of the relationship (1). We provide the trajectory and the autocorrelation function of this model on figure 1. We observe switches on the trajectory and slow decay of the sample ACF, eventhough this model is classified as short memory process for the values of parameters used here.

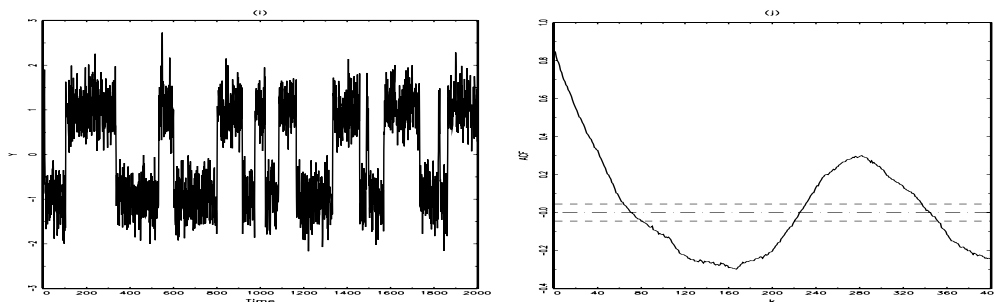


Figure 1: Trajectory and ACF of the Threshold Auto-Regressive model defined by equation (2) with $T = 2000$, $\sigma_\varepsilon^2 = 0.2$ and $\mu_0 = -\mu_1 = -1$.

In their paper, Granger and Terasvirta (1999), discuss this "spurious" long memory observed for specific samples stemming from the size and the persistence of the shifts in their mean. It corresponds to a specific time spend in each state whose distribution is not known.

Now if we are interested by volatility modelling, we prefer to analyse the sample $Y^\delta = (Y_1^\delta, Y_2^\delta, \dots, Y_T^\delta)$ instead of Y . Then, we build sequences of intervals such that $Y^\delta = ((Y_1^{\delta(1)}, \dots, Y_{T_{q_1}}^{\delta(1)}), \dots, (Y_{T_{q_{r-1}}+1}^{\delta(r)}, \dots, Y_{T_{q_r}}^{\delta(r)}))$, assuming existence of distinct stationary ergodic process on each of them, then:

Lemma 3.2 *Let be r subsamples $(Y_1^{\delta(i)}), \dots, (Y_{T_{q_i}}^{\delta(i)})$, $i = 1, \dots, r$ and $\delta \in R^+$, coming from the sample Y^δ , each subsample corresponding to a stationary distinct ergodic process with finite second order moments, whose sample covariance is equal to $\tilde{\gamma}_{(Y^{\delta(i)})}(h)$, then the sample autocorrelation function $\tilde{\gamma}_{Y^\delta}$ of the sample Y^δ is such that:*

$$\tilde{\gamma}_{Y^\delta}(h) \rightarrow \sum_{i=1}^r p_i \tilde{\gamma}_{(Y^{\delta(i)})}(h) + \sum_{1 \leq i \leq j \leq r} p_i p_j (E(Y^{\delta(j)}) - E(Y^{\delta(i)}))^2, \quad h \rightarrow \infty. \quad (3)$$

Under the property of stationary ergodicity, the ACF of each process has an exponential decay. Thus, the sample Y^δ has its sample ACF $\tilde{\gamma}_{Y^\delta}(h)$ that decays quickly for the first lags and then approach positive constants given by $\sum_{1 \leq i \leq j \leq r} p_i p_j (E(Y^{(j)})^\delta - E(Y^{(i)})^\delta)^2$.

This last term explains some long memory effect observed on the ACF of the series when we analyse them with the sample Y^δ . This last term shows how shifts in the variances (modelled using Y^δ could explain long memory effect inside the sample ACF. Mikosch and Starica (2004) have already illustrate this kind of behavior using Y^2 and $|Y|$ samples for modelling the variance of the log returns. Here, we extend their approach in a more general setting. Great lines for the proof of the proposition 3.1, which can be easily extended for this lemma is postponed at the end of the paper.

In summary, it appears that samples exhibiting shifts in the mean will produce kind of long memory behavior which can be explained by the proposition 3.1; if the shifts appear in the variance, the lemma 3.2 can explain some specific persistence in the data set. Presence of seasonals not removed by filtering creates some kind of long memory and this behavior is explained by the first term in (1). Finally in presence of distorsion or explosions the variance being not defined the use of covariance function has no sense. These previous results clarify some mis-understandings concerning the existence of long memory behavior regarding at the sample autocovariance function.

