

Article

Not peer-reviewed version

Robust Estimations from Distribution Structures: II. Central Moments

[Tuobang Li](#) *

Posted Date: 15 February 2024

doi: 10.20944/preprints202402.0843.v1

Keywords: moments; invariant; unimodal; U -statistics; generalized L-statistics



Preprints.org is a free multidiscipline platform providing preprint service that is dedicated to making early versions of research outputs permanently available and citable. Preprints posted at Preprints.org appear in Web of Science, Crossref, Google Scholar, Scilit, Europe PMC.

Copyright: This is an open access article distributed under the Creative Commons Attribution License which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited.

Article

Robust Estimations from Distribution Structures: II, Central Moments

Tuobang Li ^{1,2,3},

¹ Technion-Israel Institute of Technology, Haifa 32000, Israel; tl@biomathematics.org

² Guangdong Technion-Israel Institute of Technology, Shantou 515063, China

³ University of California, Berkeley, CA 94720

Abstract: In descriptive statistics, U -statistics arise naturally in producing minimum-variance unbiased estimators. In 1984, Serfling considered the distribution formed by evaluating the kernel of the U -statistics and proposed generalized L -statistics which includes Hodges-Lehmann estimator and Bickel-Lehmann spread as special cases. However, the structures of the kernel distributions remain unclear. In 1954, Hodges and Lehmann demonstrated that if X and Y are independently sampled from the same unimodal distribution, $X - Y$ will exhibit symmetrical unimodality with its peak centered at zero. Building upon this foundational work, the current study delves into the structure of the kernel distribution. It is shown that the k th central moment kernel distributions ($k > 2$) derived from a unimodal distribution exhibit location invariance and is also nearly unimodal with the mode and median close to zero. This article provides an approach to study the general structure of kernel distributions.

Keywords: moments; invariant; unimoda; U -statistics; generalized L -statistics

Significance Statement In nonparametric statistics, the focus is on the relative differences of robust estimators, which is considered more crucial than their precise values. This principle implies that if the underlying distribution's parameters shift, then all corresponding nonparametric estimates are expected to uniformly and asymptotically adjust in a consistent direction, provided they target the same characteristic of the distribution. This article discusses the validity of this fundamental principle of nonparametrics in various scenarios. It is found that for the k th central moment, kernel distributions generally follow this principle.

The most popular robust scale estimator currently, the median absolute deviation, was popularized by Hampel (1974) [1], who credits the idea to Gauss in 1816 [2]. It can be seen as evaluating the median of a pseudo-sample formed by the absolute deviations of all values related to the sample median. The pseudo-sample size is n . Indeed, most scale estimators can be transformed in such ways. For example, range or interquartile range can be seen as evaluating the mean of a pseudo-sample with two values, and they belong to the class of scale estimators called quantile differences. In 1976, in their landmark series *Descriptive Statistics for Nonparametric Models*, Bickel and Lehmann [3] generalized a class of estimators as measures of the dispersion of a symmetric distribution around its center of symmetry. Median absolute deviation, sample variance, and average absolute deviation are all belonging to this class. In 1979, the same series, they [4] proposed a class of estimators referred to as measures of spread, which consider the pairwise differences of a random variable, irrespective of its symmetry, throughout its distribution, rather than focusing on dispersion relative to a fixed point. In the final section [4], they explored a version of the trimmed standard deviation based on n^2 pairwise differences, which is modified here for comparison,

$$\left[n \left(\frac{1}{2} - \epsilon \right) \right]^{-\frac{1}{2}} \left[\sum_{i=\frac{n}{2}}^{n(1-\epsilon)} [X_i - X_{n-i+1}]^2 \right]^{\frac{1}{2}}, \quad (1)$$

and

$$\left[\binom{n}{2} (1 - \epsilon_0 - \gamma \epsilon_0) \right]^{-\frac{1}{2}} \left[\sum_{i=\binom{n}{2}\gamma\epsilon_0}^{\binom{n}{2}(1-\epsilon_0)} (X_{i_1} - X_{i_2})^2 \right]^{\frac{1}{2}}, \quad (2)$$

where $(X_{i_1} - X_{i_2})_1 \leq \dots \leq (X_{i_1} - X_{i_2})_{\binom{n}{2}}$ are the order statistics of the pseudosample, $X_{i_1} - X_{i_2}$, $i_1 < i_2$, provided that $\binom{n}{2}\gamma\epsilon_0 \in \mathbb{N}$ and $\binom{n}{2}(1 - \epsilon_0) \in \mathbb{N}$. They showed that, when $\epsilon_0 = 0$, the result obtained using [2] is equal to $\sqrt{2}$ times the sample standard deviation. The paper ended with, "We do not know a fortiori which of the measures is preferable and leave these interesting questions open."

Two examples of the impacts of that series are as follows. Oja (1981, 1983) [5,6] provided a more comprehensive and generalized examination of these concepts, and integrated the measures of location, dispersion, and spread as proposed by Bickel and Lehmann [3,4,7], along with van Zwet's convex transformation order of skewness and kurtosis (1964) [8] for univariate and multivariate distributions, resulting a greater degree of generality and a broader perspective on these statistical constructs. Rousseeuw and Croux proposed a popular efficient scale estimator based on separate medians of pairwise differences taken over i_1 and i_2 [9] in 1993. However the importance of tackling the symmetry assumption has been greatly underestimated, as will be discussed later.

Here, their open question is addressed in two different aspects [4]. First, since the estimation of scale can be transformed into the location estimation of a pseudo-sample, according to the principle of the central limit theorem, the variances of such scale estimators should be linearly dependent on the standard deviation of the pseudo-sample and inversely dependent on the square root of the pseudo-sample size. Then, [2] based on n^2 pairwise differences is obviously better than [1] since the ratio of its pseudo-sample size over that of [1] is n . So if just considering the size, the variance of [2] is $\frac{1}{\sqrt{n}}$ of the variance of [1]. Another factor that needs to be considered is the standard deviation of the pseudo-sample. However, the standard deviation of the pseudo-sample is generally independent of different pseudo-sample sizes. So, no matter how different the standard deviation of the pseudo-samples of [1] and [2] is, as the sample size increases, the variance of [1] will always dominate that of [2]. Second, the nomenclature used in this series is introduced as follows:

Nomenclature. Given a robust estimator, $\hat{\theta}$, which has an adjustable breakdown point, ϵ , that can approach zero asymptotically, the name of $\hat{\theta}$ comprises two parts: the first part denotes the type of estimator, and the second part represents the population parameter θ , such that $\hat{\theta} \rightarrow \theta$ as $\epsilon \rightarrow 0$. The abbreviation of the estimator combines the initial letters of the first part and the second part. If the estimator is symmetric, the upper asymptotic breakdown point, ϵ , is indicated in the subscript of the abbreviation of the estimator, with the exception of the median. For an asymmetric estimator based on quantile average, the associated γ follows ϵ .

