

Financial Data: I.I.D. Modeling

Portfolio Optimization — Chapter 3

Daniel P. Palomar (2025). *Portfolio Optimization: Theory and Application*.
Cambridge University Press.

portfoliooptimizationbook.com

Outline

- 1 I.I.D. Model
- 2 Sample Estimators
- 3 Location Estimators
- 4 Gaussian ML Estimators
- 5 Heavy-Tailed ML Estimators
- 6 Prior Information
 - Shrinkage
 - Factor Models
 - Black-Litterman Model
- 7 Summary

Executive Summary

- The efficient-market hypothesis suggests that a security's price reflects its intrinsic value, incorporating all available information.
- Then, prices can be modeled as a random walk with returns being independent and identically distributed (i.i.d.) random variables.
- These slides explore various methods to characterize the multivariate i.i.d. distribution of returns.
- Methods range from simple sample estimators to more advanced robust non-Gaussian estimators.
- Advanced estimators incorporate prior information through:
 - Shrinkage techniques
 - Factor modeling
 - Prior views
- Reference: (Palomar 2025, chap. 3)

Outline

- 1 I.I.D. Model
- 2 Sample Estimators
- 3 Location Estimators
- 4 Gaussian ML Estimators
- 5 Heavy-Tailed ML Estimators
- 6 Prior Information
 - Shrinkage
 - Factor Models
 - Black-Litterman Model
- 7 Summary

Introduction to Financial Data Modeling

- Financial data modeling is crucial for understanding and predicting market behaviors.
- It involves N securities or assets from various classes (e.g., bonds, equities, commodities).

Random Returns Representation

- Random returns of assets at time t are denoted by $\mathbf{x}_t$ or $\mathbf{r}_t$.
- The time index t can represent different periods (minutes to years).

I.I.D. Model for Returns

- Returns are modeled as:

$$\mathbf{x}_t = \boldsymbol{\mu} + \boldsymbol{\epsilon}_t.$$

- $\boldsymbol{\mu}$ represents the expected return, and $\boldsymbol{\epsilon}_t$ is the residual with zero mean.
- $\boldsymbol{\Sigma}$ denotes the covariance matrix of residuals.

Efficient-Market Hypothesis

- The i.i.d. model is motivated by the efficient-market hypothesis (EMH).
- Eugene F. Fama, a proponent of EMH, won the Nobel Prize in 2013.

Random Walk Model

- The i.i.d. model corresponds to the random walk model on log-prices $\mathbf{y}_t$.
- Log-prices are defined as: $\mathbf{y}_t \triangleq \log \mathbf{p}_t$.
- This leads to the i.i.d. model when considering log-returns: $\mathbf{x}_t = \mathbf{y}_t - \mathbf{y}_{t-1}$.

Limitations of the I.I.D. Model

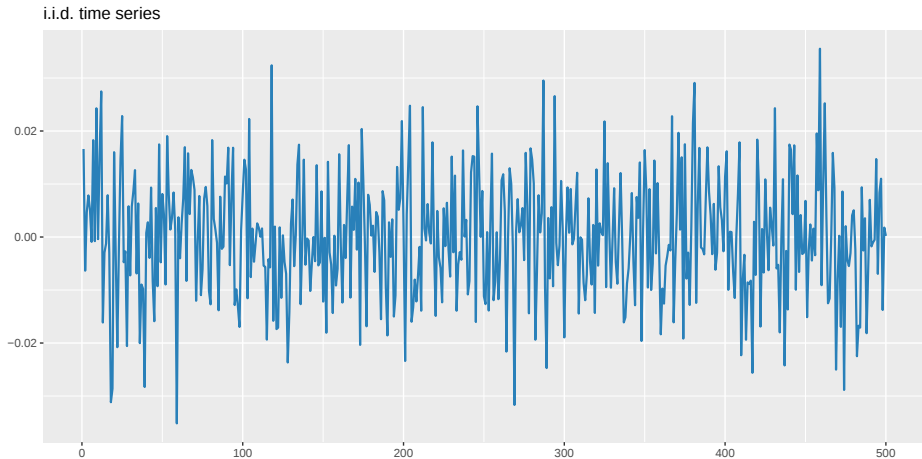
- It ignores temporal structure or dependency in financial data.

Sophisticated Time Series Models

- There is a myriad of different time series models developed over the past seven decades that attempt to capture the temporal structure.
- Recommended textbooks for financial data modeling are: (Meucci 2005; Tsay 2010, 2013; Ruppert and Matteson 2015; Lütkepohl 2007).

I.I.D. Model

Example of a synthetic Gaussian i.i.d. time series:



Outline

- 1 I.I.D. Model
- 2 **Sample Estimators**
- 3 Location Estimators
- 4 Gaussian ML Estimators
- 5 Heavy-Tailed ML Estimators
- 6 Prior Information
 - Shrinkage
 - Factor Models
 - Black-Litterman Model
- 7 Summary

Sample Estimators

Estimating I.I.D. Model Parameters

- The parameters (μ, Σ) are estimated using historical data $\mathbf{x}_1, \dots, \mathbf{x}_T$.
- It utilizes T past observations for estimation.

Sample Mean Estimator:

$$\hat{\mu} = \frac{1}{T} \sum_{t=1}^T \mathbf{x}_t.$$

- It represents the average of past observations.

Sample Covariance Matrix Estimator:

$$\hat{\Sigma} = \frac{1}{T-1} \sum_{t=1}^T (\mathbf{x}_t - \hat{\mu})(\mathbf{x}_t - \hat{\mu})^T.$$

- It measures variability around the sample mean.

Unbiasedness of Estimators

- The sample mean and covariance estimators are unbiased.
- The expected values of $\hat{\mu}$ and $\hat{\Sigma}$ equal the true values μ and Σ .

Bias in Sample Covariance With $1/T$

- Using $1/T$ instead of $1/(T - 1)$ in covariance estimation introduces bias.
- It results in an underestimate: $\mathbb{E}[\hat{\Sigma}] = \left(1 - \frac{1}{T}\right) \Sigma$.

Consistency of Estimators

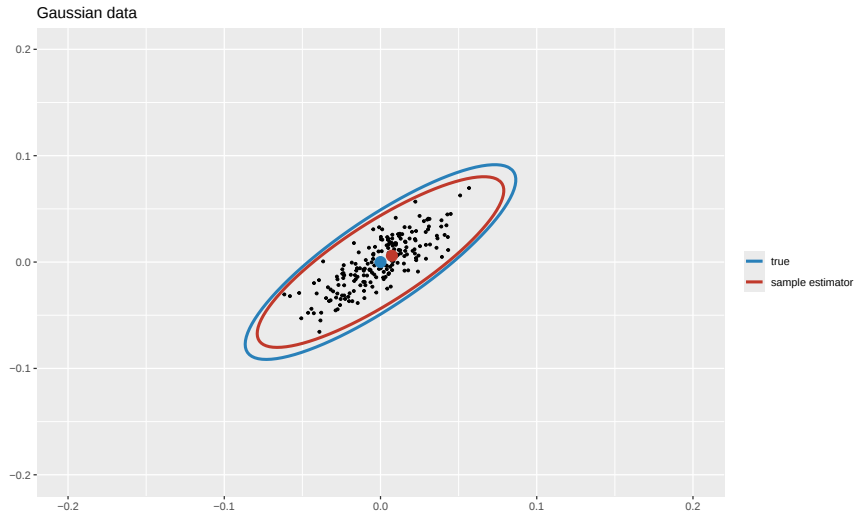
- Both estimators are consistent, converging to true values as $T \rightarrow \infty$.
- This is supported by the law of large numbers.

Estimation Error Reduction

- The estimation error decreases with increasing T .
- This is illustrated next through synthetic Gaussian data with $N = 100$.
- The normalized error approaches zero as the sample size grows.

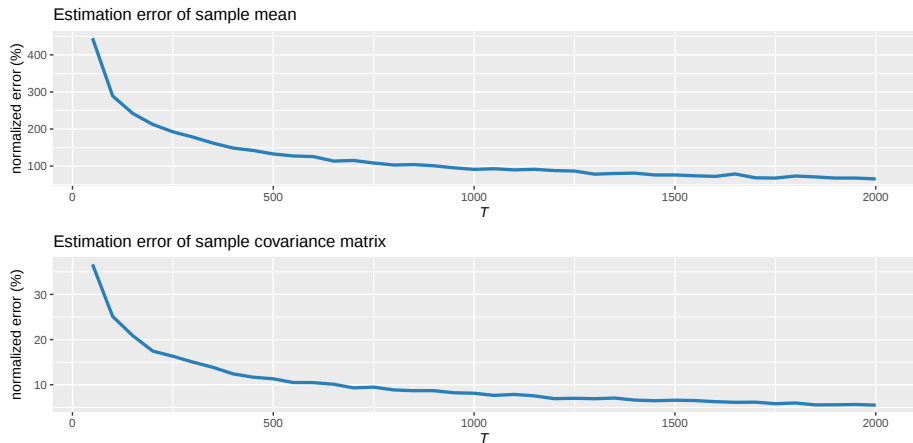
Sample Estimators

Illustration of sample mean and sample covariance matrix:



Sample Estimators

Estimation error of sample estimators versus number of observations (for Gaussian data with $N = 100$):



Sample Estimators

Challenges With Sample Estimators

- Simple and cost-effective but require a large number of observations T for accuracy.
- The sample mean $\hat{\mu}$ is particularly inefficient, leading to noisy estimates.