Thus, these simple evidences show that it is not possible to work in empirics using directly the theory. Indeed, it appears fundamental to change our working assumption, questioning the stationary assumption in practice, coming back to the assertion of Mandelbrot (1963): "price records do not "look" stationary, and statistical expressions such as the sample variance take very different values at different times: this non-stationarity seems to put a precise statistical model of price change out of the question."

4 Local stationarity and meta-distribution

In the previous section, we show that the existence of non-stationarity inside data sets pollute the behaviors of some statistics which are popular in modelling, namely the sample ACF. Thus, it appears interesting to look after the local properties of a data set and to use them in order to define new strategies in risk management and forecasting. Here, we are interested to approximate the non-stationarity data locally by stationary models. Thus, we build a sequence of models with invariant distributions and we propose a meta-distribution characterizing the whole sample.

4.1 A new test

Statistical analysis of stock prices variations Y mainly focused on the second order properties (mean, variance, covariance), but a lot of features can affect also higher order moments, and mainly the moments of order three and four which characterize the shewness and the kurtosis of these data sets. Thus, it seems natural to consider Y as a locally stationary time series: this means that we can define time intervals of varying size which are "approximately stationary". This idea is to build processes whose stationarity is based on the behavior of the first four orders moments. Hence a qualitative characterization of such locally stationary processes could be: on each interval $[t_1, t_2]$, the cumulants c_k (up to order four, for instance) compute using observations belonging to this interval may be well approximated by a function depending only on $t_2 - t_1$ as soon as these time points are close enough:

$$c_k(t_1, t_2) \sim c_k(t_2 - t_1) \text{ if } |t_2 - t_1| < l(k)/2,$$

and becomes a deterministic function of time t when the length between the time points considered is larger than a certain threshold $d(k)$, measuring somehow the "stationary rate" of the time series

$$c_k(t_1, t_2) \sim f(t) \text{ if } |t_2 - t_1| > d(k)/2.$$

Following this idea, we introduce a test which permits to carry out the local stationary behavior of a data set based on the behavior of the first four sample moments. We compute the cumulants associated to the sample up to the order k , and the spectral density of cumulant of order k , denoted $f_{c_k, Y}$. We define its estimate by $I_{c_k, Y, T}$. Let be the following statistic:

$$\tilde{T}_k(Y) = \sup_{\lambda \in [-\pi, \pi]} \left| \int_{[-\pi, \pi]^{k-1}} \left(\frac{I_{c_k, Y, T}(z)}{f_{c_k, Y}} - \frac{\tilde{c}_k}{c_k} \right) dz \right|, \quad (4)$$

where $\tilde{c}_k$ is an estimate of c_k . It can be shown that - under the null that the cumulants of order k are invariant on the subsamples - this statistic $\tilde{T}_k(Y)$ converges in distribution to $\frac{(2\pi)^{k-1}}{c_k} B(\sum_{j=1}^{k-1} \lambda_j)$, when $T \rightarrow \infty$, $-\pi < \lambda < \pi$, and $B(\cdot)$ is the Brownian bridge. This result is an extension of a result of Kluppelberg and Mikosch (1996) and is developed in a companion paper. For definition of cumulants we refer to Priestley (1981).

Now, following the same idea developed by Starica and Granger (2005), we can build homogeneity intervals using this test, for $k = 1, 2, 3, 4$, with the critical values of the previous statistic. First, we consider a subset $(Y_{m_1}, \dots, Y_{m_2})$, $\forall m_1, m_2 \in N$, on which we apply the previous test and we build confidence intervals. Second, using moving windows, we define another subset, for some $p \in N$, $(Y_{m_2+1}, \dots, Y_{m_2+p})$ on which we apply again the test and verify if the value of the statistics belongs to the confidence interval previously built, or not. If it belongs to the previous confidence interval, we continue with a new subset; if not, we consider $(Y_{m_1}, \dots, Y_{m_2+p})$ as an homogeneity interval and analyse the next subset $(X_{m_2+p+1}, \dots, X_{m_2+2p})$ defining a new confidence interval, and so on. At the end we obtain a sequence of intervals on which we can estimate a process stationary, up to four moments. Thereby, we have determined a sequence of distinct processes, each characterized by an approximately invariant distribution function.

In the following, we use this sequence of "stationary" intervals characterized by their invariant measures to build a distribution which characterizes the whole sample. This last one will be made using the concept of copulas that we recall briefly.

