

Chapter 2

VaR, MS and Extreme Quantile Estimation

From equations (1.1) and (1.3) in Chapter 1, we observe that VaR and MS of an asset turn out to be extreme quantiles of the marginal distribution of the asset return time series. In this chapter, we study the accuracy of several nonparametric and extreme value theory based estimators of extreme quantiles and compare their known properties. We compare their finite sample performance using Monte Carlo simulation. A new quantile estimator is proposed which exhibits encouraging finite sample performance while estimating extreme quantile in the right tail region.

2.1 Nonparametric quantile estimators

In this section we review some known nonparametric and extreme value theory based quantile estimators.

2.1.1 Sample & kernel quantile estimators

A natural estimator of Q_p is obtained as follows

$$\hat{Q}_p = \inf\{x : \hat{F}(x) \geq 1 - p\},$$

where $\hat{F}$ is a distribution function estimator of F based on $X_1, \dots, X_n$. Let $X_{(1)}, \dots, X_{(n)}$ denote the corresponding order statistics. Also let $I(\cdot)$ be the indicator function, with $I(S)$ equal to 0 or 1 according as the statement S is false or true. If $\hat{F}(x) = \frac{1}{n} \sum_{i=1}^n I(X_i \leq x)$ i.e. $\hat{F}$ is the empirical distribution function, $\hat{Q}_p$ equals $X_{([n(1-p)]+1)}$, where $[x]$ denotes the integral part of x . It is the $(1-p)$ th sample quantile (we call it SQ_p). Asymptotic properties of the sample quantile are well known under i.i.d. assumption (see [92]). Recently, asymptotic properties of the SQ_p has been studied extensively under various dependence assumptions.

(See for instance, [99], [106], [68], [104] among most recent). Under strong mixing dependence (with polynomial mixing rate) assumption, Wang et al. [104] obtained Bahadur representation of sample quantile which provides insight into the rate of strong convergence of the SQ_p under this dependence assumption. Under similar dependence assumption, Lahiri and Sun [68] obtained an upper bound of the accuracy of normal approximation to the sampling distribution of the SQ_p . These results provide insight into the accuracy of the SQ_p under strong mixing type dependence assumption. These asymptotic properties are obtained under the condition that p is kept fixed, as n is increased. The following Theorem provides the necessary conditions for strong convergence of the sample quantile SQ_p , under the condition that $p = o(1)$ as $n \rightarrow \infty$.

Theorem 2.1.1. *Let $X_1, \dots, X_n$ be i.i.d. random variables with a continuous density f satisfying $f(x) > 0 \forall x$ and $f(x) \rightarrow 0$ as $|x| \rightarrow \infty$. For every $\delta > 0$, $\exists x_0 > 0$ such that $\left| \frac{f(x \pm y)}{f(x)} - 1 \right| < \delta$ for all $x > x_0$ and $0 < y < 1$. Further let*

$$p = o(1) \text{ and } \epsilon_n = \frac{\sqrt{2 \log(n)}}{\sqrt{n} f(Q_p)} = o(1) \text{ as } n \rightarrow \infty.$$

Then as $n \rightarrow \infty$

$$SQ_p - Q_p = O(\epsilon_n), \text{ a.s.}$$

Proof. Under the conditions stated in Theorem 2.1.1, the distribution function F is continuous and strictly increasing, $0 < F(x) < 1 \forall x$ and $F(x) \rightarrow 1$, as $x \rightarrow \infty$. Therefore Q_p is the unique solution of $F(x-) < 1 - p \leq F(x) \forall 0 < p < 1$ and $Q_p \rightarrow \infty$ as $p = o(1)$.

Then by Theorem 2.3.2 in page 75, Serfling(1980), we see that for every $\epsilon > 0$, there exists a $\delta = \min\{F(Q_p + \epsilon) - q, q - F(Q_p - \epsilon)\}$, where $q = 1 - p$, such that

$$P(|SQ_p - Q_p| > \epsilon) \leq 2 \exp(-2n\delta^2). \quad (2.1)$$

The above inequality holds for every n and $0 < p < 1$. Therefore even under the condition $p = o(1)$ as $n \rightarrow \infty$, the above inequality continues to hold for each fixed n . Under the stated conditions $F(Q_p) = q$. Therefore, under the conditions on f , F is differentiable and we have

$$\delta = \epsilon \min\{f(Q_p + \theta\epsilon), f(Q_p - \theta'\epsilon)\}, \quad 0 < \theta, \theta' < 1.$$

Let us define, $\epsilon_n = \frac{\sqrt{2 \log(n)}}{\sqrt{n} f(Q_p)}$. Under the stated conditions, $\epsilon_n > 0$, $\forall n$ and $\epsilon_n = o(1)$ as $n \rightarrow \infty$. We now apply the inequality (2.1) with ϵ_n and under the stated conditions on f the corresponding δ_n satisfies

$$\delta_n = \frac{\sqrt{2 \log(n)}}{\sqrt{n} f(Q_p)} \min\{f(Q_p + \theta\epsilon_n), f(Q_p - \theta'\epsilon_n)\} \geq \frac{\sqrt{\log(n)}}{\sqrt{n}}, \text{ for all sufficiently large } n.$$

Therefore from the inequality (2.1),

$$P(|SQ_p - Q_p| > \epsilon_n) \leq \frac{2}{n^2}, \text{ for all sufficiently large } n.$$

Now by Borel Cantelli Lemma, $P(|SQ_p - Q_p| > \epsilon_n \text{ i.o.}) = 0$. Therefore there exists an event A with $P(A) = 1$, such that for each $\omega \in A$, the sequence $\frac{|SQ_p(\omega) - Q_p|}{\epsilon_n}$ is bounded. This completes the proof. $\square$

Remark 1. 1. If p is kept fixed, then the condition $\epsilon_n = o(1)$ is trivial and Theorem 2.1.1 essentially reduces to the Lemma B in page 96, in Serfling [92].

2. If $p = o(1)$ as $n \rightarrow \infty$, the stated conditions on F and f ensure that $Q_p \rightarrow \infty$ and $0 < f(Q_p) = o(1)$ as $n \rightarrow \infty$. The condition $\epsilon_n = o(1)$, ensures that $f(Q_p)$ does not converge to zero too fast as $p \rightarrow 0$ with increase in the sample size n .

3. If the marginal distribution is standard GPD with location parameter $\mu = 0$, scale $\sigma = 1$, shape parameter ξ , then $f(Q_p) = p^{\xi+1}$ and Theorem 2.1.1 holds for $p = o(1)$, provided $\sqrt{\frac{n}{\log(n)}} p^{\xi+1} \rightarrow \infty$ as $n \rightarrow \infty$.

A wide variety of other nonparametric distribution function estimators are available in the literature. See for instance [34] for a detailed review and comparison of these estimators. Using these distribution function estimators in $\hat{Q}_p$ we get different versions of $\hat{Q}_p$. For instance, one can use $\hat{F}(x)$ equal to a kernel based distribution function estimator defined as follows

$$\hat{F}(x) = \frac{1}{nb} \sum_{i=1}^n \int_{-\infty}^x w\left(\frac{t - X_i}{b}\right) dt.$$

where $b > 0$ is the bandwidth and w is a probability density function with zero mean and finite variance, known as the kernel. b depends on n and $b \rightarrow 0$ as $n \rightarrow \infty$. In the kernel based method the main problem lies with the selection of bandwidth. Azzalini, Bowman and Chen and Tang provide some choice of the bandwidth parameter (see [7], [16], [23]). Chen and Tang have obtained the asymptotic bias, variance and the rate of almost sure convergence of their version of $\hat{Q}_p$, under the assumption that $\{X_t\}$ is a stationary geometric α -mixing process (see [23]). The authors suggested the following choice for the optimal value of b ,

$$b_{opt1} = \left\{ \frac{2f^3(Q_p)b}{\sigma^4(f^{(1)}(Q_p))^2} \right\}^{1/3} n^{-1/3},$$

where $b = \int uw(u)G(u)du$, and $\sigma^2 = \int u^2w(u)du$. $G(\cdot)$ is the distribution function of the distribution with density w . b_{opt1} involves unknown constants Q_p , f and its derivative $f^{(1)}$ at Q_p . Chen and Tang suggested to approximate Q_p in b_{opt1} by the corresponding sample quantile (see [23]). The authors suggested to approximate f and $f^{(1)}$ by the density and

the first derivative of the Generalized Pareto distribution. We denote the Chen and Tang's quantile estimator by C-T_p.

Alemaný et al. proposed another bandwidth suitable for kernel based estimation of Var_p for p close to 0, using Epanechnikov kernel (see [3], proposition 2). Their proposal was based on minimizing a weighted mean integrated squared error $WISE = E\{\int [\hat{F}(x) - F(x)]^2 x^2 dx\}$ which assigns more weight in the tail region than the ordinary mean integrated squared error. The minimization of this criterion leads to the following optimal bandwidth

$$b_{opt2} = \left(\frac{E(X_1^2) \int G(u)(1 - G(u))du}{\int (f^{(1)}(x))^2 x^2 dx (\int t^2 w(t) dt)^2} \right)^{1/3} n^{-1/3}.$$

The authors suggested to compute the unknown functionals of f by assuming that f is a normal density with mean 0 and variance σ^2 . This leads to $b_{opt2} = \sigma^{5/3}(8/3)^{1/3}n^{-1/3}$. σ is estimated from the data. Let Al_p denote the quantile estimator by Alemany et al. (see [3]).

Remark 2. *The bandwidth proposed by Alemany et al. was obtained under i.i.d. assumption (see [3]). It remains to be seen how the resulting quantile estimator performs in the presence of α -mixing type dependence.*

Empirical versus kernel estimator

1. Under the assumption that $\{X_t\}$ is a stationary geometric α -mixing process, Chen and Tang [23] proved that as $n \rightarrow \infty$

$$C-T_p - Q_p = o(n^{-1/2} \log(n)), \text{ a.s.},$$

Under similar assumptions, using the Bahadur type representation of SQ_p in Wang et al. and a Bernstein type inequality for strongly mixing processes in Merlevede et al. we see that ([104], [79])

$$SQ_p - Q_p = o(n^{-1/2}(\log(n))^{3/4}), \text{ a.s..}$$

Therefore under similar conditions the sample quantile seem to possess faster rate of strong convergence, as n is increased.