In REDS I [10], it was shown that the bias of a robust estimator with an adjustable breakdown point is often monotonic with respect to the breakdown point in a semiparametric distribution. Naturally, the estimator's name should reflect the population parameter that it approaches as $\epsilon \rightarrow 0$. If multiplying all pseudo-samples by a factor of $\frac{1}{\sqrt{2}}$, then [2] is the trimmed standard deviation adhering to this nomenclature, since $\psi_2(x_1, x_2) = \frac{1}{2}(x_1 - x_2)^2$ is the kernel function of the unbiased estimation of the second central moment by using U -statistic [11]. This definition should be preferable, not only because it is the square root of a trimmed U -statistic, which is closely related to the minimum-variance unbiased estimator (MVUE), but also because the second γ -orderliness of the second central moment kernel distribution is ensured by the next exciting theorem.

Theorem 1. *The second central moment kernel distribution generated from any unimodal distribution is second γ -ordered, provided that $\gamma \geq 0$.*

Proof. In 1954, Hodges and Lehmann established that if X and Y are independently drawn from the same unimodal distribution, $X - Y$ will be a symmetric unimodal distribution peaking at zero [12]. Given the constraint in the pairwise differences that $X_{i_1} < X_{i_2}$, $i_1 < i_2$, it directly follows

from Theorem 1 in [12] that the pairwise difference distribution (Ξ_Δ) generated from any unimodal distribution is always monotonic increasing with a mode at zero. Since $X - X'$ is a negative variable that is monotonically increasing, applying the squaring transformation, the relationship between the original variable $X - X'$ and its squared counterpart $(X - X')^2$ can be represented as follows: $X - X' < Y - Y' \implies (X - X')^2 > (Y - Y')^2$. In other words, as the negative values of $X - X'$ become larger in magnitude (more negative), their squared values $(X - X')^2$ become larger as well, but in a monotonically decreasing manner with a mode at zero. Further multiplication by $\frac{1}{2}$ also does not change the monotonicity and mode, since the mode is zero. Therefore, the transformed pdf becomes monotonically decreasing with a mode at zero. In REDS I [10], it was proven that a right-skewed distribution with a monotonic decreasing pdf is always second γ -ordered, which gives the desired result. $\square$

In REDS I [10], it was shown that any symmetric distribution is ν th U -ordered, suggesting that ν th U -orderliness does not require unimodality, e.g., a symmetric bimodal distribution is also ν th U -ordered. In the SI Text of REDS I [10], an analysis of the Weibull distribution showed that unimodality does not assure orderliness. Theorem 1 uncovers a profound relationship between unimodality, monotonicity, and second γ -orderliness, which is sufficient for γ -trimming inequality and γ -orderliness.

On the other hand, while robust estimation of scale has been intensively studied with established methods [3,4], the development of robust measures of asymmetry and kurtosis lags behind, despite the availability of several approaches [13–17]. The purpose of this paper is to demonstrate that, in light of previous works, the estimation of all central moments can be transformed into a location estimation problem by using U -statistics and the central moment kernel distributions possess desirable properties.

Robust Estimations of the Central Moments

In 1928, Fisher constructed $\mathbf{k}$ -statistics as unbiased estimators of cumulants [18]. Halmos (1946) proved that a functional θ admits an unbiased estimator if and only if it is a regular statistical functional of degree $\mathbf{k}$ and showed a relation of symmetry, unbiasedness and minimum variance [19]. Hoeffding, in 1948, generalized U -statistics [20] which enable the derivation of a minimum-variance unbiased estimator from each unbiased estimator of an estimable parameter. In 1984, Serfling pointed out the speciality of Hodges-Lehmann estimator, which is neither a simple L -statistic nor a U -statistic, and considered the generalized L -statistics and trimmed U -statistics [21]. Given a kernel function $h_{\mathbf{k}}$ which is a symmetric function of $\mathbf{k}$ variables, the LU -statistic is defined as:

$$LU_{h_{\mathbf{k}}, \mathbf{k}, \epsilon, \gamma, n} := LL_{k, \epsilon_0, \gamma, n} \left(\text{sort} \left((h_{\mathbf{k}}(X_{N_1}, \dots, X_{N_{\mathbf{k}}}))_{N=1}^{\binom{n}{\mathbf{k}}} \right) \right),$$

where $\epsilon = 1 - (1 - \epsilon_0)^{\frac{1}{\mathbf{k}}}$ (proven in Subsection ?? in REDS III [22]), $X_{N_1}, \dots, X_{N_{\mathbf{k}}}$ are the n choose $\mathbf{k}$ elements from the sample, $LL_{k, \epsilon_0, \gamma, n}(Y)$ denotes the LL -statistic with the sorted sequence $\text{sort} \left((h_{\mathbf{k}}(X_{N_1}, \dots, X_{N_{\mathbf{k}}}))_{N=1}^{\binom{n}{\mathbf{k}}} \right)$ serving as an input. In the context of Serfling's work, the term 'trimmed U -statistic' is used when $LL_{k, \epsilon_0, \gamma, n}$ is $TM_{\epsilon_0, \gamma, n}$ [21].

In 1997, Heffernan [11] obtained an unbiased estimator of the $\mathbf{k}$ th central moment by using U -statistics and demonstrated that it is the minimum variance unbiased estimator for distributions with the finite first $\mathbf{k}$ moments. The weighted H-L $\mathbf{k}$ th central moment ($2 \leq \mathbf{k} \leq n$) is thus defined as,

$$\text{WHLM}_{k, \epsilon, \gamma, n} := LU_{h_{\mathbf{k}} = \psi_{\mathbf{k}}, \mathbf{k}, \epsilon, \gamma, n},$$

where $\text{WHLM}_{k, \epsilon_0, \gamma, n}$ is used as the $LL_{k, \epsilon_0, \gamma, n}$ in LU , $\psi_{\mathbf{k}}(x_1, \dots, x_{\mathbf{k}}) = \sum_{j=0}^{\mathbf{k}-2} (-1)^j \binom{1}{\mathbf{k}-j} \sum \left(x_{i_1}^{\mathbf{k}-j} x_{i_2} \dots x_{i_{j+1}} \right) + (-1)^{\mathbf{k}-1} (\mathbf{k}-1) x_1 \dots x_{\mathbf{k}}$, the second summation is over $i_1, \dots, i_{j+1} = 1$ to $\mathbf{k}$ with $i_1 \neq i_2 \neq \dots \neq i_{j+1}$ and $i_2 < i_3 < \dots < i_{j+1}$ [11]. Despite the complexity, the following theorem offers an approach to infer the general structure of such kernel distributions.

