

On the combination of naive and mean-variance portfolio strategies

Nathan Lassance

UCLouvain (BE), LFIN

Rodolphe Vanderveken

UCLouvain (BE), LFIN

Frédéric Vrins

UCLouvain (BE), LFIN

BFRF 2023

April 2023

Motivation



Efficient mean-variance portfolios

- Markowitz (1952) defines efficient **mean-variance portfolios**:

$$\mathbf{w}^* = \underset{\mathbf{w}}{\operatorname{argmax}} \underbrace{\mathbf{w}^\top \boldsymbol{\mu}}_{\text{portfolio mean}} - \frac{\gamma}{2} \overbrace{\mathbf{w}^\top \boldsymbol{\Sigma} \mathbf{w}}^{\text{portfolio risk}} = \frac{1}{\gamma} \boldsymbol{\Sigma}^{-1} \boldsymbol{\mu}$$

- $\mathbf{w}$: vector of weights on **risky** assets ($1 - \mathbf{1}'\mathbf{w}$ is invested in the **risk-free asset**)
- $\boldsymbol{\mu}$ and $\boldsymbol{\Sigma}$: vector of mean returns and covariance matrix of returns
- γ : risk-aversion coefficient

Efficient mean-variance portfolios

- Markowitz (1952) defines efficient **mean-variance portfolios**:

$$\mathbf{w}^* = \underset{\mathbf{w}}{\operatorname{argmax}} \underbrace{\mathbf{w}^\top \boldsymbol{\mu}}_{\text{portfolio mean}} - \frac{\gamma}{2} \overbrace{\mathbf{w}^\top \boldsymbol{\Sigma} \mathbf{w}}^{\text{portfolio risk}} = \frac{1}{\gamma} \boldsymbol{\Sigma}^{-1} \boldsymbol{\mu}$$

- $\mathbf{w}$: vector of weights on **risky** assets ($1 - \mathbf{1}'\mathbf{w}$ is invested in the **risk-free asset**)
 - $\boldsymbol{\mu}$ and $\boldsymbol{\Sigma}$: vector of mean returns and covariance matrix of returns
 - γ : risk-aversion coefficient
-
- Challenge:** Investors must estimate $\boldsymbol{\mu}$ and $\boldsymbol{\Sigma}$ and rely on

$$\hat{\mathbf{w}}^* = \frac{1}{\gamma} \hat{\boldsymbol{\Sigma}}^{-1} \hat{\boldsymbol{\mu}}$$

Estimation risk and poor performance

- **Problem:** $\hat{\mathbf{w}}^*$ performs poorly out-of-sample. The naive equally weighted portfolio ($\mathbf{w}_{ew} = \mathbf{1}/N$) often outperforms it.
- See DeMiguel et al. (2009, RFS).
- The performance problem notably stems from [estimation risk](#)

How to deal with estimation risk?

- Two extreme alternatives: plug-in or naive

How to deal with estimation risk?

- Two extreme alternatives: plug-in or naive
- Plug-in: relying on $\hat{\mathbf{w}}^*$. It is **theoretically optimal**, but subject to **heavy estimation risk**
 - Approaches exist to limit estimation errors on parameters (e.g. LW shrinkage, norm-constrained portfolios, etc.)

How to deal with estimation risk?

- Two extreme alternatives: plug-in or naive
- Plug-in: relying on $\hat{\mathbf{w}}^*$. It is **theoretically optimal**, but subject to **heavy estimation risk**
 - Approaches exist to limit estimation errors on parameters (e.g. LW shrinkage, norm-constrained portfolios, etc.)
- Naive: relying on $\mathbf{w}_{ew}$. It is **insensitive** to **estimation risk** but it is also **suboptimal**!

How to deal with estimation risk?

- Two extreme alternatives: plug-in or naive
- Plug-in: relying on $\hat{\mathbf{w}}^*$. It is **theoretically optimal**, but subject to **heavy estimation risk**
 - Approaches exist to limit estimation errors on parameters (e.g. LW shrinkage, norm-constrained portfolios, etc.)
- Naive: relying on $\mathbf{w}_{ew}$. It is **insensitive** to **estimation risk** but it is also **suboptimal**!
- Let's try to find the middle ground ...