High Estimation Error With Limited Data

- For $N = 100$ and $T = 500$, the normalized error of $\hat{\mu}$ can exceed 100%.
- In words: the error magnitude is as large as the true value of μ .

Practical Limitations for Large T

- Lack of available historical data: Ideal data span (e.g., 20 years for $N = 500$) often exceeds available records.
- Lack of stationarity: Financial markets evolve, making long-term historical data less relevant.

Implications for Portfolio Optimization

- Limited data leads to noisy estimates of $\hat{\mu}$ and $\hat{\Sigma}$.
- Estimation noise undermines the reliability of portfolio designs.
- It challenges the practical adoption of Markowitz's portfolio theory.

Outline

- 1 I.I.D. Model
- 2 Sample Estimators
- 3 Location Estimators
- 4 Gaussian ML Estimators
- 5 Heavy-Tailed ML Estimators
- 6 Prior Information
 - Shrinkage
 - Factor Models
 - Black-Litterman Model
- 7 Summary

Location Estimators

Interpreting μ in the I.I.D. Model

- μ represents the central location of the distribution of random points.

Estimating the Central Location

- Various methods exist beyond classical sample estimators.
- Sample estimators are sensitive to extreme values and missing data.

Robust Multivariate Location Estimators

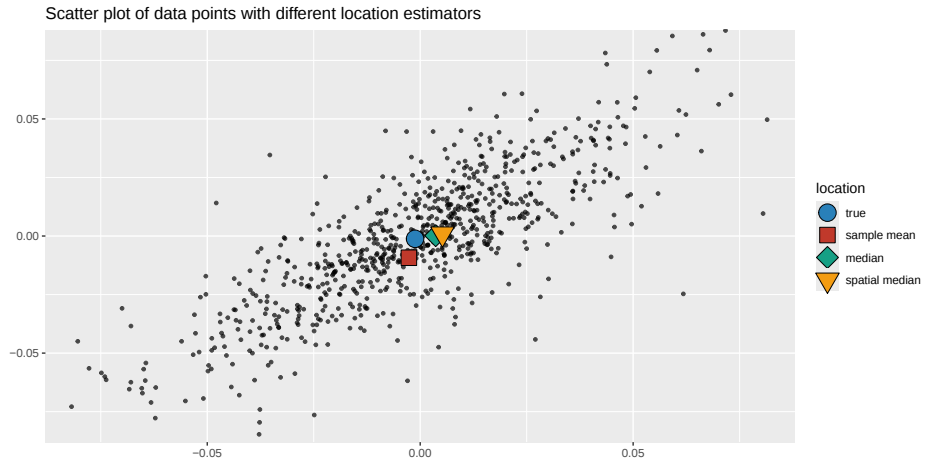
- Developed due to the limitations of sample estimators and least squares.
- Aim to be less affected by outliers and missing values.
- Historical research dates back to the 1960s.

Some Methods to Estimate This Center

- Classical approach: least squares (LS).
- Median estimator.
- Spatial median estimator.

Location Estimators

Illustration of different location estimators:



Least Squares (LS) Estimator

Least Squares (LS) Estimator

- It originates from Gauss's work in 1795 on planetary motions.
- It involves minimizing the squared difference between observed and predicted values.
- It is formulated as:

$$\underset{\mathbf{x}}{\text{minimize}} \quad \|\mathbf{y} - \mathbf{A}\mathbf{x}\|_2^2$$

- The closed-form solution is: $\mathbf{x}^* = (\mathbf{A}^T \mathbf{A})^{-1} \mathbf{A}^T \mathbf{y}$.

Application to the I.I.D. Model

- Estimating μ in the i.i.d. model can be seen as an LS problem:

$$\underset{\mu}{\text{minimize}} \quad \sum_{t=1}^T \|\mathbf{x}_t - \mu\|_2^2$$

- The solution coincides with the sample mean:

$$\hat{\mu} = \frac{1}{T} \sum_{t=1}^T \mathbf{x}_t.$$

Challenges With the Sample Mean

- It lacks robustness against outliers and heavy-tailed distributions.
- The vulnerability to contaminated points affects the reliability of $\hat{\mu}$.

Robust Estimation Needs

- Due to the limitations of LS and the sample mean, alternative robust estimators have been explored.
- They aim to improve reliability and accuracy in the presence of outliers and non-normal distributions.
- The study of robust multivariate location estimators dates back to the 1960s.

Median Estimator

Median as a Robust Estimator

- The median separates the higher half from the lower half of a sample.
- It is considered the “middle” value, providing a typical representation of the data.
- It is unaffected by extreme values, unlike the mean.
- It provides robustness against outliers.
- It represents a more “typical” value of the dataset.

Multivariate Median

- There are multiple ways to extend the median to a multivariate setting.
- Elementwise median is a straightforward extension.

Elementwise Median in the I.I.D. Model

- It can be formulated as the optimization problem:

$$\underset{\mu}{\text{minimize}} \quad \sum_{t=1}^T \|\mathbf{x}_t - \mu\|_1$$

- It uses the ℓ_1 -norm to measure error, offering robustness against outliers.

Spatial Median Estimator

Spatial or Geometric Median

- An extension of the univariate median to multivariate data.
- It is formulated as the optimization problem:

$$\underset{\mu}{\text{minimize}} \quad \sum_{t=1}^T \|\mathbf{x}_t - \mu\|_2$$

- It uses the ℓ_2 -norm (Euclidean distance) as the measure of error.

Characteristics of the Spatial Median

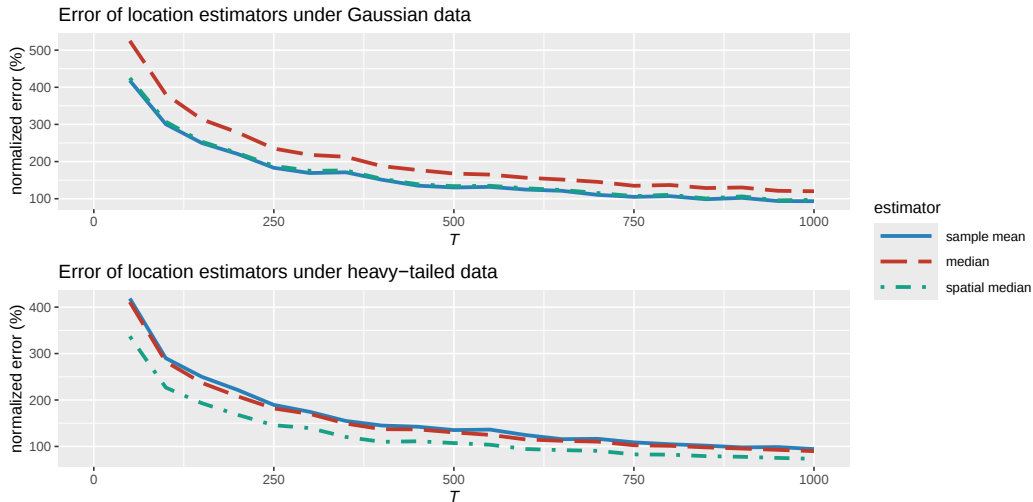
- The estimator for each element is not independent of other elements.
- For $N = 1$, it coincides with the univariate median.

Solving for the Spatial Median

- The problem is a convex second-order cone problem (SOCP).
- It can be solved using SOCP solvers or iterative algorithms.
- Efficient iterative algorithms use the majorization-minimization (MM) method (Sun, Babu, and Palomar 2017).

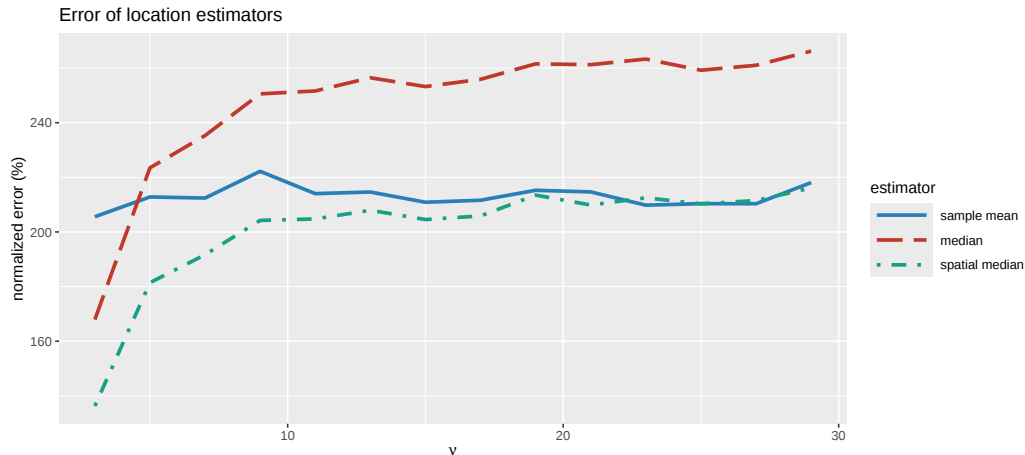
Numerical Experiments

Estimation error of location estimators versus number of observations (with $N = 100$):



Numerical Experiments

Estimation error of location estimators versus degrees of freedom in a t distribution (with $T = 200$ and $N = 100$):



Outline

- 1 I.I.D. Model
- 2 Sample Estimators
- 3 Location Estimators
- 4 Gaussian ML Estimators**
- 5 Heavy-Tailed ML Estimators
- 6 Prior Information
 - Shrinkage
 - Factor Models
 - Black-Litterman Model
- 7 Summary

Maximum Likelihood Estimation (MLE)

- A fundamental technique in estimation theory.
- It involves selecting the parameter vector θ that maximizes the likelihood of observing a given set of data.