Theorem 2. Define a set T comprising all pairs $(\psi_{\mathbf{k}}(\mathbf{v}), f_{X,\dots,X}(\mathbf{v}))$ such that $\psi_{\mathbf{k}}(\mathbf{v}) = \psi_{\mathbf{k}}(Q(p_1), \dots, Q(p_k))$ with $Q(p_1) < \dots < Q(p_k)$ and $f_{X,\dots,X}(\mathbf{v}) = \mathbf{k}! f(Q(p_1)) \dots f(Q(p_k))$ is the probability density of the $\mathbf{k}$ -tuple, $\mathbf{v} = (Q(p_1), \dots, Q(p_k))$ (a formula drawn after a modification of the Jacobian density theorem). T_{Δ} is a subset of T , consisting all those pairs for which the corresponding $\mathbf{k}$ -tuples satisfy that $Q(p_1) - Q(p_k) = \Delta$. The component quasi-distribution, denoted by ξ_{Δ} , has a quasi-pdf $f_{\xi_{\Delta}}(\bar{\Delta}) = \sum_{(\psi_{\mathbf{k}}(\mathbf{v}), f_{X,\dots,X}(\mathbf{v})) \in T_{\Delta}} f_{X,\dots,X}(\mathbf{v})$, i.e., sum over all $f_{X,\dots,X}(\mathbf{v})$ such that the pair $(\psi_{\mathbf{k}}(\mathbf{v}), f_{X,\dots,X}(\mathbf{v}))$ is $\bar{\Delta} = \psi_{\mathbf{k}}(\mathbf{v})$ in the set T_{Δ} and the first element of the pair, $\psi_{\mathbf{k}}(\mathbf{v})$, is equal to $\bar{\Delta}$. The $\mathbf{k}$ th, where $\mathbf{k} > 2$, central moment kernel distribution, labeled $\Xi_{\mathbf{k}}$, can be seen as a quasi-mixture distribution comprising an infinite number of component quasi-distributions, ξ_{Δ} s, each corresponding to a different value of Δ , which ranges from $Q(0) - Q(1)$ to 0 . Each component quasi-distribution has a support of $\left(-\left(\frac{\mathbf{k}}{3+(-1)^{\mathbf{k}}}\right)^{-1} (-\Delta)^{\mathbf{k}}, \frac{1}{\mathbf{k}} (-\Delta)^{\mathbf{k}} \right)$.

Proof. The support of ξ_{Δ} is the extrema of the function $\psi_{\mathbf{k}}(Q(p_1), \dots, Q(p_k))$ subjected to the constraints, $Q(p_1) < \dots < Q(p_k)$ and $\Delta = Q(p_1) - Q(p_k)$. Using the Lagrange multiplier, the only critical point can be determined at $Q(p_1) = \dots = Q(p_k) = 0$, where $\psi_{\mathbf{k}} = 0$. Other candidates are within the boundaries, i.e., $\psi_{\mathbf{k}}(x_1 = Q(p_1), x_2 = Q(p_k), \dots, x_k = Q(p_k))$, $\dots$, $\psi_{\mathbf{k}}(x_1 = Q(p_1), \dots, x_i = Q(p_1), x_{i+1} = Q(p_k), \dots, x_k = Q(p_k))$, $\dots$, $\psi_{\mathbf{k}}(x_1 = Q(p_1), \dots, x_{k-1} = Q(p_1), x_k = Q(p_k))$. $\psi_{\mathbf{k}}(x_1 = Q(p_1), \dots, x_i = Q(p_1), x_{i+1} = Q(p_k), \dots, x_k = Q(p_k))$ can be divided into $\mathbf{k}$ groups. The g th group has the common factor $(-1)^{g+1} \frac{1}{\mathbf{k}-g+1}$, if $1 \leq g \leq \mathbf{k}-1$ and the final $\mathbf{k}$ th group is the term $(-1)^{\mathbf{k}-1} (\mathbf{k}-1) Q(p_1)^i Q(p_k)^{\mathbf{k}-i}$. If $\frac{\mathbf{k}+1-i}{2} \leq j \leq \frac{\mathbf{k}-1}{2}$ and $j+1 \leq g \leq \mathbf{k}-j$, the g th group has $i \binom{i-1}{g-j-1} \binom{\mathbf{k}-i}{j}$ terms having the form $(-1)^{g+1} \frac{1}{\mathbf{k}-g+1} Q(p_1)^{\mathbf{k}-j} Q(p_k)^j$. If $\frac{\mathbf{k}+1-i}{2} \leq j \leq \frac{\mathbf{k}-1}{2}$ and $\mathbf{k}-j+1 \leq g \leq i+j$, the g th group has $i \binom{i-1}{g-j-1} \binom{\mathbf{k}-i}{j} + (\mathbf{k}-i) \binom{\mathbf{k}-i-1}{j-\mathbf{k}+g-1} \binom{i}{\mathbf{k}-j}$ terms having the form $(-1)^{g+1} \frac{1}{\mathbf{k}-g+1} Q(p_1)^{\mathbf{k}-j} Q(p_k)^j$. If $0 \leq j < \frac{\mathbf{k}+1-i}{2}$ and $j+1 \leq g \leq i+j$, the g th group has $i \binom{i-1}{g-j-1} \binom{\mathbf{k}-i}{j}$ terms having the form $(-1)^{g+1} \frac{1}{\mathbf{k}-g+1} Q(p_1)^{\mathbf{k}-j} Q(p_k)^j$. If $\frac{\mathbf{k}}{2} \leq j \leq \mathbf{k}$ and $\mathbf{k}-j+1 \leq g \leq j$, the g th group has $(\mathbf{k}-i) \binom{\mathbf{k}-i-1}{j-\mathbf{k}+g-1} \binom{i}{\mathbf{k}-j}$ terms having the form $(-1)^{g+1} \frac{1}{\mathbf{k}-g+1} Q(p_1)^{\mathbf{k}-j} Q(p_k)^j$. If $\frac{\mathbf{k}}{2} \leq j \leq \mathbf{k}$ and $j+1 \leq g \leq j+i < \mathbf{k}$, the g th group has $i \binom{i-1}{g-j-1} \binom{\mathbf{k}-i}{j} + (\mathbf{k}-i) \binom{\mathbf{k}-i-1}{j-\mathbf{k}+g-1} \binom{i}{\mathbf{k}-j}$ terms having the form $(-1)^{g+1} \frac{1}{\mathbf{k}-g+1} Q(p_1)^{\mathbf{k}-j} Q(p_k)^j$. So, if $i+j = \mathbf{k}$, $\frac{\mathbf{k}}{2} \leq j \leq \mathbf{k}$, $0 \leq i \leq \frac{\mathbf{k}}{2}$, the summed coefficient of $Q(p_1)^i Q(p_k)^{\mathbf{k}-i}$ is $(-1)^{\mathbf{k}-1} (\mathbf{k}-1) + \sum_{g=i+1}^{\mathbf{k}-1} (-1)^{g+1} \frac{1}{\mathbf{k}-g+1} (\mathbf{k}-i) \binom{\mathbf{k}-i-1}{g-i-1} + \sum_{g=\mathbf{k}-i+1}^{\mathbf{k}-1} (-1)^{g+1} \frac{1}{\mathbf{k}-g+1} i \binom{i-1}{g-\mathbf{k}+i-1} = (-1)^{\mathbf{k}-1} (\mathbf{k}-1) + (-1)^{\mathbf{k}+1} + (\mathbf{k}-i) (-1)^{\mathbf{k}} + (-1)^{\mathbf{k}} (i-1) = (-1)^{\mathbf{k}+1}$. The summation identities are $\sum_{g=i+1}^{\mathbf{k}-1} (-1)^{g+1} \frac{1}{\mathbf{k}-g+1} (\mathbf{k}-i) \binom{\mathbf{k}-i-1}{g-i-1} = (\mathbf{k}-i) \int_0^1 \sum_{g=i+1}^{\mathbf{k}-1} (-1)^{g+1} \binom{\mathbf{k}-i-1}{g-i-1} t^{\mathbf{k}-g} dt = (\mathbf{k}-i) \int_0^1 \left((-1)^i (t-1)^{\mathbf{k}-i-1} - (-1)^{\mathbf{k}+1} \right) dt = (\mathbf{k}-i) \left(\frac{(-1)^{\mathbf{k}}}{i-\mathbf{k}} + (-1)^{\mathbf{k}} \right) = (-1)^{\mathbf{k}+1} + (\mathbf{k}-i) (-1)^{\mathbf{k}}$ and $\sum_{g=\mathbf{k}-i+1}^{\mathbf{k}-1} (-1)^{g+1} \frac{1}{\mathbf{k}-g+1} i \binom{i-1}{g-\mathbf{k}+i-1} = \int_0^1 \sum_{g=\mathbf{k}-i+1}^{\mathbf{k}-1} (-1)^{g+1} i \binom{i-1}{g-\mathbf{k}+i-1} t^{\mathbf{k}-g} dt = \int_0^1 \left(i (-1)^{\mathbf{k}-i} (t-1)^{i-1} - i (-1)^{\mathbf{k}+1} \right) dt = (-1)^{\mathbf{k}} (i-1)$. If $0 \leq j < \frac{\mathbf{k}+1-i}{2}$ and $i = \mathbf{k}$, $\psi_{\mathbf{k}} = 0$. If $\frac{\mathbf{k}+1-i}{2} \leq j \leq \frac{\mathbf{k}-1}{2}$ and $\frac{\mathbf{k}+1}{2} \leq i \leq \mathbf{k}-1$, the summed coefficient of $Q(p_1)^i Q(p_k)^{\mathbf{k}-i}$ is $(-1)^{\mathbf{k}-1} (\mathbf{k}-1) + \sum_{g=\mathbf{k}-i+1}^{\mathbf{k}-1} (-1)^{g+1} \frac{1}{\mathbf{k}-g+1} i \binom{i-1}{g-\mathbf{k}+i-1} + \sum_{g=i+1}^{\mathbf{k}-1} (-1)^{g+1} \frac{1}{\mathbf{k}-g+1} (\mathbf{k}-i) \binom{\mathbf{k}-i-1}{g-i-1}$, the same as above. If $i+j < \mathbf{k}$, since $\binom{i}{\mathbf{k}-j} = 0$, the related terms can be ignored, so, using the binomial theorem and beta function, the summed coefficient of $Q(p_1)^{\mathbf{k}-j} Q(p_k)^j$ is $\sum_{g=j+1}^{i+j} (-1)^{g+1} \frac{1}{\mathbf{k}-g+1} i \binom{i-1}{g-j-1} \binom{\mathbf{k}-i}{j} = i \binom{\mathbf{k}-i}{j} \int_0^1 \sum_{g=j+1}^{i+j} (-1)^{g+1} \binom{i-1}{g-j-1} t^{\mathbf{k}-g} dt = \binom{\mathbf{k}-i}{j} i \int_0^1 \left((-1)^j t^{\mathbf{k}-j-1} \left(\frac{t}{t-1} \right)^{i-1} \right) dt = \binom{\mathbf{k}-i}{j} i \frac{(-1)^{j+i+1} \Gamma(i) \Gamma(\mathbf{k}-j-i+1)}{\Gamma(\mathbf{k}-j+1)} = \frac{(-1)^{j+i+1} i! (\mathbf{k}-j-i)! (\mathbf{k}-i)!}{(\mathbf{k}-j)! j! (\mathbf{k}-j-i)!} = (-1)^{j+i+1} \frac{i! (\mathbf{k}-i)!}{\mathbf{k}!} \frac{\mathbf{k}!}{(\mathbf{k}-j)! j!} = \binom{\mathbf{k}}{i}^{-1} (-1)^{1+i} \binom{\mathbf{k}}{j} (-1)^j$.

