

A comparison of mean-variance efficiency tests*

Dante Amengual

*Department of Economics, Princeton University,
Fisher Hall, Princeton, NJ 08544-1021, USA*
<amengual@princeton.edu>

Enrique Sentana

CEMFI, Casado del Alisal 5, E-28014 Madrid, Spain
<sentana@cemfi.es>

Revised: May 2009

Abstract

We analyse the asymptotic properties of mean-variance efficiency tests based on generalised methods of moments, and parametric and semiparametric likelihood procedures that assume elliptical innovations. We study the trade-off between efficiency and robustness, and prove that the parametric estimators provide asymptotically valid inferences when the conditional distribution of the innovations is elliptical but possibly misspecified and heteroskedastic. We compare the small sample performance of the alternative tests in a Monte Carlo study, and find some discrepancies with their asymptotic properties. Finally, we present an empirical application to US stock returns, which rejects the mean-variance efficiency of the market portfolio.

Keywords: Adaptivity, Elliptical Distributions, Financial Returns, Portfolio choice, Semiparametric Estimators.

JEL: C12, C13, C14, C16, G11, G12

*We would like to thank Manuel Arellano, Christian Bontemps, Marcelo Fernandes, Gabriele Fiorentini, Javier Mencía, Nour Meddahi, Francisco Peñaranda and Kevin Sheppard, as well as audiences at Queen Mary, the Imperial College Financial Econometrics Conference (London, 2007), the XV Finance Forum (Majorca), and the XXXII Symposium on Economic Analysis (Granada) for useful comments and discussions. The suggestions of an associate editor and two anonymous referees have also greatly improved the exposition. Of course, the usual caveat applies. Financial support from the Spanish Ministry of Science and Innovation through grant ECO 2008-00280 (Sentana) is gratefully acknowledged.

1 Introduction

Mean-variance analysis is widely regarded as the cornerstone of modern investment theory. Despite its simplicity, and the fact that more than five and a half decades have elapsed since Markowitz published his seminal work on the theory of portfolio allocation under uncertainty (Markowitz (1952)), it remains the most widely used asset allocation method. A portfolio with excess returns r_{Mt} is mean-variance efficient with respect to a given set of N assets with excess returns $\mathbf{r}_t$ if it is not possible to form another portfolio of those assets and r_{Mt} with the same expected return as r_{Mt} but a lower variance, or more appropriately, with the same variance but a higher expected return. Despite the simplicity of this definition, testing for mean-variance efficiency is of paramount importance in many practical situations, such as mutual fund performance evaluation (see De Roon and Nijman (2001) for a recent survey), gains from portfolio diversification (Errunza, Hogan and Hung (1999)), or tests of linear factor asset pricing models, including the capital asset pricing model and arbitrage pricing theory, as well as other empirically oriented asset pricing models (see e.g. Campbell, Lo and MacKinlay (1997) or Cochrane (2001) for textbook treatments).

As is well known, r_{Mt} will be mean-variance efficient with respect to $\mathbf{r}_t$ in the presence of a riskless asset if and only if the intercepts in the theoretical least squares projection of $\mathbf{r}_t$ on a constant and r_{Mt} are all 0 (see Jobson and Korkie (1982), Gibbons, Ross and Shanken (1989) and Huberman and Kandel (1987)). Therefore, it is not surprising that this early literature resorted to ordinary least squares (OLS) to test those theoretical restrictions empirically. If the distribution of $\mathbf{r}_t$ conditional on r_{Mt} (and their past) were multivariate normal, with a linear mean $\mathbf{a} + \mathbf{b}r_{Mt}$ and a constant covariance matrix $\mathbf{\Omega}$, then OLS would produce efficient estimators of the regression intercepts $\mathbf{a}$, and consequently, optimal tests of the mean-variance efficiency restrictions $H_0 : \mathbf{a} = \mathbf{0}$. In addition, it is possible to derive an F version of the test statistic whose sampling distribution in finite samples is known under exactly the same restrictive distributional assumptions (see Gibbons, Ross and Shanken (1989)). In this sense, this F -test generalises the t -test proposed by Black, Jensen and Scholes (1972) in univariate contexts.

However, many empirical studies with financial time series data indicate that the distribution of asset returns is usually rather leptokurtic. For that reason, MacKinlay and Richardson (1991) proposed alternative tests based on the generalised method of moments (GMM) that are robust to non-normality, unlike traditional OLS test statistics.

More recently, Hodgson, Linton, and Vorkink (2002; hereinafter HLV) developed a semipara-

metric estimation and testing methodology that enabled them to obtain optimal mean-variance efficiency tests under the assumption that the distribution of $\mathbf{r}_t$ conditional on r_{Mt} (and their past) is elliptically symmetric. Specifically, HLTV showed that their proposed estimators of $\mathbf{a}$ and $\mathbf{b}$ are adaptive under the aforementioned assumptions of linear conditional mean and constant conditional variance, which means that they are as efficient as *infeasible* maximum likelihood (ML) estimators that use the correct parametric elliptical density with full knowledge of its shape parameters. Elliptical distributions are attractive in this context because they relate mean-variance analysis with expected utility maximisation (see e.g. Chamberlain (1983), Owen and Rabinovitch (1983) and Berk (1997)). Moreover, they generalise the multivariate normal distribution, but at the same time they retain its analytical tractability irrespective of the number of assets.

Nevertheless, the finite sample performance of such semiparametric inference procedures may not be well approximated by the first-order asymptotic theory that justifies them. For that reason, an alternative approach worth considering is an unrestricted ML estimator based on the correct elliptical distribution, but which includes the unknown shape parameters as additional arguments in the maximisation algorithm (see e.g. Kan and Zhou (2006)). However, unless we are careful, this last approach may provide misleading inferences if the relevant conditional distribution does not coincide with the assumed one, even if both are elliptical. The same applies to elliptically-based restricted maximum likelihood estimators that keep the shape parameters fixed to some a priori values, even if the assumed conditional distribution is correct, unless the chosen values either imply multivariate normality, in which case such restricted estimators will reduce to the OLS-GMM ones, or they happened to coincide with the true values, in which case those restricted estimators would be identical to the infeasible ML estimators. Similarly, the HLTV approach may also lead to erroneous inferences if the true conditional distribution is either heteroskedastic or asymmetric.

Although at first sight these considerations may only seem interesting for theoretically inclined econometricians, they are also relevant for applied researchers because in practice the substantive conclusions about the mean-variance efficiency of a candidate portfolio can be rather sensitive to the distributional assumptions made, as our empirical results confirm.

In this context, the purpose of our paper is to shed some light on such efficiency-consistency trade-offs in the context of mean-variance efficiency tests. To do so, we will first exploit the results in Fiorentini and Sentana (2007) to derive the asymptotic properties of the estimators of the regression intercepts, $\mathbf{a}$, and slopes, $\mathbf{b}$, based on GMM, HLTV and elliptically-based parametric ML procedures under correct specification. Then, we will extend our results to characterise

how those asymptotic properties change under some specific forms of misspecification that are potentially relevant in practice in view of some observed characteristics of asset returns, which we will take into consideration in our empirical application. In particular, we study those situations in which the distribution of the innovations is:

- (i) *i.i.d.* elliptical but different from the parametric one assumed for estimation purposes, which will often be chosen for convenience or familiarity,
- (ii) elliptical but conditionally heteroskedastic, which arises when the *joint* distribution of excess returns for the N assets $\mathbf{r}_t$ and the reference portfolio, r_{Mt} , is elliptical, and
- (iii) not elliptically symmetric.

In addition, given that it is far from trivial to obtain exact finite sample distributions once we abandon the Gaussianity assumption, we also analyse the reliability of the usual asymptotic approximations by Monte Carlo methods.¹

Our main asymptotic results are:

1. Under correct specification, not only the HLTV procedure but also the unrestricted parametric estimators are adaptive, in the sense that they are as efficient as if one had full knowledge of the true conditional distribution, including its shape parameters.
2. Pseudo-ML (PML) estimators of the regression intercepts and slopes based on the Student t remain consistent when the conditional distribution is *i.i.d.* elliptical but not t irrespective of whether the degrees of freedom are estimated or fixed a priori. In addition, the restricted estimator will still be consistent when the true conditional distribution is t but with a number of degrees of freedom different from the one assumed a priori. Both these PML estimators are also consistent when the conditional distribution is elliptical but conditionally heteroskedastic. In all these cases, we provide correct expressions for the asymptotic covariance matrices of the regression coefficients, and explain how applied researchers can robustify their inferences in practice. The HLTV procedure also seems to yield consistent estimators in a conditionally heteroskedastic elliptical context, which confirms related results by Hodgson (2000) in a univariate framework.
3. The t -based PML estimators seem to be systematically more efficient than the GMM estimators when the conditional distribution of $\mathbf{r}_t$ given r_{Mt} is elliptical, irrespective of whether or not it is t or conditionally homoskedastic. In addition, estimating the degrees of freedom parameter instead of fixing its value a priori typically leads to efficiency gains.
4. Only the GMM estimator of the regression intercepts provides reliable inferences in the

¹See Beaulieu, Dufour and Khalaf (2007b) for a method to obtain the exact distribution of the Gibbons, Ross and Shanken (1989) F -statistic conditional on the full sample path of r_{Mt} when the innovations are *i.i.d.*

presence of asymmetries.

Although our Monte Carlo results are broadly in line with these theoretical conclusions, they also point out two interesting facts. First, we find that the HLV tests typically have much larger size distortions in finite samples than the other tests. Secondly, they have smaller size-adjusted power than the t -based PML tests, although the differences are very small when the latter are asymptotically suboptimal.

Finally, we apply those different procedures to test the mean-variance efficiency of the US aggregate stock market portfolio with respect to industry portfolios, and the book-to-market sorted portfolios popularised by Fama and French (1993). We do so using monthly data over the period July 1962 to June 2007. The results that we obtain for industry portfolios indicate that the Student t -based test clearly rejects the efficiency of the market portfolio, while the GMM test is borderline, and the HLV based test fails to reject. Given our Monte Carlo results, this contradicting behaviour is partly due to the lack of reliability of the nonparametric estimates of the asymptotic covariance matrix implicit in the HLV procedure, even though we use the improved procedure recommended by Fiorentini and Sentana (2007). In contrast, all three tests reject the mean-variance efficiency of the market portfolio relative to the book-to-market sorted portfolios of Fama and French (1993).

Importantly, we also assess the adequacy of our parametric assumptions by computing specification tests against heteroskedasticity, asymmetries, and departures from the t distribution in higher order moments. We find that while the assumption of Gaussianity is overwhelmingly rejected in both data sets, the evidence against a multivariate t distribution is weak. Nevertheless, we find quite strong evidence against conditional homoskedasticity, which confirms the usefulness of our robust asymptotic covariance expressions. In view of the trade-offs between efficiency and consistency that we characterise in our theoretical analysis, these empirical results suggest that it is probably worth using the multivariate t distribution for the purposes of testing mean-variance efficiency, as long as empirical researchers bear in mind that such a distributional assumption may be wrong, and robustify their inferences accordingly.

The rest of the paper is organised as follows. In section 2, we introduce the model and the three aforementioned estimation procedures, obtain their asymptotic distributions under the assumption that the innovations are *i.i.d.* elliptical, and discuss the testing implications of those results. Then in section 3 we derive the asymptotic properties of those estimators in alternative misspecified contexts. An extensive Monte Carlo evaluation of the different parameter estimators and testing procedures can be found in section 4, while section 5 reports our empirical results.

Finally, we present our conclusions and suggestions for future work in section 6. Proofs and auxiliary results are gathered in the appendices.

2 Econometric methods

2.1 Model description

Consider the following multivariate, conditionally homoskedastic, linear regression model

$$\mathbf{r}_t = \mathbf{a} + \mathbf{b}r_{Mt} + \mathbf{u}_t = \mathbf{a} + \mathbf{b}r_{Mt} + \boldsymbol{\Omega}^{1/2}\boldsymbol{\varepsilon}_t^*, \quad (1)$$

where $\boldsymbol{\Omega}^{1/2}$ is an $N \times N$ “square root” matrix such that $\boldsymbol{\Omega}^{1/2}\boldsymbol{\Omega}^{1/2} = \boldsymbol{\Omega}$, $\boldsymbol{\varepsilon}_t^*$ is a standardised vector martingale difference sequence satisfying $E(\boldsymbol{\varepsilon}_t^* | r_{Mt}, I_{t-1}; \boldsymbol{\gamma}_0, \boldsymbol{\omega}_0) = \mathbf{0}$ and $V(\boldsymbol{\varepsilon}_t^* | r_{Mt}, I_{t-1}; \boldsymbol{\gamma}_0, \boldsymbol{\omega}_0) = \mathbf{I}_N$, $\boldsymbol{\gamma}' = (\mathbf{a}', \mathbf{b}')$, $\boldsymbol{\omega} = \text{vech}(\boldsymbol{\Omega})$, the subscript 0 refers to the true values of the parameters, and I_{t-1} denotes the information set available at $t - 1$, which contains at least past values of r_{Mt} and $\mathbf{r}_t$. To complete the conditional model, we need to specify the distribution of $\boldsymbol{\varepsilon}_t^*$. We shall initially assume that conditional on r_{Mt} and I_{t-1} , $\boldsymbol{\varepsilon}_t^*$ is independent and identically distributed as some particular member of the elliptical family with a well defined density, or $\boldsymbol{\varepsilon}_t^* | r_{Mt}, I_{t-1}; \boldsymbol{\gamma}_0, \boldsymbol{\omega}_0, \boldsymbol{\eta}_0 \sim i.i.d. s(\mathbf{0}, \mathbf{I}_N, \boldsymbol{\eta}_0)$ for short, where $\boldsymbol{\eta}$ are some q additional parameters that determine the shape of the distribution of $\varsigma_t = \boldsymbol{\varepsilon}_t^{*\prime} \boldsymbol{\varepsilon}_t^*$.² The most prominent example is the spherical normal distribution, which we denote by $\boldsymbol{\eta} = \mathbf{0}$. Another popular and more empirically realistic example is a standardised multivariate t with ν_0 degrees of freedom, or $i.i.d. t(\mathbf{0}, \mathbf{I}_N, \nu_0)$ for short. As is well known, the multivariate Student t approaches the multivariate normal as $\nu_0 \rightarrow \infty$, but has generally fatter tails. For that reason, we define η as $1/\nu$, which will always remain in the finite range $[0, 1/2)$ under our assumptions. Following Zhou (1993), we also consider two other illustrative examples: a Kotz distribution and a discrete scale mixture of normals.

The original Kotz distribution (see Kotz (1975)) is such that ς_t is a gamma random variable with mean N and variance $N[(N + 2)\kappa_0 + 2]$, where

$$\kappa = E(\varsigma_t^2 | \boldsymbol{\eta}) / [N(N + 2)] - 1$$

is the coefficient of multivariate excess kurtosis of $\boldsymbol{\varepsilon}_t^*$ (see Mardia (1970)). The Kotz distribution nests the multivariate normal distribution for $\kappa = 0$, but it can also be either platykurtic ($\kappa < 0$) or leptokurtic ($\kappa > 0$). Although such a nesting provides an analytically convenient generalisation

²If $\boldsymbol{\varepsilon}_t^*$ is distributed as a spherically symmetric multivariate random vector, then we can write $\boldsymbol{\varepsilon}_t^* = e_t \mathbf{u}_t$, where $\mathbf{u}_t$ is uniformly distributed on the unit sphere surface in $\mathbb{R}^N$, and $e_t = \sqrt{\boldsymbol{\varepsilon}_t^{*\prime} \boldsymbol{\varepsilon}_t^*}$ is a nonnegative random variable that is independent of $\mathbf{u}_t$. Assuming that $E[e_t^2] < \infty$, then $\boldsymbol{\varepsilon}_t^*$ can be standardised by setting $E[e_t^2] = N$, so that $E[\boldsymbol{\varepsilon}_t^*] = \mathbf{0}$ and $V[\boldsymbol{\varepsilon}_t^*] = \mathbf{I}_N$.

of the multivariate normal, the density of a leptokurtic Kotz distribution has a pole at 0, which is a potential drawback from an empirical point of view.

For that reason, we also consider a standardised version of a two-component scale mixture of multivariate normals,³ which can be generated as

$$\boldsymbol{\varepsilon}_t^* = \frac{s_t + (1 - s_t)\sqrt{\varkappa}}{\sqrt{\pi + (1 - \pi)\varkappa}} \cdot \boldsymbol{\varepsilon}_t^\circ, \quad (2)$$

where $\boldsymbol{\varepsilon}_t^\circ$ is a spherical multivariate normal, s_t is an independent Bernoulli variate with $P(s_t = 1) = \pi$ and $\varkappa$ is the variance ratio of the two components. Not surprisingly, $\boldsymbol{\varepsilon}_t^*$ will be a two-component scale mixture of χ_N^{2t} s. As all scale mixtures of normals, the distribution of $\boldsymbol{\varepsilon}_t^*$ is leptokurtic, so that

$$\kappa = \frac{\pi(1 - \pi)(1 - \varkappa)^2}{[\pi + (1 - \pi)\varkappa]^2} \geq 0,$$

with equality if and only if either $\varkappa = 1$, $\pi = 1$ or $\pi = 0$, when it reduces to the spherical normal.⁴ In this sense, a noteworthy property of all discrete mixtures of normals is that their density and moments are always bounded.

Figure 1 plots the densities of a normal, a Student t , a platykurtic Kotz distribution and a discrete scale mixture of normals in the bivariate case. Although they all have concentric circular contours because we have standardised and orthogonalised the two components, their densities can differ substantially in shape, and in particular, in the relative importance of the centre and the tails. They also differ in the degree of cross-sectional “tail dependence” between the components, the normal being the only example in which lack of correlation is equivalent to stochastic independence. Allowing for dependence beyond correlation is particularly important in the context of multiple financial assets, in which the probability of the joint occurrence of several extreme events is regularly underestimated by the multivariate normal distribution.

2.2 Parameter estimation

The purpose of this section is to derive the asymptotic variances of the three estimators of the regression intercepts, $\mathbf{a}$, and slopes, $\mathbf{b}$, mentioned in the introduction (namely, OLS-GMM, as well as elliptically symmetric parametric and semiparametric procedures) under the assumption that the conditional distribution of the innovations $\boldsymbol{\varepsilon}_t^*$ is indeed *i.i.d.* spherical.

³The extension of our analytical results to discrete scale mixtures of normals with multiple components would be fairly straightforward. As is well known, multiple component mixtures can arbitrarily approximate the more empirically realistic continuous mixtures of normals such as symmetric versions of the hyperbolic, normal inverse Gaussian, normal gamma mixtures, Laplace, etc.

