

Robust and Efficient Estimation for the Block-Basu Bivariate Exponential Distribution

Sanjay Kumar^{1, a}, Debasis Kundu^{2, a}, Sharmishtha Mitra^{3, a*}

^aDepartment of Mathematics and Statistics, Indian Institute of Technology Kanpur, Kanpur, Pin 208016, India.

*Corresponding author(s). E-mail(s): 3smitra@iitk.ac.in;
Contributing authors: 1sanjaymk@iitk.ac.in; 2kundu@iitk.ac.in;

Abstract

The Block-Basu bivariate exponential (BBBE) distribution, see [1], is one of the most popular variants of an absolutely continuous bivariate distribution. The traditional maximum likelihood estimation method is highly efficient, yet the maximum likelihood estimators (MLEs) can be significantly affected by the presence of a few outliers in the data. This article focuses on introducing a robust estimation procedure in the presence of outliers. We explore the implementation of the minimum density power divergence estimator (MDPDE) for robust and efficient parameter estimation of the BBBE distribution. The MDPDE is indexed by a single tuning parameter that controls the trade-off between robustness and efficiency. We show that the MDPDE for the BBBE model has a bounded influence function, that ensures the robustness of MDPDE. We compare the asymptotic covariance matrix of MDPDE with the Fisher information matrix. We study the asymptotic relative efficiency of MDPDE. We discussed about the data-driven selection of optimal tuning parameter. The simulation results indicate that the MDPDE is more robust and efficient than MLE for the BBBE distribution when outliers are present in a data set. A real data set has been analyzed for illustration purposes.

Keywords: Block-Basu bivariate exponential distribution, Maximum likelihood, M-estimation, Minimum density power divergence estimation, Robustness and efficiency, Influence function.

1 Introduction

Marshall and Olkin [2] introduced a widely used bivariate exponential distribution, known as the Marshall-Olkin bivariate exponential (MOBE) distribution. However, this distribution is not absolutely continuous, as it has both absolutely continuous and singular components. To address this, Block and Basu [1] derived an absolutely continuous bivariate exponential distribution, referred to as the BBBE distribution, by removing the singular component from the MOBE distribution and retaining only its absolutely continuous part. Although the marginals of the BBBE distribution are not exponential, unlike those of the MOBE distribution, the BBBE model has been used in various applications where simultaneous occurrences of observations are rare. It is flexible because its marginals can be expressed as mixtures of exponential distributions. Consequently, it has received considerable attention, and extensive research has been conducted on it over the past few decades.

Parameter estimation plays a crucial role in statistical modeling. Traditional maximum likelihood estimators (MLEs) are widely used due to their desirable theoretical properties, such as consistency, asymptotic normality, and efficiency. However, while MLEs are the most efficient under ideal conditions, they are known to be highly sensitive to outliers. The primary motivation for this study arises from the analysis of a real-world data set related to kidney infections, where the naturally presence of outliers is evident and unavoidable. The MLE is significantly affected by the presence of outliers, leading to biased and inefficient estimates, and consequently, poor model performance. This sensitivity can compromise the accuracy and validity of statistical inferences. Therefore, there is a strong need for robust estimation techniques that can yield reliable parameter estimates even in the presence of outliers and other forms of data irregularities.

In recent years, there has been growing interest in the development of robust estimation methods, such as the minimum density power divergence (MDPD) approach, which minimizes a data-driven divergence measure to produce robust parameter estimates. The MDPD estimation method, along with its properties in general, was first introduced by Basu et al. [3]. In recent years, several authors have applied this method in the univariate context (see, for example, [4], [5]). The difficulty level increases manifold while shifting from the univariate (where it is already studied in the literature) to the bivariate structure, which has not been previously explored. This contribution is further non-trivial due to the specific structure of the model, which poses unique challenges in deriving the asymptotic and robustness properties. However, to the best of our knowledge, this approach has not yet been implemented for bivariate or multivariate models in the literature. This work aims to fill that gap by applying the MDPD estimation methodology to a bivariate complex model, specifically focusing on robust estimation techniques for the BBBE distribution. The MDPD method offers several advantages over traditional estimation techniques, including robustness to outliers and improved efficiency in parameter estimation. Our findings are useful for developing other bivariate distributions.

The main advantage of the MDPDE is its ability to balance robustness and efficiency in the presence of outliers by incorporating a user-specified tuning parameter, γ . In real data analysis, selecting an optimal γ is crucial. This has been extensively

studied in the univariate framework by many authors. Our contribution in this article is to extend this approach to a complex bivariate model. Specifically, we implement the ECVm method, as discussed in Section 5, to select the optimal γ in the bivariate BBBE model.

Specifically, the derivation of the MDPDE for this bivariate model and the establishment of its asymptotic properties require model-specific adaptations that are not straightforward extensions of existing univariate results. The influence function analysis and robustness properties also had to be carefully derived in the context of this particular bivariate setting. The theoretical properties of the MDPDE in the context of the BBBE model are not straightforward to establish, but they follow the same line as in the univariate case. Another key contribution is the development of a data-driven approach to select the optimal tuning parameter, which balances robustness and efficiency, thereby controlling the influence of outliers. Our results demonstrate that, in the presence of contamination, the MDPD estimate significantly outperforms the MLE.

Basu et al. [3] introduced a new family of density-based divergence measures known as density power divergences (DPDs). The family of DPDs between the underlying (true) density g and the parametric density f is defined as follows,

$$d_\gamma(g, f) = \begin{cases} \int \{f^{1+\gamma}(x) - (1 + \frac{1}{\gamma})g(x)f^\gamma(x) + \frac{1}{\gamma}g^\gamma(x)\}dx & \text{for } \gamma > 0 \\ \lim_{\gamma \rightarrow 0} d_\gamma(g, f) = d_0(g, f) = \int g(x) \log(\frac{g(x)}{f(x)})dx & \text{for } \gamma = 0. \end{cases} \quad (1)$$

The quantity defined in equation (1) is non-negative for all g and f , and is equal to zero if and only if g and f are equal almost everywhere, as stated in Theorem 1 of Basu et al. [3]. This family is indexed by a single non-negative parameter γ , referred to as the tuning parameter, which controls the trade-off between robustness and asymptotic efficiency of the estimators.

Although this family is not directly related to the power-divergence family of Cressie and Read [6], some indirect connections may be drawn. The family includes the Kullback-Leibler (KL) divergence [7] when, $\gamma = 0$ and L_2 -distance when $\gamma = 1$. Therefore the minimum density power divergence estimator (MDPDE) can be thought as bridge between MLE (efficient but non-robust) and L_2 -estimator (robust but inefficient). The degree of robustness increases with increasing γ . However, larger values of γ are also associated with more inefficient estimators. Unlike other divergence-based approaches, such as those in [8–10], the main advantage of the MDPD method lies in the fact that it does not require any non-parametric smoothing or related complications for any value of γ . The MDPD estimation method is the robust generalization of the MLE and it is similar to the MLE when $\gamma = 0$.

In this article, we implement the MDPD estimation method for robust and efficient parameter estimates of the BBBE distribution. This method aims to address the limitations of existing estimation techniques by providing robust and efficient estimates of the model parameters. By minimizing the density-based divergence measure, our proposed estimator is able to mitigate the effects of outliers and data irregularities, thereby improving the accuracy and reliability of parameter estimation. We

discuss the theoretical properties of MDPDE, where we derive the asymptotic distribution of the MDPDEs of the parameters of the BBBE distribution. We observe that the without outlier case the asymptotic covariance of MDPDEs with $\gamma \rightarrow 0$ become asymptotic covariance of the MLEs. We discussed the asymptotic relative efficiency of MLEs over MDPDEs. Further, we study the behavior of the influence function (IF) for unknown parameters at different choices of the tuning parameter γ and we have seen that the IF for MDPDE is bounded that ensure the MDPDEs are robust. Then we study how the MDPDEs for the BBBE model varies across different γ and describe data-driven optimal tuning parameter selection by minimizing the empirical Cramer-von Mises (ECVM) distance following [11]. Subsequently, we perform a simulation study to explore the effectiveness of MDPDEs for the BBBE distribution at low through high levels of contamination. Here, we compare MDPDE with some alternative parameter estimation strategies like MLE. An extensive simulations perform and see that the MDPDEs with small γ are similar to MLEs without outlier case. In the presence of outlier the MDPDEs are robust with small loss of efficiency as compared to the MLEs. We analyzed a real data set for the illustrative purposes and see its applicability in real scenarios.

The rest of paper is organized as follows. In Section 2, we provide a brief overview of the BBBE model. In Section 3, we present robust and efficient estimation approach based on MDPD for the BBBE distribution. In Section 4, we discussed robustness and asymptotic properties of MDPDE. Selection of optimal tuning parameter is discussed in Section 5. In Section 6, we conduct a comprehensive simulation study to evaluate the performance of our proposed method and compare it with MLE. In Section 7, we analyzed a real data set to demonstrate its practical applications. Finally, we conclude the paper in Section 8.

2 Block-Basu Bivariate Exponential Distribution

The probability density function (PDF) and the corresponding survival function (SF) of an exponential distribution with scale parameter $\alpha > 0$, denoted as $T \sim Exp(\alpha)$, are given by $f_{ED}(t; \alpha) = \alpha e^{-\alpha t}$ and $S_{ED}(t; \alpha) = e^{-\alpha t}$ for $t > 0$, respectively. Suppose $U_i \sim Exp(\alpha_i)$ for $i = 1, 2, 3$, and that they are independent random variables. Define $X_j = \min\{U_j, U_3\}$ for $j = 1, 2$, then (X_1, X_2) has a MOBE distribution, denoted by $(X_1, X_2) \sim MOBE(\alpha_1, \alpha_2, \alpha_3)$. The joint SF of (X_1, X_2) is $S_{(X_1, X_2)}(x_1, x_2) = e^{-\alpha_1 x_1} e^{-\alpha_2 x_2} e^{-\alpha_3 \max\{x_1, x_2\}}$; $x_1 > 0$, $x_2 > 0$ and the corresponding PDF is

$$f_{(X_1, X_2)}(x_1, x_2) = \begin{cases} \alpha_1(\alpha_2 + \alpha_3)e^{-\alpha_1 x_1} e^{-(\alpha_2 + \alpha_3)x_2} & \text{if } x_1 < x_2 \\ (\alpha_1 + \alpha_3)\alpha_2 e^{-(\alpha_1 + \alpha_3)x_1} e^{-\alpha_2 x_2} & \text{if } x_1 > x_2 \\ \alpha_3 e^{-(\alpha_1 + \alpha_2 + \alpha_3)x} & \text{if } x_1 = x_2 = x. \end{cases} \quad (2)$$

The first two components of equation (2) are absolutely continuous parts, while the third is the singular part. The MOBE distribution is not absolutely continuous.

The Block and Basu [1] have been proposed an interesting absolutely continuous bivariate exponential distribution related to the MOBE distribution, known as the BBBE distribution. It is obtained from the MOBE distribution by removing

the singular part and keeping only the absolutely continuous part. The derivation of the BBBE distribution is immediate from the MOBE distribution as follows if $(X_1, X_2) \sim MOBE(\alpha_1, \alpha_2, \alpha_3)$ and $(Y_1, Y_2) = (X_1, X_2)$ given $X_1 \neq X_2$, then (Y_1, Y_2) has a BBBE distribution with parameters α_1, α_2 and α_3 , we will denote it by $(Y_1, Y_2) \sim BBBE(\alpha_1, \alpha_2, \alpha_3)$. The joint PDF of BBBE distribution is obtained as,

$$f_{BB}(y_1, y_2; \theta) = \begin{cases} f_1(y_1, y_2; \theta) & \text{if } y_1 < y_2 \\ f_2(y_1, y_2; \theta) & \text{if } y_1 > y_2, \end{cases} \quad (3)$$

$$\begin{aligned} \text{where } f_1(y_1, y_2; \theta) &= c\alpha_1(\alpha_2 + \alpha_3)e^{-\alpha_1 y_1}e^{-(\alpha_2 + \alpha_3)y_2}, \\ f_2(y_1, y_2; \theta) &= c(\alpha_1 + \alpha_3)\alpha_2e^{-(\alpha_1 + \alpha_3)y_1}e^{-\alpha_2 y_2} \end{aligned}$$

and $c = (\alpha_1 + \alpha_2 + \alpha_3)/(\alpha_1 + \alpha_2)$ is the normalizing constant and $\theta = (\alpha_1, \alpha_2, \alpha_3)$. The corresponding joint SF is given as,

$$S_{BB}(y_1, y_2) = \begin{cases} ce^{-\alpha_1 y_1}e^{-(\alpha_2 + \alpha_3)y_2} - \frac{\alpha_3}{\alpha_1 + \alpha_2}e^{-(\alpha_1 + \alpha_2 + \alpha_3)y_2} & \text{if } y_1 < y_2 \\ ce^{-(\alpha_1 + \alpha_3)y_1}e^{-\alpha_2 y_2} - \frac{\alpha_3}{\alpha_1 + \alpha_2}e^{-(\alpha_1 + \alpha_2 + \alpha_3)y_1} & \text{if } y_1 > y_2. \end{cases} \quad (4)$$

The marginal SF and PDF of Y_i , $i = 1, 2$, for $y_i > 0$, are given, respectively

$$\begin{aligned} S_{Y_i}(y_i) &= ce^{-(\alpha_i + \alpha_3)y_i} - \frac{\alpha_3}{\alpha_1 + \alpha_2}e^{-(\alpha_1 + \alpha_2 + \alpha_3)y_i}, \\ f_{Y_i}(y_i) &= cf_{ED}(y_i; \alpha_i + \alpha_3) - \frac{\alpha_3}{\alpha_1 + \alpha_2}f_{ED}(y_i; \alpha_1 + \alpha_2 + \alpha_3). \end{aligned} \quad (5)$$

The marginals of the BBBE model are not exponential distributions, unlike those of the MOBE model. However, $\min\{Y_1, Y_2\} \sim \text{Exp}(\alpha_1 + \alpha_2 + \alpha_3)$. If $\alpha_3 = 0$, the marginals Y_1 and Y_2 of the BBBE distribution are independent and each follows an exponential distribution. For $\alpha_3 > 0$, Y_1 and Y_2 are dependent, with a positive correlation between them.

3 Minimum Density Power Divergence Estimation

In this section, we discuss the computation of robust and efficient estimates, along with the MLEs, for the unknown parameters α_1, α_2 and α_3 of the BBBE distribution. MLE is the most commonly used parameter estimation method due to its attractive theoretical properties, such as consistency and efficiency. However, it is highly sensitive to the presence of even a single outlier and can be significantly affected by either a wrong data entry or an extreme observation (see [12, 13]), leading to unreliable estimates. In this work, the robust parameter estimates are obtained by minimizing a data-driven density power divergence (DPD) over the parameter space.

