



Applications of Mathematics. – *Extending the Gini index to higher dimensions via whitening processes*, by GENNARO AURICCHIO, PAOLO GIUDICI and GIUSEPPE TOSCANI, communicated on 14 February 2025.

ABSTRACT. – Measuring the degree of inequality expressed by a multivariate statistical distribution is a challenging problem in many fields of science and engineering. In this paper, we propose to extend the well-known univariate Gini coefficient to multivariate distributions by maintaining most of its properties. Our extension is based on applying whitening processes that possess the property of scale stability.

KEYWORDS. – multivariate inequality measures, Gini index, whitening process, Mahalanobis distance.

MATHEMATICS SUBJECT CLASSIFICATION 2020. – 91B82 (primary); 62P20, 94A17 (secondary).

1. INTRODUCTION

The Lorenz curve [22] and the Gini index [13, 14] are still the most important tools to measure the inequality (mutual variability) expressed by a statistical distribution, such as the distribution of income or wealth in a country [8, 11, 12, 27]. However, they are univariate instruments, so that, for a given random n -dimensional vector $\mathbf{X}$ of scalar components X_i , $i = 1, 2, \dots, n$, they are suitable to analyze the variables X_i individually, ignoring the dependence structure they have as components of $\mathbf{X}$. In reason of their importance in economical applications, there had been several efforts to extend the notions of Lorenz curve and Gini index to the multivariate case. The earliest approach, by means of methods of differential geometry, is due to Taguchi [25, 26]. Further suggestions came by Arnold [2], Arnold and Sarabia [3], Gajdos and Weymark [10], Grothe, Kächele, and Schmid [16], Koshevoy and Mosler [19, 20], and Sarabia and Jorda [24]. Unfortunately, as outlined in [3], all these multivariate extensions are essentially determined by elegant mathematical considerations but often lack applicability and interpretability. In addition, these proposals do not possess some of the fundamental properties which are required to inequality measures, properties satisfied by the univariate Gini index.

In a recent paper [15], the possibility of measuring the inequality of multidimensional statistical distributions by resorting to Fourier transform [4, 5, 27] has been investigated.

There, one of the key properties that a multivariate Gini index should possess has been identified in the *scaling invariance* property on components [17], which is essential when trying to recover the value of the inequality index in a multivariate phenomenon composed by different quantities, possibly measured in a different unit of measurement.

By resorting to the Mahalanobis distance [23], in place of the Euclidean distance, in [15], a new version of the multivariate Gini coefficient satisfying the scaling invariance property was proposed and studied. Furthermore, it was shown that, for multivariate Gaussian distributions, the value of the proposed multivariate Gini index is related to the coefficient of variation introduced by Voinov and Nikulin [28], as it does for the univariate case. Owing to the fact that Mahalanobis distance is closely related to the process of *whitening* of a random vector, we will here extend the methods in [15], leading to some generalizations of the Mahalanobis distance. Among them, we will extract one which is particularly well suited to define a new multivariate Gini index, which appears to be easy to handle and interpret.

Whitening is a linear transformation which, given a random n -dimensional vector $\mathbf{X} = (X_1, \dots, X_n)^T$, of mean value $\mathbf{m} = (m_1, \dots, m_n)^T$ and covariance matrix Σ , returns a new random vector $\mathbf{X}^*$ whose entries are orthonormal; that is, the variance of each X_i^* is 1 and the covariance of any X_i^* and X_j^* is null, whenever $i \neq j$. Considering that orthonormality among random variables greatly simplifies multivariate data analysis, both from a computational and a statistical standpoint, whitening is a critically important tool, most often employed in pre-processing. In essence, whitening is a generalization of standardization, a transformation that is carried out by

$$(1.1) \quad \mathbf{X}^* = V^{-1/2} \mathbf{X},$$

where the diagonal matrix $V = \text{diag}(\text{var}(X_1), \text{var}(X_2), \dots, \text{var}(X_n))$ contains the variances of X_i , $i = 1, 2, \dots, n$. This results in a new random vector, namely, $\mathbf{X}^*$, whose components have unitary variance, that is, $\text{var}(X_i^*) = 1$, for every $i = 1, 2, \dots, n$. Notice, however, that this transformation does not remove any correlation that the original entries of the vector $\mathbf{X}$ possess. As we shall see, most whitening processes lack the *scale stable* property, which ensures that the whitened random vector remains unchanged if the components of the original random vector $\mathbf{X}$ are scaled by a positive quantity. This property is essential to obtain a multivariate inequality index that is scale invariant.

The content of this paper is as follows. In Section 2, we describe in detail the whitening process. Furthermore, we discuss the properties which are important in connection with a good definition of a multivariate inequality index. Then, in view of applications, in Section 3, we introduce the new multivariate Gini-type index, by discussing its main properties. Section 4 presents an application of the multivariate index to the study of market economic inequality.

2. THE WHITENING PROCESS

In what follows, we denote with $\mathcal{P}(\mathbb{R}^n)$ the set of n -dimensional probability measures. Since every random vector is identified by its associated probability distribution, with a slight abuse of notation, we use the random vector $\mathbf{X}$ and its associated probability measure μ interchangeably. Moreover, we denote with $\mathcal{P}_{\text{Id}}(\mathbb{R}^n)$ the subset of $\mathcal{P}(\mathbb{R}^n)$ containing the probability measures whose covariance matrix is the identity matrix. In its most generic form, a *whitening process* is a map that, given an n -dimensional random vector $\mathbf{X}$, returns a new n -dimensional random vector whose covariance matrix is the identity, that is, $\mathcal{S} : \mathcal{P}(\mathbb{R}^n) \rightarrow \mathcal{P}_{\text{Id}}(\mathbb{R}^n)$. We say that a whitening process $\mathcal{S}$ is linear if, for any given $\mathbf{X} \in \mathcal{P}(\mathbb{R}^n)$, there exists an $n \times n$ square matrix that depends on $\mathbf{X}$ through its distribution, namely, W_μ , such that

$$(2.1) \quad \mathbf{X}^* = \mathcal{S}(\mathbf{X}) = W_\mu \mathbf{X}.$$

The matrix W_μ is also known as whitening matrix (associated with $\mathcal{S}$) [18]. If the covariance matrix of $\mathbf{X}$, namely, Σ , is invertible, then the whitening matrix in (2.1) must satisfy the identity $W_\mu \Sigma W_\mu^T = I$ and thus $W_\mu (\Sigma W_\mu^T W_\mu) = W_\mu$, which boils down to

$$(2.2) \quad W_\mu^T W_\mu = \Sigma^{-1}.$$

We remark that, given an n -dimensional random vector $\mathbf{X}$ whose covariance matrix is Σ , condition (2.2) does not determine uniquely the linear application that sends $\mathbf{X}$ to a whitened vector. Indeed, the identity (2.2) does not fully identify W_μ but allows for rotational freedom. For example, given a whitening matrix W_μ , any $\tilde{W}_\mu$ of the form

$$\tilde{W}_\mu = Z W_\mu$$

is a whitening matrix as long as Z is an orthogonal matrix, i.e., $Z^T Z = I$, since

$$(\tilde{W}_\mu)^T \tilde{W}_\mu = W_\mu^T Z^T Z W_\mu = W_\mu^T W_\mu = \Sigma^{-1};$$

hence, $\tilde{W}_\mu$ satisfies (2.2) regardless of the choice of Z . Consequentially, there are multiple ways to whiten a random vector, even if we restrict our attention to linear whitening processes [21].