Combination of naive and mean-variance portfolio strategies

- Tu and Zhou (2011, JFE) introduce a **portfolio combination**:

$$\hat{\mathbf{w}}(\boldsymbol{\kappa}) = \kappa_1 \hat{\mathbf{w}}^* + \kappa_2 \mathbf{w}_{ew} \quad \text{with} \quad \mathbf{c} := \kappa_1 + \kappa_2 = 1.$$

- $\boldsymbol{\kappa} = (\kappa_1, \kappa_2)$ are the **combination coefficients**.

Combination of naive and mean-variance portfolio strategies

- Tu and Zhou (2011, JFE) introduce a **portfolio combination**:

$$\hat{\mathbf{w}}(\boldsymbol{\kappa}) = \kappa_1 \hat{\mathbf{w}}^* + \kappa_2 \mathbf{w}_{ew} \quad \text{with} \quad \mathbf{c} := \kappa_1 + \kappa_2 = 1.$$

- $\boldsymbol{\kappa} = (\kappa_1, \kappa_2)$ are the **combination coefficients**.
- As in Kan And Zhou (2007, JFQA), they are found by maximizing the **expected out-of-sample utility (EU)**

$$\begin{aligned} \boldsymbol{\kappa}^c &= \underset{\boldsymbol{\kappa}}{\operatorname{argmax}} \quad \mathbb{E} \left[\underbrace{\hat{\mathbf{w}}(\boldsymbol{\kappa})' \boldsymbol{\mu} - \frac{\gamma}{2} \hat{\mathbf{w}}(\boldsymbol{\kappa})' \boldsymbol{\Sigma} \hat{\mathbf{w}}(\boldsymbol{\kappa})}_{\text{out-of-sample utility}} \right] \quad \text{s.t.} \quad \mathbf{c} \\ &= \underset{\boldsymbol{\kappa}}{\operatorname{argmax}} \quad \mathbf{EU} \quad \text{s.t.} \quad \mathbf{c} \end{aligned}$$

Combination of naive and mean-variance portfolio strategies

- Tu and Zhou (2011, JFE) introduce a **portfolio combination**:

$$\hat{\mathbf{w}}(\boldsymbol{\kappa}) = \kappa_1 \hat{\mathbf{w}}^* + \kappa_2 \mathbf{w}_{ew} \quad \text{with} \quad \mathbf{c} := \kappa_1 + \kappa_2 = 1.$$

- $\boldsymbol{\kappa} = (\kappa_1, \kappa_2)$ are the **combination coefficients**.
- As in Kan And Zhou (2007, JFQA), they are found by maximizing the **expected out-of-sample utility (EU)**

$$\begin{aligned} \boldsymbol{\kappa}^c &= \underset{\boldsymbol{\kappa}}{\operatorname{argmax}} \quad \mathbb{E} \left[\underbrace{\hat{\mathbf{w}}(\boldsymbol{\kappa})' \boldsymbol{\mu} - \frac{\gamma}{2} \hat{\mathbf{w}}(\boldsymbol{\kappa})' \boldsymbol{\Sigma} \hat{\mathbf{w}}(\boldsymbol{\kappa})}_{\text{out-of-sample utility}} \right] \quad \text{s.t.} \quad \mathbf{c} \\ &= \underset{\boldsymbol{\kappa}}{\operatorname{argmax}} \quad \mathbf{EU} \quad \text{s.t.} \quad \mathbf{c} \end{aligned}$$

- This combined rule significantly outperforms $\hat{\mathbf{w}}^*$ and may outperform $\mathbf{w}_{ew}$.

This seems all well and good, but ...

- Tu and Zhou (2011) force combination coefficients to sum to one:

$$\hat{\mathbf{w}}(\kappa) = \kappa_1 \hat{\mathbf{w}}^* + \kappa_2 \mathbf{w}_{ew} \quad \text{with} \quad \mathbf{c} := \kappa_1 + \kappa_2 = 1.$$

- The constraint is needed in combinations of fully invested portfolios

This seems all well and good, but ...

- Tu and Zhou (2011) force combination coefficients to sum to one:

$$\hat{\mathbf{w}}(\kappa) = \kappa_1 \hat{\mathbf{w}}^* + \kappa_2 \mathbf{w}_{ew} \quad \text{with} \quad \mathbf{c} := \kappa_1 + \kappa_2 = 1.$$

- The constraint is needed in combinations of fully invested portfolios
- **Problem:** $\hat{\mathbf{w}}^*$ is **not** a fully invested portfolio: $\mathbf{1}'\hat{\mathbf{w}}^* \neq 1$.
 - It invests $\pi_{rf} = 1 - \mathbf{1}'\hat{\mathbf{w}}^*$ in the risk-free asset

This seems all well and good, but ...