Suppose $\mathcal{F} = \{f_{BB}(y_1, y_2; \theta), (y_1, y_2) \in \mathcal{Y} : \theta \in \Theta \subset \mathbb{R}_+^3\}$ be denotes a parametric family of BBBE densities, where $\mathcal{Y} = \{(y_1, y_2) : f_{BB}(y_1, y_2; \theta) > 0\}$ and $\theta = (\alpha_1, \alpha_2, \alpha_3)$. This family is identifiable in the sense that the set $\{(y_1, y_2) \in \mathcal{Y} :$

$f_{BB}(y_1, y_2; \theta_1) \neq f_{BB}(y_1, y_2; \theta_2)$ has positive Lebesgue measure if $\theta_1 \neq \theta_2$. We are interested in the estimation of θ based on a random sample of (2-dimensional) observations $\{(y_{11}, y_{21}), (y_{12}, y_{22}), \dots, (y_{1n}, y_{2n})\}$ from the true distribution $G(\cdot)$ with density $g(\cdot)$. We use the following notations for further development, $I_1 = \{i : y_{1i} < y_{2i}\}$ and $I_2 = \{i : y_{1i} > y_{2i}\}$. Let $n_j = |I_j|$ denotes the number of observations in the set I_j , $j = 1, 2$ such that $n = n_1 + n_2$. In the DPD expression given in (1), we replace the density f with the BBDE density $f_{BB}(y_1, y_2)$. For any given tuning parameter $\gamma \geq 0$, the MDPD functional (trivariate) $T_\gamma(\cdot)$ at G is defined as the point in the parameter space $\Theta \subset \mathbb{R}_+^3$ corresponding to the element in $\mathcal{F}$ closest to g such that

$$d_\gamma(g, f_{BB}(y_1, y_2; T_\gamma(G))) = \min_{\theta \in \Theta} d_\gamma(g, f_{BB}(y_1, y_2; \theta)). \quad (6)$$

This MDPD functional $\theta^g = T_\gamma(G)$ is the target parameter vector that provides the best fit under G . In practice, however, the true density g is never known and hence to obtain the minimizer of DPD, we need to use an estimate of g . Unlike other robust divergence measures, the particular DPD family proposed by [3] and has a major advantage because of not requiring any non-parametric smoothing for estimating g . This is because we can rewrite the DPD for $\gamma > 0$,

$$d_\gamma(g, f(y_1, y_2; \theta)) = \int_0^\infty \int_0^\infty f_{BB}^{1+\gamma}(y_1, y_2; \theta) dy_1 dy_2 - \left(1 + \frac{1}{\gamma}\right) E_g[f_{BB}^\gamma(Y_1, Y_2; \theta)] + \frac{1}{\gamma} E_g[g^\gamma(Y_1, Y_2)] \quad (7)$$

and for $\gamma = 0$, $d_0(g, f_\theta) = E_g[\log g(Y_1, Y_2)] - E_g[\log f_{BB}(Y_1, Y_2; \theta)]$, where $E_g[\cdot]$ denotes the expectation with respect to the true density g . Here, the two terms $E_g[g^\gamma(Y_1, Y_2)]$ and $E_g[\log(g(Y_1, Y_2))]$ are independent of θ , and hence can be ignored in the optimization of (6). The other two expectations $E_g[f_{BB}^\gamma(Y_1, Y_2; \theta)]$ and $E_g[\log(f_{BB}(Y_1, Y_2; \theta))]$ in (7) can be estimated directly through empirical means and does not require any non-parametric smoothing for estimating g . Therefore, the MDPDE of θ for fixed tuning parameter γ is given by $\hat{\theta}_{\gamma,n} = T_\gamma(G_n) \in \Theta$, where G_n is the empirical distribution function associated with the observed random sample. That is, $\hat{\theta}_{\gamma,n}$ minimizes $d_\gamma(g_n, f_\theta)$ with respect to the θ . For given tuning parameter γ , the MDPDE $\hat{\theta}_{\gamma,n}$ is given by

$$\hat{\theta}_{\gamma,n} = \arg \min_{\theta \in \Theta} H_{\gamma,n}(\theta), \quad (8)$$

where for $\gamma > 0$

$$H_{\gamma,n}(\theta) = \int_0^\infty \int_0^\infty f_{BB}^{1+\gamma}(y_1, y_2; \theta) dy_1 dy_2 - \left(1 + \frac{1}{\gamma}\right) \frac{1}{n} \sum_{i=1}^n f_{BB}^\gamma(y_{1i}, y_{2i}; \theta). \quad (9)$$

Here, for $\gamma = 0$, the MDPD function $H_{\gamma,n}(\theta) = -1/n \sum_{i=1}^n \log f_{BB}(y_{1i}, y_{2i}; \theta)$ is the negative of $1/n$ times log-likelihood function, and hence the MDPDE $\hat{\theta}_{0,n} = T_0(G_n)$

coincides with the MLE. On the other hand, for $\gamma = 1$, the MDPDE $\hat{\theta}_{1,n} = T_1(G_n)$ corresponds to the L_2 -square distance estimator. We are interested to obtain the MDPDE only for $0 \leq \gamma \leq 1$. There is no strong reason to avoid the case for $\gamma > 1$ except the loss of efficiency. The explicit expression of $H_{\gamma,n}(\theta)$ is provided in the following theorem.

Theorem 1. Suppose $(y_{11}, y_{21}), (y_{12}, y_{22}), \dots, (y_{1n}, y_{2n})$ are n observations from the underlying true distribution G . For $\gamma > 0$, the MDPD objective function $H_{\gamma,n}(\theta)$ for the BBBE($\alpha_1, \alpha_2, \alpha_3$) distribution is

$$\begin{aligned} H_{\gamma,n}(\theta) = & \frac{c^{1+\gamma} [\alpha_1^{1+\gamma} (\alpha_2 + \alpha_3)^\gamma + (\alpha_1 + \alpha_3)^\gamma \alpha_2^{1+\gamma}]}{(1 + \gamma)^2 (\alpha_1 + \alpha_2 + \alpha_3)} \\ & - \left(1 + \frac{1}{\gamma}\right) \frac{c^\gamma}{n} \left[\alpha_1^\gamma (\alpha_2 + \alpha_3)^\gamma \sum_{i \in I_1} e^{-\gamma \alpha_1 y_{1i}} e^{-\gamma (\alpha_2 + \alpha_3) y_{2i}} \right] \\ & - \left(1 + \frac{1}{\gamma}\right) \frac{c^\gamma}{n} \left[(\alpha_1 + \alpha_3)^\gamma \alpha_2^\gamma \sum_{i \in I_2} e^{-\gamma (\alpha_1 + \alpha_3) y_{1i}} e^{-\gamma \alpha_2 y_{2i}} \right]. \end{aligned} \quad (10)$$

Proof. For $\gamma > 0$, we have

$$H_{\gamma,n}(\theta) = \int_0^\infty \int_0^\infty f_{BB}^{1+\gamma}(y_1, y_2; \theta) dy_1 dy_2 - \left(1 + \frac{1}{\gamma}\right) \frac{1}{n} \sum_{i=1}^n f_{BB}^\gamma(y_{1i}, y_{2i}; \theta).$$

Now we obtain the first term of $H_{\gamma,n}(\theta)$ as

$$\begin{aligned} & \int_0^\infty \int_0^\infty f_{BB}^{1+\gamma}(y_1, y_2; \theta) dy_1 dy_2 \\ &= \int_0^\infty \int_{y_1}^\infty f_1^{1+\gamma}(y_1, y_2; \theta) dy_2 dy_1 + \int_0^\infty \int_{y_2}^\infty f_2^{1+\gamma}(y_1, y_2; \theta) dy_1 dy_2 \\ &= c^{1+\gamma} \alpha_1^{1+\gamma} (\alpha_2 + \alpha_3)^{1+\gamma} \int_0^\infty \int_{y_1}^\infty e^{-\alpha_1 (1+\gamma) y_1} e^{-(\alpha_2 + \alpha_3) (1+\gamma) y_2} dy_2 dy_1 \\ &+ c^{1+\gamma} (\alpha_1 + \alpha_3)^{1+\gamma} \alpha_2^{1+\gamma} \int_0^\infty \int_{y_1}^\infty e^{-(\alpha_1 + \alpha_3) (1+\gamma) y_1} e^{-\alpha_2 (1+\gamma) y_2} dy_1 dy_2 \\ &= \frac{c^{1+\gamma} \alpha_1^{1+\gamma} (\alpha_2 + \alpha_3)^\gamma}{(1 + \gamma)^2 (\alpha_1 + \alpha_2 + \alpha_3)} + \frac{c^{1+\gamma} \alpha_2^{1+\gamma} (\alpha_1 + \alpha_3)^\gamma}{(1 + \gamma)^2 (\alpha_1 + \alpha_2 + \alpha_3)}. \end{aligned}$$

The second term is obtained by simply taking the γ^{th} power of the PDF in (3) of the BBBE distribution, that is,

$$\begin{aligned} \sum_{i=1}^n f_{BB}^\gamma(y_{1i}, y_{2i}; \theta) &= \sum_{i \in I_1} f_1^\gamma(y_{1i}, y_{2i}; \theta) + \sum_{i \in I_2} f_2^\gamma(y_{1i}, y_{2i}; \theta) \\ &= c^\gamma \alpha_1^\gamma (\alpha_2 + \alpha_3)^\gamma \sum_{i \in I_1} e^{-\alpha_1 (1+\gamma) y_{1i}} e^{-(\alpha_2 + \alpha_3) (1+\gamma) y_{2i}} + c^\gamma (\alpha_1 + \alpha_3)^\gamma \alpha_2^\gamma \sum_{i \in I_2} e^{-\alpha_1 (1+\gamma) y_{1i}} e^{-\alpha_2 (1+\gamma) y_{2i}}. \end{aligned}$$

This completes the proof. $\square$

The expression given in equation (10) is also known by the negative modified likelihood function. The MDPDEs of the unknown parameters α_1, α_2 and α_3 are obtained by minimizing the $H_{\gamma,n}(\theta)$, as given in (10). Clearly it is observed that the MDPDEs are not obtained in closed forms. Thus it become three dimensional optimization problem. The MDPDEs are obtained numerically using standard optimization algorithms, such as the Nelder-Mead algorithm, which is available in R through the *optim()* function.

4 Properties of MDPDE for the BBBE distribution

The MDPDE for the BBBE distribution provides a robust and flexible estimation method, effective in handling outliers and model misspecification. With asymptotic properties such as consistency and asymptotic normality, it forms a strong basis for statistical inference. The tuning parameter γ allows a balance between efficiency and robustness, making the MDPDE a powerful tool for complex BBBE model. These estimators are particularly valuable for statistical inference due to their robustness and asymptotic properties, striking a balance between robustness and efficiency.

4.1 Unbiased estimating equation

The unbiased estimating equation for any $\gamma \geq 0$ is derived by differentiating $H_{\gamma,n}(\theta)$, as defined in equation (9), with respect to θ . The resulting equation for the BBBE distribution is given by

$$U_{\theta,n} = \frac{1}{n} \sum_{i=1}^n S_{\theta}(y_{1i}, y_{2i}) f_{BB}^{\gamma}(y_{1i}, y_{2i}; \theta) - \int_0^{\infty} \int_0^{\infty} S_{\theta}(y_1, y_2) f_{BB}^{\gamma+1}(y_1, y_2; \theta) dy_1 dy_2 = 0, \quad (11)$$

Here $S_{\theta}(y_1, y_2) = \nabla \log[f_{BB}(y_1, y_2; \theta)]$ is the maximum likelihood score vector of the BBBE model $f_{BB}(y_1, y_2; \theta)$, given in Appendix A. The unbiased estimating equation of the MDPDE for the MOBE distribution is a system of three equations. The respective equations for the BBBE parameters α_1, α_2 and α_3 are given as for $j = 1, 2, 3$,

$$U_{k,\theta,n} = \frac{1}{n} \sum_{i=1}^n S_{k,\theta}(y_{1i}, y_{2i}) f_{BB}^{\gamma}(y_{1i}, y_{2i}; \theta) - \int_0^{\infty} \int_0^{\infty} S_{k,\theta}(y_1, y_2) f_{BB}^{\gamma+1}(y_1, y_2; \theta) dy_1 dy_2 = 0,$$

where $S_{k,\theta}(y_{1i}, y_{2i})$ is the k^{th} component of $S_{\theta}(y_{1i}, y_{2i})$. At $\gamma = 0$, the equation (11) boils down to the usual score equations for the MLE, as given $U_{\theta,n} = 1/n \sum_{i=1}^n S_{\theta}(y_{1i}, y_{2i}) = 0$. However, for $\gamma > 0$, the MDPDE yields weighted score equations, where the weights are given by $f_{BB}^{\gamma}(y_{1i}, y_{2i}; \theta)$ for (y_{1i}, y_{2i}) . These weights are smaller for outlying observations, effectively downweighting their influence and resulting in robust estimates. The equation (11) can also be expressed as

$$\sum_{i=1}^n \psi(y_{1i}, y_{2i}, \theta) = 0, \quad (12)$$

where

$$\psi(y_1, y_2, \theta) = S_\theta(y_1, y_2) f_{BB}^\gamma(y_1, y_2; \theta) - \int_0^\infty \int_0^\infty S_\theta(y_1, y_2) f_{BB}^{\gamma+1}(y_1, y_2; \theta) dy_1 dy_2$$

is a vector-valued function of three dimension same as that of the parameter space of the BBBE distribution. From equation (12), it is evident that the MDPDEs for the BBBE model are the M-estimators, for further details see Huber [14], and Hampel [15].

4.2 Robustness properties

In this section, we evaluate the robustness of the MDPDE estimators using the widely recognized influence function (IF) framework introduced by Hampel et al. [15]. The robustness of the MDPDE in the context of the BBBE distribution is one of its key advantages. The ability to control for the impact of outliers and model misspecification is crucial in many applications. Under the model conditions, i.e., $G = F_{BB}(\cdot; \theta^0)$, the IF of the DPD functional $T_\gamma(\cdot)$ with fixed $\gamma > 0$ for the BBBE model is

$$IF(y_1, y_2; T_\gamma, F_{BB}(\cdot; \theta^0)) = J_\gamma^{-1}(\theta^*) [S_{\theta^0}(y_1, y_2) f_{BB}^\gamma(y_1, y_2; \theta^0) - \xi_\gamma(\theta^0)], \quad (13)$$

where $S_{\theta^0}(y_1, y_2)$ is the score vector given in Appendix A. For $\gamma = 0$ (equivalently $\gamma \rightarrow 0$), the IF is $IF(y_1, y_2) = I_n(\theta^0)^{-1} S_{\theta^0}(y_1, y_2)$. The IF (13) is bounded whenever the $S_{\theta^0}(y_1, y_2) f_{BB}^\gamma(y_1, y_2; \theta^0)$ is bounded. The following theorem demonstrates that for $\gamma > 0$, the term $S_{\theta^0}(y_1, y_2) f_{BB}^\gamma(y_1, y_2; \theta^0)$ is bounded. Consequently, the IF is also bounded. However, when $\gamma = 0$, corresponding to the MLE case, the IF becomes unbounded.