The *Zero-phase Components Analysis* whitening transformation, also known as Mahalanobis whitening [23], is characterized by the matrix

$$(2.3) \quad W_\mu^{\text{Maha}} = \Sigma^{-1/2},$$

where $\Sigma^{-1/2}$ is defined as follows. Since Σ is symmetric and positive-definite, we can decompose it as

$$\Sigma = Z \Theta Z^T,$$

where Z is the eigenmatrix associated with Σ , and Θ is the diagonal matrix whose diagonal contains the positive eigenvalues of Σ . Moreover, Θ is diagonal and its diagonal contains only positive values; we then have that $\Sigma = (\Theta^{1/2} Z)^T (\Theta^{1/2} Z)$, and thus we set

$$(2.4) \quad \Sigma^{-1/2} = Z^T \Theta^{-1/2} Z.$$

EXAMPLE 1. Let $\mathbf{X}$ be an n -dimensional Gaussian random vector, of mean $\mathbf{m}$ and positive-definite $n \times n$ covariance matrix Σ , with associated probability density function

$$f_{\mathbf{X}}(\mathbf{x}) = \frac{1}{(2\pi)^{n/2} (\det \Sigma)^{1/2}} \exp \left\{ -\frac{1}{2} (\mathbf{x} - \mathbf{m})^T \Sigma^{-1} (\mathbf{x} - \mathbf{m}) \right\}.$$

Since the whitening process returns a new n -dimensional Gaussian random vector $\mathbf{X}^*$ of the same dimension n and with unit diagonal *white* covariance, the components of $\mathbf{X}^*$ are uncorrelated. In the Gaussian case, this is equivalent to independence since we can express the joint probability density function as a product of the marginals. Hence, $\mathbf{X}^*$ is an n -dimensional Gaussian random vector, of mean $\mathbf{m}^* = W_{\mu} \mathbf{m}$ and unit diagonal covariance, with independent components. Indeed,

$$\begin{aligned} f_{\mathbf{X}^*}(\mathbf{x}) &= \frac{1}{(2\pi)^{n/2}} \exp \left\{ -\frac{1}{2} (\mathbf{x} - \mathbf{m}^*)^T (\mathbf{x} - \mathbf{m}^*) \right\} \\ &= \prod_{i=1}^n \frac{1}{(2\pi)^{1/2}} \exp \left\{ -\frac{1}{2} (x_i - m_i^*)^2 \right\}. \end{aligned}$$

When it comes to measure the inequality of a probability distribution, not all whitening processes are, however, equal. For example, only a few of the known whitening processes are *Scale Stable*, i.e. such that the random vector $\mathbf{X}^*$ obtained via the whitening process does not change if we multiply one or more entries of the pre-whitening vector $\mathbf{X}$ by a positive constant.

DEFINITION 1 (Scale Stability). A whitening process $\mathcal{S}$ is said to be *Scale Stable* if

$$\mathcal{S}(\mathbf{X}) = \mathcal{S}(Q\mathbf{X}),$$

for every random vector $\mathbf{X}$ and for any diagonal matrix $Q = \text{diag}(q_1, q_2, \dots, q_n)$ whose diagonal elements are positive constants.

Scale Stability is an essential property for any sparsity index whose definition relies on whitened data, as it is connected to the scale invariance of the index itself. For this reason, we now show that linear *Scale Stable* whitening processes always exist and identify two of them.

The Choleski whitening

Choleski whitening is a process based on the Choleski factorization of a positive-definite matrix. In this case, the whitening matrix is

$$(2.5) \quad W_{\mu}^{\text{Chol}} = L^T,$$

where L is the unique lower triangular matrix with positive diagonal values which satisfies (2.2). Owing to the triangular structure of W_{μ}^{Chol} , the whitening process it induces is Scale Stable, as the following result shows.

THEOREM 1. *The Choleski whitening process is Scale Stable.*

PROOF. Let $Q = \text{diag}(q_1, q_2, \dots, q_n)$ be a diagonal matrix such that $q_i > 0$. Given a random vector $\mathbf{X}$, let us denote with Σ its covariance matrix. We then have that the covariance matrix of $\mathbf{Y} = Q\mathbf{X}$ is $Q\Sigma Q$.

Let L be the Choleski factorization of Σ^{-1} , i.e. $\Sigma^{-1} = L^T L$. It is easy to see that the inverse matrix of $Q\Sigma Q$ is $Q^{-1}\Sigma^{-1}Q^{-1}$. Hence, we have

$$(2.6) \quad (Q\Sigma Q)^{-1} = (LQ^{-1})^T (LQ^{-1}).$$

Since L is lower triangular, LQ^{-1} is lower triangular as well since the i -th column of LQ^{-1} is equal to the i -th column of L multiplied by q_i^{-1} . Owing to the uniqueness of the Choleski factorization, we conclude that LQ^{-1} is the Choleski factorization associated with $(Q\Sigma Q)^{-1}$. Finally, notice that

$$LQ^{-1}(Q\mathbf{X}) = L\mathbf{X},$$

which concludes the proof. ■

The correlation whitening

The correlation whitening, also known as *Zero-Components Analysis* (ZCA-cor), employs a whitening matrix which derives from the correlation matrix. In this case, given a random vector $\mathbf{X}$, the whitening matrix is defined as

$$(2.7) \quad W_{\mu}^{\text{ZCA}} = P^{-1/2} V^{-1/2},$$

where P is the correlation matrix of $\mathbf{X}$, and V is the diagonal matrix introduced in (1.1). Again, notice that $P^{-\frac{1}{2}}$ in (2.7) is not defined uniquely; thus, there are multiple ZCA-cor matrices associated with the same $\mathbf{X} \in \mathcal{P}(\mathbb{R}^n)$. Since the correlation matrix P is scale invariant, it is easy to prove that the ZCA-cor whitening process is Scale Stable.