- Tu and Zhou (2011) force combination coefficients to sum to one:

$$\hat{\mathbf{w}}(\boldsymbol{\kappa}) = \kappa_1 \hat{\mathbf{w}}^* + \kappa_2 \mathbf{w}_{ew} \quad \text{with} \quad \mathbf{c} := \kappa_1 + \kappa_2 = 1.$$

- The constraint is needed in combinations of fully invested portfolios
- **Problem:** $\hat{\mathbf{w}}^*$ is **not** a fully invested portfolio: $\mathbf{1}'\hat{\mathbf{w}}^* \neq 1$.
 - It invests $\pi_{rf} = 1 - \mathbf{1}'\hat{\mathbf{w}}^*$ in the risk-free asset
- We actually invest in three portfolios: the sample tangent, the risk-free asset, and the equally-weighted portfolio...
- ... But under $\mathbf{c}$, we only have one coefficient to control the weight allocated to each one of them !

This paper

- We propose to relax the constraint on combination coefficients:

$$\hat{\mathbf{w}}(\kappa) = \kappa_1 \hat{\mathbf{w}}^* + \kappa_2 \mathbf{w}_{ew} \quad \text{with } \kappa \in \mathcal{C}$$

- We find the **unconstrained combination coefficients** that maximize the **expected out-of-sample utility** of the portfolio combination:

$$\kappa^u = \underset{\kappa}{\operatorname{argmax}} \quad \mathbf{EU} \quad \text{s.t. } \kappa \in \mathcal{C}$$

Contributions



Contributions

Theoretical results:

1. We derive the **unconstrained combination coefficients**
2. We highlight analytically the **key benefits** of relaxing the constraint
3. We introduce a **shrinkage strategy** alleviating estimation risk

Empirical results:

1. The unconstrained strategy always delivers a **positive net utility**
2. The constrained strategy may deliver **negative net utility**
3. The **shrinkage strategy** accomplishes its mission

Unconstrained vs. constrained
combination



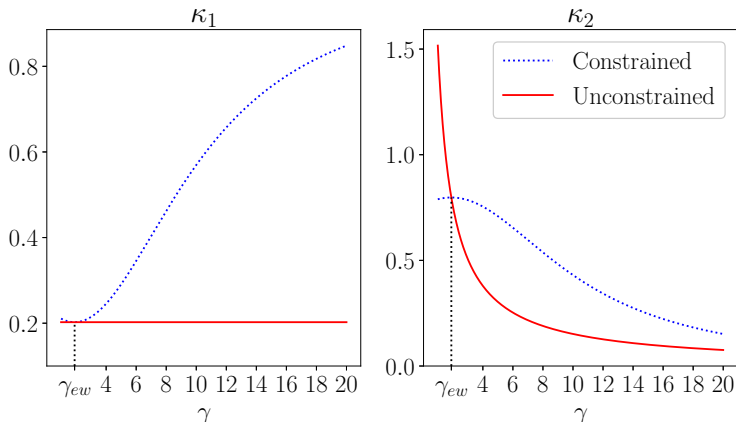
Benefits of relaxing the constraint

Key benefits of relaxing the constraint;

1. Optimal exposition to the three funds
2. Outperformance relative to the risk-free asset
3. Less extreme portfolio weights
4. Outperformance in expected out-of-sample utility (EU)
5. Outperformance relative to the two-fund rule
6. Outperformance in expected out-of-sample Sharpe Ratio

