

Robust Gini covariance matrix estimation for portfolio selection based on a factor model

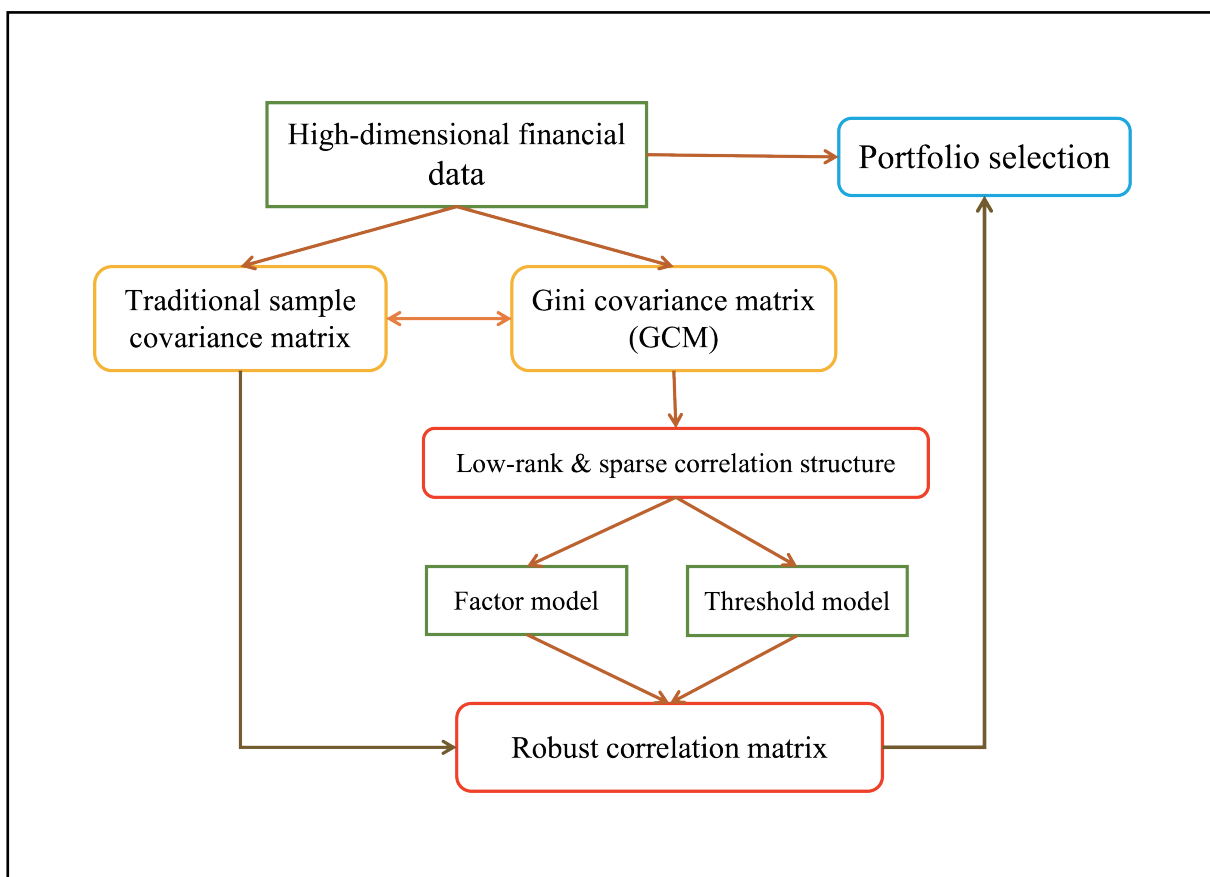
Yongda Zhu, and Lei Shu 

Department of Statistics and Finance, School of Management, University of Science and Technology of China, Hefei 230026, China

 Correspondence: Lei Shu, E-mail: sl2018@ustc.edu.cn

© 2025 The Author(s). This is an open access article under the CC BY-NC-ND 4.0 license (<http://creativecommons.org/licenses/by-nc-nd/4.0/>).

Graphical abstract



A robust portfolio selection method based on a low-rank and sparse structured Gini covariance matrix.

Public summary

- This article proposes a Gini covariance matrix estimation method based on factor models, which demonstrates greater robustness than traditional sample covariance matrices when handling high-dimensional financial data.
- By integrating factor models to capture low-rank structures and employing thresholding rules to achieve the final estimation, the study establishes the consistency of its estimator.
- Through extensive simulation experiments and empirical analysis of S&P 100 stock data, this study illustrates the practical application value of the proposed method in constructing high-performance portfolios.

Robust Gini covariance matrix estimation for portfolio selection based on a factor model

Yongda Zhu, and Lei Shu 

Department of Statistics and Finance, School of Management, University of Science and Technology of China, Hefei 230026, China

 Correspondence: Lei Shu, E-mail: sl2018@ustc.edu.cn

© 2025 The Author(s). This is an open access article under the CC BY-NC-ND 4.0 license (<http://creativecommons.org/licenses/by-nc-nd/4.0/>).



Cite This: *JUSTC*, 2025, 55(8): 0806 (8pp)



Read Online



Supporting Information

Abstract: Portfolio theory has been extensively studied and applied in finance. To determine the optimal portfolio weight under the global minimum variance strategy, it is necessary to estimate both the covariance matrix and its inverse. However, the high dimensionality and heavy-tailed nature of financial data pose significant challenges to this estimation. In this study, we propose a method to estimate the Gini covariance matrix by introducing a low-rank and sparse correlation structure, as an alternative to the traditional sample covariance matrix. Our approach employs a factor model to capture the low-rank structure, combined with thresholding rules to achieve the final estimation. We demonstrate the consistency of our estimators and validate our approach through simulation experiments and empirical portfolio analyses. Simulation results show that our method is highly applicable across a variety of distributional scenarios. Furthermore, empirical portfolio analysis indicates that our method can construct portfolios with superior performance.

Keywords: elliptical distribution; factor model; Gini covariance matrix; portfolio selection

CLC number: F064.1

Document code: A

2020 Mathematics Subject Classification: 91G10

1 Introduction

Portfolio theory, encompassing portfolio selection and asset allocation, has become a cornerstone of modern financial analysis with extensive practical applications. The foundation of modern portfolio theory was established by Markowitz^[1], who pioneered the integration of mathematical statistics into portfolio analysis. This seminal work has subsequently led to the development of numerous portfolio selection methodologies, including the minimum variance approach^[1], equal weight strategy^[2], and maximum diversification technique^[3]. Among these, the minimum variance portfolio, which focuses on minimizing volatility risk, represents a fundamental concept in portfolio theory and serves as the primary focus of our investigation. The minimum variance portfolio can be mathematically formulated via the covariance matrix of stock returns (Σ) as follows:

$$\mathbf{w}_{\text{opt}} = \arg \min_{\mathbf{1}^\top \mathbf{w} = 1} \mathbf{w}^\top \Sigma \mathbf{w} = \frac{\Sigma^{-1} \mathbf{1}}{\mathbf{1}^\top \Sigma^{-1} \mathbf{1}}, \quad (1)$$

where $\mathbf{1}$ represents a vector of ones with a dimension of $N \times 1$ and $\mathbf{w}_{\text{opt}}$ denotes the optimal weights of the stocks. Thus, accurately estimating the covariance matrix and its inverse is vital for deriving optimal portfolio weights as Eq. (1).