⁴In general, though, we require at least sixth moments to globally identify $\boldsymbol{\eta} = (\pi, \varkappa)'$. Since the labels of the components are arbitrary, we also need to impose either $0 \leq \varkappa \leq 1$ or $\pi \geq \frac{1}{2}$.

2.2.1 Maximum likelihood estimators

Let $\phi = (\gamma', \omega', \eta)' \equiv (\theta', \eta)'$ denote the $2N + N(N + 1)/2 + q$ parameters of interest, which we assume variation free. The log-likelihood function of a sample of size T based on a particular parametric spherical assumption will take the form $L_T(\phi) = \sum_{t=1}^T l_t(\phi)$, with $l_t(\phi) = d_t(\theta) + c(\eta) + g[\varsigma_t(\theta), \eta]$, where $d_t(\theta) = -\frac{1}{2} \ln |\Omega|$ corresponds to the Jacobian, $c(\eta)$ to the constant of integration of the assumed density, and $g[\varsigma_t(\theta), \eta]$ to its kernel, where $\varsigma_t(\theta) = \varepsilon_t^{*'}(\theta) \varepsilon_t^*(\theta)$, $\varepsilon_t^*(\theta) = \Omega^{-1/2} \varepsilon_t(\theta)$ and $\varepsilon_t(\theta) = \mathbf{y}_t - \mathbf{a} - \mathbf{b}r_{Mt}$.⁵

Let $\mathbf{s}_t(\phi)$ denote the score function $\partial l_t(\phi)/\partial \phi$, and partition it into three blocks, $\mathbf{s}_{\gamma t}(\phi)$, $\mathbf{s}_{\omega t}(\phi)$, and $\mathbf{s}_{\eta t}(\phi)$, whose dimensions conform to those of γ , ω and η , respectively. A straightforward application of expression (2) in Fiorentini and Sentana (2007) implies that

$$\mathbf{s}_{\gamma t}(\phi) = \begin{pmatrix} 1 \\ r_{Mt} \end{pmatrix} \otimes \delta[\varsigma_t(\theta), \eta] \Omega^{-1} \varepsilon_t(\theta), \quad (3)$$

$$\mathbf{s}_{\omega t}(\phi) = \frac{1}{2} \mathbf{D}'_N [\Omega^{-1} \otimes \Omega^{-1}] \text{vec} \{ \delta[\varsigma_t(\theta), \eta] \varepsilon_t(\theta) \varepsilon_t'(\theta) - \Omega \}, \quad (4)$$

where $\mathbf{D}_N$ is the duplication matrix of order N such that $\text{vec}(\Omega) = \mathbf{D}_N \text{vech}(\Omega)$ (see Magnus and Neudecker (1988)), while the scalar

$$\delta[\varsigma_t(\theta), \eta] = -2 \partial g[\varsigma_t(\theta), \eta] / \partial \varsigma$$

reduces to

$$(N\eta + 1) / [1 - 2\eta + \eta \varsigma_t(\theta)]$$

in the Student t case, to

$$[N(N + 2) \kappa \varsigma_t^{-1}(\theta) + 2] / [(N + 2) \kappa + 2]$$

in the case of the Kotz distribution, to

$$[\pi + (1 - \pi) \varkappa] \cdot \frac{\pi + (1 - \pi) \varkappa^{-(N/2+1)} \exp \left[-\frac{[\pi + (1 - \pi) \varkappa](1 - \varkappa)}{2\varkappa} \varsigma_t(\theta) \right]}{\pi + (1 - \pi) \varkappa^{-N/2} \exp \left[-\frac{[\pi + (1 - \pi) \varkappa](1 - \varkappa)}{2\varkappa} \varsigma_t(\theta) \right]} \quad (5)$$

for the two-component mixture, and to 1 under Gaussianity.⁶

Given correct specification, the results in Crowder (1976) imply that the score vector $\mathbf{s}_t(\phi)$ evaluated at the true parameter values has the martingale difference property. His results also imply that, under suitable regularity conditions, which typically require that both r_{Mt} and r_{Mt}^2

⁵Fiorentini, Sentana and Calzolari (2003) provide expressions for $c(\eta)$ and $g_t[\varsigma_t(\theta), \eta]$ in the multivariate Student case, which under normality collapse to $-(N/2) \log \pi$ and $-\frac{1}{2} \varsigma_t(\theta)$, respectively.

⁶See Fiorentini, Sentana and Calzolari (2003) for numerically reliable expressions for $\mathbf{s}_{\theta t}(\phi)$ and $\mathbf{s}_{\eta t}(\phi)$ in the multivariate t case.

are strictly stationary processes with absolutely summable autocovariances, the asymptotic distribution of the unrestricted ML estimator will be given by the following expression

$$\sqrt{T} \left(\hat{\phi}_{ML} - \phi_0 \right) \longrightarrow N \left[\mathbf{0}, \mathcal{I}^{-1}(\phi_0) \right]$$

where $\mathcal{I}(\phi_0) = E[\mathcal{I}_t(\phi_0)|\phi_0]$,

$$\mathcal{I}_t(\phi) = V[\mathbf{s}_t(\phi)|r_{Mt}, I_{t-1}; \phi] = -E[\mathbf{h}_t(\phi)|r_{Mt}, I_{t-1}; \phi],$$

and $\mathbf{h}_t(\phi)$ denotes the Hessian function $\partial \mathbf{s}_t(\phi)/\partial \phi' = \partial^2 l_t(\phi)/\partial \phi \partial \phi'$. These expressions adopt particularly simple forms for our model of interest:

Proposition 1 *If $\varepsilon_t^*|r_{Mt}, I_{t-1}; \phi$ in (1) is i.i.d. $s(\mathbf{0}, \mathbf{I}_N, \boldsymbol{\eta})$ with density $\exp[c(\boldsymbol{\eta}) + g(\varsigma_t, \boldsymbol{\eta})]$, then the only non-zero elements of $\mathcal{I}_t(\phi_0)$ will be:*

$$\begin{aligned} \mathcal{I}_{\gamma\gamma t}(\phi) &= \mathbf{M}_{ll}(\boldsymbol{\eta}) \begin{pmatrix} 1 & r_{Mt} \\ r_{Mt} & r_{Mt}^2 \end{pmatrix} \otimes \boldsymbol{\Omega}^{-1}, \\ \mathcal{I}_{\omega\omega t}(\phi) &= \frac{\mathbf{M}_{ss}(\boldsymbol{\eta})}{2} \mathbf{D}'_N [\boldsymbol{\Omega}^{-1} \otimes \boldsymbol{\Omega}^{-1}] \mathbf{D}_N + \frac{\mathbf{M}_{ss}(\boldsymbol{\eta}) - 1}{4} \mathbf{D}'_N [\text{vec}(\boldsymbol{\Omega}^{-1}) \text{vec}'(\boldsymbol{\Omega}^{-1})] \mathbf{D}_N, \\ \mathcal{I}_{\omega\eta t}(\phi) &= \frac{1}{2} \mathbf{M}_{sr}(\boldsymbol{\eta}) \mathbf{D}'_N \text{vec}(\boldsymbol{\Omega}^{-1}), \\ \mathcal{I}_{\eta\eta t}(\phi) &= V[\mathbf{s}_{\eta t}(\phi)|\phi] = -E[\mathbf{h}_{\eta t}(\phi)|\phi], \end{aligned}$$

where

$$\begin{aligned} \mathbf{M}_{ll}(\boldsymbol{\eta}) &= E \left\{ \delta^2[\varsigma_t(\boldsymbol{\theta}), \boldsymbol{\eta}] \frac{\varsigma_t(\boldsymbol{\theta})}{N} \middle| \phi \right\} = E \left\{ \frac{2\partial\delta[\varsigma_t(\boldsymbol{\theta}), \boldsymbol{\eta}]}{\partial\varsigma} \frac{\varsigma_t(\boldsymbol{\theta})}{N} + \delta[\varsigma_t(\boldsymbol{\theta}), \boldsymbol{\eta}] \middle| \phi \right\}, \\ \mathbf{M}_{ss}(\boldsymbol{\eta}) &= \frac{N}{N+2} \left[1 + V \left\{ \delta[\varsigma_t(\boldsymbol{\theta}), \boldsymbol{\eta}] \frac{\varsigma_t}{N} \middle| \phi \right\} \right] = E \left\{ \frac{2\partial\delta[\varsigma_t(\boldsymbol{\theta}), \boldsymbol{\eta}]}{\partial\varsigma} \frac{\varsigma_t^2(\boldsymbol{\theta})}{N(N+2)} \middle| \phi \right\} + 1, \\ \mathbf{M}_{sr}(\boldsymbol{\eta}) &= E \left[\left\{ \delta[\varsigma_t(\boldsymbol{\theta}), \boldsymbol{\eta}] \frac{\varsigma_t(\boldsymbol{\theta})}{N} - 1 \right\} \mathbf{e}'_{rt}(\phi) \middle| \phi \right] = -E \left\{ \frac{\varsigma_t(\boldsymbol{\theta})}{N} \frac{\partial\delta[\varsigma_t(\boldsymbol{\theta}), \boldsymbol{\eta}]}{\partial\boldsymbol{\eta}'} \middle| \phi \right\}. \end{aligned}$$

In the multivariate standardised Student t case, in particular:

$$\begin{aligned} \mathbf{M}_{ll}(\eta) &= \frac{\nu(N+\nu)}{(\nu-2)(N+\nu+2)}, \\ \mathbf{M}_{ss}(\eta) &= \frac{(N+\nu)}{(N+\nu+2)}, \\ \mathbf{M}_{sr}(\eta) &= -\frac{2(N+2)\nu^2}{(\nu-2)(N+\nu)(N+\nu+2)}, \end{aligned}$$

which under normality reduce to 1, 1 and 0, respectively (see Fiorentini, Sentana and Calzolari (2003)). As for the Kotz distribution, we can combine the moments of the gamma and reciprocal gamma random variables to show that

$$\mathbf{M}_{ll}(\kappa) = \frac{1}{[(N+2)\kappa+2]^2} \left\{ \frac{N(N+2)^2\kappa^2}{N - [(N+2)\kappa+2]} + 4[(N+2)\kappa+1] \right\}, \quad (6)$$

as long as $\kappa < (N - 2)/(N + 2)$ when $\kappa \neq 0$,

$$M_{ss}(\kappa) = \frac{1}{[(N + 2)\kappa + 2]^2} \left\{ (N + 2)^2 \kappa^2 + \frac{4}{N} [N + (N + 2)\kappa + 2] + 4(N + 2)\kappa \right\},$$

and $M_{sr}(\kappa) = 0 \forall \kappa$, as in the Gaussian case. Finally, we provide the relevant expressions for the case of the two-component scale mixture of normals in Supplemental Appendix D.

The next result follows directly from Proposition 1:

Proposition 2 *If $\varepsilon_t^*|r_{Mt}, I_{t-1}; \phi_0$ in (1) is i.i.d. $s(\mathbf{0}, \mathbf{I}_N, \boldsymbol{\eta}_0)$ with density $\exp[c(\boldsymbol{\eta}) + g(\varsigma_t, \boldsymbol{\eta})]$ such that $M_{ll}(\boldsymbol{\eta}_0) < \infty$, and both r_{Mt} and r_{Mt}^2 are strictly stationary processes with absolutely summable autocovariances, then*

$$\sqrt{T}(\hat{\boldsymbol{\gamma}}_{ML} - \boldsymbol{\gamma}_0) \rightarrow N[\mathbf{0}, \mathcal{I}_{\boldsymbol{\gamma}\boldsymbol{\gamma}}^{-1}(\phi_0)],$$

where

$$\mathcal{I}_{\boldsymbol{\gamma}\boldsymbol{\gamma}}^{-1}(\phi) = \frac{1}{M_{ll}(\boldsymbol{\eta})} \begin{pmatrix} (1 + \mu_M^2/\sigma_M^2) & -\mu_M/\sigma_M^2 \\ -\mu_M/\sigma_M^2 & 1/\sigma_M^2 \end{pmatrix} \otimes \boldsymbol{\Omega}, \quad (7)$$

$\mu_M = E(r_{Mt}|\phi)$ and $\sigma_M^2 = V(r_{Mt}|\phi)$, so that μ_M/σ_M can be interpreted as the Sharpe ratio of the reference portfolio.

Importantly, expression (7) is valid regardless of whether or not the shape parameters $\boldsymbol{\eta}$ are fixed to their true values $\boldsymbol{\eta}_0$, as in the infeasible ML estimator, $\hat{\mathbf{a}}_{IML}$ say, or jointly estimated with $\boldsymbol{\theta}$, as in the unrestricted one, $\hat{\mathbf{a}}_{UML}$ say. The reason is that the scores corresponding to the mean parameters, $\mathbf{s}_{\boldsymbol{\gamma}t}(\phi_0)$, and the scores corresponding to variance and shape parameters, $\mathbf{s}_{\boldsymbol{\omega}t}(\phi_0)$ and $\mathbf{s}_{\boldsymbol{\eta}t}(\phi_0)$, respectively, are asymptotically uncorrelated under our sphericity assumption in view of Proposition 1.

2.2.2 GMM estimators

MacKinlay and Richardson (1991) developed a robust test of mean-variance efficiency by using Hansen's (1982) GMM methodology. If we call $\mathbf{R}_t' \equiv (r_{Mt}, \mathbf{r}_t')$, the orthogonality conditions that they considered are

$$\begin{aligned} E[\mathbf{m}_U(\mathbf{R}_t; \boldsymbol{\gamma})] &= \mathbf{0}, \\ \mathbf{m}_U(\mathbf{R}_t; \boldsymbol{\gamma}) &= \begin{pmatrix} 1 \\ r_{Mt} \end{pmatrix} \otimes \boldsymbol{\varepsilon}_t(\boldsymbol{\gamma}). \end{aligned} \quad (8)$$

The advantage of working within a GMM framework is that under fairly weak regularity conditions inference can be made robust to departures from the assumption of normality, conditional homoskedasticity, serial independence or identity of distribution. But since the above moment conditions exactly identify $\boldsymbol{\gamma}$, the unrestricted GMM estimators coincide with the Gaussian

pseudo ML estimators, which in turn coincide with the equation by equation OLS estimators in the regression of each element of $\mathbf{r}_t$ on a constant and r_{Mt} . An alternative way of reaching the same conclusion is by noticing that the influence function $\mathbf{m}_U(\mathbf{R}_t; \gamma)$ is a full-rank linear transformation with time-invariant weights of the Gaussian pseudo-score $\mathbf{s}_{\gamma t}(\boldsymbol{\theta}, \boldsymbol{\eta} = \mathbf{0})$.⁷

It is convenient to derive an expression for the asymptotic covariance matrix of $\hat{\gamma}_{GMM}$ under *i.i.d.* innovations:

Proposition 3 *If $\boldsymbol{\varepsilon}_t^* | r_{Mt}, I_{t-1}; \phi$ in (1) is i.i.d. $(\mathbf{0}, \mathbf{I}_N)$ with density function $f(\boldsymbol{\varepsilon}_t^*; \boldsymbol{\varrho})$, where $\boldsymbol{\varrho}$ are some shape parameters, and both r_{Mt} and r_{Mt}^2 are strictly stationary processes with absolutely summable autocovariances, then*

$$\sqrt{T}(\hat{\gamma}_{GMM} - \gamma_0) \rightarrow N[\mathbf{0}, \mathcal{C}_{\gamma\gamma}(\phi_0)], \quad (9)$$

where

$$\begin{aligned} \mathcal{C}_{\gamma\gamma}(\phi) &= \mathcal{A}_{\gamma\gamma}^{-1}(\phi) \mathcal{B}_{\gamma\gamma}(\phi) \mathcal{A}_{\gamma\gamma}^{-1}(\phi), \\ \mathcal{A}_{\gamma\gamma}(\phi) &= -E[\mathbf{h}_{\gamma\gamma t}(\boldsymbol{\theta}, \mathbf{0}) | \phi] = E[\mathcal{A}_{\gamma\gamma t}(\phi) | \phi], \\ \mathcal{A}_{\gamma\gamma t}(\phi) &= -E[\mathbf{h}_{\gamma\gamma t}(\boldsymbol{\theta}; \mathbf{0}) | r_{Mt}, I_{t-1}; \phi] = \begin{pmatrix} 1 & r_{Mt} \\ r_{Mt} & r_{Mt}^2 \end{pmatrix} \otimes \boldsymbol{\Omega}^{-1}, \\ \mathcal{B}_{\gamma\gamma}(\phi) &= V[\mathbf{s}_{\gamma t}(\boldsymbol{\theta}, \mathbf{0}) | \phi] = E[\mathcal{B}_{\gamma\gamma t}(\phi) | \phi], \\ \mathcal{B}_{\gamma\gamma t}(\phi) &= V[\mathbf{s}_{\gamma t}(\boldsymbol{\theta}; \mathbf{0}) | r_{Mt}, I_{t-1}; \phi] = \mathcal{A}_{\gamma\gamma t}(\phi), \end{aligned}$$

so that

$$\mathcal{C}_{\gamma\gamma}(\phi_0) = \begin{pmatrix} (1 + \mu_{M0}^2/\sigma_{M0}^2) & -\mu_{M0}/\sigma_{M0}^2 \\ -\mu_{M0}/\sigma_{M0}^2 & 1/\sigma_{M0}^2 \end{pmatrix} \otimes \boldsymbol{\Omega}_0. \quad (10)$$

Importantly, note that $\mathcal{C}_{\gamma\gamma}(\phi_0)$ does not depend on the specific distribution for the innovations that we are considering, regardless of whether or not the conditional distribution of $\boldsymbol{\varepsilon}_t^*$ is spherical, as long as it is *i.i.d.*⁸

2.2.3 HLTV elliptically symmetric semiparametric estimators

HLTV proposed a semiparametric estimator of multivariate linear regression models that updates $\hat{\boldsymbol{\theta}}_{GMM}$ (or any other root- T consistent estimator) by means of a single scoring iteration

⁷The obvious GMM estimator of $\boldsymbol{\omega}$ is given by $\hat{\boldsymbol{\Omega}}_{GMM} = \frac{1}{T} \sum_{t=1}^T \boldsymbol{\varepsilon}_t(\hat{\gamma}_{GMM}) \boldsymbol{\varepsilon}_t'(\hat{\gamma}_{GMM})$, which is the sample analogue to the residual covariance matrix.