2. The sampling distributions of both sample quantile and the kernel quantile estimators can be approximated by normal distributions. However, the following result is known only for the sample quantile. Under the assumption that $\{X_t\}$ is a stationary α -mixing process with $\alpha(n) < \frac{c}{n^d}$ and $c > 1$, $d > 12$, Lahiri and Sun [68] proved that there exists a constant $C \geq 1$ such that for all $n \geq 1$,

$$\sup_{x \in \mathbb{R}} |P(\sqrt{n}(SQ_p - Q_p) \leq x) - \Psi(x)| \leq \frac{C}{\sqrt{n}}. \quad (2.2)$$

This result provides insight into the accuracy of the normal approximation to the sampling distribution of SQ_p under the stated dependence assumption. Such a result does not seem to be known for the kernel based quantile estimators.

3. The kernel method depends crucially on the choice of the bandwidth b . Under strongly mixing dependence assumption Chen and Tang obtained asymptotically optimal value of b , which depends on unknown constants (see [23]). So the optimal b has to be again estimated from the data. Even with this optimal choice of b the difference in accuracy between the kernel based estimator and SQ_p can be quite small. For example, from the simulation study in Chen and Tang we see that for the ARCH(1) model and $p = 0.99$, the improvement in the standard deviation and the root mean squared error of their kernel quantile estimator is less than three percent the same for the SQ_p (see [23]).

2.1.2 L -estimator

The sample quantile $X_{([n(1-p)]+1)}$ is a natural estimator of the population quantile. But it is effected by the variability of individual order statistics. An obvious way of improving the efficiency of sample quantiles is to reduce this variability by forming a weighted average of all the order statistics, using an appropriate weight function (see [95]). Such an estimator is commonly called an L -estimator. A popular class of L -estimators is a kernel quantile estimator defined as follows ([95])

$$\hat{Q}_L = \sum_{i=1}^n \left[\int_{\frac{i-1}{n}}^{\frac{i}{n}} w\left(\frac{t-p}{b}\right) dt \right] X_{(i)},$$

where w is a density function called the kernel, $b > 0$ and $b \rightarrow 0$ as $n \rightarrow \infty$. b is called the bandwidth. Sheather and Marron provide a detailed theoretical analysis of the asymptotic properties of $\hat{Q}_L$ and a data based method for choice of b (which appears to be very complicated) (see [95]). The authors conclude that one can expect only modest improvement (upto 15 percent) over the sample quantile, even with the best possible L -estimator. Given this limited improvement, the effort involved in data based choice of b and the contrasting ease with which one can compute $X_{([n(1-p)]+1)}$ and its known asymptotic properties, the later seems to be a more reasonable choice.

2.1.3 Harrell-Davis (1982)

Harrell-Davis [52] introduced a quantile estimator (we call it H-D estimator) which is a weighted linear combination of order statistics defined as follows

$$\begin{aligned} \text{H-D}_p &= \sum_{i=1}^n w_i X_{(i)} \\ w_i &= I_{i/n}((1-p)(n+1), p(n+1)) - I_{(i-1)/n}((1-p)(n+1), p(n+1)), \quad i = 1, \dots, n, \end{aligned}$$

where $I_x(a, b)$ denotes incomplete beta function.

Based on simulation study Harrel-Davis suggested that their estimator has much to offer over the sample quantile especially for extreme quantiles (see page 639, [52]). The H-D estimator is the limit of a bootstrap average as the number of bootstrap resamples becomes infinitely large. It is available in R software (see `hdquantile` function in `Hmisc` package in R software for statistical computing).

2.1.4 SV estimators

Sfakianakis and Verginis introduced three L -statistics type estimators, $SV1_p$, $SV2_p$ and $SV3_p$ (see [93]). Among these estimators $SV3_p$ seems to be the appropriate estimator for Q_p , especially for p close to zero. It is defined as follows

$$SV3_p = \sum_{i=1}^n B(i, n, 1-p) X_{(i)} + (2X_{(1)} - X_{(2)}) B(0, n, 1-p),$$

where $B(i, n, 1-p)$ is the probability mass function of the Binomial distribution with parameters n and $1-p$.

2.1.5 Quantile estimation based on Extreme Value Theory (EVT)

In this approach the idea is to let the tails speak for themselves, that is, use merely the largest returns for the estimation of the extreme quantiles (see [33]). Estimation of quantiles for values of $1-p$ close to 1 by extreme value theory is related to Pickands-Balkema-de Haan theorem (see [10]). Pickands-Balkema-de Haan theorem claims that if F is in the domain of attraction of the Generalized Extreme Value (GEV) (we denote it by $F \in D(\text{GEV})$), the conditional distribution of $X_1 - u$, given that $X_1 > u$, can be well approximated by Generalized Pareto distribution (GPD) with tail parameter ξ and with some shape parameter $\beta(u)$, for u large enough (see [21]). Based on this theorem, a GPD distribution fitted to the k largest observations in the sample to approximate the tail of the conditional loss distribution, given that the loss exceeds some threshold value. Let $\hat{\xi}$ and $\hat{\beta}$ denote the

maximum likelihood estimates of the GPD (with threshold $X_{(n-k)}$) based on $\{X_{(n-k+1)} - X_{(n-k)}, \dots, X_{(n)} - X_{(n-k)}\}$. Under the assumption that $\xi > 0$, the $(1-p)$ th quantile of the loss distribution is approximated by the $(1-p)$ th quantile of the fitted GPD distribution. The resulting estimator is given by

$$EVT1_p = X_{(n-k)} + \frac{\hat{\beta}_k}{\hat{\xi}_k} \left(\left[\frac{n}{k}(1-p) \right]^{-\hat{\xi}_k} - 1 \right), \quad (2.3)$$

For small k (i.e. for large threshold), the GPD approximation of the tail of F is more accurate (assuming $F \in D(\text{GEV})$), but lesser observations in the sample are available for fitting the GPD. In contrast for large k , more data are available for fitting the GPD distribution, but the GPD approximation to the tail of F is biased. Consequently, an important issue in this approach is the choice of k (see [56]). Usually, k is a function of n , satisfying $k \rightarrow \infty$ and $\frac{k}{n} = o(1)$, as $n \rightarrow \infty$. From extensive simulation we find that $k = [np] + 1$ works well for i.i.d. as well as GARCH(1,1) time series model, especially for p close to zero ($k = [np] + 1$ satisfies the condition $\frac{k}{n} = o(1)$, provided $p \rightarrow 0$ as $n \rightarrow \infty$).

Under i.i.d. assumption the condition that F is in the domain of attraction of the GEV is equivalent to the condition that there exists $a_n > 0$, $b_n \in \mathbb{R}$ such that $\frac{X_{(n)} - b_n}{a_n}$ converges in law to a GEV distribution, with parameters ξ and β , as n is increased. Drees has extended the extreme value theory based estimation of extreme quantiles in the presence of β -mixing¹ type dependence (which cover a broad class of time series models, including those considered in this chapter) (see [33]). The author assumed that the common distribution function F satisfies the property that as $\lambda \rightarrow 0$

$$\frac{F^{-1}(1 - \lambda t)}{F^{-1}(1 - \lambda)} \rightarrow \frac{1}{t^\xi}, \quad t > 0,$$

for some $\xi > 0$. Under this assumption one can argue that for small (positive) ξ

$$Q_p \equiv Q_{k_n/n} \left(\frac{k_n}{np} \right)^\xi,$$

where $c \equiv d$ means the ratio c/d is close to one. Above approximation naturally leads to the following estimator

$$EVT2_p = X_{(n-k_n)} \left(\frac{k_n}{np} \right)^{\hat{\gamma}_n}, \quad (2.4)$$

¹The series $\{X_t\}_{t \in N}$ is said to be β -mixing if

$$\beta(l) := \sup_{m \in N} E \left(\sup_{A \in \mathcal{B}_{m+l+1}^\infty} |P(A|\mathcal{B}_1^m) - P(A)| \right) \rightarrow 0$$

as $l \rightarrow \infty$, where $\mathcal{B}_1^m$ and $\mathcal{B}_{m+l+1}^\infty$ denote the σ -fields generated by $\{X_t, 1 \leq t \leq m\}$ and $\{X_t, m+l+1 \leq t\}$ (see [33]).

here $1 \leq k_n < n$ and $k_n \rightarrow \infty$, $\frac{k_n}{n} \rightarrow 0$, as $n \rightarrow \infty$. $\hat{\gamma}_n$ is a suitable estimator of the tail index $\hat{\gamma}$, say the Hill estimator

$$\hat{\gamma}_n = \frac{1}{k_n} \sum_{i=1}^{k_n} \log \frac{X_{(n-i+1)}}{X_{(n-k_n)}}.$$

Drees have studied extreme quantile estimation in the presence of β -mixing type dependence (see [33]). For instance, under the assumption that as $n \rightarrow \infty$, $p \rightarrow 0$ in such a way that $\frac{\log(n(1-p))}{\sqrt{k_n}} \rightarrow 0$, $\frac{n(1-p)}{k_n} \rightarrow 0$,

$$\frac{\sqrt{k_n}}{\log(k_n/np)} \left(\frac{EVT2_p}{Q_p} - 1 \right) \xrightarrow{d} N(0, \sigma_{T,\gamma}^2),$$

with $\sigma_{T,\gamma}^2$ as defined in Drees [33]. From Remark 2.5, in page 627 in Drees [33] we see that for a broad class of β -mixing processes and $k = [np] + 1$ his EVT estimator $EVT2_p$ is a consistent estimator with relative estimation error of order $\frac{1}{\sqrt{np}}$, where $p \rightarrow 0$ and $np \rightarrow \infty$ as n is increased.

