

Lux, Thomas; Kaizoji, Taisei

Working Paper

Forecasting volatility and volume in the Tokyo stock market: Long memory, fractality and regime switching

Economics Working Paper, No. 2006-13

Provided in Cooperation with:

Christian-Albrechts-University of Kiel, Department of Economics

Suggested Citation: Lux, Thomas; Kaizoji, Taisei (2006) : Forecasting volatility and volume in the Tokyo stock market: Long memory, fractality and regime switching, Economics Working Paper, No. 2006-13, Kiel University, Department of Economics, Kiel

This Version is available at:

<https://hdl.handle.net/10419/3924>

Standard-Nutzungsbedingungen:

Die Dokumente auf EconStor dürfen zu eigenen wissenschaftlichen Zwecken und zum Privatgebrauch gespeichert und kopiert werden.

Sie dürfen die Dokumente nicht für öffentliche oder kommerzielle Zwecke vervielfältigen, öffentlich ausstellen, öffentlich zugänglich machen, vertreiben oder anderweitig nutzen.

Sofern die Verfasser die Dokumente unter Open-Content-Lizenzen (insbesondere CC-Lizenzen) zur Verfügung gestellt haben sollten, gelten abweichend von diesen Nutzungsbedingungen die in der dort genannten Lizenz gewährten Nutzungsrechte.

Terms of use:

Documents in EconStor may be saved and copied for your personal and scholarly purposes.

You are not to copy documents for public or commercial purposes, to exhibit the documents publicly, to make them publicly available on the internet, or to distribute or otherwise use the documents in public.

If the documents have been made available under an Open Content Licence (especially Creative Commons Licences), you may exercise further usage rights as specified in the indicated licence.

Forecasting Volatility and Volume in the Tokyo Stock Market: Long Memory, Fractality and Regime Switching

by Thomas Lux and Taisei Kaizoji



Christian-Albrechts-Universität Kiel

Department of Economics

Economics Working Paper

No 2006-13



Forecasting Volatility and Volume in the Tokyo Stock Market: Long Memory, Fractality and Regime Switching*

Thomas Lux ¹ and Taisei Kaizoji ²

¹ Department of Economics, University of Kiel,
Olshausenstr. 40, 24118 Kiel, Germany
lux@bwl.uni-kiel.de

² Division of Social Sciences, International Christian University
Mitaka City, Tokyo 181 – 8585, Japan
kaizoji@icu.ac.jp

Abstract: We investigate the predictability of both volatility and volume for a large sample of Japanese stocks. The particular emphasis of this paper is on assessing the performance of long memory time series models in comparison to their short-memory counterparts. Since long memory models should have a particular advantage over long forecasting horizons, we consider predictions of up to 100 days ahead. In most respects, the long memory models (ARFIMA, FIGARCH and the recently introduced multifractal model) dominate over GARCH and ARMA models. However, while FIGARCH and ARFIMA also have quite a number of cases with dramatic failures of their forecasts, the multifractal model does not suffer from this shortcoming and its performance practically always improves upon the naïve forecast provided by historical volatility. As a somewhat surprising result, we also find that, for FIGARCH and ARFIMA models, pooled estimates (i.e. averages of parameter estimates from a sample of time series) give much better results than individually estimated models.

Keywords: forecasting, long memory models, volume, volatility

JEL-classification: C22, C53, G12

May 2006

Corresponding author:

Thomas Lux, e-mail: lux@bwl.uni-kiel.de, Tel. +49 431 880 3661, Fax: +49 431 880 4383

* We are grateful to seminar participants at the University of Kiel for helpful discussions and to Laurie Davis and Philipp Sibbertsen for stimulating exchanges on long-memory volatility models. Helpful comments by Sheri Markose and Kyriakos Chourdakis are also gratefully acknowledged. We are furthermore appreciative for financial support by the Japan Society for the Promotion of Science and the European Commission (STREP contract no 516446)

Introduction

It is well known that both volatility and trading volume are characterized by a much higher degree of predictability than the returns of financial assets. In the huge literature on forecasting volatility, the vast majority of papers use variants of the GARCH family of stochastic processes, which provide an easy and convenient way to capture the basic autoregressive structure of conditional variances (see Granger and Poon, 2003, for a recent survey of the voluminous literature on volatility forecasting). However, results are not unanimously in favour of the potential of GARCH models to improve upon the forecasting performance of simpler models like the historical mean volatility or moving average or smoothed representations of it. Dimson and Marsh (1990) and Figlewski (1997), among others, find that simpler models can indeed outperform GARCH or related approaches at least when applied to low-frequency (weekly, monthly) data. On the other hand, dozens of papers investigate whether improvements over GARCH as a benchmark are possible using non-linear models or artificial intelligence techniques (e.g. West and Cho, 1995; Brailsford and Faff, 1996; Donaldson and Kamstra, 1997; Neely and Weller, 2002). Considering the overwhelming evidence for long-term dependence in volatility (i.e. hyperbolic decay of its autocorrelation function rather than the exponential decay characteristic of short-memory models), it also appears worthwhile to explore the potential value added by models sharing this feature. Long memory in volatility has been first pointed out as a stylised fact by Ding, Engle and Granger (1992). Prevalence of this feature in financial data has meanwhile been confirmed in many subsequent studies and counts now as one of the truly universal properties of asset markets (cf. Lobato and Savin, 1998).

Long memory generalizations of standard short-memory time series models are available in the ARFIMA (Granger and Joyeux, 1980) and FIGARCH models (Baillie et al., 1996). When browsing the literature on volatility forecasting, it comes as a certain surprise that these candidate models have received relatively scant attention so far. Basically, only two papers with a direct comparison between GARCH and FIGARCH forecasts appear to be available at present, Vilasuso (2002) and Zumbach (2004), both considering volatility forecasts in foreign exchange markets. Vilasuso reports relatively large reductions of both mean squared errors and mean absolute errors over forecasting horizons of 1 to 10 days with FIGARCH compared to GARCH. Zumbach's results using intra-daily data are more sobering in that he finds improvements in daily forecasts to be only of the order of one to two percent of MSE. Given that there is essentially only one study supportive of superior predictability of long-memory models, a more systematic analysis of this issue seems to be worthwhile.

Despite the unanimous finding of hyperbolic decay of autocorrelation functions of absolute and squared returns, it might also be the case that it is not due to 'genuine' long memory (like in FIGARCH or ARFIMA models), but is rather a consequence of unaccounted structural breaks. Lamoureux and Lastrapes (1990) have already shown that the high persistence of estimated GARCH models is greatly reduced when allowing for structural shifts. As shown by Granger and Teräsvirta (1999), simple regime-switching processes with short memory can easily give rise to spurious findings of fractional integration. In contrast, Ryden *et al.* (1998) had shown that a Markov switching model can replicate many of the time series properties of raw and absolute returns, but that it leads to exponential rather than hyperbolic autocorrelations. LeBaron (2001), however, indicated that a calibrated short memory stochastic volatility model with three factors leads to spurious finding of temporal scaling, while Alfarano and Lux (2006) show the same phenomenon to occur in a bi-stable behavioural asset-pricing framework. As these contributions show, the demarcation lines

between long memory and regime switching models are fuzzy the more so as one can construct variants of regime-switching models that literally display the signatures of long memory (cf. the ‘thought experiments’ of Diebold and Inoue, 2001, who construct particular regime-switching models with hyperbolic decay of their autocorrelation functions). This ambiguity motivates us to include a new type of Markov-switching model into our array of volatility models, the Markov-switching multifractal model of Calvet and Fisher (2001). Despite allowing for a large number of regimes, this model is more parsimonious in parameterization than other regime-switching models. It is furthermore well-known to give rise to apparent long memory on a bounded interval (Calvet and Fisher, 2004) and it has limiting cases in which it converges to a ‘true’ long memory process.

The exclusive focus on exchange markets also raises the question of how long-memory models would perform in other types of markets, e.g. in national stock markets. We attempt to shed light on both issues with our investigation of Japanese stocks. Our data base consists of daily prices and volume for more than 1,000 stocks traded in the first section of the Tokyo stock exchange (trading in the first section requires that certain criteria are met on outstanding shares, trading volume and dividend payments). Data are available at daily frequency starting in 1975. In the present study, we investigate daily stock prices and trading volume over the twenty-seven years period 01/01/1975 to 12/31/2001. A rather typical example of the evolution of stock price and daily trading volume is shown in *Fig. 1* for the Nippon Suisan Company. What stands out here and in most other time series is the enormous increase of stock prices during the Japanese bubble in the second half of the eighties and the decline thereafter. Another common feature of many stocks is the increase of trading volume during the bubble reaching levels that are often by far the highest over the whole sample. After the collapse of the bubble, prices gradually fell to their levels in the early eighties while volume also went down to roughly its level of the pre-bubble period. As can also be observed from *Fig. 1*, both volatility and volume share similarly heterogeneous time series behaviour and their periods of highest activity do broadly coincide. Both also have relatively slowly decaying autocorrelations with those for volume exhibiting distinctly more short-run dependency than those for volatility. For some stocks in our data base, trading stopped before 2001 because of bankruptcy of the company. To our knowledge, analyses of the volatility dynamics of Japanese stocks have so far been confined to short memory GARCH type models (e.g., Tse, 1991; Fong, 1997). An exception is Ray *et al.* (1997) who estimate ARFIMA models for the Tokyo Stock Price Index (TOPIX) and 15 individual stocks. However, they are interested in the predictability of raw returns (they find predictable components in the range between 5 and 15% of monthly variations), but do not explore the issue of predicting volatility or volume as we will do in the following.

Since analysis of all stocks appeared to be too time-consuming, we selected two subsets of one hundred entries, respectively. The first of these subsets consisted of those stocks with the largest average trading volume, the second subset was composed of 100 randomly chosen stocks. For all these stocks we estimated four time series models for their volatility dynamics: GARCH, FIGARCH, ARFIMA and the ‘Markov-switching multi-fractal model’ (MSM) recently introduced by Calvet and Fisher (2001), another model that at least allows to ‘mimic’ long-term dependence (see below for details). We included ARFIMA models to see whether a difference exists between the performance of the original ARFIMA structure applied to volatility and its embedding into a GARCH framework (i.e., the FIGARCH model). The multi-fractal model provides an alternative formalization of long-term dependence in volatility and has already been found to outperform GARCH and FIGARCH in some time series (Calvet and Fisher, 2004; Lux, 2005). In contrast to the additive structure of the

GARCH dynasty, the multi-fractal model conceives volatility as a hierarchical, multiplicative process with heterogeneous components, but, in fact, achieves this in a rather parsimonious way using (in the version applied here) only two parameters. Our overall finding is that improvements over GARCH can be achieved by alternative models which is in contrast to frequent findings of the opposite in the literature (which, however, mostly does not include long-memory models as alternatives).

Since volume is known to share the long memory property of volatility (Bollerslev and Jubinski, 1999; Lobato and Velasco, 2000) and to be strongly contemporaneously correlated with volatility, it seems to be worthwhile to also investigate its predictability along similar lines. Since GARCH type models are not applicable to volume, we use only the ARFIMA and multi-fractal models and compare their performance to ARMA models as a short-memory benchmark. Again, dominance of the alternative models is confirmed. As it also turns out, the predictable component in volume appears to be much higher on average than that in volatility judged by the improvements in mean-squared errors and mean absolute errors against naïve forecasts.

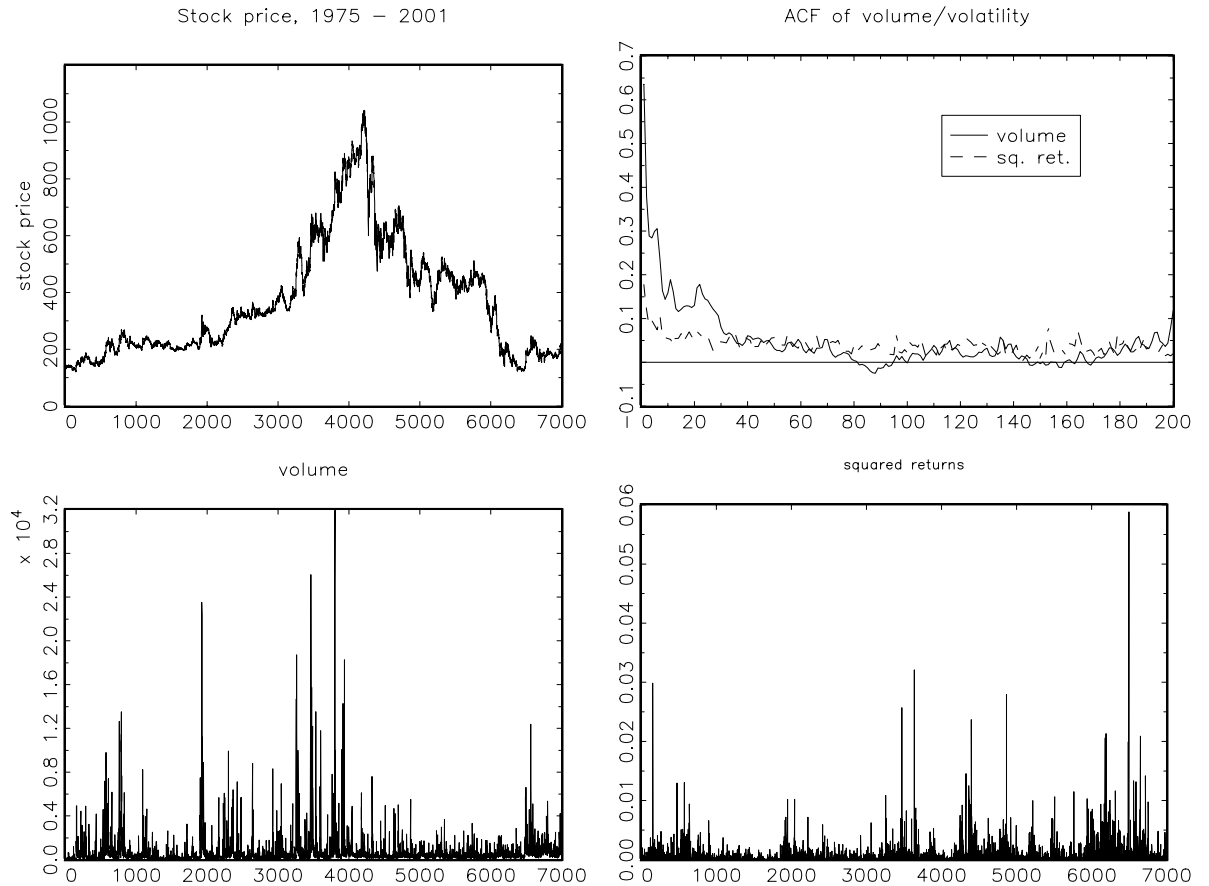


Fig. 1: Stock price, volatility (proxied by squared returns) and volume of Nippon Suisan Company (stock identification number 1332). Nippon Suisan is a company processing marine food. It has been established in 1911 and, as of September 2003, it had 1,534 employees and 36,336 shareholders. As with most other entries in our database, the Japanese bubble in the second half of the eighties has also affected the price evolution of this stock. As one typically observes, volume also reaches its maximum levels during the bubble phase.