According to the binomial theorem, the coefficient of $Q(p_1)^i Q(p_k)^{\mathbf{k}-i}$ in $\binom{\mathbf{k}}{i}^{-1} (-1)^{1+i} (Q(p_1) - Q(p_k))^{\mathbf{k}}$ is $\binom{\mathbf{k}}{i}^{-1} (-1)^{1+i} \binom{\mathbf{k}}{i} (-1)^{\mathbf{k}-i} = (-1)^{\mathbf{k}+1}$, same as the above summed coefficient of $Q(p_1)^i Q(p_k)^{\mathbf{k}-i}$, if $i+j = \mathbf{k}$. If $i+j < \mathbf{k}$, the coefficient of $Q(p_1)^{\mathbf{k}-j} Q(p_k)^j$ is

$\binom{k}{i}^{-1}(-1)^{1+i}\binom{k}{j}(-1)^j$, same as the corresponding summed coefficient of $Q(p_1)^{k-j}Q(p_k)^j$. Therefore, $\psi_k(x_1 = Q(p_1), \dots, x_i = Q(p_1), x_{i+1} = Q(p_k), \dots, x_k = Q(p_k)) = \binom{k}{i}^{-1}(-1)^{1+i}(Q(p_1) - Q(p_k))^k$, the maximum and minimum of ψ_k follow directly from the properties of the binomial coefficient.