2.1.6 Quantile estimation based on transformation

In this approach we first construct a quantile estimator based on the transformed data $Y_i = T(X_i)$, $i = 1, \dots, n$, where $T : \mathbb{R} \rightarrow [0, 1]$ is a monotonic increasing invertible function. The quantile estimator based on $X_1, \dots, X_n$ is obtained by back transform, i.e. $\hat{Q}_p(X_1, \dots, X_n) = T^{-1}(\hat{Q}_p(Y_1, \dots, Y_n))$. This is due to the fact that if Q_p is the $(1-p)$ th quantile of X_1 , $T(Q_p)$ is the $(1-p)$ th quantile of Y_1 and vice versa. T can be chosen to be a continuous distribution function estimator based on the original data or the distribution function of a suitable continuous distribution fitted to the data $X_1, \dots, X_n$.

Kernel based transform

Swanepoel and Grann introduced the idea of distribution function estimation based on a nonparametric transformation of the data (see [100]). Their suggestion was based on the fact that if $X_1, \dots, X_n$ are identically distributed with distribution function F , with density f then

$$S_n(x) = \frac{1}{n} \sum_{i=1}^n K \left(\frac{F(x) - F(X_i)}{h} \right),$$

is an unbiased estimator of $F(x)$, where K is a known distribution function with a symmetric density supported on $[-1, 1]$. For proof, see the calculations following equation (25) in page 560 in Swanepoel and Grann (see [100]). This result implies that the bias of a kernel based distribution function estimator can be eliminated by transforming the data. In practice F is

unknown, and the authors suggested to replace F in S_n by another kernel based distribution function estimator with smoothing parameter say g . This leads to the following distribution function estimator

$$\tilde{F}_n(x) = \frac{1}{n} \sum_{i=1}^n K \left(\frac{F_{n,g}(x) - F_{n,g}(X_i)}{h} \right),$$

where $F_{n,g}(x) = \frac{1}{n} \sum_{i=1}^n K \left(\frac{x-X_i}{g} \right)$, $g = ch^\alpha$, $1 \leq \alpha < 3$. In Theorem 1, page 554, Swanepoel and Grann [100] obtained the asymptotic bias and variance of $\tilde{F}_n(x)$. See equations (14) and (15) in Swanepoel and Grann ([100]). Swanepoel and Grann suggested to use

$$g = h = \left[\frac{375\sqrt{3}}{28\pi} \right]^{1/7} \sigma^{-4/7} n^{-1/7},$$

where $\sigma = \min\{S, IQR/1.349\}$, S and IQR are the sample standard deviation and inter quartile range respectively (see [100]). The authors claim that using such choice of g and h , and Epanechnikov kernel considerable bias and MISE (mean integrated squared error) reduction are achieved.

We hope that an improved distribution function estimator can provide better quantile estimate. Hence we define a kernel estimator based on Swanepoel and Grann's distribution function estimator as follows

$$S-G_p = \inf\{x : \tilde{F}_n(x) \geq 1 - p\}. \quad (2.5)$$

where $\tilde{F}_n(x)$ is the Swanepoel and Grann's distribution function estimator defined above. First we state an asymptotic property of the estimator $\tilde{F}_n$.

Lemma 2.1.2. (Dutta [37]) *Let $X_1, \dots, X_n$ be i.i.d. random variables with distribution function F and density f which has a bounded derivative. The kernel distribution function K is differentiable, with a bounded kernel density k with zero mean and finite variance. The bandwidth sequences g, h satisfies that $g = h = o(1)$ as $n \rightarrow \infty$. Let $b_n = o(1)$ such that $nb_n^2 \rightarrow \infty$, as $n \rightarrow \infty$ and $\sum_{n=1}^{\infty} \exp(-\frac{1}{4\|k\|} nb_n^2 h^2) < \infty$, as $n \rightarrow \infty$, then*

$$\sum_{n=1}^{\infty} P \left[\|\tilde{F}_n - F\| > b_n \right] < \infty.$$

Repeating the arguments similar to those used in the proof of the Theorem and the above Lemma we get the following asymptotic property of the S- G_p estimator.

Theorem 2.1.3. *Let $X_1, \dots, X_n$ be i.i.d. random variables with a continuous density f satisfying $f(x) > 0 \forall x$ and $f(x) \rightarrow 0$ as $|x| \rightarrow \infty$. For every $\delta > 0$, $\exists x_0 > 0$ such that*

$\left| \frac{f(x \pm y)}{f(x)} - 1 \right| < \delta$ for all $x > x_0$ and $0 < y < 1$. The kernel distribution function K is differentiable, with a bounded kernel density k with zero mean and finite variance. Further let $g = h = Cn^{-1/7}$, $C > 0$ and

$$p = o(1) \text{ and } \epsilon_n = \frac{\sqrt{2 \log(n)}}{\sqrt{n} f(Q_p) h} = \frac{\sqrt{2 \log(n)}}{n^{5/14} f(Q_p)} = o(1) \text{ as } n \rightarrow \infty.$$

Then as $n \rightarrow \infty$

$$|S - G_p - Q_p| = O(\epsilon_n), \text{ a.s..}$$

Proof.

$$\begin{aligned} P(S - G_p - Q_p > \epsilon) &= P(1 - p > \tilde{F}_n(Q_p + \epsilon)) \\ &= P(F(Q_p + \epsilon) - \tilde{F}_n(Q_p + \epsilon) > F(Q_p + \epsilon) - (1 - p)) \\ &\leq P(\|\tilde{F}_n - F\| > \delta_1), \text{ where} \\ \delta_1 &= F(Q_p + \epsilon) - (1 - p) = \epsilon f(Q_p + \theta\epsilon), \quad 0 < \theta < 1. \end{aligned}$$

And similarly

$$P(Q_p - S - G_p > \epsilon) \leq P(\|\tilde{F}_n - F\| > \delta_2), \quad \delta_2 = (1 - p) - F(Q_p - \epsilon) = \epsilon f(Q_p - \theta'\epsilon), \quad 0 < \theta' < 1$$

Therefore

$$P(|Q_p - S - G_p| > \epsilon) \leq P(\|\tilde{F}_n - F\| > \delta), \quad \delta = \min\{\delta_1, \delta_2\}. \quad (2.6)$$

Now let $\epsilon = \epsilon_n = \frac{2\sqrt{\log(n)}}{\sqrt{n} f(Q_p) h}$. Under the stated conditions on f , the corresponding δ_n satisfies the following inequality.

$$\delta_n = \frac{2\sqrt{\log(n)}}{\sqrt{n} f(Q_p) h} \min\{f(Q_p + \theta\epsilon), f(Q_p - \theta'\epsilon)\} \geq \frac{\sqrt{2 \log(n)}}{h\sqrt{n}}, \text{ for all sufficiently large } n.$$

We now apply Lemma 2.1.2, with $g = h = Cn^{-1/7}$, $C > 0$, and $b_n = \frac{\sqrt{2 \log(n)}}{\sqrt{n} h}$. Then $b = o(1)$ and the condition $\sum_{n=1}^{\infty} \exp(-\frac{1}{4\|k\|} n b_n^2 h^2) < \infty$ is satisfied. And therefore under these conditions

$$\sum_{n=1}^{\infty} P(\|\tilde{F}_n - F\| > \delta_n) < \infty. \quad (2.7)$$

And therefore under the stated conditions on f, K, p , and assuming $g = h = Cn^{-1/7}$, $C > 0$, and $\epsilon_n = \frac{2\sqrt{\log(n)}}{\sqrt{n} f(Q_p) h} = o(1)$, as $n \rightarrow \infty$, using the equations (2.6) and (2.7), we see that

$$\sum_{n=1}^{\infty} P(|Q_p - S - G_p| > \epsilon_n) < \infty.$$

Therefore under these conditions

$$|Q_p - S-G_p| = O(\epsilon_n), \text{ a.s.}$$

□

2.2 Simulation & real data analysis

We compare the performance of eight nonparametric quantile estimators, viz. SQ_p , $H-D_p$, $SV3_p$, the kernel based estimators $C-T_p$, Al_p , $S-G_p$ and the two EVT based estimators $EVT1_p$ and $EVT2_p$ for different sample size n and different choices of p . It is difficult to obtain the exact bias and the MSE of these estimators. Therefore we use Monte Carlo simulation to approximate the bias, the standard deviation and the MSE of each of these estimators.

To approximate the bias or the MSE of a statistic T_n using Monte Carlo simulation we draw m random samples of size n from a test distribution or stochastic process. From each of the m samples we compute the value of the statistic. Let T_{ni}^* , $i = 1, \dots, m$, be the values. The bias, variance and the MSE of T_n are approximated by $\frac{1}{m} \sum_{i=1}^m T_{ni}^* - T_n$, variance of T_{ni}^* , $i = 1, \dots, m$, and $\frac{1}{m} \sum_{i=1}^m (T_{ni}^* - T_n)^2$ respectively. In this simulation study we use $m = 1000$.

In general the stochastic process generating the observed data is not known. However in a Monte Carlo simulation study we can compute the Monte Carlo estimate assuming some test distribution or data generating process. In this simulation study we consider the following ten time series models

- (i) $\{X_i\}_{i=1,2,\dots}$ is an i.i.d. process, marginal distribution GPD with $\xi = 1/3$.
- (ii) $\{X_i\}_{i=1,2,\dots}$ is an i.i.d. process, marginal distribution Student's with 4 df.
- (iii) $\{X_i\}_{i=1,2,\dots}$ is an i.i.d. process, marginal distribution $N(0, 1)$.

To study the effect of dependence on the above mentioned quantile estimators consider the following ARMA(1,1) models in Drees [33]

$$X_i - \phi X_{i-1} = Z_i + \theta Z_{i-1},$$

$$(iv) \phi = 0.95, \theta = -0.6,$$

$$(v) \phi = 0.95, \theta = -0.9,$$

$$(vi) \phi = 0.3, \theta = 0.9.$$

In addition the following GARCH(1,1) models are also considered

$$\begin{aligned}
X_t &= \sigma_t Z_t, \\
(vii) \quad \sigma_t^2 &= 0.0001 + 0.9X_{t-1}^2, \\
(viii) \quad \sigma_t^2 &= 0.0001 + 0.4X_{t-1}^2 + 0.5\sigma_{t-1}^2, \\
(ix) \quad \sigma_t^2 &= 0.0751X_{t-1}^2 + 0.9194\sigma_{t-1}^2.
\end{aligned}$$

The first two models are motivated by empirical observations by Cont regarding the extent of tail heaviness of the marginal asset return distributions (see [26]). Cont mentioned that when sample moments based on asset return data are plotted against sample size, the sample variance seems to stabilize with increase in sample size (see [26]). But the behavior of the fourth order sample moment seems to be erratic as n is increased. This feature is also exhibited by the sample moments based on i.i.d. draws from the Student's t-distribution with four degrees of freedom, which displays a tail behavior similar to many asset return distributions. Cont also mentioned that the daily return distributions of stocks, market indices and exchange rates seem to exhibit power law tail with exponent α satisfying, $\xi = 1/\alpha$ varying between 0.2 and 0.4 (see [26]). These observations motivate choice of the marginal distributions in (i) and (ii). The third model (iii) represents the classical Black-Scholes assumption on the return model.

