

ESTIMATING THE SKEWNESS PARAMETER OF A LÉVY-STABLE

HADDA SAIDANE¹, DJAMEL MERAGHNI^{2*} and LOUIZA SOLTANE³

ABSTRACT. Stable distributions consist in probability models that are described by four parameters for the shape, location, scale and skewness. These parameters were estimated by several classical methods. Recently, the extreme value theory were exploited to propose estimators to the first three. In this paper, we use the same technique to introduce an estimator to the skewness parameter. We carry out a simulation study to evaluate its performance by computing estimation biases and mean squared errors and providing bootstrap confidence bounds as well.

Keywords. Extreme values, Hill estimator, Skewness parameter, Stable distribution

2020 Mathematics Subject Classification. Primary 46L55; Secondary 44B20

1. INTRODUCTION

The class of Lévy-stable laws was introduced in the early 1920's by Paul Lévy [6] when he was interested in the behavior of sums of independent random variables. In the last century, Mandelbrot [7] studied stock market fluctuations, for which it was quite clear that the Gaussian model was not suitable. He concluded that the Lévy-stable law (also called stable Paretian, alpha-stable or simply stable) realistically describes the variation of prices in stock markets and Fama [2] would validate this conclusion a bit later. These probability distributions have many important mathematical properties and statistical applications as they allow skewness and tail thickness. They represent a good candidate to fit skewed heavy-tailed data and provide reliable risk measures (based on the distribution tails) in several real-life case studies such as finance, environmental science, signal processing,... see for instance, [7] and [2]. In addition to this empirical evidence, there is a theoretical strength to stable modeling, namely the generalized central limit theorem which says that stable laws are the only possible limit distributions for sums of independent and identically distributed (iid) random variables (rv's), see, for instance [10, page 20].

Date: Received: Aug 26, 2025; Accepted: Sep 20, 2025.

* Corresponding author

© The Author(s) 2025. This article is licensed under a Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License. To view a copy of the licence, visit <https://creativecommons.org/licenses/by-nc-nd/4.0/>.

Denoted, following Samorodnitsky and Taqqu [13], by $S_\alpha(\sigma, \beta, \mu)$ this class is characterized by four parameters, namely the shape or stability parameter $0 < \alpha \leq 2$, the location parameter $-\infty < \mu < \infty$, the scale parameter $\sigma > 0$ and the skewness parameter $-1 \leq \beta \leq 1$. The case $\alpha = 2$ corresponds to the Gaussian distribution. For $\alpha < 2$, the variance of a rv belonging to $S_\alpha(\sigma, \beta, \mu)$ is infinite whereas for $0 < \alpha \leq 1$, the expectation does not exist. When $\alpha > 1$, the location parameter μ is equal to the distribution mean which is finite in this case. More generally, the k th moment of a Lévy-stable rv is finite if and only if $k < \alpha$. In other words, if X is a rv with probability law $S_\alpha(\sigma, \beta, \mu)$, then, for $k > 0$, we have

$$E(|X|^k) : \begin{cases} < \infty, & \text{for } 0 < k < \alpha, \\ = \infty, & \text{for } k \geq \alpha. \end{cases}$$

The major drawback of the stable class $S_\alpha(\sigma, \beta, \mu)$ is that it does not have closed explicit forms for its probability density and cumulative distribution functions (pdf and cdf) except in three cases, namely the distributions of Gauss $S_2(\sigma, 0, \mu)$, Cauchy $S_1(\sigma, 0, \mu)$ and Lévy $S_{0.5}(\sigma, 1, \mu)$. The most common way to describe the stable distribution is through its characteristic function with many representations, among which the most famous is defined, for $t \in \mathbb{R}$, by

$$\varphi(t) = \begin{cases} \exp \left\{ i\mu t - \sigma^\alpha |t|^\alpha \left[1 - i\beta \operatorname{sign}(t) \tan \left(\frac{\pi\alpha}{2} \right) \right] \right\}, & \alpha \neq 1, \\ \exp \left\{ i\mu t - \sigma |t| \left[1 + i\beta \operatorname{sign}(t) \frac{2}{\pi} \ln |t| \right] \right\}, & \alpha = 1, \end{cases}$$

where

$$i^2 = -1 \text{ and } \operatorname{sign}(t) := \begin{cases} 1 & \text{if } t > 0, \\ 0 & \text{if } t = 0, \\ -1 & \text{if } t < 0. \end{cases}$$

The lack of analytic expression for the pdf severely hampers the very challenging task of estimating the stable parameters. However, several useful numerical procedures, which are more or less accurate and more or less fast, have been proposed via different approaches. Indeed, estimates are calculated on the basis of sample quantiles, sample cdf and maximum likelihood. A detailed inventory of these methods may be found in, for instance, [3].