In a final exercise, we use pooled parameters estimates (averages of the parameter estimates obtained by any particular model over the whole sample of 100 stocks) for forecasting of future volatility. Counterintuitively, discarding information about individual time series in this way leads to vast improvements of average forecast quality for volatility and, albeit to a lesser extent, for volume as well.

Our study proceeds along the following lines: sec. 2 deals with volatility forecasts. In sec. 2.1 we introduce the models and estimation methods, while results are presented in secs. 2.2 and 2.3. Similarly, models and results for volume are found in secs. 3.1 and 3.2./3.3. Sec. 4, then, considers pooled estimates. Our conclusions are to be found in sec. 5.

1. Forecasting Volatility of Japanese Stocks

2.1. Models

In our analysis of forecastability of volatility, the standard benchmark is the GARCH (1,1) model, which we expand – like all other models – by allowing for a constant and first-order autoregressive component in raw returns (r_t):

$$r_t = \mu + \rho \cdot r_{t-1} + \sqrt{h_t} \cdot \varepsilon_t \quad \text{with } \varepsilon_t \sim N(0, 1) \quad (1)$$

with volatility dynamics being governed by:

$$h_t = \omega + \alpha_1 \varepsilon_{t-1}^2 + \beta_1 h_{t-1}, \quad \omega > 0, \alpha_1, \beta_1 \geq 0. \quad (2)$$

The fractionally integrated extension of the GARCH model (FIGARCH) expands the variance equation by considering fractional differences. Like with the baseline GARCH model, we restrict attention to one lag in both the autoregressive and moving average terms, i.e., FIGARCH(1,d,1), which can be written as:

$$h_t = \omega + \beta_1 h_{t-1} + (1 - \beta_1 L - (1 - \phi_1 L)(1 - L)^d) \varepsilon_t^2. \quad (3)$$

As is well known, the Binomial expansion of the fractional difference operator introduces an infinite number of past lags with hyperbolically decaying coefficients. In practice, the number of lags considered in estimating a FIGARCH model has to be truncated. We used lag truncation at 1,000 steps.¹ Because of the time needed for FIGARCH estimation, we only consider FIGARCH (1,d,1). Both GARCH and FIGARCH are estimated via the standard MLE procedures.

Since FIGARCH adopts the ARFIMA approach for modelling the dynamics of conditional volatility, one may ask whether one could not also use the original ARFIMA model as a

¹ We also tried a moving lag length using all available past data plus 1,000 presample entries following the recommendation given in Chung (2002). However, this computationally even more burdensome practice produced virtually the same results.

possible data generating mechanism for financial volatility. The general ARFIMA model reads:

$$\Phi(L) (1-L)^d y_t = \Theta(L) \eta_t \quad (4)$$

with $\Phi(L)$ and $\Theta(L)$ the AR and MA polynomials, respectively, d the parameter of fractional differentiation and η_t white noise. In our present application, y_t is given by squared residuals after filtering of linear dependence according to eq. (1). Like with GARCH and FIGARCH, we restrict ourselves to a maximum of one autoregressive and one MA term (i.e., $p \leq 1$ and $q \leq 1$). In contrast to FIGARCH, we also tried somewhat more parsimonious variants of the model allowing for $p = 0$ or $q = 0$ (this was possible because the computational burden for ARFIMA estimation is only a fraction of that necessary to obtain FIGARCH estimates). However, the specification $p = q = 1$ was almost always preferred. Estimation has first been tried via Fox and Taqqu's frequency domain maximum likelihood approach. However, when estimating the whole set of the parameters in this way, preliminary analysis of a smaller sample of time series showed extremely volatile and often very extreme results. We, therefore, resorted to estimating the fractional differentiation parameter via the Geweke and Porter-Hudak (1983) periodogram regression and, then, estimated the remaining parameters via the method of Fox and Taqqu assuming lag polynomials with roots strictly greater than 1 in modulus (which, in fact, seems to be the most popular method for estimating ARFIMA models in applied work). For the ARFIMA model we allowed for non-stationary variants by estimating the ARFIMA model with differenced data when the initial GPH estimate of the fractional differencing parameter d exceeded the benchmark 0.5. Forecasting, then, is performed by integrating the forecasts of the differenced series.

Finally, the fourth model is the Markov-Switching Multi-Fractal Model (MSM) introduced by Calvet and Fisher (2001). Since this is a very recent addition to the array of volatility models, a few more words on its genesis and motivation are in order. The main distinguishing feature of the MSM models is the assumption of a hierarchical, multiplicative structure of heterogeneous volatility components which is fundamentally different from the linear structure of the above volatility models. Introduction of the MSM model was motivated by findings of *multi-scaling* or *multi-fractality* of financial data – different degrees of long-term dependence in different powers of returns (Ding, Engle and Granger, 1993). While traditional models in financial econometrics (i.e., the GARCH family) are of a uni-fractal nature and are, therefore, unable to accommodate this stylized fact, the literature on turbulent fluids in statistical physics has developed a class of models that exhibit just such a heterogeneous scaling behaviour (cf., Mandelbrot, 1974). In a seminal series of working papers, Mandelbrot, Fisher and Calvet (1997), Fisher, Calvet and Mandelbrot (1997), and Calvet, Fisher and Mandelbrot (1997) have proposed an adaptation of these multi-fractal processes to financial data via a compound process in which the multi-fractal component enters as a time transformation. While this model had been shown to generate scaling behaviour of moments in accordance with empirical observations (Calvet and Fisher, 2002), its attractiveness for applied research had been limited due to the combinatorial nature of the underlying multi-fractal time transformation and the non-stationarity of the resulting compound process. These limitations have been overcome by the introduction of the closely related *Markov-Switching Multifractal Model* (MSM) in Calvet and Fisher (2001). Although MSM keeps the main feature of a multiplicative hierarchy of volatility components, its Markovian structure guarantees stationarity and existence of moments and allows applying standard estimation techniques. To this end, Calvet and Fisher (2004) have proposed maximum likelihood

estimation and have demonstrated the successful performance of the model in forecasting exchange rate volatility. However, maximum likelihood is only applicable in the case of a discrete distribution for the volatility components and also becomes computationally unfeasible for too large a number of components. As an alternative, Lux (2005) has introduced GMM (generalized method of moments) estimation which is computationally much less demanding and is applicable to a broader range of specifications of the MSM model.

In MSM, financial returns are modelled along the lines of stochastic volatility models:

$$x_t = \sigma_t \cdot u_t \quad (5)$$

with σ_t instantaneous volatility being determined from a Markov-switching process with a heterogeneous ‘cascade’ of volatility components and innovations u_t drawn from a standard Normal distribution $N(0, 1)$. The volatility process is characterized by a multiplicatively connected hierarchy of k random variables $m_t^{(i)}$ of various mean life-times:

$$\sigma_t = \sigma_0 \cdot 2^k \prod_{i=1}^k m_t^{(i)}. \quad (6)$$

Adopting the perhaps most popular choice of specification of the distribution of the multipliers, we follow Mandelbrot (1974) and Mandelbrot, Fisher and Calvet (1997) by assuming that the components $m_t^{(i)}$ are drawn from a Lognormal distribution, $m_t^{(i)} \sim \text{LN}\left(-\lambda \ln(2), s^2 \ln(2)^2\right)$ in which the second parameter s can be determined by the normalization of the mean value of each component to 0.5, $E[m_t^{(i)}] = 0.5$.² The factor 2^k in

(6), then, serves to normalize the product of volatility components so that $E\left[2^k \prod_{i=1}^k m_t^{(i)}\right] = 1$

and σ_0 is a scale factor. The key feature responsible for the ‘multi-fractal’ properties of this model is that the renewal of volatility components follows a stochastic process with different probabilities depending on their rank in the hierarchy of multipliers. Here we assume a very simple structure of these transitions probabilities, i.e. $\text{Prob}(\text{new } m_t^{(i)}) = 2^{-(k-i)}$ and that replacement of a multiplier j implies simultaneous renewal of all multipliers of higher rank $i > j$ as well. This relatively simple and parsimonious construction allows for a variety of active components in current instantaneous volatility with mean live-time extending from $2^{(k-1)}$ days for the first multiplier down to $2^{(k-k)} = 1$ day for the k -th component. The MSM, therefore, combines long term dependency (via constant components which might remain active for long time horizons) with the potential of sudden changes that is rather characteristic of regime-switching models.

² Note, however, that a comparison of the Lognormal MSM models with a discrete specification with underlying Binomial distribution yields almost identical results in goodness-of-fit and forecasting capacity for a sample of exchange rates (cf., Lux, 2005).

Again, returns have been corrected for a constant mean and first-order serial dependence prior to estimation of this volatility model (i.e., x_t in eq. 5 is equivalent to $r_t - \mu - \rho r_{t-1}$ in the notation of eq. 1 and the squared returns y_t in the ARFIMA model are identical to x_t^2). A somewhat problematic feature is the choice of the number k of volatility components in the multifractal model which, in principle, amounts to a specification test of non-nested alternatives. In the absence of a convenient algorithm for the choice of the number of components, we resort to a simple heuristic advice: we estimate the parameter λ for $k = 1$ through 20 and determine the relevant value of k as that from which onward the estimate of λ did practically not change any more (did not change by more than 0.001). For both estimation of the Lognormal multi-fractal model and its use for forecasting purposes, we follow Lux (2005) by implementing the GMM estimator devised in this paper with the same moment conditions (first and second moments of log increments at various lags).

In order to derive forecasts of future volatility (future squared returns) from the above models, different algorithms have to be used. While it is possible to explicitly derive conditional expectations for GARCH, FIGARCH and ARFIMA models which, then, give the most efficient forecasts, this is not possible for the Lognormal MSM model with its highly nonlinear structure. In the later case, we, therefore, resort to best linear forecasts computed via the Levinson-Durbin algorithm (cf. Brockwell and Davis (1991, chap. 5) on the base of the autocovariances for which closed form solutions can be obtained (cf. Lux, 2005).³ Note that, although the complexity of MSM in the sense of the number of volatility states increases hyperbolically with k , the model itself remains very parsimonious. It, therefore, enables one to model a very heterogeneous structure of volatility using just three parameters, λ , σ_0 and k .

2.2. Estimated Models

The parameter estimates of the GARCH, FIGARCH, ARFIMA and MSM models are exhibited in *Tables 1 to 4*. From the roughly 1,200 stocks represented in the data base we have selected two subsamples: one consisting of the 100 companies with largest average trading volume and another with a random sample of 100 firms. The appendix lists the names and length of the available time series for all those stocks (typically from 1975 to 2001 except for firms that were liquidated over the nineties). The length of the time series used for in-sample estimation of the parameters of the various models has been restricted to the 10 year period from 1975 to the end of 1984. Our main aim in restricting the in-sample period to roughly 40 percent of the data was to leave a relatively large sample for assessment of the forecasting quality of our models which then could be investigated over more than 15 years. Assuming a stationary volatility process according to one of our models, one may argue that ten years of data should be sufficient to estimate the model parameters with sufficient reliability.

³ For the simpler Binomial MSM model of Calvet and Fisher (2004) explicit conditional expectations can be derived. However, this involves Bayesian updating of the conditional probabilities of the 2^k volatility states which is computationally burdensome. A comparison between both forecasting algorithms (Lux, 2005) shows that the performance of best linear forecasts is only slightly inferior to Bayesian forecasting.

Table 1: GARCH parameter estimates

Large volume						
	mean	std	min	max	GARCH preferred	
ω	0.345	0.320	0.004	1.978	AIC	BIC
α_1	0.769	0.165	0.000	0.983	15	21
β_1	0.151	0.135	0.016	0.999		
$\alpha_1 + \beta_1$	0.920	0.081	0.474	0.999		
Random sample						
	mean	std	min	max	GARCH preferred	
ω	0.376	0.315	0.002	1.978	AIC	BIC
α_1	0.801	0.160	0.000	0.983	14	28
β_1	0.121	0.137	0.000	0.999		
$\alpha_1 + \beta_1$	0.923	0.083	0.572	0.999		

When inspecting the distribution of parameter estimates (whose mean, standard deviation, minimum and maximum across the pertinent subsamples are given in *Tables 1 to 4*), one observes a relatively large variability. For example, both the parameters α_1 and β_1 of the GARCH model as well as the parameter of fractional differencing d in the FIGARCH model have values spread over the entirety of their admissible range $[0,1]$. Similarly high variability is observed for the ARFIMA's d although in the later case, we have not restricted the range of admissible values to $d < 1$.⁴ *Table 1* also indicates how often the GARCH model would be preferred over FIGARCH on the base of the Akaike and Schwartz information criteria (AIC and BIC). As it turns out, FIGARCH is preferred by about two 70 to 85 percent of all cases and more so under AIC. This squares with the usual observation that BIC favours more parsimonious models. *Table 3* also reports the order of the ARFIMA models (p,d,q) with $p \in \{0,1\}$ and $q \in \{0,1\}$ estimated by the AIC criterion. The $(1,d,1)$ model is the preferred one without one single exception (results with BIC are only marginally different).

Table 2: FIGARCH parameter estimates

Large Volume					Random sample			
	mean	std	min	max	mean	std	min	max
ω	0.408	0.375	0.000	2.091	0.519	0.481	0.002	2.288
β_1	0.456	0.350	-0.583	0.987	0.572	0.328	-0.499	0.986
ϕ_1	0.318	0.349	-0.646	0.874	0.370	0.352	-0.621	0.972
d	0.340	0.215	0.001	0.999	0.367	0.288	0.001	0.999

⁴ There is hardly any other choice than allowing for unrestricted d in the ARFIMA case since we cannot guarantee that the Geweke and Porter-Hudak estimates would fall into any prescribed interval. In contrast, the usual MLE estimation procedure for FIGARCH presupposes $0 \leq d < 1$.