⁸The assumption of constant conditional third and fourth moments implicit in the assumption of *i.i.d.* innovations also implies that the optimal GMM estimators of Meddahi and Renault (1998) do not offer any asymptotic efficiency gains over $\hat{\mathbf{a}}_{GMM}$.

without line searches. The crucial ingredient of their method is the so-called elliptically symmetric semiparametric efficient score (see e.g. Proposition 7 in Fiorentini and Sentana (2007)):

$$\hat{\mathbf{s}}_{\theta t}(\phi_0) = \mathbf{s}_{\theta t}(\phi_0) - \mathbf{W}_s(\phi_0) \left\{ \left[\delta[\zeta_t(\boldsymbol{\theta}_0), \boldsymbol{\eta}_0] \frac{\zeta_t(\boldsymbol{\theta}_0)}{N} - 1 \right] - \frac{2}{(N+2)\kappa_0 + 2} \left[\frac{\zeta_t(\boldsymbol{\theta}_0)}{N} - 1 \right] \right\},$$

where

$$\mathbf{W}'_s(\phi) = \begin{bmatrix} \mathbf{0} & \mathbf{0} & \frac{1}{2} \text{vec}'(\boldsymbol{\Omega}^{-1}) \mathbf{D}_N \end{bmatrix}$$

in the case of model (1). In fact, the special structure of $\mathbf{W}_s(\phi)$ implies that we can update the GMM estimator of $\boldsymbol{\gamma}$ by means of the following simple Berndt, Hall, Hall and Hausman (1974) (BHHH) correction:

$$\left[\sum_{t=1}^T \mathbf{s}_{\gamma t}(\phi_0) \mathbf{s}'_{\gamma t}(\phi_0) \right]^{-1} \sum_{t=1}^T \mathbf{s}_{\gamma t}(\phi_0), \quad (11)$$

which does not require the computation of $\hat{\mathbf{s}}_{\omega t}(\phi_0)$. In practice, of course, $\mathbf{s}_{\gamma t}(\phi_0)$ has to be replaced by a semiparametric estimate obtained from the joint density of $\boldsymbol{\varepsilon}_t^*$. However, the elliptical symmetry assumption allows one to obtain such an estimate from a nonparametric estimate of the univariate density of ζ_t , $h(\zeta_t; \boldsymbol{\eta})$, avoiding in this way the curse of dimensionality (see HLV and appendix B1 in Fiorentini and Sentana (2007) for details).

Proposition 7 in Fiorentini and Sentana (2007) shows that the elliptically symmetric semiparametric efficiency bound will be given by:

$$\hat{\mathcal{S}}(\phi_0) = \mathcal{I}_{\theta\theta}(\phi_0) - \mathbf{W}_s(\phi_0) \mathbf{W}'_s(\phi_0) \cdot \left\{ \left[\frac{N+2}{N} M_{ss}(\boldsymbol{\eta}_0) - 1 \right] - \frac{4}{N[(N+2)\kappa_0 + 2]} \right\},$$

which implies that $\hat{\mathcal{S}}_{\gamma\gamma}(\phi_0) = \mathcal{I}_{\gamma\gamma}(\phi_0)$ in our case in view of the structure of $\mathbf{W}_s(\phi_0)$. This result confirms that the HLV estimator of $\boldsymbol{\gamma}$ is adaptive.⁹

2.3 Relative efficiency of estimators and test procedures under correct specification

Let $\hat{\mathbf{a}}$ denote any of the asymptotically normal, root- T estimators of $\mathbf{a}$ analysed in the previous section, and denote its asymptotic covariance matrix by $V(\hat{\mathbf{a}})$. To test $H_0 : \mathbf{a} = \mathbf{0}$, we can in principle use any of the trinity of classical hypothesis tests, namely, Wald (W_T), Lagrange Multiplier (LM_T) and Likelihood Ratio/Distance Metric test (LR_T). For the sake of concreteness, though, we shall centre our discussion around the Wald test, which examines whether the

⁹HLV also consider alternative estimators that iterate the semiparametric adjustment (11) until it becomes negligible. However, since they have the same asymptotic distribution, we shall not discuss them separately.

homogeneity constraints imposed by H_0 are approximately satisfied by $\hat{\mathbf{a}}$.¹⁰ More formally,

$$W_T = T \cdot \hat{\mathbf{a}}' V^{-1}(\hat{\mathbf{a}}) \hat{\mathbf{a}}.$$

As is well known, W_T will be asymptotically distributed as a χ^2 with N degrees of freedom under the null, and as a non-central χ^2 with the same degrees of freedom and non-centrality parameter $\boldsymbol{\delta}' V^{-1}(\hat{\mathbf{a}}) \boldsymbol{\delta}$ under the Pitman sequence of local alternatives $H_l : \mathbf{a} = \boldsymbol{\delta} / \sqrt{T}$ (see Newey and MacFadden (1994)). In contrast, W_T will diverge to infinity for fixed alternatives of the form $H_f : \mathbf{a} = \boldsymbol{\delta}$, which makes it a consistent test. In that case, we can use Theorem 1 in Geweke (1981) to show that

$$p \lim \frac{1}{T} W_T = \boldsymbol{\delta}' V^{-1}(\hat{\mathbf{a}}) \boldsymbol{\delta}$$

coincides with Bahadur's (1960) definition of the approximate slope of the Wald test. This expression differs from the non-centrality parameter in that the covariance matrix is no longer evaluated under the null. However, since $V(\hat{\mathbf{a}})$ does not depend on $\mathbf{a}$ when the true distribution is elliptical for any of the estimators considered in the previous section, both comparison criteria coincide.

In addition, since $V(\hat{\mathbf{a}}_{GMM}) = \mathcal{C}_{\mathbf{aa}}(\boldsymbol{\phi}_0)$ in view of (9), while $V(\hat{\mathbf{a}}_{IML}) = V(\hat{\mathbf{a}}_{FML}) = V(\hat{\mathbf{a}}_{HLV}) = \mathbf{M}_U^{-1}(\boldsymbol{\eta}_0) \mathcal{C}_{\mathbf{aa}}(\boldsymbol{\phi}_0)$ in view of (7), we can use $\mathbf{M}_U(\boldsymbol{\eta}_0)$ to measure the relative efficiency of the GMM-based test procedure regardless of the value of $\boldsymbol{\delta}$. In fact, since the proportionality applies not only to $\mathbf{a}$ but also to $\mathbf{b}$, we can also use $\mathbf{M}_U(\boldsymbol{\eta}_0)$ to measure the relative efficiency of the estimators of both regression intercepts and slopes in other contexts.

We know from Proposition 9 in Fiorentini and Sentana (2007) that $\mathbf{M}_U(\boldsymbol{\eta}_0) = 1$ if and only if the true conditional distribution is indeed normal. Otherwise, $0 \leq \mathbf{M}_U^{-1}(\boldsymbol{\eta}_0) < 1$. This means that while there is no asymptotic efficiency loss in estimating $\boldsymbol{\eta}$ when the true conditional distribution is Gaussian, the efficiency gains could be potentially very large for other elliptical distributions. In the multivariate Student t case with $\nu_0 > 2$, in particular, the relative efficiency ratio becomes $(\nu_0 - 2)(\nu_0 + N + 2) / [\nu_0(\nu_0 + N)]$. For any given N , this ratio is monotonically increasing in ν_0 , and approaches 1 from below as $\nu_0 \rightarrow \infty$, and 0 from above as $\nu_0 \rightarrow 2^+$. At the same time, this ratio is decreasing in N for a given ν_0 , which reflects the fact that the Student t information matrix is “increasing” in N . Figure 2a presents a plot of this efficiency ratio as a function of η for several values of N . Similarly, Figure 2b presents the efficiency ratio as a function of κ for different values of N in the case of the Kotz distribution, where we have obtained $\mathbf{M}_U^{-1}(\kappa)$

¹⁰Another advantage of the Wald test, shared with the LM test, is that it is easy to robustify with respect to misspecification, unlike the LR test.

from (6). In this sense, it is worth mentioning that the excess kurtosis coefficient of any elliptical distribution is bounded from below by $-2/(N + 2)$, which is the excess kurtosis of a random vector that is uniformly distributed on the unit sphere. This explains why the lower limit of admissible values for κ gets closer and closer to 0 from below as N increases. Finally, Figure 2c contains the corresponding efficiency ratios for a two-component scale mixture of normals in which $\pi = \frac{1}{2}$ as a function of the relative variance parameter $\varkappa$. As expected, the GMM and ML/HLV estimators are equally efficient for $\varkappa = 1$, since in that case the mixture of normals is itself normal. Once again, though, the relative efficiency of the ML/HLV estimators increases as we move away from normality, the more so the bigger N is.

We can assess the power implications of such efficiency gains by computing the probability of rejecting the null hypothesis when it is false as a function of $\mathbf{a}$ under the assumption that the asymptotic non-central chi-square distributions of the Wald tests implied by (7) or (9) provide reliable rejection probabilities in finite samples. The results for $T = 500$ at the usual 5% level are plotted in Figure 3 under the fairly innocuous assumptions that $\mathbf{\Omega} = \mathbf{I}_N$, $\sqrt{12}\mu_M/\sigma_M = \frac{1}{2}$ and $\mathbf{a} = a\ell_N$, with $\ell'_N = (1, \dots, 1)'$ and $a \in [0, .2]$. We consider two examples of elliptical distributions whose $M_U(\boldsymbol{\eta})$ correspond to those of a Student t with 8 and 20 degrees of freedom, respectively. Not surprisingly, the power of all tests increases as we depart from the null. Similarly, their power also increases with the number of series due to the lack of cross-sectional correlation of the regression residuals. More importantly, the power of the efficient tests is always larger than the power of the GMM tests, although the differences are unsurprisingly small when the true distribution is not too far away from the normal.

In empirical applications, it is customary to pay attention not only to the joint Wald test of $H_0 : \mathbf{a} = \mathbf{0}$, but also to individual tests of the form $H_0 : a_i = 0$ for some i between 1 and N . Given that the asymptotic power of such partial tests under either local or fixed alternatives will depend on the non-centrality parameter $a_i^2/V(\hat{a}_i)$, the discussion in the previous paragraphs applies directly to those individual Wald tests too (see Sentana (2008) for a discussion on the advantages and disadvantages of joint versus individual tests on these contexts).

3 Misspecification analysis

In section 2.2 we obtained the asymptotic covariance matrix of the estimators of the regression intercepts, $\mathbf{a}$, and slopes, $\mathbf{b}$, under the assumption that the model used for estimation purposes in the parametric maximum likelihood procedure and the data generation process coincide. The

main purpose of this section is to study how our earlier results change in situations in which the conditional distribution assumed for estimation purposes differs from the true one. In those cases in which the parametric maximum likelihood estimators remain consistent, we will provide their asymptotic variances, compare them to the full information parametric efficiency bounds and the asymptotic variances of the GMM estimators, and explain how to robustify inference in practice. We omit a discussion of the testing implications of the relative efficiency of the estimators because the analysis is entirely analogous to the one in section 2.3. Given that the estimated model will be incorrect, we consider a restricted parametric ML estimator that fixes the shape parameters to some arbitrary value $\bar{\eta} \neq \mathbf{0}$ in place of the infeasible ML estimator discussed in section 2.2.1.

3.1 Misspecified elliptical distributions for the innovations

We begin by deriving the asymptotic distribution of the unrestricted and restricted ML estimators when the true conditional distribution of $\mathbf{r}_t$ given r_{Mt} and their past is *i.i.d.* elliptical, but does not coincide with the distribution assumed for estimation purposes. For the sake of concreteness, we assume in what follows that those parametric (pseudo) ML estimators are based on the erroneous assumption that $\varepsilon_t^*|r_{Mt}, I_{t-1}; \boldsymbol{\theta}, \eta \sim i.i.d. t(\mathbf{0}, \mathbf{I}_N, \nu)$. Nevertheless, our results can be trivially extended to any other spherically-based likelihood estimators, as the only advantage of the Student t likelihood for our purposes is the fact that its limiting relationship to the Gaussian distribution can be made explicit. In this context, the restricted t -based PML estimator should be understood as the one that fixes the parameter η to some $\bar{\eta}$ between 0 and $\frac{1}{2}$.

For simplicity, we shall also define the pseudo-true values of $\boldsymbol{\theta}$ and η as consistent roots of the expected t pseudo log-likelihood score, which under appropriate regularity conditions will maximise the expected value of the t pseudo log-likelihood function. Specifically, if we define the pseudo-true values of $\boldsymbol{\phi}$ as the values of $\mathbf{a}$, $\mathbf{b}$, $\boldsymbol{\Omega}$, and η that will set to zero the expected value of the score vector, $\mathbf{s}_t(\boldsymbol{\phi}_0)$, where the expected value is taken with respect to the true distribution of the data, then we can derive the following result, which particularises to our context Proposition 15 in Fiorentini and Sentana (2007):

Proposition 4 *If $\varepsilon_t^*|r_{Mt}, I_{t-1}; \boldsymbol{\varphi}_0$ in (1) is *i.i.d.* $s(\mathbf{0}, \mathbf{I}_N, \boldsymbol{\varrho}_0)$ but not t and $\kappa_0 \leq 0$, where $\boldsymbol{\varphi}_0 = (\gamma_0, \boldsymbol{\omega}_0, \boldsymbol{\varrho}_0)$, then:*

1. *The pseudo-true value of the unrestricted Student t -based ML estimator of $\boldsymbol{\phi} = (\gamma, \boldsymbol{\omega}, \eta)'$, $\boldsymbol{\phi}_\infty$, is such that γ_∞ and $\boldsymbol{\omega}_\infty$ are equal to their corresponding true values γ_0 and $\boldsymbol{\omega}_0$,*

respectively, and $\eta_\infty = 0$.

$$2. \sqrt{T}(\hat{\gamma}_{UML} - \hat{\gamma}_{GMM}) = o_p(1).$$

Intuitively, the reason is that since η must be estimated subject to the non-negativity restriction $\eta \geq 0$, the most platykurtic Student t distribution that one can obtain is the normal distribution, in which case the unrestricted Student t -based PML estimator coincides with the GMM one.

The following result derives the asymptotic distribution of the unrestricted t -based PML estimator of θ in the more realistic case of leptokurtic disturbances. To keep the algebra simple, we will reparametrise Ω as $\tau \Upsilon(\mathbf{v})$, so that $\vartheta = (\gamma, \mathbf{v}, \tau)$, where $\mathbf{v}$ are $N(N+1)/2 - 1$ parameters that ensure that $|\Upsilon(\mathbf{v})| = 1 \forall \mathbf{v}$. In other words, our reparametrisation will be such that

$$\tau = |\Omega|^{1/N} \quad (12)$$

and

$$\Upsilon(\mathbf{v}) = \Omega / |\Omega|^{1/N}. \quad (13)$$

Nevertheless, the t -based ML estimator of γ will be unaffected by this change.

Proposition 5 *If $\varepsilon_t^* | r_{Mt}, I_{t-1}; \varphi_0$ is i.i.d. $s(\mathbf{0}, \mathbf{I}_N, \varrho_0)$ but not t with $\kappa_0 > 0$, where $\varphi_0 = (\gamma_0, \mathbf{v}_0, \tau_0, \varrho_0)$, then:*

1. *The pseudo-true value of the unrestricted Student- t based ML estimator of $\phi = (\gamma, \mathbf{v}, \tau, \eta)'$, ϕ_∞ , is such that γ_∞ and $\mathbf{v}_\infty$ are equal to their corresponding true values γ_0 and $\mathbf{v}_0$.*
2. *$\mathcal{O}(\phi_\infty; \varphi_0) = E[\mathcal{O}_t(\phi_\infty; \varphi_0) | \varphi_0]$ and $\mathcal{H}(\phi_\infty; \varphi_0) = E[\mathcal{H}_t(\phi_\infty; \varphi_0) | \varphi_0]$ will be block diagonal between $(\gamma, \mathbf{v})$ and (τ, η) , where both*

$$\mathcal{O}_t(\phi_\infty; \varphi_0) = V[\mathbf{s}_t(\phi_\infty) | r_{Mt}, I_{t-1}; \varphi_0]$$

and

$$\mathcal{H}_t(\phi_\infty; \varphi_0) = -E[\mathbf{h}_t(\phi_\infty) | r_{Mt}, I_{t-1}; \varphi_0]$$

will share the structure of $\mathcal{I}_t(\phi_\infty; \varphi_0)$ in Proposition 1, with $\mathcal{O}_{\eta\eta}(\phi; \varphi) = V[s_{\eta t}(\phi) | \varphi]$,

$$\mathcal{H}_{\eta\eta}(\phi; \varphi) = -E[h_{\eta\eta t}(\phi) | \varphi],$$

$$M_{ll}^O(\phi; \varphi) = E \{ \delta^2[\varsigma_t(\boldsymbol{\vartheta}), \eta] \cdot [\varsigma_t(\boldsymbol{\vartheta})/N] | \varphi \} \quad (14)$$

$$M_{ss}^O(\phi; \varphi) = N(N+2)^{-1} [1 + V \{ \delta[\varsigma_t(\boldsymbol{\gamma}, \boldsymbol{v}, \tau), \eta] \cdot [\varsigma_t(\boldsymbol{\vartheta})/N] | \varphi \}],$$

$$M_{sr}^O(\phi; \varphi) = E [\{ \delta[\varsigma_t(\boldsymbol{\vartheta}), \eta] \cdot [\varsigma_t(\boldsymbol{\vartheta})/N] - 1 \} s_{\eta t}(\phi) | \varphi],$$

$$M_{ll}^H(\phi; \varphi) = E \{ 2\partial\delta[\varsigma_t(\boldsymbol{\vartheta}), \eta] / \partial\varsigma \cdot [\varsigma_t(\boldsymbol{\vartheta})/N] + \delta[\varsigma_t(\boldsymbol{\vartheta}), \eta] | \varphi \}, \quad (15)$$

$$M_{ss}^H(\phi; \varphi) = E \{ 2\partial\delta[\varsigma_t(\boldsymbol{\vartheta}), \eta] / \partial\varsigma \cdot \varsigma_t^2(\boldsymbol{\vartheta}) / [N(N+2)] | \varphi \} + 1,$$

$$M_{sr}^H(\phi; \varphi) = -E \{ [\varsigma_t(\boldsymbol{\vartheta})/N] \cdot \partial\delta[\varsigma_t(\boldsymbol{\vartheta}), \eta] / \partial\eta | \varphi \}.$$

Intuitively, what the first part of this proposition shows is that

$$E \{ \mathbf{s}_{\gamma t}[\boldsymbol{\gamma}_0, \boldsymbol{v}_0, \tau_\infty(\eta), \eta] | \boldsymbol{\gamma}_0, \boldsymbol{v}_0, \tau_0, \boldsymbol{\varrho}_0 \} = \mathbf{0},$$

$$E \{ \mathbf{s}_{vt}[\boldsymbol{\gamma}_0, \boldsymbol{v}_0, \tau_\infty(\eta), \eta] | \boldsymbol{\gamma}_0, \boldsymbol{v}_0, \tau_0, \boldsymbol{\varrho}_0 \} = \mathbf{0},$$

for any elliptical distribution for the innovations, which implies in particular that the t -based PML estimators of $\mathbf{a}$ and $\mathbf{b}$ will be consistent. In contrast, when $\kappa > 0$ we cannot find any distribution for $\boldsymbol{\varepsilon}_t^*$ other than the multivariate t for which

$$E [s_{\tau t}(\phi) | \boldsymbol{\varphi}_0] = 0,$$

$$E [s_{\eta t}(\phi) | \boldsymbol{\varphi}_0] = 0,$$

which means that the overall scale parameter τ will be inconsistently estimated.