4.2 The copula concept

Consider a general random vector $Z = (X, Y)^T$ and assume that it has a joint distribution function $F(x, y) = \mathbb{P}[X \leq x, Y \leq y]$ and that each random variable X and Y has a continuous marginal distribution function respectively denoted F_X and F_Y . It has been shown by Sklar (1959) that every 2-dimensional distribution function F with margins F_X and F_Y can be written as $F(x, y) = C(F_X(x), F_Y(y))$ for an unique (because the marginals are continuous) function C that is known as the copula of F (this result is also true in the r -dimensional setting). Generally a copula will depend almost on one parameter, then we denote it C_α and we have the following relationship:

$$F(x, y) = C_\alpha(F_X(x), F_Y(y)). \quad (5)$$

Here, a copula C_α is a bivariate distribution with uniform marginals and it has the important property that it does not change under strictly increasing transformations of the random variables X and Y . Moreover, it makes sense to interpret C_α as the dependence structure of the vector Z . In the literature, this function has been called "dependence function", "uniform representation" and "copula". We keep this denomination here, Sklar (1959).

Practically, to get the joint distribution F of the random vector $Z = (X, Y)^T$ given the margins, we have to choose a copula that we apply to these margins. There exists different families of copulas: the elliptical copulas, the archimedean copulas, the meta-copulas, etc., Nielsen (1999). In order to choose the best copula adjusted for a pair of random variables, we need to estimate the parameters of the copula and to estimate the copulas. We quickly review the different methods, and for more details, we refer to Cherubini et al. (2001), and Caillault and Guégan (2005). Concerning the parameters of the copulas, we can use the Kendall's tau because it exists a fairly relationship between this coefficient and the parameters of the elliptical copulas and the Archimedean copulas. A classical maximum likelihood approach is also possible as soon as the random variables with which we work have been whitened. To determine the copulas, several criteria can be used, the pseudo loglikelihood, the AIC criteria, or a diagnosis based on the L^2 distance:

$$D_2 = \sum_{m=0}^T \sum_{n=0}^T \left| C_{\hat{\alpha}}(\hat{F}_X(x_m), \hat{F}_Y(y_n)) - \hat{F}(m/T, n/T) \right|^2,$$

where $\hat{F}_X$, and $\hat{F}_Y$ correspond to the margins, and $\hat{F}(\cdot, \cdot)$ is the empirical distribution function associated to the sample. The copula $C_{\hat{\alpha}}$ for which we will get the minimum distance D_2 will be chosen as the best approximation to link $\hat{F}_X$ and $\hat{F}_Y$, in that sense.

4.3 A meta distribution

We can use the previous method in a more general setting, and apply it to our situation. recall that we observe $Y_1, \dots, Y_T$, and we want to determine the joint distribution function $F_Y = P[Y_1 \leq y_1, \dots, Y_T \leq y_T]$ under invariance assumptions.

On figure 2, we illustrate the previous procedure, which has permitted to build a sequence of r homogeneity intervals on which we define stationary processes characterized by an invariant distribution function $F_{Y^{(i)}}$, $i = 1, \dots, r$. On this

figure, we have identified four homogeneity intervals characterized by changes in mean or in variance, and we have estimated different distribution functions on each subsample denoted $F_Y(1), F_Y(2), F_Y(3), F_Y(4)$. We formalize now this example.

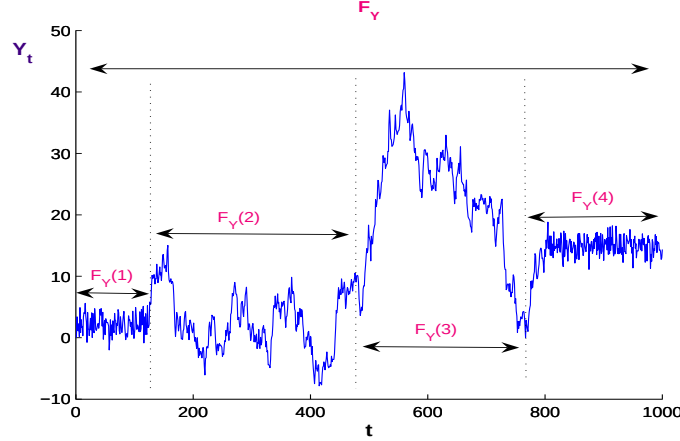


Figure 2: Example of a sequence of invariant distribution functions

Having built previously stationary margins $F_{Y^{(i)}}$, $i = 1, \dots, r$ on each subsample, we know that it exists an unique copula C_α linking this sequence of invariant distribution functions, such that:

$$F(Y_1 \leq y_1, \dots, Y_T \leq y_T) = C_\alpha(F_1(Y^{(1)}), \dots, F_r(Y^{(r)})). \quad (6)$$

The expression (6) provides a new way to characterize the joint distribution of the sample $Y_1, \dots, Y_T$ using a sequence of invariant distributions. We call this distribution a meta-distribution and we will see that we can use it for different purposes.