Table 3: ARFIMA parameter estimates: volatility

Large volume							
Chosen models based on AIC				Estimates of d			
(1,d,1)	(1,d,0)	(0,d,1)	(0,d,0)	mean	std	min	max
100	0	0	0	0.235	0.104	-0.040	0.532
Random sample							
Chosen models based on AIC				Estimates of d			
(1,d,1)	(1,d,0)	(0,d,1)	(0,d,0)	mean	std	min	max
100	0	0	0	0.200	0.114	-0.093	0.479

Lastly, turning to our parameter estimates for the MSM model, we again see a large variation of parameter values. Note that the Lognormal distribution parameter λ is restricted to the open half line $[1, \infty)$. Estimates $\lambda = 1$ make the volatility cascade collapse to a constant value which leads to the benchmark case of Normally distributed returns. The mean values of the number of cascade steps k are about 14 and 12, for the ‘large volume’ and ‘random sample’ cases, respectively. It might be noted that, unlike ARFIMA and FIGARCH, the MSM model is not a ‘true’ long-memory model, but only mimics hyperbolic decline of the autocorrelation function over about 2^k time lags after which one encounters an exponential drop-off of the ACF. Note that our average estimates of k , therefore, amount to slow decline of memory over up to 16,000 time steps – much more than used for estimating the model. The maximum estimate $k = 19$ even amounts to hyperbolic decline of the autocorrelation function over roughly half a million days, i.e. 2,000 years of daily trading. It is, therefore, obvious that the deviation of the MSM model from ‘true’ long memory can be arbitrarily small. On the other hand, the minimum $k = 2$ has a range of hyperbolic scaling of just 4 days which makes it a rather clear-cut short-memory model.

Table 4: Multi-fractal parameter estimates: volatility

Large volume							
Estimate of λ				Estimates of k			
mean	std	min	max	mean	std	min	max
1.320	0.298	1.000	2.583	12.270	2.867	2	18
Random sample							
Estimate of λ				Estimates of k			
mean	std	min	max	mean	std	min	max
1.591	0.442	1.000	4.221	14.290	2.090	2	19

2.3 Forecasting Performance

Now turn to the results of our horse race for forecasting volatility: our estimated models have been tested out-of-sample for the 16-year period 1986 to 2001. Forecasting horizons start at the daily level and proceed via 5 day and 10 day forecasts up to 100 days ahead. Note that we have used only one set of parameter estimates and have not re-estimated the parameters within the out-of-sample period. The reason for not using rolling estimates is the computational burden of the maximum likelihood estimation of the parameters of the FIGARCH model – with the other models (including GMM estimation of the MSM model) it

would have been feasible. We have also looked at the performance in subsamples (1986-1990, 1991-1995, and 1996-2001), but to our surprise found no remarkable differences. As these periods cover quite diverse financial and economic conditions in Japan (including the stock market bubble, its crash and the subsequent stagnation) the homogeneity of the results speaks in favour of very *regular structure* in the volatility dynamics despite large changes in the *level* of volatility over time.

In order to compare the performance of the four candidate models, we apply the traditional concepts of mean squared error (MSE) and mean absolute error (MAE). However, since we want to have a meaningful measure allowing to compare the performance across stocks we have to standardize these statistics. We do so by reporting *relative* MSE and MAE obtained after division by the pertinent mean squared error and mean absolute error of the naïve predictor using historical volatility (i.e., the sample mean of squared returns over the period 1975 to 1984). Note that in order not to compound errors in the mean equations and in the volatility dynamics, we have also first filtered out linear dependence so that the naïve MSE and MAE are also computed from the squared residuals.

Our criteria for comparing predictive accuracy are, thus:

$$\text{relative MSE} = \frac{1}{N} \sum_{t=1}^N (h_{t,j} - \varepsilon_t^2)^2 \bigg/ \frac{1}{N} \sum_{t=1}^N (h_{t,n} - \varepsilon_t^2)^2, \quad (7)$$

$$\text{relative MAE} = \frac{1}{N} \sum_{t=1}^N |h_{t,j} - \varepsilon_t^2| \bigg/ \frac{1}{N} \sum_{t=1}^N |h_{t,n} - \varepsilon_t^2|; \quad (8)$$

with $t = 1, \dots, N$ the out-of-sample observations, $j = \{ \text{GARCH, FIGARCH, ARFIMA, MSM} \}$ the estimates from the candidate time series models, the subscript n denoting the naïve predictions using historical volatility and ε_t the residuals obtained after linear filtering of returns (using the in-sample means and first-order autocorrelations).

Table 5 compares the average relative MSEs and MAEs of our four models for the ‘large volume’ of 100 stocks. Since results obtained for the ‘random sample’ cases are very similar in most respects, we have not reproduced them here for the sake of brevity, but are prepared to supply them upon request.

The winner in terms of average MSE reduction is the ARFIMA model (which so far has seldomly been considered as a model of volatility dynamics) followed by MSM lagging behind ARFIMA by at most one percentage point. Except for short horizons, the average forecasting quality of GARCH and FIGARCH is quite disappointing. Interestingly, GARCH performs worse than most other models *even over relatively short horizons*. The average improvement compared to naïve forecasts are in the range of up to about 8 percent at daily horizons, 4 percent at ten days and still 2.5 percent at 100 day horizons.

To provide more details, *Fig. 2* shows box plots of the distribution of MSEs and MAEs over all 100 stocks, for all methods and time horizons. A glance at the range of results for the different methods reveals some interesting tendencies. It particularly shows that for all long-memory models the inter-quartile range is below the benchmark of one for all time horizons. In contrast to GARCH, we, therefore, do find an improvement against the naïve prediction in the *majority* of cases with these models so that mean values above 1 in *Table 5* are due to

very large entries in the upper end (the median is always < 1 for all methods). At the lower end, we find that for the one day horizon, MSE can be reduced against the naïve forecast by more than 30 percent in one particularly successful application of the FIGARCH model. At the long horizon end (100 day forecasts), ARFIMA and FIGARCH provide the best cases with 10 percent improvements over the naïve prediction (with MSM and GARCH only having slightly worse ‘best cases’).

In terms of the upper end (the worst prediction within the sample), MSM is best with a maximum MSE that except for one marginal case at the one-period horizon *never rises above unity* (i.e., that is never worse than that of the naïve forecasts). Note that it, therefore, can be described as the *least dangerous* method. This is in contrast to all other models with which the user always appears to face the danger of forecasts that are worse than the most naïve ones. Note in particular how ‘dangerous’ GARCH and FIGARCH forecasts can be in certain cases! It is interesting to note that the most disastrous outcome under both the GARCH and FIGARCH models stems from the same stock. In this case, estimated GARCH parameters are close to the non-stationary case (i.e., $\alpha_1 + \beta_1$ close to 1) while with FIGARCH the parameter of fractional differentiation approaches its upper boundary 0.999 (imposed for maximum likelihood estimation). However, inspection of other cases reveals that (closeness to) non-stationarity does not always come along with poor forecasts (there are quite some similar cases in the estimated GARCH and FIGARCH models with satisfactory forecasting success as well as truly non-stationary ARFIMA cases with similarly satisfactory performance).

Lastly, it appears noteworthy that ARFIMA and MSM have very small variability of their performance which also decreases with time horizon, while the fluctuations across assets in the (FI)GARCH models rather show a tendency for increasing dispersion at long horizons.

Table 5: Forecasting Volatility

relative MSE					relative MAE			
horizon	GARCH	FIGARCH	ARFIMA	MSM	GARCH	FIGARCH	ARFIMA	MSM
1	0.939	0.944	0.915	0.921	1.060	1.088	1.049	1.026
5	0.984	0.966	0.947	0.956	1.079	1.111	1.063	1.038
10	1.005	0.982	0.957	0.965	1.096	1.133	1.068	1.042
20	1.028	1.003	0.963	0.971	1.123	1.167	1.068	1.042
30	1.056	1.031	0.968	0.976	1.153	1.199	1.068	1.041
40	1.085	1.063	0.969	0.977	1.182	1.229	1.066	1.039
50	1.124	1.103	0.971	0.979	1.215	1.261	1.064	1.038
60	1.170	1.152	0.973	0.981	1.248	1.293	1.063	1.037
70	1.220	1.207	0.974	0.982	1.279	1.324	1.062	1.036
80	1.279	1.271	0.975	0.982	1.312	1.355	1.060	1.035
90	1.342	1.339	0.975	0.983	1.343	1.385	1.058	1.033
100	1.413	1.418	0.975	0.983	1.375	1.417	1.057	1.032

Note: the ‘winners’ under each criterion are marked by bold numbers.

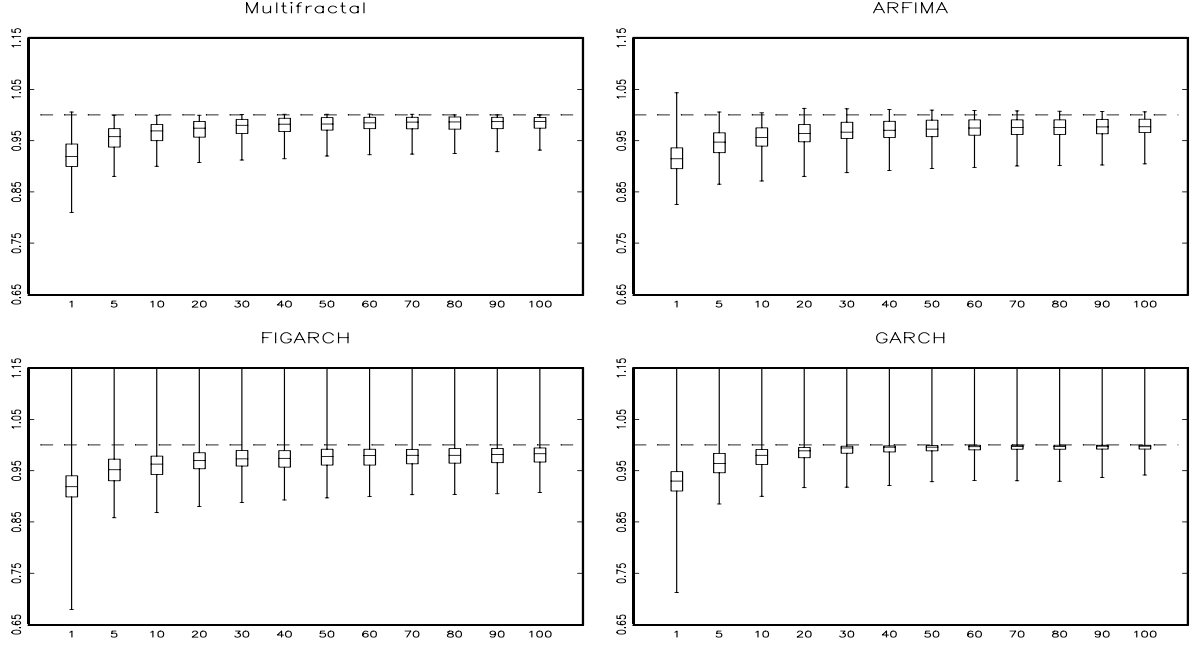


Fig. 2a: Distribution of MSEs of volatility predictions on the base of individual parameter estimates. The boxes show the median of the distribution surrounded by a box that spans the centre half of the data set (the inter-quartile range). The whiskers give the full range spanned by all 100 cases. For better comparability, we have chosen the same scale for all four box plots. The plots for the FIGARCH and GARCH results, therefore, do not show their respective maximum MSE which extends from 1.70 at lag 1 to 30.42 at lag 100 for FIGARCH (1.37 at lag 1 to 26.97 at lag 100 for GARCH).

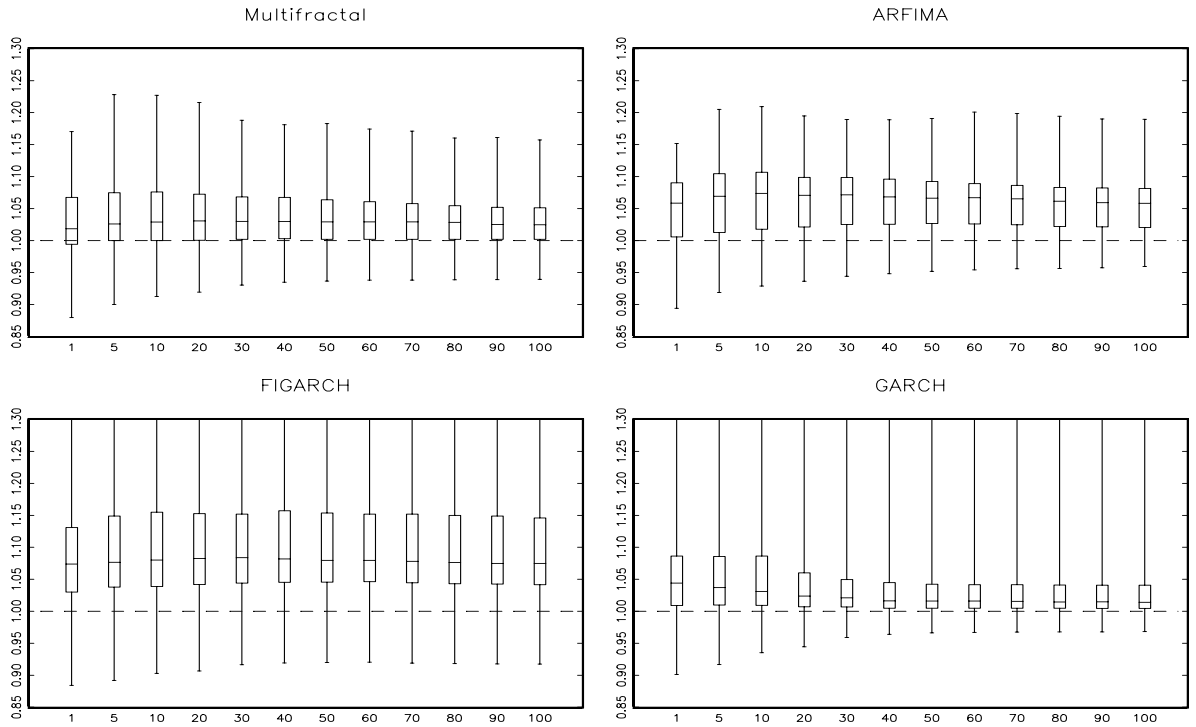


Fig. 2b: Distribution of MAEs of volatility predictions on the base of individual parameter estimates. For better comparability, we have chosen the same scale for all four box plots. The plots for the FIGARCH and GARCH results, therefore, do not show their respective maximum MAE which extends from 1.53 at lag 1 to 16.61 at lag 100 for FIGARCH (1.45 at lag 1 to 15.65 at lag 100 for GARCH).

With respect to MAE, the multi-fractal model is the winner over all time horizons. However, as a grain of salt, average performance of all models is worse than that of naive forecasts. The largest reductions of MAE achieved ranges from about 12 percent (1 day) to 8 percent (100 days). Otherwise, results are comparable to those for the MSE criterion with a narrow range of entries for MSM and a wide variation for FIGARCH and GARCH. Note also that MSM comes closest to at least generating ‘neutral’ results under this criterion while all other methods have the inter-quartile range above the benchmark value of 1 and, therefore, lead to deterioration against naive forecasts in the majority of cases.