The asymptotic distribution of the unrestricted t -based PML estimator of $\boldsymbol{\gamma}$ follows immediately from Proposition 5:

Corollary 1 *If $\boldsymbol{\varepsilon}_t^* | r_{Mt}, I_{t-1}; \boldsymbol{\varphi}_0$ is i.i.d. $s(\mathbf{0}, \mathbf{I}_N, \boldsymbol{\varrho}_0)$ but not t with $\kappa_0 > 0$, where $\boldsymbol{\varphi}_0 = (\boldsymbol{\gamma}_0, \boldsymbol{v}_0, \lambda_0, \boldsymbol{\varrho}_0)$, and both r_{Mt} and r_{Mt}^2 are strictly stationary processes with absolutely summable autocovariances, then:*

$$\sqrt{T} (\hat{\boldsymbol{\gamma}}_{UML} - \boldsymbol{\gamma}_0) \rightarrow N \left[\mathbf{0}, \frac{M_{ll}^O(\phi_\infty; \boldsymbol{\varphi}_0)}{[M_{ll}^H(\phi_\infty; \boldsymbol{\varphi}_0)]^2} \cdot \frac{1}{\lambda_\infty} \mathcal{C}_{\boldsymbol{\gamma}\boldsymbol{\gamma}}(\boldsymbol{\varphi}_0) \right], \quad (16)$$

where $\lambda_\infty = \tau_0 / \tau_\infty$.

In practice, it is trivial to obtain consistent estimators of the robust asymptotic covariance matrix in (16) because $\mathcal{C}_{\boldsymbol{\gamma}\boldsymbol{\gamma}}(\hat{\boldsymbol{\varphi}}_{UML})$ converges in probability to $\lambda_\infty^{-1} \mathcal{C}_{\boldsymbol{\gamma}\boldsymbol{\gamma}}(\boldsymbol{\varphi}_0)$ in view of (10) and (13), and both $M_{ll}^O(\phi_\infty; \boldsymbol{\varphi}_0)$ and $M_{ll}^H(\phi_\infty; \boldsymbol{\varphi}_0)$ can be consistently estimated by using the sample analogues of (14) and (15), respectively, evaluated at $\hat{\boldsymbol{\varphi}}_{UML}$.

The analysis of the restricted t -based PML estimator is entirely analogous, except for the fact that the pseudo-true value of τ becomes $\tau_\infty(\bar{\eta})$, as opposed to $\tau_\infty = \tau_\infty(\eta_\infty)$.

A natural question in this context is a comparison of the efficiency of the t -based pseudo ML estimator and the GMM estimator when the distribution is elliptical but not t . We answer this question by assuming that the conditional distribution is either normal, Kotz, or the two-component scale mixture of normals discussed in section 2.1. It turns out that in all three cases we can obtain analytical expressions for the efficiency ratio $M_{ll}^O(\phi_\infty; \varphi_0) / \{\lambda_\infty [M_{ll}^H(\phi_\infty; \varphi_0)]^2\}$ (see Supplemental Appendix B).

The panels of Figure 4 present the relative efficiency of these two estimators of γ as a function of $\bar{\eta}$ for five cross-sectional dimensions. In addition, the vertical straight lines indicate the position of the pseudo-true values when we also estimate η , while the horizontal lines on the right indicate the relative efficiency of the correct ML estimator. As expected, if the true conditional distribution is Gaussian (Figure 4a), then the restricted ML estimator that makes the erroneous assumption that it is a Student t with $\bar{\eta}^{-1}$ degrees of freedom is inefficient relative to the GMM estimator, the more so the larger the value of $\bar{\eta}$. Nevertheless, this inefficiency becomes smaller and less sensitive to $\bar{\eta}$ as the number of assets increases. But of course $\eta_\infty = 0$ in this case in view of Proposition 4, which suggests that estimating η is clearly beneficial under misspecification. In fact, the restricted t -based PML estimator seems to be strictly more efficient than the GMM one at the pseudo-true value of η when the true conditional distribution is leptokurtic. This is indeed true for any value of $\bar{\eta}$ for a Kotz distribution with $\kappa_0 = 1/8$ (Figure 4b), which is equal to the excess kurtosis of a t with 20 degrees of freedom, as well as for a two-component mixture of normals with $\pi = 1/2$ and $\kappa_0 = 1/4$ (Figure 4c), which coincides with the excess kurtosis of the more empirically realistic t distribution with 12 degrees of freedom. It is noteworthy that as N increases the restricted t -based PML estimator tends to achieve the full efficiency of the ML estimator for any $\bar{\eta} > 0$.¹¹ Whether such efficiency gains always accrue at the pseudo true value of η is left for future research.¹²

¹¹The values corresponding to $N = \infty$ in Figures 4 and 5 are intended to reflect the maximum efficiency gains that could be obtained by increasing the number of series; and hence, they are derived under sequential limits, i.e., T converges to infinity with a fixed N and then N converges to infinity. In this sense, $\lim_{N \rightarrow \infty} M_{ll}(\boldsymbol{\eta}) = (1 - \kappa)^{-1}$ in the case of both Kotz innovations and discrete scale mixture of normals innovations.

¹²Another pending issue is whether η_∞ is always larger than $\max(0, \kappa_0) / [4 \max(0, \kappa_0) + 2]$, which is the value of η that matches the excess kurtosis of the t distribution with the excess kurtosis of the true distribution, as Figures 4a and 4b seem to suggest.

3.2 Elliptical distributions for returns

In this section we explicitly study the framework analysed by MacKinlay and Richardson (1991) and Kan and Zhou (2006), who considered a *joint* distribution of excess returns for the N assets $\mathbf{r}_t$ and the reference portfolio, r_{Mt} . When the joint distribution of $\mathbf{R}_t$ is *i.i.d.* Gaussian, the distribution of $\mathbf{r}_t$ conditional on r_{Mt} must also be normal, with a mean $\mathbf{a} + \mathbf{b}r_{Mt}$ that is a linear function of r_{Mt} , and a covariance matrix $\mathbf{\Omega}$ that does not depend on r_{Mt} . However, while the linearity of the conditional mean will be preserved when $\mathbf{R}_t$ is elliptically distributed but non-Gaussian, the conditional covariance matrix will no longer be independent of r_{Mt} . For instance, if we assume that $\mathbf{\Sigma}^{-1/2}(\boldsymbol{\rho})[\mathbf{R}_t - \boldsymbol{\mu}(\boldsymbol{\rho})] \sim i.i.d. \ t(\mathbf{0}, \mathbf{I}_{N+1}, \eta)$, where

$$\boldsymbol{\mu}(\boldsymbol{\rho}) = \begin{pmatrix} \mu_M \\ \mathbf{a} + \mathbf{b}\mu_M \end{pmatrix}, \quad (17)$$

$$\mathbf{\Sigma}(\boldsymbol{\rho}) = \begin{pmatrix} \sigma_M^2 & \sigma_M^2 \mathbf{b}' \\ \sigma_M^2 \mathbf{b} & \sigma_M^2 \mathbf{b}\mathbf{b}' + \mathbf{\Omega} \end{pmatrix}, \quad (18)$$

and $\boldsymbol{\rho}' = (\mathbf{a}', \mathbf{b}', \boldsymbol{\omega}', \mu_M, \sigma_M^2)$, then

$$\begin{aligned} E[\mathbf{r}_t | r_{Mt}; \boldsymbol{\rho}, \eta] &= \mathbf{a} + \mathbf{b}r_{Mt}, \\ V[\mathbf{r}_t | r_{Mt}; \boldsymbol{\rho}, \eta] &= \left(\frac{\nu - 2}{\nu - 1} \right) \left[1 + \frac{(r_{Mt} - \mu_M)^2}{(\nu - 2)\sigma_M^2} \right] \mathbf{\Omega} \equiv \boldsymbol{\Psi}_t(\boldsymbol{\rho}, \eta), \end{aligned}$$

which means that model (1) will be misspecified due to contemporaneous, conditionally heteroskedastic innovations. In other words, the variances and covariances of the regression residuals will be a function of the regressor.

As MacKinlay and Richardson (1991) pointed out, the GMM estimator of $\boldsymbol{\gamma}$ remains consistent in this case. In addition, we know from Lemma D3 in Peñaranda and Sentana (2008) that if $\mathbf{R}_t$ is independently and identically distributed as an elliptical random vector with mean $\boldsymbol{\mu}(\boldsymbol{\rho})$, covariance matrix $\mathbf{\Sigma}(\boldsymbol{\rho})$, and bounded fourth moments, then the asymptotic covariance matrix of $\sqrt{T}\bar{\mathbf{m}}_U(\mathbf{R}_t; \boldsymbol{\gamma}_0)$ will be given by

$$\mathbf{S}_U(\boldsymbol{\gamma}_0) = \begin{bmatrix} 1 & \mu_{M0} \\ \mu_{M0} & (\kappa_0 + 1)\sigma_{M0}^2 + \mu_{M0}^2 \end{bmatrix} \otimes \mathbf{\Omega}_0,$$

where $\bar{\mathbf{m}}_U(\mathbf{R}_t; \boldsymbol{\gamma}_0)$ is the sample mean of $\mathbf{m}_U(\mathbf{R}_t; \boldsymbol{\gamma}_0)$ in (8). Hence,

$$V(\hat{\boldsymbol{\gamma}}_{GMM}) = \begin{bmatrix} 1 + (1 + \kappa_0)(\mu_{M0}^2/\sigma_{M0}^2) & -(1 + \kappa_0)(\mu_{M0}/\sigma_{M0}^2) \\ -(1 + \kappa_0)(\mu_{M0}/\sigma_{M0}^2) & (1 + \kappa_0)(1/\sigma_{M0}^2) \end{bmatrix} \otimes \mathbf{\Omega}_0. \quad (19)$$

In this sense, note that the only difference with respect to (10) is the factor $(1 + \kappa_0)$. Although in principle one could use the sample analogue of (19), in practice, we will typically estimate

$V(\hat{\gamma}_{GMM})$ by using White (1980) heteroskedastic robust standard errors. Specifically, we would use the sandwich expression $\mathcal{C}_{\gamma\gamma}(\phi) = \mathcal{A}_{\gamma\gamma}^{-1}(\phi)\mathcal{B}_{\gamma\gamma}(\phi)\mathcal{A}_{\gamma\gamma}^{-1}(\phi)$, but this time with

$$\hat{\mathcal{B}}_{\gamma\gamma}(\phi) = \frac{1}{T} \sum_{t=1}^T \mathbf{s}_{\gamma t}(\boldsymbol{\theta}; \mathbf{0}) \mathbf{s}_{\gamma t}'(\boldsymbol{\theta}; \mathbf{0}) \quad (20)$$

while we will continue to use

$$\hat{\mathcal{A}}_{\gamma\gamma}(\phi) = \frac{1}{T} \sum_{t=1}^T \begin{pmatrix} 1 & r_{Mt} \\ r_{Mt} & r_{Mt}^2 \end{pmatrix} \otimes \boldsymbol{\Omega}^{-1}. \quad (21)$$

At the other extreme of the efficiency range, we can consider the joint ML estimator that makes the correct assumption that $\boldsymbol{\Sigma}^{-1/2}(\boldsymbol{\rho})[\mathbf{R}_t - \boldsymbol{\mu}(\boldsymbol{\rho})] \sim i.i.d. s(\mathbf{0}, \mathbf{I}_{N+1}, \boldsymbol{\eta})$, whose asymptotic distribution can be obtained from the following result:

Proposition 6 *Let $\boldsymbol{\epsilon}_t^*(\boldsymbol{\rho}) = \boldsymbol{\Sigma}^{-1/2}(\boldsymbol{\rho})\boldsymbol{\epsilon}_t(\boldsymbol{\rho})$, where $\boldsymbol{\epsilon}_t(\boldsymbol{\rho}) = \mathbf{R}_t - \boldsymbol{\mu}(\boldsymbol{\rho})$, $\boldsymbol{\mu}(\boldsymbol{\rho})$ and $\boldsymbol{\Sigma}(\boldsymbol{\rho})$ are defined in (17) and (18), respectively, and $\boldsymbol{\rho}' = (\mathbf{a}', \mathbf{b}', \boldsymbol{\omega}', \mu_M, \sigma_M^2)$. If $\boldsymbol{\epsilon}_t^*(\boldsymbol{\rho}_0)|I_{t-1}; \boldsymbol{\rho}_0, \boldsymbol{\eta}_0 \sim i.i.d. s(\mathbf{0}, \mathbf{I}_{N+1}, \boldsymbol{\eta}_0)$ with density $\exp[c_{N+1}(\boldsymbol{\eta}) + g_{N+1}(\varsigma_t, \boldsymbol{\eta})]$, then the only non-zero elements of the information matrix other than $\mathcal{I}_{\eta\eta}(\phi) = V[\mathbf{s}_{\eta t}(\phi)|\phi] = -E[\mathbf{h}_{\eta\eta t}(\phi)|\phi]$ will be:*

$$\begin{aligned} \mathcal{I}_{\gamma\gamma}(\phi) &= \left[M_{ll}(\boldsymbol{\eta}) \begin{pmatrix} 1 & \mu_M \\ \mu_M & \mu_M^2 \end{pmatrix} + M_{ss}(\boldsymbol{\eta}_0) \begin{pmatrix} 0 & 0 \\ 0 & \sigma_M^2 \end{pmatrix} \right] \otimes \boldsymbol{\Omega}^{-1}, \\ \mathcal{I}_{\omega\omega}(\phi) &= \frac{M_{ss}(\boldsymbol{\eta})}{2} \mathbf{D}'_N [\boldsymbol{\Omega}^{-1} \otimes \boldsymbol{\Omega}^{-1}] \mathbf{D}_N + \frac{M_{ss}(\boldsymbol{\eta}) - 1}{4} \mathbf{D}'_N [\text{vec}(\boldsymbol{\Omega}^{-1}) \text{vec}'(\boldsymbol{\Omega}^{-1})] \mathbf{D}'_N, \\ \mathcal{I}_{\mu_M \mu_M}(\phi) &= \frac{M_{ll}(\boldsymbol{\eta})}{\sigma_M^2}, \quad \mathcal{I}_{\sigma_M^2 \sigma_M^2}(\phi) = \frac{3M_{ss}(\boldsymbol{\eta}) - 1}{4\sigma_M^4}, \\ \mathcal{I}_{\omega \sigma_M^2}(\phi) &= \frac{M_{ss}(\boldsymbol{\eta}) - 1}{4\sigma_M^2} \mathbf{D}'_N \text{vec}(\boldsymbol{\Omega}^{-1}), \quad \mathcal{I}_{\omega\eta}(\phi) = \frac{M_{sr}(\boldsymbol{\eta})}{2} \mathbf{D}'_N \text{vec}(\boldsymbol{\Omega}^{-1}), \quad \mathcal{I}_{\sigma_M^2 \eta}(\phi) = \frac{M_{sr}(\boldsymbol{\eta})}{2\sigma_M^2}, \end{aligned}$$

where $M_{ll}(\boldsymbol{\eta}_0)$, $M_{ss}(\boldsymbol{\eta}_0)$ and $M_{sr}(\boldsymbol{\eta}_0)$ for this $(N+1)$ -dimensional distribution are defined analogously to Proposition 1.

We can use this Proposition to extend the result in equation (31) in Kan and Zhou (2006) and show that

$$V(\hat{\gamma}_{JML}) = \begin{bmatrix} M_{ll}^{-1}(\boldsymbol{\eta}_0) + M_{ss}^{-1}(\boldsymbol{\eta}_0)(\mu_{M0}^2/\sigma_{M0}^2) & -M_{ss}^{-1}(\boldsymbol{\eta}_0)(\mu_{M0}/\sigma_{M0}^2) \\ -M_{ss}^{-1}(\boldsymbol{\eta}_0)(\mu_{M0}/\sigma_{M0}^2) & M_{ss}^{-1}(\boldsymbol{\eta}_0)(1/\sigma_{M0}^2) \end{bmatrix} \otimes \boldsymbol{\Omega}_0, \quad (22)$$

where $\hat{\boldsymbol{\theta}}_{JML}$ denotes the joint ML estimator that makes the correct assumption that $\boldsymbol{\Sigma}^{-1/2}(\boldsymbol{\rho})[\mathbf{R}_t - \boldsymbol{\mu}(\boldsymbol{\rho})] \sim i.i.d. s(\mathbf{0}, \mathbf{I}_{N+1}, \boldsymbol{\eta})$, and both $M_{ll}(\boldsymbol{\eta}_0)$ and $M_{ss}(\boldsymbol{\eta}_0)$ correspond to this $(N+1)$ -dimensional distribution. However, $\hat{\gamma}_{JML}$ assumes omniscience on the part of the researcher, which is unrealistic.