□

The component quasi-distribution, ζ_Δ , is closely related to Ξ_Δ , which is the pairwise difference distribution, since $\sum_{\bar{\Delta} = -(\frac{k}{2}(-\Delta)^k)^{-1}(-\Delta)^k}^{\frac{1}{k}(-\Delta)^k} f_{\zeta_\Delta}(\bar{\Delta}) = f_{\Xi_\Delta}(\Delta)$. Recall that Theorem 1 established that

$f_{\Xi_\Delta}(\Delta)$ is monotonic increasing with a mode at zero if the original distribution is unimodal, $f_{\Xi_{-\Delta}}(-\Delta)$ is thus monotonic decreasing with a mode at zero. In general, if assuming the shape of ζ_Δ is uniform, Ξ_k is monotonic left and right around zero. The median of Ξ_k also exhibits a strong tendency to be close to zero, as it can be cast as a weighted mean of the medians of ζ_Δ . When $-\Delta$ is small, all values of ζ_Δ are close to zero, resulting in the median of ζ_Δ being close to zero as well. When $-\Delta$ is large, the median of ζ_Δ depends on its skewness, but the corresponding weight is much smaller, so even if ζ_Δ is highly skewed, the median of Ξ_k will only be slightly shifted from zero. Denote the median of Ξ_k as mkm , for the five parametric distributions here, $|mkm|$ s are all $\leq 0.1\sigma$ for Ξ_3 and Ξ_4 , where σ is the standard deviation of Ξ_k (SI Dataset S1). Assuming $mkm = 0$, for the even ordinal central moment kernel distribution, the average probability density on the left side of zero is greater than that on the right side, since $\frac{\frac{1}{2}}{\binom{k}{2}^{-1}(Q(0)-Q(1))^k} > \frac{\frac{1}{2}}{\frac{1}{k}(Q(0)-Q(1))^k}$. This means that, on average, the inequality $f(Q(\epsilon)) \geq f(Q(1-\epsilon))$ holds. For the odd ordinal distribution, the discussion is more challenging since it is generally symmetric. Just consider Ξ_3 , let $x_1 = Q(p_i)$ and $x_3 = Q(p_j)$, changing the value of x_2 from $Q(p_i)$ to $Q(p_j)$ will monotonically change the value of $\psi_3(x_1, x_2, x_3)$, since $\frac{\partial \psi_3(x_1, x_2, x_3)}{\partial x_2} = -\frac{x_1^2}{2} - x_1x_2 + 2x_1x_3 + x_2^2 - x_2x_3 - \frac{x_3^2}{2} - \frac{3}{4}(x_1 - x_3)^2 \leq \frac{\partial \psi_3(x_1, x_2, x_3)}{\partial x_2} \leq -\frac{1}{2}(x_1 - x_3)^2 \leq 0$. If the original distribution is right-skewed, ζ_Δ will be left-skewed, so, for Ξ_3 , the average probability density of the right side of zero will be greater than that of the left side, which means, on average, the inequality $f(Q(\epsilon)) \leq f(Q(1-\epsilon))$ holds. In all, the monotonic decreasing of the negative pairwise difference distribution guides the general shape of the k th central moment kernel distribution, $k > 2$, forcing it to be unimodal-like with the mode and median close to zero, then, the inequality $f(Q(\epsilon)) \leq f(Q(1-\epsilon))$ or $f(Q(\epsilon)) \geq f(Q(1-\epsilon))$ holds in general. If a distribution is ν th γ -ordered and all of its central moment kernel distributions are also ν th γ -ordered, it is called completely ν th γ -ordered.

Another crucial property of the central moment kernel distribution, location invariant, is introduced in the next theorem.

Theorem 3. $\psi_k(x_1 = \lambda x_1 + \mu, \dots, x_k = \lambda x_k + \mu) = \lambda^k \psi_k(x_1, \dots, x_k)$.

Proof. Recall that for the k th central moment, the kernel is $\psi_k(x_1, \dots, x_k) = \sum_{j=0}^{k-2} (-1)^j \binom{1}{k-j} \sum \left(x_{i_1}^{k-j} x_{i_2} \dots x_{i_{j+1}} \right) + (-1)^{k-1} (k-1) x_1 \dots x_k$, where the second summation is over $i_1, \dots, i_{j+1} = 1$ to k with $i_1 \neq i_2 \neq \dots \neq i_{j+1}$ and $i_2 < i_3 < \dots < i_{j+1}$ [11].

ψ_k consists of two parts. The first part, $\sum_{j=0}^{k-2} (-1)^j \binom{1}{k-j} \sum \left(x_{i_1}^{k-j} x_{i_2} \dots x_{i_{j+1}} \right)$, involves a double summation over certain terms. The second part, $(-1)^{k-1} (k-1) x_1 \dots x_k$, carries an alternating sign $(-1)^{k-1}$ and involves multiplication of the constant $k-1$ with the product of all the x variables, $x_1 x_2 \dots x_k$. Consider each multiplication cluster $(-1)^j \binom{1}{k-j} \sum \left(x_{i_1}^{k-j} x_{i_2} \dots x_{i_{j+1}} \right)$ for j ranging from 0 to $k-2$ in the first part. Let each cluster form a single group. The first part can be divided into $k-1$ groups. Combine this with the second part $(-1)^{k-1} (k-1) x_1 \dots x_k$. Together, the terms of ψ_k can be divided into a total of k groups. From the 1st to $k-1$ th group, the g th group has $\binom{k}{g} \binom{g}{1}$ terms having the form $(-1)^{g+1} \frac{1}{k-g+1} x_{i_1}^{k-g+1} x_{i_2} \dots x_{i_g}$. The final k th group is the term $(-1)^{k-1} (k-1) x_1 \dots x_k$.

There are two ways to divide ψ_k into k groups according to the form of each term. The first choice is, if $k \neq g$, the g th group of ψ_k has $\binom{k-l}{g-l}$

terms having the form $(-1)^{g+1} \frac{1}{k-g+1} x_{i_1}^{k-g+1} x_{i_2} \cdots x_{i_l} x_{i_{l+1}} \cdots x_{i_g}$, where $x_{i_1}, x_{i_2}, \dots, x_{i_l}$ are fixed, $x_{i_{l+1}}, \dots, x_{i_g}$ are selected such that $i_{l+1}, \dots, i_g \neq i_1, i_2, \dots, i_l$ and $i_{l+1} \neq \dots \neq i_g$. Define another function $\Psi_k(x_{i_1}, x_{i_2}, \dots, x_{i_l}, x_{i_{l+1}}, \dots, x_{i_g}) = (\lambda x_{i_1} + \mu)^{k-g+1} (\lambda x_{i_2} + \mu) \cdots (\lambda x_{i_l} + \mu) (\lambda x_{i_{l+1}} + \mu) \cdots (\lambda x_{i_g} + \mu)$, the first group of Ψ_k is $\lambda^k x_{i_1} \cdots x_{i_l} x_{i_{l+1}} \cdots x_{i_g}$, the h th group of Ψ_k , $h > 1$, has $\binom{k-g+1}{k-h-l+2}$ terms having the form $\lambda^{k-h+1} \mu^{h-1} x_{i_1}^{k-h-l+2} x_{i_2} \cdots x_{i_l}$. Transforming ψ_k by Ψ_k , then combining all terms with $\lambda^{k-h+1} \mu^{h-1} x_{i_1}^{k-h-l+2} x_{i_2} \cdots x_{i_l}$, $k-h-l+2 > 1$, the summed coefficient is $S1_l = \sum_{g=l}^{h+l-1} (-1)^{g+1} \frac{1}{k-g+1} \binom{k-g+1}{k-h-l+2} \binom{k-l}{g-l} = \sum_{g=l}^{h+l-1} (-1)^{g+1} \frac{(k-l)!}{(h+l-g-1)!(k-h-l+2)!(g-l)!} = 0$, since the summation is starting from l , ending at $h+l-1$, the first term includes the factor $g-l=0$, the final term includes the factor $h+l-g-1=0$, the terms in the middle are also zero due to the factorial property.