There exists different ways to build this meta-distribution. Indeed, the expression (6) provides one approach but we can complexify the building.

1. We can assume that the parameter α of the copula evolves in time, then (6) becomes:

$$F(Y_1 \leq y_1, \dots, Y_T \leq y_T) = C_{\alpha_t}(F_1(Y^{(1)}), \dots, F_r(Y^{(r)})). \quad (7)$$

We get a dynamical copula. It has already been investigated and estimated in different papers, Dias and Embrechts (2004), Fermanian (2005), Caillault and Guégan (2009) and Zhang and Guégan (2009).

2. We can build a sequence of copulas, with different parameters permitting to link successively two or three margins, in order, for instance, to have a parameter α which stays constant. In case of two margins, we get the sequence: $C_{\alpha_{12}}(F_1(Y^{(1)}), F_2(Y^{(2)})), C_{\alpha_{34}}(F_3(Y^{(3)}), F_4(Y^{(4)})), \dots, C_{\alpha_{r,r-1}}(F_{r-1}(Y^{(r-1)}), F_r(Y^{(r)}))$. Then, the meta distribution will be defined thanks to another copula C_β . Indeed, we are going to consider each expression $C_{\alpha_{ij}}(F_i(Y^{(i)}), F_j(Y^{(j)}))$ as a margin, and then C_β will link all these margins:

$$F(Y_1 \leq y_1, \dots, Y_T \leq y_T) = C_\beta(C_{\alpha_{12}}, C_{\alpha_{34}}, \dots, C_{\alpha_{r,r-1}}). \quad (8)$$

An estimation procedure has to be defined for this latter approach, extending the previous cited works.

3. Now, we can use more than two margins. In case of three margins, we can use Archimedean copulas under specific restrictions (Nielsen (1999)). Then, we define a copula C_ξ linking the margins $C_{\alpha_{ijk}}(F_i(Y^{(i)}), F_j(Y^{(j)}), F_k(Y^{(k)}))$.

$$F(Y_1 \leq y_1, \dots, Y_T \leq y_T) = C_\xi(C_{\alpha_{123}}, \dots). \quad (9)$$

4. To work with more than three margins, an interesting approach based on the vines can be conducted, Aas et al. (2009) and Guégan and Maugis (2008).

The methodology that we have proposed here permits to give new openings concerning several econometric issues.

1. To forecast with a non-stationary sample.

- We can use one of the previous linking copulas C_α , C_{α_t} , C_β or C_ξ to get a suitable forecast for the process $(Y_t)_t$ assuming the knowledge of the whole information set $I_T = \sigma(Y_t, t < T)$. Then, for instance, we will compute $E_{C_\alpha}[Y_{t+h}|I_T]$.
- We can also decide to forecast using a smaller information set, based on one or several homogeneity intervals, and then we will associate the distribution function which corresponds to this choice.
 - (a) If we consider the last homogeneity interval, then the predictor will be computed using $E_{F_Y^{(r)}}[Y_{t+h}|I_r]$, where $I_r = \sigma(Y_{T_{q_{r-1}+1}}^{(r)}, \dots, Y_T^{(r)})$, and $F_Y^{(r)}$ is the margin which characterizes this subset.

- (b) If we consider the last two homogeneity intervals, we will compute the expectation using the copula linking the two margins corresponding to each subsample and the information set will be the reunion of the two corresponding subsamples. Then we get $E_{C_\alpha}[Y_{t+h}|I_{r-1} \cup I_r]$, and C_α links $F_Y^{(r-1)}$ and $F_Y^{(r)}$.
- (c) If we look at the Figure 2, we can decide to use the intervals 1 and 4 to do forecast. In that case, we will compute: $E_{C_\alpha}[Y_{t+h}|I(1) \cup I(4)]$ and C_α will link $F_Y^{(1)}$ and $F_Y^{(4)}$.

We can expect that working with these approaches will provide superior forecasts.