The GARCH(1,1) processes are known to model the volatility clustering observed in financial time series data. The first two GARCH models are used in the simulation study in Drees (see [33]). The GARCH model (ix) is the GARCH model fitted to the CNX Nifty daily loss data for the duration 1st April 2009 to 31st March 2013 (sample size is 995). The data are obtained from the daily closing values CNX Nifty index during the above mentioned period. Source http://www.nseindia.com/products/content/equities/indices/historical_index_data.htm.

We also consider a small-scale experiment to compare performance of the estimators of VaR and MS under netting agreements. The term netting is used to describe the process of offsetting mutual positions or obligations between two parties (see [41]). Suppose a trader borrows money from a broker, takes a long position on a certain equity and also buys a put option (short position) of the market index future to hedge against any random fall in the stock market. The trader can adopt two strategies. In the event of any unforeseen downward movement in the market, he may cover the gains in the put option and take delivery of the stocks by paying remaining dues to the broker in cash. Otherwise the trader can exit both the long and short positions at market price, and return the dues to the broker. In this example a sudden downward market movement is the event that causes default. The first strategy is not netted, as only positions with positive gains are used to meet the default

obligation. The second strategy involves netting, where overall portfolio gain is used to meet the traders obligation to the broker. Our model (x) represents the loss in the second strategy at time t . To study the effect of netting on VaR estimation let us consider a simple portfolio made of a long position in one asset and a short position in another one with the same counter party.

Let E_{1t} , E_{2t} denote the gains in the long and short positions respectively. The vector (E_{1t}, E_{2t}) is assumed to be Gaussian. m_i and σ_i are mean and standard deviation of E_{it} , $i = 1, 2$ and ρ is the correlation coefficient. Since E_{1t} and E_{2t} are long and short position gains, we assume that ρ is negative. Let D_t be a Bernoulli random variable, independent of (E_{1t}, E_{2t}) , such that $D_t = 1$ represents a credit event that causes default at time t (and hence initiation of a netting agreement). In case of default, without any netting arrangement, the loss at time t equals $E_{1t}^+ + E_{2t}^+$. However under netting arrangement, the loss due to default at time t equals $(E_{1t} + E_{2t})^+$ (see Fermanian and Scaillet[41], page 937). Therefore under this netting arrangement, the loss at time t equals $I(D_t = 1)(E_{1t} + E_{2t})^+ - I(D_t = 0)(E_{1t} + E_{2t})$. This motivates model (x) in our simulation study

$$(x) \quad X_t = I(D_t = 1)(E_{1t} + E_{2t})^+ - I(D_t = 0)(E_{1t} + E_{2t}),$$

where $\{(E_{1t}, E_{2t})\}$ is an i.i.d. Gaussian process, with $m_1 = 10$, $m_2 = -1$, $\rho = 0.89$ and $\sigma_i = 1, i = 1, 2$. And we take $P(D_t = 1) = 0.20$, i.e. the chance of default is assumed to be twenty percent.

Let MSE1, MSE2 denote the Monte Carlo estimates of the MSE of $EVT1_p$ and $EVT2_p$. MSE3, MSE4 and MSE5 denote the same for the estimators C- T_p , H- D_p and $SV3_p$ respectively. MSE6 denotes the Monte Carlo MSE estimate of the quantile estimator Al_p using Epachnikov kernel and bandwidth b_{opt2} . MSE7 and MSE8 are the same for the estimators S- G_p and the sample quantile SQ_p .

2.3 Findings

From each of the above models (i)-(x), we have the following simulation results

1. *Comparison of MSEs.* None of the nonparametric estimators can be claimed to be uniformly the best in terms of their MSE for all the ten models. However, the EVT based estimator $EVT2_p$ by Drees [33] seems to outperform the empirical estimator for $p = 0.01$ and $30 < n \leq 500$, based on i.i.d. data (see Table 2.1). For $n \leq 500$ and $p \leq 0.01$, Swanepoel and Grann s' estimator S- G_p performs best under ARMA model (See Table 2.8).

The MSE of $EVT2_p$ is also substantially smaller than the same of $S-G_p$, for the GARCH models (vii) and (viii) for $p = 0.01$ and $30 < n \leq 500$ (see Table 2.3). For the GARCH model (ix), the MSE of $EVT2_p$ and $SV3_p$ seems to be only slightly smaller (less than a percent) than the MSE of the empirical estimator SQ_p . SQ_p outperforms the other nonparametric estimators for this model. Under i.i.d. assumption, the H-D estimator performs well for $n = 30$ and data generated by the GPD distribution. But in general, the accuracy of H-D_p does not seem to be substantially superior to the SQ_p especially in the presence of dependence.

2. For $n \leq 250$ and $k = [np] + 1$, the $EVT1_p$ is not defined for $p = 0.01$. This is due to the fact that for $n \leq 250$ and $p = 0.01$ (or less), k is small i.e. there are not enough observations to fit the GPD distribution by maximum likelihood method. So $EVT1_p$ is not recommended for small n and p close to zero. For large sample size, $EVT1_p$ provides slight improvement over the empirical quantile estimator SQ_p , for p close to zero and F equal to a heavy tailed distribution function (See Table 2.1).
3. *Comparison of bias.* For $n \geq 100$ and $p = 0.01$, the sample quantile seems to have least bias per unit standard deviation under i.i.d. and GARCH models and the bias is positive (see Tables 2.2 and 2.5). From the Tables 2.2, 2.4 and 2.5, we see that $EVT2_p$ estimator is negatively biased irrespective of the underlying model. Hence, $EVT2_p$ seems to under estimate a quantile in the right tail of F . The sample quantile SQ_p seems to exhibit least bias per unit standard deviation (see Tables 2.2, 2.4 and 2.5).
4. *Small n and p close to zero.* For $n \leq 500$ and $p \leq 0.01$, the MSE and also the bias (per unit standard deviation) of the $S-G_p$ estimator seem to be substantially smaller than the same of the empirical quantile estimator SQ_p for ARMA models (see Tables 2.8 and 2.9).

For $n \leq 500$ and $p = 0.001$ the MSE and the bias of the $S-G_p$ estimator seems to be substantially smaller than the same of the empirical quantile estimator SQ_p for the GARCH models (see Tables 2.10 and 2.11). Under similar conditions the MSE of the $EVT2_p$ estimator is also substantially smaller than that of the $(1 - p)$ th sample quantile for i.i.d. model with heavy tailed marginal distributions (see Table 2.6).

5. *Large n and p close to zero.* From the Table 2.10 we see that for $n \geq 500$ and $p = 0.001$ the MSE of the $SV3$ estimators is substantially smaller than that of the $(1 - p)$ th sample quantile SQ_p for the GARCH models (vii) and (viii). However for the GARCH (ix) model, no other estimator seem to produce substantially improved quantile estimate than the SQ_p for any value of p , especially for large n (see Tables 2.10 and 2.11).

6. *CNX Nifty data analysis.* For the GARCH model (ix), the empirical quantile seems to be the optimal estimator for $n \geq 500$ and any value of p . As mentioned above, this model is a good fit to the CNX Nifty daily loss data for the duration 1st April 2009 to 31st March 2013. In this data the number of observations exceed 500. Therefore based on the above simulation results we recommend the sample quantile SQ_p as the appropriate estimator of VaR and MS based on the CNX Nifty data (see Table 2.10).

The sample quantile based estimates of the 99 percent VaR and MS of the Nifty index during 1st April 2009 to 31st March 2013 are equal to -0.013 and -0.015 respectively. The MS value implies that during 1st April 2009 to 31st March 2013, the chance that the return of the Nity index on a trading day was less than -1.45 percent seems to be less than 0.5 percent.

7. *Effect of netting.* Under model (x), we see that the $S-G_p$ estimator outperforms the empirical estimator for all choices of n and $p = 0.01$ (see Table A.12 of Appendix A). We also observe that $S-G_p$ estimator outperforms the empirical estimator for $n \leq 500$ and $p = 0.001$. We observe that SV_3 estimator outperforms the empirical estimator for $n \geq 100$ and $p = 0.01$ (see Table 2.12). We also observe that the $C-T_p$ estimator seems to be close to the empirical estimator (see Table 2.14). From Table 2.15, we observe that SQ_p estimator seems to have least bias per unit standard deviation and its bias is positive. From Table 2.15, we see that SV_3 estimator is negatively biased irrespective of the underlying model. Hence, SV_3 seems to under estimate a quantile in the right tail of F .

2.4 Concluding remarks

Remark 3. 1. For $n \leq 500$ and $p = 0.001$, the proposed $S-G_p$ estimator performs very well under all the different time series models considered in our simulation study. We recommend this estimator for estimation of extreme quantiles in the right tail of F , especially for small n and p . Therefore this estimator appears to be useful for estimation of VaR and MS based on short term (less than one financial year) asset return data. Under the assumption that $p \rightarrow 0$ as $n \rightarrow \infty$, investigating the theoretical properties of $S-G_p$ seems to be an interesting but challenging problem.

2. The kernel based quantile estimator using the bandwidth by Alemany et al. [3] was developed under i.i.d. assumptions. But our simulations reveal that the MSE values of the Al_p and sample quantile are close for ARMA and GARCH models, especially for large n . For ARMA model, Al_p outperforms the sample quantile for small n and p . These obser-

variations naturally motivate an investigation of the properties of Al_p under such dependence assumptions.