Concept of MLE

- It is based on the probability distribution function (pdf) f of the random variable $\mathbf{x}$.
- For T independent observations $\mathbf{x}_1, \dots, \mathbf{x}_T$, the likelihood is $f(\mathbf{x}_1) \times \dots \times f(\mathbf{x}_T)$.
- MLE chooses the parameter θ that maximizes this product for a family of distributions f_θ .

Optimization Problem in MLE

- It is formulated as:

$$\underset{\theta}{\text{maximize}} \quad f_\theta(\mathbf{x}_1) \times \dots \times f_\theta(\mathbf{x}_T).$$

- Equivalently, maximizing the log-likelihood:

$$\underset{\theta}{\text{maximize}} \quad \sum_{t=1}^T \log f_\theta(\mathbf{x}_t).$$

Theoretical Properties of MLE

- Asymptotically unbiased: The MLE's bias diminishes as $T \rightarrow \infty$.
- Asymptotically efficient: It achieves the Cramer-Rao bound, representing the lowest variance for an unbiased estimator.

Practical Considerations

- The effectiveness of MLE's asymptotic properties depends on the size of T .
- Determining how large T needs to be for these properties to hold is crucial in practice.

PDF for I.I.D. Model with Gaussian Residuals

- Assuming residuals follow a multivariate normal distribution:

$$f(\mathbf{x}) = \frac{1}{\sqrt{(2\pi)^N |\Sigma|}} \exp \left(-\frac{1}{2} (\mathbf{x} - \boldsymbol{\mu})^\top \Sigma^{-1} (\mathbf{x} - \boldsymbol{\mu}) \right).$$

- The model parameters are $\boldsymbol{\theta} = (\boldsymbol{\mu}, \Sigma)$.

MLE Formulation for Gaussian I.I.D. Model

$$\underset{\boldsymbol{\mu}, \Sigma}{\text{minimize}} \quad \log \det(\Sigma) + \frac{1}{T} \sum_{t=1}^T (\mathbf{x}_t - \boldsymbol{\mu})^\top \Sigma^{-1} (\mathbf{x}_t - \boldsymbol{\mu}).$$

Deriving MLE for μ and Σ

- Set the gradient with respect to μ and Σ to zero.
- This results in the estimators:

$$\hat{\mu} = \frac{1}{T} \sum_{t=1}^T \mathbf{x}_t$$

$$\hat{\Sigma} = \frac{1}{T} \sum_{t=1}^T (\mathbf{x}_t - \mu)(\mathbf{x}_t - \mu)^T.$$

Comparison With Sample Estimators

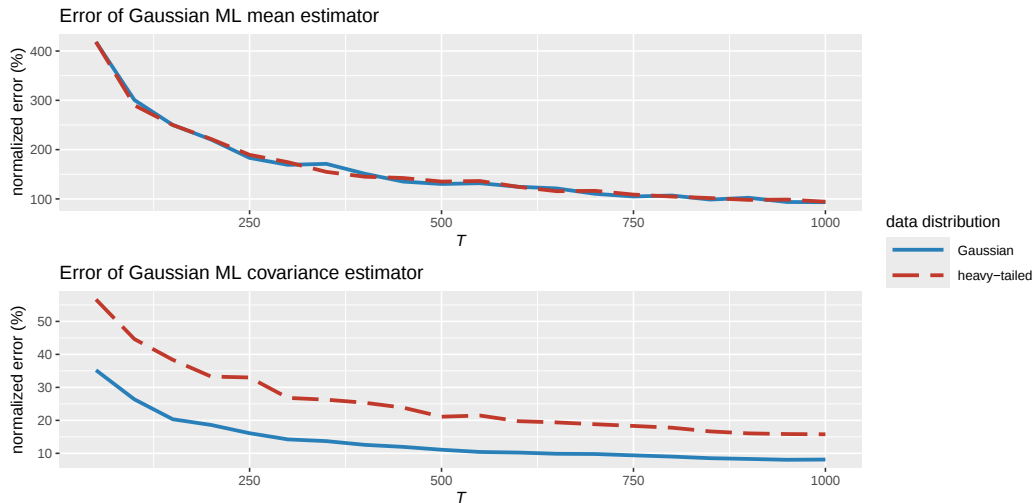
- The MLE coincides with the sample mean and covariance estimators, except for the factor $1/T$ instead of $1/(T-1)$.
- The MLE of the covariance matrix is biased but asymptotically unbiased as $T \rightarrow \infty$.

Implications of MLE

- Sample estimators are optimal for Gaussian-distributed data.
- For non-Gaussian distributions, the optimal ML estimators will be different.

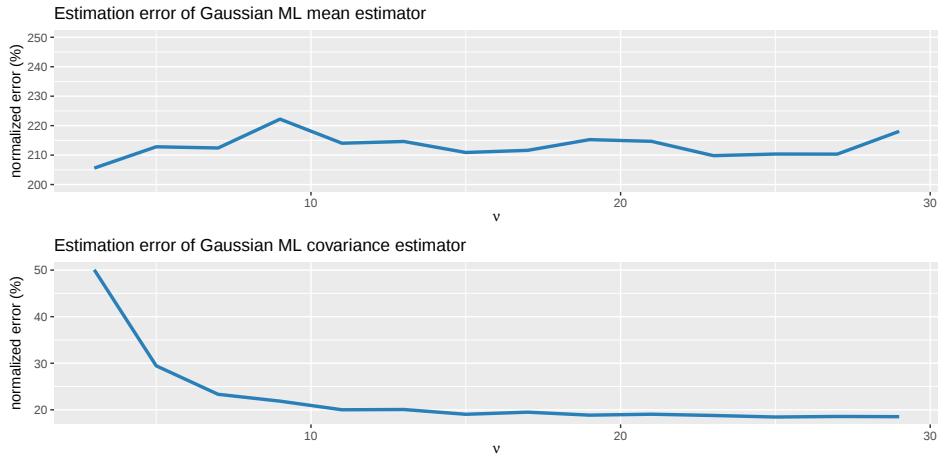
Numerical Experiments

Estimation error of Gaussian ML estimators versus number of observations (with $N = 100$):



Numerical Experiments

Estimation error of Gaussian ML estimators versus degrees of freedom in a t distribution (with $T = 200$ and $N = 100$):



Outline

- 1 I.I.D. Model
- 2 Sample Estimators
- 3 Location Estimators
- 4 Gaussian ML Estimators
- 5 Heavy-Tailed ML Estimators**
- 6 Prior Information
 - Shrinkage
 - Factor Models
 - Black-Litterman Model
- 7 Summary

Gaussian ML Estimators and Heavy-Tailed Distributions

- Optimal for Gaussian-distributed data, but financial data often exhibit heavy tails.
- There is a need to assess the impact of heavy-tailed distributions on these estimators.

Properties of Gaussian ML Estimators

- They coincide with sample estimators for Gaussian data.
- They are unbiased and consistent, which are beneficial characteristics.

Evaluating Estimator Performance With Heavy Tails

- Despite being unbiased and consistent, they may not be the best choice for heavy-tailed data.
- It is important to consider the efficiency and robustness of estimators under such conditions.
- There is potential for improved estimators that better handle the peculiarities of financial data distributions.

The Failure of Gaussian ML Estimators

Impact of Heavy Tails on Estimation

- Heavy tails significantly affect covariance matrix estimation but not mean estimation.
- The error varies with tail heaviness: smaller ν means heavier tails and larger estimation error.

Visualizing Detrimental Effects

- The following scatter plots illustrate the impact of heavy tails and outliers on estimation.
- With $T = 10$ and $N = 2$, the sample covariance matrix is sensitive to these effects.

Consequences for Gaussian ML Estimators

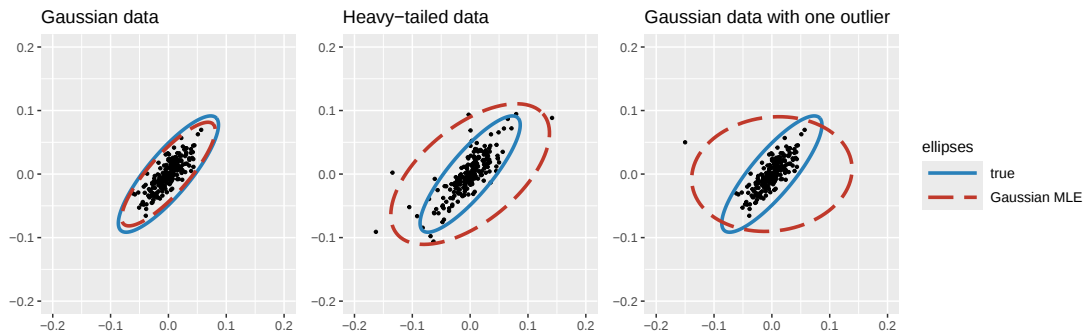
- They perform well for Gaussian data but poorly with outliers or heavy-tailed data.
- Due its simplicity, the sample covariance matrix is widely used by practitioners.