THEOREM 2. *The ZCA-cor whitening process is Scale Stable, regardless of how the square root of P^{-1} is selected.*

PROOF. Without loss of generality, we show this result for a specific square root of P^{-1} as our proof can be generalized to any square root of P^{-1} . Let $Q = \text{diag}(q_1, q_2, \dots, q_n)$ be a diagonal matrix such that $q_i > 0$. Given a random vector $\mathbf{X}$, let us denote with Σ its covariance matrix; then, we have that the covariance matrix of $\mathbf{Y} = Q\mathbf{X}$ is $Q\Sigma Q$.

Let P be the correlation matrix associated with $\mathbf{X}$. Since P is positive-definite and symmetric, we decompose P as $P = O^T \Lambda O$, where Λ is a diagonal matrix containing the eigenvalues of P and O is the matrix containing the eigenvectors associated with P . Notice that $\Sigma = V^{\frac{1}{2}} P V^{\frac{1}{2}}$, where $V = \text{diag}(\text{var}(X_1), \text{var}(X_2), \dots, \text{var}(X_n))$. Moreover, the correlation matrix P is scale invariant; thus, the correlation matrix induced by $Q\mathbf{X}$ is still P . Let us now consider $\Lambda^{-\frac{1}{2}} O^T V^{-\frac{1}{2}}$ and define $\mathbf{X}^* = \Lambda^{-\frac{1}{2}} O^T V^{-\frac{1}{2}} \mathbf{X}$. It is easy to see that the covariance matrix induced by $\mathbf{X}^*$ is the identity matrix. Let us now consider $\mathbf{Y} = Q\mathbf{X}$. The variance of each Y_i is equal to q_i^2 times the variance of X_i , that is, $\text{var}(Y_i) = q_i^2 \text{var}(X_i)$. Since the correlation matrix of $\mathbf{Y}$ is the same as the correlation matrix of $\mathbf{X}$, and since we have that the ZCA-cor whitening matrix induced by $\mathbf{Y}$ is $\Lambda^{-\frac{1}{2}} O^T V_Q^{-\frac{1}{2}}$, where $V_Q = \text{diag}(q_1^2 \text{var}(X_1), q_2^2 \text{var}(X_2), \dots, q_n^2 \text{var}(X_n)) = Q V Q$, we infer that

$$\Lambda^{-\frac{1}{2}} O^T V_Q^{-\frac{1}{2}} \mathbf{Y} = \Lambda^{-\frac{1}{2}} O^T V^{-\frac{1}{2}} Q^{-1} \mathbf{Y} = \Lambda^{-\frac{1}{2}} O^T V^{-\frac{1}{2}} \mathbf{X},$$

which concludes the proof for the ZCA-cor whitening. ■

Counter example

We remark that not all the whitening processes are Scale Stable. Consider for example the Principal Components Analysis (PCA) whitening, a well-known statistical pre-processing method, whose whitening matrix is defined as

$$(2.8) \quad W^{\text{PCA}} = \Theta^{-1/2} Z^T,$$

where Θ is the diagonal matrix containing the eigenvalues of the covariance matrix Σ , and Z is the corresponding (orthogonal) eigenvector matrix (e.g. [9]). The PCA transformation first rotates the variables using the eigenvector matrix of Σ . This results in orthogonal components, but with different variances. To obtain whitened components, the rotated variables are then scaled by the square root of the eigenvalues via the matrix $\Theta^{-1/2}$. Note that, due to the sign ambiguity of the eigenvectors Z , the PCA whitening matrix given by (2.8) is not unique. However, adjusting the column signs in Z such that the elements on the diagonal of Θ are positive, all diagonal elements positive, results in a unique PCA whitening transformation with positive diagonal cross-covariance and cross-correlation.

Notice that this procedure is different from the one defining the ZCA-cor whitening process since the ZCA-cor first scales the entries, then rotates the variables, and then scales the variables according to the eigenvalues of the correlation matrix. Despite its similarities with the ZCA-cor, the PCA whitening is not Scale Stable, as the next example shows.

Let us consider a Gaussian random vector $\mathbf{X}$, and let us set ν its probability measure. Its mean is $m_{\mathbf{X}} = (1, 1)$ and covariance matrix is

$$\Sigma_{\nu} = \begin{pmatrix} 4 & -2 \\ -2 & 3 \end{pmatrix}.$$

Since the eigenvalues of Σ_{ν} are $\theta_1 = 5.56$ and $\theta_2 = 1.44$ and their associated eigenvectors are $v_1 = (0.78, 0.61)$ and $v_2 = (-0.61, 0.78)$, respectively, from (2.8), we infer that

$$W_{\nu}^{\text{PCA}} = \begin{pmatrix} \frac{1}{\sqrt{5.56}} & 0 \\ 0 & \frac{1}{\sqrt{1.44}} \end{pmatrix} \begin{pmatrix} 0.78 & -0.61 \\ 0.61 & 0.78 \end{pmatrix} = \begin{pmatrix} 0.33 & -0.26 \\ 0.51 & 0.66 \end{pmatrix}.$$

Therefore, the random vector $W_{\nu}^{\text{PCA}}\mathbf{X}$ is a Gaussian vector whose covariance is the identity matrix and its mean is $W_{\nu}^{\text{PCA}}m_{\mathbf{X}} = (0.07, 1.17)$.

Let us now consider the Gaussian random vector $\mathbf{Y} = (2X_1, X_2)$, that is, the random vector obtained by multiplying the first entry of $\mathbf{X}$ by two. Let us denote by $\tilde{\nu}$ its probability measure. It is easy to see that $\mathbf{Y}$ is still a Gaussian random vector, whose mean is $m_{\mathbf{Y}} = (2, 1)$ and whose covariance matrix is

$$\Sigma_{\tilde{\nu}} = \begin{pmatrix} 16 & -4 \\ -4 & 3 \end{pmatrix}.$$

In this case, the eigenvalues of $\Sigma_{\tilde{\nu}}$ are $\eta_1 = 17.13$ and $\eta_2 = 1.87$ and their associated eigenvectors are $w_1 = (0.96, 0.27)$ and $w_2 = (-0.27, 0.96)$, respectively. In particular, the PCA whitening matrix associated with $\mathbf{Y}$ is

$$W_{\tilde{\nu}}^{\text{PCA}} = \begin{pmatrix} 0.23 & -0.06 \\ 0.20 & 0.71 \end{pmatrix}.$$

Therefore, the random vector $W_{\tilde{\nu}}^{\text{PCA}}\mathbf{Y}$ follows a Gaussian distribution whose covariance matrix is the identity and whose mean is $W_{\tilde{\nu}}^{\text{PCA}}m_{\mathbf{Y}} = (0.40, 1.11)$. We then conclude that $W_{\nu}^{\text{PCA}}\mathbf{X} \neq W_{\tilde{\nu}}^{\text{PCA}}\mathbf{Y}$ so that the PCA whitening is not Scale Stable.