Another possible choice is the g th group of ψ_k has $(k-h) \binom{h-1}{g-k+h-1}$ terms having the form

$(-1)^{g+1} \frac{1}{k-g+1} x_{i_1} x_{i_2} \cdots x_{i_j}^{k-g+1} \cdots x_{i_{k-h+1}} x_{i_{k-h+2}} \cdots x_{i_g}$, provided that $k \neq g$, $2 \leq j \leq k-h+1$, where $x_{i_1}, \dots, x_{i_{k-h+1}}$ are fixed, $x_{i_j}^{k-g+1}$ and $x_{i_{k-h+2}}, \dots, x_{i_g}$ are selected such that $i_{k-h+2}, \dots, i_g \neq i_1, i_2, \dots, i_{k-h+1}$ and $i_{k-h+2} \neq \dots \neq i_g$. Transforming these terms by $\Psi_k(x_{i_1}, x_{i_2}, \dots, x_{i_j}, \dots, x_{i_{k-h+1}}, x_{i_{k-h+2}}, \dots, x_{i_g}) = (\lambda x_{i_1} + \mu) (\lambda x_{i_2} + \mu) \cdots (\lambda x_{i_j} + \mu)^{k-g+1} \cdots (\lambda x_{i_{k-h+1}} + \mu) (\lambda x_{i_{k-h+2}} + \mu) \cdots (\lambda x_{i_g} + \mu)$, then there are $k-g+1$ terms having the form $\lambda^{k-h+1} \mu^{h-1} x_{i_1} x_{i_2} \cdots x_{i_{k-h+1}}$. Transforming the final k th group of ψ_k by $\Psi_k(x_1, \dots, x_k) = (\lambda x_1 + \mu) \cdots (\lambda x_k + \mu)$, then, there is one term having the form $(-1)^{k-1} (k-1) \lambda^{k-h+1} \mu^{h-1} x_1 x_2 \cdots x_{k-h+1}$. Another possible combination is that the g th group of ψ_k contains $(g-k+h-1) \binom{h-1}{g-k+h-1}$ terms having the form $(-1)^{g+1} \frac{1}{k-g+1} x_{i_1} x_{i_2} \cdots x_{i_{k-h+1}} x_{i_{k-h+2}} \cdots x_{i_j}^{k-g+1} \cdots x_{i_g}$. Transforming these terms by $\Psi_k(x_{i_1}, x_{i_2}, \dots, x_{i_{k-h+1}}, x_{i_{k-h+2}}, \dots, x_{i_j}, \dots, x_{i_g}) =$

$(\lambda x_{i_1} + \mu) (\lambda x_{i_2} + \mu) \cdots (\lambda x_{i_{k-h+1}} + \mu) (\lambda x_{i_{k-h+2}} + \mu) \cdots (\lambda x_{i_j} + \mu)^{k-g+1} \cdots (\lambda x_{i_g} + \mu)$, then there is only one term having the form $\lambda^{k-h+1} \mu^{h-1} x_{i_1} x_{i_2} \cdots x_{i_{k-h+1}}$. The above summation $S1_l$ should also be included, i.e., $x_{i_1}^{k-h-l+2} = x_{i_1}$, $k = h+l-1$. So, combining all terms with $\lambda^{k-h+1} \mu^{h-1} x_{i_1} x_{i_2} \cdots x_{i_{k-h+1}}$, according to the binomial theorem, the summed coefficient is $S2_l = \sum_{g=k-h+1}^{k-1} (-1)^{g+1} \binom{h-1}{g-k+h-1} \left(k-h+1 + \frac{g-k+h-1}{k-g+1} \right) + (-1)^{k-1} (k-1) = (k-h+1) \sum_{g=k-h+1}^{k-1} (-1)^{g+1} \binom{h-1}{g-k+h-1} + \sum_{g=k-h+1}^{k-1} (-1)^{g+1} \binom{h-1}{g-k+h-1} \left(\frac{g-k+h-1}{k-g+1} \right) + (-1)^{k-1} (k-1) = (-1)^k (k-h+1) + (h-2)(-1)^k + (-1)^{k-1} (k-1) = 0$. The summation identities required are $\sum_{g=k-h+1}^{k-1} (-1)^{g+1} \binom{h-1}{g-k+h-1} = (-1)^k$ and $\sum_{g=k-h+1}^{k-1} (-1)^{g+1} \binom{h-1}{g-k+h-1} \left(\frac{g-k+h-1}{k-g+1} \right) = (h-2)(-1)^k$. These two summation identities are proven in Lemma 4 and 5 in the SI Text.

Thus, no matter in which way, all terms including μ can be canceled out. The proof is complete by noticing that the remaining part is $\lambda^k \psi_k(x_1, \dots, x_k)$.

□

A direct result of Theorem 3 is that, $\text{WHL}k_m$ after standardization is invariant to location and scale. So, the weighted H-L standardized k th moment is defined to be

$$\text{WHL}skm_{e=\min(\epsilon_1, \epsilon_2), k_1, k_2, \gamma_1, \gamma_2, n} := \frac{\text{WHL}k_{m_{k_1, \epsilon_1, \gamma_1, n}}}{(\text{WHL}var_{k_2, \epsilon_2, \gamma_2, n})^{k/2}}.$$

To avoid confusion, it should be noted that the robust location estimations of the kernel distributions discussed in this paper differ from the approach taken by Joly and Lugosi (2016) [23], which is computing the median of all U -statistics from different disjoint blocks. Compared to bootstrap

median U -statistics, this approach can produce two additional kinds of finite sample bias, one arises from the limited numbers of blocks, another is due to the size of the U -statistics (consider the mean of all U -statistics from different disjoint blocks, it is definitely not identical to the original U -statistic, except when the kernel is the Hodges-Lehmann kernel). Laforgue, Clemencon, and Bertail (2019)'s median of randomized U -statistics [24] is more sophisticated and can overcome the limitation of the number of blocks, but the second kind of bias remains unsolved.