For a complete description with full details on this class of distributions, one refers to, for instance, [13, 15] and [10]. The rest of the paper is organized as follows. Section 2 is devoted to a reminder of the estimators of the shape, location and scale parameters, which were constructed via the extreme value theory (EVT) approach. In Section 3, we define our estimator to the skewness parameter β . In Section 4, we carry out an intensive simulation study to evaluate the performance of the proposed estimator, with respect to (absolute) bias and mean squared error (mse), and provide bootstrap confidence intervals as well. Finally, the simulation results are gathered in Section 5.

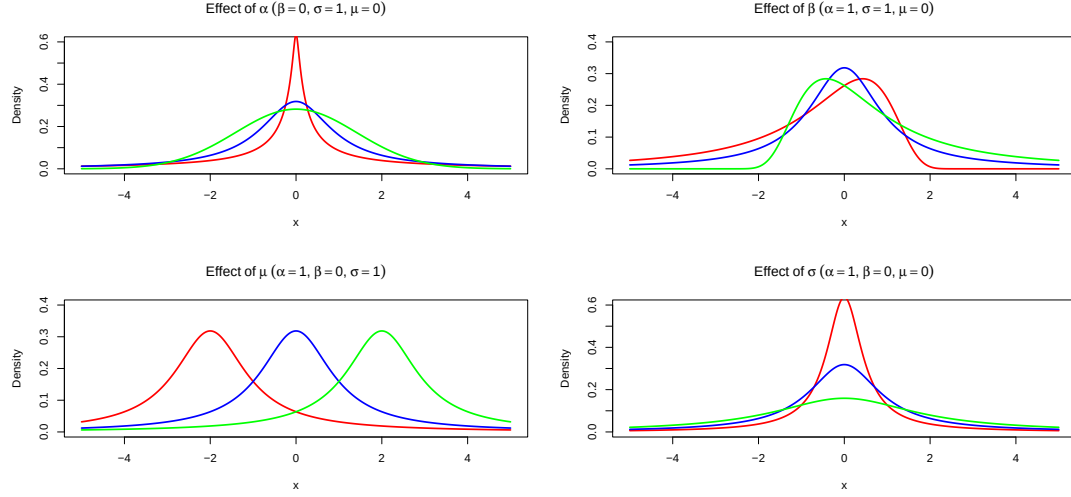


FIGURE 1. Probability density function of the Lévy-stable distribution with three different values for each one of the four parameters.

2. EVT-BASED ESTIMATION

Given the heavy-tailed nature of the distribution, EVT offers a valuable alternative to these classical procedures. Indeed, EVT does not require a parametric assumption for the entire distribution, it only focuses on tail behavior. Moreover, it provides explicit-form estimators with desirable properties such as asymptotic normality. For the very crucial shape parameter α , there are a variety of EVT-based estimators including the most widely used Hill estimator [4]:

$$\hat{\alpha}_n := \left(\frac{1}{k} \sum_{i=1}^k \log^+(Z_{n-i+1,n}) - \log^+(Z_{n-k,n}) \right)^{-1}, \quad (2.1)$$

where $Z_{1,n} \leq \dots \leq Z_{n,n}$ are the order statistics pertaining to a sample $Z_1, \dots, Z_n$, $n \geq 1$, from the rv $Z = |X|$ and $k = k_n$ is an integer sequence satisfying $k \rightarrow \infty$ and $k/n \rightarrow 0$ as $n \rightarrow \infty$. The notation $\log^+(x)$ stands for $\log(x \vee 1)$. For a discussion on this estimator, see [15] and those interested in an overview on estimating the extreme value index (in general) are referred to, for instance, [1].