Table 6: Rank Correlations of Volatility Predictions Across Assets

Large volume sample: relative RMSE						
lead	GARCH-FIGARCH	GARCH - ARFIMA	GARCH-MSM	FIGARCH-ARFIMA	FIGARCH - MSM	ARFIMA-MSM
1	0.938	0.629	0.307	0.674	0.360	0.487
5	0.902	0.678	0.296	0.834	0.355	0.479
10	0.810	0.568	0.319	0.816	0.330	0.501
20	0.608	0.419	0.207	0.817	0.282	0.460
30	0.524	0.344	0.166	0.830	0.293	0.446
40	0.498	0.315	0.138	0.828	0.265	0.423
50	0.509	0.293	0.140	0.789	0.235	0.420
60	0.498	0.279	0.129	0.811	0.251	0.418
70	0.499	0.275	0.125	0.805	0.243	0.430
80	0.482	0.264	0.117	0.786	0.220	0.416
90	0.489	0.272	0.098	0.765	0.176	0.400
100	0.491	0.274	0.116	0.768	0.194	0.406
Large volume sample: relative MAE						
lead	GARCH-FIGARCH	GARCH - ARFIMA	GARCH-MSM	FIGARCH-ARFIMA	FIGARCH - MSM	ARFIMA-MSM
1	0.943	0.693	0.541	0.717	0.563	0.555
5	0.893	0.608	0.500	0.680	0.524	0.533
10	0.851	0.515	0.423	0.645	0.497	0.499
20	0.733	0.346	0.327	0.601	0.451	0.449
30	0.674	0.238	0.234	0.569	0.401	0.402
40	0.643	0.178	0.202	0.555	0.385	0.408
50	0.616	0.116	0.176	0.533	0.370	0.405
60	0.609	0.098	0.149	0.527	0.364	0.400
70	0.602	0.074	0.139	0.519	0.340	0.392
80	0.594	0.050	0.116	0.514	0.333	0.385
90	0.593	0.042	0.103	0.511	0.333	0.385
100	0.582	0.028	0.091	0.502	0.326	0.380

Note: At a significance level of 95 percent, the null hypothesis of no correlation in the performance of different methods would have to be rejected for absolute entries beyond 0.197 ($=1.96 \sqrt{(n-1)}$) with $n = 100$ in our case).

A typical question arising in comparative studies of alternative predictors is whether the models under investigation use different information or not. The interesting consequence is that combinations of forecasts could improve results if the various models would not rely on the same information, whereas no such improvement appears feasible if their differences in performance are explained by different success in exploitation of the same underlying information. Typically one would use encompassing tests (Chong and Hendry, 1986) in order

to shed light on this issue. However, our large sample of stocks renders this approach somewhat unpractical. Instead, we explore this question by computing the rank correlation of the forecasting success across all assets for each pair of methods. A high entry would suggest that two methods use virtually the same information so that the difference in their relative MSEs and MAEs is mainly to be explained by difference in the accuracy of the conditional expectations. Low rank correlation, on the other hand, might suggest room for improvement via forecast combination.

Table 6 gives the rank correlation across assets for all pairs of two methods for both relative MSE and relative MAE. If all methods would have the same ranking of MSEs and MAEs across assets, rank correlations would be 1. This is not the case: although a relatively large rank correlation exists at small horizons, different methods are more or less successful in predicting the volatility of individual assets. This implies that they are not simply using the same information more or less efficiently, but that they might perform differently on different assets. Combination of forecasts, therefore, might still improve the overall results. Furthermore, the highest correlations exist between FIGARCH and GARCH at small forecasting horizons and between FIGARCH and ARFIMA at longer horizons pointing to the built-in similarities in their behaviour for short and long time horizons, respectively.

3. Forecasting Volume

3.1. Models

The empirical finance literature has mainly concentrated on trying to predict returns and volatility, but has hardly paid any attention to volume: a search for pertinent contributions in the literature has brought about only one single entry, Kaastra and Boyd (1995). These authors use neural networks and ARIMA models to forecast monthly futures trading volume for the Winnipeg Commodity Exchange. They emphasize the practical implications of volume predictions for the operation of the exchange. Besides its importance for forecasting transaction fees and liquidity, we may add that volume forecasting is also interesting in view of the similarity of the time series properties of both volatility and volume. Given the evidence on similar long term dependence in both series (Bollerslev and Jubinski, 1999; Lobato and Velasco, 2000; Ray and Tsay, 2000) it seems interesting to explore whether models with this feature are similarly capable of predicting both future volume and volatility.

We, therefore, continue our study by also using the volume entries in our data base to investigate the forecastability of transaction volume. To enhance comparability with the results obtained for volatility, we use again the same sample of stocks represented in the ‘large volume’ selection as well as the second subsample of 100 randomly chosen stocks. Since results are again quite similar, we only exhibit those for the ‘high volume’ cases and provide additional results for the randomly selected stocks upon request. As before, we estimate models on the base of the eleven year period 1975 through 1985 and test the forecasting performance of the models for the remaining sixteen years (sometimes less) 1986 to 2001 thus allowing an assessment of their success over a relatively long time horizon. When investigating volume data of stock exchanges, researchers typically find that these are non-stationary and have to be detrended first before they can be used to shed light on volatility and return dynamics. Interestingly, considering the 27 year period from 1975 to 2001 as a whole, trends in trading volume are practically non existent in the Japanese market. This is readily apparent from the quite typical behaviour of the volume of the Nippon Suisan Kaisha share exhibited in *Fig. 1*: while long subsets of the data from 1975 to about 1990 would have given rise to the impression of a positive trend, the later development suggests that the increase of volume in the second half of the eighties should be interpreted as an intermittent episode rather than the signature of a secular trend. Given this absence of clear trends, we use the *raw volume data without any correction or detrending* in our subsequent forecasting exercise.

Because of the lack of applicability of the GARCH family, only three models have been estimated for the volume time series: As a short-memory benchmark we estimate an ARMA(p,q) model: in order to see whether a moderate number of lags suffices to capture the time dependence in volume records, we select an ARMA(p,q) model for forecasting within the range $p \leq 5$ and $q \leq 5$ via maximum likelihood and use the Akaike criterion for selection of one of these alternatives. We deliberately chose AIC rather than the typically more parsimonious BIC criterion in order to allow for a sizable number of lags which could suffice for modelling the dynamics without having to resort to genuine long memory models. However, despite BIC’s tendency towards more parsimonious models, results with respect to forecasting quality turned out to be similar.

Our second model is the fractionally integrated ARFIMA(p,d,q) model. Because of the higher computational burden and also because longer lags should be captured by the fractional differentiation term, we restrict ourselves again to a maximum of one autoregressive and one MA term (i.e., $p \leq 1$ and $q \leq 1$). Estimation proceeds along the same lines as in the application of ARFIMA to volatility. For both ARMA and ARFIMA models, estimation is restricted to lag polynomials with roots strictly greater than 1 in modulus. For the ARFIMA models we allowed for non-stationary variants by estimating the ARFIMA model with differenced data when the initial GPH estimate of the fractional differencing parameter d exceeded the benchmark 0.5. Forecasting, then, is performed by integrating the forecasts of the differenced series.

Third, we also apply the multi-fractal cascade process as a model for volume. Since volume has a structure similar to volatility, we simply adopt the volatility cascade part depicted in eq. (5):

$$\theta_t = 2^k \prod_{i=1}^k m_t^{(i)}, \quad (9)$$

but skip the incremental Normal distribution introduced in eq. (6) which in the volatility model mainly serves to randomise the sign of returns. All that is needed to use this as a model of the volume dynamics is an additional scaling factor to capture the different size of mean volume in each stock. Hence, volume can be written as $\text{vol}_t = \theta_t \cdot \tilde{\sigma}_0$ where $\tilde{\sigma}_0$ is the scaling factor for asset i . Estimation is again based on GMM with appropriate moment conditions adapted from Lux (2005) and the number of cascade steps k is determined along the lines of the previous application of MSM to volatility.

Since our short-memory benchmark is now an ARMA(p,q) model, we also do *not* subject the data to filtering out linear dependency before estimating the MSM parameters. The multi-fractal cascade, therefore, is applied to volume in the format of the raw data (eq. 9) without any additional adjustments. An advantage of MSM against ARMA and ARFIMA models might be seen in the fact that by its very definition, it allows for positive entries only while negative realizations cannot be excluded in the two alternative models. We, therefore, used zero as the lower bound for our forecasts from ARMA and ARFIMA models (which, however, was hardly ever a binding constraint).

In our initial tests, we also estimated ARIMA(p,1,q) models. However, since it turned out that their forecasting performance was almost always far worse than that of the alternative models, we dropped them from our final design of this forecasting exercise.

3.2. Parameter Estimates

Now turn to the estimation results: *Table 7*, first, shows that ARMA estimation tends to favour models with many parameters at least under the AIC criterion. Comparing the AIC and BIC model selection criteria for the preferred ARMA and ARFIMA models, we see that AIC would prefer the ARMA over the ARFIMA specification in 98 out of 100 cases for both the large volume and random sample. Hence, the long-term dependence seems not to be able to compensate for the admission of more AR and MA components in our ARMA design. The

BIC, however, produces fewer cases of preferred ARMA models confirming the well-known finding that it typically favours more parsimonious models than AIC.

Table 7: ARMA parameter estimates

Large volume							
Chosen models						ARMA preferred	
(5,5)	(5,4)	(5,3)	(4,5)	(3,5)	other	AIC	BIC
26	19	10	16	6	23	98	58
Random sample							
Chosen models						ARMA preferred	
(5,5)	(5,4)	(5,3)	(4,5)	(3,5)	other	AIC	BIC
35	14	6	14	10	21	98	58

Table 8 shows that the estimated parameter of fractional differentiation from the ARFIMA models indicates a somewhat higher average d for the ‘large volume’ sample than for the random sample of firms. The preferred type of model is overwhelmingly the (1,d,1) variant modulating the prevalent long-term dependence via additional AR and MA components. Due to the wide variability of the AR and MA components, their statistics are not shown but are available upon request.

Table 8: ARFIMA parameter estimates: Volume

Large volume							
Chosen models				Estimate of d			
(1,d,1)	(1,d,0)	(0,d,1)	(0,d,0)	mean	std	min	max
91	5	4	0	0.344	0.121	0.074	0.639
Random sample							
Chosen models				Estimate of d			
(1,d,1)	(1,d,0)	(0,d,1)	(0,d,0)	mean	std	min	max
97	1	2	0	0.291	0.135	0.002	0.671

Table 9 exhibits information about the parameters of the estimated multi-fractal model. One can infer that the estimates of the key parameter λ are all within the interval between 1.00 and about 1.2 (1.00 being the lower limit for this parameter). Estimated λ 's are somewhat higher on average for the random sample of stocks signalling larger bursts of activity. This might be explained by a larger increase of trading volume during the bubble for the average firm compared to large firms which already had relatively high trading volume before the bubble episode. Compared with Table 4, we see that both the estimated λ 's as well as the number of cascade steps are smaller on average for volume than for volatility.

Table 9: Multi-fractal parameter estimates: volume

Large volume							
Estimate of λ				Estimate of k			
mean	std	min	max	mean	std	min	max
1.089	0.027	1.050	1.195	8.950	1.617	6	12
Random sample							
Estimate of λ				Estimate of k			
mean	std	min	max	mean	std	min	max
1.118	0.039	1.000	1.221	9.250	2.153	2	13

3.3. Forecasting Results

Table 10 shows the forecasting performance of the ARMA, ARFIMA, and MSM models over forecasting horizons of 1, 5, 10, 20 etc. up to 100 days for the out-of-sample period 1986 to 2001. We resort again to the criteria of relative mean squared error (MSE) and relative mean absolute error (MAE) for our assessment of the forecasting performance. The table shows the mean relative MSEs and MAEs over the 100 time series from stocks with the highest average trading volume (results for the random sample of stocks are again quite similar and can be obtained upon request). The results are perplexing: in both categories, the multi-fractal model has lowest average MSEs and MAEs over almost all time horizons. Furthermore, these means are all smaller than in the case of volatility signalling a sizable average gain in forecasting performance against the naïve model (i.e., the in-sample mean value of the time series). Roughly, MSM achieves an average improvement of 53 percent (MSE) or 33 percent (MAE) for one-day horizons and even over a forecasting horizon of 100 days has a performance that is by about 6 -8 percent better in both criteria than the naïve model. The ARFIMA forecasts mostly reach second rank for short horizon forecasts, but falls back behind ARMA to third rank for the longer forecast horizons (with both models being worse than naïve forecasts on average beyond the 1 day or 5 day horizon anyway).

A glance at the box plots in Fig. 3 also reveals that the MSM model has again the smallest standard deviation of its forecast errors showing that the success of this method is more uniform than that of ARFIMA and ARMA models. While ARFIMA is comparable to MSM in the lower end and in its inner-quartile range, it has much more extreme cases of ‘failure’ (the upper end of the whiskers). MSM’s maxima, in contrast, do at most marginally exceed the benchmark value of 1, so that here the danger of getting worse forecasts than with the naïve model is almost non existent. As with GARCH and FIGARCH in the case of volatility forecasting, ARMA models of volume produce very unsatisfactory results. Not only does one face the danger of extremely poor entries (with the record being set by an MSE up to 480 times that of the naïve model at the 10 day horizon in one case), but rather the whole ensemble of 100 forecasting exercises performs quite poorly. Overall, ARMA only produces an improvement against the naïve forecasts for the one day horizon and does hugely worse thereafter both in terms of the mean and the inner-quartile range. It is particularly astonishing that even for the one day horizon, ARMA is dominated by both ARFIMA and MSM throughout and that it performs the worst at the five day horizon although it typically uses 5

lags in either the MA or AR component or both.⁵ A closer inspection of the extreme cases of failure reveals that they occur for different stocks under the ARMA and ARFIMA models. The worst performing case of ARFIMA is one with an estimated d close to non-stationarity (the point estimate is 0.49). However, similarly as with GARCH and FIGARCH, the non-stationary cases (remember that we allow for $d > 0.5$) are not necessarily those with a poor performance out-of-sample. With ARFIMA, for instance, satisfactory forecasts are obtained with some of the non-stationary estimates.

The most astonishing feature is, however, the success of the MSM model which even for small horizons is better than its competitors – although one estimates only two parameters and unlike in the AR(FI)MA classes there are no parameters available in this model for fine-tuning of short-term dependence. Nevertheless, the MSM model mostly produces better short-term forecasts than the naïve benchmark prediction and even in its ‘bad’ cases has relative MSEs and MAEs only slightly above one while the ARFIMA and particularly the ARMA models can be far off the mark.