Theorem 2. *If the true parameters α_1^0 , α_2^0 and α_3^0 of the BBBE model, then the each component of the IF, $IF(y_1, y_2; T_\gamma, F_{BB}(\cdot; \theta^0))$ as given in (13) is bounded for $\gamma > 0$ and is unbounded for $\gamma = 0$, the MLE case.*

Proof. Since $J_\gamma(\theta^0)$ and $\xi_\gamma(\theta^0)$ in the influence function $IF(y_1, y_2; T_\gamma, F_{BB}(\cdot; \theta^0))$ are independent of (y_1, y_2) . Thus, only the vector $S_{\theta^0}(y_1, y_2) f_{BB}^\gamma(y_1, y_2; \theta^0)$ determines the boundedness. For $IF(y_1, y_2; T_\gamma, F_{BB}(\cdot; \theta^0))$ to be bounded, all four components of $S_{\theta^0}(y_1, y_2) f_{Y_1, Y_2}^\gamma(y_1, y_2; \theta^0)$ must be bounded.

Case 1. When $\gamma > 0$

The first component of $\psi(y_1, y_2; \theta^0) = S_{\theta^0}(y_1, y_2) f_{BB}^\gamma(y_1, y_2; \theta^0)$ is

$$\psi_1(y_1, y_2; \theta^0) = \begin{cases} d_1(c_1 - y_1) e^{-\gamma \alpha_1^0 y_1} e^{-\gamma(\alpha_2^0 + \alpha_3^0) y_2} & \text{if } 0 < y_1 < y_2 \\ d_2(c_2 - y_1) e^{-\gamma(\alpha_1^0 + \alpha_2^0) y_1} e^{-\gamma \alpha_2^0 y_2} & \text{if } 0 < y_2 < y_1, \end{cases}$$

where $d_i > 0$ and $c_i > 0$, for $i = 1, 2$, are constants. These constants do not affect the boundedness of $\psi_1(y_1, y_2; \theta^0)$. For $y_1 < y_2$, since the exponential term $e^{-\gamma \alpha_1^0 y_1}$, where $\alpha_1^0, \gamma > 0$, dominates the growth of $(c_1 - y_1)$, forcing the entire function to tend to 0, as $y_1 \rightarrow \infty$. Similarly, for $y_1 > y_2$, the exponential term $e^{-\gamma \alpha_2^0 y_2}$, where $\alpha_2^0, \gamma > 0$, decays to 0 exponentially fast as $y_2 \rightarrow \infty$. Consequently, as $y_1 \rightarrow \infty$ and $y_2 \rightarrow \infty$, the first component $\psi_1(y_1, y_2; \theta^0)$ tend to 0, indicating that it is bounded. Similarly, the

second and third components $\psi_i(y_1, y_2; \theta^0)$, for $i = 2, 3$, are also bounded. Therefore, $S_{\theta^0}(y_1, y_2)f_{Y_1, Y_2}^\gamma(y_1, y_2; \theta^0)$ is bounded for all $y_1, y_2 > 0$ and hence the function $IF(y_1, y_2; T_\gamma, F_{BB}(\cdot; \theta^0))$ is bounded.

Case 2. When $\gamma = 0$, we have

$$\psi_1(y_1, y_2; \theta^0) = \begin{cases} (c_1 - y_1) & \text{if } 0 < y_1 < y_2 \\ (c_2 - y_1) & \text{if } 0 < y_2 < y_1. \end{cases}$$

As $y_1 \rightarrow \infty$ and $y_2 \rightarrow \infty$, $|\psi_1(y_1, y_2; \theta^0)| \rightarrow \infty$. Similarly, the same behavior can be observed for the remaining components. Hence, when $\gamma = 0$, corresponding to the MLE case, $IF(y_1, y_2; T_\gamma, F_{BB}(\cdot; \theta^0))$ is unbounded $\square$

Theorem 3. For fixed (y_1, y_2) , the components of the IF, $IF(y_1, y_2; T_\gamma, F_{BB}(\cdot; \theta^0))$ as given in (13) are decreasing function of $\gamma > 0$.

Proof. The proof is the usual routine, so we avoid it. $\square$

To study the robustness, Hampel [16] introduced three important summary values of the IF apart from its expected square. The first and most important one is the supremum of the absolute value, therefore the gross-error-sensitivity (GES) of T_γ at G

$$GES = \sup_{(y_1, y_2)} \|IF(y_1, y_2; T_\gamma, G)\|, \quad (14)$$

where $\|\dots\|$ is the Euclidean norm and the supremum being taken over all observations (y_1, y_2) where $IF(y_1, y_2; T_\gamma, G)$ exists. If GES is infinite/finite the MDPDE is said to be sensitivity/B-robust to outliers.

4.3 Asymptotic Properties

Asymptotically, MDPDEs possess desirable properties that make them suitable for large-sample inference, derived from the large-sample behavior of the density power divergence. The asymptotic properties of the MDPDE for the BBBE distribution are key to understanding its long-term performance and are derived from M-estimation theory, as discussed in [14, 15].

We further explore some theoretical properties of the MDPDE $\hat{\theta}_{\gamma, n}$. Let $\theta^g = T_\gamma(G)$ is the best fitting value of the parameter, in the sense of minimizing the discrepancy $d_\gamma(g, f)$, i.e., is the best fitting MDPD functional at G corresponding to tuning parameter γ . We assume that the MDPD functional $T_\gamma(\cdot)$ is Fisher consistent, see Kallinpur and Rao [17], $T_\gamma(F_{BB}(\cdot; \theta)) = \theta \forall \theta \in \Theta$, which means at the model the estimators asymptotically measures the right quantity. The notion of Fisher consistency is more suitable and elegant for functionals than usual consistency or asymptotic unbiasedness. Consider the score function $S_\theta(y_1, y_2)$ and the Hessian matrix $H_\theta(y_1, y_2)$ of the BBBE model, their expressions are given in Appendix A. Define the 3×3 symmetric

matrices $J_\gamma(\theta)$ and $K_\gamma(\theta)$ as follows,

$$\begin{aligned} J_\gamma(\theta) &= \int_0^\infty \int_0^\infty S_\theta(y_1, y_2) S_\theta(y_1, y_2)^T f_{BB}^{1+\gamma}(y_1, y_2; \theta) dy_1 dy_2 \\ &\quad + \int_0^\infty \int_0^\infty \{I_\theta(y_1, y_2) - \gamma S_\theta(y_1, y_2) S_\theta(y_1, y_2)^T\} [g(y_1, y_2) \\ &\quad - f_{BB}(y_1, y_2; \theta)] f_{BB}^\gamma(y_1, y_2; \theta) dy_1 dy_2 \end{aligned} \quad (15)$$

where $I_\theta(y_1, y_2) = -H_\theta(y_1, y_2) = -\nabla[S_\theta(y_1, y_2)]$, the so-called information function of the model, which is positive definite for all required θ and S^T indicates the transpose of a matrix S , and

$$K_\gamma(\theta) = \int_0^\infty \int_0^\infty S_\theta(y_1, y_2) S_\theta(y_1, y_2)^T f_{BB}^{2\gamma}(y_1, y_2; \theta) g(y_1, y_2) dy_1 dy_2 - \xi_\gamma(\theta) \xi_\gamma(\theta)^T \quad (16)$$

$$\text{with} \quad \xi_\gamma(\theta) = \int_0^\infty \int_0^\infty S_\theta(y_1, y_2) f_{BB}^\gamma(y_1, y_2; \theta) g(y_1, y_2) dy_1 dy_2.$$

In the following theorem, we discuss the asymptotic properties of the MDPDE for the BBBE model. This is a bivariate extension of the univariate results discussed by Hazra [5] and Juarez and Rez [4].

Theorem 4. *For fixed $\gamma > 0$ and under the regularity conditions (D1) – (D5) of Basu et al. (page 304) [18], the MDPDEs have the following asymptotic properties, as $n \rightarrow \infty$,*

- (a) *The MDPD estimating equation (11) has consistent sequence of roots $\hat{\theta}_{\gamma,n}$, that is, $\hat{\theta}_{\gamma,n}$ is consistent for θ^g , and*
- (b) *The quantity $\sqrt{n}(\hat{\theta}_{\gamma,n} - \theta^g)$ has an asymptotic trivariate normal distribution with mean vector $\mu = (0, 0, 0)^T$ and variance-covariance matrix $\Sigma_{\theta^g} = J_\gamma(\theta^g)^{-1} K_\gamma(\theta^g) J_\gamma(\theta^g)^{-1}$, where $J_\gamma(\theta)$ and $K_\gamma(\theta)$ are as defined in equations (15) and (16) respectively.*

Proof. The BBBE distribution satisfied all the regularity conditions (D1) – (D5) of Basu et al. (page 304) [18]. The proof is along the same line as the proof of Theorem 9.2., discussed by Basu et al. [18] (page 304). Therefore, we omit the proof. $\square$

Now we provide the expression of the matrices $J_\gamma(\theta)$, $K_\gamma(\theta)$ and $\xi_\gamma(\theta)$ when the true distribution G belongs to the parametric family of BBBE distributions. Suppose, we assume that the model is correctly specified so that the data are generated from the true distribution function, $G := F_{BB}(\cdot; \theta^0)$, where $\theta^0 = (\alpha_1^0, \alpha_2^0, \alpha_3^0) \in \Theta \subset \mathbb{R}_+^3$. Then, under certain regularity conditions, see [3], we have

$$\begin{aligned} J_\gamma(\theta) &= \int_0^\infty \int_0^\infty S_\theta(y_1, y_2) S_\theta(y_1, y_2)^T f_{BB}^{1+\gamma}(y_1, y_2; \theta) dy_1 dy_2, \\ K_\gamma(\theta) &= J_{2\gamma}(\theta) - \xi_\gamma(\theta) \xi_\gamma(\theta)^T \end{aligned}$$

where,
$$\xi_\gamma(\theta) = \int_0^\infty \int_0^\infty S_\theta(y_1, y_2) f_{BB}^{1+\gamma}(y_1, y_2; \theta) dy_1 dy_2.$$

For fixed $\gamma > 0$ and the integrability conditions (A1) - (A3), given in Appendix A, are satisfied for the BBBE model. The MDPDEs $\hat{\theta}_{\gamma,n}$ is consistent for θ^0 and $\hat{\theta}_{\gamma,n} = \sqrt{n}(\hat{\theta}_{\gamma,n} - \theta^0) \sim AN(\mu, \Sigma_{\theta^0})$, where $\mu = (0, 0, 0)^T$ and $\Sigma_{\theta^0} = J_\gamma(\theta^0)^{-1} K_\gamma(\theta^0) J_\gamma(\theta^0)^{-1}$. Note that when $g = f_{BB}(\cdot; \theta^0)$ and $\gamma = 0$ (i.e., $\gamma \rightarrow 0$), one gets $J_0(\theta^0) = K_0(\theta^0) = F_n(\theta^0)$, the Fisher information matrix, so that the asymptotic covariance matrix Σ_{θ^0} is simply the inverse of $F_n(\theta^0)$.

The elements of the symmetric matrices $J_\gamma(\theta)$ and $K_\gamma(\theta)$ are derived in Appendix B. The final expressions for elements of the covariance matrix Σ_{θ^0} are functions of the model parameters α_1, α_2 and α_3 , and the MDPD tuning parameter γ . We plot $\Sigma_{\theta^0}^{(i,j)}$, the $(i, j)^{th}$ element of Σ_{θ^0} for $i, j = 1, 2, 3$ for different values of tuning parameter γ , given in Figure 1. The term $\Sigma_{\theta^0}^{(i,i)}$ is the asymptotic variance of $\sqrt{n}\hat{\alpha}_i$, increases as the tuning parameter γ increases. The asymptotic covariance $\Sigma_{\theta^0}^{(1,2)}$ between $\sqrt{n}\hat{\alpha}_1$ and $\sqrt{n}\hat{\alpha}_2$ increases when γ increases, and the asymptotic covariance $\Sigma_{\theta^0}^{(1,3)}$ between $\sqrt{n}\hat{\alpha}_1$ and $\sqrt{n}\hat{\alpha}_3$, and $\Sigma_{\theta^0}^{(2,3)}$ between $\sqrt{n}\hat{\alpha}_2$ and $\sqrt{n}\hat{\alpha}_3$ decrease when γ increases. Thus the MDPDE with higher γ is more robust to outliers but less efficient.

The consistent estimates of the matrices J_γ and K_γ are obtained by substituting $\hat{\theta}_{\gamma,n}$ in these matrices, we obtain $n^{-1}\hat{J}_\gamma(\hat{\theta})^{-1}\hat{K}_\gamma(\hat{\theta})\hat{J}_\gamma(\hat{\theta})^{-1}$ as a consistent estimates of the covariance matrix of the MDPDE $\hat{\theta}_{\gamma,n}$. The asymptotic confidence interval of the parameter is given as $\hat{\theta}_{i,\gamma,n} \pm z_{\alpha/2} \frac{\sigma_{ii}}{\sqrt{n}}$ where $z_{\alpha/2}$ is the upper $(\alpha/2)^{th}$ percentile of the standard normal distribution and σ_{ii}^2 is the i^{th} diagonal element of the covariance matrix Σ_{θ^0} , $i = 1, 2, 3$.

4.4 Asymptotic Relative Efficiency

One can verify that the asymptotic variance of the MDPDE is minimum when $\gamma = 0$, i.e., the case of MLE. Among two unbiased estimators, we generally prefer an estimator that has smaller variance, and hence, we need to study the asymptotic relative efficiency (ARE), the ratio of the asymptotic variance of the MLE over that of the MDPDE, assuming no outlier is present in the data set. If ARE is close to one, standard errors of the estimators for both the methods are comparable and hence, robustness can be achieved with only a small compromise in uncertainty. By definition, the ARE of the estimator in MDPDE is one if $\gamma = 0$. The elements the matrices $J_\gamma(\theta)$ and $K_\gamma(\theta)$, given in Appendix II, and hence $\Sigma_{\hat{\theta}}$ can be simplified only up to in terms of polynomial functions. From Figure 2 and Figure 3, we observe that all AREs for the estimators of α_1, α_2 and α_3 decrease as γ increases; however, the loss is not high for smaller γ . Thus, under the no outlier scenario, the asymptotic variances of the MDPDEs are comparable with those of the MLEs at least for smaller γ values and so in case MDPDE provides high robustness against outlier, it is justified to use MDPDE over MLE.

Fig. 1: Variation of the elements of the covariance matrix $\Sigma_{\hat{\theta}}$ for different values of the tuning parameter γ .