For $1 < \alpha < 2$, Peng (2001) [11] introduced an asymptotically normal estimator $\hat{\mu}$ to the location parameter μ :

$$\hat{\mu}_n := \hat{\mu}_n^{(-)} + \tilde{\mu}_n + \hat{\mu}_n^{(+)},$$

where

$$\hat{\mu}_n^{(-)} := \frac{kX_{k,n}}{n} \frac{\hat{\alpha}_n^{(-)}}{\hat{\alpha}_n^{(-)} - 1}, \quad \tilde{\mu}_n := \frac{1}{n} \sum_{i=k+1}^{n-k} X_{i,n} \quad \text{and} \quad \hat{\mu}_n^{(+)} := \frac{kX_{n-k+1,n}}{n} \frac{\hat{\alpha}_n^{(+)}}{\hat{\alpha}_n^{(+)} - 1},$$

with

$$\hat{\alpha}_n^{(-)} := \left(\frac{1}{k} \sum_{i=1}^k \log^+(-X_{i,n}) - \log^+(-X_{k,n}) \right)^{-1},$$

and

$$\hat{\alpha}_n^{(+)} := \left(\frac{1}{k} \sum_{i=1}^k \log^+(X_{n-i+1,n}) - \log^+(X_{n-k,n}) \right)^{-1},$$

being two consistent versions of Hill's estimator of α where $X_{1,n} \leq \dots \leq X_{n,n}$ denote the order statistics pertaining to a sample $X_1, \dots, X_n$ from the rv X .

For the scale parameter σ , Meraghni and Necir (2007) [8] proposed an estimator $\hat{\sigma}_n$ enjoying an asymptotic normality property:

$$\hat{\sigma}_n := Z_{n-k,n} \left(\frac{k\pi}{2n\Gamma(\hat{\alpha}) \sin(\pi\hat{\alpha}/2)} \right)^{1/\hat{\alpha}_n}. \quad (2.2)$$

In his paper, Weron [15] has discussed the performance of Hill's estimator $\hat{\alpha}_n$. He has arrived to the conclusion that for small values of α , the estimation is quite reasonable but as α gets larger, there is a significant overestimation when considering samples of typical size. For such values of α , very large samples are needed in order to obtain acceptable estimates and avoid misleading inference, because the true tail behavior of stable distributions is only visible for extremely large datasets. Fortunately, this kind of data have recently become available in practice and their treatment has been made possible by a very sophisticated technology. The same remark could be said about the other two estimators $\hat{\mu}_n$ and $\hat{\sigma}_n$ and the one that we will introduce in the next section, as they fundamentally depend on $\hat{\alpha}_n$.

3. DEFINING THE SKEWNESS PARAMETER ESTIMATOR

The upper and lower tails of a Lévy-distribution asymptotically exhibit a Pareto-like behavior (they fall off like a power function). The rate of decay is governed by the characteristic index α : the smaller the latter, the slower the decay and so the heavier the tails (see the top left panel in Figure 1). As the value of α approaches 2, the impact of the skewness parameter β becomes negligible, causing the distribution to converge toward the normal distribution. More precisely, the following results hold for $0 < \alpha < 2$ (see, for instance, [10, page 13]).

$$1 - F(x) := P(X > x) \sim (1 + \beta) C_\alpha \sigma^\alpha x^{-\alpha}, \text{ as } x \rightarrow \infty, \quad (3.1)$$

and

$$F(-x) := P(X < -x) \sim (1 - \beta) C_\alpha \sigma^\alpha x^{-\alpha}, \text{ as } x \rightarrow \infty, \quad (3.2)$$

where

$$C_\alpha := \left(\int_0^\infty x^{-\alpha} \sin x dx \right)^{-1} = \frac{\Gamma(\alpha)}{\pi} \sin \frac{\pi\alpha}{2},$$

with Γ being the gamma function defined, for $\theta > 0$, by $\Gamma(\theta) = \int_0^\infty x^{\theta-1} e^{-x} dx$.