These results are also interesting from the perspective of the theoretical literature on forecasting on the base of ARMA and ARFIMA models. Most perplexingly, the results exhibited for MSM in Table 10 (which roughly also correspond to those in the inner-quartile range of estimated ARFIMA models) are close to (if not better than) what one could expect to obtain for an ARFIMA process with *known* parameter values. Beran (1944, chap. 4) computes MSE improvements over the (known) unconditional variance from best linear forecasts with an infinite number of past observations. In his Table 8.8. we find for $d = 0.1$ and $d = 0.4$ improvements by 1.91 percent and 48.82 percent (one step ahead), 0.22 percent and 27.06 percent (ten steps), and 0.03 percent and 14.75 percent (hundred steps). A rough interpolation between these extremes with our mean estimates for d of about 0.34 from Table 8, in fact, indicates that our volume forecasts are pretty much in line with if not better than what could be expected with a ‘true’ underlying ARFIMA model even without accounting for parameter uncertainty.

⁵ ARFIMA performs worse for the large volume than for the random sample but inspection shows that the difference is actually only due to one extreme outlier in the large volume sample. Eliminating this entry, results for large volume become almost the same as for the random sample.

Table 10: Forecasting Volume

relative MSE				relative MAE		
horizon	ARMA	ARFIMA	MSM	ARMA	ARFIMA	MSM
1	0.824	0.474	0.477	0.933	0.670	0.667
5	3.538	2.336	0.722	1.286	0.993	0.835
10	7.757	2.804	0.779	1.382	1.051	0.871
20	2.527	2.825	0.829	1.221	1.084	0.899
30	3.272	2.849	0.857	1.246	1.106	0.916
40	2.386	2.846	0.875	1.191	1.116	0.924
50	2.252	2.846	0.885	1.180	1.123	0.929
60	1.985	2.847	0.892	1.153	1.126	0.931
70	1.936	2.847	0.898	1.140	1.130	0.934
80	1.781	2.844	0.903	1.126	1.131	0.935
90	1.640	2.843	0.910	1.109	1.136	0.938
100	1.540	2.845	0.915	1.105	1.138	0.940

Note: the ‘winners’ under each criterion are marked by bold numbers.

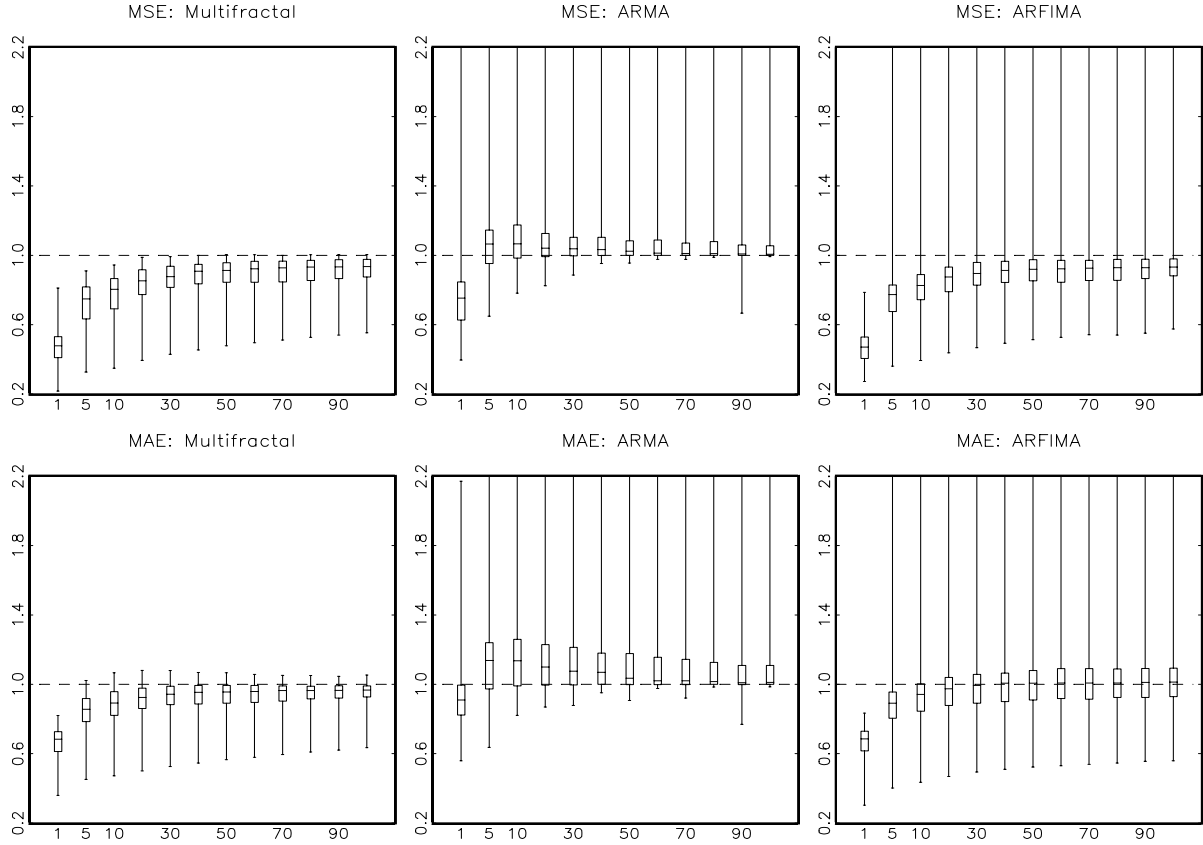


Fig. 3: Distribution of MSEs and MAEs of volume predictions on the base of individual parameter estimates.

The boxes show the median of the distribution surrounded by a box that spans the centre half of the data set (the inter-quartile range). The whiskers give the full range spanned by all 100 cases. For better comparability, we have chosen the same scale for all four box plots. The plots for the ARMA and ARFIMA results, therefore, do not show their respective maximum values which extend from 4.06 at lag 1 to 50.94 at lag 100 for MSEs (2.179 at lag 1 to 5.23 at lag 100 for MAEs) for ARMA and from 154.73 at lag 5 to 188.44 at lag 100 for MSEs (11.58 at lag 5 to 12.58 at lag 100 for MAEs) in the case of ARFIMA. Interestingly, the MSEs and MAEs of the estimated ARMA models exhibit an inverted U shape in most cases with maximum errors at the 5 and 10 day forecasting horizon (the maxima over all stocks at the 10 day horizon are 479.75 for MSE and 15.17 for MAE).

From this perspective, however, the poor performance of the ARMA class is surprising as several papers show that suitably adapted ARMA models can produce forecasts comparable to that of ‘true’ underlying ARFIMA models (Basak et al., 2001; Man, 2003). This divergence might be explained by different factors: on the one hand, estimated fractional differentiation parameters of our data are relatively high so that it is hard to cover the persistency of the data by short-memory models. In fact, it has also been shown that the approximation of long-memory models by ARMA structures works best for small values of d and becomes less satisfactory for strongly dependent processes (Brodsky and Hurvich, 1999; Crato and Ray, 1996). Furthermore, our ARMA models have been chosen by the usual AIC (or BIC) criteria and, therefore, are not those optimally adapted for forecasting an assumed underlying long-memory process. In any case, the results illustrate that the choice between short-memory and long-memory processes can crucially affect forecasting performance (even over short horizons).

Again, we try to provide an overall assessment of the degree of complementarity between methods. To this end, Spearman’s coefficients of rank correlation are exhibited in *Table 11* for both the MSE and MAE values achieved at various forecasting horizons by the various time series models. Interestingly, the correlation between the success of the MSM and ARFIMA models is highly significant over all forecasting horizons. This means that if MSM produces a high (low) reduction of MSE and MAE against the naïve model, the same is also likely to happen for the ARFIMA model. Hence, exploitation of information from past variables is quite uniform with both models: cases in which one model is particularly good while the other performs very poorly occur relatively seldom. Mostly, good results obtained from MSM coincide with relatively good forecasting performance of ARFIMA as well. In contrast, correlation between the long-memory models and the ARMA model is significant only for the one-period horizon in most cases and is uniformly much smaller than the MSM-ARFIMA correlation.

Table 11: Rank Correlations of Volume Forecasts Across Assets

lead	MSE			MAE		
	ARMA- ARFIMA	ARMA- MSM	ARFIMA- MSM	ARMA- ARFIMA	ARMA- MSM	ARFIMA- MSM
1	0.416	0.365	0.945	0.465	0.450	0.970
5	0.135	0.238	0.820	0.181	0.168	0.891
10	-0.007	0.078	0.829	0.138	0.111	0.887
20	-0.020	0.071	0.776	0.101	0.069	0.856
30	-0.063	-0.048	0.717	0.145	0.036	0.820
40	-0.104	-0.100	0.707	0.078	-0.062	0.812
50	-0.073	-0.084	0.693	0.122	-0.038	0.762
60	-0.042	-0.041	0.689	0.065	-0.053	0.736
70	-0.012	-0.066	0.679	0.146	-0.058	0.715
80	-0.120	-0.111	0.665	0.040	-0.134	0.685
90	-0.087	-0.083	0.640	0.100	-0.083	0.663
100	-0.093	-0.086	0.628	0.093	-0.111	0.640

Note: At a significance level of 95 percent, the null hypothesis of no correlation in the performance of different methods would have to be rejected for absolute entries beyond 0.197 ($= 1.96 \sqrt{(n-1)}$) with $n = 100$ in our case).

4. Forecasting with Pooled Estimates

Inspection of parameter estimates and forecasting results for GARCH, FIGARCH, ARMA and ARFIMA models across stocks shows that some of the worst results are obtained with extreme parameter estimates (although, as pointed out above, such ‘extreme’ parameters do in no way always coincide with poor forecast performance). For example, in volatility forecasting GARCH and FIGARCH performance is worst for some cases of nearly integrated processes, i.e. $\alpha_1 + \beta_1 \approx 1$ in the GARCH and $d \approx 1$ in the FIGARCH model, respectively. Similarly, the performance of the ARMA models for volume often becomes extremely poor when one of the roots approaches one. For ARFIMA, we also encounter relatively poor results for high estimates of d . Interestingly, the problem of extreme failures in some individual stocks seems non-existent for the multifractal model which in this important sense appears to be much more robust – in both its application to volatility and volume – than all the more traditional models.

The big failures of some methods with some series could have quite different sources: first, if some of the underlying time series were, in fact, non-stationary or almost non-stationary, it might simply be that their degree of forecastability is lower than for some other series. Along a similar line of argument, they might simply possess some large outliers (remember that we included the bubble period) or other particularities, which could have affected our forecasting results. Interestingly, our comparative investigation of alternative forecasts seems to allow to safely excluding this possibility: in all cases there had at least been one method whose forecasts did not perform too badly for exactly the same series (i.e., the MSM model).

An alternative explanation would, therefore, have to look for the fault in the parameter estimates of the poorly performing models. The poor forecasts might then be attributed to the variability of parameter estimates with extreme failures being due to rather large random deviations between estimated and ‘true’ parameters.⁶

In order to see whether restricting the variability of parameter estimates allows us to avoid the defective results in some cases, we designed another forecasting experiment using the *mean parameter estimates* for each model obtained across our 100 stocks.

These average estimates are taken from *Tables 1 to 4* for the volatility models and from *Tables 7 to 9* for the models used to forecast volume. To account for varying scale of the fluctuations across stocks, the following adjustments have to be made:

For the GARCH models, for instance, we now forecast on the base of the average parameters $\bar{\alpha}_1 = 0.769$ and $\bar{\beta}_1 = 0.151$. Since the remaining parameter, ω , gives the scale of fluctuations (with the unconditional variance equal to $\frac{\omega}{1-\alpha_1-\beta_1}$) it would hardly be useful to average this coefficient across stocks. Instead we compute ω from the unconditional sample variance of each (linearly filtered) return series, $\hat{\sigma}^2$, using the average estimates of the remaining parameters: $\omega = (1 - \bar{\alpha}_1 - \bar{\beta}_1) \cdot \hat{\sigma}^2$. Alternatively, we used averages for the dynamic

⁶ To be precise, we cannot really speak of ‘true’ parameters in a comparative study of various models of which none will be the ‘true’ data generating mechanism. One might, therefore, rather think of the ‘true’ parameters as those that best represent the particular class of models for a certain purpose (forecasting).

parameters α_1 and β_1 but kept the previous stock-specific estimate of ω which yielded practically identical results.

For the FIGARCH model, the mean volatility level (the unconditional variance) is not defined. However, in practice one has to approximate the fractional difference operator on the RHS of eq. (3) by a finite approximation of its expansion. The chosen cut-off of the infinite sum, then, in fact guarantees existence of the unconditional variance which is given by $\frac{\omega}{\delta_d(1)}$

with $\delta_d(L) = (1-L)^d \approx \sum_{k=0}^{k_{\max}} \delta_{d,k} \cdot L^k$ where the summation up to the cut-off $k_{\max}$ instead of

the infinite theoretical sum leads to a non-vanishing $\delta_d(1) \neq 0$ (Chung, 2002). Taking this implication of the practical approach to FIGARCH modelling into account, we can fix the parameter ω similarly as for the GARCH model in order to capture the different scales of fluctuations for individual assets. Alternatively, we also tried the average ω (along with average values of α_1, β_1 and d) for all stocks which produced practically identical results).

For the multifractal model we simply took the mean estimate of the crucial parameter λ together with the mean of the number of cascade steps rounded to the nearest integer. Since the parameters σ_0 and $\tilde{\sigma}_0$ define the scale of the process, the stock-specific sample standard deviation of returns and the sample mean of volume are used in order to maintain the different scales of fluctuations of different assets. Analogously, we also kept the scale parameters of the ARMA and ARFIMA models, but averaged over the remaining parameters. Since we allowed for flexible choice of the number of lags for these models, the mean estimates were computed for the maximum number of lags. The coefficients of the average estimates are, then, the means over the one hundred individual samples with cases of more parsimonious models contributing a zero value for their missing coefficients.

The results of our exercise are quite striking: overall, we mostly see an improvement of forecasting performance when using average instead of individually optimised parameters. *Table 12* details our results for volatility: as one can see, *under the MSE criterion, we find improvements for all models under almost all perspectives*. In particular, the mean MSE is always smaller than with the individual parameter estimates with the improvement being most pronounced for FIGARCH and GARCH whose performance was lacking behind ARFIMA and MSM when using individual parameter estimates. As a result, the three long memory models are now practically head to head with FIGARCH and ARFIMA heading the field and MSM only very slightly behind the pooled FIGARCH and ARFIMA models. Again, for all lags (even the smallest ones) GARCH despite its improvement falls clearly behind the long memory models. What is more, a look at *Fig. 4* shows that their better average performance does not come at the price of deterioration of the best cases (the lower part of the whiskers shows little variation between *Figs. 2* and *4*). It, therefore, appears that the whole distribution of forecasting results seems to shift to the left. Overall, under the MSE criterion, averaging appears almost unambiguously superior as it not only improves forecasting performance in good cases but also appears to minimize the risk of poor predictions (all the maximum entries are now close to 1).