Traditional portfolio selection methods typically rely on the sample covariance matrix estimator. However, this approach faces two fundamental limitations in modern financial applications: ① sensitivity to heavy-tailed return distributions and ② dimensionality challenges in high-frequency

trading environments. Our proposed solution using the Gini covariance matrix (GCM) addresses both limitations through robust statistical estimation and structural decomposition. The conventional assumption of Gaussian returns proves inadequate for modeling financial data, where heavy-tailed distributions and extreme events occur frequently^[4,5]. While researchers^[6–8] have adopted the elliptical distribution family (encompassing both normal and Student's t -distributions) as a more flexible framework, practical implementation remains a challenge due to outlier sensitivity. The GCM, introduced by Sang et al.^[9] and Dang et al.^[10], resolves this through its rank-based formulation. This transformation achieves outlier resistance by down-weighting extreme observations through its bounded influence function while maintaining validity across all elliptical distributions.

When the dimensionality N (number of assets) approaches the sample size T (historical returns), traditional covariance estimation becomes ill-posed: The number of parameters ($O(N^2)$) exceeds the number of data points ($O(NT)$); and although there is no singularity theoretically, the singularity occurs when $N > T$. Our solution decomposes the GCM within a factor model structure: “Low-rank” + “Sparse”, where latent common factors (“Low-rank”) capture systematic risk through rank constraints, whereas idiosyncratic components (“Sparse”) admit sparse estimation. This dual structure enables consistent estimation even when $N/T \rightarrow c \in (0, \infty)$, overcoming the curse of dimensionality through three mechanisms: The factor structure reduces the effective parameter space, the sparsity constraint stabilizes the idiosyncratic

estimates, and matrix invertibility is maintained to optimize the portfolio.

Factor models, which are widely utilized in economics and finance^[11–13], describe the latent common component of a low-rank structure. The latent common return x_i in these models is expressed as:

$$x_i = l_{i1}f_1 + \cdots + l_{im}f_m, \quad i = 1, \dots, N, \quad (2)$$

where $f_1, \dots, f_m$ denote the excess returns of m latent factors, with l_{ij} ($i = 1, \dots, N$, $j = 1, \dots, m$) as the unknown factor loadings. The factor model (2) significantly reduces the GCM estimation parameters if a few factors effectively capture cross-sectional risk. Additionally, under the assumption of a sparse structure for the GCM of the idiosyncratic errors, the removal of common components substantially reduces parameter dimensionality. Consequently, this study estimates the GCM and analyzes portfolios in terms of a low-rank factor structure and a sparse idiosyncratic error structure, aligning with previous studies such as Refs. [14–17].

To underscore the practical value of our methods, we initiate an extensive simulation study, evaluating the empirical performance of the methods across a range of scenarios. This simulation study provides valuable insights into the effectiveness of the GCM in contrast to the sample covariance matrix and the sample Kendall τ matrix within a factor structure framework. Next, in our empirical analysis, we turn to S&P 100 stock data spanning from 2020 to 2022 to examine the performance of portfolios constructed via the Gini method compared with the PCA method and the Kendall τ method. Our findings consistently demonstrate that the Gini method yields superior portfolio performance, reinforcing the merits of our approach. We also establish the consistency property of factor estimation in the elliptical distribution framework under some mild assumptions, providing theoretical guarantees of the methodology.

The rest of this paper is organized as follows. Section 2 introduces the portfolio optimization problem on the basis of global minimum variance, for which we propose a multi-factor Gini covariance strategy and present the corresponding algorithm. In Section 3, we outline the theoretical guarantees of the strategy and present the consistency of the latent common component estimates under some basic assumptions. Section 4 presents the results of our simulation study. We then delve into an analysis of a real financial dataset, exploring optimal portfolio allocation and portfolio returns in Section 5. Finally, Section 6 provides our concluding remarks. Supporting information presents several lemmas and proofs of Proposition 1 and Theorem 1.

Notations. For any vector $\mu = (\mu_1, \dots, \mu_n)^T \in \mathbb{R}^n$, let $\|\mu\| = \left(\sum_{i=1}^n \mu_i^2\right)^{1/2}$, $\|\mu\|_\infty = \max_i |\mu_i|$. For a matrix M , let M_{ij} be the entry ij of M , M^T the transpose of M , $\text{diag}(M)$ a vector consisting of the diagonal elements of M if M is a square matrix. Let $\|M\|_\infty = \max \{|M_{ij}|\}$ be the matrix entry-wise maximum value and $\|M\|_F = \left(\sum_{i,j} M_{ij}^2\right)^{1/2}$ be the Frobenius norm of M . Denote $\lambda_j(M)$ as the j th largest eigenvalue of a non-negative definite matrix M .

2 Methodology

2.1 Background on portfolio optimization and elliptical distribution

The minimum variance portfolio is an important concept in modern portfolio theory. Consider the following estimation problem for a global minimum variance portfolio and assume that there are no short-selling restrictions. The formulation of this problem is as follows:

$$\min_w w^T \Sigma w \quad \text{subject to } w^T \mathbf{1} = 1, \quad (3)$$

where $\Sigma \in \mathbb{R}^{N \times N}$ denotes the covariance matrix of the N stock returns. The primary objective of this optimization problem (3) is to determine an asset allocation strategy that minimizes the overall variance, subject to the constraint that the total weight of the portfolio is fixed at 1. This problem can be solved by Eq. (1). In practical applications, it is common to replace the unknown Σ with the estimate $\hat{\Sigma}$, yielding a feasible portfolio.

$$\hat{w}_{\text{opt}} = \frac{\hat{\Sigma}^{-1} \mathbf{1}}{\mathbf{1}^T \hat{\Sigma}^{-1} \mathbf{1}}.$$

Effective covariance matrix estimation is thus crucial in solving the portfolio problem. However, as previously noted, directly estimating the covariance matrix poses a significant challenge. In the following section, we explore a multi-factor strategy in which the GCM is used to address this issue.

Moreover, the elliptical distribution is crucial in portfolio theory, as it is a generalization of the normal distribution, including the heavy-tailed distributions common in finance, such as the t distribution. The elliptical distribution shares several properties with the normal distribution. Specifically, within the elliptical distribution, both the marginal and conditional distributions are elliptical, which is similar to the normal distribution. The definition of the elliptical distribution is as follows:

Definition 1. (Elliptical distribution) Let $\mu \in \mathbb{R}^p$ and $\Sigma \in \mathbb{R}^{p \times p}$ with $\text{rank}(\Sigma) = q \leq p$. A p -dimensional random vector Y is said to have an elliptical distribution, denoted as $Y \sim \text{ED}(\mu, \Sigma, \zeta)$, if it can be represented stochastically as:

$$Y \stackrel{d}{=} \mu + \zeta AU,$$

where U is a random vector uniformly sampled from the unit sphere S^{q-1} in $\mathbb{R}^q$ and $\zeta \geq 0$ is a non-negative random variable independent of U . Additionally, $A \in \mathbb{R}^{p \times q}$ is a deterministic matrix satisfying $AA^T = \Sigma$, where Σ is known as the scatter matrix.