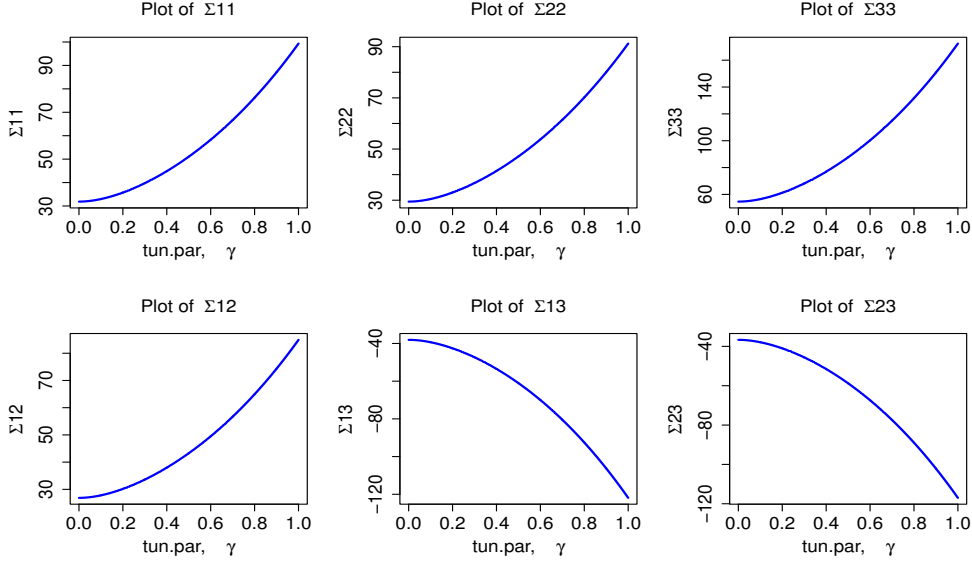
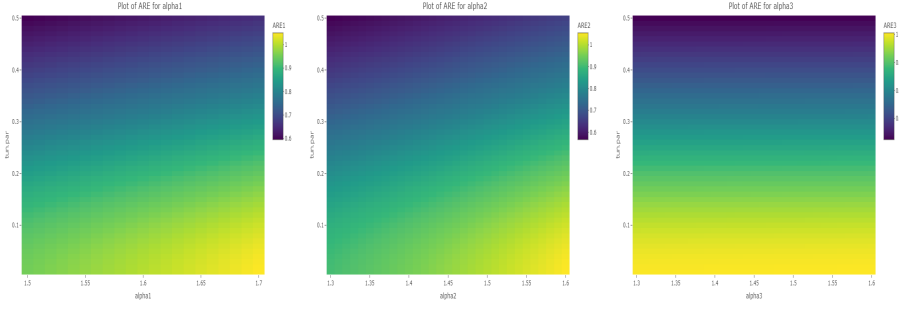


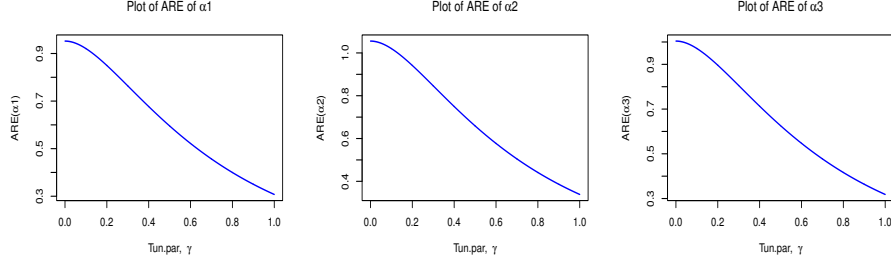
Fig. 2: Variation of ARE for the estimates of α_1, α_2 and α_3 for different values of the tuning parameter γ .



5 Selection of optimal tuning parameter

The MDPDE is a robust and efficient variant with a small γ which has already been suggested and studied. We wish to use a large γ if we hope for improved in robustness. However, the higher value of γ would reduce the efficiency, as $\gamma = 0$ represents the MLE with full asymptotic efficiency. An ideal tuning parameter for optimizing robustness and efficiency will exist. The γ in the MDPDE balance the robustness and efficiency of the underlying estimators. Thus, γ needs to be ideally chosen depending on the proportion of outliers in the data. If there are no outliers in the data, the best choice is $\gamma = 0$ (MLE). A bigger γ should be appropriate for a higher proportion of outliers,

Fig. 3: Variation of ARE for the estimates of α_1, α_2 and α_3 for different values of the tuning parameter γ .



and vice versa. However, the outlier proportion is not known in practice, and thus, a data-driven approach is needed to select an optimal value of γ . Here, we look at the divergence based numerical technique as follow the approach proposed by [11], the empirical Cramer-von Mise (ECVM) distance can robustly be approximated by

$$ECVM = \frac{1}{n} \sum_{i=1}^n \left[\frac{n_i}{n+1} - F_{BB}(y_{1i}, y_{2i}; \hat{\theta}_\gamma) \right]^2, \quad (17)$$

where $F_{BB}(y_1, y_2; \theta)$ denotes the CDF of the BBBE model. Taking the idea of cross-validation, the optimal γ is chosen by minimizing the ECVM distance

$$\hat{\gamma} = \underset{\gamma}{\operatorname{argmin}} \frac{1}{n} \sum_{i=1}^n \left[\frac{n_i}{n+1} - F_{BB}(y_{1i}, y_{2i}; \hat{\theta}_\gamma^{(-i)}) \right]^2, \quad (18)$$

where n_i is the number of observations (y_{1j}, y_{2j}) such that $y_{1j} \leq y_{1i}$ and $y_{2j} \leq y_{2i}$; $j = 1, 2, \dots, n$, i.e., $n_i = \sum_{j=1}^n \mathbb{1}(y_{1j} \leq y_{1i}, y_{2j} \leq y_{2i})$, $\mathbb{1}(\cdot)$ is the indicator function, and $\hat{\theta}_\gamma^{(-i)}$ is the MDPDE with tuning parameter γ obtained after removing the i^{th} , $i = 1, 2, 3, \dots, n$, observation (y_{1i}, y_{2i}) from the sample.

6 Simulation Study

In this section, we present some simulation results to evaluate the effectiveness of the proposed MDPDE for different choices of the tuning parameter γ , and compare its performance with that of the MLE across different sample sizes, different proportion of outliers and different parameter values of the BBBE model. The MDPDEs of three unknown parameters α_1, α_2 and α_3 of the BBBE model are obtained numerically by existing standard optimize technique, namely Nelder-Mead method (employing under “optim()” function in the R environment) and the MLEs obtained by the EM algorithm discussed by Kundu and Gupta [19]. First we perform the simulation without considering the outlier, i.e., under the true data scenario (0% outlier). We generate 1000 random samples of size n , namely $n = 25, 50, 75, 100, 150, 250, 500, 1000$,

from the true data-generating distribution $BBBE(1.5, 1.6, 1.4)$. For each of the 1000 replications, we compute the MDPDEs for various values of the tuning parameter $\gamma = 0.01, 0.1, 0.3, 0.5, 0.7, 0.9$, as well as the MLEs of α_1, α_2 and α_3 . We then calculate the average biases (ABs) and mean squared errors (MSEs) of the resulting estimates. The simulation results are summarized in Table 1. From Table 1, we observe that the ABs and MSEs for all the estimation methods decrease as the sample size increases. These indicate that the MDPDEs and the MLEs are consistent, as expected. Under the true data scenario most of the cases the MDPDEs and MLEs have negative biases for α_1 and α_2 and positive biases for α_3 . The ABs and MSEs for MDPDEs with small values of γ (close to zero, e.g., 0.01) are similar to those of the MLEs, confirming that MDPDEs converge to the MLEs as $\gamma \rightarrow 0$. Additionally, the ABs and MSEs of the MDPDEs increase with larger values of γ , especially for larger sample sizes. These results further support the consistency of the MDPDEs and their equivalence to MLEs when the tuning parameter γ is small. The corresponding asymptotic variances and lengths of the 95% confidence intervals (LCIs) for all estimation methods are calculated and reported in Table 2. We observe that the asymptotic variances and LCIs for MDPDEs with small values of γ are similar to those of MLEs and they increase as γ increases. Thus the asymptotic covariance matrix of the MDPDEs tend to the inverse of the Fisher information matrix as γ tend to zero. They decrease as the sample size increases. We also experimented with other parameter values (results not reported here), and similar trends were observed.

Furthermore, we conducted an extensive simulation study to evaluate and compare the performance of the proposed MDPDE, for various choices of γ , with that of the MLE in the presence of outliers in the data. Now we consider five different contamination or outlier proportions 1%, 5%, 10%, 20% and 30%, respectively. The samples with outliers are generated from a mixture of the BBBE distribution and a relevant contaminating bivariate distribution $H(\cdot)$, i.e., $(1 - \pi)BBBE(\alpha_1, \alpha_2, \alpha_3) + \pi H(\cdot)$, where $\pi \in [0, 1]$ is the outlier proportion. Here we take the Block-Basu bivariate Weibull (BBBW) distribution with four parameters, denoted as $BBBW(\eta_1, \eta_2, \eta_3, \lambda)$. The average bias (AB) and the mean square error (MSE) give us an idea about the bias, consistency and robustness of the estimators. In each case, we compute the ABs and MSEs over the 10^3 replications. For the samples of size $n = 100$ generated from the $(1 - \pi)BBBE(1.5, 1.6, 1.4) + \pi BBBW(0.001, 0.001, 0.001, 1.5)$ with different values of $\pi = 0, 0.01, 0.05, 0.1, 0.15, 0.20, 0.30$, the results are reported in Table 3. We have also tried for the different outlier model as $BBBG(10, 10, 10, 5)$, for different sample sizes and different sets of parameter values, but not reported here. From Table 3, it is evident that the ABs and MSEs of both the MDPDEs and the MLEs increase as the proportion of outliers increases, which is expected under contamination. Although the MDPDEs become less efficient in absolute terms with higher levels of contamination, they still perform better than the MLEs, demonstrating greater robustness. In this set up, the MDPDEs with $\gamma = 0.1$ performs better than both the MDPDEs with other values of γ and the MLE in terms of bias and MSE. Thus, the MDPDE with an optimal tuning parameter γ is more robust than the MLE. It effectively reduces the influence of outliers, resulting in more stable and reliable estimates under data contamination. Based on the experimental results, we observe that the MDPD estimation

Table 1: Average bias (AB) and MSE of estimates for BBBE(1.5, 1.6, 1.4).

Method	θ		sample sizes, $n \rightarrow$							
			25	50	75	100	150	250	500	1000
MLE $\gamma = 0$	α_1	AB	-0.2108	-0.1407	-0.0734	-0.0384	-0.0259	-0.0070	0.0015	-0.0012
		MSE	0.5937	0.4260	0.2954	0.2533	0.1784	0.1158	0.0582	0.0294
	α_2	AB	-0.1971	-0.1397	-0.0571	-0.0431	-0.0415	-0.0089	-0.0017	0.0012
		MSE	0.7337	0.4626	0.3487	0.2719	0.2026	0.1211	0.0631	0.0305
	α_3	AB	0.4237	0.2662	0.1582	0.1077	0.0839	0.0222	0.0037	0.0026
		MSE	1.2323	0.8105	0.5747	0.4685	0.3497	0.2085	0.1084	0.0530
MDPDE $\gamma = 0.01$	α_1	AB	-0.2154	-0.1440	-0.0757	-0.0398	-0.0271	-0.0083	0.0017	-0.0010
		MSE	0.5880	0.4216	0.2934	0.2517	0.1775	0.1152	0.0579	0.0290
	α_2	AB	-0.2016	-0.1425	-0.0591	-0.0441	-0.0426	-0.0099	-0.0016	0.0014
		MSE	0.7261	0.4582	0.3465	0.2702	0.2015	0.1206	0.0628	0.0301
	α_3	AB	0.4258	0.2677	0.1599	0.1083	0.0848	0.0233	0.0034	0.0021
		MSE	1.2195	0.8014	0.5714	0.4652	0.3473	0.2078	0.1079	0.0522
MDPDE $\gamma = 0.1$	α_1	AB	-0.2533	-0.1706	-0.0979	-0.0565	-0.0400	-0.0181	0.0054	-0.0035
		MSE	0.5607	0.4000	0.2866	0.2485	0.1791	0.1175	0.0604	0.0304
	α_2	AB	-0.2368	-0.1651	-0.0800	-0.0593	-0.0552	-0.0192	-0.0025	-0.0018
		MSE	0.6918	0.4361	0.3405	0.2671	0.2032	0.1235	0.0652	0.0312
	α_3	AB	0.4532	0.2880	0.1836	0.1251	0.0994	0.0352	0.0079	0.0064
		MSE	1.1705	0.7643	0.5660	0.4596	0.3482	0.2135	0.1130	0.0544
MDPDE $\gamma = 0.3$	α_1	AB	-0.3033	-0.2231	-0.1473	-0.0955	-0.0646	-0.0381	-0.0092	-0.0074
		MSE	0.5830	0.4025	0.3072	0.2722	0.2026	0.1375	0.0725	0.0374
	α_2	AB	-0.2764	-0.2102	-0.1265	-0.0969	-0.0805	-0.0373	-0.0081	-0.0073
		MSE	0.7271	0.4395	0.3686	0.2936	0.2307	0.1460	0.0784	0.0381
	α_3	AB	0.5117	0.3515	0.2536	0.1770	0.1357	0.0640	0.0154	0.0149
		MSE	1.2243	0.7792	0.6229	0.5076	0.3943	0.2570	0.1368	0.0667
MDPDE $\gamma = 0.5$	α_1	AB	-0.3239	-0.2664	-0.1984	-0.1430	-0.0919	-0.0575	-0.0122	-0.0116
		MSE	0.6629	0.4517	0.3602	0.3221	0.2455	0.1729	0.0917	0.0493
	α_2	AB	-0.2894	-0.2449	-0.1740	-0.1440	-0.1096	-0.0552	-0.0169	-0.0129
		MSE	0.8384	0.5014	0.4342	0.3543	0.2829	0.1855	0.1000	0.0502
	α_3	AB	0.5563	0.4187	0.3347	0.2486	0.1816	0.0946	0.02993	0.0246
		MSE	1.3397	0.8971	0.7388	0.6200	0.4877	0.3308	0.1754	0.0886
MDPDE $\gamma = 0.7$	α_1	AB	-0.3285	-0.3043	-0.2513	-0.1944	-0.1258	-0.0783	-0.0242	-0.0170
		MSE	0.7818	0.5356	0.4375	0.4038	0.3066	0.2226	0.1178	0.0657
	α_2	AB	-0.2905	-0.2717	-0.2233	-0.1946	-0.1455	-0.0752	-0.0288	-0.0199
		MSE	1.0170	0.6134	0.5345	0.4543	0.3576	0.2428	0.1300	0.0675
	α_3	AB	0.5910	0.4827	0.4219	0.3290	0.2387	0.1281	0.0494	0.0361
		MSE	1.4539	1.0877	0.9120	0.8044	0.6260	0.4343	0.2289	0.1200
MDPDE $\gamma = 0.9$	α_1	AB	-0.3197	-0.3359	-0.3024	-0.2513	-0.1647	-0.1011	-0.0397	-0.0238
		MSE	1.0154	0.6678	0.5576	0.5273	0.3932	0.2881	0.1518	0.0866
	α_2	AB	-0.2789	-0.2904	-0.2688	-0.2502	-0.1867	-0.0968	-0.0440	-0.0279
		MSE	1.2816	0.8027	0.6894	0.5986	0.4632	0.3221	0.1695	0.0905
	α_3	AB	0.6167	0.5413	0.5074	0.4191	0.3045	0.1647	0.0739	0.0491
		MSE	1.6579	1.3735	1.1843	1.0741	0.8245	0.5739	0.2998	0.1618

method is quite effective in producing robust and efficient estimates of the parameters α_1, α_2 and α_3 of the BBBE model, especially in the presence of outliers.