Congruent Distribution

In the realm of nonparametric statistics, the relative differences, or orders, of robust estimators are of primary importance. A key implication of this principle is that when there is a shift in the parameters of the underlying distribution, all nonparametric estimates should asymptotically change in the same direction, if they are estimating the same attribute of the distribution. If, on the other hand, the mean suggests an increase in the location of the distribution while the median indicates a decrease, a contradiction arises. It is worth noting that such contradiction is not possible for any LL -statistics in a location-scale distribution, as explained in Theorem 2 and 18 in REDS I. However, it is possible to construct counterexamples to the aforementioned implication in a shape-scale distribution. In the case of the Weibull distribution, its quantile function is $Q_{Wei}(p) = \lambda(-\ln(1-p))^{1/\alpha}$, where $0 \leq p \leq 1$, $\alpha > 0$, $\lambda > 0$, λ is a scale parameter, α is a shape parameter, $\ln$ is the natural logarithm function. Then, $m = \lambda \sqrt[\alpha]{\ln(2)}$, $\mu = \lambda \Gamma\left(1 + \frac{1}{\alpha}\right)$, where Γ is the gamma function. When $\alpha = 1$, $m = \lambda \ln(2) \approx 0.693\lambda$, $\mu = \lambda$, when $\alpha = \frac{1}{2}$, $m = \lambda \ln^2(2) \approx 0.480\lambda$, $\mu = 2\lambda$, the mean increases as α changes from 1 to $\frac{1}{2}$, but the median decreases. In the last section, the fundamental role of quantile average was demonstrated by using the method of classifying distributions through the signs of derivatives. To avoid such scenarios, this method can also be used. Let the quantile average function of a parametric distribution be denoted as $QA(\epsilon, \gamma, \alpha_1, \dots, \alpha_i, \dots, \alpha_k)$, where α_i represent the parameters of the distribution, then, a distribution is γ -congruent if and only if the sign of $\frac{\partial QA}{\partial \alpha_i}$ remains the same for all $0 \leq \epsilon \leq \frac{1}{1+\gamma}$. If $\frac{\partial QA}{\partial \alpha_i}$ is equal to zero or undefined, it can be considered both positive and negative, and thus does not impact the analysis. A distribution is completely γ -congruent if and only if it is γ -congruent and all its central moment kernel distributions are also γ -congruent. Setting $\gamma = 1$ constitutes the definitions of congruence and complete congruence. Replacing the QA with $\gamma mHLM$ gives the definition of γ - U -congruence. Chebyshev's inequality implies that, for any probability distributions with finite second moments, as the parameters change, even if some LL -statistics change in a direction different from that of the population mean, the magnitude of the changes in the LL -statistics remains bounded compared to the changes in the population mean. Furthermore, distributions with infinite moments can be γ -congruent, since the definition is based on the quantile average, not the population mean.

The following theorems show the conditions that a distribution is congruent or γ -congruent.

Theorem 4. *A symmetric distribution is always congruent and U -congruent.*

Proof. As shown in Theorem 2 and Theorem 18 in REDS I, for any symmetric distribution, all quantile averages and all $\gamma mHLM$ s coincide. The conclusion follows immediately. $\square$

Theorem 5. *A positive definite location-scale distribution is always γ -congruent.*

Proof. As shown in Theorem 2, for a location-scale distribution, any quantile average can be expressed as $\lambda QA_0(\epsilon, \gamma) + \mu$. Therefore, the derivatives with respect to the parameters λ or μ are always positive. By application of the definition, the desired outcome is obtained. $\square$

For the Pareto distribution, $\frac{\partial Q}{\partial \alpha} = \frac{x_m(1-p)^{-1/\alpha} \ln(1-p)}{\alpha^2}$. Since $\ln(1-p) < 0$ for all $0 < p < 1$, $(1-p)^{-1/\alpha} > 0$ for all $0 < p < 1$ and $\alpha > 0$, so $\frac{\partial Q}{\partial \alpha} < 0$, and therefore $\frac{\partial QA}{\partial \alpha} < 0$, the Pareto distribution is γ -congruent. It is also γ - U -congruent, since $\gamma mHLM$ can also express as a function of $Q(p)$.

For the lognormal distribution, $\frac{\partial Q_A}{\partial \sigma} = \frac{1}{2} \left(\sqrt{2} \operatorname{erfc}^{-1}(2\gamma\epsilon) \left(-e^{\frac{\sqrt{2}\mu - 2\sigma \operatorname{erfc}^{-1}(2\gamma\epsilon)}{\sqrt{2}}} \right) + \left(-\sqrt{2} \right) \operatorname{erfc}^{-1}(2(1 - \epsilon)) e^{\frac{\sqrt{2}\mu - 2\sigma \operatorname{erfc}^{-1}(2(1-\epsilon))}{\sqrt{2}}} \right)$. Since the inverse complementary error function is positive when the input is smaller than 1, and negative when the input is larger than 1, and symmetry around 1, if $0 \leq \gamma \leq 1$, $\operatorname{erfc}^{-1}(2\gamma\epsilon) \geq -\operatorname{erfc}^{-1}(2 - 2\epsilon)$, $e^{\mu - \sqrt{2}\sigma \operatorname{erfc}^{-1}(2 - 2\epsilon)} > e^{\mu - \sqrt{2}\sigma \operatorname{erfc}^{-1}(2\gamma\epsilon)}$. Therefore, if $0 \leq \gamma \leq 1$, $\frac{\partial Q_A}{\partial \sigma} > 0$, the lognormal distribution is γ -congruent. Theorem 4 implies that the generalized Gaussian distribution is congruent and U -congruent. For the Weibull distribution, when α changes from 1 to $\frac{1}{2}$, the average probability density on the left side of the median increases, since $\frac{\frac{1}{2}}{\lambda \ln(2)} < \frac{\frac{1}{2}}{\lambda \ln^2(2)}$, but the mean increases, indicating that the distribution is more heavy-tailed, the probability density of large values will also increase. So, the reason for non-congruence of the Weibull distribution lies in the simultaneous increase of probability densities on two opposite sides as the shape parameter changes: one approaching the bound zero and the other approaching infinity. Note that the gamma distribution does not have this issue, Numerical results indicate that it is likely to be congruent.

The next theorem shows an interesting relation between congruence and the central moment kernel distribution.

Theorem 6. *The second central moment kernel distribution derived from a continuous location-scale unimodal distribution is always γ -congruent.*

Proof. Theorem 3 shows that the central moment kernel distribution generated from a location-scale distribution is also a location-scale distribution. Theorem 1 shows that it is positively definite. Implementing Theorem 12 in REDS 1 yields the desired result. $\square$

Although some parametric distributions are not congruent, as shown in REDS 1. In REDS 1, Theorem 12 establishes that γ -congruence always holds for a positive definite location-scale family distribution and thus for the second central moment kernel distribution generated from a location-scale unimodal distribution as shown in Theorem 6. Theorem 2 demonstrates that all central moment kernel distributions are unimodal-like with mode and median close to zero, as long as they are generated from unimodal distributions. Assuming finite moments and constant $Q(0) - Q(1)$, increasing the mean of a distribution will result in a generally more heavy-tailed distribution, i.e., the probability density of the values close to $Q(1)$ increases, since the total probability density is 1. In the case of the k th central moment kernel distribution, $k > 2$, while the total probability density on either side of zero remains generally constant as the median is generally close to zero and much less impacted by increasing the mean, the probability density of the values close to zero decreases as the mean increases. This transformation will increase nearly all symmetric weighted averages, in the general sense. Therefore, except for the median, which is assumed to be zero, nearly all symmetric weighted averages for all central moment kernel distributions derived from unimodal distributions should change in the same direction when the parameters change.