Table 12: Forecasting Volatility: Pooled Estimates

relative MSE					relative MAE			
horizon	GARCH	FIGARCH	ARFIMA	MSM	GARCH	FIGARCH	ARFIMA	MSM
1	0.915*	0.906*	0.905*	0.909*	1.034*	1.079*	1.051	1.028
5	0.959*	0.945*	0.941*	0.947*	1.031*	1.102*	1.065	1.038*
10	0.973*	0.954*	0.952*	0.958*	1.019*	1.112*	1.069	1.040*
20	0.985*	0.958*	0.959*	0.966*	1.004*	1.116*	1.069	1.039*
30	0.994*	0.963*	0.965*	0.971*	1.001*	1.119*	1.068*	1.038*
40	0.997*	0.964*	0.967*	0.973*	1.000*	1.118*	1.065*	1.035*
50	0.999*	0.966*	0.969*	0.975*	1.000*	1.118*	1.063*	1.034*
60	1.000*	0.968*	0.971*	0.977*	1.000*	1.119*	1.062*	1.033*
70	1.000*	0.969*	0.972*	0.979*	1.000*	1.119*	1.060*	1.032*
80	1.000*	0.969*	0.973*	0.979*	1.000*	1.118*	1.058*	1.030*
90	1.000*	0.969*	0.973*	0.980*	1.000*	1.117*	1.056*	1.029*
100	1.000*	0.969*	0.974*	0.981*	1.000*	1.117*	1.055*	1.028*

Note: the ‘winners’ under each criterion are marked by bold numbers. The asterisks indicate an improvement in average MSE and MAE against the forecasts with individual parameter estimates in sec. 2. Note that pooled estimates lead to improvements for all models under the MSE criterion and for all but MSM (and FIGARCH at 10 day horizon) under the MAE criterion.

Results are somewhat different under the MAE criterion: here we find a slight deterioration for MSM and ARFIMA over short horizons, but again improvements for GARCH and FIGARCH over all horizons and ARFIMA and MSM at medium to long horizons. The winner for pooled parameter estimates is the GARCH model (except for one day forecasts) while we had a clear dominance of the MSM model for individually optimised estimates. However, with the transition from the formerly winning MSM to GARCH and ARFIMA as the best performing alternatives in the pooled estimation exercise no real gain is achieved under the MAE criterion since the average GARCH prediction essentially coincides with the naïve forecast for horizons of 20 days and more. Therefore, the major gain consists in a reduction of the relative MSE of the worst performing cases.

One remarkable difference between the results for the one hundred stocks with the largest trading volume and the second sample of randomly selected stocks is that the performance of pooled MSM estimates is much weaker in the later case than in the former (details are available upon request). Closer inspection revealed that the higher average estimate of λ in the random sample (1.59) generates less predictability in squared returns due to the larger variability of the volatility components drawn for a Lognormal distribution with a higher λ . Interestingly, using the average $\lambda = 1.32$ from the large volume stocks for the random sample as well, we recover results almost identical to those reported in Table 12. In our view, this rather underlines the power of averaging over estimates as the large average λ is apparently due to a few extreme realizations of individual estimates (maybe due to infrequent trading of some assets) and relying on a pool of estimates for more typical assets (the more frequently traded ones of our first sample) overcomes their detrimental effect.

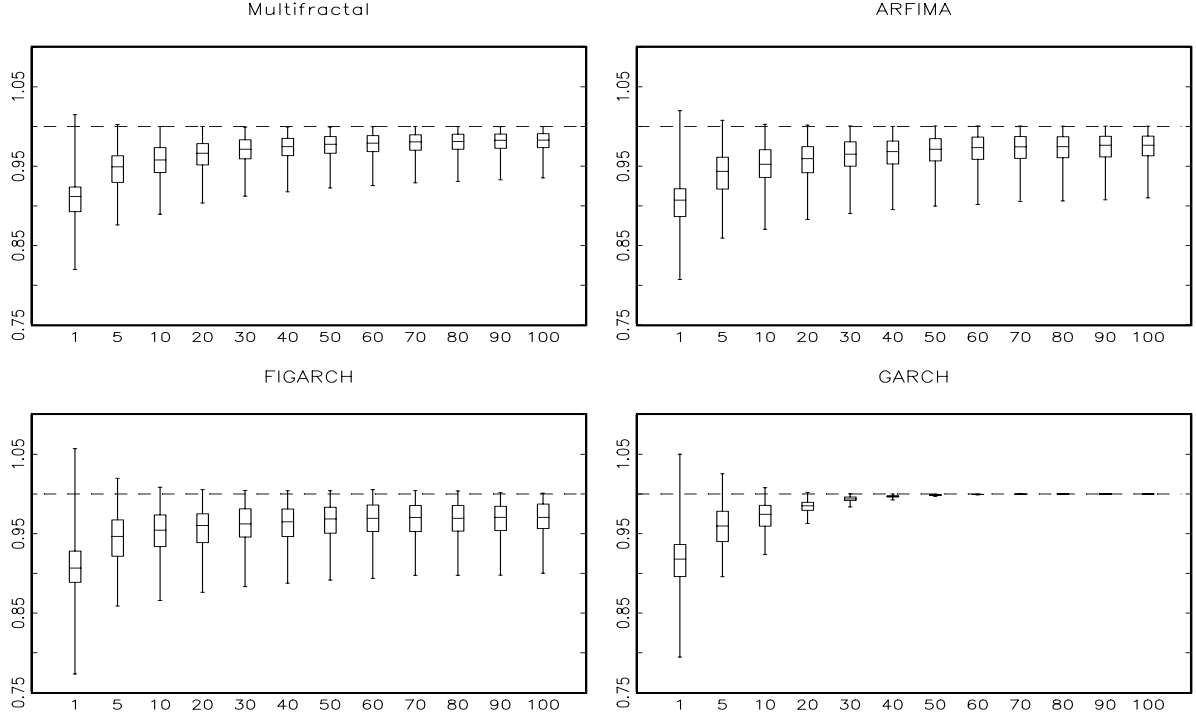


Fig. 4a: Distribution of MSEs of volatility predictions on the base of pooled parameter estimates. For the construction of the box plot, cf. the legend of Fig. 2. Apparently, the danger of poor predictions is dramatically reduced.

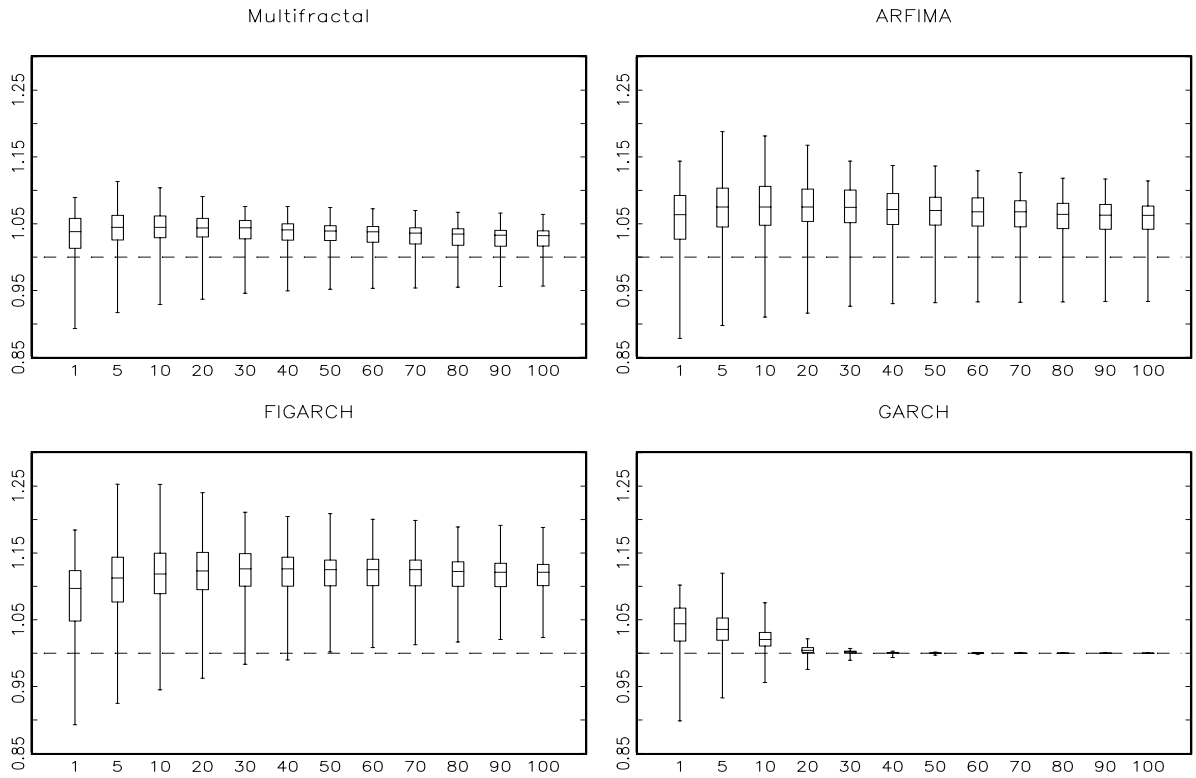


Fig. 4b: Distribution of MAEs of volatility predictions on the base of pooled parameter estimates. For the construction of the box plot, cf. the legend of Fig. 2.

Now turn to volume: again we find improvements throughout under the MSE criterion and the same applies under the MAE criterion when replacing the individual parameter estimates by pooled estimates (*Table 13*). Improvements are more spectacular for ARFIMA and particularly so for ARMA, which had a relatively poor performance with individual estimates. For the MSM model, improvements are also consistently observed at all time horizons but the magnitude of changes of MSE and MAE is relatively small. The improvement for the ARFIMA models are so pronounced that its pooled estimates even dominate over the pooled MSM forecasts. However, as can be seen from *Fig. 5*, the overall performance of MSM and ARFIMA is now very close for the MSE criterion, while MSM still has a sizable advantage compared to ARFIMA under the MAE criterion. Similarly as with volatility the most striking finding is the highly reduced danger of poor performance, particularly so for ARMA and ARFIMA models.

Table 13: Forecasting Volume: Pooled Estimates

horizon	relative MSE			relative MAE		
	ARMA	ARFIMA	MSM	ARMA	ARFIMA	MSM
1	0.590*	0.464*	0.477*	0.757*	0.665*	0.666*
5	0.958*	0.715*	0.721*	0.979*	0.853*	0.833*
10	0.962*	0.769*	0.776*	0.973*	0.896*	0.867*
20	0.989*	0.814*	0.825*	0.989*	0.929*	0.893*
30	0.994*	0.839*	0.853*	0.993*	0.949*	0.909*
40	0.997*	0.853*	0.870*	0.995*	0.959*	0.916*
50	0.999*	0.860*	0.881*	0.997*	0.964*	0.920*
60	0.999*	0.864*	0.888*	0.997*	0.966*	0.922*
70	1.000*	0.868*	0.895*	0.998*	0.969*	0.924*
80	1.000*	0.872*	0.901*	0.998*	0.971*	0.926*
90	1.000*	0.876*	0.907*	0.998*	0.974*	0.929*
100	1.000*	0.880*	0.913*	0.998*	0.976*	0.931*

Note: the ‘winners’ under each criterion are marked by bold numbers. The asterisks indicate an improvement in average MSE and MAE against the forecasts with individual parameter estimates in sec. 2. Note that pooled estimates lead to improvements for practically all models and horizons under the MSE criterion and for all but ARFIMA under the MAE criterion.

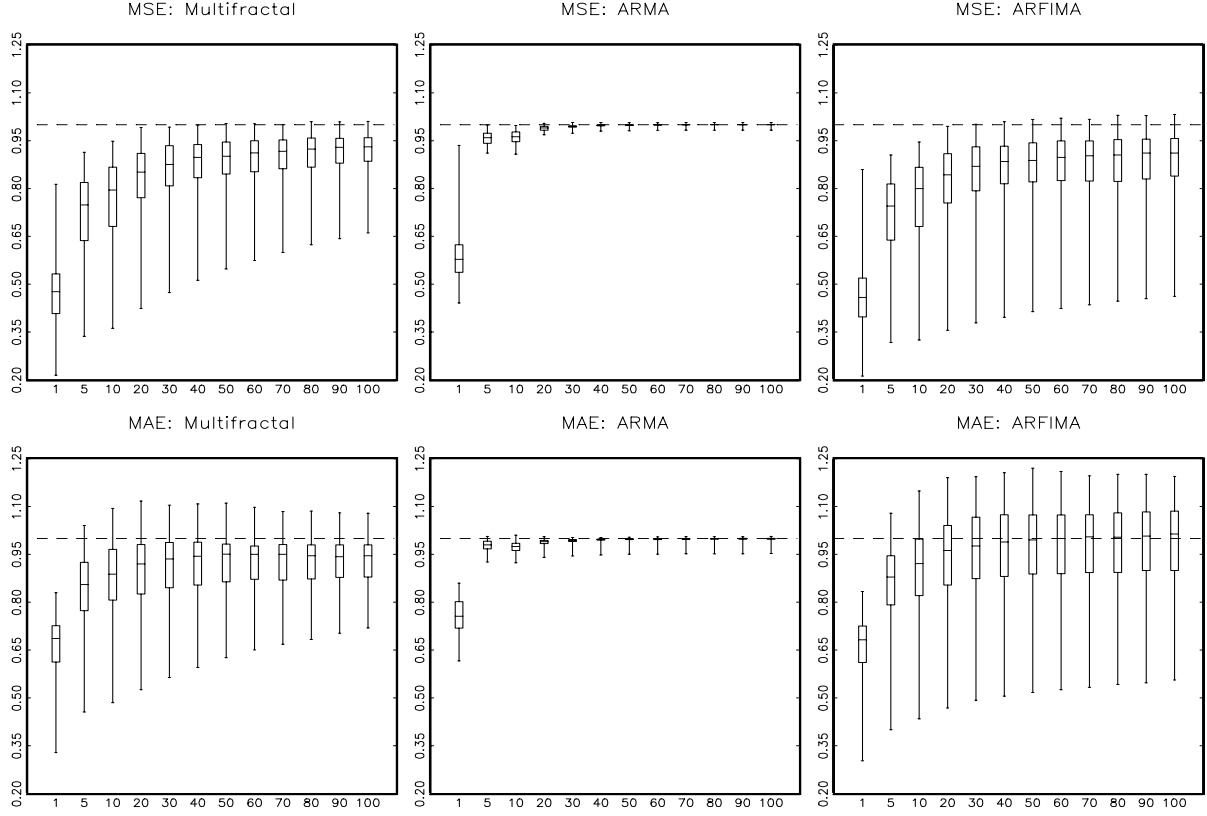


Fig. 5: Distribution of MSEs of MAEs of volume predictions on the base of pooled parameter estimates. For the construction of the box plot, cf. the legend of Fig. 2. Apparently, the danger of poor predictions is dramatically reduced.