7 Real Data Analysis

In this section, we examine a real data set for illustrative purposes. We consider the kidney infection data set, initially analyzed by McGilchrist and Aisbett [20]. This study recorded the infection recurrence times of kidney patients using a portable dialysis

Table 2: Asymptotic variance (Var) and length of 95% confidence interval (LCI) of estimates for BBBE(1.5, 1.6, 1.4).

Method	θ		sample sizes, $n \rightarrow$							
			25	50	75	100	150	250	500	1000
MLE $\gamma = 0$	α_1	Var	1.1811	0.5935	0.3959	0.2981	0.1991	0.1180	0.0590	0.0294
		LCI	4.1306	2.9704	2.4429	2.1237	1.7409	1.3427	0.9507	0.6718
	α_2	Var	1.2702	0.6373	0.4281	0.3210	0.2128	0.1274	0.0637	0.0319
		LCI	4.3092	3.0962	2.5466	2.2104	1.8023	1.3968	0.9889	0.7002
	α_3	Var	2.1768	1.1057	0.7340	0.5513	0.3678	0.2180	0.1091	0.0546
		LCI	5.6732	4.0874	3.3428	2.8998	2.3718	1.8283	1.2940	0.9156
MDPDE $\gamma = 0.01$	α_1	Var	1.0927	0.5625	0.3845	0.2914	0.1968	0.1168	0.0588	0.0294
		LCI	3.9869	2.8983	2.4102	2.1017	1.7312	1.3364	0.9491	0.6714
	α_2	Var	1.1872	0.6155	0.4184	0.3149	0.2097	0.1264	0.0634	0.0318
		LCI	4.1762	3.0459	2.5192	2.1904	1.7893	1.3916	0.9866	0.6993
	α_3	Var	2.0265	1.0594	0.7161	0.5402	0.3631	0.2162	0.1086	0.0545
		LCI	5.4868	4.0045	3.3025	2.8713	2.3566	1.8205	1.2913	0.9147
MDPDE $\gamma = 0.1$	α_1	Var	1.1006	0.5716	0.3943	0.2990	0.2024	0.1204	0.0608	0.0304
		LCI	3.9936	2.9195	2.4397	2.1282	1.7555	1.3564	0.9650	0.6825
	α_2	Var	1.1975	0.6291	0.4297	0.3239	0.2160	0.1305	0.0656	0.0329
		LCI	4.1890	3.0780	2.5522	2.2210	1.8157	1.4135	1.0030	0.7109
	α_3	Var	2.0484	1.0827	0.7375	0.5563	0.3745	0.2233	0.1123	0.0563
		LCI	5.5099	4.0468	3.3507	2.9132	2.3931	1.8501	1.3130	0.9301
MDPDE $\gamma = 0.3$	α_1	Var	1.2816	0.6693	0.4691	0.3557	0.2426	0.1443	0.0732	0.0366
		LCI	4.2776	3.1491	2.6559	2.3174	1.9192	1.4840	1.0589	0.7492
	α_2	Var	1.4092	0.7464	0.5135	0.3874	0.2590	0.1570	0.0790	0.0397
		LCI	4.5169	3.3445	2.7858	2.4258	1.9870	1.5501	1.1009	0.7805
	α_3	Var	2.3927	1.2833	0.8873	0.6677	0.4509	0.2693	0.1355	0.0680
		LCI	5.9297	4.3985	3.6710	3.1892	2.6245	2.0310	1.4419	1.0217
MDPDE $\gamma = 0.5$	α_1	Var	1.6068	0.8430	0.5949	0.4510	0.3100	0.1843	0.0938	0.0470
		LCI	4.7467	3.5179	2.9817	2.6037	2.1664	1.6755	1.1981	0.8485
	α_2	Var	1.7940	0.9496	0.6549	0.4935	0.3311	0.2016	0.1014	0.0510
		LCI	5.0563	3.7591	3.1387	2.7326	2.2440	1.7552	1.2467	0.8844
	α_3	Var	2.9881	1.6302	1.1388	0.8566	0.5802	0.3468	0.1745	0.0876
		LCI	6.5939	4.9458	4.1527	3.6086	2.9752	2.3039	1.6359	1.1597
MDPDE $\gamma = 0.7$	α_1	Var	2.0631	1.0842	0.7664	0.5821	0.4026	0.2393	0.1220	0.0612
		LCI	5.3296	3.9696	3.3726	2.9506	2.4646	1.9072	1.3659	0.9685
	α_2	Var	2.3236	1.2315	0.8493	0.6393	0.4301	0.2632	0.1323	0.0666
		LCI	5.7056	4.2641	3.5646	3.1030	2.5542	2.0040	1.4235	1.0106
	α_3	Var	3.7840	2.1070	1.4841	1.1185	0.7600	0.4548	0.2288	0.1150
		LCI	7.3805	5.6075	4.7331	4.1186	3.4027	2.6372	1.8729	1.3285
MDPDE $\gamma = 0.9$	α_1	Var	2.6707	1.4003	0.9884	0.7522	0.5230	0.3109	0.1587	0.0798
		LCI	6.0020	4.4865	3.8154	3.3444	2.8034	2.1713	1.5564	1.1053
	α_2	Var	3.0202	1.6015	1.1047	0.8300	0.5588	0.3439	0.1727	0.0870
		LCI	6.4469	4.8425	4.0531	3.5260	2.9068	2.2887	1.6256	1.1549
	α_3	Var	4.7944	2.7220	1.9323	1.4617	0.9954	0.5967	0.3001	0.1510
		LCI	8.2496	6.3531	5.3902	4.7016	3.8911	3.0195	2.1449	1.5222

machine. Recurrence times are the times from infection until the subsequent infection. For each patient, two such recurrence times are recorded: the first recurrence time (FRT) in days, denoted by Y_1 , and the second recurrence time (SRT) in days, by Y_2 . Now, many questions arise, e.g., what is the effect of portable dialysis on infection in the kidney? Were both recurrence times coming from the same distribution? Are FRT and SRT independently distributed? Now we are giving the answer to all such questions by fitting the BBBE model to the data set by using the robust and efficient

Table 3: Average bias (AB) and MSE for BBBE(1.5, 1.6, 1.4) when outliers coming from BBBW(0.001, 0.001, 0.001, 1.5).

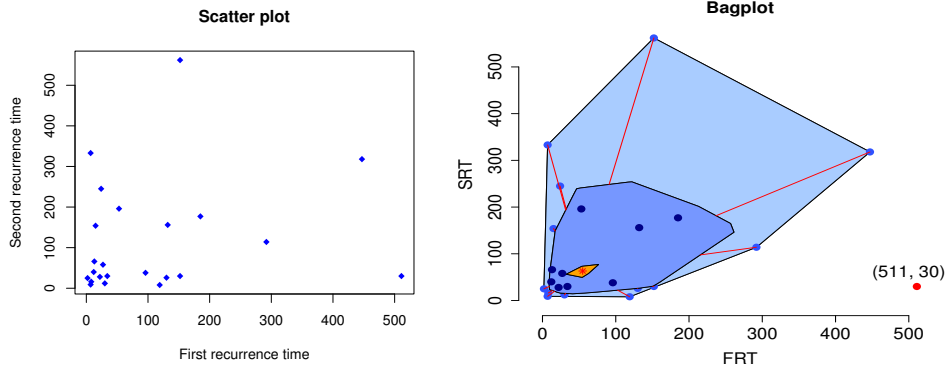
Method	θ		outlier proportion, $\pi \rightarrow$						
			0%	1%	5%	10%	15%	20%	30%
MLE	α_1	AB	-0.0384	-0.6398	-1.3767	-1.4639	-1.4825	-1.4854	-1.4877
		MSE	0.2533	0.9871	1.9615	2.1483	2.1990	2.2109	2.2203
	α_2	AB	-0.0431	-0.6974	-1.4755	-1.5620	-1.5830	-1.5853	-1.5868
		MSE	0.2719	1.1342	2.2554	2.4469	2.5068	2.5174	2.5253
	α_3	AB	0.1077	-0.3089	-0.9944	-1.1819	-1.2407	-1.2771	-1.3093
		MSE	0.4685	0.8352	1.1135	1.4132	1.5471	1.6423	1.7296
MDPDE 0.01	α_1	AB	-0.0398	-0.3725	-1.2412	-1.4450	-1.4804	-1.4884	-1.4951
		MSE	0.2517	0.5481	1.7281	2.1051	2.1932	2.2157	2.2353
	α_2	AB	-0.0441	-0.3939	-1.3330	-1.5434	-1.5809	-1.5880	-1.5948
		MSE	0.2702	0.6159	1.9741	2.4018	2.5006	2.5223	2.5436
	α_3	AB	0.1083	0.1937	-0.5637	-1.0988	-1.2168	-1.2623	-1.2872
		MSE	0.4652	0.5745	0.7658	1.2629	1.4891	1.5991	1.6690
MDPDE 0.1	α_1	AB	-0.0565	-0.1118	-0.1466	-0.1181	-0.1470	-0.1765	-0.2168
		MSE	0.2485	0.2627	0.2839	0.2871	0.2749	0.3185	0.3548
	α_2	AB	-0.0593	-0.1184	-0.1592	-0.1276	-0.1566	-0.1885	-0.2213
		MSE	0.2671	0.2865	0.3030	0.3217	0.3081	0.3398	0.3955
	α_3	AB	0.1251	0.2001	0.2376	0.1658	0.1899	0.2144	0.1969
		MSE	0.4596	0.5089	0.5713	0.5611	0.5087	0.5752	0.6281
MDPDE 0.3	α_1	AB	-0.0955	-0.1617	-0.2111	-0.2012	-0.2433	-0.2982	-0.3605
		MSE	0.2722	0.3015	0.3206	0.3266	0.3091	0.3710	0.4091
	α_2	AB	-0.0969	-0.1692	-0.2275	-0.2135	-0.2530	-0.3189	-0.3772
		MSE	0.2936	0.3337	0.3432	0.3650	0.3538	0.4046	0.4506
	α_3	AB	0.1770	0.2692	0.2896	0.2046	0.1984	0.2140	0.1190
		MSE	0.5076	0.6111	0.6393	0.6021	0.5178	0.5782	0.5606
MDPDE 0.5	α_1	AB	-0.1430	-0.2166	-0.2744	-0.2767	-0.3272	-0.3933	-0.4637
		MSE	0.3221	0.3690	0.3903	0.3960	0.3812	0.4561	0.4971
	α_2	AB	-0.1440	-0.2248	-0.2946	-0.2922	-0.3360	-0.4221	-0.4886
		MSE	0.3543	0.4144	0.4215	0.4462	0.4388	0.5053	0.5494
	α_3	AB	0.2486	0.3541	0.3565	0.2618	0.2317	0.2359	0.0898
		MSE	0.6200	0.7762	0.7670	0.7062	0.5961	0.6466	0.5763
MDPDE 0.7	α_1	AB	-0.1944	-0.2701	-0.3328	-0.3432	-0.3929	-0.4615	-0.5303
		MSE	0.4038	0.4606	0.4876	0.4825	0.4774	0.5522	0.5831
	α_2	AB	-0.1946	-0.2781	-0.3562	-0.3608	-0.4000	-0.4965	-0.5596
		MSE	0.4543	0.5212	0.5325	0.5488	0.5477	0.6179	0.6487
	α_3	AB	0.3290	0.4400	0.4314	0.3340	0.2819	0.2782	0.0999
		MSE	0.8044	0.9912	0.9519	0.8590	0.7344	0.7646	0.6338
MDPDE 0.9	α_1	AB	-0.2513	-0.3157	-0.3823	-0.3984	-0.4430	-0.5084	-0.5759
		MSE	0.5273	0.5782	0.6049	0.5866	0.5895	0.6555	0.6688
	α_2	AB	-0.2502	-0.3238	-0.4069	-0.4179	-0.4495	-0.5472	-0.6086
		MSE	0.5986	0.6527	0.6644	0.6692	0.6702	0.7352	0.7461
	α_3	AB	0.4191	0.5176	0.5030	0.4094	0.3391	0.3249	0.1335
		MSE	1.0741	1.2552	1.1778	1.0623	0.9208	0.9199	0.7216

MDPD estimation procedure. The data are recorded for $n = 23$ patients, given in Table 4. In this data set we have $n_1 = 13$ patients for $Y_1 < Y_2$ and $n_2 = 10$ patients for $Y_1 > Y_2$. Note that the Spearman's rho and Kendall's tau between FRT and SRT are given by 0.266 and 0.184, respectively. This observation demonstrates a clear positive dependence between them. [Therefore, the BBBE model may be appropriate for analyzing the data.](#)

Table 4: Kidney infection data of 23 patients

S.N.	FRT	SRT	S.N.	FRT	SRT	S.N.	FRT	SRT	S.N.	FRT	SRT
1	8	16	2	22	28	3	447	318	4	30	12
5	24	245	6	7	9	7	511	30	8	53	196
9	15	154	10	7	333	11	96	38	12	185	177
13	292	114	14	152	562	15	13	66	16	12	40
17	132	156	18	34	30	19	2	25	20	130	26
21	27	58	22	152	30	23	119	8			

We use the bagplot to get an idea about the existence of outliers in data. A bagplot is a bivariate boxplot. This plot can read like a boxplot: the orange central region is the bivariate median, the dark blue middle region “the bag” is the bivariate interquartile range (IQR) (it contains the 50% most central points), and the light blue outer region “the fence” contains the points that are further away (but not enough that they would be considered outliers). The data points outside the fence are outliers as far as the bagplot is concerned; for more details, see [21]. The scatter plot and bagplot of the data are given in Figure 4. From the bagplot, see Figure 4, there is one outlier point (511, 30) in the data set. Based on some prior calculations we have seen that

Fig. 4: Scatter plot and bagplot of the kidney infection data

the BBBE model is not suitable to fit the kidney infection data. That means the data does not follow the BBBE distribution. Thus before analyzing the data we need to transform all the data points. After some prior analysis we observe that the transformation $Y_1^* = (Y_1/100)^{0.75}$ and $Y_2^* = (Y_2/100)^{0.75}$ is more appropriate for the analysis of data by using the BBBE distribution. We make this transformation mainly for computational and illustration purposes.