The p -Mahalanobis metrics

Given an n -dimensional random vector $\mathbf{X}$, we denote with $\mathbf{m} = (m_1, m_2, \dots, m_n)^T$ its mean and with Σ_{μ} its positive-definite $n \times n$ covariance matrix. The Mahalanobis

metric is then defined as

$$(2.9) \quad \mathbf{m}_2(\mathbf{X}) = \sqrt{\mathbf{m}^T \Sigma_\mu^{-1} \mathbf{m}} = \sqrt{(W_\mu \mathbf{m})^T (W_\mu \mathbf{m})} = \|W_\mu \mathbf{m}\|_2,$$

where W_μ is any whitening matrix. Since for any whitening matrix W_μ we have that

$$\mathbf{m}^T W_\mu^T W_\mu \mathbf{m} = \mathbf{m}^T \Sigma_\mu^{-1} \mathbf{m},$$

the Euclidean norm of $W_\mu \mathbf{m}$ does not depend on W_μ . This property, however, is lost if we consider other norms of the vector $W_\mu \mathbf{m}$. In this case, the choice of the whitening matrix W_μ affects the value of the norm, and it may not be scale invariant. To overcome this issue, it suffices to consider a Scale Stable whitening.

DEFINITION 2 (The l_p Mahalanobis norm). Let $\mathcal{S}$ be a *Scale Stable* whitening process and let $\mathbf{X}$ be a random vector. For any $p \geq 1$, we define the l_p -Mahalanobis norm induced by $\mathcal{S}$ of $\mathbf{X}$ as follows:

$$(2.10) \quad M_p(\mathbf{X}) = \|\mathbb{E}[\mathcal{S}(\mathbf{X})]\|_p,$$

where $\mathbb{E}[\mathcal{S}(\mathbf{X})]$ is the vector containing the mean of $\mathcal{S}(\mathbf{X})$.

REMARK 1. From a constructive viewpoint, note that, owing to Theorems 1 and 2, both Choleski and correlation whitening processes allow us to define a generalized Mahalanobis norm that is scale invariant, i.e.

$$M_p(\mathbf{X}) = M_p(Q\mathbf{X}),$$

for every diagonal matrix Q whose diagonal contains strictly positive values.

3. A NEW MULTIVARIATE GINI-TYPE INDEX

In this section, we show how to define a scaling invariant Gini index for multivariate distributions using the l_p -Mahalanobis norm introduced in Definition 2. For the sake of simplicity, from now on we consider only the ZCA-cor whitening.

DEFINITION 3. For any $\mathbf{X}$ random vector, let W_μ^{ZCA} be the ZCA-cor whitening process associated with $\mathbf{X}$. Then, we define

$$(3.1) \quad G_p(\mathbf{X}) = \frac{1}{2M_p(\mathbf{X})} \int_{\mathbb{R}^n \times \mathbb{R}^n} \|W_\mu^{\text{ZCA}}(\mathbf{x} - \mathbf{y})\|_p \mu(d\mathbf{x}) \mu(d\mathbf{y}),$$

where, for every $p \geq 1$, $M_p(\mathbf{X})$ is the Mahalanobis metric (see Definition 2).

Depending on the whitening process, i.e. on the function $\mathbf{X} \rightarrow W_\mu^{\text{ZCA}}$ that maps a random vector into its whitening matrix, the index G_p satisfies the defining properties of an inequality index. For example, if the components of $\mathbf{X}$ are non-negative, the whitened vector $W_\mu^{\text{ZCA}}\mathbf{X}$ might not be non-negative. Consequentially, the multivariate Gini indices preserve the well-known properties that the one-dimensional Gini index possesses for positive quantities if and only if there exists a whitening matrix whose entries are positive, as it maps non-negative vectors into non-negative vectors.

LEMMA 1. *Given $\mathbf{X}$ a non-negative random vector whose correlation matrix P is invertible, let P^{-1} denote the inverse of P . Moreover, let us decompose P^{-1} as $P^{-1} = O\Lambda O^T$ where O is an orthogonal matrix and Λ is the diagonal matrix containing the eigenvalues of P^{-1} . Finally, let V be the diagonal matrix containing the variances of $\mathbf{X}$ and R an orthogonal matrix. Then, the random vector*

$$(3.2) \quad \mathbf{X}^* := W_\mu^{\text{ZCA}}\mathbf{X} := R O \Lambda^{-\frac{1}{2}} O^T V^{-\frac{1}{2}} \mathbf{X}$$

is non-negative whenever R is such that $(R O \Lambda^{-\frac{1}{2}} O^T)_{i,j} \geq 0$ for every $i, j = 1, \dots, n$.

PROOF. It follows from the fact that W_μ^{ZCA} has only positive entries; thus, $W_\mu^{\text{ZCA}}\mathbf{X}$ is non-negative whenever $\mathbf{X}$ is non-negative. ■

For the sake of simplicity, given a random vector $\mathbf{X}$, from now on we consider the ZCA correlation whitening process induced by $R = \text{Id}$, thus, $P^{-\frac{1}{2}} = O \Lambda^{-\frac{1}{2}} O^T$, and set

$$(3.3) \quad W_\mu^{\text{ZCA}} = O \Lambda^{-\frac{1}{2}} O^T V^{-\frac{1}{2}},$$

so that if $\mathbf{X}$ is a random vector whose covariance matrix is the identity, then $W_\mu^{\text{ZCA}} = \text{Id}$. We will then consider the family of multidimensional Gini indices induced by W_μ^{ZCA} , so that

$$G_p(\mathbf{X}) = \frac{1}{2M_p(\mathbf{X})} \int_{\mathbb{R}^n \times \mathbb{R}^n} \|W_\mu^{\text{ZCA}}(\mathbf{x} - \mathbf{y})\|_p \mu(d\mathbf{x}) \mu(d\mathbf{y}).$$

When $p = 1$, the latter identity becomes particularly interesting as, in this case, we can express G_1 as a convex combination of the one-dimensional Gini indices, which we denote with G , of the components of the vector $\mathbf{X}^* = W_\mu^{\text{ZCA}}\mathbf{X}$.