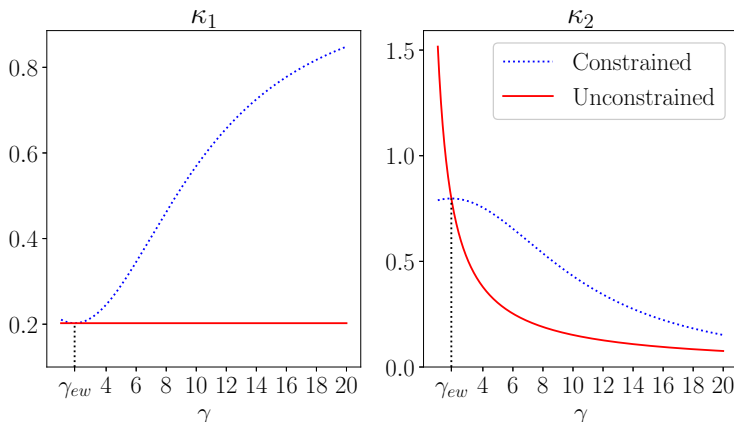
Combination coefficients

$$\hat{\mathbf{w}}(\kappa) = \kappa_1 \hat{\mathbf{w}}^* + \kappa_2 \mathbf{w}_{ew}$$



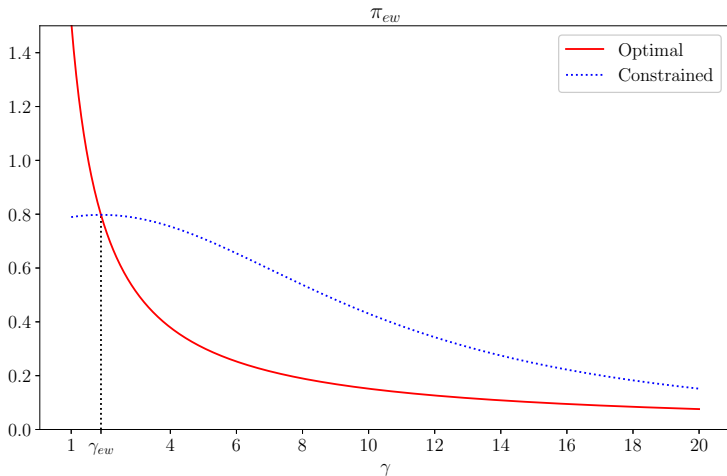
Combination coefficients

$$\hat{\mathbf{w}}(\kappa) = \kappa_1 \hat{\mathbf{w}}^* + \kappa_2 \mathbf{w}_{ew}$$



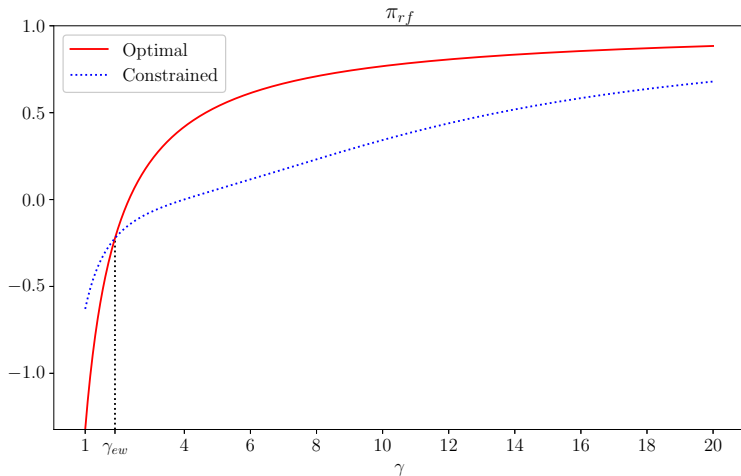
- Key facts: $\kappa_1^c \geq \kappa_1^u$ and equality iff $\gamma = \gamma_{ew} = \mu_{ew}/\sigma_{ew}^2$
- We use the 25 portfolios of stocks sorted on size and book-to-market (25SBTM) dataset spanning 07/26 to 12/21.

Benefit 1: optimal exposure to the equally-weighted portfolio



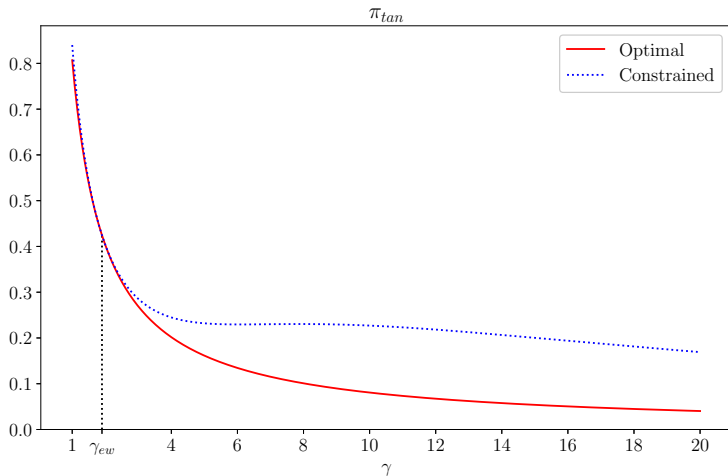
When $\gamma = 10$, $T = 120$, and $N = 25$, $\pi_{ew}^c = 43\%$ and $\pi_{ew}^u = 15\%$