3. VaR and MS estimation in presence of netting. In presence of netting, our new proposed quantile estimator $S-G_p$ performs very well for all choices of n and $p = 0.01$. But for $p = 0.001$ the $S-G_p$ performs very well for $n \leq 500$. Hence, we recommend this estimator for small n and p . Also for all choices of n and $p = 0.01$ this estimator is recommended.

Table 2.1: Ratios estimated at 99% for i.i.d. cases.

Dist	n	$\frac{MSE1}{MSE8}$	$\frac{MSE2}{MSE8}$	$\frac{MSE3}{MSE8}$	$\frac{MSE4}{MSE8}$	$\frac{MSE5}{MSE8}$	$\frac{MSE6}{MSE8}$	$\frac{MSE7}{MSE8}$
GPD	30	NaN	2.247	0.873	0.882	0.689	1.188	1.010
	100	NaN	0.213	0.854	1.406	0.548	1.375	1.455
	250	NaN	0.691	1.222	2.435	1.156	0.998	0.974
	500	0.886	0.814	1.269	1.251	0.939	1.004	0.994
	1000	0.916	0.872	1.040	1.008	0.887	0.994	0.972
	2500	0.979	0.961	1.154	0.992	0.939	0.991	0.976
Student's t	30	NaN	1.512	1.817	0.926	0.821	0.999	0.921
	100	NaN	0.356	1.297	1.299	0.591	1.264	1.232
	250	NaN	0.751	1.095	1.702	0.931	0.961	0.900
	500	0.922	0.817	1.003	1.112	0.876	0.982	0.911
	1000	0.934	0.878	0.985	1.007	0.897	0.970	0.896
	2500	0.964	0.950	1.116	0.981	0.930	0.966	0.921
N(0,1)	30	NaN	1.167	0.822	0.997	1.037	0.858	0.996
	100	NaN	0.863	0.997	1.042	0.861	1.024	1.531
	250	NaN	0.873	0.937	0.941	0.760	0.890	1
	500	1.170	0.936	0.976	0.915	0.831	0.942	0.886
	1000	1.096	0.983	1.022	0.910	0.864	0.943	0.947
	2500	0.995	0.952	0.981	0.916	0.893	0.938	1.072

Table 2.2: Bias/sd estimated at 99% for i.i.d. cases.

Dist	n	$\frac{Bias}{sd(EV'T1_p)}$	$\frac{Bias}{sd(EV'T2_p)}$	$\frac{Bias}{sd(C-T_p)}$	$\frac{Bias}{sd(H-D_p)}$	$\frac{Bias}{sd(SV3_p)}$	$\frac{Bias}{sd(AL_p)}$	$\frac{Bias}{sd(S-G_p)}$	$\frac{Bias}{sd(SQ_p)}$
GPD	30	NA	0.035	-0.169	-0.225	-0.395	-0.082	-0.107	-0.158
	100	NA	-0.136	0.234	0.280	0.086	0.250	0.251	0.265
	250	NA	-0.239	0.123	0.459	0.190	0.087	0.084	0.068
	500	-0.458	-0.185	0.089	0.370	0.148	0.071	0.077	0.066
	1000	-0.337	-0.167	0.078	0.236	0.072	0.021	0.028	0.014
	2500	-0.214	-0.115	0.114	0.149	0.041	0.008	0.024	0.001
Student's t	30	NA	-0.075	-0.130	-0.381	-0.579	-0.142	-0.101	-0.296
	100	NA	-0.231	0.202	0.238	-0.033	0.192	0.217	0.205
	250	NA	-0.309	0.041	0.420	0.085	0.036	0.104	0.008
	500	-0.503	-0.212	0.069	0.314	0.085	0.057	0.125	0.046
	1000	-0.358	-0.175	0.012	0.207	0.043	0.022	0.116	0.004
	2500	-0.172	-0.063	0.005	0.182	0.075	0.062	0.193	0.0469
N(0,1)	30	NA	-0.222	-0.343	-0.605	-0.855	-0.211	-0.483	-0.489
	100	NA	-0.433	0.028	0.071	-0.350	0.075	0.138	0.043
	250	NA	-0.378	-0.013	0.278	-0.113	0.032	-0.019	-0.019
	500	-0.590	-0.301	-0.001	0.185	-0.070	0.013	0.465	-0.026
	1000	-0.393	-0.211	-0.001	0.136	-0.036	0.035	0.550	-0.001
	2500	-0.249	-0.139	0.050	0.101	-0.006	0.043	0.680	0.008

Table 2.3: Ratios estimated at 99% for dependent cases.

Model	n	$\frac{MSE1}{MSE8}$	$\frac{MSE2}{MSE8}$	$\frac{MSE3}{MSE8}$	$\frac{MSE4}{MSE8}$	$\frac{MSE5}{MSE8}$	$\frac{MSE6}{MSE8}$	$\frac{MSE7}{MSE8}$
ARMA (0.95,-0.9)	30	NaN	1.0633	0.7585	1.0209	1.0959	0.8317	0.6788
	100	NaN	0.957	0.922	1.023	0.946	0.998	0.833
	250	NaN	0.950	0.905	0.922	0.821	0.910	0.769
	500	1.247	0.985	1.092	0.885	0.870	0.930	0.795
	1000	1.184	1.031	0.981	0.903	0.913	0.949	0.853
	2500	1.245	1.014	0.985	0.901	0.917	0.937	0.876
ARMA (0.95,-0.6)	30	NaN	NaN	0.923	1.029	1.090	0.878	0.7636
	100	NaN	NaN	1.028	0.990	1.108	0.973	0.883
	250	NaN	1.051	0.929	0.927	0.981	0.958	0.876
	500	1.139	1.053	1.218	0.963	0.995	0.979	0.913
	1000	1.032	1.002	1.026	0.975	0.979	0.978	0.922
	2500	1.005	0.997	1.156	0.985	0.982	0.987	0.957
ARMA (0.3,0.9)	30	NaN	0.955	0.880	1.032	1.107	0.771	0.672
	100	NaN	0.965	0.954	0.983	0.958	0.948	0.871
	250	NaN	0.929	1.131	0.958	0.843	0.891	0.819
	500	1.101	0.942	1.016	0.949	0.886	0.934	0.862
	1000	1.090	0.993	0.990	0.944	0.923	0.932	0.868
	2500	1.021	0.987	1.058	0.954	0.943	0.946	0.932
GARCH (0.9)	30	NaN	1.218	34.031	0.942	0.854	1.015	1.442
	100	NaN	0.620	15.103	1.017	0.703	0.996	1.081
	250	NaN	0.607	12.866	1.340	0.884	1	1.169
	500	0.571	0.714	14.285	1.407	1.045	1.062	1.411
	1000	0.939	0.848	15.151	1.462	1.254	1.052	2.024
	2500	0.856	0.909	15.454	1.090	1.023	1.006	3.087
GARCH (0.4,0.5)	30	NaN	1.093	51.631	0.989	0.981	1.005	1.626
	100	NaN	0.811	19.722	1.003	0.862	1.049	1.331
	250	NaN	0.751	13	1.104	0.897	1.000	1.342
	500	0.749	0.799	11.714	1.122	0.941	1.030	1.467
	1000	0.780	0.820	8.604	1.171	1.041	1.011	1.704
	2500	0.959	0.987	9.230	1.067	1.026	0.996	2.512
GARCH (0.075, 0.919)	30	NaN	0.976	2219.411	0.989	1.060	0.995	11.197
	100	NaN	0.998	1053.511	1.003	0.999	0.991	8.233
	250	NaN	0.997	644.549	1.104	0.979	0.999	6.465
	500	0.996	0.989	462.904	1.123	0.992	0.993	5.787
	1000	0.974	0.984	290.94	1.171	0.996	1.002	4.837
	2500	0.976	0.986	205.076	1.067	0.997	1.002	4.800

Table 2.4: Bias/sd estimated for ARMA models at 99%.

Model	n	$\frac{Bias}{sd(EV'T1_p)}$	$\frac{Bias}{sd(EV'T2_p)}$	$\frac{Bias}{sd(C-T_p)}$	$\frac{Bias}{sd(H-D_p)}$	$\frac{Bias}{sd(SV'3_p)}$	$\frac{Bias}{sd(AL_p)}$	$\frac{Bias}{sd(S-G_p)}$	$\frac{Bias}{sd(SQ_p)}$
ARMA(0.95,-0.9)	30	NA	-0.342	-0.451	-0.712	-0.933	-0.340	-0.061	-0.605
	100	NA	-0.526	-0.051	-0.057	-0.450	-0.057	0.150	-0.087
	250	NA	-0.446	-0.141	0.148	-0.204	-0.079	0.244	-0.128
	500	-0.702	-0.437	-0.051	0.022	-0.212	-0.133	0.255	-0.170
	1000	0.557	-0.383	-0.106	-0.040	-0.202	-0.139	0.320	-0.165
	2500	-0.400	-0.302	-0.093	-0.105	-0.202	-0.151	0.389	-0.178
ARMA(0.95,-0.6)	30	NA	NA	-0.985	-1.114	-1.200	-0.900	-0.795	-1.067
	100	NA	NA	-0.605	-0.629	-0.803	-0.592	-0.523	-0.636
	250	NA	-0.535	-0.400	-0.293	-0.447	-0.363	-0.269	-0.409
	500	-0.550	-0.434	-0.218	-0.240	-0.341	-0.284	-0.175	-0.315
	1000	-0.325	-0.252	-0.165	-0.119	-0.184	-0.144	-0.019	-0.172
	2500	-0.171	-0.127	-0.031	-0.034	-0.076	-0.046	0.106	-0.0708
ARMA(0.3,0.9)	30	NA	-0.625	-0.683	-0.888	-1.038	-0.411	-0.378	-0.808
	100	NA	-0.487	-0.094	-0.136	-0.439	-0.075	0.002	-0.141
	250	NA	-0.367	-0.058	0.138	-0.164	-0.002	0.137	-0.079
	500	-0.545	-0.306	-0.060	0.100	-0.102	-0.004	0.179	-0.068
	1000	-0.428	-0.272	-0.032	0.032	-0.108	-0.031	0.201	-0.087
	2500	-0.251	-0.161	-0.024	0.023	-0.062	0.001	0.300	-0.052

Table 2.5: Bias/sd estimated for GARCH models at 99%.