DEFINITION 4 (l_1 Gini index, ZCA). Let $\mathbf{X}$ be a random vector whose mean vector is $\mathbf{m} = (m_1, m_2, \dots, m_n)^T$ and whose covariance matrix Σ is positive-definite. Denoted with μ the probability measure associated with $\mathbf{X}$, we define

$$G_1(\mathbf{X}) = \frac{1}{2 \sum_{i=1}^n |(W_\mu^{\text{ZCA}}\mathbf{m})_i|} \int_{\mathbb{R}^n \times \mathbb{R}^n} \sum_{i=1}^n |(W_\mu^{\text{ZCA}}(\mathbf{x} - \mathbf{y}))_i| \mu(d\mathbf{x}) \mu(d\mathbf{y}).$$

THEOREM 3. *Let $\mathbf{X}$ be a random vector of mean $\mathbf{m} = (m_1, m_2, \dots, m_n)^T$ and positive-definite $n \times n$ covariance matrix Σ . Then, we have*

$$(3.4) \quad G_1(\mathbf{X}) = \sum_{i=1}^n \frac{|m_i^*|}{\sum_{i=1}^n |m_i^*|} G((W_\mu^{\text{ZCA}} \mathbf{X})_i),$$

where $m_i^* = (W_\mu^{\text{ZCA}} \mathbf{m})_i$ and $G(X_i^*)$ is the one-dimensional Gini index of the i -th component of $\mathbf{X}^*$. Furthermore, if the components of $\mathbf{X}$ are non-negative and $(W_\mu^{\text{ZCA}})_{i,j} \geq 0$ for every $i, j = 1, \dots, n$, we have that

$$0 \leq G_1(\mathbf{X}) \leq 1.$$

Equation (3.4) is the most important result of this paper and establishes that the higher-dimensional Gini index induced by the l_1 Mahalanobis norm is a convex combination of the 1-dimensional Gini indices of the random variable $\mathbf{X}^* = W_\mu^{\text{ZCA}} \mathbf{X}$.

PROOF. First, notice that the mean vector of $\mathbf{X}^*$ is, by definition, $\mathbf{m}^* := W_\mu^{\text{ZCA}} \mathbf{m}$, so that

$$\sum_{i=1}^n |m_i^*| = \sum_{i=1}^n |(W_\mu^{\text{ZCA}} \mathbf{m})_i|.$$

By definition of the one-dimensional Gini index G , we have that

$$G(X_i^*) = \frac{1}{2|m_i^*|} \int_{\mathbb{R}^n \times \mathbb{R}^n} |x_i^* - y_i^*| (N_i)_\# \mu(d\mathbf{x}_*) (N_i)_\# \mu(d\mathbf{y}_*),$$

where μ is the probability measure associated with $\mathbf{X}$ and $N_i : \mathbb{R}^n \rightarrow \mathbb{R}^n$ defined as $N_i : \mathbf{x} \rightarrow (W_\mu^{\text{ZCA}} \mathbf{x})_i$. By a change of variables, we have that

$$G(X_i^*) = \frac{1}{2|m_i^*|} \int_{\mathbb{R}^n \times \mathbb{R}^n} |(W_\mu^{\text{ZCA}}(\mathbf{x} - \mathbf{y}))_i| \mu(d\mathbf{x}) \mu(d\mathbf{y}).$$

By plugging the value of every $G(X_i^*)$ in the right-hand side of (3.4), we conclude the first part of the thesis.

To conclude the proof, notice that since $W_\mu^{\text{ZCA}} \mathbf{X}$ is a non-negative random vector, we have that $G(X_i^*) \in [0, 1]$ and thus $G_1(\mathbf{X}) \in [0, 1]$ since $G_1(\mathbf{X})$ is a convex combination of values in $[0, 1]$. ■

It is insightful to interpret the result presented in Theorem 3 in the context of what a whitening process does to the multivariate data. When analyzing a multidimensional statistical distribution derived from a set of data, the measured quantities are typically interdependent. Consequently, the inequality expressed by the distribution, which measures how far apart are the individual multidimensional observations from each other, cannot be assessed as a function of the one-dimensional Gini indices obtained by

considering each one-dimensional component of the distribution at the time. However, by whitening the data, we can express the same inequality as a function of the standard one-dimensional Gini indices, applied to the whitened one-dimensional components.

Indeed, Theorem 3 indicates a natural method to combine the one-dimensional Gini indices by means of a convex combination. Specifically, the weight assigned to the Gini index of the i -th component of $\mathbf{X}^*$ is proportional to the absolute value of its mean, normalized by the sum of all mean values. We formalize the properties of the G_1 inequality measure in the following corollary.

COROLLARY 1. *Let $\mathbf{X}$ be a random vector of mean $\mathbf{m} = (m_1, m_2, \dots, m_n)^T$ and positive-definite $n \times n$ covariance matrix Σ_μ . Then, the following properties hold:*

- (1) *For any $\varepsilon > 0$, there exists a random variable $\mathbf{X}_\varepsilon$ such that $(W_{\mu_\varepsilon}^{\text{ZCA}})_{i,j} \geq 0$ for every $i, j = 1, \dots, n$ and $G_1(\mathbf{X}_\varepsilon) \leq \varepsilon$.*
- (2) *For any $\varepsilon > 0$, there exists a random variable $\mathbf{X}_\varepsilon$ such that $(W_{\mu_\varepsilon}^{\text{ZCA}})_{i,j} \geq 0$ for every $i, j = 1, \dots, n$ and $G_1(\mathbf{X}_\varepsilon) \geq 1 - \varepsilon$.*
- (3) *The G_1 is Scale Invariant, that is, $G_1(\mathbf{X}) = G_1(\mathbf{X}_Q)$ for any diagonal matrix $Q = \text{diag}(q_1, q_2, \dots, q_n)$ with $q_j > 0$. Moreover, if $(W_{\mu}^{\text{ZCA}})_{i,j} \geq 0$ for every $i, j = 1, \dots, n$ and $\mathbf{X}$ is non-negative, then the Rising Tide property, that is,*

$$G_1(\mathbf{X} + \mathbf{c}) \leq G_1(\mathbf{X}),$$

holds for any positive vector $\mathbf{c} \in \mathbb{R}_+^n$.

- (4) *Let us assume that $|m_1^*|$ is much larger than $\sum_{i=2}^n |m_i^*|$. Then,*

$$\lim_{|(m_1)_*| \rightarrow \infty} G_1(\mathbf{X}) = G(X_1^*)$$

that is, $G_1(\mathbf{X}) \sim G(X_1^)$.*

This property tells us that if one of the uncorrelated components of $\mathbf{X}^$ dominates the others meanwise, then the overall inequality is mostly determined by how unequal the dominant component is.*

PROOF. We divide the proof into four parts.