Overall, it appears that using pooled estimates is more useful in improving predictive power for volatility models than for volume prediction, but for both volume and volatility it greatly reduces the danger of arriving at very poorly performing models. It is also instructive to compare the rank correlations between methods for pooled estimates (*Tables 14 and 15*) to those computed for our original estimates (*Tables 6 and 11*). In contrast to sec. 3 and 4, we now find a much higher correlation among the long memory models over all time horizons which even at a forecasting horizon of one-hundred days remains mostly above 90 percent. This indicates that pooled estimates of different models exploit the same features of the data. Combination of forecasts would, then, surely not improve upon the performance of the best model. In contrast, rank correlation between long memory (FIGARCH, ARFIMA and MSM) and short memory models (GARCH and ARMA) are decreasing much faster with increasing time horizon. Rank correlations between ARIMA and other models become insignificant for the volume forecasts while those of GARCH vis-à-vis other models for volatility show a somewhat disturbing coincidence of significantly positive correlation under MSE and simultaneously significantly negative correlation under MAE.

How can we explain the superiority of pooled forecasts and the good performance of the MSM model for both individual parameter estimates and pooled ones vis-à-vis the dismal behaviour of some other models? We try to shed some light on the possible origin of our findings via the illustration of forecasts over various time horizons from pooled estimates exhibited in Fig. 6. The chosen stock (Kawasaki Steel) and the time horizon (a snapshot of 500 days from the out-of-sample series) is accidental: the features we wish to highlight can be

found in all stocks, and both for volatility and volume forecasts. Comparing the forecasts over 1 day, 10 days, 50 days and 100 days, we first see in the upper left-hand panel, that forecasts from all four models at the one-day horizon are hard to distinguish by the naked eye. Obviously, all models closely track the empirical development of volatility and the subtle differences in their behaviour are certainly not obvious from a visual inspection of the plot without a more detailed statistical analysis. Some remarkable differences emerge when one proceeds to longer forecasting horizons. First, and in accordance with our expectation, GARCH forecasts successively flatten out: while they show a few remaining spikes after bursts of volatility at the 50 period horizon, they show practically no variation any more at the 100 day horizon. The same happens for ARMA forecasts of volume and both findings are in line with the short-memory properties of these models.

All other models, however, show quite some variation even at the longest horizon. The most perplexing feature is perhaps that their reactions are almost perfectly synchronous which, of course, explains their almost identical performance in Tables 12 and 13 and the high rank correlation across stocks. There is nevertheless a certain hierarchy in the strength of their reaction to changes of volatility: ARFIMA's reaction to an increase of contemporaneous volatility is most pronounced, followed by FIGARCH and MSM (in volume, ARFIMA's reaction is also stronger than that of MSM). Note that the forecasts of MSM also look smoother than those of ARFIMA and FIGARCH at all horizons. The likely reason for this weaker reaction to changes of volatility is the regime-switching nature of MSM. Since regime-switching allows for sudden changes of volatility due to renewal of multipliers, it makes the model react less on contemporaneous changes of volatility. Similarly, FIGARCH may have less volatile forecasts than ARFIMA because of the compounding of the volatility process with Normally distributed increments. The absence of both additional sources of stochasticity in the latter model may explain why its forecasts are the most volatile. As the best average performance from pooled parameter estimates is obtained for the FIGARCH model in Table 12, one is tempted to conclude that forecasts from MSM were too smooth and those from ARFIMA somewhat too volatile. However, the ranking of the three models depends on very tiny differences so that such a conclusion might not be warranted (after all, the mean MSE and MAE will also depend on the average parameter values which will change somewhat with the composition of the pool of stocks so that the ranking might be very sensitive to the number of stocks included).

However, our finding of the tendency of MSM to generate relatively smooth forecasts (due to its 'awareness' of potential shifts in volatility) might also explain its nice performance for single stocks. We conjecture that it is this feature which makes it avoid extreme reactions on changes in contemporaneous volatility and volume, which occur for other models at least in the case of some extreme parameter estimates in the vicinity of non-stationarity.

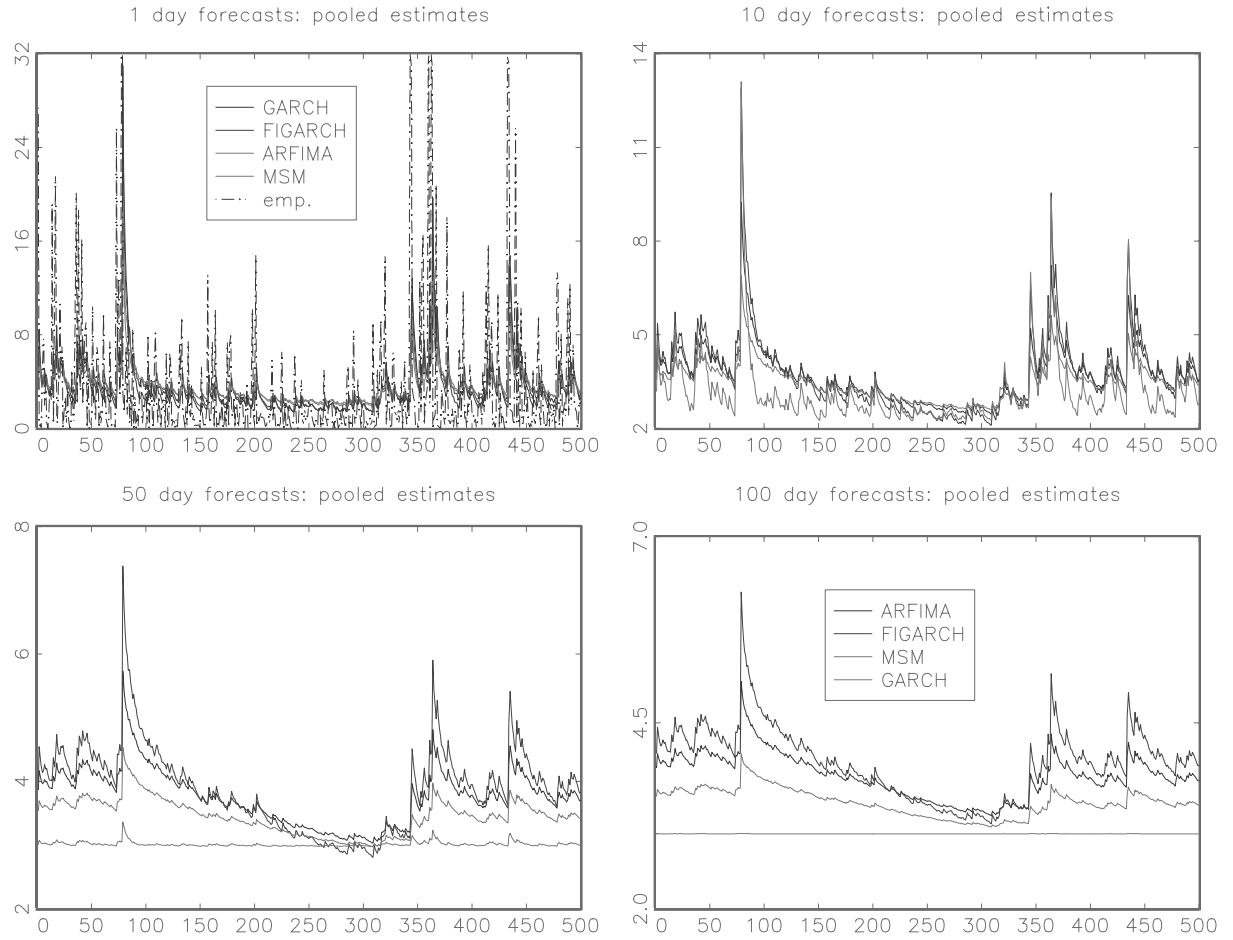


Fig. 6: Forecasts over 1, 10, 50 and 100 days from GARCH, FIGARCH, ARFIMA and MSM models with pooled parameter estimates. The underlying stock is Kawasaki Steel, but similar patterns are found for all other stocks. The upper left-hand panel exhibits squared returns alongside with its estimates from the various models, the other panels compare forecasts over longer horizons

Table 14: Rank Correlations Across Assets: Volatility Forecasts

	MSE					
lead	GARCH-FIGARCH	GARCH - ARFIMA	GARCH-MSM	FIGARCH-ARFIMA	FIGARCH -MSM	ARFIMA-MSM
1	0.974	0.976	0.983	0.961	0.977	0.995
5	0.981	0.957	0.936	0.992	0.982	0.996
10	0.958	0.936	0.915	0.995	0.986	0.997
20	0.929	0.907	0.893	0.995	0.989	0.998
30	0.939	0.927	0.919	0.996	0.992	0.998
40	0.936	0.923	0.917	0.996	0.993	0.999
50	0.943	0.931	0.927	0.996	0.992	0.999
60	0.959	0.951	0.947	0.995	0.992	0.999
70	0.949	0.943	0.938	0.995	0.993	0.999
80	0.941	0.936	0.933	0.995	0.992	0.999
90	0.937	0.923	0.922	0.994	0.991	0.999
100	0.908	0.892	0.890	0.995	0.992	0.999
	MAE					
lead	GARCH-FIGARCH	GARCH - ARFIMA	GARCH-MSM	FIGARCH-ARFIMA	FIGARCH -MSM	ARFIMA-MSM
1	0.967	0.951	0.968	0.989	0.975	0.984
5	0.947	0.924	0.939	0.991	0.971	0.984
10	0.896	0.881	0.912	0.991	0.963	0.978
20	0.708	0.692	0.768	0.983	0.954	0.974
30	0.586	0.561	0.668	0.985	0.961	0.970
40	0.438	0.417	0.562	0.976	0.954	0.965
50	0.289	0.186	0.346	0.966	0.950	0.962
60	0.341	0.205	0.376	0.961	0.946	0.961
70	0.281	0.165	0.349	0.955	0.948	0.957
80	-0.002	-0.171	0.018	0.948	0.948	0.958
90	-0.353	-0.567	-0.393	0.934	0.945	0.958
100	-0.427	-0.657	-0.483	0.922	0.942	0.956

Note: At a significance level of 95 percent, the null hypothesis of no correlation in the performance of different methods would have to be rejected for absolute entries beyond $0.197 (= 1.96 \sqrt{(n-1)})$ with $n = 100$ in our case).

Table 15: Rank Correlations Across Assets: Forecasts of Volume

	MSE			MAE		
lead	ARMA-ARFIMA	ARMA-MSM	ARFIMA-MSM	ARMA-ARFIMA	ARMA-MSM	ARFIMA-MSM
1	0.940	0.906	0.992	0.887	0.896	0.994
5	0.724	0.731	0.998	0.575	0.680	0.982
10	0.891	0.894	0.999	0.682	0.786	0.980
20	0.841	0.845	0.999	0.669	0.772	0.965
30	0.691	0.697	0.998	0.605	0.657	0.951
40	0.599	0.603	0.998	0.421	0.417	0.940
50	0.378	0.379	0.998	0.259	0.221	0.927
60	0.247	0.251	0.997	0.143	0.060	0.906
70	0.158	0.168	0.996	0.055	-0.058	0.891
80	0.106	0.107	0.995	0.020	-0.110	0.874
90	0.056	0.054	0.994	-0.014	-0.154	0.861
100	0.026	0.031	0.994	-0.019	-0.177	0.843

Note: At a significance level of 95 percent, the null hypothesis of no correlation in the performance of different methods would have to be rejected for absolute entries beyond $0.197 (= 1.96 \sqrt{(n-1)})$ with $n = 100$ in our case).

In any case, taking into account the better performance of the long-memory models, our result underscore that the data exhibit features, which are solely detected and exploited by models with long-term dependence. Since pooled GARCH and ARMA forecasts are mostly undistinguishable from naïve forecasts at long horizons, this finding speaks in favour of the either ‘true’ presence of long correlations or ‘apparent’ long-term dependence due to regime-switching in both the volume and volatility data.

5. Conclusions

This paper has examined the potential of time series models with long memory (FIGARCH, ARFIMA, multifractal) to improve upon the forecasts derived from short-memory models (GARCH for volatility, ARMA for volume). In order to get a broad picture, we have used a large data-base applying the competing models to long forecasting horizons for a long out-of-sample period. A number of interesting results emerged from this exercise: first, as concerns volatility, our selection of long-memory models performs better in most cases than the naive sample variance and GARCH forecasts.

However, this potential improvement against short-memory models is overshadowed by occasional dramatic failures particularly by the FIGARCH model and to a lesser extent by ARFIMA. Interestingly, the newly proposed multifractal approach seems not to suffer at all from this problem. Remarkably, results are better throughout for the MSE than the MAE criterion (some trial runs with other data suggest that this is not a particularity of the Japanese market). Time series methods, thus, seem to be better suited for forecasting large realizations of volatility rather than small or medium ones.

Second, as concerns volume, we find a much higher degree of forecastability than with volatility (both under the MSE and MAE criterion) and again a dominance of long-memory models (ARFIMA and MSM). As with volatility, MSM also provides much safer forecasts which hardly ever rise above the benchmark of unity under the relative MSE criterion.

Third, our observation of different degrees of the variability of performance of different methods motivated an analysis of the forecasting quality of pooled estimates (i.e. mean values of estimated parameters over the one hundred selected stocks). Astonishingly, nothing was lost by discarding stock-specific estimates, but results improved under practically all perspectives. In particular, the formerly more ‘dangerous’ methods with some extremely poorly performing cases now also became as safe as the multi-fractal model. Using pooled estimates, we also saw an even more clear-cut difference between all long-memory models and their short-memory counterparts: as can be seen in *Figs. 4* and *5*, pooled GARCH and ARMA quickly converge to the behaviour of naïve forecasts for increasing forecasting horizon yielding uniform relative MSE and relative MAE equal to unity from horizons of about forty days onward. In contrast, the long-memory models all have a uniformly better performance at least with respect to the MSE criterion. It is also remarkable, that rank correlations over markets in the pooled estimation cases are close to unity for the long-memory models showing that they all extract similar information. Overall, these results suggest the following interpretation: volatility and volume are characterized by processes which have strong persistency. While it is not obvious whether the origin of this persistent component is a ‘true’ long-memory feature or whether it is ‘apparent’ long memory stemming from a hierarchical regime-switching process, the phenomenology of the data can be captured to some degree by different time series models which have built-in long correlations. The

improvement from pooled estimates indicates that persistency is similar across stocks so that one gets a better assessment of the dependence structure by increasing the data size via merging all stocks rather than by the more common fine-tuning of individual estimates. Work in progress, in fact, confirms this view: in preliminary experiments we applied our average estimates from the Japanese market to data from other countries and again found a better performance than with individually estimated parameters. A more systematic exploration of these findings is left for future research.