Challenges in Practice

- The prevalence of heavy-tailed distributions in financial data complicates estimation.
- A single outlier can lead to significantly skewed estimations of the covariance matrix.
- There is a need for more robust estimation techniques that can handle outliers and non-Gaussian distributions effectively.

The Failure of Gaussian ML Estimators

Effect of heavy tails and outliers in the Gaussian ML covariance matrix estimator:



Heavy-Tailed ML Estimation

MLE for Heavy-Tailed Distributions

- Gaussian MLE is not optimal for heavy-tailed data, which is common in finance.
- The Student t distribution is used to model heavy tails with parameter ν .

PDF for Multivariate t Distribution

- The probability density function is:

$$f(\mathbf{x}) = \frac{\Gamma((\nu + N)/2)}{\Gamma(\nu/2)\sqrt{(\nu\pi)^N|\Sigma|}} \left(1 + \frac{1}{\nu}(\mathbf{x} - \boldsymbol{\mu})^\top \Sigma^{-1}(\mathbf{x} - \boldsymbol{\mu})\right)^{-(\nu+N)/2}$$

- The parameters are $\boldsymbol{\mu}$ (location), Σ (scatter matrix), and ν (degrees of freedom).

MLE Formulation With t Distribution

- For fixed $\nu = 4$, the MLE problem simplifies to:

$$\underset{\boldsymbol{\mu}, \Sigma}{\text{minimize}} \quad \log \det(\Sigma) + \frac{\nu + N}{T} \sum_{t=1}^T \log \left(1 + \frac{1}{\nu}(\mathbf{x}_t - \boldsymbol{\mu})^\top \Sigma^{-1}(\mathbf{x}_t - \boldsymbol{\mu})\right)$$

Heavy-Tailed ML Estimation

Deriving MLE for μ and Σ

- Set the gradient with respect to μ and Σ to zero.
- This results in the fixed-point equations for μ and Σ :

$$\mu = \frac{\frac{1}{T} \sum_{t=1}^T w_t(\mu, \Sigma) \times \mathbf{x}_t}{\frac{1}{T} \sum_{t=1}^T w_t(\mu, \Sigma)}$$
$$\Sigma = \frac{1}{T} \sum_{t=1}^T w_t(\mu, \Sigma) \times (\mathbf{x}_t - \mu)(\mathbf{x}_t - \mu)^\top$$

where the weights $w_t(\mu, \Sigma)$ are defined as

$$w_t(\mu, \Sigma) = \frac{\nu + N}{\nu + (\mathbf{x}_t - \mu)^\top \Sigma^{-1} (\mathbf{x}_t - \mu)}.$$

Advantages of Heavy-Tailed MLE

- Provides a more robust estimation for datasets with heavy tails.
- Weights observations differently, reducing the influence of outliers.

Heavy-Tailed ML Estimation

MM-based method to solve the heavy-tailed ML fixed-point equations

Initialization:

- Choose initial point (μ^0, Σ^0) .
- Set iteration counter $k \leftarrow 0$.

Repeat (k th iteration):

- 1 Iterate the weighted sample mean and sample covariance matrix as

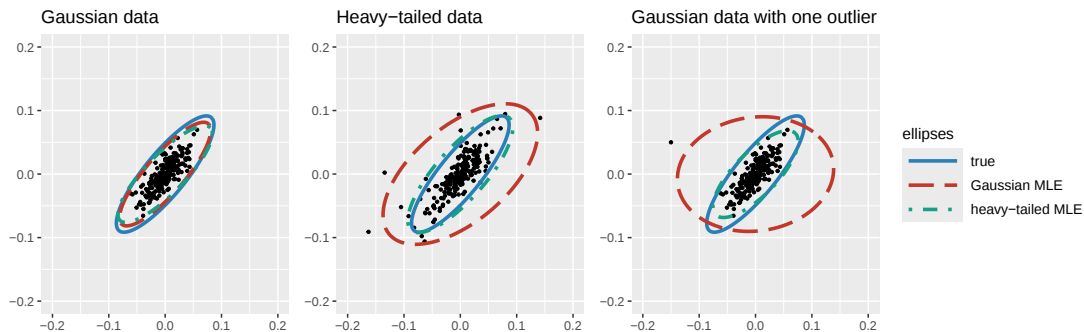
$$\mu^{k+1} = \frac{\frac{1}{T} \sum_{t=1}^T w_t(\mu^k, \Sigma^k) \times \mathbf{x}_t}{\frac{1}{T} \sum_{t=1}^T w_t(\mu^k, \Sigma^k)}$$
$$\Sigma^{k+1} = \frac{1}{T} \sum_{t=1}^T w_t(\mu^{k+1}, \Sigma^k) \times (\mathbf{x}_t - \mu^{k+1})(\mathbf{x}_t - \mu^{k+1})^\top$$

- 2 $k \leftarrow k + 1$

Until: convergence

Numerical Experiments

Effect of heavy tails and outliers in heavy-tailed ML covariance matrix estimator:



Robust Estimators Overview

- Estimators resilient to outliers and deviations from the assumed distribution.
- They are reliable under non-ideal conditions.
- References: (Huber 1964, 2011; Maronna 1976; Maronna, Martin, and Yohai 2006; Wiesel and Zhang 2014; Zoubir et al. 2018, chap. 4; Palomar 2025, chap. 3).

Measuring Robustness

- Influence function: Assesses the impact of deviations on the estimator.
- Breakdown point: Minimum fraction of contaminated data that compromises the estimator, with higher values indicating better robustness.

Sensitivity of Gaussian ML Estimators

- Gaussian-based estimators lack robustness and are highly sensitive to distribution tails.
- A single outlier can hinder the sample mean or covariance (breakdown point of $1/T$).

Robustness of Median and Heavy-Tailed ML Estimators

- The median offers greater robustness with a breakdown point of approximately 0.5.
- Heavy-tailed ML estimators effectively handle deviations, enhancing robustness.

Robust Estimators: M -Estimators*

Introduction to M -Estimators

- The term dates back to the 1960s, introduced by Huber (1964).
- They are a generalization of maximum likelihood estimators.
- They are defined by fixed-point equations for robust estimation of location and scatter.

Fixed-Point Equations for M -Estimators of μ and Σ :

$$\frac{1}{T} \sum_{t=1}^T u_1 \left(\sqrt{(\mathbf{x}_t - \mu)^T \Sigma^{-1} (\mathbf{x}_t - \mu)} \right) (\mathbf{x}_t - \mu) = \mathbf{0}$$

$$\frac{1}{T} \sum_{t=1}^T u_2 \left((\mathbf{x}_t - \mu)^T \Sigma^{-1} (\mathbf{x}_t - \mu) \right) (\mathbf{x}_t - \mu)(\mathbf{x}_t - \mu)^T = \Sigma$$

- The weight functions $u_1(\cdot)$ and $u_2(\cdot)$ must satisfy certain conditions.

Properties and Robustness of M -Estimators

- They are weighted sample mean and covariance matrix.
- They have a bounded influence function for robustness.
- The breakdown point is relatively low, despite robustness.

Robust Estimators: M -Estimators*

Other Robust Estimators (with higher breakdown points)

- Minimum volume ellipsoid (MVE).
- Minimum covariance determinant (MCD).

Gaussian ML Estimators as M -Estimators

- Trivial M -estimators with weight functions:

$$u_1(s) = u_2(s) = 1.$$

Relation to Heavy-Tailed ML Estimators

- M -estimators relate to heavy-tailed ML estimators with the choice:

$$u_1(s) = u_2(s^2) = \frac{\nu + N}{\nu + s^2}.$$

Robust Estimators: Tyler's Estimator

Tyler's Estimator for Scatter Matrix

- It was introduced by Tyler (1987) for heavy-tailed distributions.
- It is the most robust version of an M -estimator.

Elliptical Distribution Assumption

- It assumes a zero-mean elliptical distribution for $\mathbf{x}$.
- If the mean is not zero, it must be estimated and subtracted from the observations.