The elliptical distribution allows portfolios to be characterized entirely by location and scale when asset returns follow these distributions. This means that portfolios with the same location and scale parameters have identical return distributions. This property is consistent across portfolio analysis aspects, including mutual fund separation theorems and the capital asset pricing model, as demonstrated in Refs. [6, 8, 18].

2.2 Gini covariance matrix estimation

Denote by T the number of historical returns for each stock

and by N the number of stocks, with observations $\mathbf{y}_t = (y_{1t}, \dots, y_{Nt})^\top$, $1 \leq t \leq T$. As mentioned earlier, traditional sample covariance matrix estimation may not be robust in an elliptical distribution framework, particularly with heavy-tailed effects and outliers. Therefore, we suggest using the Gini covariance matrix (GCM) as a robust alternative to the covariance matrix. Our initial step is to use the observations $\mathbf{Y} = (\mathbf{y}_1, \dots, \mathbf{y}_T)$ to estimate the GCM, which was pioneered by Dang et al.^[10] and is a spatial rank-based measure of covariance among random vectors, as shown in Definition 2.

Definition 2. (Gini covariance matrix (GCM)) Given a random vector $\mathbf{Y} \in \mathbb{R}^p$, the GCM of $\mathbf{Y}$ was defined as

$$\mathbf{G} = 2\mathbb{E}[\mathbf{Y}\mathbf{r}^\top(\mathbf{Y})],$$

where $\mathbf{r}(\mathbf{Y})$ represents the spatial rank function defined as $\mathbf{r}(\mathbf{Y}) = \mathbb{E}_{\tilde{\mathbf{Y}}}[\mathbf{s}(\mathbf{Y} - \tilde{\mathbf{Y}})]$. $\tilde{\mathbf{Y}}$ is an independent copy of $\mathbf{Y}$ and $\mathbf{s}(\mathbf{y}) = \mathbf{y}/\|\mathbf{y}\|$ denotes the spatial sign function. For straightforward algebraic operations, we have

$$\begin{aligned} \mathbf{G} &= 2\mathbb{E}[\mathbf{Y}\mathbb{E}[\mathbf{s}^\top(\mathbf{Y} - \tilde{\mathbf{Y}}) | \mathbf{Y}]] = 2\mathbb{E}[\mathbf{Y}\mathbf{s}^\top(\mathbf{Y} - \tilde{\mathbf{Y}})] = \\ &= \mathbb{E} \frac{(\mathbf{Y} - \tilde{\mathbf{Y}})(\mathbf{Y} - \tilde{\mathbf{Y}})^\top}{\|\mathbf{Y} - \tilde{\mathbf{Y}}\|}. \end{aligned}$$

Examining the GCM definition reveals its incorporation of a rank function, which is a characteristic that enhances its robustness, particularly against outliers. In essence, GCM exhibits greater resilience in the face of individual data points that exhibit substantial deviations from the majority of the data. This attribute renders it a valuable choice for exploring cross-sectional correlations, especially in scenarios characterized by heavy-tailed distributions. The estimation of $\mathbf{G}$ can be achieved through a second-order U-statistic:

$$\tilde{\mathbf{G}}_y = \frac{2}{T(T-1)} \sum_{t < t'} \frac{(\mathbf{y}_t - \mathbf{y}_{t'}) (\mathbf{y}_t - \mathbf{y}_{t'})^\top}{\|\mathbf{y}_t - \mathbf{y}_{t'}\|}. \quad (4)$$

Furthermore, to improve the accuracy of the estimation in the high-dimensional case, we consider a high-dimensional latent component model for the dataset $\{\mathbf{y}_{it}, 1 \leq i \leq N, 1 \leq t \leq T\}$,

$$\mathbf{y}_{it} = \mathbf{x}_{it} + \boldsymbol{\epsilon}_{it}, \quad i \leq N, \quad t \leq T, \quad \text{or} \quad \mathbf{y}_t = \mathbf{x}_t + \boldsymbol{\epsilon}_t \quad \text{in vector form,} \quad (5)$$

where $\mathbf{y}_t = (y_{1t}, \dots, y_{Nt})^\top$ represents the observed stock returns at time t , $\mathbf{x}_t = (x_{1t}, \dots, x_{Nt})^\top$ is the common component of $\mathbf{y}_t$, and $\boldsymbol{\epsilon}_t = (\epsilon_{1t}, \dots, \epsilon_{Nt})^\top$ represents the idiosyncratic errors. We maintain the assumption that $\mathbb{E}(\boldsymbol{\epsilon}_t | \mathbf{x}_t) = 0$, and that $\boldsymbol{\epsilon}_t$ and $\mathbf{x}_t$ are independent.

Under the model (5), we can see that $\mathbf{G}_y$, which represents the correlation of returns between different assets, is calculated as:

$$\begin{aligned} \mathbf{G}_y &= \mathbb{E} \frac{(\mathbf{y}_t - \tilde{\mathbf{y}}_t)(\mathbf{y}_t - \tilde{\mathbf{y}}_t)^\top}{\|\mathbf{y}_t - \tilde{\mathbf{y}}_t\|} = \\ &= \mathbb{E} \frac{(\mathbf{x}_t - \tilde{\mathbf{x}}_t)(\mathbf{x}_t - \tilde{\mathbf{x}}_t)^\top}{\|\mathbf{y}_t - \tilde{\mathbf{y}}_t\|} + \mathbb{E} \frac{(\boldsymbol{\epsilon}_t - \tilde{\boldsymbol{\epsilon}}_t)(\boldsymbol{\epsilon}_t - \tilde{\boldsymbol{\epsilon}}_t)^\top}{\|\mathbf{y}_t - \tilde{\mathbf{y}}_t\|} = \mathbf{G}_x + \mathbf{G}_\epsilon, \end{aligned}$$

where $\mathbf{G}_x$ and $\mathbf{G}_\epsilon$ denote the correlations embedded in the latent common component and the idiosyncratic error, respectively. Within the framework (5), the return GCM is divided

into two separate unknown components. The combined estimates of these components may serve as an alternative to directly estimating the return GCM. However, this division still poses the challenge of dimensional catastrophe in high-dimensional contexts. To mitigate the issue of excessive parameters, we propose imposing a structure on $\mathbf{G}_x$ and $\mathbf{G}_\epsilon$. Specifically, the latent common components are assumed to conform to a low-rank structure, essentially a factor model, as follows:

$$\mathbf{x}_{it} = \mathbf{l}_i^\top \mathbf{f}_t, \quad i \leq N, \quad t \leq T, \quad \text{or} \quad \mathbf{x}_t = \mathbf{L} \mathbf{f}_t, \quad (6)$$

where $\mathbf{f}_t \in \mathbb{R}^m$ (m is the number of the factors) is the latent factor, $\mathbf{l}_i \in \mathbb{R}^m$ represents the factor loading of the i th stock and $\mathbf{L} = (\mathbf{l}_1, \dots, \mathbf{l}_N)^\top$ is the factor loading matrix. Therefore, we will split into two components, estimating $\mathbf{G}_x$ and $\mathbf{G}_\epsilon$ separately to obtain the final $\mathbf{G}_y$ estimate.