The following proposition shows the consistency of the t -based estimators which make the erroneous assumption that $V[\mathbf{r}_t|r_{Mt}] = \tau \boldsymbol{\Upsilon}(\mathbf{v})$, where τ and $\boldsymbol{\Upsilon}(\mathbf{v})$ are defined in (12) and (13),

and provides expressions for the conditional variance of the score and expected Hessian matrix under such misspecification:

Proposition 7 *If $\Sigma^{-1/2}(\boldsymbol{\rho})[\mathbf{R}_t - \boldsymbol{\mu}(\boldsymbol{\rho})]|I_{t-1}; \boldsymbol{\varphi}_0 \sim i.i.d. s(\mathbf{0}, \mathbf{I}_{N+1}, \boldsymbol{\varrho}_0)$ with $\kappa_0 > 0$, where $\boldsymbol{\mu}(\boldsymbol{\rho})$ and $\Sigma(\boldsymbol{\rho})$ are defined in (17) and (18) respectively, $\boldsymbol{\rho}' = (\mathbf{a}', \mathbf{b}', \boldsymbol{\omega}', \mu_M, \sigma_M^2)$ and $\boldsymbol{\varphi} = (\boldsymbol{\rho}', \boldsymbol{\varrho}')'$, then:*

1. *The pseudo-true value of the unrestricted Student-t based PML estimator of $\boldsymbol{\phi} = (\boldsymbol{\gamma}, \mathbf{v}, \tau, \eta)'$, $\boldsymbol{\phi}_\infty$, is such that $\boldsymbol{\gamma}_\infty$ and $\mathbf{v}_\infty$ are equal to their corresponding true values $\boldsymbol{\gamma}_0$ and $\mathbf{v}_0$.*
2. *$\mathcal{O}(\boldsymbol{\phi}_\infty; \boldsymbol{\varphi}_0) = E[\mathcal{O}_t(\boldsymbol{\phi}_\infty; \boldsymbol{\varphi}_0)|\boldsymbol{\varphi}_0]$ and $\mathcal{H}(\boldsymbol{\phi}_\infty; \boldsymbol{\varphi}_0) = E[\mathcal{H}_t(\boldsymbol{\phi}_\infty; \boldsymbol{\varphi}_0)|\boldsymbol{\varphi}_0]$ will be block diagonal between $(\boldsymbol{\gamma}, \mathbf{v})$ and (τ, η) , where both*

$$\mathcal{O}_t(\boldsymbol{\phi}_\infty; \boldsymbol{\varphi}_0) = V[\mathbf{s}_t(\boldsymbol{\phi}_\infty)|r_{Mt}, I_{t-1}; \boldsymbol{\varphi}_0]$$

and

$$\mathcal{H}_t(\boldsymbol{\phi}_\infty; \boldsymbol{\varphi}_0) = -E[\mathbf{h}_t(\boldsymbol{\phi}_\infty)|r_{Mt}, I_{t-1}; \boldsymbol{\varphi}_0]$$

will share the structure of $\mathcal{I}_t(\boldsymbol{\phi}_\infty; \boldsymbol{\varphi}_0)$ in Proposition 1, with $\mathcal{O}_{\eta\eta}(\boldsymbol{\phi}; \boldsymbol{\varphi}) = V[s_{\eta t}(\boldsymbol{\phi})|\boldsymbol{\varphi}]$, $\mathcal{H}_{\eta\eta}(\boldsymbol{\phi}; \boldsymbol{\varphi}) = -E[h_{\eta\eta t}(\boldsymbol{\phi})|\boldsymbol{\varphi}]$,

$$\begin{aligned} M_{ll}^O(\boldsymbol{\phi}; \boldsymbol{\varphi}) &= E\{\delta^2[\varsigma_t(\boldsymbol{\rho}), \eta] \cdot [\varsigma_t(\boldsymbol{\rho})/N]|\boldsymbol{\varphi}\} \\ M_{ss}^O(\boldsymbol{\phi}; \boldsymbol{\varphi}) &= N(N+2)^{-1}[1 + V\{\delta[\varsigma_t(\boldsymbol{\rho}), \eta] \cdot [\varsigma_t(\boldsymbol{\rho})/N]|\boldsymbol{\varphi}\}], \\ M_{sr}^O(\boldsymbol{\phi}; \boldsymbol{\varphi}) &= E[\{\delta[\varsigma_t(\boldsymbol{\rho}), \eta] \cdot [\varsigma_t(\boldsymbol{\rho})/N] - 1\} s_{\eta t}(\boldsymbol{\phi})|\boldsymbol{\varphi}], \\ M_{ll}^H(\boldsymbol{\phi}; \boldsymbol{\varphi}) &= E\{2\partial\delta[\varsigma_t(\boldsymbol{\rho}), \eta]/\partial\varsigma \cdot [\varsigma_t(\boldsymbol{\rho})/N] + \delta[\varsigma_t(\boldsymbol{\theta}), \eta]|\boldsymbol{\varphi}\}, \\ M_{ss}^H(\boldsymbol{\phi}; \boldsymbol{\varphi}) &= E\{2\partial\delta[\varsigma_t(\boldsymbol{\rho}), \eta]/\partial\varsigma \cdot \varsigma_t^2(\boldsymbol{\rho})/[N(N+2)]|\boldsymbol{\varphi}\} + 1, \\ M_{sr}^H(\boldsymbol{\phi}; \boldsymbol{\varphi}) &= -E\{[\varsigma_t(\boldsymbol{\rho})/N] \cdot \partial\delta[\varsigma_t(\boldsymbol{\rho}), \eta]/\partial\eta|\boldsymbol{\varphi}\}. \end{aligned}$$

The asymptotic distribution of the unrestricted t -based PML estimator of $\boldsymbol{\gamma}$ follows immediately from Proposition 7:

Corollary 2 *If $\Sigma^{-1/2}(\boldsymbol{\rho})[\mathbf{R}_t - \boldsymbol{\mu}(\boldsymbol{\rho})]|I_{t-1}; \boldsymbol{\varphi}_0 \sim i.i.d. s(\mathbf{0}, \mathbf{I}_{N+1}, \boldsymbol{\varrho}_0)$ with $\kappa_0 > 0$, where $\boldsymbol{\mu}(\boldsymbol{\rho})$ and $\Sigma(\boldsymbol{\rho})$ are defined in (17) and (18) respectively, $\boldsymbol{\rho}' = (\mathbf{a}', \mathbf{b}', \boldsymbol{\omega}', \mu_M, \sigma_M^2)$ and $\boldsymbol{\varphi} = (\boldsymbol{\rho}', \boldsymbol{\varrho}')'$, then:*

$$\sqrt{T}(\hat{\boldsymbol{\gamma}}_{UML} - \boldsymbol{\gamma}_0) \rightarrow N\left[\mathbf{0}, \frac{M_{ll}^O(\boldsymbol{\phi}_\infty; \boldsymbol{\varphi}_0)}{[M_{ll}^H(\boldsymbol{\phi}_\infty; \boldsymbol{\varphi}_0)]^2} \cdot \frac{1}{\lambda_\infty} \mathcal{C}_{\boldsymbol{\gamma}\boldsymbol{\gamma}}(\boldsymbol{\varphi}_0)\right], \quad (23)$$

where $\lambda_\infty = \tau_0/\tau_\infty$.

As we mentioned after Corollary 1, it is very easy to obtain consistent estimators of the robust asymptotic covariance matrix in (23), and the same is true in the case of the restricted estimators.

Figure 5 presents the efficiency of the t -based PML estimators of γ in relation to the corresponding GMM estimator as a function of $\bar{\eta}$ when $\mathbf{R}_t$ is distributed as a multivariate t with 8 degrees of freedom ($\eta_0 = .125$) for three cross-sectional dimensions. In addition, the vertical straight lines in the top panel indicate the position of the pseudo-true values η_∞ when we also estimate this parameter, while the horizontal ones describe the efficiency of the joint ML estimator of γ in (22) relative to GMM estimator in (19).¹³ As in Figures 4a and 4b, the restricted t -based PML estimator of γ is more efficient than the GMM estimator for all values of $\bar{\eta}$, the more so the larger N is. Furthermore, the unrestricted t -based PML estimator that also estimates η gets close to achieving the full efficiency of the joint ML estimator, especially for large N . Finally, another noteworthy fact is the very small asymptotic bias of the t -based PML estimator of η .

In principle, Proposition 7, and in particular the block diagonal structure of $\mathcal{O}(\phi_\infty; \varphi_0)$ and $\mathcal{H}(\phi_\infty; \varphi_0)$ will continue to hold if we replace the unrestricted t -based ML estimator by any other estimator based on a specific *i.i.d.* elliptical distribution for $\mathbf{r}_t|r_{Mt}$. But since the HLV estimator is asymptotically equivalent to a parametric estimator that uses a flexible elliptical distribution as we increase the number of shape parameters, Proposition 7 suggests that the HLV estimator of γ will continue to be consistent. In fact, an argument analogous to the one made by Hodgson (2000) in a closely related univariate context would imply that the HLV estimator is as efficient as the parametric estimator that used the true *unconditional* distribution of the innovations $\varepsilon_t = \mathbf{r}_t - \mathbf{a}_0 - \mathbf{b}_0 r_{Mt}$. Nevertheless, inferences about $\mathbf{a}$ and $\mathbf{b}$ would have to be adjusted to reflect the contemporaneous conditional heteroskedasticity of ε_t , which is not straightforward.

3.3 Asymmetric distributions

To focus our discussion, we assume in this section that ε_t^* is distributed as an *i.i.d.* multivariate asymmetric t . Following Mencía and Sentana (2009), if we choose

$$\varepsilon_t^* = \beta [\xi_t^{-1} - c(\beta, \eta)] + \sqrt{\frac{\zeta_t}{\xi_t}} \Xi^{1/2} \mathbf{u}_t, \quad (24)$$

where $\mathbf{u}_t$ is uniformly distributed on the unit sphere in $\mathbb{R}^N$, ζ_t is a χ^2 random variable with N degrees of freedom, ξ_t is Gamma random variable with parameters $(2\eta)^{-1}$ and $\delta^2/2$ with

¹³These graphs are based on the expressions in Proposition 7, with the relevant expectations computed by Monte Carlo integration with 10^6 drawings.

$\delta = (1 - 2\eta)\eta^{-1}c(\boldsymbol{\beta}, \eta)$, $\boldsymbol{\beta}$ is a $N \times 1$ parameter vector, and $\boldsymbol{\Xi}$ is a $N \times N$ positive definite matrix given by

$$\boldsymbol{\Xi} = \frac{1}{c(\boldsymbol{\beta}, \eta)} \left[I_N + \frac{c(\boldsymbol{\beta}, \eta) - 1}{\boldsymbol{\beta}'\boldsymbol{\beta}} \boldsymbol{\beta}\boldsymbol{\beta}' \right],$$

with

$$c(\boldsymbol{\beta}, \eta) = \frac{-(1 - 4\eta) + \sqrt{(1 - 4\eta)^2 + 8\boldsymbol{\beta}'\boldsymbol{\beta}(1 - 4\eta)\eta}}{4\boldsymbol{\beta}'\boldsymbol{\beta}\eta},$$

then $E[\boldsymbol{\varepsilon}_t^*] = \mathbf{0}$ and $V[\boldsymbol{\varepsilon}_t^*] = \mathbf{I}_N$. In this sense, note that $\lim_{\boldsymbol{\beta}'\boldsymbol{\beta} \rightarrow 0} c(\boldsymbol{\beta}, \eta) = 1$, so that the above distribution collapses to the usual multivariate symmetric t when $\boldsymbol{\beta} = \mathbf{0}$. Therefore, we allow for asymmetries by introducing the vector of parameters $\boldsymbol{\beta}$.

To study the consistency of the symmetric t -based PML estimator when the data generating process (DGP) is asymmetric, it is once again convenient to look at its score. Specifically, given the definition of (24), we can write

$$\mathbf{s}_{\text{at}}(\gamma_0, \omega_0, \eta) = \boldsymbol{\Omega}_0^{-1/2} \frac{N\eta + 1}{1 - 2\eta + \eta(\zeta_t/\xi_t)} \left\{ \boldsymbol{\beta} [\xi_t^{-1} - c(\boldsymbol{\beta}, \eta)] + \sqrt{\frac{\zeta_t}{\xi_t}} \boldsymbol{\Xi}^{1/2} \mathbf{u}_t \right\}. \quad (25)$$

The expected value of $\boldsymbol{\varepsilon}_t^*$ in (24) is clearly zero by construction. Similarly, the expected value of (25) is also zero when $\boldsymbol{\beta}_0 = \mathbf{0}$ since $\mathbf{u}_t$ and (ζ_t/ξ_t) are independent. But when $\boldsymbol{\beta}_0 \neq \mathbf{0}$, the expected value of (25) will be generally different from zero because ξ_t^{-1} appears both in the numerator and denominator. Consequently, the mean parameters $\mathbf{a}$ will be inconsistently estimated. In contrast, $\mathbf{b}$ will be consistently estimated precisely because the estimator of $\mathbf{a}$ will fully mop up the bias in the mean. More formally, re-write model (1) as

$$\mathbf{r}_t = \boldsymbol{\Omega}^{1/2} \boldsymbol{\delta} + \mathbf{b}r_{Mt} + \boldsymbol{\varepsilon}_t,$$

where $\boldsymbol{\Omega}^{-1/2} \mathbf{a} = \boldsymbol{\delta}$. This homeomorphic reparametrisation satisfies the conditions of Proposition 17 in Fiorentini and Sentana (2007), which implies the consistency of $\mathbf{b}$. Unfortunately, mean-variance efficiency tests are based on $\mathbf{a}$, not $\mathbf{b}$.

For analogous reasons, the HLV estimator of $\mathbf{a}$ also becomes inconsistent under asymmetry. Intuitively, the problem is that it will not be true any more that the N -dimensional density of $\boldsymbol{\varepsilon}_t^*$ could be written as a function of $\varsigma_t = \boldsymbol{\varepsilon}_t^{*'} \boldsymbol{\varepsilon}_t^*$ alone. Therefore, a semiparametric estimator of $\mathbf{s}_{\gamma t}(\phi_0)$ that combines the elliptical symmetry assumption with a non-parametric specification for $\delta[\varsigma_t(\boldsymbol{\theta}), \boldsymbol{\eta}]$ will be contaminated by the skewness of the data.

In contrast, the GMM estimator always yields a consistent estimator of $\mathbf{a}$, on the basis of which we can develop a GMM-based Wald test with the correct asymptotic size since (9) remains valid under asymmetry.

4 Monte Carlo analysis

In this section we assess the finite sample size and power properties of the GMM, HLV and unrestricted t -based ML test statistics of the joint null hypothesis $H_0 : \mathbf{a} = \mathbf{0}$ for five different distributional assumptions for the innovations, namely Gaussian, Student- t with 4 degrees of freedom, Kotz with $\kappa = 1/8$, two-component scale mixture of normals with the same kurtosis, and asymmetric- t innovations.¹⁴ We also consider a t_4 distributional assumption for the returns, $\mathbf{R}_t$.¹⁵ In all cases, we carry out 10,000 replications with $T = 500$, $N = 5$, $\mathbf{\Omega} = 4\sigma_M^2 \times \mathbf{I}_5$, $\sqrt{12}\mu_M/\sigma_M = \frac{1}{2}$ and $\mathbf{b} = \mathbf{0}$ both under the null hypothesis, and under the alternative that $\mathbf{a} = 4\mu_M \times \ell_5$.¹⁶

We sample Gaussian and Student t random numbers using standard MATLAB routines. To sample the Kotz innovations, we exploit the fact that $\boldsymbol{\varepsilon}_t^* = \sqrt{\xi_t} \mathbf{u}_t$, where ξ_t is a univariate Gamma with mean N and variance $N[(N+2)\kappa+2]$. Similarly, we use (2) to sample the discrete mixture of normals. Finally, to draw asymmetric t innovations we first generate a univariate Gamma and N independent standard Gaussian variates, and then use the decomposition presented in (24).

As mentioned in section 2.2.2, the GMM estimators of $\boldsymbol{\gamma}$ coincide with the equation by equation coefficient estimates in the OLS regression of r_{it} on a constant and r_{Mt} . Similarly, a GMM estimator of $\mathbf{\Omega}$ can be easily obtained from the covariance matrix of the OLS regression residuals, as explained in footnote 7. We use the expressions in Proposition 3 to compute its covariance matrix under the maintained assumption of *i.i.d.* innovations. In contrast, we combine (20) with (21) to obtain heterokedasticity robust standard errors.

Following Fiorentini, Sentana and Calzolari (2003), we obtain a consistent estimator of the reciprocal degrees of freedom parameter η on the basis of the GMM estimators as

$$\hat{\eta}_{SMM} = \frac{\max[0, \bar{\kappa}_T(\hat{\boldsymbol{\theta}}_{GMM})]}{4 \max[0, \bar{\kappa}_T(\hat{\boldsymbol{\theta}}_{GMM})] + 2}, \quad (26)$$

where

$$\bar{\kappa}_T(\hat{\boldsymbol{\theta}}_{GMM}) = \frac{T^{-1} \sum_{t=1}^T \varsigma_t^2(\hat{\boldsymbol{\theta}}_{GMM})}{N(N+2)} - 1$$

is Mardia's (1970) sample coefficient of multivariate excess kurtosis of the estimated standardised residuals. Then, we use $\hat{\eta}_{SMM}$ as initial value to obtain the sequential ML estimator of η proposed

¹⁴In these cases, a sample of r_{Mt} is drawn from a Gaussian distribution for each replication.

¹⁵Amengual and Sentana (2008) present additional Monte Carlo results for other distributions, as well as for the individual t tests of $H_0 : a_i = 0$.

¹⁶The value of $\mathbf{b}$ does not affect the asymptotic distribution of the different estimators of $\mathbf{a}$ and the corresponding test procedures, while the value of σ_M simply scales up or down all the return series, and consequently $\mathbf{\Omega}$, μ_M and $\mathbf{a}$.

by Fiorentini and Sentana (2007), $\hat{\eta}_{SML}$ say, which maximises the t -based log-likelihood function with respect to η keeping $\boldsymbol{\theta}$ fixed at $\hat{\boldsymbol{\theta}}_{GMM}$.

Having obtained $\hat{\boldsymbol{\theta}}_{GMM}$ and $\hat{\eta}_{SML}$, we compute a one-step ML estimator of $\boldsymbol{\theta}$ by means of the BHHH correction

$$\left[\sum_{t=1}^T \mathbf{s}_{\gamma t}(\boldsymbol{\theta}) \mathbf{s}_{\gamma t}'(\boldsymbol{\theta}) \right]^{-1} \sum_{t=1}^T \mathbf{s}_{\gamma t}(\boldsymbol{\theta}), \quad (27)$$

with the analytical expressions for the t -score derived in section 2.2.1.¹⁷ Next, we carry out a few EM iterations over $\boldsymbol{\theta}$ using this one-step ML estimator as initial value (see Supplemental Appendix C), and finally switch to a quasi-Newton procedure until convergence. The (non-robust) asymptotic covariance matrix is computed using the expressions in Proposition 1, while for the robust standard errors we use the expressions in Corollaries 1 and 2.

As for the HLV estimator and its asymptotic covariance matrix, we follow the computational approach described in Appendix B1 of Fiorentini and Sentana (2007). Specifically, for the purposes of computing reliable standard errors they recommend a simple average of the sample analogue of the outer product of the score expression for $M_u(\boldsymbol{\eta})$ in Proposition 1, and an alternative estimator based on the following expression:

$$M_u(\boldsymbol{\eta}) = cov \left\{ \delta[\varsigma_t(\boldsymbol{\theta}), \boldsymbol{\eta}], \delta[\varsigma_t(\boldsymbol{\theta}), \boldsymbol{\eta}] \frac{\varsigma_t}{N} \middle| \boldsymbol{\eta} \right\} + (N-2)E[\varsigma^{-1}(\boldsymbol{\theta})|\boldsymbol{\eta}],$$

which is valid as long as $E[\varsigma^{-1}|\boldsymbol{\eta}_0]$ is bounded, which in the Gaussian case, for instance, requires $N \geq 3$.