The statement $u(x) \sim v(x)$, as $x \rightarrow x_0$ means $\lim_{x \rightarrow x_0} u(x)/v(x) = 1$. Combining the tail approximations (3.1) and (3.2) yields what is called balance conditions,

namely, as $x \rightarrow \infty$, we have

$$\frac{P(X > x)}{P(|X| > x)} \rightarrow \frac{1 + \beta}{2} \text{ and } \frac{P(X < -x)}{P(|X| > x)} \rightarrow \frac{1 - \beta}{2} \text{ as } x \rightarrow \infty.$$

From the relation (3.1), we have

$$\beta \sim \left(\frac{x}{\sigma}\right)^\alpha \frac{\pi}{\Gamma(\alpha) \sin(\pi\alpha/2)} P(X > x) - 1, \text{ as } x \rightarrow \infty,$$

We set $s = P(X > x)$, then $x = F^{\leftarrow}(1-s)$, where $F^{\leftarrow}(s) := \inf\{x : F(x) \geq s, 0 < s < 1\}$ is the generalized inverse or quantile function of cdf F . Thus, we have

$$\beta \sim s \left(\frac{F^{\leftarrow}(1-s)}{\sigma}\right)^\alpha \frac{\pi}{\Gamma(\alpha) \sin(\pi\alpha/2)} - 1, \text{ as } s \rightarrow 0.$$

To construct an estimator $\hat{\beta}_n$ to the skewness parameter β , we take $s = k/n$, which tends to zero as n goes to infinity, and we replace $F^{\leftarrow}(1 - k/n)$ by its empirical counterpart $X_{n-k,n}$. Hence, we get

$$\hat{\beta}_n = \left(\frac{X_{n-k,n}}{\hat{\sigma}_n}\right)^{\hat{\alpha}_n} \left(\frac{k\pi}{n\Gamma(\hat{\alpha}_n) \sin(\pi\hat{\alpha}_n/2)}\right) - 1, \quad (3.3)$$

where $\hat{\alpha}_n$ and $\hat{\sigma}_n$ are those defined in (2.1) and (2.2) respectively.

This newly proposed estimator $\hat{\beta}_n$ has the advantage to have an explicit formula unlike the classical estimators of the skewness parameter which were obtained through approximation and optimization numerical procedures (sometimes very complicated). This is valid for the other three EVT-based estimators $\hat{\alpha}_n$, $\hat{\mu}_n$ and $\hat{\sigma}_n$.

4. SIMULATION STUDY

In this section, we use the statistical software R [5] to evaluate the performance and finite sample behavior of the newly proposed estimator $\hat{\beta}_n$ through simulated stable datasets with different values of the shape, scale and skewness parameters $\alpha \in \{0.25, 0.5, 0.75, 1, 1.2\}$, $\sigma \in \{0.1, 0.5, 1\}$ and $\beta \in \{-0.5, 0, 0.5, 0.8\}$. The location parameter is taken to be zero. For each experiment, we generate 2000 samples of size $n = 1000$. We compute the estimate β -value $\hat{\beta}_n$ and the corresponding absolute bias (bias) and mean squared error (mse). We also construct 95%-confidence intervals (conf int) for β by the bootstrap percentile interval method (see, for instance, [14, page 34]) which is implemented in the software R and we assess their precision by their lengths and coverage probabilities (cov prob). Our final results, obtained by averaging over all replications, are graphically illustrated in Figures 2 and 3 (where we plot the estimator $\hat{\beta}_n$ versus the number k of top statistics) and numerically summarized in Tables 1, 2, 3 and 4, all gathered in the appendix.

It is clear that EVT-based estimators of the tail index and related quantities essentially rely on the number k of top statistics and so their behaviors are affected by this number. Indeed, using too many data, in the estimation process, results in a substantial bias whereas with too few observations one gets a large variance. In

the literature, there exist several approaches to treat the thorny issue of selecting the optimal k -value, called optimal sample fraction (which locates where the distribution tail really begins), in order to guarantee the best possible estimates. In addition to prior graphical exploration, the idea is to analytically minimize the asymptotic mean squared error (amse) which is equal to the sum of the squared asymptotic bias and the asymptotic variance. The k -value where the amse reaches its minimum represents an optimal choice for the number of largest statistics that should be used in the estimation. To this end, there are many research works proposing numerical adaptive procedures, among which we can cite the algorithm of Reiss and Thomas (see [12, page 121]) that we use in our simulation study. If the tail index to be estimated is γ (for $S_\alpha(\sigma, \beta, \mu)$, we have $\gamma = 1/\alpha$) then Reiss and Thomas define the optimal number of extremes needed in estimate computation, by

$$k_{opt} := \arg \min_{1 \leq k \leq n} \frac{1}{k} \sum_{i=1}^k i^\delta |\hat{\gamma}_i - \text{median} \{\hat{\gamma}_1, \dots, \hat{\gamma}_k\}|, \quad 0 \leq \delta \leq 1/2,$$

where $\hat{\gamma}_i$ is the γ -estimator based on the i upper order statistics. In this study, we set $\delta = 0.3$, as it provides better performance in terms of bias and mse (see [9]).