Component 1: Estimation of $\mathbf{G}_x$. In practice, we can use the estimation of factors and factor loadings to derive the common components $\hat{\mathbf{x}}_t = \hat{\mathbf{L}} \hat{\mathbf{f}}_t$, which gives us an estimate of $\mathbf{G}_x$ as

$$\hat{\mathbf{G}}_x = \frac{2}{T(T-1)} \sum_{t < t'} \frac{(\hat{\mathbf{x}}_t - \hat{\mathbf{x}}_{t'}) (\hat{\mathbf{x}}_t - \hat{\mathbf{x}}_{t'})^\top}{\|\mathbf{y}_t - \mathbf{y}_{t'}\|}. \quad (7)$$

Due to the complexity of simultaneously estimating $\mathbf{l}_i$ and $\mathbf{f}_t$, as described in Remark 1, we adopt a sequential estimation approach.

Initially, we use an eigen decomposition-based approach to estimate the factor loadings. Following the GCM estimation $\tilde{\mathbf{G}}_y$ in (4), we identify its leading m eigenvectors, $\hat{\boldsymbol{\xi}}_1, \dots, \hat{\boldsymbol{\xi}}_m$, to form the matrix $\hat{\mathbf{F}} = (\hat{\boldsymbol{\xi}}_1, \dots, \hat{\boldsymbol{\xi}}_m)$. We use the scaled matrix $\hat{\mathbf{L}} = (\hat{\mathbf{l}}_1, \dots, \hat{\mathbf{l}}_N)^\top = \sqrt{N} \hat{\mathbf{F}}$ as the estimated factor loading matrix $\mathbf{L}$. Remark 2 elaborates that this approach is theoretically sound within the elliptical distribution framework.

A least squares regression approach is subsequently used to derive the corresponding latent factors. After obtaining $\hat{\mathbf{L}}$, we estimate $\mathbf{f}_t$ by regressing $\mathbf{y}_t$ on $\hat{\mathbf{L}}$. The estimation is formulated as:

$$\hat{\mathbf{f}}_t = \arg \min_{\boldsymbol{\beta}_t \in \mathbb{R}^m} \sum_{i=1}^N (y_{it} - \hat{\mathbf{l}}_i^\top \boldsymbol{\beta}_t)^2, \quad \text{for each } t = 1, \dots, T. \quad (8)$$

Let $\mathbf{F} = (\mathbf{f}_1, \dots, \mathbf{f}_T)^\top$. Then, the explicit solution for $\hat{\mathbf{F}}$ is $\hat{\mathbf{F}} = (\hat{\mathbf{L}}^\top \hat{\mathbf{L}})^{-1} \hat{\mathbf{L}}^\top \mathbf{Y}$. Throughout this study, we treat m as a pre-determined and fixed parameter. For cases where m requires estimation, the methodologies introduced by Refs. [19, 20] provide reliable estimation frameworks.

Remark 1. The factor loading matrix $\mathbf{L}$ and the factor matrix $\mathbf{F}$ remain inherently unobservable, thereby rendering them nonseparately identifiable. In fact, for any arbitrary invertible $m \times m$ matrix $\mathbf{H}$, alternative representations can be easily constructed as $\mathbf{L}^* = \mathbf{L}\mathbf{H}^\top$ and $\mathbf{F}^* = \mathbf{F}\mathbf{H}^{-1}$, while still satisfying the equality $\mathbf{L}^* \mathbf{F}^{*\top} = \mathbf{L} \mathbf{F}^\top$. To ensure identifiability in light of this, we impose specific constraints. Specifically, we stipulate that:

$$\frac{1}{N} \sum_{i=1}^N \sum_{j=1}^N \mathbf{l}_i \mathbf{l}_j^\top = \mathbf{I}_m.$$

This is the rationale behind our utilization of the eigenvectors associated with the first m eigenvalues in the formation of a matrix, which serves as an estimation of the factor loadings.

Remark 2. The procedure for estimating factor loadings utilizing the GCM enjoys theoretical guarantees. The assertion is supported by Theorem 2.1 in Ref. [10], which establishes a fundamental property concerning the eigenvector space. If $\mathbf{Y}$ is elliptically distributed with the first moment and the scatter parameter Σ has the spectral decomposition $\mathbf{V}\mathbf{A}\mathbf{V}^\top$, then we have $\mathbf{G} = \mathbf{V}\mathbf{A}_g\mathbf{V}^\top$ with

$$\mathbf{A}_g = \text{diag}(\lambda_{g,1}, \dots, \lambda_{g,q}) = c \cdot \mathbb{E} \left[\frac{\mathbf{A}^{1/2} \mathbf{U} \mathbf{U}^\top \mathbf{A}^{1/2}}{\sqrt{\mathbf{U}^\top \mathbf{A} \mathbf{U}}} \right].$$

Consequently, we obtain the expressions for $\lambda_{g,i}$ as follows:

$$\lambda_{g,i} = c \cdot \mathbb{E} \left[\frac{\lambda_i u_i^2}{\sum_{j=1}^q \lambda_j u_j^2} \right].$$

Here, $\mathbf{U} = (u_1, \dots, u_q)^\top$ is uniformly distributed on the unit sphere and λ_i represents the eigenvalues of Σ , where $\text{rank}(\Sigma) = q$ and c is a constant associated with the distribution.

The foremost revelation derived from this analysis attests to the alignment between the matrices Σ and $\mathbf{G}$, both of which share a congruent orthogonal matrix $\mathbf{V}$. In essence, this assertion signifies that the eigenvectors are consistently projected onto the identical subspace, a fundamental tenet underpinning the factor model. This establishes the theoretical groundwork for ensuring consistency in both the estimated factors and their corresponding loadings.

Component 2: Estimation of $\mathbf{G}_\epsilon$. To address idiosyncratic errors, we apply a specific regularization structure to the $\mathbf{G}_\epsilon$ matrix, utilizing sparsity constraints, diagonal matrix settings, and related methods. Assuming a sparse structure for $\mathbf{G}_\epsilon$, we employ a thresholding technique to estimate this matrix as follows:

$$\begin{aligned} \widehat{\mathbf{G}}_\epsilon &= \text{Threshold}(\mathbf{M}, \tau), \\ \mathbf{M} &= \frac{2}{T(T-1)} \sum_{t < t'} \frac{(\widehat{\epsilon}_t - \widehat{\epsilon}_{t'}) (\widehat{\epsilon}_t - \widehat{\epsilon}_{t'})^\top}{\|\mathbf{y}_t - \mathbf{y}_{t'}\|}, \end{aligned} \quad (9)$$

where $\text{Threshold}(\mathbf{M}, \tau)$ is a truncation operator for a matrix $\mathbf{M}$ and a constant τ to obtain the sparse structure. It is specifically defined as

$$[\text{Threshold}(\mathbf{M}, \tau)]_{ij} = \begin{cases} M_{ij}, & \text{if } |M_{ij}| \geq \tau; \\ 0, & \text{otherwise.} \end{cases}$$

This definition was put forward by Bickel and Levina^[21] for the sparse estimation of the idiosyncratic error covariance matrix. The approach involves compressing the elements of the matrix with absolute values less than a certain threshold τ to 0.

Finally, using the estimated common component covariance matrix $\widehat{\mathbf{G}}_x$ and the idiosyncratic error covariance matrix $\widehat{\mathbf{G}}_\epsilon$, an estimate of $\mathbf{G}_y$ was derived as $\widehat{\mathbf{G}}_y = \widehat{\mathbf{G}}_x + \widehat{\mathbf{G}}_\epsilon$. Algorithm 1 summarizes the entire process of solving for optimal portfolio weights as described above.