Model	n	$\frac{Bias}{sd(EV'T1_p)}$	$\frac{Bias}{sd(EV'T2_p)}$	$\frac{Bias}{sd(C-T_p)}$	$\frac{Bias}{sd(H-D_p)}$	$\frac{Bias}{sd(SV'3_p)}$	$\frac{Bias}{sd(AL_p)}$	$\frac{Bias}{sd(S-GP)}$	$\frac{Bias}{sd(SQ_p)}$
GARCH(0.9)	30	NA	-0.016	10.853	-0.306	-0.393	-0.257	0.535	0.148
	100	NA	-0.150	4.259	0.091	-0.069	0.073	0.600	-0.264
	250	NA	-0.141	5.693	0.305	0.156	0.1461	0.778	0.072
	500	-0.195	-0.040	7.541	0.309	0.203	0.150	0.960	0.146
	1000	-0.031	0.024	7.704	0.236	0.164	0.134	1.002	0.128
	2500	0.015	0.058	8.797	0.204	0.148	0.119	1.279	0.116
GARCH(0.4,0.5)	30	NA	-0.449	-0.683	-0.676	-0.799	-0.599	0.795	-0.610
	100	NA	-0.276	-0.094	-0.131	-0.324	-0.124	0.849	-0.069
	250	NA	-0.174	-0.058	0.115	-0.041	-0.022	0.921	-0.005
	500	-0.274	-0.134	-0.060	0.154	0.048	0.048	1.042	0.013
	1000	-0.128	-0.048	-0.032	0.119	0.051	0.042	1.091	0.051
	2500	-0.130	-0.080	-0.024	0.043	-0.004	-0.015	1.366	-0.024
GARCH(0.0751,0.9194)	30	NA	0.797	121.641	-0.970	-1.182	-0.960	2.149	-1.041
	100	NA	-0.675	117.609	-0.497	-0.647	-0.487	1.848	-0.486
	250	NA	-0.549	111.665	-0.444	-0.474	-0.444	1.895	-0.355
	500	-0.377	-0.450	104.537	-0.377	-0.384	-0.376	1.916	-0.320
	1000	-0.220	-0.264	85.002	-0.220	-0.221	-0.217	1.916	-0.186
	2500	-0.1264	-0.107	56.894	-0.083	-0.083	-0.082	2.141	-0.06

Table 2.6: Ratios estimated for i.i.d. cases with varying p .

Dist	n	$1 - p$	$\frac{MSE1}{MSE8}$	$\frac{MSE2}{MSE8}$	$\frac{MSE3}{MSE8}$	$\frac{MSE4}{MSE8}$	$\frac{MSE5}{MSE8}$	$\frac{MSE6}{MSE8}$	$\frac{MSE7}{MSE8}$
GPD	250	0.95	0.903	0.901	0.994	1.009	0.865	1.217	0.936
		0.97	0.834	0.825	0.995	1.031	0.878	1.064	0.933
		0.99	NaN	0.691	0.999	2.237	1.156	0.998	0.982
		0.999	NaN	5.416	0.993	0.917	0.751	1.025	0.992
	500	0.95	0.964	0.954	0.996	0.979	0.918	1.025	0.943
		0.97	0.904	0.890	1.000	1.018	0.920	0.990	0.961
		0.99	0.886	0.814	1.046	1.298	0.939	1.004	1.009
		0.999	NaN	0.402	0.999	0.733	0.471	0.999	0.991
	1000	0.95	0.957	0.952	0.997	0.971	0.930	0.965	0.954
		0.97	0.974	0.963	0.999	0.982	0.922	0.970	0.962
		0.99	0.916	0.872	1.012	1.052	0.887	0.994	0.979
		0.999	NaN	0.305	1.680	1.354	0.563	1.414	1.318
Student's t	250	0.95	0.972	0.957	0.936	0.943	0.897	0.932	0.921
		0.97	0.920	0.891	0.968	1.014	0.862	0.915	0.875
		0.99	NaN	0.751	0.990	1.761	0.931	0.961	0.866
		0.999	NaN	4.130	0.989	0.932	0.812	0.980	0.974
	500	0.95	1.027	0.995	0.978	0.955	0.930	0.931	0.953
		0.97	0.985	0.946	0.983	0.974	0.882	0.932	0.867
		0.99	0.922	0.817	1.026	1.207	0.876	0.982	0.929
		0.999	NaN	0.494	0.999	0.780	0.527	0.996	0.993
	1000	0.95	0.996	0.984	0.980	0.957	0.936	0.941	0.978
		0.97	0.992	0.964	0.984	0.954	0.918	0.940	0.910
		0.99	0.934	0.878	0.996	1.065	0.897	0.970	0.886
		0.999	NaN	0.468	1.225	1.231	0.619	1.317	1.103
N(0,1)	250	0.95	1.016	0.981	0.913	0.912	0.856	0.907	0.978
		0.97	1.024	0.960	0.913	0.861	0.837	0.903	0.884
		0.99	NaN	0.873	0.933	0.942	0.760	0.890	0.771
		0.999	NaN	1.248	0.872	1.011	1.053	0.825	0.611
	500	0.95	1.028	0.982	0.941	0.909	0.912	0.932	0.954
		0.97	1.048	0.970	0.944	0.945	0.878	0.917	0.945
		0.99	1.170	0.936	0.965	0.935	0.831	0.942	0.893
		0.999	NaN	0.829	0.976	0.901	0.849	0.960	0.753
	1000	0.95	1.013	0.990	0.933	0.933	0.937	0.944	1.059
		0.97	1.039	1.004	0.943	0.922	0.911	0.934	1.053
		0.99	1.096	0.983	0.960	0.921	0.864	0.943	0.920
		0.999	NaN	0.824	0.996	1.048	0.828	1.182	0.823

Table 2.7: Bias/sd estimated for i.i.d. cases with varying p .

Dist	n	$1-p$	$\frac{Bias}{sd(EV.T1_p)}$	$\frac{Bias}{sd(EV.T2_p)}$	$\frac{Bias}{sd(C.T_p)}$	$\frac{Bias}{sd(H.D_p)}$	$\frac{Bias}{sd(SV3_p)}$	$\frac{Bias}{sd(AL_p)}$	$\frac{Bias}{sd(S.G_p)}$	$\frac{Bias}{sd(SQ_p)}$
GPD	250	0.95	-0.164	-0.098	0.046	0.207	0.105	0.132	0.129	0.058
		0.97	-0.230	-0.129	-0.015	0.275	0.142	0.123	0.067	0.088
		0.99	NA	-0.239	0.072	0.441	0.190	0.087	0.073	0.068
		0.999	NA	0.057	-0.335	-0.354	-0.447	-0.220	-0.274	-0.246
	500	0.95	-0.183	-0.094	-0.002	0.165	0.069	0.100	0.096	0.039
		0.97	-0.233	-0.105	0.052	0.219	0.097	0.091	0.073	0.060
		0.99	-0.458	-0.185	0.026	0.370	0.148	0.071	0.090	0.066
		0.999	NA	-0.219	0.011	0.009	-0.229	0.087	0.109	0.087
	1000	0.95	-0.135	-0.073	0.001	0.108	0.037	0.074	0.199	0.023
		0.97	-0.181	-0.099	0.006	0.156	0.036	0.036	0.123	0.007
		0.99	-0.337	-0.167	0.038	0.258	0.072	0.021	0.061	-0.115
		0.999	NA	-0.172	0.206	0.237	0.042	0.273	0.249	0.266
Student's t	250	0.95	-0.234	-0.163	0.028	0.136	0.020	0.094	0.362	0.013
		0.97	-0.311	-0.200	0.044	0.219	0.034	0.057	0.231	-0.000
		0.99	NA	-0.309	0.012	0.427	0.085	0.036	0.097	0.008
		0.999	NA	-0.005	-0.308	-0.434	-0.622	-0.325	-0.166	-0.376
	500	0.95	-0.275	-0.179	0.022	0.129	-0.032	0.031	0.368	-0.035
		0.97	-0.297	-0.168	0.039	0.176	0.005	0.035	0.237	-0.013
		0.99	-0.503	-0.212	0.059	0.333	0.085	0.057	0.173	0.046
		0.999	NA	-0.288	0.041	-0.117	-0.335	0.043	0.092	0.040
	1000	0.95	-0.173	-0.108	0.015	0.124	-0.009	0.052	0.400	-0.008
		0.97	-0.225	-0.136	-0.032	0.155	-0.006	0.025	0.295	-0.018
		0.99	-0.358	-0.175	0.019	0.202	0.043	0.022	0.094	0.004
		0.999	NA	-0.181	0.266	0.257	-0.045	0.194	0.187	0.209
N(0,1)	250	0.95	-0.272	-0.198	0.017	0.095	-0.045	0.053	0.491	-0.024
		0.97	-0.335	-0.225	0.023	0.123	-0.044	0.055	0.420	0.001
		0.99	NA	-0.378	-0.056	0.240	-0.113	0.032	0.309	-0.019
		0.999	NA	-0.223	-0.548	-0.794	-0.987	-0.460	-0.014	-0.669
	500	0.95	-0.283	-0.185	0.049	0.070	-0.048	0.031	0.500	-0.028
		0.97	-0.348	-0.212	0.037	0.089	-0.052	0.030	0.561	-0.018
		0.99	-0.590	-0.301	-0.010	0.183	-0.070	0.013	0.461	-0.026
		0.999	NA	-0.460	-0.144	-0.355	-0.616	-0.074	0.133	-0.104
	1000	0.95	-0.200	-0.134	-0.026	0.068	-0.031	0.037	0.601	-0.013
		0.97	-0.228	-0.138	-0.014	0.060	-0.025	0.040	0.642	-0.007
		0.99	-0.393	-0.211	0.024	0.119	-0.036	0.035	0.511	-0.001
		0.999	NA	-0.362	0.002	0.038	-0.310	0.082	0.225	0.066

Table 2.8: Ratios estimated for ARMA model with varying p .