Benefit 1: optimal exposure to the risk-free asset



When $\gamma = 10$, $T = 120$, and $N = 25$, $\pi_{rf}^c = 34\%$ and $\pi_{rf}^u = 77\%$

Benefit 1: optimal exposure to the tangent portfolio



When $\gamma = 10$, $T = 120$, and $N = 25$, $\pi_{tan}^c = 23\%$ and $\pi_{tan}^u = 8\%$

Benefit 2: Outperformance relative to the risk-free asset

Problem: With estimation risk, the constrained combination can underperform the risk-free asset for investors with a risk-aversion higher than a given threshold: γ_{neg}^1

¹ $\gamma_{neg} = 5.41$ with the 25SBTM dataset.

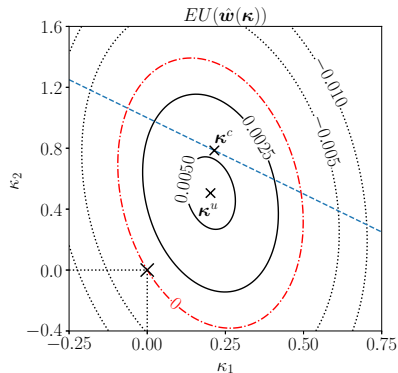
Benefit 2: Outperformance relative to the risk-free asset

Problem: With estimation risk, the constrained combination can **underperform the risk-free asset** for investors with a risk-aversion higher than a given threshold: γ_{neg}^1

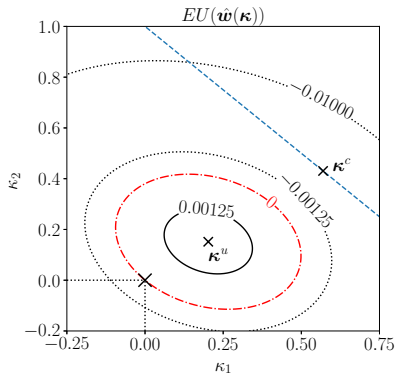
Intuition: The combination should theoretically outperform each of its components, but it might not be possible because of the constraint.

¹ $\gamma_{neg} = 5.41$ with the 25SBTM dataset.

Benefit 2: Outperformance relative to the risk-free asset



(a) $\gamma = 3 < \gamma_{neg}$



(b) $\gamma = 10 > \gamma_{neg}$

Shrinkage strategy & empirical results



Shrinkage strategy

- **Practical concern:** $\hat{\kappa}_2^u$ is unbounded and very sensitive to estimation errors in $\hat{\mu}$

Shrinkage strategy

- **Practical concern:** $\hat{\kappa}_2^u$ is unbounded and very sensitive to estimation errors in $\hat{\mu}$
- Estimation errors have a substantial impact on out-of-sample performance. They hurt most κ^u relative to κ^c when $\gamma = \gamma_{ew}$

Shrinkage strategy

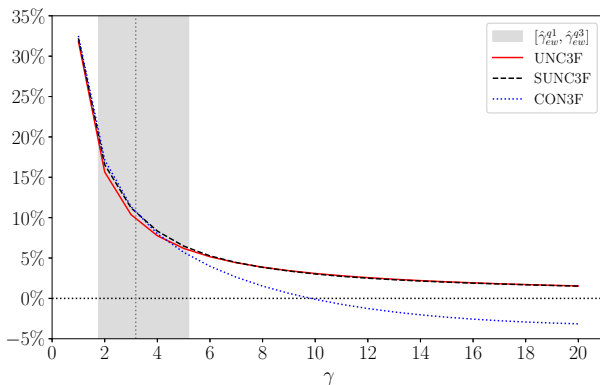
- **Practical concern:** $\hat{\kappa}_2^u$ is unbounded and very sensitive to estimation errors in $\hat{\mu}$
- Estimation errors have a substantial impact on out-of-sample performance. They hurt most κ^u relative to κ^c when $\gamma = \gamma_{ew}$
- $\hat{\kappa}_1^u$ is bounded in $[0, 1]$ and not very sensitive to estimation errors
- We define a **shrinkage strategy** in which we shrink $\hat{\kappa}_2^u$ toward the target $1 - \kappa_1^u$:

$$\hat{\kappa}_2^u(\delta) = (1 - \delta)\hat{\kappa}_2^u + \delta(1 - \kappa_1^u), \quad (1)$$

where $\delta \in [0, 1]$ is the shrinkage intensity and is calibrated to minimize the MSE of $\hat{\kappa}_2^u(\delta)$

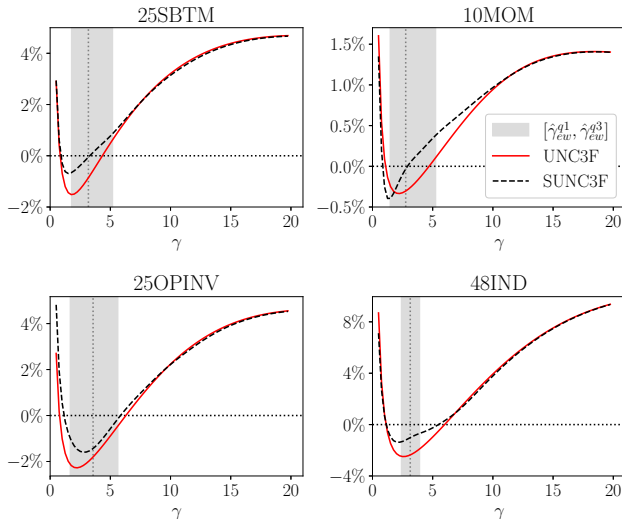
Empirical results

Net out-of-sample utility delivered by portfolio combinations (25SBTM)



Empirical results

Difference in net out-of-sample utility relative to the constrained strategy



Takeaways



Takeaways

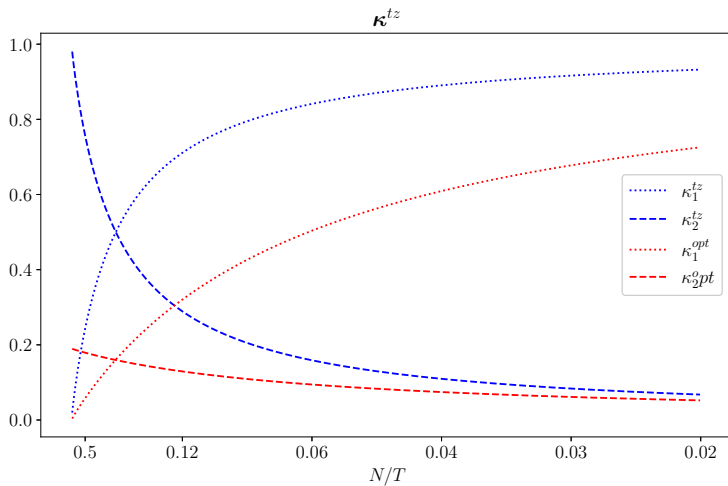
1. We derive the **unconstrained combination coefficients**
2. We highlight analytically the **key benefits** of relaxing the constraint
3. We introduce a **shrinkage strategy** to alleviate estimation risk
4. The unconstrained strategy always delivers a **positive net utility**
5. The constrained strategy may deliver **negative net utility**
6. The **shrinkage strategy** accomplishes its mission

Thank you!