Algorithm 1: Portfolio optimization.

Input: The sample return data $\mathbf{Y} = (\mathbf{y}_1, \dots, \mathbf{y}_T)$, the number of factors m .

Output: The optimal portfolio weight $\widehat{\mathbf{w}}_{\text{opt}}$.

- 1: Calculate the GCM based on (4).
- 2: Calculate the leading m eigenvectors of $\widehat{\mathbf{G}}_y$ and obtain the estimated factor loading matrix $\widehat{\mathbf{L}} = (\widehat{\mathbf{l}}_1, \dots, \widehat{\mathbf{l}}_m)^\top$.
- 3: **for** t in $1:T$ **do**
- 4: Use the least squares regression to obtain $\widehat{\mathbf{f}}_t$ based on (8).
- 5: Estimate the common component as $\widehat{x}_{it} = \widehat{\mathbf{l}}_i^\top \widehat{\mathbf{f}}_t$ and the idiosyncratic error as $\widehat{\epsilon}_{it} = y_{it} - \widehat{x}_{it}$.
- 6: Estimate the correlations embedded in the latent common component and the idiosyncratic error based on (7) and (9).
- 7: Calculate the final return GCM as $\widehat{\mathbf{G}}_y = \widehat{\mathbf{G}}_x + \widehat{\mathbf{G}}_\epsilon$.
- 8: Output the optimal portfolio weight

$$\widehat{\mathbf{w}}_{\text{opt}} = \frac{\widehat{\mathbf{G}}_y^{-1} \mathbf{1}}{\mathbf{1}^\top \widehat{\mathbf{G}}_y^{-1} \mathbf{1}}.$$

3 Theoretical analysis

This section discusses theoretical guarantees within the covariance matrix estimation process, which are primarily set in the elliptical distribution framework. Initially, we assume a low-rank structure for latent common components and apply factor modeling for estimation. During this, we extract the eigenvectors of the sample covariance matrix to estimate factor loadings. We now analyze the theoretical properties of estimating latent common components on the basis of certain fundamental assumptions.

Assumption 1. We assume that

$$\begin{pmatrix} \mathbf{f}_t \\ \epsilon_t \end{pmatrix} = \zeta_t \begin{pmatrix} \mathbf{I}_m & \mathbf{0} \\ \mathbf{0} & \mathbf{A} \end{pmatrix} \frac{\mathbf{g}_t}{\|\mathbf{g}_t\|},$$

where ζ_t 's represent independent samples from ζ , and $\mathbf{g}_t$'s are independent realizations of $\mathbf{g} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_{m+N})$ with fixed m . Additionally, ζ and $\mathbf{g}$ are mutually independent, and the condition $\zeta / \sqrt{N} = O_p(1)$ holds as $N \rightarrow \infty$. This implies that the pairs $(\mathbf{f}_t^\top, \epsilon_t^\top)^\top$ are independent realizations of $\text{ED}(\mathbf{0}, \Sigma_0, \zeta)$. Here we use ED to represent the elliptical distribution. Elliptical distributions are crucial for $t = 1, \dots, T$, with $\Sigma_0 = \begin{pmatrix} \mathbf{I}_m & \mathbf{0} \\ \mathbf{0} & \Sigma_\epsilon \end{pmatrix}$ and $\Sigma_\epsilon = \mathbf{A}\mathbf{A}^\top$. To guarantee the identifiability of the model, we stipulate that $\|\text{diag}(\Sigma_0)\|_\infty = 1$.

Assumption 2. We assume that $\mathbf{L}^\top \mathbf{L} / N \rightarrow \mathbf{V}$ as $N \rightarrow \infty$, with $\mathbf{V}$ being positive definite. Additionally, we posit the existence of positive constants c_1 and c_2 , bounding the eigenvalues of $\mathbf{V}$: $c_1 \leq \lambda_m(\mathbf{V}) < \dots < \lambda_1(\mathbf{V}) \leq c_2$.

Assumption 3. We assume that the eigenvalues of Σ_ϵ are bounded by two positive constants c_3 and c_4 : $c_3 \leq \lambda_{\min}(\Sigma_\epsilon) \leq \lambda_{\max}(\Sigma_\epsilon) \leq c_4$.

Assumption 1 states that the vectors $(\mathbf{f}_t^\top, \epsilon_t^\top)^\top$ are identically and independently distributed (i.i.d.) and adhere to an elliptical distribution. Consequently, the vectors $\mathbf{y}_t$ are also i.i.d. and elliptically distributed, with a scatter matrix given by $\Sigma_y = \mathbf{L}\mathbf{L}^\top + \Sigma_\epsilon$. We assume the scatter matrix for $\mathbf{f}_t$ is the identity matrix; however, this can extend to any symmetric

positive definite matrix Σ_f . Specifically, with $L^{\text{new}} = L\Sigma_f^{1/2}$ and $f_i^{\text{new}} = \Sigma_f^{-1/2}f_i$, the model adapts this generalization. Assumption 2 assumes that $L^T L/N$ converges to a fixed positive definite matrix with bounded eigenvalues. The distinct eigenvalues assumption in Assumption 2 ensures the identifiability of the corresponding eigenvectors. Lastly, Assumption 3 requires that the eigenvalues of Σ_ϵ are bounded within specific limits. This ensures that the idiosyncratic error is negligible compared with common components.

Next, We introduce a key property encompassing the consistency of factor and factor loading estimations, along with a principal theorem for the consistency of common component estimation. These form the theoretical basis for estimating portfolio weights in our multi-factor strategy.

Proposition 1. Within the framework of Assumptions 1–3, there exists a sequence of matrices $\widehat{H}$ such that $\widehat{H}^T V \widehat{H} \xrightarrow{p} I_m$ as $\min\{N, T\} \rightarrow \infty$. Furthermore, it has been established:

(i) The factor loading matrix estimation satisfies

$$\frac{1}{N} \|\widehat{L} - L\widehat{H}\|_F^2 = O_p\left(\frac{1}{N} + \frac{1}{T}\right);$$

(ii) for any $t \leq T$, the factor estimation satisfies

$$\|\widehat{H}f_t - f_t\|^2 = O_p\left(\frac{1}{N} + \frac{1}{T}\right).$$

Proposition 1 establishes the convergence rates for both factor and factor loading estimates without imposing moment constraints on the factors or idiosyncratic error terms, which is made possible by means of the GCM.

Theorem 1. Under the framework of Assumption 1–3 and $\min\{N, T\} \rightarrow \infty$, for any $t \leq T$, we have

$$\frac{1}{N} \|\widehat{x}_t - x_t\|^2 = O_p\left(\frac{1}{N} + \frac{1}{T}\right).$$

Theorem 1 demonstrates that the estimated common component consistently approximates the actual component, thereby substantiating our ability to estimate the covariance matrix via the estimated common component.