Normalization and ML Estimation

- Observations are normalized as

$$\mathbf{s}_t = \frac{\mathbf{x}_t}{\|\mathbf{x}_t\|_2}$$

- ML estimation is based on the normalized points, with the pdf (angular distribution) given by

$$f(\mathbf{s}) \propto \frac{1}{\sqrt{|\Sigma|}} \left(\mathbf{s}^T \Sigma^{-1} \mathbf{s} \right)^{-N/2}$$

Robust Estimators: Tyler's Estimator

MLE Formulation

- Given T observations, the MLE is formulated as

$$\underset{\Sigma}{\text{minimize}} \quad \log \det(\Sigma) + \frac{N}{T} \sum_{t=1}^T \log \left(\mathbf{x}_t^T \Sigma^{-1} \mathbf{x}_t \right)$$

- This leads to the fixed-point equation

$$\Sigma = \frac{1}{T} \sum_{t=1}^T w_t(\Sigma) \times \mathbf{x}_t \mathbf{x}_t^T,$$

with weights given by

$$w_t(\Sigma) = \frac{N}{\mathbf{x}_t^T \Sigma^{-1} \mathbf{x}_t}.$$

Robustness and Weights

- The weights enhance robustness by down-weighting outliers.
- A solution exists if $T > N$.

Comparison of Estimators for Mean and Covariance Matrix

- Gaussian MLE
- Tyler's estimator for covariance matrix (paired with spatial median for location)
- Heavy-tailed MLE

Observations

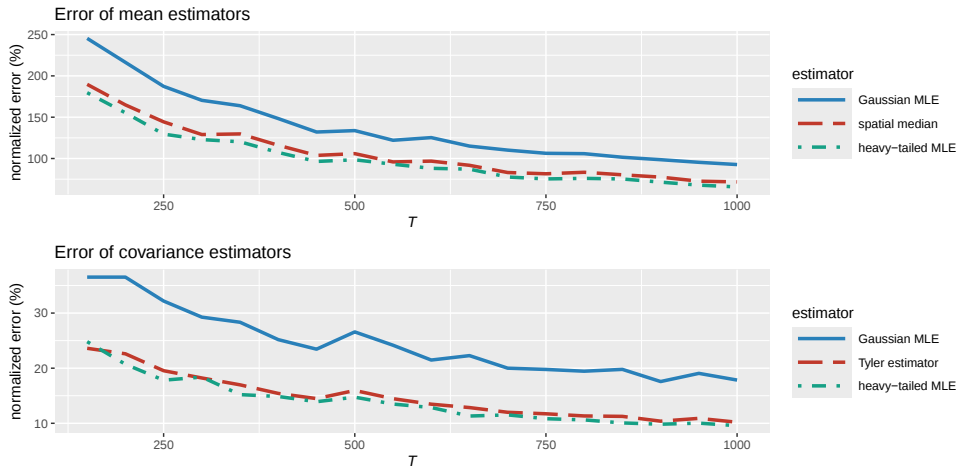
- Gaussian MLE and heavy-tailed MLE perform similarly for Gaussian tails.
- As tails become heavier (smaller ν), heavy-tailed MLE significantly outperforms Gaussian MLE.
- Tyler's estimator (with spatial median for the mean) is also not bad.

Conclusion

- Historical data-based mean vector μ estimation errors can be substantial.
- Financial data's heavy-tailed nature necessitates robust heavy-tailed ML estimators.
- The computational cost of robust estimators is comparable to traditional sample estimators, with convergence in $3 \sim 5$ iterations.
- Practitioners often use factors from data providers for μ estimation or opt for portfolio designs that bypass μ estimation, such as GMVP or RPP.

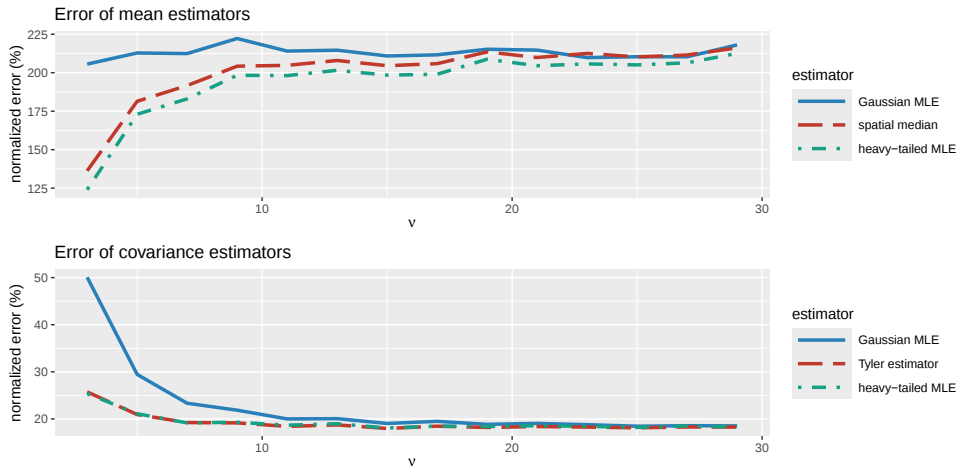
Numerical Experiments

Estimation error of different ML estimators versus number of observations (for t -distributed heavy-tailed data with $\nu = 4$ and $N = 100$):



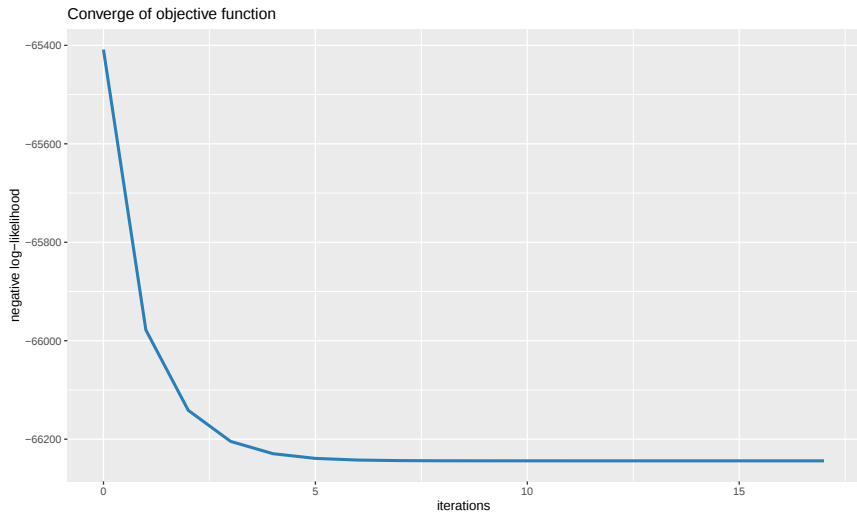
Numerical Experiments

Estimation error of different ML estimators versus degrees of freedom in a t distribution (with $T = 200$ and $N = 100$):



Numerical Experiments

Convergence of robust heavy-tailed ML estimator:



Outline

- 1 I.I.D. Model
- 2 Sample Estimators
- 3 Location Estimators
- 4 Gaussian ML Estimators
- 5 Heavy-Tailed ML Estimators
- 6 Prior Information**
 - Shrinkage
 - Factor Models
 - Black-Litterman Model
- 7 Summary

Challenges With Historical Data

- A limited number of observations (T) can lead to high estimation errors.
- Practical settings often lack sufficient data for accurate parameter estimation.

Improving Estimators With Prior Information

- Researchers and practitioners have developed methods to integrate prior knowledge to enhance estimators.

Three Popular Methods to Incorporate Prior Information

- **Shrinkage** integrates prior knowledge through parameter targets and aims to improve estimation by pulling estimates towards a target.
- **Factor modeling** utilizes structural information about the data, which helps in reducing dimensionality and improving parameter estimation.
- The **Black-Litterman** approach merges historical data with subjective views, balancing empirical data with investor-specific insights.

Outline

- 1 I.I.D. Model
- 2 Sample Estimators
- 3 Location Estimators
- 4 Gaussian ML Estimators
- 5 Heavy-Tailed ML Estimators
- 6 Prior Information**
 - Shrinkage
 - Factor Models
 - Black-Litterman Model
- 7 Summary

Shrinkage

Introduction to Shrinkage

- This is a technique to reduce estimation error by introducing bias.
- It originated with Stein in 1955 and has been popular in finance for covariance matrix shrinkage since the early 2000s.

Bias-Variance Trade-off

- The mean squared error (MSE) of an estimator is the sum of its variance and squared bias.
- Small sample sizes lead to high variance, while larger samples may see bias dominate the error.

Stein's Seminal Contribution (Stein 1955)

- Demonstrated the benefit of introducing bias for overall error reduction.
- Shrinkage involves moving an estimator towards a target value to minimize error.

Shrinkage Estimator Formula

$$\hat{\theta}^{\text{sh}} = (1 - \rho) \hat{\theta} + \rho \theta^{\text{tgt}}$$

where ρ is the shrinkage factor, $\hat{\theta}$ denotes the original estimator, and θ^{tgt} is the target.

Implementation Considerations

- The **choice of target** (θ^{tgt}) represents prior information or market views.
- The **choice of shrinkage factor** (ρ) is critical for balancing the weight of the target in the estimator.

Choosing the Shrinkage Factor

- An **empirical choice** is based on cross-validation.
- An **analytical choice** utilizes advanced mathematical techniques (e.g., random matrix theory).

Application in Finance

- The parameter θ could be the mean vector μ or the covariance matrix Σ .
- Estimators like the sample mean or covariance matrix can be adjusted using shrinkage for better accuracy.