Discussion

In the light of the simulation results, we may conclude the following:

In both Figures 2 and 3, we notice, as expected, significant volatility for small values of k . This is characteristic of EVT-based estimation in general.

Figure 3 indicates that $\hat{\beta}_n$ trajectories for both scale values ($\sigma = 0.1$ and $\sigma = 1$) are broadly similar in shape meaning that the changes in scale do not significantly distort the convergence pattern, whereas Figure 2 shows that the behavior of $\hat{\beta}_n$ is influenced by the value of the stability index α . Indeed, the estimator of β exhibits better long-run performance for $\alpha = 0.25$ than for $\alpha = 0.75$. In both panels of Figure 2, the black curves ($\alpha = 0.25$) stabilize around the red line (true value of β) more consistently than the blue curves ($\alpha = 0.75$).

The note above is confirmed by the numerical results in all four tables. When α is small (e.g., $\alpha = 0.25$ or $\alpha = 0.5$), the estimator achieves good accuracy, characterized by low absolute bias and mse. Additionally, the corresponding confidence intervals exhibit high coverage probabilities and relatively short lengths.

Finally, it is noteworthy that, after a great number of trials, we found that, as the value of α approaches and moves beyond 1, our estimation procedure gives unsatisfactory results. In this case, the estimation performance degrades considerably as the bias increases significantly (exceeding 0.5 in some cases), the mse becomes large and the coverage probabilities of the confidence intervals drops below acceptable levels, indicating unreliable inference. This is clearly checked in the bottom parts ($\alpha = 1$ and $\alpha = 1.2$) of each one of the four tables in the appendix. This drawback is probably due to the fact that our estimator $\hat{\beta}_n$

heavily relies on Hill's estimator $\hat{\alpha}_n$ which can be inaccurate for large values of α with small to average sample sizes, as discussed at the end of Section 2.

Our overall conclusion is that, compared to α , the parameter σ has minor effect on the behavior of $\hat{\beta}_n$. That's to say that $\hat{\beta}_n$ is highly sensitive to the distribution tail thickness controlled by α and more robust to variations in scale. In other words, the performance of the estimator of the skewness parameter β of $S_\alpha(\sigma, \beta, \mu)$ is strongly influenced by the stability index α , while being less sensitive to variations in the scale parameter σ .

Acknowledgement. The authors gratefully acknowledge the reviewers for their insightful comments and constructive suggestions, which have significantly improved the paper.

Declaration. The authors have no conflict of interests to disclose.

REFERENCES

1. Csörgő, S., & Viharos, L. (1998). Estimating the tail index. In Asymptotic methods in probability and statistics (pp. 833-881). North-Holland. <https://doi.org/10.1016/B978-044450083-0/50055-1>
2. Fama, E. F. (1965). The behavior of stock-market prices. The journal of Business, 38(1), 34-105. <https://doi.org/10.1086/294743>
3. Garcia, R., Renault, E., & Veredas, D. (2011). Estimation of stable distributions by indirect inference. Journal of Econometrics, 161(2), 325-337. <https://doi.org/10.1016/j.jeconom.2010.12.007>
4. Hill, B. M. (1975). A simple general approach to inference about the tail of a distribution. The annals of statistics, 1163-1174. <https://doi.org/10.1214/aos/1176343247>
5. Ihaka, R., & Gentleman, R. (1996). R: a language for data analysis and graphics. Journal of computational and graphical statistics, 5(3), 299-314. <https://doi.org/10.1080/10618600.1996.10474713>
6. Lévy, P. (1925). Calcul des probabilités. Gauthier-Villars. <https://doi.org/10.2307/3602880>
7. Mandelbrot, B. (1963). The variation of certain speculative prices. Journal of business, 36(4), 394. <https://doi.org/10.1086/294632>
8. Meraghni, D., & Necir, A. (2007). Estimating the scale parameter of a Lévy-stable distribution via the extreme value approach. Methodology and Computing in Applied Probability, 9(4), 557-572. <https://doi.org/10.1007/s11009-007-9021-y>
9. Neves, C., & Alves, M. F. (2004). Reiss and Thomas' automatic selection of the number of extremes. Computational statistics & data analysis, 47(4), 689-704. <https://doi.org/10.1016/j.csda.2003.11.011>
10. Nolan, J. P. (2020). Univariate Estimation. In Univariate Stable Distributions: Models for Heavy Tailed Data (pp. 159-222). Cham: Springer International Publishing. https://doi.org/10.1007/978-3-030-52915-4_4
11. Peng, L. (2001). Estimating the mean of a heavy tailed distribution. Statistics & Probability Letters, 52(3), 255-264. [https://doi.org/10.1016/S0167-7152\(00\)00203-0](https://doi.org/10.1016/S0167-7152(00)00203-0)
12. Reiss, R. D., & Thomas, M. (2007). Statistical analysis of extreme values: with applications to insurance, finance, hydrology and other fields. Basel: Birkhäuser Basel. <https://doi.org/10.1007/s001800000035>
13. Samorodnitsky, G., & Taqqu, M. S. (1994). Stable non-Gaussian random processes: stochastic models with infinite variance (Vol. 1). CRC press. <https://doi.org/10.1201/9780203738818>
14. Wasserman, L. (2006). All of nonparametric statistics. New York, NY: Springer New York. <https://doi.org/10.1177/0049124110371313>

