

DISCUSSION OF: TESTING OUT-OF-SAMPLE PORTFOLIO PERFORMANCE

by Ekaterina Kazak and Winfried Pohlmeier

Robert A. Korajczyk

Northwestern University

15 September 2017

Outline

- 1 Overview
- 2 Benchmarks
- 3 The paper improves on the plug-in portfolio strategies
- 4 Other ways to improve on the plug-in portfolio strategies?

Difficult to beat naïve ($1/N$) diversification through portfolio optimization.

- ① Optimization with known parameters.
- ② Plugging in mean and covariance estimates for true parameters.
- ③ Out of sample randomness.

Outline

- 1 Overview
- 2 Benchmarks
- 3 The paper improves on the plug-in portfolio strategies
- 4 Other ways to improve on the plug-in portfolio strategies?

Normal i.i.d. world

$$\mathbb{E}[r_t] = \mu, \text{ Var}[r_t] = \Sigma. \quad (1)$$

- ① Global mean-variance efficient portfolio (GMVP): $\omega(g) = \frac{\Sigma^{-1} \mathbf{1}}{\mathbf{1}' \Sigma^{-1} \mathbf{1}}$.
- ② Mean-variance efficient for risk aversion γ : $\omega(m) = \frac{\Sigma^{-1} \mu}{\gamma}$.
- ③ Equally weighted: $\omega(e) = \frac{\mathbf{1}}{N}$.

Use “plug-In” portfolio weights

- ① Global mean-variance efficient portfolio (GMVP): $\hat{\omega}(g) = \frac{\hat{\Sigma}^{-1} \mathbf{1}}{\mathbf{1}' \hat{\Sigma}^{-1} \mathbf{1}}$.
- ② Mean-variance efficient for risk aversion γ : $\hat{\omega}(m) = \frac{\hat{\Sigma}^{-1} \hat{\mu}}{\gamma}$.
- ③ Equally weighted: $\hat{\omega}(e) = \frac{\mathbf{1}}{N}$.

Comparing two portfolio strategies, s and $\tilde{s}$

Differences in certainty equivalent

$$CE(\omega(s)) = \mu_p(s) - \frac{\gamma}{2} \sigma_p^2(s) \quad (2)$$

$$\Delta_o(s, \tilde{s}) = CE(\omega(s)) - CE(\omega(\tilde{s}))$$

$$\Delta_{op}(s, \tilde{s}) = \Delta_o(s, \tilde{s}) - A$$

where A is due to the volatility induced by the fact that $\hat{\omega}(s) \neq \omega(s)$ and that post-formation returns are random.

The $1/N$ approach does not optimize $(-)$ but it is not subject to the estimation error in the portfolio weights $(+)$.

The paper shows convincingly that the need to estimate parameters has a very large effect on out of sample portfolio performance, particularly as N grows,

Outline

- 1 Overview
- 2 Benchmarks
- 3 The paper improves on the plug-in portfolio strategies
- 4 Other ways to improve on the plug-in portfolio strategies?

Paper: Alternative weights - Shrink $\hat{\Sigma}$ to I

Choose shrinkage parameter as is Ledoit and Wolf (2003).

Ledoit and Wolfe (2003) shrink $\hat{\Sigma}$ but we are shrinking $\hat{\Sigma}^{-1}$. Does the same shrinkage factor work?

Outline

- 1 Overview
- 2 Benchmarks
- 3 The paper improves on the plug-in portfolio strategies
- 4 Other ways to improve on the plug-in portfolio strategies?

Alternative weights - Correct the biases

The plug-in weights are biased, due to Jensen's inequality.

$$\mathbb{E}[\hat{\Sigma}^{-1}] = \frac{T}{T - N - 2} \Sigma^{-1}. \quad (3)$$

For the tangency portfolio this suggests a simple adjustment:

$$\hat{\omega}_a(m) = \frac{N-T-2}{T} \frac{\hat{\Sigma}^{-1} \hat{\mu}}{\gamma}.$$

For the empirical sample sizes used in Table 4, this adjustment ranges from 0.94 ($N = 5$, $T = 120$; 6% bias) to 0.13 ($N = 50$, $T = 60$; 650% bias).

Similar to a shrinkage estimator, except shrinking to zero to correct the bias and letting the bias determine the shrinkage factor.

Alternative weights - Correct the biases

The correction is more complicated for the GMVP since we have two terms subject to Jensen's inequality:

① $\hat{\Sigma}^{-1}_{\ell}$

② $\frac{1}{\ell' \hat{\Sigma}^{-1}_{\ell}}$

Alternative weights - Impose constraints on $\hat{\omega}$

For example, one could impose a ban on short positions: $\hat{\omega} \geq 0$

- ① Not theoretically correct: Green and Hollifield (*JF* 1992)
- ② Works well in practice: Jagannathan and Ma (*JF* 2002)
- ③ Can be combined with any of the alternative weights: plug-in, bias-corrected, shrinkage.

Alternative weights - Impose a factor structure

- ① $\hat{\Sigma}$ has $N \times T$ observations and $\frac{N(N+1)}{2}$ parameters, for $\frac{2T}{N+1}$ observations per parameter. Its inverse is badly behaved when N is large, relative to T .
- ② If we impose a K -factor structure, with $K \ll T$ the observations per parameter improve to $\frac{T}{K+1}$.

Alternative weights - Shrink to a factor structure

Shrink a bias-corrected estimate of Σ^{-1} toward an estimate based on a factor model $(B'\Sigma_F B + \Sigma_\epsilon)^{-1}$.

- ① Observable factor model (e.g., CAPM, Fama-French factors).
- ② Latent Factor model (e.g., Connor and Korajczyk (1986)).