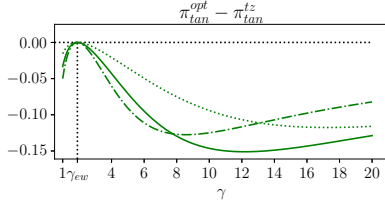
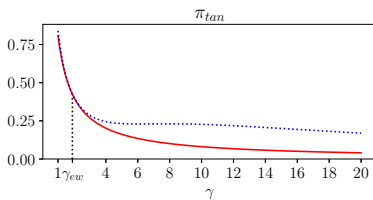
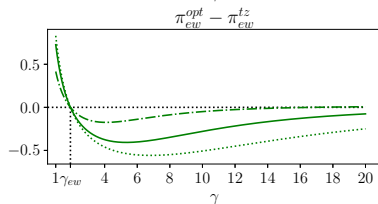
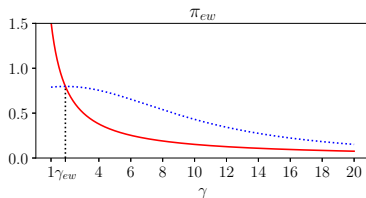
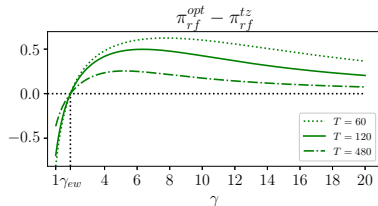
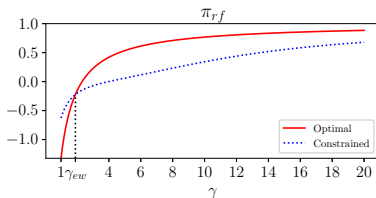
Please send comments to

`rodolphe.vanderveken@uclouvain.be`

Impact of estimation risk



optimal exposure to the three funds + estimation risk



Combination coefficients

The **constrained combination coefficients** are

$$\kappa_1^c = \frac{\psi^2 + \sigma_{ew}^2(\gamma - \gamma_{ew})^2}{\psi^2 + \sigma_{ew}^2(\gamma - \gamma_{ew})^2 + d} \in [0, 1] \quad \text{and} \quad \kappa_2^c = 1 - \kappa_1^c \in [0, 1].$$

The **unconstrained combination coefficients** are

$$\kappa_1^u = \frac{\psi^2}{\psi^2 + d} \in [0, 1] \quad \text{and} \quad \kappa_2^u = \frac{\gamma_{ew}}{\gamma}(1 - \kappa_1^u) \in \mathbb{R},$$

d is a parameter that increases with estimation risk N/T

γ_{neg} threshold definition

$$\gamma > \gamma_{ew} \left(1 + \sqrt{1 + \frac{\theta^4 / \theta_{ew}^2}{d - \theta^2}} \right) = \gamma_{neg}.$$

Notation

- Sample mean-variance portfolio (SMV) (from sample of size T)

$$\hat{\mathbf{w}}^* = \frac{1}{\gamma} \hat{\Sigma}^{-1} \hat{\mu}$$

Notation

- Sample mean-variance portfolio (SMV) (from sample of size T)

$$\hat{\mathbf{w}}^* = \frac{1}{\gamma} \hat{\Sigma}^{-1} \hat{\boldsymbol{\mu}}$$

- We define $\mu_{ew} = \mathbf{w}'_{ew} \boldsymbol{\mu}$ and $\sigma_{ew}^2 = \mathbf{w}'_{ew} \boldsymbol{\Sigma} \mathbf{w}_{ew}$. Moreover, we define $\gamma_{ew} = \mu_{ew} / \sigma_{ew}^2$

Notation

- Sample mean-variance portfolio (SMV) (from sample of size T)

$$\hat{\mathbf{w}}^* = \frac{1}{\gamma} \hat{\Sigma}^{-1} \hat{\boldsymbol{\mu}}$$

- We define $\mu_{ew} = \mathbf{w}'_{ew} \boldsymbol{\mu}$ and $\sigma_{ew}^2 = \mathbf{w}'_{ew} \boldsymbol{\Sigma} \mathbf{w}_{ew}$. Moreover, we define $\gamma_{ew} = \mu_{ew} / \sigma_{ew}^2$
- We define $\theta^2 = \boldsymbol{\mu}' \boldsymbol{\Sigma}^{-1} \boldsymbol{\mu}$ as the maximum squared Sharpe ratio, $\theta_{ew}^2 = \mu_{ew}^2 / \sigma_{ew}^2$ as the squared Sharpe ratio of the EW portfolio, and ψ^2 as the difference between the two:

$$\psi^2 = \theta^2 - \theta_{ew}^2 \geq 0.$$

Sample estimators

$$\hat{\mu} = \frac{1}{T} \sum_{t=1}^T R_t, \quad \hat{\Sigma} = \frac{1}{T - N - 2} \sum_{t=1}^T (R_t - \hat{\mu})(R_t - \hat{\mu})'.$$

Efficient frontier - Short reminder