4 Simulations

In this section, we conduct simulation studies to examine the empirical performance of the proposed method in finite samples. The foundation of these simulation exercises is rooted in the following model:

$$y_{it} = x_{it} + \sqrt{\theta}e_{it}, \quad x_{it} = \sum_{k=1}^m l_{ik}f_{kt},$$

$$e_{it} = \rho e_{it-1} + (1-\beta)v_{it} + \sum_{l=1}^{i+J} \beta v_{lt}, \quad i = 1, \dots, N, \quad t = 1, \dots, T.$$

Notably, similar data-generating models have been employed in previous studies, including Refs. [17, 19, 20, 22], where they have exhibited favorable behavior in simulation exercises. In this study, we explore dimensions for (N, T) , focusing on (150, 100) and (250, 150), with the number of factors set to $m = 3$. We also examine three parameter configurations and three distinct data-generating scenarios.

Setting 1. The vectors $(f_i^T, v_i^T)^T$ are generated as independent and identically distributed (i.i.d.) samples from a multivariate Gaussian distribution $\mathcal{N}(\mathbf{0}, I_{N+m})$, with the factor loadings $l_{ik} \sim \mathcal{N}(0, 1)$.

Setting 2. The vectors $(f_i^T, v_i^T)^T$ are generated as i.i.d. samples from a multivariate centralized t -distribution $t_2(\mathbf{0}, I_{N+m})$, with factor loadings $l_{ik} \sim \mathcal{N}(0, 1)$.

Setting 3. The factor f_i is generated as i.i.d. samples from a multivariate skewed- t_2 distribution, with the components of v_i independently sampled from $\mathcal{N}(\mathbf{0}, I_N)$. The factor loadings l_{ik} are drawn from a uniform distribution $U(0, 1)$.

Scenario A. The case of pure serial correlation is obtained when the cross-sectional correlation parameter $\beta = 0$. The following parameter values are used: $\beta = 0$, $\rho = 0.5$, $J = 0$, and $\theta = 1$.

Scenario B. The scenario of pure cross-sectional dependence is realized when $\rho = 0$. The following parameter values are employed: $\beta = 0.2$, $\rho = 0$, $J = \max\{N/20, 10\}$, and $\theta = 1$.

Scenario C. This scenario introduces serial and cross-sectional correlation settings. We set the parameters as follows: $\beta = 0.2$, $\rho = 0.5$, $J = \max\{N/20, 10\}$, and $\theta = 1$.

To evaluate the empirical performance of different methods, we employ the following performance indicators. MEAN_FL: Mean deviation in estimating the factor loading matrices; MEAN_FS: Mean deviation in estimating the factor matrices; MED_CC: Median deviation in normalized estimates of the common components. These indicators serve as quantitative measures to assess the effectiveness of the methods under consideration. Specifically,

$$\text{MEAN_FL} = \frac{1}{R} \sum_{i=1}^R \mathcal{D}(\widehat{L}_i, L),$$

$$\text{MEAN_FS} = \frac{1}{R} \sum_{i=1}^R \mathcal{D}(\widehat{F}_i, F),$$

$$\text{MED_CC} = \text{median} \left\{ \frac{\|\widehat{L}_i \widehat{F}_i^T - L F^T\|_F^2}{\|L F^T\|_F^2}, i = 1, \dots, R \right\},$$

where R represents the number of replications, and $\widehat{L}_i$ and $\widehat{F}_i$ correspond to the estimated factor loading and factor matrices from the i th replication, respectively. Given two orthogonal matrices M_1 and M_2 with dimensions $N \times p_1$ and $N \times p_2$, respectively, we define the distance metric $\mathcal{D}(M_1, M_2)$ as:

$$\mathcal{D}(M_1, M_2) = \left(1 - \frac{1}{\max(p_1, p_2)} \text{tr}(M_1 M_1^T M_2 M_2^T)\right)^{1/2}.$$

If M_1 and M_2 are not column orthogonal initially, we can use the Gram–Schmidt orthogonal transformation procedure. The metric $\mathcal{D}(M_1, M_2)$ quantifies the distance between the column spaces of M_1 and M_2 , with values ranging from 0 to 1. A value of 0 indicates that the column spaces of M_1 and M_2 are identical, whereas a value of 1 signifies orthogonality between them.

Table 1 presents a summary of the simulation outcomes of the three scenarios. As described in this paper, “Gini” refers to the factor estimation method using the GCM, “PCA” denotes the method based on the sample covariance matrix, and “KEN” represents the factor estimation method employing

Table 1. Simulation results for Scenarios A, B, and C.

		Method	$(N, T) = (150, 100)$			$(N, T) = (250, 150)$		
			MEAN_FL	MEAN_FS	MED_CC	MEAN_FL	MEAN_FS	MED_CC
Scenario A	Setting 1	Gini	0.12(0.01)	0.10(0.01)	0.02(0.00)	0.10(0.01)	0.07(0.01)	0.01(0.00)
		PCA	0.12(0.01)	0.10(0.01)	0.02(0.00)	0.10(0.01)	0.07(0.01)	0.01(0.00)
		KEN	0.12(0.01)	0.10(0.01)	0.02(0.00)	0.10(0.01)	0.07(0.01)	0.01(0.00)
	Setting 2	Gini	0.24(0.10)	0.23(0.15)	0.05(0.07)	0.19(0.12)	0.18(0.15)	0.03(0.03)
		PCA	0.33(0.21)	0.32(0.22)	0.07(0.36)	0.29(0.21)	0.27(0.22)	0.04(0.26)
		KEN	0.19(0.06)	0.18(0.10)	0.05(0.04)	0.15(0.04)	0.13(0.07)	0.03(0.02)
	Setting 3	Gini	0.12(0.02)	0.08(0.02)	0.01(0.01)	0.09(0.01)	0.06(0.01)	0.01(0.01)
		PCA	0.11(0.02)	0.08(0.02)	0.01(0.01)	0.08(0.01)	0.06(0.01)	0.01(0.01)
		KEN	0.15(0.02)	0.09(0.02)	0.03(0.01)	0.12(0.01)	0.06(0.01)	0.02(0.01)
Scenario B	Setting 1	Gini	0.15(0.02)	0.13(0.02)	0.04(0.01)	0.13(0.01)	0.10(0.01)	0.02(0.01)
		PCA	0.15(0.02)	0.13(0.02)	0.04(0.01)	0.12(0.01)	0.10(0.01)	0.02(0.01)
		KEN	0.15(0.02)	0.13(0.02)	0.04(0.01)	0.13(0.01)	0.10(0.01)	0.03(0.01)
	Setting 2	Gini	0.33(0.13)	0.33(0.16)	0.12(0.21)	0.25(0.12)	0.25(0.15)	0.06(0.08)
		PCA	0.45(0.20)	0.43(0.21)	0.26(0.63)	0.41(0.21)	0.39(0.22)	0.15(0.57)
		KEN	0.27(0.07)	0.27(0.11)	0.10(0.10)	0.21(0.05)	0.19(0.07)	0.06(0.04)
	Setting 3	Gini	0.16(0.05)	0.12(0.04)	0.03(0.02)	0.12(0.03)	0.08(0.02)	0.02(0.01)
		PCA	0.12(0.03)	0.10(0.03)	0.02(0.02)	0.10(0.02)	0.08(0.02)	0.01(0.01)
		KEN	0.23(0.07)	0.14(0.06)	0.05(0.03)	0.17(0.03)	0.09(0.03)	0.03(0.01)
Scenario C	Setting 1	Gini	0.18(0.02)	0.16(0.02)	0.06(0.02)	0.15(0.02)	0.12(0.02)	0.04(0.01)
		PCA	0.18(0.02)	0.15(0.02)	0.06(0.02)	0.15(0.02)	0.12(0.02)	0.04(0.01)
		KEN	0.19(0.02)	0.16(0.02)	0.06(0.02)	0.15(0.02)	0.12(0.02)	0.04(0.01)
	Setting 2	Gini	0.49(0.15)	0.48(0.17)	0.37(0.57)	0.40(0.16)	0.41(0.18)	0.25(0.48)
		PCA	0.54(0.17)	0.52(0.18)	0.48(0.75)	0.51(0.20)	0.50(0.20)	0.42(0.77)
		KEN	0.47(0.13)	0.46(0.15)	0.34(0.48)	0.37(0.12)	0.35(0.14)	0.20(0.21)
	Setting 3	Gini	0.21(0.08)	0.16(0.08)	0.04(0.04)	0.15(0.04)	0.10(0.03)	0.03(0.02)
		PCA	0.17(0.04)	0.13(0.04)	0.03(0.03)	0.13(0.03)	0.09(0.02)	0.02(0.02)
		KEN	0.31(0.11)	0.21(0.11)	0.09(0.08)	0.22(0.06)	0.13(0.05)	0.05(0.03)