Model Coeff.	n	$1 - p$	$\frac{MSE1}{MSE8}$	$\frac{MSE2}{MSE8}$	$\frac{MSE3}{MSE8}$	$\frac{MSE4}{MSE8}$	$\frac{MSE5}{MSE8}$	$\frac{MSE6}{MSE8}$	$\frac{MSE7}{MSE8}$
(0.95, -0.9)	250	0.95	1.0572	1.0207	0.9516	0.9157	0.9245	0.9267	0.9118
		0.97	1.068	1.004	0.947	0.943	0.888	0.910	0.851
		0.99	NaN	0.950	0.938	0.939	0.825	0.910	0.799
		0.999	NaN	1.185	0.868	1.017	1.062	0.835	0.623
	500	0.95	1.111	1.054	0.974	0.948	0.954	0.938	1.080
		0.97	1.168	1.070	0.967	0.941	0.936	0.932	0.828
		0.99	1.247	0.985	0.960	0.957	0.870	0.930	0.840
		0.999	NaN	0.886	0.967	0.908	0.924	0.952	0.736
	1000	0.95	1.063	1.028	0.977	0.978	0.971	0.949	1.180
		0.97	1.107	1.050	0.955	0.955	0.968	0.955	1.011
		0.99	1.184	1.031	0.966	0.944	0.913	0.949	0.862
		0.999	NaN	0.891	1.020	1.039	0.889	1.102	0.944
(0.95, -0.6)	250	0.95	1.027	1.018	0.991	0.915	0.989	0.987	0.951
		0.97	1.051	1.031	0.990	0.944	0.985	0.973	0.914
		0.99	NaN	1.051	0.983	0.939	0.981	0.958	0.879
		0.999	NaN	0.921	0.921	1.017	1.077	0.845	0.672
	500	0.95	1.019	1.009	0.994	0.991	0.986	0.9804	0.9442
		0.97	1.038	1.017	0.992	0.976	0.988	0.982	0.919
		0.99	1.139	1.053	0.989	0.960	0.995	0.979	0.896
		0.999	NaN	1.099	0.986	1.056	1.156	0.941	0.787
	1000	0.95	1.012	1.009	0.997	0.989	0.995	0.991	0.974
		0.97	1.019	1.011	0.997	0.988	0.993	0.989	0.991
		0.99	1.032	1.002	0.990	0.976	0.974	0.978	0.915
		0.999	NaN	1.119	1.004	0.996	1.102	1.003	0.855
(0.3, 0.9)	250	0.95	1.015	0.991	0.961	0.946	0.915	0.918	0.882
		0.97	1.039	0.997	0.957	0.961	0.900	0.919	0.890
		0.99	NaN	0.929	0.972	0.993	0.843	0.891	0.869
		0.999	NaN	1.051	0.913	1.015	1.072	0.773	0.601
	500	0.95	1.045	1.014	0.976	0.966	0.948	0.938	0.971
		0.97	1.040	0.991	0.983	0.982	0.926	0.934	0.845
		0.99	1.101	0.942	0.985	1.010	0.886	0.929	0.871
		0.999	NaN	0.914	0.986	0.927	0.936	0.925	0.799
	1000	0.95	1.032	1.014	0.975	0.981	0.956	0.933	0.881
		0.97	1.055	0.991	0.982	0.985	0.947	0.930	1.002
		0.99	1.090	0.942	0.985	0.978	0.923	0.932	0.788
		0.999	NaN	0.914	0.998	1.018	0.886	0.993	0.875

Table 2.9: Bias/sd estimated for ARMA model with varying p .

Model Coeff.	n	$1 - p$	$\frac{Bias}{sd(EV T_1 p)}$	$\frac{Bias}{sd(C-T_1 p)}$	$\frac{Bias}{sd(H-D_1 p)}$	$\frac{Bias}{sd(SV_2 p)}$	$\frac{Bias}{sd(AL_1 p)}$	$\frac{Bias}{sd(SG_1 p)}$	$\frac{Bias}{sd(SQ_1 p)}$
(0.95, -0.9)	250	0.95	-0.361	-0.040	0.092	-0.189	-0.108	0.233	-0.152
		0.97	-0.422	-0.002	0.117	-0.187	-0.095	0.280	-0.138
		0.99	NA	-0.037	0.230	-0.204	-0.079	0.279	-0.128
		0.999	NA	-0.465	-0.811	-0.997	-0.500	-0.015	-0.698
	500	0.95	-0.438	-0.015	0.098	-0.254	-0.189	0.545	-0.225
		0.97	-0.501	0.032	0.091	-0.251	-0.180	0.181	-0.222
		0.99	-0.702	-0.437	0.173	-0.212	-0.133	0.281	-0.170
		0.999	NA	-0.049	-0.374	-0.714	-0.162	0.107	-0.189
	1000	0.95	-0.387	0.022	0.181	-0.268	-0.213	0.644	-0.246
		0.97	-0.434	-0.039	0.160	-0.267	-0.209	0.422	-0.242
		0.99	-0.557	-0.052	0.189	-0.202	-0.139	0.310	-0.165
		0.999	NA	0.238	0.046	-0.416	0.003	0.476	-0.019
(0.95, -0.6)	250	0.95	-0.315	-0.204	-0.204	-0.251	-0.205	-0.084	-0.233
		0.97	-0.411	-0.273	-0.260	-0.312	-0.261	-0.242	-0.297
		0.99	NA	-0.476	-0.341	-0.447	-0.363	-0.363	-0.409
		0.999	NA	-1.014	-1.216	-1.184	-0.824	-0.725	-1.026
	500	0.95	-0.253	-0.153	-0.113	-0.189	-0.153	-0.102	-0.179
		0.97	-0.336	-0.234	-0.147	-0.238	-0.196	-0.208	-0.226
		0.99	-0.550	-0.378	-0.229	-0.341	-0.284	-0.213	-0.315
		0.999	NA	-0.669	-0.815	-0.957	-0.600	-0.408	-0.642
	1000	0.95	-0.165	-0.173	-0.048	-0.116	-0.086	0.032	-0.109
		0.97	-0.209	-0.033	-0.056	-0.139	-0.105	0.135	-0.132
		0.99	-0.325	-0.404	-0.104	-0.184	-0.144	-0.065	-0.172
		0.999	NA	-0.438	-0.421	-0.644	-0.378	-0.321	-0.403
(0.3, 0.9)	250	0.95	-0.255	-0.005	0.103	-0.083	0.0332	0.1560	-0.0472
		0.97	-0.317	0.075	0.173	-0.111	0.015	0.201	-0.073
		0.99	NA	-0.083	0.187	-0.164	-0.002	0.191	-0.079
		0.999	NA	-0.684	-0.720	-1.001	-0.483	-0.359	-0.758
	500	0.95	-0.276	-0.099	0.167	-0.098	-0.001	0.271	-0.085
		0.97	-0.320	-0.042	0.230	-0.102	-0.004	0.106	-0.064
		0.99	-0.545	-0.253	0.208	-0.102	-0.004	0.161	-0.0682
		0.999	NA	-0.177	-0.300	-0.611	-0.147	-0.018	-0.201
	1000	0.95	-0.225	-0.100	0.238	-0.108	-0.023	0.169	-0.092
		0.97	-0.301	-0.058	0.321	-0.143	-0.057	0.379	-0.114
		0.99	-0.428	-0.079	0.209	-0.108	-0.031	0.002	-0.087
		0.999	NA	-0.118	0.110	-0.346	-0.008	0.083	-0.030

Table 2.10: Ratios estimated for GARCH model with varying p .

Model Coeff.	n	$1 - p$	$\frac{MSE1}{MSE8}$	$\frac{MSE2}{MSE8}$	$\frac{MSE3}{MSE8}$	$\frac{MSE4}{MSE8}$	$\frac{MSE5}{MSE8}$	$\frac{MSE6}{MSE8}$	$\frac{MSE7}{MSE8}$
$\alpha = 0.9$	250	0.95	0.862	0.874	502.127	1.1790	0.997	0.997	21.081
		0.97	0.797	0.831	127.887	1.498	1.026	0.999	6.554
		0.99	NaN	0.607	22.454	1.324	0.884	1.0000	1.169
		0.999	NaN	1.358	2.909	0.978	0.956	0.999	0.763
	500	0.95	0.936	0.938	619.296	1.059	0.991	0.998	36.187
		0.97	0.911	0.931	269.253	1.137	1.003	1.002	12.610
		0.99	0.550	0.624	19.244	1.511	1.045	1.062	1.411
		0.999	NaN	0.762	2.307	0.879	0.734	1	0.912
	1000	0.95	0.982	0.982	718.174	1.038	0.980	1.058	59.782
		0.97	0.956	0.961	302.152	1.067	1.002	1.026	21.052
		0.99	0.910	0.854	28.349	1.182	1.254	1.052	2.024
		0.999	NaN	0.752	1.907	1.072	0.781	0.912	0.980
$\alpha = 0.4,$ $\beta = 0.5$	250	0.95	0.929	0.951	161.733	1.085	1.026	0.999	9.344
		0.97	0.815	0.900	70.594	1.151	0.957	0.999	3.371
		0.99	NaN	0.751	16.545	1.072	0.897	1.000	1.208
		0.999	NaN	1.090	5.494	1.004	1.017	0.997	0.553
	500	0.95	0.968	0.966	184.817	1.076	0.975	0.998	14.098
		0.97	0.872	0.931	101.86	1.048	0.979	1.013	5.126
		0.99	0.710	0.799	17.538	1.057	0.941	1.030	1.186
		0.999	NaN	0.914	2.251	0.977	0.898	0.999	0.879
	1000	0.95	0.990	0.984	252.187	1.024	0.989	1.002	21.016
		0.97	0.982	0.962	116.344	1.031	1.002	1.003	10.479
		0.99	0.816	0.820	26.623	1.171	1.041	1.011	2.118
		0.999	NaN	0.784	1.211	1.024	0.824	1.006	0.894
$\alpha = 0.075,$ $\beta = 0.919$	250	0.95	0.956	0.965	1905.106	1.005	0.981	1.000	18.392
		0.97	0.979	0.989	1152.501	0.986	1.008	1.000	13.384
		0.99	NaN	0.997	1204.665	1.014	0.979	0.999	8.206
		0.995	NaN	1	413.582	0.936	0.965	0.995	3.671
	500	0.95	0.969	0.980	1016.256	1.027	0.986	1.001	13.839
		0.97	0.969	0.979	931.328	1.021	0.998	0.996	9.864
		0.99	0.9966	0.989	508.175	0.995	0.992	0.993	5.912
		0.995	NaN	0.999	334.377	0.984	0.989	1.000	3.372
	1000	0.95	0.969	0.977	602.071	1.005	0.997	1.001	19.276
		0.97	0.962	0.972	609.922	1.010	0.997	0.999	7.365
		0.99	0.974	0.984	382.135	1.017	0.996	1.002	3.478
		0.995	0.976	1	223.585	1.009	0.990	1.007	2.900
		0.999	NaN	1.050	187.984	0.994	1.041	0.998	2.181

Table 2.11: Bias/sd estimated for GARCH model with varying p .