Proof of point (1). Let $\varepsilon > 0$ be fixed. Let $\mathbf{X} = (X_1, X_2, \dots, X_n)$ be a random vector whose components are independent and identically distributed. Moreover, assume that each X_i is distributed as follows:

$$X_i = \begin{cases} 0 & \text{with probability } p = \frac{1}{2}, \\ 2 & \text{with probability } p = \frac{1}{2}. \end{cases}$$

It is easy to check that $\mathbb{E}[X_i] = 1$, $\text{Var}(X_i) = 1$ for every $i = 1, \dots, n$, and, by construction, $\text{Cor}(X_i, X_j) = 0$ if $i \neq j$; therefore, $\mathbf{X}^* = W_\mu^{\text{ZCA}} \mathbf{X} = \mathbf{X}$. Since components $\mathbf{X}^*$

are identically distributed, formula (3.4) boils down to

$$G_1(\mathbf{X}) = G(X_1) = \frac{1}{2\mathbb{E}[X_1]} = \frac{1}{2}.$$

Given $M > 0$, let us define $\mathbf{X}_M = (X_1 + M, X_2 + M, \dots, X_n + M)$. By the same argument used above, we have that $\mathbf{X}_M^* = \mathbf{X}_M$ and that

$$G_1(\mathbf{X}_M) = G(X_1 + M) = \frac{1}{2(1 + M)}.$$

It is then easy to see that if $M > \frac{1}{2\varepsilon}$, then $G_1(\mathbf{X}_M) \leq \varepsilon$.

Proof of point (2). Let $\varepsilon > 0$ be fixed. Consider $\mathbf{X} = (X_1, X_2, \dots, X_n)$ to be a random vector whose components are independent and identically distributed. Moreover, assume that each X_i is distributed as follows:

$$X_i = \begin{cases} 0 & \text{with probability } 1 - p, \\ \frac{1}{\sqrt{p(1-p)}} & \text{with probability } p, \end{cases}$$

where $p \in (0, 1)$. It is easy to see that, for every $p \in (0, 1)$, the covariance matrix of $\mathbf{X}$ is the identity matrix, thus, $\mathbf{X}^* = W_\mu^{\text{ZCA}} \mathbf{X} = \mathbf{X}$. Moreover, we have that $\mathbb{E}[X_i] = \sqrt{\frac{p}{1-p}}$ and thus

$$G_1(\mathbf{X}) = G(X_1) = \frac{\sqrt{1-p}}{2\sqrt{p}} \frac{2p(1-p)}{\sqrt{p(1-p)}} = 1 - p.$$

In particular, if $p \leq \varepsilon$, we have $G_1(\mathbf{X}) \geq 1 - \varepsilon$.

Proof of point (3). The scale invariance follows directly from Theorem 2. Let us now consider the rising tide property. Let $\mathbf{c}$ be a vector whose components are positive, that is, $\mathbf{c} = (c_1, c_2, \dots, c_n)$, with $c_i \geq 0$. Since $\mathbf{X}$ and $\mathbf{X} + \mathbf{c}$ have the same covariance matrix, it follows that $W_\mu^{\text{ZCA}} = W_{\mathbf{X}+\mathbf{c}}^{\text{ZCA}}$ and that $\mathbb{E}[\mathbf{X} + \mathbf{c}] = \mathbb{E}[\mathbf{X}] + \mathbf{c}$. In particular, we have that

$$W_{\mathbf{X}+\mathbf{c}}^{\text{ZCA}}(\mathbb{E}[\mathbf{X} + \mathbf{c}]) = W_\mu^{\text{ZCA}}(\mathbb{E}[\mathbf{X}] + \mathbf{c}) = \mathbf{m}^* + W_\mu^{\text{ZCA}} \mathbf{c}$$

and thus

$$\begin{aligned} G_1(\mathbf{X} + \mathbf{c}) &= \frac{1}{2 \sum_{i=1}^n |(\mathbf{m}^* + W_\mu^{\text{ZCA}} \mathbf{c})_i|} \int_{\mathbb{R}^n \times \mathbb{R}^n} \sum_{i=1}^n |(W_\mu^{\text{ZCA}}(\mathbf{x} + \mathbf{c} - (\mathbf{y} + \mathbf{c})))_i| \mu(d\mathbf{x}) \mu(d\mathbf{y}) \\ &= \frac{1}{2 \sum_{i=1}^n |(\mathbf{m}^* + W_\mu^{\text{ZCA}} \mathbf{c})_i|} \int_{\mathbb{R}^n \times \mathbb{R}^n} \sum_{i=1}^n |(W_\mu^{\text{ZCA}}(\mathbf{x} - \mathbf{y}))_i| \mu(d\mathbf{x}) \mu(d\mathbf{y}) \\ &\leq \frac{1}{2 \sum_{i=1}^n |\mathbf{m}_i^*|} \int_{\mathbb{R}^n \times \mathbb{R}^n} \sum_{i=1}^n |(W_\mu^{\text{ZCA}}(\mathbf{x} - \mathbf{y}))_i| \mu(d\mathbf{x}) \mu(d\mathbf{y}) = G_1(\mathbf{X}), \end{aligned}$$

where the last inequality follows from the fact that W_μ^{ZCA} maps $\{\mathbf{x} \in \mathbb{R}^n \text{ s.t. } x_i \geq 0\}$ into itself, thus, $|(\mathbf{m}^* + W_\mu^{ZCA}\mathbf{c})_i| = |\mathbf{m}_i^*| + |(W_\mu^{ZCA}\mathbf{c})_i| \geq |\mathbf{m}_i^*|$ for every $i = 1, \dots, n$.

Proof of point (4). It follows from the following identity:

$$\lim_{|(m_1)_*| \rightarrow \infty} \frac{|m_i^*|}{\sum_{j=1}^n |m_j^*|} = \begin{cases} 1 & \text{if } i = 1, \\ 0 & \text{otherwise.} \end{cases} \quad \blacksquare$$

Finally, since any affine transformation of a Gaussian random vector is a Gaussian random vector, we can explicitly express G_1 as a function of the parameters of the Gaussian distribution.

THEOREM 4. *If $\mathbf{X}$ is a Gaussian random vector whose mean is non-null, that is, $\mathbf{m} \neq 0$, then also $\mathbf{X}^*$ is a Gaussian random vector with non-null mean. Moreover, we have that*

$$G_1(\mathbf{X}) = \frac{n}{\sqrt{\pi} \sum_{i=1}^n |(W_\mu^{ZCA}\mathbf{m})_i|}.$$

PROOF. It follows from the fact that any whitened multivariate Gaussian distribution is a Gaussian distribution with independent components and whose covariance matrix is the identity. $\blacksquare$

This shows that the G_1 index of a Gaussian distribution is proportional to the inverse of the $\|\mathbf{m}^*\|_1$. However, the result in Corollary 1 does not hold, as Gaussian vectors take value on $\mathbb{R}^n$. In this case, Theorem 4 is to be interpreted as giving the explicit expression of a multidimensional coefficient of variation, which exhibits the same properties of that of Voinov and Nikulin [1, 28], which is obtained by evaluating the multivariate Gini index of a Gaussian random vector with respect to the 2-Mahalanobis Metric [15], and it is related to the G_2 Gini index [6].