We begin with exploratory analyses to assess whether the BBBE distribution adequately fits the data after removing the outliers. To fit the BBBE distribution to the data set, we first identify potential outliers using the widely adopted bagplot method. The model parameters are then estimated using the MDPDE on the full data set, including outliers, with the optimal tuning parameter selected as outlined

in Section 5. In the context of the pure BBBE distribution, these outliers are considered as contamination and are therefore naturally excluded during the goodness-of-fit (GoF) assessment. Before analyzing the data set without outliers using the BBBE distribution, we first fit the marginal distributions to Y_1^* , Y_2^* and $\min\{Y_1^*, Y_2^*\}$ separately. In addition to model checking, this also helps in selecting initial values. The Kolmogorov-Smirnov (KS) distances with associated p -value (in braket) for Y_1^* , Y_2^* and $\min\{Y_1^*, Y_2^*\}$ are 0.1321 (0.8372), 0.1545 (0.6699), and 0.1698 (0.5499), respectively. Based on these results, the marginal distributions and their minimum of the BBBE model cannot be rejected to fit marginlas Y_1^* and Y_2^* , and minimum $\min\{Y_1^*, Y_2^*\}$ of the data set. This indicates that the BBBE model is suitable for fitting the kidney infection data set.

After removing the outliers, the BBBE model fits the data set (Y_1^*, Y_2^*) , and the MLEs of α_1 , α_2 and α_3 are obtained as 0.4664, 0.3111, and 1.0232, respectively. To evaluate model adequacy, we employ a bivariate GoF test based on statistical equivalence blocks (SEBs), as described in [22], along with marginal GoF assessments using the KS test. The p -value for the bivariate GoF test is obtained via a parametric bootstrap procedure. The KS distance for the bivariate GoF test is 0.1345, with an associated p -value is 0.3333. The results of the marginal GoF tests are provided in Table 5. Based on the p -values from both the bivariate and marginal tests, we can safely assume that, after removing the outliers, the Kidney infection data set can be reasonably modeled by the BBBE distribution.

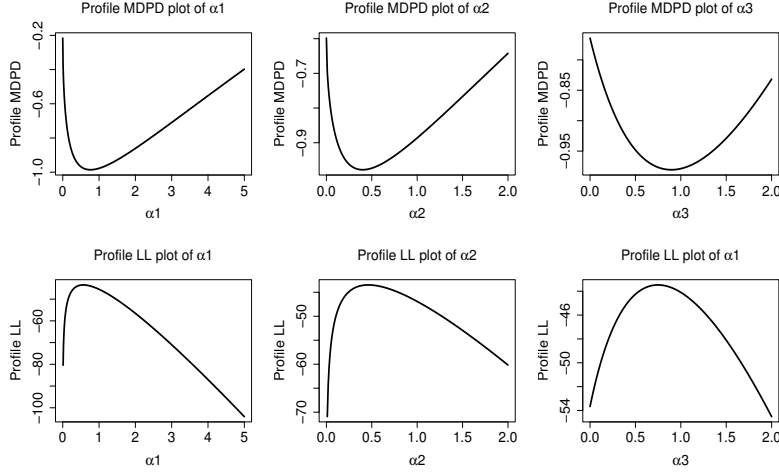
Table 5: Marginal GoF test

Var.	KS distance	p -value
X_1	0.1483	0.7183
X_2	0.1518	0.6912
$\min\{X_1, X_2\}$	0.1719	0.5341

We now fit the BBBE model to the full data set, including the outliers. To find the MLEs and MDPDEs we need the initial values of the parameters and optimal tuning parameter. To get initial estimates of α_1 , α_2 and α_3 , we fit the exponential distribution with parameter $\alpha = \alpha_1 + \alpha_2 + \alpha_3$ to the $\min\{Y_1^*, Y_2^*\}$, and the estimate is obtained as $\hat{\alpha} = 1.7920$. The initial value for α_1 is $\alpha_1^{(0)} = 1.7920/3 = 0.5973$ when $n_1 > n_2$ and for α_2 is obtained by using the relation $P(Y_1 < Y_2) = \alpha_1/(\alpha_1 + \alpha_2)$ and $n_1/(n_1 + n_2)$ then we have $\alpha_2^{(0)} = (\alpha_1^{(0)} - n_1/n)/(n_1/n) = 0.4595$, where $n_1/n = 0.5652$ and hence the initial value of α_3 is $\alpha_3^{(0)} = 1.7620 - 0.5973 - 0.4595 = 0.7353$. Using these initial values (0.5973, 0.4595, 0.7353) and the idea of cross-validation into the empirical Cramer-von Mises (ECVM) distance, given in equation (18), the optimal γ is obtained as $\gamma^* = 0.549$ and corresponding ECVM distance is 0.0043.

The plots of profile log-likelihood and profile log modified likelihood function (i.e., MDPD function) for each parameter are given in Figure 5. These plots ensure that there are unique estimates of the parameters. The MLEs and MDPDEs with optimal and other tuning parameters of the model parameters for the data are obtained. The standard error (SE) and total variation (Tvar) of the estimates are calculated to

Fig. 5: Plots of profile minimum density power divergence (MDPD) and log-likelihood (LL) for the kidney infection data.



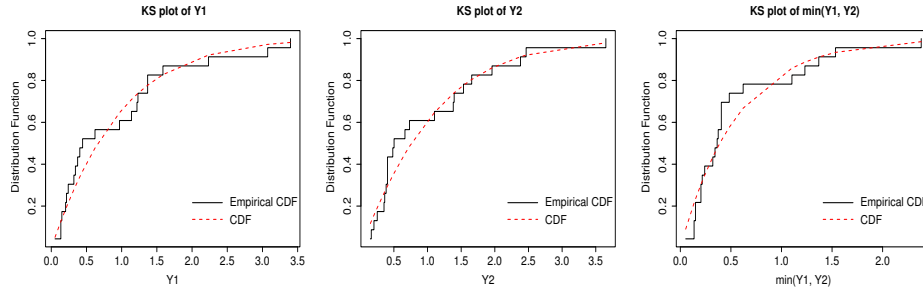
see their efficiency. The GES of estimates are obtained to check their robustness. The asymptotic 95% confidence intervals and ECVM distance, see equation (17), and associated bootstrap p -value are calculated to check the goodness of fit. The p -value for the ECVM are obtained by the parametric bootstrap method, proposed by Bradley Efron [23]. These empirical results are given in the Table 6. The estimation method with smallest ECVM distance is better fit the data than other methods. Thus from Table 6 we have seen that the MDPDEs with optimal tuning parameter $\gamma^* = 0.549$ that have smallest ECVM distance is better fit the kidney infection data set than other estimation methods. Also we have seen the MDPDEs with $\gamma^* = 0.549$ have smallest GES but they have Tvar larger than that of the MLEs. Therefore the MDPDEs with $\gamma^* = 0.549$ is more robust than MDPDEs with other γ and MLEs but less efficient than MLEs. From Table 6 we observed that the MDPDEs with γ which is close to zero are similar to MLEs. Now the natural question is whether the BBBE model using the MDPDE with optimal γ^* fits the data well or not. First we consider the bivariate goodness-of-fit (GoF) test using the statistically equivalence blocks (SEBs), see [24] [25]. Using the SEBs, the bivatiante GoF test transformed into univariate KS GoF test. The p -value and critical value are obtained by using the parametric bootstrap, see Bradley Efron [26]. For the bivariate GoF test, the KS distance is 0.1384 and the corresponding critical value is 0.3165 and the associated p -value is 0.6360. It indicate that the BBBE model using robust and efficient MDPD estimation method with the optimal tuning parameter $\gamma^* = 0.549$ fits the data well. Now we consider the univariate GoF test for marginals and their minimum, the KS distances and the corresponding p -values are reported in Table 7. It indicates that the marginals can be used for analyzing Y_1 , Y_2 and $\min\{Y_1, Y_2\}$, also see the KS distance plots 6 and percentile-percentile (PP) plot 7. But it does not guarantee that (Y_1, Y_2) will have BBBE model, it just give an indication that the BBBE model may be used to analyze this bivariate data set.

Table 6: Estimates with SE, 95% CI, ECVM with p -value, GES and Tvar

Methods	Estimates	SE	95% CI	ECVM	p -value	GES	Tvar
MLE (EM) $\gamma = 0$	$\hat{\alpha}_1 = 0.5682$ $\hat{\alpha}_2 = 0.4651$ $\hat{\alpha}_3 = 0.7399$	0.4804 0.4376 0.6303	(0.0000, 1.5097) (0.0000, 1.3227) (0.0000, 1.9753)	0.0044	0.2542	5.2082e+6	0.8194
MLE (optim) ($\gamma = 0$)	$\hat{\alpha}_1 = 0.5681$ $\hat{\alpha}_2 = 0.4650$ $\hat{\alpha}_3 = 0.7400$	0.4803 0.4374 0.6301	(0.0000, 1.5095) (0.0000, 1.3224) (0.0000, 1.9750)	0.0044	0.2547	5.2082e+6	0.8191
MDPDE ($\gamma = 0.01$)	$\hat{\alpha}_1 = 0.5668$ $\hat{\alpha}_2 = 0.4619$ $\hat{\alpha}_3 = 0.7436$	0.4848 0.4172 0.6196	(0.0000, 1.5169) (0.0000, 1.2796) (0.0000, 1.9580)	0.0044	0.2476	111.5516	0.7929
MDPDE ($\gamma = 0.1$)	$\hat{\alpha}_1 = 0.5594$ $\hat{\alpha}_2 = 0.4292$ $\hat{\alpha}_3 = 0.7890$	0.4991 0.4081 0.6256	(0.0000, 1.5377) (0.0000, 1.2290) (0.0000, 2.0151)	0.0042	0.2057	14.3922	0.8070
MDPDE (optimal) ($\gamma^* = 0.549$)	$\hat{\alpha}_1 = 0.3013$ $\hat{\alpha}_2 = 0.1656$ $\hat{\alpha}_3 = 1.3412$	0.7172 0.3992 0.8257	(0.0000, 1.7071) (0.0000, 0.9480) (0.0000, 2.9595)	0.0036	0.2907	5.1944	1.355
MDPDE ($\gamma = 0.6$)	$\hat{\alpha}_1 = 0.2478$ $\hat{\alpha}_2 = 0.1325$ $\hat{\alpha}_3 = 1.4233$	0.7534 0.4049 0.8656	(0.0000, 1.7245) (0.0000, 0.9261) (0.0000, 3.1199)	0.0037	0.2931	5.2615	1.4808
MDPDE ($\gamma = 0.8$)	$\hat{\alpha}_1 = 0.0237$ $\hat{\alpha}_2 = 0.0117$ $\hat{\alpha}_3 = 1.7340$	0.9128 0.4505 1.0532	(0.0000, 1.8127) (0.0000, 0.8946) (0.0000, 3.7982)	0.0039	0.3388	5.3655	2.1453

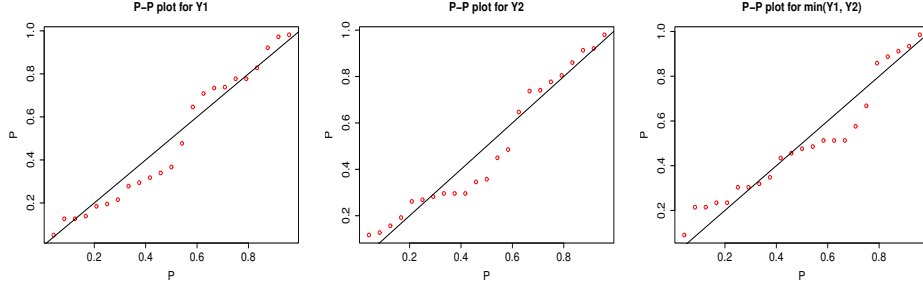
Table 7: Univariate GoF for MDPDE with optimal $\gamma^* = 0.549$

	Marginal Y_1	Marginal Y_2	$\min(Y_1, Y_2)$
KS	0.1614	0.1797	0.1762
p -value	0.5871	0.4475	0.4734

Fig. 6: KS plots of the kidney infection data for MDPDE with optimal $\gamma^8 = 0.549$ 

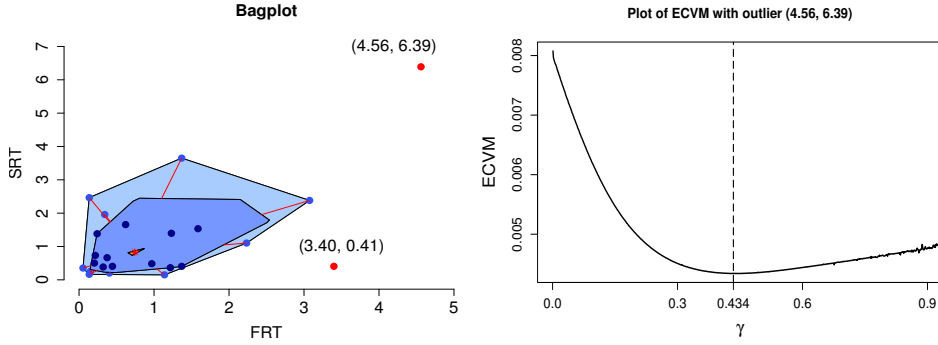
Now by adding one outlier (4.56, 6.39) in the transformed data (y_1^*, y_2^*) we see the performance of MLEs and MDPDEs. Basically we take the outlier from $BBBW(0.1, 0.1, 0.1, \beta^0)$. Then we have $n = 24$, $n_1 = 14$ and $n_2 = 10$. For this, the initial values of $(\alpha_1, \alpha_2, \alpha_3)$ are obtained as (0.4599, 0.3285, 0.5913). Using the initial

Fig. 7: PP plots of the kidney infection data for MDPDE with optimal $\gamma^* = 0.549$



values (0.4599, 0.3285, 0.5913) the optimal tuning parameter is obtained as $\gamma^* = 0.434$ and the corresponding ECVm distance is = 0.0043, bagplot and the plot of ECVm against γ are given in Figure 8. The estimate with SE, 95% CI, ECVm with p -value, GES and Tvar are reported in the Table 8. From the Table 8, we have seen that the

Fig. 8: Boxplot and ECVm plots of the kidney infection data by with outlirs (4.56, 6.39)



MDPDEs with optimal tuning parameter is good fit the data and they are robust to the presence of outliers but it is less efficient because they have larger asymptotic variances than those of the MLEs. From the bivariate GoF test using SEBs and univariate GoF test given in Table 9, 10 we have seen that the BBBE model using MDPDEs with $\gamma^* = 0.434$ fits the data well, also see the KS distance plots and P-P plots for marginals and their minimum are given in Figure 9.