15. Weron, R. (2001). Levy-stable distributions revisited: tail index > 2 does not exclude the Levy-stable regime. International Journal of Modern Physics C, 12(02), 209-223. <https://doi.org/10.1142/s0129183101001614>

5. APPENDIX: SIMULATION RESULTS

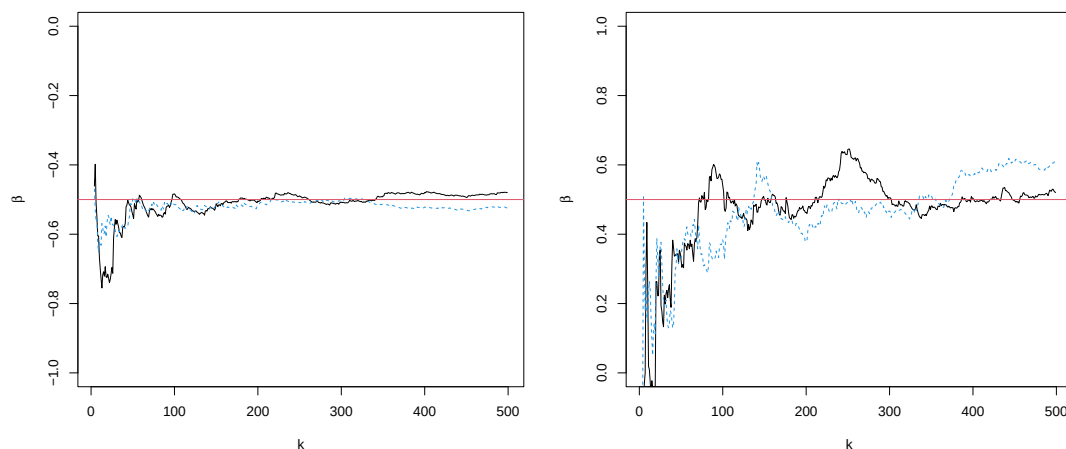


FIGURE 2. Estimator of the skewness parameter β , based on 2000 samples of size 1000, of a Levy-stable distribution vs. the number k of upper order statistics for $\alpha = 0.25$ (black line), $\alpha = 0.75$ (blue line) and $\sigma = 0.5$. The horizontal lines represent the true values $\beta = -0.5$ (left) and $\beta = 0.5$ (right).