2. To compute a risk measure. The same discussion as before can be done to compute for instance the classical Value-at-Risk measure (VaR_α), which is simply the maximum loss that is exceeded over a specified period with a level of confidence $1 - \alpha$ for a given α . For a random variable Y with distribution F_Y , it is defined by

$$F_Y(VaR_\alpha) = Pr[Y \leq VaR_\alpha] = \alpha. \quad (10)$$

Then, to make this computation, we can decide to work with the distribution which appears the more appropriate: this means that it could be one of the previous meta-distributions defined in (6), (7), (8) or (9). For instance, in case of our example (Figure 2), we can decide to use the two subsamples whose variability is the more important to compute the VaR, then the function F_Y will be the copula permitting to link $F_Y(2)$ and $F_Y(3)$, then in (10), we will use $F_Y = C_\alpha(F_Y(2), F_Y(3))$. As soon as we consider copulas of copulas to get this measure, extensions of known works has to be done to estimate it, Fermanian (2005), Caillault and Guégan, (2005) and Guégan (2009)

3. To determine the unconditional distribution of non-linear models. Indeed, this approach could permit to solve this open problem. In particular, it is possible to obtain the distribution function of any Markov switching model, this work being the purpose of a companion paper

5 Conclusion

In this paper, we have discussed deeply the influence of non-stationarity on the stylized facts observed on the data sets and on specific statistics. For these statistics, a lack of robustness is observed in presence of non stationarity, and this work emphasizes the fact that the theoretical concept of

autocovariance function $\gamma_Y(t, h)$ and the concept of sample autocovariance $\tilde{\gamma}_Y(h)$ are totally different as soon as some specific features are detected in the sample. This evidence is illustrated through an example built using a simple switching process¹. It is important to notice that this work does not concern asymptotic theory but discusses a new working framework to analyse non-stationary time series.

Then a new methodology is proposed in order to take the non-stationarity into account, building a sequence of invariant "homogeneity" intervals up to order four, and considering a new way to associate to a sample a joint distribution which can be used to compute forecast or risks. This new methodology opens the routes to solve a certain number of technical unsolved problems as the computation for instance of the unconditional distribution of some non-linear processes.

Some extended researches can be also considered, we cite some for illustrations.

- The use of the change point theory could be used to verify the date at which we determine the beginning of an homogeneity interval. This could be a nice task. Indeed, most of the works concerning the change point theory concern detection of breaks in mean or in volatility. These works have to be reexamined taking into account the fact that breaks can provoke spurious long memory. Indeed, in that latter case, the use of the covariance matrix can be a problem in the sense that we cannot observe change point in the covariance matrix.
- The time spend in each state when we observe breaks is a challenging problem. We do not consider here the approach done in the ACD models, we are interested to characterize the distribution function permitting to quantify the time spend in a state. This last random variable is important in order to characterize the existence of states, and it can be connected to the creation of the long memory behavior. It is known that for a Markov switching process, this law is geometric depending on the transition probabilities. More deep work is necessary to understand its rule in the creation of "spurious" long memory. One way is to study the behavior of the autocovariance function when, for example, we assume that the time spend on a regime follows a specific law like the Poisson law or any classical continuous law.

¹Other examples can be provided under request.

- The discussion of models taking into account sharp switches and time varying parameters. A theory has to be developed to answer to a lot of questions coming from practitioners. If the model proposed by Hyung and Franses (2005) appears interesting in that context, because it nests several related models by imposing certain parameter restrictions (AR, ARFI, STOPBREAK, models for instance, etc..), more identification theory concerning this class of models need to be done to understand how it can permit to give some answers to the problematic developed in this paper.
- Another approach which can be linked to the previous work concerns the test theory to detect "spurious" long memory behavior when this one is created by non-stationarity, some interesting references being Sibbersten and Kruse (2009), and Ohanassian, Russell and Tsay (2008). Another way could be, using the previous work approach to adjust an FI(d) on each interval, and to test the null that $d_1 = d_2 = \dots = d_r = d$, where d is the value of the fractional differencing parameter obtained with the whole sample. Some preliminary empirical discussions on this approach have been done by Charffedine and Guégan (2008).

Acknowledgments. Part of this work has been done during an opportunity offered to visit CREATES in Aarhus University in Denmark. The author wants also to thank the participants to the different seminars where the paper has been presented, namely in the Universities of Hong Kong, QUT in Brisbane, Monash in Melbourne and University of Adelaide, in Australia. These very interesting discussions and suggestions have permitted to finalize this work.

References

- [1] Aas K., Frigessi A., Bakken H. (2009) Pair-copula constructions of multiple dependence, *Forthcoming in Insurance: Mathematics and Economics*.
- [2] Bec F., Ben Salem M., Carrasco M. (2004) Tests for unit-root versus threshold specification with an application to the purchasing power parity relationship, *Journal of Business and Economic Statistics*, 22, 382 - 395.
- [3] Brockwell P.J., Davis R.A. (1996). *Introduction to time series analysis*, Springer Verlag, New York.