Discussion

Moments, including raw moments, central moments, and standardized moments, are the most common parameters that describe probability distributions. Central moments are preferred over raw moments because they are invariant to translation. In 1947, Hsu and Robbins proved that the arithmetic mean converges completely to the population mean provided the second moment is finite [25]. The strong law of large numbers (proven by Kolmogorov in 1933) [26] implies that the k th sample central moment is asymptotically unbiased. Recently, fascinating statistical phenomena regarding Taylor's law for distributions with infinite moments have been discovered by Drton and Xiao (2016) [27], Pillai and Meng (2016) [28], Cohen, Davis, and Samorodnitsky (2020) [29], and Brown, Cohen, Tang, and Yam (2021) [30]. Lindquist and Rachev (2021) raised a critical question in their inspiring comment

to Brown et al's paper [30]: "What are the proper measures for the location, spread, asymmetry, and dependence (association) for random samples with infinite mean?" [31]. From a different perspective, this question closely aligns with the essence of Bickel and Lehmann's open question in 1979 [4]. They suggested using median, interquartile range, and medcouple [32] as the robust versions of the first three moments. While answering this question is not the focus of this paper, it is almost certain that the estimators proposed in this paper will have a place. Since the estimation of central moments can be transformed into the location estimation of a pseudosample, according to the general principle of central limit theorem, the optimal estimator should always has a combinatorial pseudosample size, which explains, in another aspect, why the theory of U -statistics allows a minimum-variance unbiased estimator to be derived from each unbiased estimator of an estimable parameter. Similar to the robust version of L-moment [33] being trimmed L-moment [17], central moments now also have their robust nonparametric version, weighted Hodges-Lehmann central moments, based on the complete U -congruence of the underlying distribution.

Author Contributions: T.L. designed research, performed research, analyzed data, and wrote the paper.

Conflicts of Interest: The author declares no competing interest.

Software Availability: The codes used to compute the weighted H-L k th central moment have been deposited in [GitHub](#).

References

1. FR Hampel, The influence curve and its role in robust estimation. *Journal of the american statistical association* **69**, 383–393 (1974).
2. CF Gauss, Bestimmung der genauigkeit der beobachtungen. *Ibidem* pp. 129–138 (1816).
3. PJ Bickel, EL Lehmann, Descriptive statistics for nonparametric models. iii. dispersion in *Selected works of EL Lehmann*. (Springer), pp. 499–518 (2012).
4. PJ Bickel, EL Lehmann, Descriptive statistics for nonparametric models iv. spread in *Selected Works of EL Lehmann*. (Springer), pp. 519–526 (2012).
5. H Oja, On location, scale, skewness and kurtosis of univariate distributions. *Scandinavian Journal of statistics* pp. 154–168 (1981).
6. H Oja, Descriptive statistics for multivariate distributions. *Statistics & Probability Letters* **1**, 327–332 (1983).
7. PJ Bickel, EL Lehmann, Descriptive statistics for nonparametric models ii. location in *selected works of EL Lehmann*. (Springer), pp. 473–497 (2012).
8. W van Zwet, Convex transformations: A new approach to skewness and kurtosis in *Selected Works of Willem van Zwet*. (Springer), pp. 3–11 (2012).
9. PJ Rousseeuw, C Croux, Alternatives to the median absolute deviation. *Journal of the American Statistical association* **88**, 1273–1283 (1993).
10. T Li, Robust estimations from distribution structures: Mean (2023).
11. PM Heffernan, Unbiased estimation of central moments by using u -statistics. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)* **59**, 861–863 (1997).
12. J Hodges, E Lehmann, Matching in paired comparisons. *The Annals of Mathematical Statistics* **25**, 787–791 (1954).
13. AL Bowley, *Elements of statistics*. (King) No. 8, (1926).
14. WR van Zwet, *Convex Transformations of Random Variables: Nebst Stellingen*. (1964).
15. RA Groeneveld, G Meeden, Measuring skewness and kurtosis. *Journal of the Royal Statistical Society: Series D (The Statistician)* **33**, 391–399 (1984).
16. J SAW, Moments of sample moments of censored samples from a normal population. *Biometrika* **45**, 211–221 (1958).
17. EA Elamir, AH Seheult, Trimmed l-moments. *Computational Statistics & Data Analysis* **43**, 299–314 (2003).
18. RA Fisher, Moments and product moments of sampling distributions. *Proceedings of the London Mathematical Society* **2**, 199–238 (1930).
19. PR Halmos, The theory of unbiased estimation. *The Annals of Mathematical Statistics* **17**, 34–43 (1946).

20. W Hoeffding, A class of statistics with asymptotically normal distribution. *The Annals of Mathematical Statistics* **19**, 293–325 (1948).
21. RJ Serfling, Generalized l-, m-, and r-statistics. *The Annals of Statistics* **12**, 76–86 (1984).
22. T Li, Robust estimations from distribution structures: Invariant moments. *Zenodo* (2023).
23. E Joly, G Lugosi, Robust estimation of u-statistics. *Stochastic Processes and their Applications* **126**, 3760–3773 (2016).
24. P Laforgue, S Cléménçon, P Bertail, On medians of (randomized) pairwise means in *International Conference on Machine Learning*. (PMLR), pp. 1272–1281 (2019).
25. PL Hsu, H Robbins, Complete convergence and the law of large numbers. *Proceedings of the national academy of sciences* **33**, 25–31 (1947).
26. A Kolmogorov, Sulla determinazione empirica di una lgge di distribuzione. *Inst. Ital. Attuari, Giorn.* **4**, 83–91 (1933).
27. M Drton, H Xiao, Wald tests of singular hypotheses. *Bernoulli* **22**, 38–59 (2016).
28. NS Pillai, XL Meng, An unexpected encounter with cauchy and lévy. *The Annals of Statistics* **44**, 2089–2097 (2016).
29. JE Cohen, RA Davis, G Samorodnitsky, Heavy-tailed distributions, correlations, kurtosis and taylor’s law of fluctuation scaling. *Proceedings of the Royal Society A* **476**, 20200610 (2020).
30. M Brown, JE Cohen, CF Tang, SCP Yam, Taylor’s law of fluctuation scaling for semivariances and higher moments of heavy-tailed data. *Proceedings of the National Academy of Sciences* **118**, e2108031118 (2021).
31. WB Lindquist, ST Rachev, Taylor’s law and heavy-tailed distributions. *Proceedings of the National Academy of Sciences* **118**, e2118893118 (2021).
32. G Brys, M Hubert, A Struyf, A robust measure of skewness. *Journal of Computational and Graphical Statistics* **13**, 996–1017 (2004).
33. JR Hosking, L-moments: Analysis and estimation of distributions using linear combinations of order statistics. *Journal of the Royal Statistical Society: Series B (Methodological)* **52**, 105–124 (1990).

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.