Shrinking the Mean Vector

Sample Mean and Its Properties

- The sample mean $\hat{\mu}$ is an unbiased estimator of the mean vector μ .
- Its distribution is characterized by $\hat{\mu} \sim \mathcal{N}\left(\mu, \frac{1}{T}\Sigma\right)$.
- The mean squared error (MSE) is $\mathbb{E}\left[\|\hat{\mu} - \mu\|^2\right] = \frac{1}{T}\text{Tr}(\Sigma)$.

Stein's Insight on Bias and MSE

- Stein's 1955 paper (Stein 1955) showed that allowing bias can reduce the overall MSE.
- Shrinkage introduces bias towards a target to achieve a lower MSE.

James-Stein Estimator

$$\hat{\mu}^{\text{JS}} = (1 - \rho) \hat{\mu} + \rho \mu^{\text{tgt}}$$

- It improves the MSE over the sample mean for any target μ^{tgt} with a properly chosen ρ :

$$\rho = \frac{(N + 2)}{(N + 2) + T \times (\hat{\mu} - \mu^{\text{tgt}})^{\text{T}} \Sigma^{-1} (\hat{\mu} - \mu^{\text{tgt}})}$$

Shrinking the Mean Vector

Adaptability of ρ

- $\rho \rightarrow 0$ as T increases, favoring the original sample mean.
- $\rho \rightarrow 0$ if the target significantly differs from the sample mean, acting as a safety mechanism.

Choosing the Target μ^{tgt}

- The choice of target is flexible, but MSE improvement depends on the target's informativeness.
- Common choices include:
 - Zero: $\mu^{\text{tgt}} = \mathbf{0}$.
 - Grand mean: $\mu^{\text{tgt}} = \frac{\mathbf{1}^\top \hat{\mu}}{N} \times \mathbf{1}$.
 - Volatility-weighted grand mean: $\mu^{\text{tgt}} = \frac{\mathbf{1}^\top \hat{\Sigma}^{-1} \hat{\mu}}{\mathbf{1}^\top \hat{\Sigma}^{-1} \mathbf{1}} \times \mathbf{1}$.

Shrinking the Covariance Matrix

Shrinkage in Covariance Matrix Estimation

- Introduces bias to reduce estimation error.
- The shrinkage estimator formula is:

$$\hat{\Sigma}^{\text{sh}} = (1 - \rho) \hat{\Sigma} + \rho \Sigma^{\text{tgt}},$$

where Σ^{tgt} is the target and ρ is the shrinkage factor.

Historical Context

- The concept was used in the 1980s in wireless communications as “diagonal loading”.
- It gained popularity in finance in the early 2000s through the work of Ledoit and Wolf (Ledoit and Wolf 2003, 2004).

Common Targets for Covariance Matrix

- Scaled identity matrix: $\Sigma^{\text{tgt}} = \frac{1}{N} \text{Tr}(\hat{\Sigma}) \times I$.
- Diagonal matrix: $\Sigma^{\text{tgt}} = \text{Diag}(\hat{\Sigma})$.
- Equal-correlation matrix, where the off-diagonal elements are equal to the average cross-correlation.

Shrinking the Covariance Matrix

Determining the Shrinkage Factor ρ

- Can be empirical (via cross-validation) or analytical: random matrix theory (RMT).
- Ledoit and Wolf popularized the RMT approach to minimize $\mathbb{E} \left[\|\hat{\Sigma}^{\text{sh}} - \Sigma\|_F^2 \right]$.
- The asymptotic RMT formula for ρ is (Ledoit and Wolf 2003, 2004):

$$\rho = \min \left(1, \frac{\frac{1}{T} \sum_{t=1}^T \|\hat{\Sigma} - \mathbf{x}_t \mathbf{x}_t^T\|_F^2}{\|\hat{\Sigma} - \Sigma^{\text{tgt}}\|_F^2} \right).$$

Extension to Heavy-Tailed Distributions

- Shrinkage factor ρ can be adapted for heavy-tailed distributions, enhancing robustness.

Alternative Error Measures

- Error in terms of the inverse covariance matrix: $\mathbb{E} \left[\|(\hat{\Sigma}^{\text{sh}})^{-1} - \Sigma^{-1}\|_F^2 \right]$.
- Choose ρ to maximize directly the achieved Sharpe ratio.

Nonlinear Shrinkage

- Extends the idea to the eigenvalues of the covariance matrix.
- Requires increased mathematical sophistication.

Shrinkage Estimators and Observation Size

- Estimation error analysis for shrinkage estimators with synthetic Gaussian data.
- Clear improvement in mean vector estimation; modest improvement in covariance matrix estimation.
- Shrinkage benefits decrease as the number of observations (T) increases.

Shrinkage to Zero

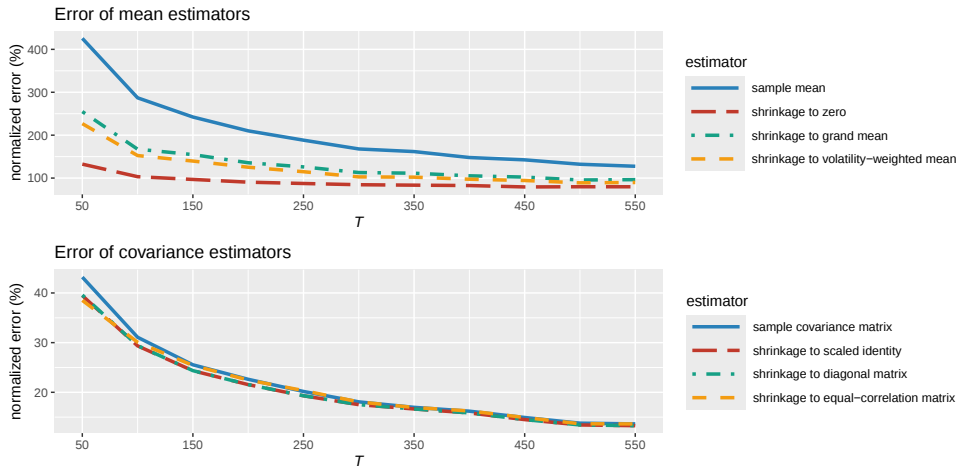
- Shrinkage to zero return vector yields the best results.
- Aligns with the efficient-market hypothesis, suggesting current prices reflect all available information.

Consideration of MSE in Portfolio Optimization

- Numerical results are based on the MSE of estimators.
- MSE may not be the optimal error measure in portfolio optimization contexts.
- The importance of using more appropriate error measures is highlighted.

Numerical Experiments

Estimation error of different shrinkage estimators versus number of observations (for Gaussian data with $N = 100$):



Outline

- 1 I.I.D. Model
- 2 Sample Estimators
- 3 Location Estimators
- 4 Gaussian ML Estimators
- 5 Heavy-Tailed ML Estimators
- 6 Prior Information**
 - Shrinkage
 - **Factor Models**
 - Black-Litterman Model
- 7 Summary

Factor Modeling in Finance

- Incorporates prior information into asset return models.
- Found in numerous finance textbooks.

Single-Factor Model

- The simplest form of factor modeling.
- The equation is:

$$\mathbf{x}_t = \alpha + \beta f_t^{\text{mkt}} + \epsilon_t$$

- α and β represent asset-specific alpha and beta.
- f_t^{mkt} is the market factor and ϵ_t is the zero-mean residual.

Connection to CAPM

- The single-factor model is related to the Capital Asset Pricing Model (CAPM).
- CAPM assumes zero alpha and relates expected excess returns to market beta.

Multi-Factor Modeling

- A generalization of the single-factor model.
- The equation is:

$$\mathbf{x}_t = \alpha + \mathbf{B}\mathbf{f}_t + \epsilon_t$$

- $\mathbf{f}_t$ contains K factors and $\mathbf{B}$ has factor loadings.
- Dynamic factor models include time-dependency in the factors.

Idiosyncratic Component

- The residual term ϵ_t is assumed to have a diagonal covariance matrix Ψ .
- It captures asset-specific noise not explained by the factors.

Factor Models

Mean and Covariance Matrix:

$$\mu = \alpha + B\mu_f$$

$$\Sigma = B\Sigma_f B^T + \Psi$$

- Decomposes the covariance into low-rank and full-rank diagonal components.
- For a single-factor model, the covariance matrix is: $\Sigma = \sigma_f^2 \beta \beta^T + \Psi$.

Parameter Reduction

- Factor models reduce the number of parameters that need to be estimated.
- For example, for $N = 500$ assets and $K = 3$ factors, the number of parameters is reduced from 125,750 to 1,503.

Types:

- **Macroeconomic factor models** use observable economic factors and have unknown loadings.
- **Fundamental factor models** use loadings from asset characteristics and have unknown factors.
- **Statistical factor models** have both unknown factors and loadings.

Macroeconomic Factor Models

Macroeconomic Factor Models Overview

- Allows for the integration of economic indicators into the analysis of asset returns.
- Utilizes observable economic time series as factors (e.g., market index, GDP growth rate, interest rates, inflation rates).
- Factors are often proprietary, derived from complex analyses of various data sources.
- Investment funds may pay high premiums for access to these factors, whereas small investors might use publicly available data.