References:

- Alfarano, S. and T. Lux (2006) A Noise Trader Model as a Generator of Apparent Financial Power Laws and Long Memory, *Macroeconomic Dynamics* (in press).
- Baillie, R.T., T. Bollerslev and H.O. Mikkelsen (1996) Fractionally Integrated Generalized Autoregressive Conditional Heteroskedasticity, *Journal of Econometrics* 74, 3 – 30.
- Basak, G., N. Chan and W. Palma (2001) The Approximation of Long-Memory Processes by an ARMA Model, *Journal of Forecasting* 20, 367-389.
- Beran, J. (1994) *Statistics for Long-Memory Processes*. New York: Chapman & Hall.
- Bollerslev, T. and D. Jubinski (1999) Equity Trading Volume and Volatility: Latent Information Arrivals and Common Long-Run Dependencies, *Journal of Business & Economic Statistics* 17, 9-21.
- Brailsford, T. and R. Faff (1996) An Evaluation of Volatility Forecasting Techniques, *Journal of Banking and Finance* 20, 419-438
- Brockwell, P.J. and R.A. Davis (1991) *Time Series: Theory and Methods*. 2nd ed., New York, Springer.
- Brodsky, J. and C. Hurvich (1999) Multi-Step Forecasting for Long-Memory Processes, *Journal of Forecasting* 18, 59-75.
- Calvet, L., A. Fisher and B. Mandelbrot (1997) *Large Deviations and the Distribution of Price Changes*. Cowles Foundation Discussion Paper no. 1165, Yale University.
- Calvet, L. and A. Fisher (2001) Forecasting Multifractal Volatility, *Journal of Econometrics* 105, 27 – 58.
- Calvet, L. and A. Fisher (2002) Multi-Fractality in Asset Returns: Theory and Evidence, *Review of Economics & Statistics* 84, 381 – 406.
- Calvet, L. and A. Fisher (2004) Regime-Switching and the Estimation of Multifractal Processes, *Journal of Financial Econometrics* 2, 49 – 83.
- Chong, Y. and D. Hendry (1986) Econometric Evaluation of Linear Macroeconomic Models, *Review of Economic Studies* 53, 671 – 90.
- Chung, C.-F. (2002) *Estimating the Fractionally Integrated GARCH Model*. Mimeo: Academia Sinica.

- Crato, N. and B. Ray (1996) Model Selection and Forecasting for Long-Range Dependent Processes, *Journal of Forecasting* 15, 107-125.
- Diebold, F. and A. Inoue (2001) Long Memory and Regime Switching, *Journal of Econometrics* 105, 131-159.
- Dimson, E. and P. Marsh (1990) Volatility Forecasting without Data-Snooping, *Journal of Banking and Finance* 14, 399-421.
- Ding, Z., C.W.J. Granger and R.F. Engle (1993) A Long Memory Property of Stock Market Returns and a New Model, *Journal of Empirical Finance* 1, 83 - 106.
- Donaldson, G. and M. Kamstra (1997) An Artificial Neural Network-GARCH Model for International Stock Return Volatility, *Journal of Empirical Finance* 4, 17 – 46.
- Figlewski, S. (1997) Forecasting Volatility, *Financial Markets, Institutions & Investment* 6, 1-88.
- Fisher, A., L. Calvet and B. Mandelbrot (1997) *Multifractality of Deutschemark/US Dollar Exchange Rates*. Cowles Foundation Discussion Paper no. 1166, Yale University.
- Fong, W. (1997) Volatility Persistence and Switching ARCH in Japanese Stock Returns, *Financial Engineering and the Japanese Markets* 4, 37-57.
- Granger, C. and R. Joyeux (1980) An Introduction to Long-Memory Time Series Models and Fractional Differencing, *Journal of Time Series Analysis* 1, 1980, 15 – 39.
- Granger, C. and T. Teräsvirta (1999) A Simple Nonlinear Time Series Model with Misleading Linear Properties, *Economics Letters* 62, 161-165.
- Kaasta, I. And M. Boyd (1995) Forecasting Futures Trading Volume Using Neural Networks, *Journal of Futures Markets* 15, 1995, 953 – 970.
- Lamoureux, C. and W. Lastrapes (1990) Persistence in Variance, Structural Change and the GARCH Model, *Journal of Business and Economics Statistics* 8, 225-234.
- LeBaron, B. (2001) Stochastic Volatility as a Simple Generator of Apparent Financial Power Laws and Long memory, *Quantitative Finance* 1, 621-631.
- Lobato, I. and N. Savin (1998) Real and Spurious Long-Memory Properties of Stock Market Data, *Journal of Business and Economics Statistics* 16, 261 – 283.
- Lobato, I. and C. Velasco (2000) Long-Memory in Stock Market Trading Volume, *Journal of Business and Economics Statistics* 18, 410 - 427.
- Lux, T. (2005) *The Markov-Switching Multi-Fractal Model of Asset Returns: GMM Estimation and Linear Forecasting of Volatility*. Working paper, University of Kiel.
- Man, K. (2003) Long Memory Time Series and Short Term Forecasts, *International Journal of Forecasting* 19, 477-491.
- Mandelbrot, B. (1974) Intermittent Turbulence in Self Similar Cascades: Divergence of High Moments and Dimension of the Carrier, *Journal of Fluid Mechanics* 62, 331 – 358.
- Mandelbrot, B., A. Fisher and L. Calvet (1997) *A Multifractal Model of Asset Returns*. Cowles Foundation Discussion Paper no. 1164, Yale University.

- Neely, C. and P. Weller (2002) Predicting Exchange Rate Volatility: Genetic Programming versus GARCH and RiskMetrics, *Federal Reserve Bank of St. Louis Review*, May/June, 43 – 54
- Poon, S.-H. and C. Granger (2003) Forecasting Volatility in Financial Markets: A Review, *Journal of Economic Literature* 41, 468 - 539.
- Ray, B., S. Chen and J. Jarrett (1997) Identifying Permanent and Temporary Components in Daily and Monthly Japanese Stock Prices, *Financial Engineering and the Japanese Markets* 4, 233-256.
- Ray, B. and R. Tsay (2000) Long-range Dependence in Daily Stock Volatilities, *Journal of Business and Economics Statistics* 18, 254 – 262.
- Ryden, T., T. Teräsvirta and S. Asbrink (1998) Stylized Facts of Daily Return Time Series and the Hidden Markov Model, *Journal of Applied Econometrics* 13, 217-244.
- Tse, Y. (1991) Stock Returns Volatility in the Tokyo Stock Exchange, *Japan and the World Economy* 3, 285-298.
- Vilasuso, J. (2002) Forecasting Exchange Rate Volatility, *Economics Letters* 76, 59 – 64.
- West, K. and D. Cho (1995) The Predictive Ability of Several Models of Exchange Rate Volatility, *Journal of Econometrics* 69, 367-391.
- Zumbach, G. (2004) Volatility Processes and Volatility Forecast with Long Memory, *Quantitative Finance* 4, 70 – 86.

Appendix: Stocks used in the analysis

The 100 stocks with the highest average trading volume are (stock identification numbers in parentheses):

Teikoku Oil (1601), Taisei Corp. (1801), Obayashi Corp. (1802), AC Real Estate Corp. (1806), Kajima Corp. (1812), Kumagai Gumi (1861), Aoki Corp. (1886), Daiwa House Industry (1925), Kirin Brewery (2503), Toyobo (3101), Unitika (3103), Teijin (3401), Toray Industries (3402), Mitsubishi Rayon (3404), Asahi Kasei Corp. (3407), Sangoku Pulp (3702, until 07/92), Oji Paper (3861), Mitsui Toatsu Chemical (4001, until 12/96), Showa Denko K.K. (4004), Sumitomo Chemical (4005), Mitsubishi Chemical Corp. (4010), Ishihara Sangyo Kaisha (4028), Tosoh Corp. (4042), Denki Kagaku Kogyo Kabushiki Kai (4061), Ube Industries (4208), Takeda Chemical Industries (4502), Dainippon Ink and Chemicals (4631), Fuji Photo Film (4901), Nippon Oil Corp. (5001), Mitsubishi Oil (5004, until 03/99), Cosmo Oil (5007), Nippon Sheet Glass (5202), Taiheiyo Cement Corp. (5233), Nippon Steel Corp. (5401), Kawasaki Steel (5403), NKK Corp. (5404), Sumitomo Metal Industries (5405), Kobe Steel (5406), Nisshin Steel (5407), The Japan Steel Works (5631), Nippon Light Metal Company (5701), Mitsui Mining and Smelting (5706), Mitsubishi Materials Corp. (5711), Nippon Mining (5712, until 07/91), Sumitomo Metal Mining (5713), Dowa Mining (5714), The Furukawa Electric (5801), Sumitomo Electric Industries (5802), Fujikura (5803), Komatsu (6301), Sumitomo Heavy Industries (6302), Kubota Corp. (6326), Hitachi (6501), Toshiba Corp. (6502), Mitsubishi Electric Corp. (6503), Fuji Electric (6504), Nec Corp. (6701), Fujitsu (6702), Oki Electric Industry (6703), Matsushita Electric Industrial (6752), Sharp Corp. (6753), Sony Corp. (6758), Sanyo Electric (6764), Mitsui Engineering & Shipbuilding (7003), Hitachi Zosen Corp. (7004), Mitsubishi Heavy Industries (7011), Kawasaki Heavy Industries (7012), Ishikawajima-Harima Heavy Indust. (7013), Nissan Motor (7201), Isuzu Motors (7202), Toyota Motor Corp. (7203), Mazda Motor Corp. (7261), Honda Motor (7267), Fuji Heavy Industries (7270), Canon (7751), Ricoh Company (7752), Itochu Corp. (8001), Marubeni Corp. (8002), Mitsui & co. Ltd. (8031), Mitsubishi Corp. (8058), Nissho Iwai Corp. (8063), Sakura Bank (8314, until 03/2001), Bank of Tokyo-Mitsubishi (8315, until 03/2001), Sumitomo Bank (8318), Asahi Bank (8322),

Daiwa Securities Group Inc. (8601), Yamaichi (8602, until 11/97), Nikko Cordial Corp. (8603), Nomura Holdings (8604), Tokio Marine & Fire Ins. (8751), Mitsui Fudosan (8801), Mitsubishi Estate (8802), Tobu Railway (9001), Tokyu Corp. (9005), Keisei Electric Railway (9009), Mitsui O.S.K. Lines (9104), Kawasaki Kisen Kaisha (9107), The Tokyo Electric Power (9501), Tokyo Gas (9531), Osaka Gas (9532).

The ‘random sample’ consisted of the following stocks which were randomly chosen from our data base (in parentheses: stock identification number):

Nippon Suisan Kaisha (1332), Hoko Fishing (1351), Fudo Construction (1813), Tekken Corp. (1815), Nakano Corp. (1827), Toda Corp. (1860), Penta-Ocean Construction (1893), Obayashi Road Corp. (1896), Daiwa House Industry (1925), Nippon Koei (1954), Morinaga (2201), Nippon Meat Packers (2282), Itoham Foods (2284), Nichirei Corp. (2871), Kawashima Textile Manufacturers (3009), Nisshinbo Industries (3105), Ashimori Industries (3526), Showa Denko K.K. (4004), Nippon Carbide Industries (4064), Sakai Chemical Industries (4078, until 09/97), Mitsui Chemicals (4183), JSR Corp. (4185), Nippon Kayaku (4272), Sankyo (4501), Yamanouchi Pharmaceutical (4503), Daiichi Pharmaceutical (4505), Kansai Paint (4613), Tohpe Corp. (4614), Chugoku Marine Paints (4617), Toyo Ink Mfg. (4634), Takasago International Corp. (4914), Toyo Tire & Rubber (5105), Osaka Cement (5235, until 12/93), Nippon Hume Corporation (5262), Nippon Yakin Kogyo (5480), Nippon Denko (5563), Suzuki Metal Industries (5657), Nihon Seiko (5729), Sakurada (5917), Amada Machines (6107), Koike Sanso Kogyo (6137), Kioritz Corp. (6313), Meiji Machines (6334), Sintokogio (6339), Sumitomo Precision Products (6355), Nippon Gear (6356), Sakai Heavy Industries (6358), Toyo Kanetsu K.K. (6369), Tsubakimoto Chain Co. (6371), Fuso Lexel (6386), Sanjo Machine Works (6437), Riken Corp. (6462), Koyo Seiko (6473), Yashakaw Electric Corp. (6506), Makita Corp. (6586), Nishishiba Electric (6591), Kawaden (6648, until 09/2000), The Nippon Signal (6741), Nec Tokin Corp. (6759), Kawasaki Heavy Industries (7012), Shin Maywa Industries (7224), Tokyo Radiator Mfg. (7235), Daihatsu Motor (7262), Oval Corp. (7727), Riken Keiki (7734), Chinon Industries (7738), Pentax Corp. (7750), Canon (7751), Dantani Corp. (7910), Yamaha Corp. (7951), Takara Standard (7981), Daiwa Seiko (7990), Kanematsu Corp. (8020), Tohto Suisan (8038), Tsukiji Uoichiba (8039), Seiko Corp. (8050), Shoko (8090), Inabata (8098), GSI Creos Corp. (8101), Sinanen (8132), Matsuzakaya (8235), Maruzen (8236), Bank of Yokohama (8332), Gunma Bank (8334), Musashino Bank (8336), Hyakugo Bank (8368), Kiyo Bank (8370), Iyo Bank (8385), Oita Bank (8392), Sumitomo Insurance (8753, until 09/2001), Nipponkoa Insurance (8754), Sompo Japan Insurance Inc. (8755), Nissan Fire & Marine Ins. (8756), Nissay Dowa General Insurance (8759), Nichido Fire & Marine Ins. (8760), Taiheiyo Kaiun (9123), KDD (9431, until 09/2000), Chugoku Electric Power (9504), Toho Gas (9533), Shochiku (9601).