Model Coeff.	n	$1-p$	$\frac{Bias}{sd(EVT_p)}$	$\frac{Bias}{sd(EVT_p)}$	$\frac{Bias}{sd(C-T_p)}$	$\frac{Bias}{sd(H-D_p)}$	$\frac{Bias}{sd(SV_p)}$	$\frac{Bias}{sd(AL_p)}$	$\frac{Bias}{sd(SG_p)}$
$\alpha = 0.9$	250	0.95	-0.099	-0.056	41.199	0.264	0.089	0.060	2.599
		0.97	-0.087	-0.027	27.685	0.288	0.153	0.094	2.058
		0.99	NA	-0.040	18.529	0.250	0.156	0.146	0.778
		0.999	NA	-0.342	2.398	-0.838	-0.858	-0.676	-0.679
	500	0.95	-0.146	-0.088	43.622	0.191	0.028	-0.002	2.960
		0.97	-0.095	-0.024	36.844	0.231	0.100	0.068	2.471
		0.99	-0.175	-0.016	9.710	0.237	0.203	0.150	0.960
		0.999	NA	-0.283	1.048	-0.351	-0.329	-0.097	-0.097
	1000	0.95	-0.139	-0.099	44.741	0.211	-0.027	-0.037	2.333
		0.97	-0.093	-0.044	34.758	0.248	0.051	0.028	2.006
		0.99	-0.050	0.024	13.678	0.206	0.164	0.134	1.002
		0.999	NA	-0.019	0.604	-0.130	0.005	0.076	0.100
$\alpha = 0.4,$ $\beta = 0.5$	250	0.95	-0.099	-0.049	31.926	0.178	0.042	0.026	2.366
		0.97	-0.116	-0.043	24.513	0.158	0.068	0.054	1.719
		0.99	NA	-0.174	7.716	0.136	-0.041	-0.022	0.590
		0.999	NA	-0.550	4.925	-1.082	-1.109	-0.917	-0.301
	500	0.95	-0.111	-0.105	28.544	0.173	0.016	0.018	2.713
		0.97	-0.132	-0.081	23.607	0.124	0.048	0.037	1.643
		0.99	-0.252	-0.133	9.449	0.125	0.048	0.048	0.516
		0.999	NA	-0.547	2.067	-0.662	-0.583	-0.315	-0.315
	1000	0.95	-0.098	-0.082	28.699	0.193	-0.006	-0.010	2.926
		0.97	-0.082	-0.040	23.650	0.123	0.027	0.021	2.832
		0.99	-0.140	-0.048	14.326	0.122	0.051	0.042	1.420
		0.999	NA	-0.193	1.123	-0.311	-0.154	-0.014	0.001
$\alpha = 0.075,$ $\beta = 0.919$	250	0.95	-0.105	-0.088	128.866	0.065	-0.050	-0.044	2.756
		0.97	-0.244	-0.215	128.488	-0.032	-0.164	-0.163	2.440
		0.99	NA	-0.549	128.712	-0.303	-0.475	-0.444	1.919
		0.995	NA	-38.307	111.824	-0.591	-0.748	-0.723	1.568
	500	0.999	NA	-1.502	122.058	-1.652	-1.911	-1.763	0.511
		0.95	-0.072	-0.053	118.436	0.107	-0.028	-0.022	2.554
		0.97	-0.179	-0.151	122.764	0.024	-0.115	-0.111	2.268
		0.99	-0.528	-0.450	117.799	-0.193	-0.385	-0.376	1.928
	1000	0.995	NA	-40.208	103.091	-0.464	-0.587	-0.562	1.540
		0.999	NA	-1.426	122.99	-1.096	-1.492	-1.252	-0.738
		0.95	-0.029	-0.017	91.002	0.126	0.002	0.001	3.082
		0.97	-0.085	-0.069	110.349	0.056	-0.045	-0.042	2.211
		0.99	-0.303	-0.264	108.969	-0.119	-0.222	-0.217	1.592
		0.995	-0.491	-42.256	82.661	-0.290	-0.352	-0.338	1.480
		0.999	NA	-0.984	118.185	-0.712	-0.962	-0.801	1.177
									-0.818

Table 2.12: Ratios estimated under netting condition at 99%.

Cond.	n	$\frac{MSE1}{MSE8}$	$\frac{MSE2}{MSE8}$	$\frac{MSE3}{MSE8}$	$\frac{MSE4}{MSE8}$	$\frac{MSE5}{MSE8}$	$\frac{MSE6}{MSE8}$	$\frac{MSE7}{MSE8}$
Netted	30	NaN	NaN	0.941	1.041	1.200	1.836	0.771
	100	NaN	NaN	1.006	1.013	0.940	1.437	0.896
	250	NaN	NaN	0.982	0.856	0.778	1.232	0.799
	500	1.436	NaN	0.996	0.880	0.834	1.177	0.778
	1000	1.043	1.085	0.991	0.944	0.892	1.070	0.792
	2500	1.057	1.017	0.990	0.932	0.920	1.274	0.878

Table 2.13: Bias/sd estimated under netting condition at 99%.

Cond.	n	$\frac{Bias}{sd(EVT1_p)}$	$\frac{Bias}{sd(EVT2_p)}$	$\frac{Bias}{sd(C-T_p)}$	$\frac{Bias}{sd(H-D_p)}$	$\frac{Bias}{sd(SV3_p)}$	$\frac{Bias}{sd(AL_p)}$	$\frac{Bias}{sd(S-G_p)}$	$\frac{Bias}{sd(SQ_p)}$
Netted	30	NA	NA	-0.504	-0.635	-0.953	0.734	-0.157	-0.515
	100	NA	NA	-0.017	-0.015	-0.403	0.748	0.125	-0.035
	250	NA	NA	-0.149	0.089	-0.169	0.784	0.250	-0.200
	500	-0.817	NA	-0.006	-0.026	-0.019	0.749	0.117	-0.208
	1000	-0.335	-0.342	0.013	0.196	-0.096	0.598	0.173	0.056
	2500	-0.312	-0.214	0.015	0.036	0.232	0.243	0.270	-0.059

Table 2.14: Ratios estimated under netting condition with varying p .

Cond.	n	$1 - p$	$\frac{MSE1}{MSE8}$	$\frac{MSE2}{MSE8}$	$\frac{MSE3}{MSE8}$	$\frac{MSE4}{MSE8}$	$\frac{MSE5}{MSE8}$	$\frac{MSE6}{MSE8}$	$\frac{MSE7}{MSE8}$
Netted	250	0.95	1.130	NaN	0.981	0.897	0.912	1.068	0.875
		0.97	1.085	NaN	0.978	0.861	0.881	1.061	0.865
		0.99	NaN	NaN	0.982	0.856	0.778	1.232	0.799
		0.999	NaN	NaN	0.939	1.019	1.088	0.596	0.618
	500	0.95	1.062	NaN	0.986	0.933	0.910	1.044	0.887
		0.97	1.156	NaN	0.986	0.900	0.889	1.222	0.883
		0.99	1.436	NaN	0.996	0.880	0.834	1.177	0.778
		0.999	NaN	0.937	0.993	0.919	1.072	0.624	0.769
	1000	0.95	1.013	NaN	0.992	0.932	0.926	1.001	0.882
		0.97	1.105	NaN	0.992	0.925	0.897	1.087	0.914
		0.99	1.043	1.085	0.991	0.944	0.892	1.070	0.792
		0.999	NaN	0.857	0.990	1.020	0.805	0.946	0.906

Table 2.15: Bias/sd estimated under netting condition with varying p .

Cond.	n	$1 - p$	$\frac{Bias}{sd(EVT1_p)}$	$\frac{Bias}{sd(EVT2_p)}$	$\frac{Bias}{sd(C-T_p)}$	$\frac{Bias}{sd(H-D_p)}$	$\frac{Bias}{sd(SV3_p)}$	$\frac{Bias}{sd(AL_p)}$	$\frac{Bias}{sd(S-G_p)}$	$\frac{Bias}{sd(SQ_p)}$
Netted	250	0.95	-0.358	NA	0.020	-0.047	-0.182	0.363	0.101	-0.034
		0.97	-0.368	NA	0.063	0.031	-0.231	0.557	0.261	0.043
		0.99	NA	NA	-0.149	0.089	-0.169	0.784	0.250	-0.200
		0.999	NA	NA	-0.631	-0.862	-1.066	0.432	-0.089	-0.686
	500	0.95	-0.217	NA	0.067	0.096	-0.051	0.387	0.226	0.055
		0.97	-0.385	NA	0.135	0.050	-0.140	0.618	0.233	-0.013
		0.99	-0.817	NA	-0.006	-0.026	-0.019	0.749	0.117	-0.208
		0.999	NA	-0.714	-0.149	-0.377	-0.948	0.255	0.119	-0.167
	1000	0.95	-0.141	NA	-0.194	0.083	0.198	0.371	0.146	-0.069
		0.97	-0.329	NA	-0.012	-0.040	-0.004	0.500	0.322	0.034
		0.99	-0.335	-0.342	0.013	0.196	-0.096	0.598	0.173	0.056
		0.999	NA	-0.445	0.068	-0.003	-0.238	0.555	0.348	0.121