Parameter Estimation

- With known factors, model parameters (α and $\mathbf{B}$) can be estimated through least squares regression:

$$\underset{\alpha, \mathbf{B}}{\text{minimize}} \quad \sum_{t=1}^T \|\mathbf{x}_t - (\alpha + \mathbf{B}\mathbf{f}_t)\|_2^2,$$

- The mean vector μ and covariance matrix Σ are derived from the estimated parameters as per the factor model equation.

Fundamental Factor Models

Fundamental Factor Models Overview

- Use observable asset characteristics, known as fundamentals, to define factors.
- Common fundamentals include industry classification, market capitalization, and style classification (value, growth).

Industry Approaches

- Fama-French approach:
 - Form portfolios based on asset characteristics to derive factors $\mathbf{f}_t$.
 - Loadings $\mathbf{B}$ are estimated similarly to macroeconomic models.
 - The original model had $K = 3$ factors: firm size, book-to-market values, and excess market return (Fama and French 1992).
 - It was extended to $K = 5$ factors, including profitability and investment patterns (Fama and French 2015).
- Barra risk factor analysis approach:
 - Loadings $\mathbf{B}$ are constructed from asset characteristics.
 - Factors $\mathbf{f}_t$ are estimated via regression, opposite of macroeconomic models.
 - This approach was developed by Barra Inc. in 1975.

Statistical Factor Models Overview

- Both the factors $\mathbf{f}_t$ and the loading matrix $\mathbf{B}$ are unknown.
- They introduce structure to the covariance matrix Σ as a low-rank plus diagonal matrix.

Covariance Matrix Structure

- Σ is decomposed into $\mathbf{B}\Sigma_f\mathbf{B}^T$ (low-rank) and Ψ (diagonal).
- The factors are assumed to be zero-mean and normalized for simplification.

Heuristic Formulation

Approximate the sample covariance matrix $\hat{\Sigma}$ with the desired structure:

$$\underset{\mathbf{B}, \psi}{\text{minimize}} \quad \|\hat{\Sigma} - (\mathbf{B}\mathbf{B}^T + \text{Diag}(\psi))\|_F^2.$$

ML Estimation Under the Gaussian Assumption

- The ML estimation is formulated by imposing the covariance structure:

$$\begin{aligned} & \underset{\alpha, \Sigma, \mathbf{B}, \psi}{\text{minimize}} && \log \det(\Sigma) + \frac{1}{T} \sum_{t=1}^T (\mathbf{x}_t - \alpha)^\top \Sigma^{-1} (\mathbf{x}_t - \alpha) \\ & \text{subject to} && \Sigma = \mathbf{B}\mathbf{B}^\top + \text{Diag}(\psi). \end{aligned}$$

- The nonconvex nature of the problem makes it challenging.
- Iterative algorithms have been developed to address the nonconvex optimization challenges.

Extensions

- Heavy-tailed distributions can be accommodated.
- The structure of financial data, such as nonnegative asset correlation, can be integrated into model formulations.

Principal Component Analysis (PCA)*

PCA in a Nutshell

- It is a statistical technique for handling high-dimensional datasets by focusing on the most significant variance components.
- It facilitates more efficient data analysis and interpretation.

PCA for Dimension Reduction

- PCA identifies the directions of maximum variance in high-dimensional data.
- It simplifies data analysis by reducing dimensions to a lower-dimensional subspace.

PCA Methodology

- It maximizes the variance along a direction $\mathbf{u}$: $\text{Var}(\mathbf{u}^T \mathbf{x}) = \mathbf{u}^T \Sigma \mathbf{u}$.
- The solution is found through eigenvalue decomposition: $\Sigma \approx \mathbf{U}^{(K)} \mathbf{D}^{(K)} \mathbf{U}^{(K)T}$.
- $\mathbf{U}^{(K)}$ contains the first K eigenvectors, and $\mathbf{D}^{(K)}$ has the largest K eigenvalues.

Principal Component Analysis (PCA)*

PCA in Statistical Factor Models

- It approximates the solution to the statistical factor model by performing PCA on the sample covariance matrix.
- As a heuristic, it keeps K principal components and uses a scaled identity matrix for the diagonal component Ψ .
- The approximate solution to the factor model is:

$$\mathbf{B} = \mathbf{U}^{(K)} \text{Diag} \left(\sqrt{\lambda_1 - \kappa}, \dots, \sqrt{\lambda_K - \kappa} \right)$$

$$\Psi = \kappa \mathbf{I}$$

where κ is the average of the $N - K$ smallest eigenvalues.

PCA Estimator for the Covariance Matrix

- The PCA estimator for the covariance matrix is:

$$\hat{\Sigma} = \mathbf{U} \text{Diag} (\lambda_1, \dots, \lambda_K, \kappa, \dots, \kappa) \mathbf{U}^T.$$

- This achieves noise averaging of the smallest eigenvalues, similar to the concept of shrinkage.

Evaluation of Covariance Matrix Estimation Under a Factor Model

- Estimation accuracy depends on how well the data follows the factor model structure.
- Incorrect assumptions about the factor model structure can lead to poorer results than using the sample covariance matrix.

Importance of Model Choice

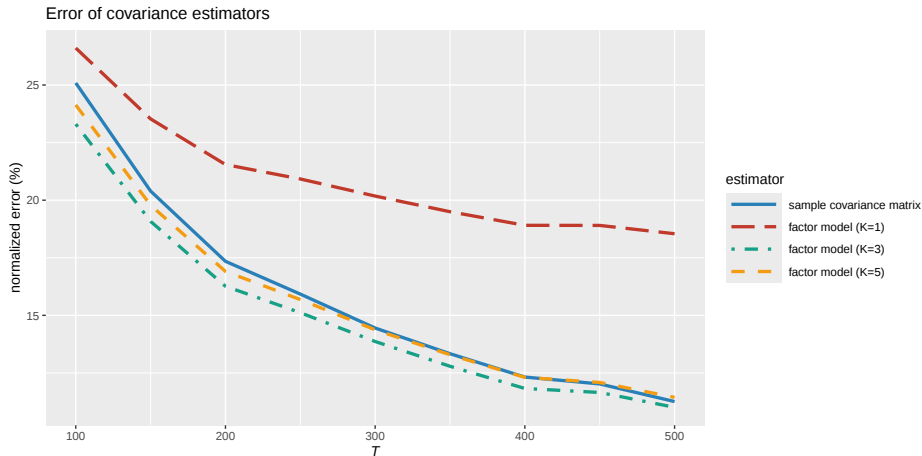
- The decision to use a factor model for covariance matrix estimation must be made cautiously.
- Factor modeling strategies and their implications on trading are elaborated in (Fabozzi, Focardi, and Kolm 2010).

Observations From Synthetic Data Analysis

- An estimation error comparison between factor model-based and sample covariance matrices is shown for synthetic Gaussian data.
- The data complies with a factor model structure having $K = 3$ principal components.
- Accurate estimation with the correct K improves results over the sample covariance matrix.
- Incorrect K values, such as $K = 1$, can significantly worsen estimation accuracy.

Numerical Experiments

Estimation error of covariance matrix under factor modeling versus number of observations
(with $N = 100$):



Outline

- 1 I.I.D. Model
- 2 Sample Estimators
- 3 Location Estimators
- 4 Gaussian ML Estimators
- 5 Heavy-Tailed ML Estimators
- 6 Prior Information**
 - Shrinkage
 - Factor Models
 - Black-Litterman Model
- 7 Summary

Black-Litterman Model

Black-Litterman Model Basics

- It combines historical data with an investor's prior information.
- It is a standard in finance, detailed in textbooks like (Fabozzi, Focardi, and Kolm 2010) and (Meucci 2005).

Components of the Black-Litterman Model

- **Market equilibrium** is an estimate of μ from the market, denoted as $\pi = \hat{\mu}$:

$$\pi = \mu + \epsilon$$

with ϵ being zero-mean and having a covariance of $\tau \Sigma$, where τ is often set to $1/T$.

- **Investor's views** consist of K views on asset returns, expressed as:

$$\mathbf{v} = \mathbf{P}\mu + \mathbf{e}$$

where $\mathbf{v}$ and $\mathbf{P}$ represent the views and their relation to asset returns, and the error term $\mathbf{e}$ is zero-mean with covariance Ω .

Quantitative and Qualitative Views

- Quantitative views specify expected returns and their uncertainties.
- Qualitative views provide directional expectations (e.g., bullish, bearish) without specific return figures.

Example of Quantitative Investor's Views Two independent views on $N = 5$ stocks:

- View 1: stock 1 expected to return 1.5% with a standard deviation of 1%.
- View 2: stock 3 expected to outperform stock 2 by 4% with a 1% standard deviation.

Mathematical Representation

- The views are expressed as:

$$\begin{bmatrix} 1.5\% \\ 4\% \end{bmatrix} = \begin{bmatrix} 1 & 0 & 0 & 0 & 0 \\ 0 & -1 & 1 & 0 & 0 \end{bmatrix} \boldsymbol{\mu} + \boldsymbol{e}$$

- The covariance of the error $\boldsymbol{e}$ is:

$$\boldsymbol{\Omega} = \begin{bmatrix} 1\%^2 & 0 \\ 0 & 1\%^2 \end{bmatrix}.$$

Merging the Market Equilibrium with the Views

Combining Market Equilibrium and Investor's Views

- Various mathematical formulations (least squares, maximum likelihood, Bayesian) yield similar solutions for integrating market equilibrium with an investor's views.