The values in parentheses are the standard deviations for MEAN_FL, MEAN_FS, and the interquartile ranges for MED_CC.

the Kendall τ matrix. Several key observations emerge from the analysis of the tables. When both the factor and error terms follow normal distributions, the performances of the three methods are similar, suggesting that satisfactory factor estimates can be obtained via either moment or rank information. In Setting 2, where both the factor and error terms are heavy-tailed, the Gini and KEN methods significantly outperform PCA. This is because, in the heavy-tailed case, the Gini and KEN methods leverage rank information, which helps mitigate the absence of higher-order moments, a limitation of PCA. However, in Setting 3, where the factor and error terms follow t -distributions and normal distributions, respectively, the Gini and PCA methods perform better. This is because, compared with the KEN method, which considers only rank information, the Gini and PCA methods also utilize moment information, allowing them to capture the data structure more efficiently. The performance of all three methods improves as the sample sizes N and T increase, demonstrating their scalability and suitability for application to larger datasets. In

conclusion, the Gini method proposed in this paper is highly applicable across a variety of distributional scenarios.

5 Portfolio analysis

This section provides an analysis of daily stock return data derived from the S&P 100 constituent companies, covering the period from January 1, 2020 to December 31, 2022. The S&P 100 index, comprising 100 large-cap blue-chip stocks traded on the New York Stock Exchange and Nasdaq, represents leading enterprises across diverse industries. These constituent companies not only serve as industry benchmarks but also reflect the developmental status and trends of crucial sectors within the U.S. economy. Given its composition, the performance of the S&P 100 index is strongly correlated with both the U.S. and global economic conditions, making it a focal point for investors, financial institutions, and analysts worldwide. In our methodology, we implement a sliding window technique to sequentially determine optimal investment

weights, ensuring dynamic portfolio adjustments in response to market fluctuations. A standardized time window of approximately two months ($s = 44$ trading days) is used, with out-of-sample testing starting on March 6, 2020. For each investment day t , we calculate $\mathbf{G}_{m,t}$ via returns from the preceding 44 trading days, where m denotes the number of factors. This estimated covariance matrix is then employed to construct the optimal portfolio. To compare the performance of models with different numbers of factors, we set m equal to 3 and 5, mirroring the Fama–French 3-factor and 5-factor models.

Specifically, for each day following March 6, 2020, we recursively compile returns over the past $s = 44$ days, resulting in a data panel of dimensions 100×44 . We then estimate $\widehat{\mathbf{G}}_{m,t}$ for day t via the methodology outlined in Algorithm 1. According to Section 2, we consider the following problem for portfolios that do not allow short selling:

$$\begin{aligned} \widehat{\mathbf{w}}_{\text{opt},t} &= \arg \min_{\mathbf{w}} \mathbf{w}^T \widehat{\mathbf{G}}_{m,t} \mathbf{w} \\ \text{subject to } & \mathbf{w}^T \mathbf{1} = 1, 0 \leq w_i \leq 1 \text{ for all } 1 \leq i \leq N, \end{aligned}$$

where $\mathbf{w} = (w_1, \dots, w_N)^T$ denotes the optimal weights. This

optimization problem incorporates a short-selling constraint into the framework of problem (3). Such a formulation constitutes a standard quadratic programming problem with boundary constraints, which can be solved via existing tools such as the R package **quadprog**. Hence, the portfolio return on day t can be expressed as $\widehat{\mathbf{w}}_{\text{opt},t}^T \mathbf{y}_t$, with $\mathbf{y}_t$ representing the original return for day t . Note that in a later analysis, we comparatively analyze the KEN method and the PCA method, denoting portfolios on the basis of the sample covariance estimates proposed by Fan et al.^[12]

Fig. 1 shows the net value curves for the three methods at $m = 3$ and 5 from 2020 to 2022, excluding transaction costs and liquidity risk. Overall, the Gini method outperformed the PCA and KEN methods across most of the period, maintaining a small but stable advantage. This suggests that the Gini method has a certain superiority in portfolio optimization. However, the cumulative returns experienced several declines, reflecting the significant fluctuations in the U.S. stock market over the past few years. Early to mid-2020: global economic stagnation and market panic triggered by the COVID-19 outbreak. Mid-2022: the Federal Reserve's interest rate hike and rising inflation triggered market concerns.