- [4] Caillault C., Guégan D. (2005) Empirical Estimation of Tail Dependence Using Copulas. Application to Asian Markets, *Quantitative Finance*, 5, 489 - 501.
- [5] Caillault C., Guégan D. (2009) Forecasting VaR and Expected shortfall using dynamical Systems : a risk Management, *Frontiers in Finance and Economics*, nov. 2009.
- [6] Charfeddine L., Guégan D. (2008) Breaks or long memory behaviour: an empirical investigation, WP-2008.22, Publications du CES Paris 1 Panthéon-Sorbonne.
- [7] Cherubini, U., Luciano E., Vecchiato W. (2004) *Copula Methods in Finance*, Wiley Finance, New York.
- [8] Dias, A., and Embrechts, P. (2004) Dynamic Copula Models for Multivariate High- Frequency Data in Finance. In: G. Szegoe (Ed.), Risk Measures for the 21st Century, Chapter 16, pp. 321-335, Wiley Finance Series,
- [9] Ding Z., Granger C.W.J. (1996). Modelling volatility persistence of speculative returns: a new approach, *Journal of Econometrics*, 73, 185 - 215.
- [10] Engle R. F. (1982). Autoregressive conditional heteroscedasticity with estimates of the variance of the United Kingdom, *Econometrica*, 50, 987 - 1007.
- [11] Fermanian, J-D. (2005) Goodness of fit tests for copulas. Journal of multivariate analysis, 95 (1): 119-152.
- [12] Granger, C.W.J. (1980). long memory relationships and the aggregation of dynamic models, *Journal of Econometrics*, 14, 227 - 238.
- [13] Granger, C.W.J., A.P. Andersen (1978). *An introduction to bilinear time series analysis*, Vandenhoeck and rupecht, Gottingen.
- [14] Gray H., N. Zhang, W. Woodward (1989). On generalized fractional processes, *Journal of Time Series Analysis*, 10, 233-257.
- [15] D. Guégan (1987) Different representations for bilinear models, *Journal of Time Series Analysis*, 8, 389 - 408.
- [16] Guégan D., Pham T.D. (1992) Power of Score Tests against Bilinear Time Series Analysis , *Statistica Sinica*, Vol. 2,1, 157-171.

- [17] Guégan D. (2000). A new model: the k -factor GIGARCH process, *Journal of Signal Processing*, 4, 265 - 271.
- [18] Guégan D. (2005) How can we define the concept of long memory? An econometric survey, *Econometric Reviews*, 24, 113 - 149.
- [19] Guégan D. (2009) VaR computation in a Non-Stationary Setting, eds. Gregoriou G.N., Hopp C., Wehn K. *Model Risk Evaluation Handbook*, Chapter 15, McGraw Hill Cie.
- [20] Guégan D., Maugis P.A. (2008) Note on new prospects on vines, WP 2008-95, Working paper, CES University Paris1 Panthéon -Sorbonne, paris, France.
- [21] Guégan D., Zhang J. (2009) Change analysis of dynamic copula for measuring dependence in multivariate financial data, *Quantitative Finance*, in press.
- [22] Mandelbrot B. (1963) The variation of certain speculative prices, *The Journal of Business*, 36, 394 - 419.
- [23] Kluppelberg C, Mikosch T. (1996) Gaussian limit fields for the integrated periodogram, *Annals of Applied Probability*, 6, 969 - 991.
- [24] Mikosch T. , Starica C. (2004). Nonstationarities in financial time series, the long range dependence and the IGARCH effects, *The Review of Economics and Statistics*, 86, 378 - 390.
- [25] Nelsen, R. (1999). *An Introduction to Copulas*, Springer, New York.
- [26] Ohanassian A., Russell J.R., Tsay R.S. (2008) True or Spurious Long Memory? A New Test?, *Journal of Business and Economic Statistics*, 26, 161 - 175.
- [27] Poskitt D.S., Chung S.H. (1996). Markov Chain Models, Time Series Analysis and Extreme Value Theory, *Advances in Applied Probability*, 28, 405-425.
- [28] Priestley M. B. (1981) *Spectral analysis and Time series*, Academic Press, New York.
- [29] Robinson P.M. (1978). Statistical inference for a random coefficient autoregressive model, *Scandinavian Journal of Statistics*, 5, 163 - 168.