¹ DEPARTMENT OF MATHEMATICS, UNIVERSITY OF KASDI MERBAH, OUARGLA, ALGERIA.
Email address: saidane.hadda@univ-ouargla.dz

² LABORATORY OF APPLIED MATHEMATICS, MOHAMED KHIDER UNIVERSITY, BISKRA, ALGERIA.
Email address: djamel.meraghni@univ-biskra.dz

³ LABORATORY OF APPLIED MATHEMATICS, MOHAMED KHIDER UNIVERSITY, BISKRA, ALGERIA.
Email address: l.soltane@univ-biskra.dz

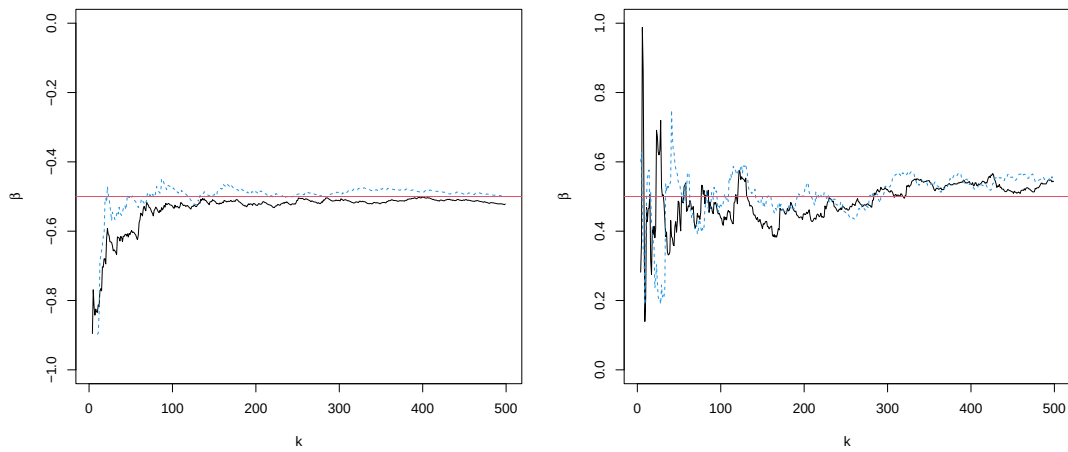


FIGURE 3. Estimator of the skewness parameter β , based on 2000 samples of size 1000, of a Levy-stable distribution vs. the number k of upper order statistics for $\sigma = 0.1$ (black line), $\sigma = 1$ (blue line) and $\alpha = 0.5$. The horizontal lines represent the true values $\beta = -0.5$ (left) and $\beta = 0.5$ (right).

TABLE 1. Estimation results for the skewness parameter $\beta = -0.5$ of a Lévy-stable distribution based on 2000 samples of 1000 observations.

$\beta = -0.5$							
α	σ	$\hat{\beta}_n$	bias	mse	conf int	cov prob	length
0.25	0.1	-0.557	0.057	0.014	-0.853, -0.322	0.84	0.531
	0.5	-0.560	0.061	0.015	-0.851, -0.325	0.79	0.526
	1	-0.555	0.055	0.015	-0.839, -0.323	0.80	0.516
0.5	0.1	-0.592	0.092	0.021	-0.900, -0.357	0.71	0.542
	0.5	-0.598	0.098	0.021	-0.906, -0.368	0.73	0.537
	1	-0.596	0.096	0.020	-0.898, -0.363	0.76	0.535
0.75	0.1	-0.741	0.241	0.079	-1.127, -0.488	0.44	0.639
	0.5	-0.764	0.264	0.086	-1.166, -0.522	0.35	0.645
	1	-0.753	0.254	0.085	-1.175, -0.507	0.37	0.668
1	0.1	0.168	0.668	0.482	-0.254, 0.663	0.24	0.916
	0.5	0.071	0.571	0.362	-0.367, 0.522	0.34	0.888
	1	-0.004	0.496	0.281	-0.464, 0.439	0.43	0.902
1.2	0.1	0.366	0.866	0.845	-0.191, 1.417	0.25	1.608
	0.5	0.154	0.654	0.492	-0.289, 0.988	0.29	1.277
	1	0.151	0.651	0.488	-0.298, 0.971	0.27	1.269

TABLE 2. Estimation results for the skewness parameter $\beta = 0$ of a Lévy-stable distribution based on 2000 samples of 1000 observations.