4.1 Sampling distribution of the different estimators

Although we are mostly interested in the test statistics, it is convenient to study first the finite sample distributions of the estimators of $\mathbf{a}$, which are not affected by the estimation of their asymptotic covariance matrices.

In this sense, Figure 6 presents box-plots of the unrestricted t -based PML, HLV and GMM estimators for eight different DGP's that we have considered. As usual, the central boxes describe the first and third quartiles of the sampling distributions, as well as their median. The maximum length of the whiskers is one interquartile range.

By and large, the behaviour of the different estimators is in accordance with what the asymptotic results would suggest. The only ‘‘surprises’’ are the fact that the dispersion of the

¹⁷This one-step ML estimator is asymptotically equivalent to $\hat{\gamma}_{ML}$. An alternative asymptotically equivalent estimator of $\hat{\gamma}_{ML}$ will update the whole of $\hat{\boldsymbol{\theta}}_{GMM}$ by means of a simple BHHH correction based on $\mathbf{s}_{\boldsymbol{\theta}t}$.

distribution of the HLV estimator is systematically larger than the distribution of the ML estimator under correct specification of the latter, and that this result continues to hold even when the innovations follow a discrete mixture of normals. The other interesting results occur when the joint distribution of $\mathbf{r}_t$ and r_{Mt} is elliptical, so that the conditional mean of $\mathbf{r}_t$ given r_{Mt} continues to be linear in r_{Mt} but the conditional variance is no longer constant. In this case not only the HLV and ML estimators of $\mathbf{a}$ remain consistent despite this misspecification, as we discussed in section 3.2, but they are also more efficient than the GMM estimator.

4.2 Sampling distribution of the associated test statistics

The first question that we need to address is whether the asymptotic distribution under the null attributed to the joint Wald test statistics introduced in section 2.3 is reliable in finite samples. To do so, we employ the p -value discrepancy plots proposed by Davidson and MacKinnon (1998). Let w_j denote the simulated value of a given test statistic, and let p_j be the asymptotic p -value of w_j , that is the probability of observing a value of the test statistic at least as large as w_j according to its asymptotic distribution under the null. Let also $\hat{F}(x)$ for $x \in (0, 1)$ be the empirical distribution function of p_j i.e. the sample proportion of p_j 's which are not greater than x . A p -value discrepancy plot is a plot of $[\hat{F}(x) - x]$ against x , i.e. a plot of the difference between actual and nominal size for a range of nominal sizes. If the candidate distribution for w_j is correct, then the p -value discrepancy should be close to zero.

The left panels of Figures 7a-7c show p -value discrepancy plots of the joint tests (“Wald statistics”) of $H_0 : \mathbf{a} = \mathbf{0}$ for the six DGPs that we have considered. The most striking fact that we find is that the HLV-based joint and individual tests have systematically the largest size distortions irrespective of whether the assumptions that justify them are correct. In contrast, the GMM tests that use expression (9) to compute the asymptotic weighting matrix have finite sample sizes that are close to their asymptotically equivalent in all cases, including when the correct expression should be (19). As for the tests that use the unrestricted t -based PML estimator, there is also little to choose between the robust and non-robust versions, which are both well behaved even when the conditional distribution is heteroskedastic. The only exception seems to be the discrete mixture of normals example (Figure 7b), in which case the non-robust test is surprisingly better behaved than the robust one. As expected, though, when the distribution of the innovations is asymmetric (Figure 7c), the HLV and ML tests present considerable size distortions.

We can complement our finite sample analysis with *size-power curves*, which is another

graphical method proposed by Davidson and MacKinnon (1998) to display the simulation evidence on the power of the different tests. We can define $\hat{F}^*(x)$ for $x \in (0, 1)$ as the empirical distribution function of the asymptotic p -values under the null when the data are generated under the alternative. A size-power curve is a plot of test power versus actual test size for a range of test sizes.

The right panels of Figures 7a-7c show size-power curves for the same six DGPs. Not surprisingly, the size-adjusted powers of the robust tests are very close to the corresponding non-robust tests in all cases. Contrary to the asymptotic results, though, GMM tests seem to have more power than the others under Gaussian innovations. In all other cases, in contrast, the HLV-based tests are more powerful than the GMM ones, but less so than the ones that use the unrestricted t -based PML estimator –except in the discrete mixture of normals example. In addition, the differences in power between HLV and t -based PML tests are very small in the case of Kotz and discrete mixture of normals innovations, despite the fact that the t -based estimator is suboptimal.

5 Empirical application

In this section we use the alternative estimators previously discussed to test the mean-variance efficiency of the US aggregate stock market portfolio using monthly data over the period July 1962 to June 2007 (540 observations). As for $\mathbf{r}_t$, we consider two different sets of $N = 5$ portfolios from Ken French’s Data Library: one grouped by industry, and another one sorted by their book-to-market ratio. Specifically, each NYSE, AMEX, and NASDAQ stock is assigned to an industry portfolio at the end of June of year t based on its four-digit SIC code at the time.¹⁸ Similarly, quintile portfolios are formed on BE/ME at the end of each June using NYSE breakpoints. The BE used in June of year t is the book equity for the last fiscal year end in $t - 1$, while ME is price times shares outstanding at the end of December of $t - 1$. The excess return on the market portfolio corresponds to the value weighted return measure on all NYSE, AMEX and NASDAQ stocks in CRSP, while the safe asset is the 1-month TBill return from Ibbotson and Associates (see <http://mba.tuck.dartmouth.edu/pages/faculty/ken.french/data_library.html> for further details).

The most obvious characteristic of these portfolios for our purposes is their leptokurtosis. The LM test of normality against the alternative of multivariate Student t proposed by Fioren-

¹⁸Industry definitions: Cnsmr: Consumer Durables, NonDurables, Wholesale, Retail, and Some Services (Laundries, Repair Shops). Manuf: Manufacturing, Energy, and Utilities. HiTec: Business Equipment, Telephone and Television Transmission. Hlth: Healthcare, Medical Equipment, and Drugs. Other: Other – Mines, Constr, BldMt, Trans, Hotels, Bus Serv, Entertainment, Finance.

tini, Sentana and Calzolari (2003) yields a value of 3173.71 for the industry portfolios residuals from (1), and 1997.83 for the book to market ones. This confirms our empirical motivation for estimation and testing procedures that exploit such a prevalent feature of the data.

Table 1a presents the parameter estimates and (asymptotic) robust standard errors for the GMM, HLV and t -based ML estimators of the intercept of model (1), while Table 1b reports the corresponding joint tests of $H_0 : \mathbf{a} = \mathbf{0}$.¹⁹ Our results for industry portfolios indicate that the Student t -based joint test clearly rejects the efficiency of the market portfolio, even though the univariate t -tests would not, which confirms the recommendation of Gibbons, Ross and Shanken (1989) to increase power by taking into account the covariance structure of the residuals in many empirically relevant situations. Our results also show that the joint GMM test is borderline, while the HLV-based test fails to reject, which is in line with the results reported by Vorkink (2003).

Given the expressions for the test statistics in sections 2 and 3, the contradicting conclusions obtained with the different tests must be due to three causes. First, the point estimates of $\mathbf{a}$ are somewhat different, the HLV being on average closer to 0 in magnitude. Second, the point estimates of the idiosyncratic covariance matrix $\mathbf{\Omega}$ also differ, although even less so. More importantly, the scalar factors that multiply $\mathbf{\Omega}^{-1}$ are noticeably different too. In particular, they are 1.87 and 1.98 for the robust and non-robust versions of the ML tests, but only 1.43 for the HLV test. Both our Monte Carlo results and the results reported in Fiorentini and Sentana (2007) indicate the unreliable nature of the non-parametric estimates of M_{ll} in finite samples.

In contrast, all three tests reject the mean-variance efficiency of the market portfolio relative to the book-to-market sorted portfolios of Fama and French (1993). Still, we also find important differences in the estimates of the scalar factors mentioned in the previous paragraph.

As we saw in section 3.3, though, both parametric and semiparametric elliptically-based procedures are sensitive to the assumption of elliptical symmetry. For that reason, we follow Mencía and Sentana (2009), and test the null hypothesis of multivariate Student t innovations against the multivariate asymmetric t distribution in (24).²⁰ The values of the test statistic and

¹⁹A full set of results is available on request.

²⁰Mencía and Sentana (2009) also propose LM tests for kurtosis against symmetric generalised hyperbolic distributions, as well as joint tests for asymmetry and kurtosis. Their kurtosis statistic tests that $(1 + \gamma)^{-1} \equiv \psi = 1$ under the maintained hypothesis of $\boldsymbol{\beta} = \mathbf{0}$, where γ is the second tail shape parameter of a generalised inverse Gaussian (GIG) distribution and $\boldsymbol{\beta}$ is the $N \times 1$ vector of coefficients that appears in (24). In effect, this amounts to testing that the tail behaviour of the multivariate t distribution adequately reflects the kurtosis in the data. In turn, the asymmetry statistic tests that $\boldsymbol{\beta} = \mathbf{0}$ under the maintained assumption that $\psi = 1$. Unfortunately, the kurtosis-based test requires finite fourth moments under the null hypothesis, while the ML estimates of η , which is the reciprocal of the degrees of freedom of the multivariate t distribution, are above .25. in the two data sets that we consider

associated p -values are 10.003 and 0.075, respectively, for the industry sorted portfolios; and 9.880 and 0.079 for the book-to-market sorted portfolios. Therefore, we cannot reject the null hypothesis that the distribution of $\mathbf{r}_t$ conditional on r_{Mt} is multivariate Student t at conventional levels.

Finally, we perform a simple conditional homoskedasticity test by regressing the squared OLS residuals from the regression of r_{it} on a constant and r_{Mt} , $\hat{\varepsilon}_{it}^2$ say, on a constant, the market excess return r_{Mt} and its squared r_{Mt}^2 for $i = 1, \dots, N$ (see White (1980)). The results in Table 2 suggest that the distribution of the innovations conditional on r_{Mt} is rather heteroskedastic, as we reject the null hypothesis at 5% significance level in almost all cases. This result confirms the need to use the robust estimates of the asymptotic covariance matrix of the t -based ML procedures in Proposition 7, as well as the problems that the HLV standard errors face, since they are based on Proposition 1 instead.

6 Summary and directions for further research

In this paper we study the efficiency-consistency trade-offs of three approaches to test the mean-variance efficiency of a candidate portfolio with returns r_{Mt} in excess of the riskless asset with respect to a set of N assets with excess returns $\mathbf{r}_t$. In particular, we consider tests based on the GMM approach advocated by MacKinlay and Richardson (1991), the elliptically symmetric semiparametric methods proposed by HLV, and an unrestricted parametric procedure that makes the assumption that, conditional on the reference portfolio, the excess returns of the original assets are independent and identically distributed as a multivariate t . We would like to emphasise, though, that most of our results apply not only to the multivariate t , but also to any other elliptically-based likelihood estimator. The main advantage of the Student t for our purposes is that we can make explicit its limiting relationship to the Gaussian distribution.

We also apply these different procedures to test the mean-variance efficiency of the US aggregate stock market portfolio using monthly data over the period July 1962 to June 2007. The results that we obtain for industry portfolios indicate that the Student t -based test clearly rejects the efficiency of the market portfolio, while the GMM test is borderline, and the HLV based test fails to reject. In contrast, all three tests reject the mean-variance efficiency of the market portfolio relative to the book-to-market sorted portfolios of Fama and French (1993). In view of the trade-offs between efficiency and consistency that we characterise in our theoretical analysis, the results of the misspecification tests that we compute suggest that it is probably worth using

the multivariate t distribution for the purposes of testing mean-variance efficiency, as long as empirical researchers bear in mind that such a distributional assumption may be wrong, and robustify their inferences accordingly.

Morales (2009) extends our results in two interesting directions. First, she considers a situation in which one wants to test the mean-variance efficiency of several reference portfolios simultaneously. She also allows both $\mathbf{a}$ and $\mathbf{b}$ to linearly depend on a vector of predictor variables known at time $t - 1$, $\mathbf{x}_{t-1}$ say, and in this way test for conditional mean variance efficiency, as discussed in Beaulieu, Dufour and Khalaf (2007a) and others.

The fact that the number of assets that we consider in our Monte Carlo experiments and in our empirical application is fairly small means that they are unlikely to be affected by the criticism raised by Gibbons, Ross and Shanken (1989) in relation to the sensitivity of the asymptotic (in T) distribution of mean-variance efficiency tests to the cross-sectional dimension N . However, situations in which N/T cannot be regarded as negligible would require different asymptotic approximations to the one used in this paper.

We could increase the efficiency of the GMM estimator of $\mathbf{a}$ discussed in section 2.2.2 and the power of the associated test procedures by including additional moment restrictions that exploit the elliptical distribution of the innovations. For instance, we could follow Renault and Sentana (2003), and consider moment conditions of the form:

$$E \{ \boldsymbol{\varepsilon}_t(\boldsymbol{\gamma}) \otimes \text{vech}[\boldsymbol{\varepsilon}_t(\boldsymbol{\gamma})\boldsymbol{\varepsilon}_t'(\boldsymbol{\gamma})] \} = \mathbf{0}. \quad (28)$$

GMM estimators that combine (8) with this moment condition will typically have a lower asymptotic variance than $\hat{\boldsymbol{\gamma}}_{GMM}$. In fact, we could regard the HLV estimator as a GMM estimator that optimally exploits the ellipticity of $\boldsymbol{\varepsilon}_t^*$, which means that in principle such augmented GMM procedures could achieve the elliptically symmetric semiparametric efficiency bound $\mathcal{I}_{\boldsymbol{\gamma}\boldsymbol{\gamma}}(\boldsymbol{\phi}_0)$. Like the HLV estimator, though, such GMM estimators will also become inconsistent if (28) does not hold, but their main advantage is that GMM integrates estimation and testing.

Importantly, we have not looked at mean-variance efficiency tests when a riskless asset is not available (as in e.g. Shanken (1986), Zhou (1991), and more recently Beaulieu, Dufour and Khalaf (2007b)). In those circumstances, it is important to distinguish between mean-variance efficiency tests on the one hand, and spanning tests on the other, in which the null hypothesis involves restrictions on both intercepts and slopes of the multivariate regression model (1) (see Huberman and Kandel (1987), and De Roan and Nijman (2001) for a recent survey, as well as Peñaranda and Sentana (2008) for a comparison of alternative GMM procedures). Another

example in which the null hypothesis involves restrictions on both intercepts and slopes of a multivariate regression would be tests of the uncovered interest parity hypothesis (see Hodgson, Linton and Vorkink (2004)).

Finally, to test the validity of the specific distributional assumption for ε_t^* made for the purposes of obtaining $\hat{\gamma}_{ML}$ in our empirical application, we have used the asymmetric LM test of Mencía and Sentana (2009), who use the generalised hyperbolic family as the nesting distribution. And although there are many other tests of ellipticity in the statistical literature (see e.g. Beran (1979)), for the purposes of testing mean-variance efficiency we could also use the Hausman specification tests proposed by Fiorentini and Sentana (2007), which compare the consistent but inefficient estimator $\hat{\mathbf{a}}_{GMM}$ with the efficient but potentially inconsistent estimators $\hat{\mathbf{a}}_{HLV}$ and $\hat{\mathbf{a}}_{UML}$. An alternative procedure would be a moment test that checks whether the information matrix equality for $\mathbf{M}_H$ implicit in Proposition 1 holds, as suggested by Fiorentini and Sentana (2007). All these issues constitute interesting avenues for further research.

Appendix

A Proofs

Proposition 1:

The result follows directly from Proposition 1 in Fiorentini and Sentana (2007) by using the fact that in the case of model (1)

$$\mathbf{Z}'_{lt}(\boldsymbol{\theta}) = \boldsymbol{\Omega}^{-1/2} \frac{\partial(\mathbf{a} + \mathbf{b}r_{Mt})}{\partial \boldsymbol{\theta}'} = \boldsymbol{\Omega}^{-1/2} \begin{bmatrix} (1, r_{Mt}) \otimes \mathbf{I}_N & \mathbf{0} \end{bmatrix} \quad (\text{A1})$$

and

$$\mathbf{Z}'_{st}(\boldsymbol{\theta}) = \frac{1}{2}(\boldsymbol{\Omega}^{-1/2} \otimes \boldsymbol{\Omega}^{-1/2}) \frac{\partial \text{vec}(\boldsymbol{\Omega})}{\partial \boldsymbol{\theta}'} = \frac{1}{2}(\boldsymbol{\Omega}^{-1/2} \otimes \boldsymbol{\Omega}^{-1/2}) \begin{pmatrix} \mathbf{0} & \mathbf{D}_N \end{pmatrix}. \quad (\text{A2})$$

Proposition 2

The asymptotic normality of the ML estimator of $\mathbf{a}$ follows from standard arguments by combining a central limit theorem for the score with a uniform law of large numbers for the Hessian matrix under the explicit assumptions that $\boldsymbol{\varepsilon}_t^*$ is *i.i.d.* and both r_{Mt} and r_{Mt}^2 are strictly stationary process with absolutely summable autocovariances. The expression for the asymptotic covariance matrix is a direct product of the partitioned inverse formula.

Proposition 3

The expressions for the matrices $\mathcal{A}_{\gamma\gamma t}(\boldsymbol{\phi})$, $\mathcal{B}_{\gamma\gamma t}(\boldsymbol{\phi})$ and $\mathcal{C}_{\gamma\gamma t}(\boldsymbol{\phi})$ follow directly from replacing (A1) and (A2) in Proposition 2 in Fiorentini and Sentana (2007). The asymptotic normality of the GMM estimator of $\boldsymbol{\gamma}$ can be obtained using the arguments in the proof of Corollary 1.

Proposition 4

The first part of the Proposition follows directly from the first part of Proposition 15 in Fiorentini and Sentana (2007). The second part of the distribution also follows directly from the second and third parts of the same proposition because mesokurtic elliptical distributions satisfy their condition (39), as Fiorentini and Sentana (2007) explain in their proof.

Proposition 5

The first part of the Proposition follows directly from the first part of Proposition 16 in Fiorentini and Sentana (2007). Specifically, let us initially keep η fixed to some positive value.