- [30] Sibbertsen P., Kruse R. (2009) Testing for a break in persistence under long-range dependencies, *Journal of Time Series Analysis*, 30, 263 - 285.
- [31] Sklar A. (1959). Fonctions de répartition à n dimensions et leurs marges, *Publications de l'Institut de Statistique de l'Université de Paris*, 8, 229 - 231.
- [32] Starica C., Granger C. (2005). Nonstationarities in stock returns, *The Review of Economics and Statistics*, 87, 3-18.
- [33] Terasvirta T. (2003) Testing the adequacy of smooth transition autoregressive models, in: P. Newbold and S.J. Leybourne (eds): *Recent Developments in Time Series*. Cheltenham: Elgar (2003).

6 Annex 1: Proof of the proposition (3.1)

Let $Y = (Y_1, Y_2, \dots, Y_T)$ be a sample size T of a process $(Y_t)_t$. We consider r subsamples consisting of distinct ergodic stationary processes with finite second moment. Let $p_j \in R^+$, $j = 1, \dots, r$ such that $p_1 + p_2 + \dots + p_r = 1$. Hence p_j is the proportion of observations from the j th subsample in the whole sample. We define $q_j = p_1 + p_2 + \dots + p_j$, $j = 1, \dots, r$. Thus the sample is written as $Y = (Y_1^{(1)}, \dots, Y_{Tq_1}^{(1)}, Y_{Tq_1+1}^{(2)}, \dots, Y_{Tq_{r-1}+1}^{(r)}, \dots, Y_T^{(r)})$. We define the sample autocovariances of the sequence $(Y_t)_t$ as follows:

$$\tilde{\gamma}_Y(h) = \frac{1}{T} \sum_{t=1}^{T-h} (Y_t - \underline{Y}_T)(Y_{t+h} - \underline{Y}_T).$$

We develop the right hand side of the previous relationship

$$\tilde{\gamma}_Y(h) = \frac{1}{T} \sum_{t=1}^{T-h} Y_t Y_{t+h} - \frac{\underline{Y}_T}{T} \sum_{t=1}^{T-h} (Y_t + Y_{t+h}) + \frac{1}{T} \sum_{t=1}^{T-h} \underline{Y}_T^2.$$

Let

$$A = \frac{1}{T} \sum_{t=1}^{T-h} Y_t Y_{t+h}$$

and

$$B = -\frac{\underline{Y}_T}{T} \sum_{t=1}^{T-h} (Y_t + Y_{t+h}) + \frac{1}{T} \sum_{t=1}^{T-h} \underline{Y}_T^2.$$

Thus $\tilde{\gamma}_Y(h) = A + B$. First, we compute A.

$$A = \frac{1}{T} \sum_{i=1}^r \sum_{t=Tq_{i-1}+1}^{Tq_i-h} Y_t^{(i)} Y_{t+h}^{(i)} + \frac{1}{T} \sum_{i=1}^r \left[\sum_{t=Tq_i-h+1}^{Tq_{i+1}-h} Y_t^{(i)} Y_{t+h}^{(i+1)} + \cdots + \sum_{t=Tq_{r-1}-h+1}^{Tq_r-h} Y_t^{(i)} Y_{t+h}^{(r)} \right].$$

Now, we know that $\text{cov}(Y_t^{(i)}, Y_t^{(j)}) = 0$ for all $i \neq j$ by building, thus

$$A = \frac{1}{T} \sum_{i=1}^r \sum_{t=Tq_{i-1}+1}^{Tq_i-h} Y_t^{(i)} Y_{t+h}^{(i)} + O(1).$$

We develop the term of the right hand of the previous relationship. Thus we get

$$\begin{aligned} \frac{1}{T} \sum_{i=1}^r \sum_{t=Tq_{i-1}+1}^{Tq_i-h} Y_t^{(i)} Y_{t+h}^{(i)} &= \sum_{i=1}^r p_i \frac{1}{T p_i} \sum_{t=Tq_{i-1}+1}^{Tq_i-h} Y_t^{(i)} Y_{t+h}^{(i)} \\ &+ \sum_{i=1}^r p_i E[Y_t^{(i)}]^2 - \sum_{i=1}^r p_i E[Y_t^{(i)}]^2. \end{aligned}$$

Thus

$$\begin{aligned} \frac{1}{T} \sum_{i=1}^r \sum_{t=Tq_{i-1}+1}^{Tq_i-h} Y_t^{(i)} Y_{t+h}^{(i)} &= \sum_{i=1}^r p_i E[Y_0^{(i)} Y_h^{(i)}] - \sum_{i=1}^r p_i E[Y_t^{(i)}]^2 + \sum_{i=1}^r p_i E[Y_t^{(i)}]^2 \\ &= \sum_{i=1}^r p_i \gamma_{Y^{(i)}}(h) + E[Y_t^{(i)}]^2. \end{aligned}$$

And, in probability, $A \rightarrow \sum_{i=1}^r p_i \gamma_{Y^{(i)}}(h) + \sum_{i=1}^r p_i E[Y_t^{(i)}]^2$.