Finally, notice that Theorem 4 can be generalized to any case in which the multivariate probability distribution allows an explicit computation of the one-dimensional Gini index for its components.

4. APPLICATION

In this section, we show the practical importance of our proposed multivariate index of inequality, by means of the example introduced in [15], which concerns the study of market economic inequality.

A country market is unequal, from an economic viewpoint, when it is concentrated, that is, when it presents a high inequality: few companies have a large size and many have a small size. To measure market inequality, we need to specify how we measure the size of a company, using publicly available data. For publicly listed companies,

we can consider, for example, the daily market capitalization, the current number of employees, and the yearly revenues. This information is publicly downloadable from the website companiesmarketcap.com, which contains, at the moment, the 8,081 largest companies in the world (by capitalization). Table 1 reports the summary statistics for all companies, in terms of Market Capitalization, Number of Employees, and Revenues.

	Mean	Standard deviation
MarketCap	$1.22 \cdot 10^{10}$	$7.15 \cdot 10^{10}$
Revenues	$6.86 \cdot 10^9$	$2.49 \cdot 10^{10}$
Employees	$1.50 \cdot 10^4$	$5.26 \cdot 10^4$

TABLE 1. Summary statistics.

Table 1 shows that, as expected, both the mean and standard deviation of Market Capitalization and Revenues are much larger than those of the Number of Employees. In addition, the variability from the mean of the Market Capitalization is about three times higher than that of the Revenues. The three variables are not much correlated with each other, as their correlation matrix in Table 2 shows.

	MarketCap	Employees	Revenues
MarketCap	1.000	0.010	0.102
Employees	0.010	1.000	0.036
Revenues	0.102	0.036	1.000

TABLE 2. Correlation matrix.

It is usually of interest to compare market inequality in different countries. This can be done comparing the value of the Gini one-dimensional indices. For the sake of illustration, and without loss of generality, here we will measure market inequality at the overall level as well as for nine of the largest economies: Canada, China, France, Germany, Italy, France, India, Japan, the United Kingdom, and the United States. Table 3 shows the calculation of the one-dimensional Gini indices, using Market Capitalization, Number of Employees, and Revenues as the metrics with which to measure the size of the company.

From Table 3, we infer that, when all countries are considered, the three Gini one-dimensional indices are very similar to each other. Differently, when individual countries are considered, there are remarkable differences. For example, if we consider market capitalization, the United States is the most concentrated country, followed by Canada, Germany, and France. However, if we measure size in terms of number of employees, the United States is followed by Canada, Germany, and India. In terms of revenues, the most concentrated country appears Canada, followed by India, the United States, and Italy.

Countries	Number of companies	Gini MarketCap	Gini Employees	Gini Revenues	G_1
United States	3652	0.886	0.845	0.851	0.856
Canada	395	0.794	0.840	0.879	0.829
France	119	0.767	0.794	0.805	0.789
Germany	220	0.777	0.838	0.777	0.793
Italy	86	0.638	0.783	0.829	0.737
United Kingdom	258	0.754	0.813	0.828	0.794
China	314	0.761	0.782	0.785	0.776
India	564	0.747	0.816	0.860	0.801
Japan	350	0.667	0.771	0.714	0.715
All	8081	0.830	0.833	0.835	0.832

TABLE 3. Unidimensional Gini coefficients (referred as MarketCap, Employee, and Revenue) and multidimensional G_1 coefficient.

We notice that we do not have a unique ranking of the countries, in terms of market inequality: it depends on how we define the size of a company: using market capitalization, number of employees, or revenues. The intuition suggests that we should take all three scales into account, to attain a reliable ranking of the countries, in terms of market inequality. Hence, a multidimensional measure of inequality is necessary. The multidimensional G_1 index fills this gap. Table 3 reports, in the first right column, the values of the G_1 index, obtained applying equation (3.4) to the whitening process defined as in (3.3). The values of G_1 show that, considering all world countries, the multidimensional Gini index is equal to 0.82, in line with the individual values.

More importantly, the multidimensional index gives a ranking of country inequality that take all three size measurements into account. The most unequal country (most concentrated market) is the United States, followed by Canada, in line with the results of the individual Gini indices for Market Capitalization and Employees. The third most concentrated market is India, owing to its high concentration of revenues. The United Kingdom, Germany, China, and France follow, close to each other. The least unequal countries are Italy and Japan, characterized by many small and medium enterprises.

For completeness, we remark that the weights attributed to the individual indices, in equation (3.4), are equal to (0.335, 0.301, 0.363), respectively, for Market Capitalization, Employees, and Revenues. This means that the whitening process gives a slightly higher weight to the inequality in Revenues, followed by that in Market Capitalization and, last, by that in Number of Employees.

5. CONCLUSIONS

In this note, we introduced and discussed a rigorous way to extend the well-known univariate Gini index to multivariate distributions by maintaining most of its one-dimensional properties. At variance with other existing proposals, our extension is

based on applying to a given random n -dimensional vector a whitening process that possesses the property of scale stability, a property which is naturally satisfied by the one-dimensional Gini index.

We tested our proposed multivariate index of inequality on a relevant example concerned with the study of market economic inequalities. However, our proposal can be fruitfully applied to all situations in which inequality has to be measured through multidimensional data.

Among others, a very important issue would be the comparison of well-being inequality across the world countries (more than 200). Traditional comparisons measure inequality in well-being measuring only income. To achieve a better result, inequality in well-being should be measured also taking other aspects into consideration, like education and health levels.

Future research involves extending the multidimensional Gini as an evaluation measurement of SAFE machine learning and artificial intelligence, see e.g. [7] and the references therein.

FUNDING. — This work has been carried out under the activities of the National Group of Mathematical Physics (GNFM). GT wishes to acknowledge partial support by IMATI, Institute for Applied Mathematics and Information Technologies “Enrico Magenes”, Pavia, Italy. The work has also been partially supported by the European Union - NextGenerationEU, in the framework of the GRINS - Growing Resilient, INclusive and Sustainable (GRINS PE00000018). The views and opinions expressed are solely those of the authors and do not necessarily reflect those of the European Union, nor can the European Union be held responsible for them.