Since $\boldsymbol{\varepsilon}_t$ is elliptical, it can be written as $\boldsymbol{\varepsilon}_t^* = \sqrt{\varsigma_t} \mathbf{u}_t$ where $\mathbf{u}_t$ is uniformly distributed on the unit sphere surface in $\mathbb{R}^N$ and ς_t is a non-negative random variable independent of $\mathbf{u}_t$. Since

$$\varsigma_t(\boldsymbol{\gamma}_0, \mathbf{v}_0, \tau) = \frac{1}{\tau} \boldsymbol{\varepsilon}_t'(\boldsymbol{\gamma}_0) \boldsymbol{\Upsilon}^{-1}(\mathbf{v}_0) \boldsymbol{\varepsilon}_t(\boldsymbol{\gamma}_0) = \frac{\tau_0}{\tau} \varsigma_t$$

where $\varsigma_t = \varsigma_t(\boldsymbol{\gamma}_0, \mathbf{v}_0, \tau_0)$, we can write the blocks of the score corresponding to $\boldsymbol{\gamma}$, $\mathbf{v}$ and τ as

$$\mathbf{s}_{\boldsymbol{\gamma}t}(\boldsymbol{\gamma}_0, \mathbf{v}_0, \tau, \eta) = \begin{pmatrix} 1 \\ r_{Mt} \end{pmatrix} \otimes \frac{1}{\sqrt{\tau}} \boldsymbol{\Upsilon}^{-1/2}(\mathbf{v}_0) \delta[(\tau_0/\tau)\varsigma_t, \eta] \sqrt{(\tau_0/\tau)} \sqrt{\varsigma_t} \mathbf{u}_t \quad (\text{A3})$$

$$\begin{aligned} \mathbf{s}_{\mathbf{v}t}(\boldsymbol{\gamma}_0, \mathbf{v}_0, \tau, \eta) &= \frac{1}{2} \frac{\partial \text{vec}'[\boldsymbol{\Upsilon}(\mathbf{v}_0)]}{\partial \mathbf{v}} \left[\boldsymbol{\Upsilon}(\mathbf{v}_0)^{-1/2} \otimes \boldsymbol{\Upsilon}(\mathbf{v}_0)^{-1/2} \right] \\ &\quad \times \text{vec} \left\{ \delta[(\tau_0/\tau)\varsigma_t, \eta] \frac{\tau_0}{\tau} \varsigma_t \mathbf{u}_t \mathbf{u}_t' - \mathbf{I}_N \right\} \end{aligned} \quad (\text{A4})$$

and

$$s_{\tau t}(\boldsymbol{\gamma}_0, \mathbf{v}_0, \tau, \eta) = \frac{1}{2\tau} \text{vec}'(\mathbf{I}_N) \text{vec} \left\{ \delta[(\tau_0/\tau)\varsigma_t, \eta] \frac{\tau_0}{\tau} \varsigma_t \mathbf{u}_t \mathbf{u}_t' - \mathbf{I}_N \right\}. \quad (\text{A5})$$

Then, it follows that $E[\mathbf{s}_{\boldsymbol{\gamma}t}(\boldsymbol{\gamma}_0, \mathbf{v}_0, \tau, \eta) | r_{Mt}, I_{t-1}; \boldsymbol{\varphi}_0] = \mathbf{0}$ regardless of τ and η because of the serial and mutual independence of ς_t and $\mathbf{u}_t$, and the fact that $E(\mathbf{u}_t) = \mathbf{0}$.

If we define $\tau_\infty(\eta)$ as the value that solves the implicit equation

$$E \left[\frac{N\eta + 1}{1 - 2\eta + \eta(\tau_0/\tau)\varsigma_t} \frac{\tau_0}{\tau} \frac{\varsigma_t}{N} - 1 \middle| \boldsymbol{\varphi}_0 \right] = 0 \quad (\text{A6})$$

then it is straightforward to show that

$$\begin{aligned} E[\mathbf{s}_{\mathbf{v}t}(\boldsymbol{\gamma}_0, \mathbf{v}_0, \tau_\infty(\eta), \eta) | r_{Mt}, I_{t-1}; \boldsymbol{\varphi}_0] &= \mathbf{0} \\ E[s_{\tau t}(\boldsymbol{\gamma}_0, \mathbf{v}_0, \tau_\infty(\eta), \eta) | r_{Mt}, I_{t-1}; \boldsymbol{\varphi}_0] &= 0. \end{aligned}$$

by using the fact that $E(\mathbf{u}_t \mathbf{u}_t') = N^{-1} \mathbf{I}_N$.

If we choose η_∞ as the solution to the implicit equation

$$E[s_{\eta t}(\boldsymbol{\gamma}_0, \mathbf{v}_0, \tau_\infty(\eta), \eta) | \boldsymbol{\varphi}_0] = 0, \quad (\text{A7})$$

then it is clear that $\mathbf{v}_0, \tau_\infty(\eta_\infty)$ and η_∞ will be the pseudo true values of the parameters.

To obtain the variance of the t -score and the expected value of the t -hessian under misspecification it is convenient to rewrite the score as

$$\mathbf{s}_{\boldsymbol{\vartheta}t}(\boldsymbol{\vartheta}, \eta) = \begin{bmatrix} \mathbf{Z}_{\boldsymbol{\gamma}t}(\boldsymbol{\vartheta}) & \mathbf{Z}_{\mathbf{v}t}(\boldsymbol{\vartheta}) \\ \mathbf{0} & \mathbf{Z}_{\tau t}(\boldsymbol{\vartheta}) \end{bmatrix} \times [\mathbf{e}_{\eta t}(\boldsymbol{\vartheta}, \eta), \mathbf{e}_{st}(\boldsymbol{\vartheta}, \eta)]$$

where

$$\begin{aligned} \mathbf{Z}_{\boldsymbol{\gamma}t}(\boldsymbol{\vartheta}) &= \begin{pmatrix} 1 \\ r_{Mt} \end{pmatrix} \otimes \frac{1}{\sqrt{\tau}} \boldsymbol{\Upsilon}^{-1/2}(\mathbf{v}) \\ \mathbf{Z}_{\mathbf{v}t}(\boldsymbol{\vartheta}) &= \frac{1}{2} \frac{\partial \text{vec}'[\boldsymbol{\Upsilon}(\mathbf{v})]}{\partial \mathbf{v}} \left[\boldsymbol{\Upsilon}(\mathbf{v})^{-1/2} \otimes \boldsymbol{\Upsilon}(\mathbf{v})^{-1/2} \right] \\ \mathbf{Z}_{\tau t}(\boldsymbol{\vartheta}) &= \frac{1}{2} \frac{1}{\tau} \text{vec}'(\mathbf{I}_N) \end{aligned}$$

and

$$\begin{aligned}\mathbf{e}_{lt}(\boldsymbol{\vartheta}, \eta) &= \delta[\varsigma_t(\boldsymbol{\vartheta}), \eta] \sqrt{\varsigma_t(\boldsymbol{\vartheta})} \mathbf{u}_t \\ \mathbf{e}_{st}(\boldsymbol{\vartheta}, \eta) &= \text{vec}\{\delta[\varsigma_t(\boldsymbol{\vartheta}), \eta] \varsigma_t(\boldsymbol{\vartheta}) \mathbf{u}_t \mathbf{u}_t' - \mathbf{I}_N\}.\end{aligned}$$

Then, we can follow exactly the same steps as in the proof of Proposition 1 in Fiorentini and Sentana (2007) by exploiting that (A6) and (A7) hold at the pseudo-true parameter values ϕ_∞ .

Finally, we can show that $\mathcal{O}(\phi_\infty; \boldsymbol{\varphi}_0)$ and $\mathcal{H}(\phi_\infty; \boldsymbol{\varphi}_0)$ will be block diagonal between $(\gamma, \mathbf{v})$ and (τ, η) if $E[\partial d_t(\boldsymbol{\vartheta})/\partial \mathbf{v} | \boldsymbol{\varphi}_0] = \mathbf{0}$. But this trivially holds in our parametrization because $|\Upsilon(\mathbf{v})| = 1$ for all $\mathbf{v}$.

Proposition 6

If we use the subscript J to denote the joint log-likelihood function of $\mathbf{R}_t$, expression (2) in Fiorentini and Sentana (2007) implies that

$$\begin{aligned}\mathbf{s}_{J\rho t}(\boldsymbol{\rho}, \eta) &= \frac{\partial \boldsymbol{\mu}_t'(\boldsymbol{\rho})}{\partial \boldsymbol{\rho}} \boldsymbol{\Sigma}_t^{-1}(\boldsymbol{\rho}) \delta_{N+1}[\boldsymbol{\epsilon}_t^{*'}(\boldsymbol{\rho}) \boldsymbol{\epsilon}_t^*(\boldsymbol{\rho}), \eta] \cdot \boldsymbol{\epsilon}_t(\boldsymbol{\rho}) \\ &\quad + \frac{1}{2} \frac{\partial \text{vec}'[\boldsymbol{\Sigma}_t(\boldsymbol{\rho})]}{\partial \boldsymbol{\rho}} [\boldsymbol{\Sigma}_t^{-1}(\boldsymbol{\rho}) \otimes \boldsymbol{\Sigma}_t^{-1}(\boldsymbol{\rho})] \\ &\quad \times \text{vec}\{\delta_{N+1}[\boldsymbol{\epsilon}_t^{*'}(\boldsymbol{\rho}) \boldsymbol{\epsilon}_t^*(\boldsymbol{\rho}), \eta] \cdot \boldsymbol{\epsilon}_t(\boldsymbol{\rho}) \boldsymbol{\epsilon}_t'(\boldsymbol{\rho}) - \boldsymbol{\Sigma}_t(\boldsymbol{\rho})\}.\end{aligned}$$

In our case,

$$\frac{\partial \boldsymbol{\mu}_t(\boldsymbol{\rho})}{\partial \boldsymbol{\rho}'} = \frac{\partial}{\partial \boldsymbol{\rho}'} \begin{pmatrix} \mu_M \\ \mathbf{a} + \mathbf{b} \mu_M \end{pmatrix} = \begin{pmatrix} 0 & \mathbf{0}' & \mathbf{0}' & 1 & 0 \\ \mathbf{I}_N & \mu_M \mathbf{I}_N & \mathbf{0} & \mathbf{0} & \mathbf{0} \end{pmatrix}.$$

As for

$$\frac{\partial \text{vec}[\boldsymbol{\Sigma}_t(\boldsymbol{\rho})]}{\partial \boldsymbol{\rho}'} = \frac{\partial \text{vec}}{\partial \boldsymbol{\rho}'} \begin{pmatrix} \sigma_M^2 & \sigma_M^2 \mathbf{b}' \\ \sigma_M^2 \mathbf{b} & \sigma_M^2 \mathbf{b} \mathbf{b}' + \boldsymbol{\Omega} \end{pmatrix},$$

it is more convenient to obtain its elements by blocks, so that

$$\frac{\partial}{\partial \boldsymbol{\rho}'} \begin{pmatrix} \sigma_M^2 \\ \sigma_M^2 \mathbf{b} \end{pmatrix} = \begin{pmatrix} 0 & \mathbf{0}' & \mathbf{0}' & 0 & 1 \\ \mathbf{0} & \sigma_M^2 \mathbf{I}_N & \mathbf{0} & \mathbf{0} & \mathbf{b} \end{pmatrix}$$

and

$$\frac{\partial \text{vec}(\sigma_M^2 \mathbf{b} \mathbf{b}' + \boldsymbol{\Omega})}{\partial \boldsymbol{\rho}'} = \begin{bmatrix} \mathbf{0} & (\mathbf{I}_{N^2} + \mathbf{K}_{NN})(\sigma_M^2 \mathbf{b} \otimes \mathbf{I}_N) & \mathbf{D}_N & \mathbf{0} & (\mathbf{b} \otimes \mathbf{b}) \end{bmatrix},$$

and then re-arrange them appropriately.

It is also easy to see that

$$\boldsymbol{\Sigma}^{-1}(\boldsymbol{\rho}) = \begin{pmatrix} \sigma_M^{-2} + \mathbf{b}' \boldsymbol{\Omega}^{-1} \mathbf{b} & -\mathbf{b}' \boldsymbol{\Omega}^{-1} \\ -\boldsymbol{\Omega}^{-1} \mathbf{b} & \boldsymbol{\Omega}^{-1} \end{pmatrix},$$

by exploiting the Cholesky decomposition of $\boldsymbol{\Sigma}(\boldsymbol{\rho})$ in (A8).

We can also tediously prove that

$$[\Sigma^{-1}(\boldsymbol{\rho}) \otimes \Sigma^{-1}(\boldsymbol{\rho})] \frac{\partial \text{vec}(\Sigma(\boldsymbol{\rho}))}{\partial \sigma_M^2} = \begin{bmatrix} 1/\sigma_M^4 \\ \mathbf{0} \end{bmatrix},$$

and

$$\frac{\partial \text{vec}'[\Sigma(\boldsymbol{\rho})]}{\partial \mathbf{b}} [\Sigma^{-1}(\boldsymbol{\rho}) \otimes \Sigma^{-1}(\boldsymbol{\rho})] = \begin{pmatrix} -2\boldsymbol{\Omega}^{-1}\mathbf{b} & \boldsymbol{\Omega}^{-1} \otimes \mathbf{e}_{1,N+1} \end{pmatrix}$$

where $\mathbf{e}_{1,N+1}$ is a vector whose first element is one and has zeros in its remaining N positions.

On this basis, we can write

$$\begin{aligned} \mathbf{s}_{J\gamma t}(\boldsymbol{\rho}, \boldsymbol{\eta}) &= \left[\begin{pmatrix} 1 \\ \mu_M \end{pmatrix} \otimes \boldsymbol{\Omega}^{-1} \right] \delta_{N+1}[\boldsymbol{\epsilon}_t^{*'}(\boldsymbol{\rho})\boldsymbol{\epsilon}_t^*(\boldsymbol{\rho}), \boldsymbol{\eta}] \begin{bmatrix} -\mathbf{b} & \mathbf{I}_N \end{bmatrix} \boldsymbol{\epsilon}_t(\boldsymbol{\rho}) \\ &+ \begin{pmatrix} \mathbf{0} & \mathbf{0} \\ -2\boldsymbol{\Omega}^{-1}\mathbf{b} & \boldsymbol{\Omega}^{-1} \otimes \mathbf{e}_{1,N+1} \end{pmatrix} \text{vec} \{ \delta_{N+1}[\boldsymbol{\epsilon}_t^{*'}(\boldsymbol{\rho})\boldsymbol{\epsilon}_t^*(\boldsymbol{\rho}), \boldsymbol{\eta}] \boldsymbol{\epsilon}_t(\boldsymbol{\rho})\boldsymbol{\epsilon}_t'(\boldsymbol{\rho}) - \Sigma(\boldsymbol{\rho}) \}. \end{aligned}$$

and

$$\mathbf{s}_{J\omega t}(\boldsymbol{\rho}, \boldsymbol{\eta}) = \frac{1}{2} \mathbf{D}'_N [\boldsymbol{\Omega}^{-1} \otimes \boldsymbol{\Omega}^{-1}] \text{vec} \{ \delta_{N+1}[\boldsymbol{\epsilon}_t^{*'}(\boldsymbol{\rho})\boldsymbol{\epsilon}_t^*(\boldsymbol{\rho}), \boldsymbol{\eta}] \boldsymbol{\epsilon}_{\mathbf{r}t}(\boldsymbol{\rho})\boldsymbol{\epsilon}_{\mathbf{r}t}'(\boldsymbol{\rho}) - \boldsymbol{\Omega} \},$$

where $\boldsymbol{\epsilon}_{\mathbf{r}t}(\boldsymbol{\rho}) = \mathbf{r}_t - \mathbf{a} - \mathbf{b}\mu_M$.

In addition,

$$s_{J\mu_M}(\boldsymbol{\rho}, \boldsymbol{\eta}) = \frac{1}{2\sigma_M^2} \delta_{N+1}[\boldsymbol{\epsilon}_t^{*'}(\boldsymbol{\rho})\boldsymbol{\epsilon}_t^*(\boldsymbol{\rho}), \boldsymbol{\eta}] \epsilon_t(\mu_M),$$

and

$$s_{J\sigma_M^2}(\boldsymbol{\rho}, \boldsymbol{\eta}) = \frac{1}{2\sigma_M^4} \{ \delta_{N+1}[\boldsymbol{\epsilon}_t^{*'}(\boldsymbol{\rho})\boldsymbol{\epsilon}_t^*(\boldsymbol{\rho}), \boldsymbol{\eta}] \epsilon_t^2(\mu_M) - \sigma_M^2 \},$$

where $\epsilon_{Mt}(\mu_M) = r_{Mt} - \mu_M$.

Finally, the result follows from Proposition 1 in Fiorentini and Sentana (2007) if we exploit the fact that

$$\frac{\partial \text{vec}'[\Sigma(\boldsymbol{\rho})]}{\partial \mathbf{b}} \text{vec}[\Sigma^{-1}(\boldsymbol{\rho})] = \mathbf{0}$$

and

$$\frac{\partial \text{vec}'[\Sigma(\boldsymbol{\rho})]}{\partial \sigma_M^2} \text{vec}[\Sigma^{-1}(\boldsymbol{\rho})] = \frac{1}{\sigma_M^2}.$$

Specifically,

$$\begin{aligned}
M_{ll}(\boldsymbol{\eta}) &= E \left\{ \delta_{N+1}^2 [\boldsymbol{\epsilon}_t^{*'}(\boldsymbol{\rho}) \boldsymbol{\epsilon}_t^*(\boldsymbol{\rho}), \boldsymbol{\eta}] \frac{\boldsymbol{\epsilon}_t^{*'}(\boldsymbol{\rho}) \boldsymbol{\epsilon}_t^*(\boldsymbol{\rho})}{N+1} \middle| \boldsymbol{\phi} \right\} \\
&= E \left\{ \frac{2\partial\delta_{N+1}[\boldsymbol{\epsilon}_t^{*'}(\boldsymbol{\rho}) \boldsymbol{\epsilon}_t^*(\boldsymbol{\rho}), \boldsymbol{\eta}]}{\partial\boldsymbol{\zeta}} \frac{\boldsymbol{\epsilon}_t^{*'}(\boldsymbol{\rho}) \boldsymbol{\epsilon}_t^*(\boldsymbol{\rho})}{N+1} + \delta_{N+1}[\boldsymbol{\epsilon}_t^{*'}(\boldsymbol{\rho}) \boldsymbol{\epsilon}_t^*(\boldsymbol{\rho}), \boldsymbol{\eta}] \middle| \boldsymbol{\phi} \right\}, \\
M_{ss}(\boldsymbol{\eta}) &= \frac{N+1}{N+3} \left[1 + V \left\{ \delta_{N+1}[\boldsymbol{\epsilon}_t^{*'}(\boldsymbol{\rho}) \boldsymbol{\epsilon}_t^*(\boldsymbol{\rho}), \boldsymbol{\eta}] \frac{\boldsymbol{\epsilon}_t^{*'}(\boldsymbol{\rho}) \boldsymbol{\epsilon}_t^*(\boldsymbol{\rho})}{N+1} \middle| \boldsymbol{\phi} \right\} \right] \\
&= E \left\{ \frac{2\partial\delta_{N+1}[\boldsymbol{\epsilon}_t^{*'}(\boldsymbol{\rho}) \boldsymbol{\epsilon}_t^*(\boldsymbol{\rho}), \boldsymbol{\eta}]}{\partial\boldsymbol{\zeta}} \frac{[\boldsymbol{\epsilon}_t^{*'}(\boldsymbol{\rho}) \boldsymbol{\epsilon}_t^*(\boldsymbol{\rho})]^2}{(N+1)(N+3)} \middle| \boldsymbol{\phi} \right\} + 1, \\
M_{sr}(\boldsymbol{\eta}) &= E \left[\left\{ \delta_{N+1}[\boldsymbol{\epsilon}_t^{*'}(\boldsymbol{\rho}) \boldsymbol{\epsilon}_t^*(\boldsymbol{\rho}), \boldsymbol{\eta}] \frac{\boldsymbol{\epsilon}_t^{*'}(\boldsymbol{\rho}) \boldsymbol{\epsilon}_t^*(\boldsymbol{\rho})}{N+1} - 1 \right\} \mathbf{e}_{rt}'(\boldsymbol{\phi}) \middle| \boldsymbol{\phi} \right] \\
&= -E \left\{ \frac{\varsigma_t(\boldsymbol{\theta})}{N+1} \frac{\partial\delta_{N+1}[\boldsymbol{\epsilon}_t^{*'}(\boldsymbol{\rho}) \boldsymbol{\epsilon}_t^*(\boldsymbol{\rho}), \boldsymbol{\eta}]}{\partial\boldsymbol{\eta}'} \middle| \boldsymbol{\phi} \right\},
\end{aligned}$$

and the subscript $N+1$ in δ emphasises the cross-sectional dimension.