Now we compute B. Using the same remark as before, B can be simplified and we get:

$$B = -\underline{Y}_T^2 + O(1).$$

Or

$$\begin{aligned} -\underline{Y}_T^2 &= -\left(\sum_{i=1}^r p_i E[Y_t^{(i)}] \right)^2 = -\sum_{i=1}^r \sum_{j=1}^r p_i p_j E[Y_t^{(i)}] E[Y_t^{(j)}] \\ &= -\sum_{i=1}^r (p_i E[Y_t^{(i)}])^2 - 2 \sum_{1 \leq i < j \leq r} p_i p_j E[Y_t^{(i)}] E[Y_t^{(j)}]. \end{aligned}$$

Moreover $p_i = p_i^2 + p_i \sum_{j \neq i, j=1}^r p_j$. Thus

$$-Y_T^2 = -\sum_{i=1}^r p_i (E[Y_t^{(i)}])^2 + \sum_{1 \leq i \leq j \leq r} p_i p_j (E[Y_t^{(i)}] - E[Y_t^{(j)}])^2.$$

Then

$$B \rightarrow -\sum_{i=1}^r p_i (E[Y_t^{(i)}])^2 + \sum_{1 \leq i \leq j \leq r} p_i p_j (E[Y_t^{(i)}] - E[Y_t^{(j)}])^2.$$

Now, using expressions found for A and B we get:

$$\begin{aligned} A+B &= \sum_{i=1}^r p_i \gamma_{Y^{(i)}}(h) + \sum_{i=1}^r p_i E[Y_t^{(i)}]^2 - \sum_{i=1}^r p_i (r E[Y_t^{(i)}])^2 + \sum_{1 \leq i \leq j \leq r} p_i p_j (E[Y_t^{(i)}] - E[Y_t^{(j)}])^2 \\ &= \sum_{i=1}^r p_i \gamma_{Y^{(i)}}(h) + \sum_{1 \leq i \leq j \leq r} p_i p_j (E[Y_t^{(i)}] - E[Y_t^{(j)}])^2. \end{aligned}$$

Hence the proposition (3.1).

- 2009-10: José Fajardo and Ernesto Mordecki: Skewness Premium with Lévy Processes
- 2009-11: Lasse Bork: Estimating US Monetary Policy Shocks Using a Factor-Augmented Vector Autoregression: An EM Algorithm Approach
- 2009-12: Konstantinos Fokianos, Anders Rahbek and Dag Tjøstheim: Poisson Autoregression
- 2009-13: Peter Reinhard Hansen and Guillaume Horel: Quadratic Variation by Markov Chains
- 2009-14: Dennis Kristensen and Antonio Mele: Adding and Subtracting Black-Scholes: A New Approach to Approximating Derivative Prices in Continuous Time Models
- 2009-15: Charlotte Christiansen, Angelo Rinaldo and Paul Söderlind: The Time-Varying Systematic Risk of Carry Trade Strategies
- 2009-16: Ingmar Nolte and Valeri Voev: Least Squares Inference on Integrated Volatility and the Relationship between Efficient Prices and Noise
- 2009-17: Tom Engsted: Statistical vs. Economic Significance in Economics and Econometrics: Further comments on McCloskey & Ziliak
- 2009-18: Anders Bredahl Kock: Forecasting with Universal Approximators and a Learning Algorithm
- 2009-19: Søren Johansen and Anders Rygh Swensen: On a numerical and graphical technique for evaluating some models involving rational expectations
- 2009-20: Almut E. D. Veraart and Luitgard A. M. Veraart: Stochastic volatility and stochastic leverage
- 2009-21: Ole E. Barndorff-Nielsen, José Manuel Corcuera and Mark Podolskij: Multipower Variation for Brownian Semistationary Processes
- 2009-22: Giuseppe Cavaliere, Anders Rahbek and A.M.Robert Taylor: Co-integration Rank Testing under Conditional Heteroskedasticity
- 2009-23: Michael Frömmel and Robinson Kruse: Interest rate convergence in the EMS prior to European Monetary Union
- 2009-24: Dominique Guégan: A Meta-Distribution for Non-Stationary Samples