REFERENCES

- [1] S. AERTS – G. HAESBROECK – C. RUWET, [Multivariate coefficients of variation: comparison and influence functions](#). *J. Multivariate Anal.* **142** (2015), 183–198. Zbl [1327.62336](#) MR [3412747](#)
- [2] B. C. ARNOLD, [Majorization and the Lorenz order: A brief introduction](#). Lect. Notes Stat. 43, Springer, Berlin, 1987. Zbl [0649.62041](#) MR [0915159](#)
- [3] B. C. ARNOLD – J. M. SARABIA, [Analytic expressions for multivariate Lorenz surfaces](#). *Sankhya A* **80** (2018), S84–S111. Zbl [1430.62097](#) MR [3968359](#)
- [4] G. AURICCHIO – A. CODEGONI – S. GUALANDI – G. TOSCANI – M. VENERONI, [The equivalence of Fourier-based and Wasserstein metrics on imaging problems](#). *Atti Accad. Naz. Lincei Rend. Lincei Mat. Appl.* **31** (2020), no. 3, 627–649. Zbl [1457.60002](#) MR [4170640](#)

- [5] G. AURICCHIO – A. CODEGONI – S. GUALANDI – L. ZAMBON, [The Fourier discrepancy function](#). *Commun. Math. Sci.* **21** (2023), no. 3, 627–639. Zbl [1518.60003](#) MR [4562448](#)
- [6] G. AURICCHIO – P. GIUDICI – G. TOSCANI, How to measure multidimensional variation? 2024, [arXiv:2411.19529v1](#).
- [7] G. BABAEI – P. GIUDICI – E. RAFFINETTI, [A rank graduation box for SAFE AI](#). *Expert Syst. Appl.* **259** (2025), article no. 125239.
- [8] I. ELIAZAR, [A tour of inequality](#). *Ann. Physics* **389** (2018), 306–332. Zbl [1386.91109](#) MR [3762022](#)
- [9] J. H. FRIEDMAN, [Exploratory projection pursuit](#). *J. Amer. Statist. Assoc.* **82** (1987), no. 397, 249–266. Zbl [0664.62060](#) MR [0883353](#)
- [10] T. GAJDOS – J. A. WEYMARK, [Multidimensional generalized Gini indices](#). *Econom. Theory* **26** (2005), no. 3, 471–496. Zbl [1084.91050](#) MR [2214171](#)
- [11] A. GHOSH – A. CHATTERJEE – J. INOUE – B. K. CHAKRABARTI, [Inequality measures in kinetic exchange models of wealth distributions](#). *Physica A* **451** (2016), 465–474.
- [12] A. GHOSH – N. CHATTOPADHYAY – B. K. CHAKRABARTI, [Inequality in societies, academic institutions and science journals: Gini and \$k\$ -indices](#). *Physica A* **410** (2014), 30–34.
- [13] C. GINI, Sulla misura della concentrazione e della variabilità dei caratteri. *Lettere ed Arti* **73** (1914), 1203–1248. English translation: *Metron* **63** (2005), 3–38.
- [14] C. GINI, [Measurement of inequality of incomes](#). *Econ. J.* **31** (1921), 124–126.
- [15] P. GIUDICI – E. RAFFINETTI – G. TOSCANI, [Measuring multidimensional inequality: a new proposal based on the Fourier transform](#). *Statistics* **59** (2025), no. 2, 330–353. Zbl [08023935](#) MR [4875868](#)
- [16] O. GROTHE – F. KÄCHELE – F. SCHMID, [A multivariate extension of the lorenz curve based on copulas and a related multivariate Gini coefficient](#). *J. Econ. Inequal.* **20** (2022), 727–748.
- [17] N. HURLEY – S. RICKARD, [Comparing measures of sparsity](#). *IEEE Trans. Inform. Theory* **55** (2009), no. 10, 4723–4741. Zbl [1367.94094](#) MR [2597572](#)
- [18] A. KESSY – A. LEWIN – K. STRIMMER, [Optimal whitening and decorrelation](#). *Amer. Statist.* **72** (2018), no. 4, 309–314. Zbl [07663954](#) MR [3878084](#)
- [19] G. KOSHEVOY – K. MOSLER, [The Lorenz zonoid of a multivariate distribution](#). *J. Amer. Statist. Assoc.* **91** (1996), no. 434, 873–882. Zbl [0869.62041](#) MR [1395754](#)
- [20] G. A. KOSHEVOY – K. MOSLER, [Multivariate Gini indices](#). *J. Multivariate Anal.* **60** (1997), no. 2, 252–276. Zbl [0873.62062](#) MR [1437736](#)
- [21] G. LI – J. ZHANG, Sphering and its properties. *Sankhyā Ser. A* **60** (1998), no. 1, 119–133. Zbl [0976.62062](#) MR [1714783](#)
- [22] M. LORENZ, [Methods of measuring the concentration of wealth](#). *Publ. Am. Stat. Assoc.* **9** (1905), no. 70, 209–219.
- [23] P. C. MAHALANOBIS, On the generalized distance in statistics. *Proc. Nat. Inst. Sci. India* **2** (1936), 49–55. Zbl [0015.03302](#)

- [24] J. M. SARABIA – V. JORDA, [Lorenz surfaces based on the Sarmanov–Lee distribution with applications to multidimensional inequality in well-being](#). *Mathematics* **8** (2020), 2095.
- [25] T. TAGUCHI, [On the two-dimensional concentration surface and extensions of concentration coefficient and Pareto distribution to the two dimensional case \(on an application of differential geometric methods to statistical analysis\)](#). I. *Ann. Inst. Statist. Math.* **24** (1972), 355–381. Zbl [0349.62033](#) MR [0345380](#)
- [26] T. TAGUCHI, [On the two-dimensional concentration surface and extensions of concentration coefficient and Pareto distribution to the two dimensional case \(on an application of differential geometric methods to statistical analysis\)](#). II. *Ann. Inst. Statist. Math.* **24** (1972), 599–619. Zbl [0329.62016](#) MR [0345381](#)
- [27] G. TOSCANI, [On Fourier-based inequality indices](#). *Entropy* **24** (2022), no. 10, article no. 1393. MR [4513975](#)
- [28] V. G. VOINOV – M. S. NIKULIN, *Unbiased estimators and their applications. Vol. 2. Multivariate case*. Math. Appl. 362, Kluwer Academic Publishers Group, Dordrecht, 1996. Zbl [0844.62039](#) MR [1422256](#)

Received 22 September 2024,
and in revised form 19 December 2024

Gennaro Auricchio
Department of Mathematics, University of Padua
Via Trieste, 63, 10623 Padova, Italy
gennaro.auricchio@unipd.it

Paolo Giudici
Department of Economics and Management, University of Pavia
Via S. Felice, 5, 27100 Pavia, Italy
paolo.giudici@unipv.it

Giuseppe Toscani (corresponding author)
Department of Mathematics, University of Pavia
Via Ferrata, 5, 27100 Pavia
Institute of Applied Mathematics and Information Technologies
Via Ferrata, 5, 27100 Pavia, Italy
giuseppe.toscani@unipv.it