Weighted Least Squares Formulation

- A compact representation is:

$$\mathbf{y} = \mathbf{X}\boldsymbol{\mu} + \mathbf{n},$$

where $\mathbf{y} = \begin{bmatrix} \boldsymbol{\pi} \\ \mathbf{v} \end{bmatrix}$, $\mathbf{X} = \begin{bmatrix} \mathbf{I} \\ \mathbf{P} \end{bmatrix}$, and the noise covariance is $\mathbf{V} = \begin{bmatrix} \tau \boldsymbol{\Sigma} & \mathbf{0} \\ \mathbf{0} & \boldsymbol{\Omega} \end{bmatrix}$.

- The problem is formulated as:

$$\underset{\boldsymbol{\mu}}{\text{minimize}} \quad (\mathbf{y} - \mathbf{X}\boldsymbol{\mu})^T \mathbf{V}^{-1} (\mathbf{y} - \mathbf{X}\boldsymbol{\mu}),$$

- The solution is:

$$\begin{aligned} \boldsymbol{\mu}_{\text{BL}} &= (\mathbf{X}^T \mathbf{V}^{-1} \mathbf{X})^{-1} \mathbf{V}^{-1} \mathbf{y} \\ &= \left((\tau \boldsymbol{\Sigma})^{-1} + \mathbf{P}^T \boldsymbol{\Omega}^{-1} \mathbf{P} \right)^{-1} \left((\tau \boldsymbol{\Sigma})^{-1} \boldsymbol{\pi} + \mathbf{P}^T \boldsymbol{\Omega}^{-1} \mathbf{v} \right) \end{aligned}$$

Merging the Market Equilibrium with the Views*

Original Bayesian Formulation of the Black-Litterman Model

- Returns $\mathbf{x}$ are assumed to follow a normal distribution: $\mathbf{x} \sim \mathcal{N}(\boldsymbol{\mu}, \boldsymbol{\Sigma})$.
- The mean $\boldsymbol{\mu}$ is also modeled as random with a Gaussian distribution, $\boldsymbol{\mu} \sim \mathcal{N}(\boldsymbol{\pi}, \tau\boldsymbol{\Sigma})$, where $\boldsymbol{\pi}$ is the best guess for $\boldsymbol{\mu}$ and $\tau\boldsymbol{\Sigma}$ represents the uncertainty.
- Views are modeled with a Gaussian distribution: $\mathbf{P}\boldsymbol{\mu} \sim \mathcal{N}(\mathbf{v}, \boldsymbol{\Omega})$.
- The posterior distribution of $\boldsymbol{\mu}$ given the views is $\boldsymbol{\mu} \mid \mathbf{v}, \boldsymbol{\Omega} \sim \mathcal{N}(\boldsymbol{\mu}_{\text{BL}}, \boldsymbol{\Sigma}_{\text{BL}})$.
- The posterior mean $\boldsymbol{\mu}_{\text{BL}}$ matches the weighted least squares solution.
- The posterior covariance $\boldsymbol{\Sigma}_{\text{BL}}$ includes the original covariance $\boldsymbol{\Sigma}$ and the covariance of the posterior mean.
- The mean estimator is:

$$\boldsymbol{\mu}_{\text{BL}} = \boldsymbol{\pi} + \tau\boldsymbol{\Sigma}\mathbf{P}^{\text{T}} \left(\tau\mathbf{P}\boldsymbol{\Sigma}\mathbf{P}^{\text{T}} + \boldsymbol{\Omega} \right)^{-1} (\mathbf{v} - \mathbf{P}\boldsymbol{\pi})$$

- The covariance matrix estimator is:

$$\boldsymbol{\Sigma}_{\text{BL}} = (1 + \tau)\boldsymbol{\Sigma} - \tau^2\boldsymbol{\Sigma}\mathbf{P}^{\text{T}} \left(\tau\mathbf{P}\boldsymbol{\Sigma}\mathbf{P}^{\text{T}} + \boldsymbol{\Omega} \right)^{-1} \mathbf{P}\boldsymbol{\Sigma}.$$

Merging the Market Equilibrium with the Views*

Alternative Bayesian Formulation

- A variation introduced by (Meucci 2005) models views on random returns as $\mathbf{v} = \mathbf{P}\mathbf{x} + \mathbf{e}$, which differs from the original formulation where views are on $\boldsymbol{\mu}$.
- This approach leads to a posterior distribution of returns with results akin to the original Black-Litterman model.

Impact of Parameter τ on the Black-Litterman Estimator

- For $\tau = 0$, the market equilibrium is considered completely accurate, leading to $\boldsymbol{\mu}_{BL} = \boldsymbol{\pi}$.
- As $\tau \rightarrow \infty$, the market equilibrium is disregarded, and the investor's views solely influence the outcome, resulting in $\boldsymbol{\mu}_{BL} = \left(\mathbf{P}^T \boldsymbol{\Omega}^{-1} \mathbf{P}\right)^{-1} \left(\mathbf{P}^T \boldsymbol{\Omega}^{-1} \mathbf{v}\right)$.
- For $0 < \tau < \infty$, $\boldsymbol{\mu}_{BL}$ represents a blend of the market equilibrium and investor's views, embodying the principle of shrinkage.

Outline

- 1 I.I.D. Model
- 2 Sample Estimators
- 3 Location Estimators
- 4 Gaussian ML Estimators
- 5 Heavy-Tailed ML Estimators
- 6 Prior Information
 - Shrinkage
 - Factor Models
 - Black-Litterman Model
- 7 Summary

Summary

Countless models for financial data exist in the literature. The i.i.d. model, while rough, is functional and widely used. Key points of the i.i.d. model for financial data include:

- **Sample estimators perform poorly** since Gaussian assumptions don't hold in practice.
- **Robust estimators are necessary** to handle outliers, like spatial median and Tyler estimator.
- **Heavy-tailed estimators suit financial data well** as they are naturally robust. Simple iterative algorithms can compute them.
- **Estimating mean vector from historical data is extremely noisy.** Practitioners use premium data provider factors for regression instead.
- **Prior information should be used when available** via shrinkage, factor modeling, or Black-Litterman information fusion.

- Fabozzi, F. J., S. M. Focardi, and P. N. Kolm. 2010. *Quantitative Equity Investing: Techniques and Strategies*. Wiley.
- Fama, E. F., and K. R. French. 1992. "The Cross-Section of Expected Stock Returns." *The Journal of Finance* 47 (2): 427–65.
- . 2015. "A Five-Factor Asset Pricing Model." *The Journal of Finance* 116 (1): 1–22.
- Huber, P. J. 1964. "Robust Estimation of a Location Parameter." *The Annals of Statistics* 53 (1): 73–101.
- . 2011. *Robust Statistics*. Springer.
- Ledoit, O., and M. Wolf. 2003. "Improved Estimation of the Covariance Matrix of Stock Returns with an Application to Portfolio Selection." *Journal of Empirical Finance* 10 (5): 603–21.

References II

- . 2004. “A Well-Conditioned Estimator for Large-Dimensional Covariance Matrices.” *Journal of Multivariate Analysis* 88 (2): 365–411.
- Lütkepohl, H. 2007. *New Introduction to Multiple Time Series Analysis*. Springer.
- Maronna, R. A. 1976. “Robust M -Estimators of Multivariate Location and Scatter.” *The Annals of Statistics* 4 (1): 51–67.
- Maronna, R. A., D. R. Martin, and V. J. Yohai. 2006. *Robust Statistics: Theory and Methods*. Wiley Series in Probability and Statistics. Wiley.
- Meucci, A. 2005. *Risk and Asset Allocation*. Springer.
- Palomar, Daniel P. 2025. *Portfolio Optimization: Theory and Application*. Cambridge University Press.
- Ruppert, D., and S. D. Matteson. 2015. *Statistics and Data Analysis for Financial Engineering: With R Examples*. 2nd ed. Springer.

References III

- Stein, C. 1955. "Inadmissibility of the Usual Estimator for the Mean of a Multivariate Normal Distribution." *Proceedings of the 3rd Berkeley Symposium on Probability and Statistics* 1: 197–206.
- Sun, Y., P. Babu, and D. P. Palomar. 2017. "Majorization-Minimization Algorithms in Signal Processing, Communications, and Machine Learning." *IEEE Trans. Signal Processing* 65 (3): 794–816.
- Tsay, R. S. 2010. *Analysis of Financial Time Series*. 3rd ed. John Wiley & Sons.
- . 2013. *Multivariate Time Series Analysis: With R and Financial Applications*. John Wiley & Sons.
- Tyler, D. E. 1987. "A Distribution-Free M -Estimator of Multivariate Scatter." *The Annals of Statistics* 15 (1): 234–51.
- Wiesel, A., and T. Zhang. 2014. *Structured Robust Covariance Estimation*. Foundations and Trends in Signal Processing, Now Publishers.

Zoubir, A. M., V. Koivunen, E. Ollila, and M. Muma. 2018. *Robust Statistics for Signal Processing*. Cambridge University Press.