8 Conclusion

Overall, the development of robust and efficient parameter estimation techniques for the BBBE is essential for advancing statistical modeling and analysis in various fields.

Table 8: Estimates with SE, 95% CI, ECVM, GES and Tvar with outlier (4.56, 6.39).

Methods	Estimates	SE	95% CI	ECVM	p -value	GES	Tvar
MLE (EM) $\gamma = 0$	$\alpha_1 = 0.2215$	0.4045	(0.0000, 1.0143)	0.0074	0.0460	2.9698e+6	0.5245
	$\alpha_2 = 0.1590$	0.2967	(0.0000, 0.7405)				
	$\alpha_3 = 0.9952$	0.5223	(0.0000, 2.0190)				
MLE optim $\gamma = 0$	$\alpha_1 = 0.2215$	0.4038	(0.0000, 1.0130)	0.0074	0.0485	2.9698e+6	0.5227
	$\alpha_2 = 0.1590$	0.2962	(0.0000, 0.7395)				
	$\alpha_3 = 0.9952$	0.5215	(0.0000, 2.0173)				
MDPDE $\gamma = 0.01$	$\alpha_1 = 0.2293$	0.4078	(0.0000, 1.0286)	0.0071	0.0519	82.8302	0.5301
	$\alpha_2 = 0.1644$	0.2976	(0.0000, 0.7478)				
	$\alpha_3 = 0.9946$	0.5246	(0.0000, 2.0227)				
MDPDE $\gamma = 0.1$	$\alpha_1 = 0.3319$	0.4404	(0.0000, 1.1951)	0.0054	0.0822	11.8759	0.6084
	$\alpha_2 = 0.2372$	0.3253	(0.0000, 0.8748)				
	$\alpha_3 = 0.9467$	0.5556	(0.0000, 2.0356)				
MDPDE $\gamma^* = 0.434$	$\alpha_1 = 0.3858$	0.6120	(0.0000, 1.5853)	0.0036	0.2229	5.6925	1.0322
	$\alpha_2 = 0.2309$	0.3785	(0.0000, 0.9728)				
	$\alpha_3 = 1.1368$	0.7172	(0.0000, 2.5426)				
MDPDE $\gamma = 0.6$	$\alpha_1 = 0.2529$	0.7155	(0.0000, 1.6552)	0.0038	0.2440	5.1852	1.3466
	$\alpha_2 = 0.1375$	0.3915	(0.0000, 0.9047)				
	$\alpha_3 = 1.3647$	0.8255	(0.0000, 2.9827)				
MDPDE $\gamma = 0.8$	$\alpha_1 = 0.0445$	0.8639	(0.0000, 1.7378)	0.0041	0.2926	5.3699	1.9342
	$\alpha_2 = 0.0224$	0.4339	(0.0000, 0.8729)				
	$\alpha_3 = 1.6547$	0.9998	(0.0000, 3.6143)				

Table 9: Bivariate GoF for MDPDE with optimal $\gamma^* = 0.434$

	KS	Critical	p -value
MDPDE	0.1364	0.2925	0.6081

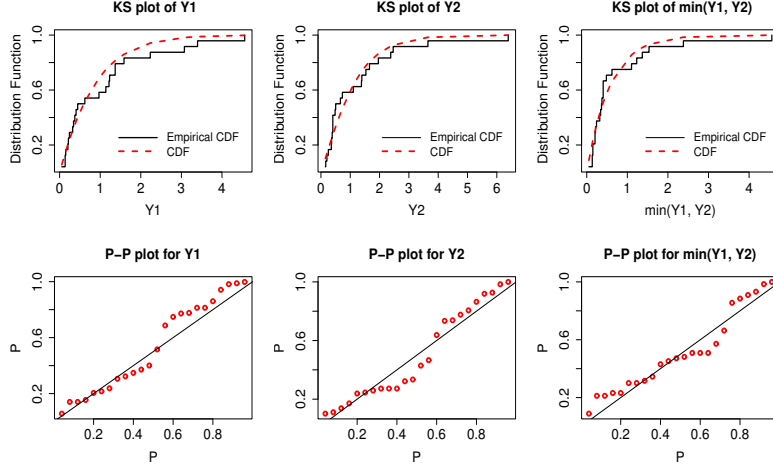
Table 10: Univariate GoF for MLE and MDPDE with optimal $\gamma^* = 0.434$

	Marginal Y_1	Marginal Y_2	$\min(Y_1, Y_2)$
KS	0.1650	0.1660	0.1706
p -value	0.5309	0.5224	0.4903

Our proposed method addresses important limitations of existing estimation techniques and offers a promising solution for improving the accuracy and reliability of statistical inference in complex data in the presence of outliers.

This paper presents the derivation of analytical expressions for the estimating equations and asymptotic distributions of the estimators. Additionally, we conduct an influence function analysis to compare MDPDE with MLE in terms of robustness, and we demonstrate the boundedness of the influence functions for MDPDE. Moreover, we analytically investigate the asymptotic relative efficiency of MDPDE for various the BBDE parameters and MDPDE tuning parameters, revealing that for a moderate tuning parameter value, the asymptotic relative efficiency approaches one. Furthermore, we apply the proposed method to both simulated datasets and real data set. The

Fig. 9: KS and P-P plots of the kidney infection data with outlier case for MDPDE with optimal $\gamma^* = 0.434$



results suggest that MDPDE outperforms than MLE in most scenarios, particularly when dealing with outliers in the data.

Acknowledgments

The authors would like to express their sincere gratitude to the editor for their careful consideration and guidance throughout the review process. We also grateful to the anonymous reviewers for their insightful comments and constructive suggestions, which have significantly improved the quality of this manuscript.

Appendix A Integrability conditions

The components of the maximum likelihood score function $S_\theta(y_1, y_2) = \nabla[\log f(y_1, y_2; \theta)]$ of the BBDE distribution, are given by the usual partial derivatives of $\log f(y_1, y_2; \theta)$, namely

$$S_{\alpha_1}(y_1, y_2) = \begin{cases} \frac{1}{\alpha_1 + \alpha_2 + \alpha_3} - \frac{1}{\alpha_1 + \alpha_2} + \frac{1}{\alpha_1} - y_1 & \text{if } y_1 < y_2 \\ \frac{1}{\alpha_1 + \alpha_2 + \alpha_3} - \frac{1}{\alpha_1 + \alpha_2} + \frac{1}{\alpha_1 + \alpha_3} - y_1 & \text{if } y_1 > y_2 \end{cases}$$

$$S_{\alpha_2}(y_1, y_2) = \begin{cases} \frac{1}{\alpha_1 + \alpha_2 + \alpha_3} - \frac{1}{\alpha_1 + \alpha_2} + \frac{1}{\alpha_2 + \alpha_3} - y_2 & \text{if } y_1 < y_2 \\ \frac{1}{\alpha_1 + \alpha_2 + \alpha_3} - \frac{1}{\alpha_1 + \alpha_2} + \frac{1}{\alpha_2} - y_2 & \text{if } y_1 > y_2 \end{cases}$$

$$S_{\alpha_3}(y_1, y_2) = \begin{cases} \frac{1}{\alpha_1 + \alpha_2 + \alpha_3} + \frac{1}{\alpha_2 + \alpha_3} - y_2 & \text{if } y_1 < y_2 \\ \frac{1}{\alpha_1 + \alpha_2 + \alpha_3} + \frac{1}{\alpha_1 + \alpha_3} - y_1 & \text{if } y_1 > y_2. \end{cases}$$

The $(i, j)^{th}$ element of the Hessian matrix, $H(y_1, y_2; \theta) = \nabla[S_\theta(y_1, y_2)]$, and the 3×3 symmetric matrix $I(y_1, y_2; \theta) = -H(y_1, y_2; \theta)$, so called the information matrix of BBBE model is given as, for $i, j = 1, 2, 3$,

$$H_{ij}(y_1, y_2; \theta) = \frac{\partial^2 \log f(y_1, y_2; \theta)}{\partial \alpha_i \partial \alpha_j} \quad \text{and} \quad I_{ij}(y_1, y_2; \theta) = -H_{ij}(y_1, y_2; \theta).$$

Then we have,

$$\begin{aligned} H_{11}(y_1, y_2; \theta) &= - \left[\frac{1}{(\alpha_1 + \alpha_2 + \alpha_3)^2} - \frac{1}{(\alpha_1 + \alpha_2)^2} + \frac{1}{\alpha_1^2} \mathbb{1}_{(y_1 < y_2)} + \frac{1}{(\alpha_1 + \alpha_3)^2} \mathbb{1}_{(y_1 > y_2)} \right] \\ H_{12}(y_1, y_2; \theta) &= - \left[\frac{1}{(\alpha_1 + \alpha_2 + \alpha_3)^2} - \frac{1}{(\alpha_1 + \alpha_2)^2} \right] \\ H_{13}(y_1, y_2; \theta) &= - \left[\frac{1}{(\alpha_1 + \alpha_2 + \alpha_3)^2} + \frac{1}{(\alpha_1 + \alpha_3)^2} \mathbb{1}_{(y_1 > y_2)} \right] \\ H_{22}(y_1, y_2; \theta) &= - \left[\frac{1}{(\alpha_1 + \alpha_2 + \alpha_3)^2} - \frac{1}{(\alpha_1 + \alpha_2)^2} + \frac{1}{(\alpha_2 + \alpha_3)^2} \mathbb{1}_{(y_1 < y_2)} + \frac{1}{\alpha_2^2} \mathbb{1}_{(y_1 > y_2)} \right] \\ H_{23}(y_1, y_2; \theta) &= - \left[\frac{1}{(\alpha_1 + \alpha_2 + \alpha_3)^2} + \frac{1}{(\alpha_2 + \alpha_3)^2} \mathbb{1}_{(y_1 < y_2)} \right] \\ H_{33}(y_1, y_2; \theta) &= - \left[\frac{1}{(\alpha_1 + \alpha_2 + \alpha_3)^2} + \frac{1}{(\alpha_2 + \alpha_3)^2} \mathbb{1}_{(y_1 < y_2)} + \frac{1}{(\alpha_1 + \alpha_3)^2} \mathbb{1}_{(y_1 > y_2)} \right]. \end{aligned}$$

The consistency and asymptotic normality of the MDPDE for the BBBE distribution follow from the integrability conditions see [18],

$$\int_0^\infty \int_0^\infty f^{2\gamma}(y_1, y_2; \theta) g(y_1, y_2) dy_1 dy_2 < \infty \quad (\text{A1})$$

$$\int_0^\infty \int_0^\infty S_{\alpha_i}(y_1, y_2) f^{2\gamma}(y_1, y_2; \theta) g(y_1, y_2) dy_1 dy_2 < \infty \quad i = 1, 2, 3 \quad (\text{A2})$$

$$\int_0^\infty \int_0^\infty \left| \frac{\partial^2 \log f(y_1, y_2; \theta)}{\partial \alpha_i \partial \alpha_j} \right| f^{2\gamma}(y_1, y_2; \theta) g(y_1, y_2) dy_1 dy_2 < \infty \quad i \neq j = 1, 2, 3. \quad (\text{A3})$$

Appendix B Covariance matrix

The components of the vector $\xi_\gamma(\theta)$ are given as

$$\begin{aligned} \xi_\gamma(\alpha_1) &= \frac{c^{1+\gamma} \alpha_1^{1+\gamma} (\alpha_2 + \alpha_3)^\gamma}{(1+\gamma)^2 \alpha} \left[\left(\frac{1}{\alpha} - \frac{1}{\alpha_1 + \alpha_2} + \frac{1}{\alpha_1} \right) - \frac{1}{(1+\gamma)\alpha} \right] \\ &\quad + \frac{c^{1+\gamma} (\alpha_1 + \alpha_3)^\gamma \alpha_2^{1+\gamma}}{(1+\gamma)^2 \alpha} \left(\frac{1}{\alpha} - \frac{1}{\alpha_1 + \alpha_2} + \frac{1}{\alpha_1 + \alpha_3} \right) \\ &\quad - \frac{c^{1+\gamma} (\alpha_1 + \alpha_3)^{1+\gamma} \alpha_2^\gamma}{(1+\gamma)^3} \left[\frac{1}{(\alpha_1 + \alpha_3)^2} - \frac{1}{\alpha^2} \right] \end{aligned}$$

$$\begin{aligned}\xi_\gamma(\alpha_2) = & \frac{c^{1+\gamma}\alpha_1^{1+\gamma}(\alpha_2+\alpha_3)^\gamma}{(1+\gamma)^2\alpha} \left(\frac{1}{\alpha} - \frac{1}{\alpha_1+\alpha_2} + \frac{1}{\alpha_2+\alpha_3} \right) \\ & - \frac{c^{1+\gamma}\alpha_1^\gamma(\alpha_2+\alpha_3)^{1+\gamma}}{(1+\gamma)^3} \left[\frac{1}{(\alpha_2+\alpha_3)^2} - \frac{1}{\alpha^2} \right] \\ & + \frac{c^{1+\gamma}(\alpha_1+\alpha_3)^\gamma\alpha_2^{1+\gamma}}{(1+\gamma)^2\alpha} \left[\left(\frac{1}{\alpha} - \frac{1}{\alpha_1+\alpha_2} + \frac{1}{\alpha_2} \right) - \frac{1}{(1+\gamma)\alpha} \right]\end{aligned}$$

$$\begin{aligned}\xi_\gamma(\alpha_3) = & \frac{c^{1+\gamma}\alpha_1^{1+\gamma}(\alpha_2+\alpha_3)^\gamma}{(1+\gamma)^2\alpha} \left(\frac{1}{\alpha} + \frac{1}{\alpha_2+\alpha_3} \right) \\ & - \frac{c^{1+\gamma}\alpha_1^\gamma(\alpha_2+\alpha_3)^{1+\gamma}}{(1+\gamma)^3} \left[\frac{1}{(\alpha_2+\alpha_3)^2} - \frac{1}{\alpha^2} \right] \\ & + \frac{c^{1+\gamma}(\alpha_1+\alpha_3)^\gamma\alpha_2^{1+\gamma}}{(1+\gamma)^2\alpha} \left(\frac{1}{\alpha} + \frac{1}{\alpha_1+\alpha_3} \right) \\ & - \frac{c^{1+\gamma}(\alpha_1+\alpha_3)^{1+\gamma}\alpha_2^\gamma}{(1+\gamma)^3} \left[\frac{1}{(\alpha_1+\alpha_3)^2} - \frac{1}{\alpha^2} \right]\end{aligned}$$