Interestingly, note that under Gaussianity $\mathcal{I}_{\omega\sigma_M^2}(\boldsymbol{\phi}) = \mathbf{0}$, which confirms that the estimators of the parameter of the marginal model for r_{Mt} and the conditional model for $\mathbf{r}_t$ will be independent.

Proposition 7

Since $\boldsymbol{\Sigma}^{-1/2}(\boldsymbol{\rho})[\mathbf{R}_t - \boldsymbol{\mu}(\boldsymbol{\rho})] | I_{t-1}; \boldsymbol{\varphi}_0 \sim i.i.d. s(\mathbf{0}, \mathbf{I}_{N+1}, \boldsymbol{\varrho}_0)$, we can write

$$\boldsymbol{\Sigma}^{-1/2}(\boldsymbol{\rho})[\mathbf{R}_t - \boldsymbol{\mu}(\boldsymbol{\rho})] = e_t \left(\frac{u_{0t}}{\sqrt{1 - u_{0t}^2}} \tilde{\mathbf{u}}_t \right)$$

where e_t is a positive random variable such that $E(e_t^2) = N+1$, u_{0t}^2 is a beta random variable with parameters $(1/2, N/2)$ and $\tilde{\mathbf{u}}_t$ is an independent uniform on the unit sphere surface in $\mathbb{R}^N$.

Given that the Cholesky decomposition of $\boldsymbol{\Sigma}(\boldsymbol{\rho})$ can be written as

$$\boldsymbol{\Sigma}^{1/2}(\boldsymbol{\rho}) = \begin{pmatrix} \sigma_M & \mathbf{0} \\ \mathbf{b}\sigma_M & \boldsymbol{\Omega}^{1/2} \end{pmatrix} \quad (\text{A8})$$

with $\boldsymbol{\Omega}^{1/2}$ denoting the Cholesky decomposition of $\boldsymbol{\Omega}$, we can write

$$\mathbf{R}_t - \boldsymbol{\mu}(\boldsymbol{\rho}) = \begin{pmatrix} \sigma_{M0} e_t u_{0t} \\ \mathbf{b}_0 \sigma_{M0} e_t u_{0t} + \boldsymbol{\Omega}_0^{1/2} e_t \sqrt{1 - u_{0t}^2} \tilde{\mathbf{u}}_t \end{pmatrix},$$

where u_{0t} is a random variable on $(-1, 1)$ with density $(1 - u_{0t}^2)^{N/2-1}/B(1/2, N/2)$. This follows from the symmetry of u_{0t} and the fact that the density of $|u_{0t}|$ is $2(1 - u_{0t}^2)^{N/2-1}/B(1/2, N/2)$ because the density of u_{0t}^2 is $(u_{0t}^2)^{-1/2}(1 - u_{0t}^2)^{N/2-1}/B(1/2, N/2)$. As a result, $\boldsymbol{\varepsilon}_t(\boldsymbol{\gamma}_0; \mathbf{R}_t) = \mathbf{r}_t - \mathbf{a}_0 - \mathbf{b}_0 r_{Mt} = \boldsymbol{\Omega}_0^{1/2} e_t \sqrt{1 - u_{0t}^2} \tilde{\mathbf{u}}_t$ and

$$\boldsymbol{\varepsilon}_t'(\boldsymbol{\gamma}_0; \mathbf{R}_t) \boldsymbol{\Omega}_0^{-1} \boldsymbol{\varepsilon}_t(\boldsymbol{\gamma}_0; \mathbf{R}_t) = e_t^2 (1 - u_{0t}^2)$$

because $\tilde{\mathbf{u}}'_t \tilde{\mathbf{u}}_t = 1$.

Let's now consider the following misspecified model

$$\mathbf{\Omega}^{-1/2}(\mathbf{r}_t - \mathbf{a} - \mathbf{b}r_{Mt})|r_{Mt}, \boldsymbol{\phi} \sim i.i.d.t(\mathbf{0}, \mathbf{I}_N, \eta)$$

and assume $\varsigma_t(\boldsymbol{\gamma}_0, \mathbf{v}_0, \tau) = \boldsymbol{\varepsilon}'_t(\boldsymbol{\gamma}_0)\tau^{-1}\boldsymbol{\Upsilon}^{-1}(\mathbf{v}_0)\boldsymbol{\varepsilon}_t(\boldsymbol{\gamma}_0) = (\tau_0/\tau)e_t^2(1 - u_{0t}^2)$. Hence, the blocks of the score corresponding to $\boldsymbol{\gamma}$, $\mathbf{v}$ and τ are given by (A3), (A4) and (A5) with $e_t^2(1 - u_{0t}^2)$ replacing ς_t . Then, the first part of this proposition can be obtained using the arguments in the proof of the first part of Proposition 4.

The proof of the second part is analogous to the proof of the second part of Proposition 4. Note, in particular, that having contemporaneous, conditionally heteroskedastic innovations is innocuous to obtain the relevant expressions since all the scalar terms $M_i^j(\boldsymbol{\phi}; \boldsymbol{\varphi}) = E \{ f_i^j[\varsigma_t(\boldsymbol{\rho})] | \boldsymbol{\varphi} \}$ appearing in $\mathcal{O}_t(\boldsymbol{\phi}_\infty; \boldsymbol{\varphi}_0)$ and $\mathcal{H}_t(\boldsymbol{\phi}_\infty; \boldsymbol{\varphi}_0)$ satisfy

$$E \{ f_i^j[\varsigma_t(\boldsymbol{\rho})] | \boldsymbol{\varphi}, r_{Mt} \} = E \{ f_i^j[\varsigma_t(\boldsymbol{\rho})] | \boldsymbol{\varphi} \}.$$

Finally, our parametrization implies that $\mathcal{O}(\boldsymbol{\phi}_\infty; \boldsymbol{\varphi}_0)$ and $\mathcal{H}(\boldsymbol{\phi}_\infty; \boldsymbol{\varphi}_0)$ will be block diagonal between $(\boldsymbol{\gamma}, \mathbf{v})$ and (τ, η) , as in Proposition 4.

References

- Abramowitz, M. and I.A. Stegun (1964). *Handbook of mathematical functions*, AMS 55, National Bureau of Standards.
- Amengual, D. and Sentana, E. (2008): “A comparison of mean-variance efficiency tests”, CEMFI Working Paper 0806.
- Bahadur, R. (1960). “Stochastic comparison of tests”, *The Annals of Mathematical Statistics*, 31, 276-295.
- Beaulieu, M.C., J.M. Dufour, and L. Khalaf (2007a). “Testing mean-variance efficiency in CAPM with possibly non-gaussian errors: an exact simulation-based approach”, *Journal of Business and Economic Statistics* 25, 398-410.
- Beaulieu, M.C., J.M. Dufour, and L. Khalaf (2007b). “Finite-sample identification-robust inference for unobservable zero-beta rates and portfolio efficiency with non-Gaussian distributions”, mimeo, McGill University.
- Beran, R. (1979). “Testing for ellipsoidal symmetry of a multivariate density”, *Annals of Statistics*, 7, 150-162.
- Berndt, E.R., B.H. Hall, R.E. Hall and J.A. Hausman (1974). “Estimation and inference in nonlinear structural models”, *Annals of Economic and Social Measurement*, 3, 653–665.
- Berk, J. (1997). “Necessary conditions for the CAPM”, *Journal of Economic Theory*, 73, 245-257.
- Black, F., M.C. Jensen, and M. Scholes (1972). “The capital asset pricing model: some empirical tests”, in Michael C. Jensen (ed.), *Studies in the Theory of Capital Markets*, Praeger Publishers.
- Campbell, J., W. Lo, and A.C. MacKinlay (1997). *The econometrics of financial markets*, Princeton: Princeton University Press.
- Chamberlain, G. (1983). “A characterization of the distributions that imply mean-variance utility functions”, *Journal of Economic Theory*, 29, 185-201.
- Cochrane, J. (2001). *Asset pricing*, Princeton: Princeton University Press.
- Cressie, N.A.C., A.S. Davis, J.L. Folks and G.E. Policello (1981). “The moment generating function and negative integer moments”, *American Statistician* 35, 148-150.
- Crowder, M.J. (1976). “Maximum likelihood estimation for dependent observations”, *Journal of the Royal Statistical Society B*, 38, 45-53.
- Davidson, R. and J.G. MacKinnon (1998). “Graphical methods for investigating the size and power of test statistics”, *The Manchester School*, 66, 1-26.

- De Roon, F.A. and T.E. Nijman (2001). “Testing for mean-variance spanning: a survey”, *Journal of Empirical Finance*, 8-2, 111-156.
- Erdélyi, A. (1981). *Higher transcendental functions, vol. 1*, Robert E. Krieger Publishing Company, Melbourne, FL.
- Errunza, V., K. Hogan and M. Hung (1999). “Can the gains from international diversification be achieved without trading abroad”, *Journal of Finance*, 54, 2075-2107.
- Fama, E.F. and French, K.R. (1993): “Common risk factors in the returns on stocks and bonds”, *Journal of Financial Economics* 33, 3-56.
- Fiorentini, G., E. Sentana, and G. Calzolari (2003). “Maximum likelihood estimation and inference on multivariate conditionally heteroscedastic dynamic regression models with Student t innovations”, *Journal of Business and Economic Statistics*, 24, 532-546.
- Fiorentini, G. and E. Sentana (2007). “On the efficiency and consistency of likelihood estimation in multivariate conditionally heteroskedastic dynamic regression models”, CEMFI Working Paper 0713.
- Geweke, J. (1981). “The approximate slopes of econometric tests”, *Econometrica*, 49, 1427-1442.
- Gibbons, M., S. Ross, and J. Shanken (1989). “A test of the efficiency of a given portfolio”, *Econometrica*, 57, 1121-1152.
- Hansen, L. (1982). “Large sample properties of generalized method of moments estimators”, *Econometrica*, 50, 1029–1054.
- Hodgson, D.J. (2000). “Unconditional pseudo-maximum likelihood and adaptive estimation in the presence of conditional heterogeneity of unknown form”, *Econometric Reviews*, 19, 175-206.
- Hodgson, D., O. Linton and K. Vorkink (2002). “Testing the capital asset pricing model efficiently under elliptical symmetry: a semiparametric approach”, *Journal of Applied Econometrics*, 17, 617-639.
- Hodgson, D., O. Linton and K. Vorkink (2004). “Testing forward exchange rate unbiasedness efficiently: a semiparametric approach”, *Journal of Applied Economics*, 7-I, 325-353.
- Huberman, G. and S. Kandel (1987). “Mean-variance spanning”, *Journal of Finance*, 42, 873-888.
- Jobson, J.D. and B. Korkie (1982). “Estimation for Markowitz portfolios”, *Journal of the American Statistical Association*, 75, 544-554.
- Kan, R. and G. Zhou (2006). “Modeling non-normality using multivariate t : implications for asset pricing”, mimeo Washington University in St.Louis.

- Kotz, S. (1975). “Multivariate distributions at a cross-road”, in G. P. Patil, S. Kotz and J.K. Ord (eds.) *Statistical Distributions in Scientific Work*, vol. I, 247-270, Reidel.
- MacKinlay, A.C. and M. Richardson (1991). “Using generalized method of moments to test mean-variance efficiency”, *Journal of Finance*, 46, 511–527.
- Magnus, J.R. and H. Neudecker (1988). *Matrix differential calculus with applications on statistics and econometrics*, Wiley: Chichester.
- Mardia, K. (1970). “Measures of multivariate skewness and kurtosis with applications”, *Biometrika*, 57, 519–530.
- Markowitz, H. (1952). “Portfolio selection”, *Journal of Finance*, 8, pp. 77-91.
- Masoom Ali, M. and S. Nadarajah (2007). “Information matrices for Normal and Laplace mixtures”, *Information Sciences* 177, 947-955.
- Meddahi, N. and E. Renault (1998). “Quadratic M-estimators for ARCH-type processes”, CIRANO Working paper 98s-29.
- Mencía, F.J. and E. Sentana (2009). “Distributional tests in multivariate dynamic models with Normal and Student t innovations”, mimeo, CEMFI.
- Morales, L. (2009): “Mean-variance efficiency tests with conditioning information: a comparison”, CEMFI Master Thesis 0902.
- Newey, W. K. and D. L. McFadden (1994). “Large sample estimation and hypothesis testing”, in R. F. Engle and (eds.), *Handbook of Econometrics*, vol. IV, 2111-2245, Elsevier.
- Owen, J. and R. Rabinovitch (1983). “On the class of elliptical distributions and their applications to the theory of portfolio choice”, *Journal of Finance*, 38, 745-752.
- Peñaranda, F. and E. Sentana (2008). “Spanning tests in return and stochastic discount factor mean-variance frontiers: a unifying approach”, CEMFI Working paper No. 0410.
- Renault, E. and E. Sentana (2003). “GMM estimation of multivariate regression models with symmetric errors”, mimeo, CEMFI.
- Sentana, E. (2008). “The econometrics of mean-variance efficiency: a survey”, CEMFI working Paper 0807.
- Shanken, J. (1986). “Testing portfolio efficiency when the zero-beta rate is unknown: A note”, *Journal of Finance*, 41, 269–276.
- Vorkink, K. (2003). “Return distributions and improved tests of asset pricing models”, *Review of Financial Studies*, 16, 845–874.
- White, H. (1980). “A heteroskedastic-consistent covariance matrix estimator and a direct test for heteroskedasticity”, *Econometrica*, 45, 817-838.

Zhou, G. (1991). “Small sample tests of portfolio efficiency”, *Journal of Financial Economics*, 30, 165–191.

Zhou, G. (1993). “Asset-pricing tests under alternative distributions”, *Journal of Finance*, 48, 1927–1942.

Table 1Table 1.a: Intercept estimates in: $\mathbf{r}_t = \mathbf{a} + \mathbf{b}r_{Mt} + \mathbf{u}_t$

Industry portfolios						
Category	GMM		HLV		t ML	
	Coeff.	Std.Err.	Coeff.	Std.Err.	Coeff.	Std.Err.
Cnsmr	0.099	0.091	0.026	0.076	0.023	0.056
Manuf	0.134	0.076	0.069	0.063	0.123	0.064
SHiTec	-0.086	0.117	-0.064	0.097	-0.146	0.099
Hlth	0.205	0.146	0.072	0.123	0.092	0.135
Other	0.088	0.087	0.016	0.073	0.003	0.085
Book-to-market sorted portfolios						
Quintile	GMM		HLV		t ML	
	Coeff.	Std.Err.	Coeff.	Std.Err.	Coeff.	Std.Err.
1	-0.108	0.062	-0.111	0.053	-0.121	0.035
2	0.040	0.060	-0.068	0.050	0.018	0.049
3	0.151	0.073	0.019	0.061	0.116	0.062
4	0.328	0.087	0.105	0.073	0.215	0.064
5	0.430	0.109	0.221	0.093	0.308	0.111

Table 1.b: Mean-variance efficiency tests ($H_0 : \mathbf{a} = \mathbf{0}$)

Industry portfolios					
	GMM	GMM robust	HLV	t ML	t ML robust
Statistic	12.056	10.911	2.194	15.226	14.365
p -value	0.034	0.053	0.822	0.009	0.013
Book-to-market sorted portfolios					
	GMM	GMM robust	HLV	t ML	t ML robust
Statistic	21.350	21.417	12.837	21.846	19.772
p -value	0.001	0.001	0.025	0.001	0.001

Notes: Sample: July:1962-June:2007. Industry definitions: Cnsmr: Consumer Durables, Non-Durables, Wholesale, Retail, and Some Services (Laundries, Repair Shops). Manuf: Manufacturing, Energy, and Utilities. HiTec: Business Equipment, Telephone and Television Transmission. Hlth: Healthcare, Medical Equipment, and Drugs. Other: Other – Mines, Constr, BldMt, Trans, Hotels, Bus Serv, Entertainment, Finance.

Table 2: Conditional heteroskedasticity test

Category	Cnsmr	Industry portfolios			
		Manuf	HiTec	Hlth	Other
Statistic	45.026	9.633	14.635	48.257	4.866
<i>p</i> -value	0.000	0.008	0.001	0.000	0.088

Quintile	Book-to-market sorted portfolios				
	1	2	3	4	5
Statistic	24.070	11.748	27.480	41.262	62.098
<i>p</i> -value	0.000	0.003	0.000	0.000	0.000

Notes: July:1962-June:2007. Based on the statistical significance of $\delta_i = (\delta_{1i}, \delta_{2i})'$ in $\hat{\varepsilon}_{it}^2 = c_i + \delta_{1i}r_{Mt} + \delta_{2i}r_{Mt}^2 + v_{it}$, where $\hat{\varepsilon}_{it}$'s are the OLS residuals from a regression of r_{it} on a constant and r_{Mt} . The test statistic, nR^2 —where R^2 is the coefficient of determination of the regression—, is distributed as a χ_2^2 under the null hypothesis of conditional homoskedasticity.