$\beta = 0$							
α	σ	$\hat{\beta}_n$	bias	mse	conf int	cov prob	length
0.25	0.1	-0.010	0.010	0.025	-0.411 , 0.393	0.89	0.803
	0.5	-0.024	0.024	0.030	-0.476 , 0.372	0.87	0.849
	1	-0.001	0.001	0.030	-0.445 , 0.395	0.85	0.840
0.5	0.1	-0.020	0.020	0.011	-0.327 , 0.260	0.90	0.586
	0.5	-0.032	0.032	0.011	-0.347 , 0.241	0.89	0.588
	1	-0.020	0.020	0.012	-0.324 , 0.270	0.90	0.594
0.75	0.1	-0.015	0.015	0.025	-0.448 , 0.394	0.90	0.842
	0.5	0.001	0.001	0.022	-0.442 , 0.412	0.93	0.854
	1	-0.015	0.015	0.025	-0.448 , 0.394	0.90	0.842
1	0.1	-0.004	0.004	0.013	-0.293 , 0.253	0.89	0.575
	0.5	-0.009	0.009	0.012	-0.323 , 0.284	0.69	0.607
	1	0.003	0.003	0.013	-0.276 , 0.261	0.59	0.536
1.2	0.1	0.010	0.010	0.023	-0.393 , 0.410	0.60	0.803
	0.5	0.005	0.005	0.025	-0.410 , 0.421	0.55	0.831
	1	0.009	0.009	0.024	-0.409 , 0.427	0.54	0.836

TABLE 3. Estimation results for the skewness parameter $\beta = 0.5$ of a Lévy-stable distribution based on 2000 samples of 1000 observations.

$\beta = 0.5$							
α	σ	$\hat{\beta}_n$	bias	mse	conf int	cov prob	length
0.25	0.1	0.503	0.003	0.021	0.052 , 0.924	0.92	0.872
	0.5	0.498	0.002	0.026	0.043 , 0.939	0.91	0.896
	1	0.512	0.012	0.022	0.061 , 0.952	0.89	0.890
0.5	0.1	0.508	0.008	0.025	0.077 , 0.928	0.88	0.851
	0.5	0.517	0.017	0.023	0.090 , 0.953	0.88	0.863
	1	0.515	0.015	0.024	0.076 , 0.934	0.89	0.858
0.75	0.1	0.571	0.071	0.024	0.162 , 0.986	0.85	0.824
	0.5	0.582	0.082	0.003	0.182 , 0.999	0.85	0.817
	1	0.574	0.074	0.029	0.174 , 0.998	0.86	0.824
1	0.1	-0.332	0.832	0.754	-0.992 , 0.118	0.22	1.110
	0.5	-0.185	0.685	0.510	-0.738 , 0.215	0.21	0.953
	1	-0.093	0.593	0.392	-0.619 , 0.323	0.31	0.942
1.2	0.1	0.043	0.457	0.300	-0.780 , 0.626	0.48	1.406
	0.5	0.039	0.461	0.292	-0.783 , 0.638	0.51	1.422
	1	0.026	0.474	0.310	-0.789 , 0.599	0.47	1.388

TABLE 4. Estimation results for the skewness parameter $\beta = 0.8$ of a Lévy-stable distribution based on 2000 samples of 1000 observations.

$\beta = 0.8$							
α	σ	$\hat{\beta}_n$	bias	mse	conf int	cov prob	length
0.25	0.1	0.794	0.006	0.013	0.469 , 1.129	0.89	0.660
	0.5	0.794	0.006	0.013	0.466 , 1.102	0.84	0.636
	1	0.784	0.016	0.018	0.459 , 1.109	0.88	0.650
0.5	0.1	0.796	0.004	0.015	0.472 , 1.123	0.85	0.652
	0.5	0.811	0.011	0.013	0.500 , 1.119	0.87	0.620
	1	0.807	0.007	0.013	0.493 , 1.144	0.87	0.651
0.75	0.1	0.839	0.039	0.014	0.558 , 1.133	0.78	0.575
	0.5	0.839	0.039	0.019	0.541 , 1.168	0.79	0.626
	1	0.833	0.033	0.015	0.527 , 1.179	0.83	0.652
1	0.1	-0.289	1.089	1.252	-0.932 , 0.217	0.12	1.149
	0.5	-0.134	0.934	0.921	-0.667 , 0.338	0.11	1.005
	1	-0.088	0.888	0.835	-0.601 , 0.371	0.15	0.972
1.2	0.1	0.162	0.638	0.563	-0.912 , 0.915	0.53	1.827
	0.5	-0.169	0.969	1.018	-0.890 , 0.445	0.19	1.335
	1	0.133	0.667	0.607	-0.975 , 0.878	0.14	1.854