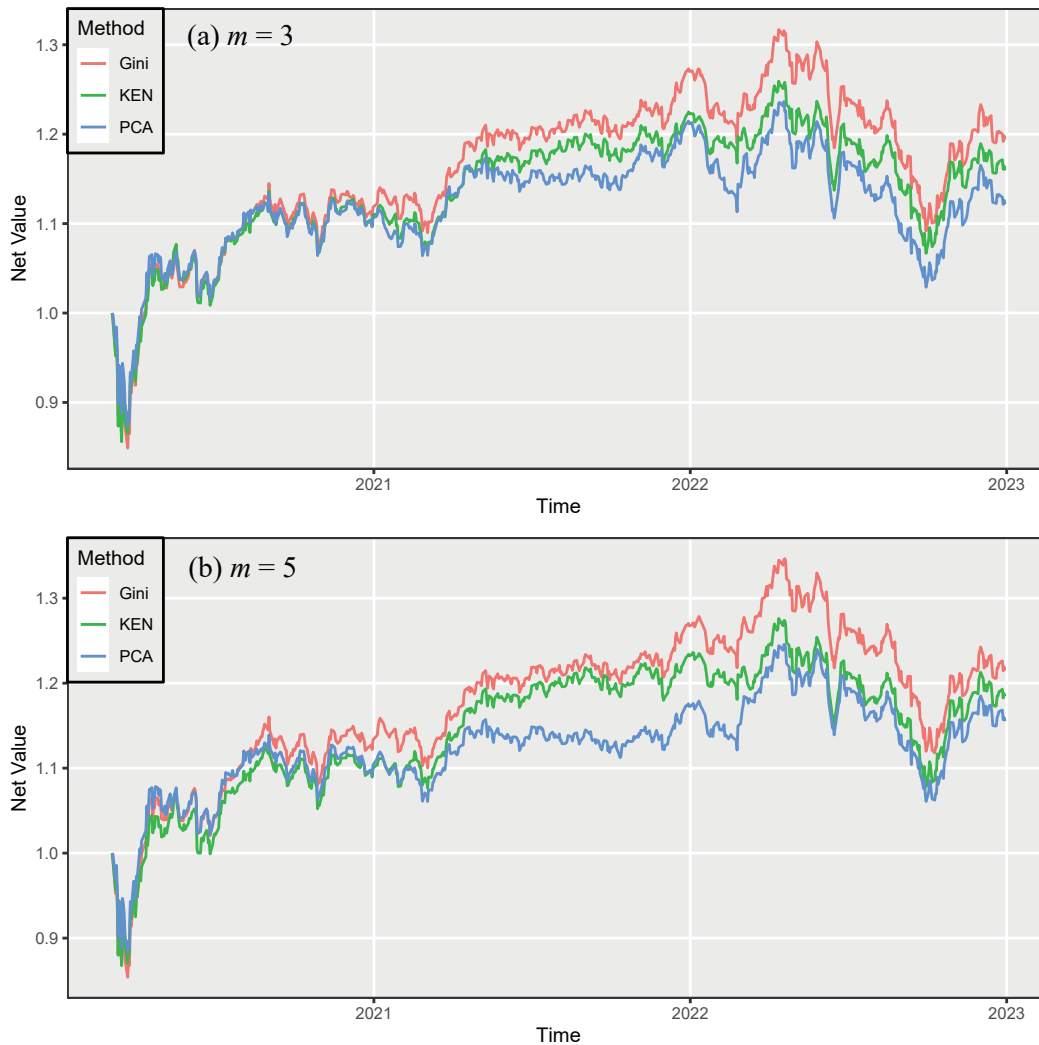


Fig. 1. Portfolio net value curves based on the Gini, PCA, and KEN methods with factors of (a) $m = 3$ and (b) $m = 5$.

As mentioned earlier, the Gini method is more adaptable to various market conditions, which explains its superior performance in the constantly fluctuating financial market.

6 Conclusions

This paper introduces a novel approach to large-dimensional portfolio analysis: the Gini covariance strategy with multiple factors, which is designed to optimize portfolios in financial markets. Specifically, first, we estimate the factor loadings and factors. Next, we obtain the estimates of the GCM for the latent common components and the idiosyncratic errors with a sparse structure. Finally, we combine these estimates to derive the overall estimate. Theoretical demonstrations of the consistency of component estimation, along with simulations and real portfolio studies, reveal that the Gini method is highly applicable across a variety of distributional scenarios. Future research will delve into the theoretical properties of the method, focusing on the estimation error bounds of the return covariance matrix.

Supporting information

The supporting information for this article can be found online at <http://doi.org/10.52396/JUSTC-2025-0003>. It includes several lemmas and proofs of Proposition 1 and Theorem 1.

Acknowledgements

This work was supported by the Postdoctoral Fellowship Program of CPSF (GZC20241651), the National Natural Science Foundation of China (12501391), and the Natural Science Foundation of Anhui Province (2408085QA005).

Conflict of interest

The authors declare that they have no conflict of interest.

Biographies

Yongda Zhu is currently a graduate student at the University of Science and Technology of China. His research mainly focuses on the factor model.

Lei Shu is a postdoctoral fellow at the University of Science and Technology of China (USTC). He received his Ph.D. degree in Statistics from USTC in 2023. His research mainly focuses on high-dimensional statistical inference.

References

- [1] Markowitz H. Portfolio selection. *The Journal of Finance*, **1952**, 7 (1): 77–91.
- [2] DeMiguel V, Garlappi L, Uppal R. Optimal versus naive diversification: How inefficient is the 1/N portfolio strategy? *The Review of Financial Studies*, **2009**, 22 (5): 1915–1953.
- [3] Choueifaty Y, Coignard Y. Toward maximum diversification. *The Journal of Portfolio Management*, **2008**, 35 (1): 40–51.
- [4] Campbell J Y, Lo A W, MacKinlay A C. *The Econometrics of Financial Markets*. Princeton, NJ, USA: Princeton University Press, **1996**.
- [5] McNeil A J, Frey R, Embrechts P. *Quantitative Risk Management: Concepts, Techniques and Tools*. Revised Edition. Princeton, NJ, USA: Princeton University Press, **2015**.
- [6] Bodnar T, Schmid W. A test for the weights of the global minimum variance portfolio in an elliptical model. *Metrika*, **2008**, 67: 127–143.
- [7] Chamberlain G. A characterization of the distributions that imply mean–variance utility functions. *Journal of Economic Theory*, **1983**, 29 (1): 185–201.
- [8] Owen J, Rabinovitch R. On the class of elliptical distributions and their applications to the theory of portfolio choice. *The Journal of Finance*, **1983**, 38 (3): 745–752.
- [9] Sang Y, Dang X, Sang H. Symmetric Gini covariance and correlation. *The Canadian Journal of Statistics*, **2016**, 44 (3): 323–342.
- [10] Dang X, Sang H, Weatherall L. Gini covariance matrix and its affine equivariant version. *Statistical Papers*, **2019**, 60 (3): 641–666.
- [11] Bai J. Inferential theory for factor models of large dimensions. *Econometrica*, **2003**, 71 (1): 135–171.
- [12] Fan J, Fan Y, Lv J. High dimensional covariance matrix estimation using a factor model. *Journal of Econometrics*, **2008**, 147 (1): 186–197.
- [13] Ledoit O, Wolf M. Improved estimation of the covariance matrix of stock returns with an application to portfolio selection. *Journal of Empirical Finance*, **2003**, 10 (5): 603–621.
- [14] Chamberlain G. Funds, factors, and diversification in arbitrage pricing models. *Econometrica*, **1983**, 51 (5): 1305–1324.
- [15] Fan J, Liu H, Wang W. Large covariance estimation through elliptical factor models. *The Annals of Statistics*, **2018**, 46 (4): 1383–1414.
- [16] De Nard G, Ledoit O, Wolf M. Factor models for portfolio selection in large dimensions: The good, the better and the ugly. *Journal of Financial Econometrics*, **2021**, 19 (2): 236–257.
- [17] He Y, Kong X, Yu L, et al. Large-dimensional factor analysis without moment constraints. *Journal of Business & Economic Statistics*, **2022**, 40 (1): 302–312.
- [18] Adcock C, Azzalini A. A selective overview of skew-elliptical and related distributions and of their applications. *Symmetry*, **2020**, 12 (1): 118.
- [19] Onatski A. Determining the number of factors from empirical distribution of eigenvalues. *The Review of Economics and Statistics*, **2010**, 92 (4): 1004–1016.
- [20] Ahn S C, Horenstein A R. Eigenvalue ratio test for the number of factors. *Econometrica*, **2013**, 81 (3): 1203–1227.
- [21] Bickel P J, Levina E. Covariance regularization by thresholding. *The Annals of Statistics*, **2008**, 36 (6): 2577–2604.
- [22] Bai J, Ng S. Determining the number of factors in approximate factor models. *Econometrica*, **2002**, 70 (1): 191–221.