The elements of the matrix $J = c^{1+\gamma}J'$ are given below

$$\begin{aligned}J'_{11} = & \frac{\alpha_1^{1+\gamma}(\alpha_2+\alpha_3)^\gamma}{(1+\gamma)^2\alpha} \left[\left(\frac{1}{\alpha} - \frac{1}{\alpha_1+\alpha_2} + \frac{1}{\alpha_1} \right)^2 \right] \\ & - \frac{2\alpha_1^{1+\gamma}(\alpha_2+\alpha_3)^\gamma}{(1+\gamma)^3\alpha^2} \left[\left(\frac{1}{\alpha} - \frac{1}{\alpha_1+\alpha_2} + \frac{1}{\alpha_1} \right) + \frac{2}{(1+\gamma)^2\alpha^2} \right] \\ & - \frac{2(\alpha_1+\alpha_3)^{1+\gamma}\alpha_2^\gamma}{(1+\gamma)^3} \left(\frac{1}{\alpha} - \frac{1}{\alpha_1+\alpha_2} + \frac{1}{\alpha_1+\alpha_3} \right) \left[\frac{1}{(\alpha_1+\alpha_3)^2} - \frac{1}{\alpha^2} \right] \\ & + \frac{(\alpha_1+\alpha_3)^\gamma\alpha_2^{1+\gamma}}{(1+\gamma)^2\alpha} \left(\frac{1}{\alpha} - \frac{1}{\alpha_1+\alpha_2} + \frac{1}{\alpha_1+\alpha_3} \right)^2 \\ & + \frac{2(\alpha_1+\alpha_3)^{1+\gamma}\alpha_2^\gamma}{(1+\gamma)^4} \left[\frac{1}{(\alpha_1+\alpha_3)^3} - \frac{1}{\alpha^3} \right]\end{aligned}$$

$$\begin{aligned}J'_{12} = & \frac{\alpha_1^{1+\gamma}(\alpha_2+\alpha_3)^\gamma}{(1+\gamma)^2\alpha} \left[\left(\frac{1}{\alpha} - \frac{1}{\alpha_1+\alpha_2} \right)^2 + \left(\frac{1}{\alpha} - \frac{1}{\alpha_1+\alpha_2} \right) \left(\frac{1}{\alpha_1} + \frac{1}{\alpha_2+\alpha_3} \right) \right] \\ & + \frac{\alpha_1^{1+\gamma}(\alpha_2+\alpha_3)^\gamma}{(1+\gamma)^2\alpha} \frac{1}{\alpha_1(\alpha_2+\alpha_3)} - \frac{\alpha_1^{1+\gamma}(\alpha_2+\alpha_3)^\gamma}{(1+\gamma)^3\alpha^2} \left(\frac{1}{\alpha} - \frac{1}{\alpha_1+\alpha_2} + \frac{1}{\alpha_2+\alpha_3} \right) \\ & - \frac{\alpha_1^\gamma(\alpha_2+\alpha_3)^{1+\gamma}}{(1+\gamma)^3} \left(\frac{1}{\alpha} - \frac{1}{\alpha_1+\alpha_2} + \frac{1}{\alpha_1} \right) \left[\frac{1}{(\alpha_2+\alpha_3)^2} - \frac{1}{\alpha^2} \right] \\ & + \frac{\alpha_1^{1+\gamma}(\alpha_2+\alpha_3)^\gamma}{(1+\gamma)^4} \left[\frac{2}{\alpha^3} + \frac{1}{(\alpha_2+\alpha_3)\alpha^2} \right] + \frac{(\alpha_1+\alpha_3)^\gamma\alpha_2^{1+\gamma}}{(1+\gamma)^2\alpha} \left(\frac{1}{\alpha} - \frac{1}{\alpha_1+\alpha_2} \right)^2 \\ & + \frac{(\alpha_1+\alpha_3)^\gamma\alpha_2^{1+\gamma}}{(1+\gamma)^2\alpha} \left[\left(\frac{1}{\alpha} - \frac{1}{\alpha_1+\alpha_2} \right) \left(\frac{1}{\alpha_1+\alpha_3} + \frac{1}{\alpha_2} \right) + \frac{1}{(\alpha_1+\alpha_3)\alpha_2} \right] \\ & - \frac{(\alpha_1+\alpha_3)^{1+\gamma}\alpha_2^\gamma}{(1+\gamma)^3} \left(\frac{1}{\alpha} - \frac{1}{\alpha_1+\alpha_2} + \frac{1}{\alpha_2} \right) \left[\frac{1}{(\alpha_1+\alpha_3)^2} - \frac{1}{\alpha^2} \right]\end{aligned}$$

$$\begin{aligned}
& - \frac{(\alpha_1 + \alpha_3)^\gamma \alpha_2^{1+\gamma}}{(1+\gamma)^3 \alpha^2} \left(\frac{1}{\alpha} - \frac{1}{\alpha_1 + \alpha_2} + \frac{1}{\alpha_1 + \alpha_3} \right) \\
& + \frac{(\alpha_1 + \alpha_3)^\gamma \alpha_2^{1+\gamma}}{(1+\gamma)^4} \left[\frac{2}{\alpha^3} + \frac{1}{(\alpha_1 + \alpha_3) \alpha^2} \right]
\end{aligned}$$

$$\begin{aligned}
J'_{13} = & \frac{\alpha_1^{1+\gamma} (\alpha_2 + \alpha_3)^\gamma}{(1+\gamma)^2 \alpha} \left[\left(\frac{1}{\alpha} - \frac{1}{\alpha_1 + \alpha_2} + \frac{1}{\alpha_1} \right) \left(\frac{1}{\alpha} + \frac{1}{\alpha_2 + \alpha_3} \right) \right] \\
& - \frac{\alpha_1^{1+\gamma} (\alpha_2 + \alpha_3)^\gamma}{(1+\gamma)^3 \alpha^2} \left(\frac{1}{\alpha} + \frac{1}{\alpha_2 + \alpha_3} \right) \\
& - \frac{\alpha_1^\gamma (\alpha_2 + \alpha_3)^{1+\gamma}}{(1+\gamma)^3} \left(\frac{1}{\alpha} - \frac{1}{\alpha_1 + \alpha_2} + \frac{1}{\alpha_1} \right) \left[\frac{1}{(\alpha_2 + \alpha_3)^2} - \frac{1}{\alpha^2} \right] \\
& + \frac{\alpha_1^{1+\gamma} (\alpha_2 + \alpha_3)^\gamma}{(1+\gamma)^4} \left[\frac{2}{\alpha^3} + \frac{1}{(\alpha_2 + \alpha_3) \alpha^2} \right] \\
& + \frac{(\alpha_1 + \alpha_3)^\gamma \alpha_2^{1+\gamma}}{(1+\gamma)^2 \alpha} \left(\frac{1}{\alpha} - \frac{1}{\alpha_1 + \alpha_2} + \frac{1}{\alpha_1 + \alpha_3} \right) \left(\frac{1}{\alpha} + \frac{1}{\alpha_1 + \alpha_3} \right) \\
& + \frac{2(\alpha_1 + \alpha_3)^{1+\gamma} \alpha_2^\gamma}{(1+\gamma)^4} \left[\frac{1}{(\alpha_1 + \alpha_3)^3} - \frac{1}{\alpha^3} \right] \\
& - \frac{(\alpha_1 + \alpha_3)^{1+\gamma} \alpha_2^\gamma}{(1+\gamma)^3} \left[\left(\frac{1}{\alpha} - \frac{1}{\alpha_1 + \alpha_2} + \frac{1}{\alpha_1 + \alpha_3} \right) \right] \left[\frac{1}{(\alpha_1 + \alpha_3)^2} - \frac{1}{\alpha^2} \right] \\
& - \frac{(\alpha_1 + \alpha_3)^{1+\gamma} \alpha_2^\gamma}{(1+\gamma)^3} \left[\left(\frac{1}{\alpha} + \frac{1}{\alpha_1 + \alpha_3} \right) \right] \left[\frac{1}{(\alpha_1 + \alpha_3)^2} - \frac{1}{\alpha^2} \right]
\end{aligned}$$

$$\begin{aligned}
J'_{22} = & \frac{\alpha_1^{1+\gamma} (\alpha_2 + \alpha_3)^\gamma}{(1+\gamma)^2 \alpha} \left(\frac{1}{\alpha} - \frac{1}{\alpha_1 + \alpha_2} + \frac{1}{\alpha_2 + \alpha_3} \right)^2 \\
& - \frac{2\alpha_1^\gamma (\alpha_2 + \alpha_3)^{1+\gamma}}{(1+\gamma)^3} \left(\frac{1}{\alpha} - \frac{1}{\alpha_1 + \alpha_2} + \frac{1}{\alpha_2 + \alpha_3} \right) \left[\frac{1}{(\alpha_2 + \alpha_3)^2} - \frac{1}{\alpha^2} \right] \\
& + \frac{2\alpha_1^\gamma (\alpha_2 + \alpha_3)^{1+\gamma}}{(1+\gamma)^4} \left[\frac{1}{(\alpha_2 + \alpha_3)^3} - \frac{1}{\alpha^3} \right] \\
& + \frac{(\alpha_1 + \alpha_3)^\gamma \alpha_2^{1+\gamma}}{(1+\gamma)^2 \alpha} \left[\left(\frac{1}{\alpha} - \frac{1}{\alpha_1 + \alpha_2} + \frac{1}{\alpha_2} \right)^2 \right] \\
& - \frac{(\alpha_1 + \alpha_3)^\gamma \alpha_2^{1+\gamma}}{(1+\gamma)^2 \alpha} \left[\frac{2}{(1+\gamma) \alpha} \left(\frac{1}{\alpha} - \frac{1}{\alpha_1 + \alpha_2} + \frac{1}{\alpha_2} \right) - \frac{2}{(1+\gamma)^2 \alpha^2} \right]
\end{aligned}$$

$$\begin{aligned}
J'_{23} = & \frac{\alpha_1^{1+\gamma} (\alpha_2 + \alpha_3)^\gamma}{(1+\gamma)^2 \alpha} \left(\frac{1}{\alpha} - \frac{1}{\alpha_1 + \alpha_2} + \frac{1}{\alpha_2 + \alpha_3} \right) \left(\frac{1}{\alpha} + \frac{1}{\alpha_2 + \alpha_3} \right) \\
& - \frac{\alpha_1^\gamma (\alpha_2 + \alpha_3)^{1+\gamma}}{(1+\gamma)^3} \left[\left(\frac{1}{\alpha} - \frac{1}{\alpha_1 + \alpha_2} + \frac{1}{\alpha_2 + \alpha_3} \right) \right] \left[\frac{1}{(\alpha_2 + \alpha_3)^2} - \frac{1}{\alpha^2} \right] \\
& - \frac{\alpha_1^\gamma (\alpha_2 + \alpha_3)^{1+\gamma}}{(1+\gamma)^3} \left[\left(\frac{1}{\alpha} + \frac{1}{\alpha_2 + \alpha_3} \right) \right] \left[\frac{1}{(\alpha_2 + \alpha_3)^2} - \frac{1}{\alpha^2} \right] \\
& + \frac{2\alpha_1^\gamma (\alpha_2 + \alpha_3)^{1+\gamma}}{(1+\gamma)^4} \left[\frac{1}{(\alpha_2 + \alpha_3)^3} - \frac{1}{\alpha^3} \right]
\end{aligned}$$

$$\begin{aligned}
& + \frac{(\alpha_1 + \alpha_3)^\gamma \alpha_2^{1+\gamma}}{(1+\gamma)^2 \alpha} \left(\frac{1}{\alpha} - \frac{1}{\alpha_1 + \alpha_2} + \frac{1}{\alpha_2} \right) \left(\frac{1}{\alpha} + \frac{1}{\alpha_1 + \alpha_3} \right) \\
& - \frac{(\alpha_1 + \alpha_3)^{1+\gamma} \alpha_2^\gamma}{(1+\gamma)^3} \left(\frac{1}{\alpha} - \frac{1}{\alpha_1 + \alpha_2} + \frac{1}{\alpha_2} \right) \left[\frac{1}{(\alpha_1 + \alpha_3)^2} - \frac{1}{\alpha^2} \right] \\
& - \frac{(\alpha_1 + \alpha_3)^\gamma \alpha_2^{1+\gamma}}{(1+\gamma)^3 \alpha^2} \left(\frac{1}{\alpha} + \frac{1}{\alpha_1 + \alpha_3} \right) + \frac{(\alpha_1 + \alpha_3)^\gamma \alpha_2^{1+\gamma}}{(1+\gamma)^4} \left[\frac{2}{\alpha^3} + \frac{1}{(\alpha_1 + \alpha_3) \alpha^2} \right] \\
j'_{33} = & \frac{\alpha_1^{1+\gamma} (\alpha_2 + \alpha_3)^\gamma}{(1+\gamma)^2 \alpha} \left(\frac{1}{\alpha} + \frac{1}{\alpha_2 + \alpha_3} \right)^2 \\
& - \frac{2\alpha_1^\gamma (\alpha_2 + \alpha_3)^{1+\gamma}}{(1+\gamma)^3} \left(\frac{1}{\alpha} + \frac{1}{\alpha_2 + \alpha_3} \right) \left[\frac{1}{(\alpha_2 + \alpha_3)^2} - \frac{1}{\alpha^2} \right] \\
& + \frac{(\alpha_1 + \alpha_3)^\gamma \alpha_2^{1+\gamma}}{(1+\gamma)^2 \alpha} \left(\frac{1}{\alpha} + \frac{1}{\alpha_1 + \alpha_3} \right)^2 \\
& - \frac{2(\alpha_1 + \alpha_3)^{1+\gamma} \alpha_2^\gamma}{(1+\gamma)^3} \left(\frac{1}{\alpha} + \frac{1}{\alpha_1 + \alpha_3} \right) \left[\frac{1}{(\alpha_1 + \alpha_3)^2} - \frac{1}{\alpha^2} \right] \\
& + \frac{2\alpha_1^\gamma (\alpha_2 + \alpha_3)^{1+\gamma}}{(1+\gamma)^4} \left[\frac{1}{(\alpha_2 + \alpha_3)^3} - \frac{1}{\alpha^3} \right] + \frac{2(\alpha_1 + \alpha_3)^{1+\gamma} \alpha_2^\gamma}{(1+\gamma)^4} \left[\frac{1}{(\alpha_1 + \alpha_3)^3} - \frac{1}{\alpha^3} \right]
\end{aligned}$$

References

- [1] Block, H.W., Basu, A.: A continuous, bivariate exponential extension. *Journal of the American Statistical Association* **69**(348), 1031–1037 (1974)
- [2] Marshall, A.W., Olkin, I.: A multivariate exponential distribution. *Journal of the American Statistical Association* **62**(317), 30–44 (1967)
- [3] Basu, A., Harris, I.R., Hjort, N.L., Jones, M.: Robust and efficient estimation by minimising a density power divergence. *Biometrika* **85**(3), 549–559 (1998)
- [4] Juárez, S.F., Schucany, W.R.: Robust and efficient estimation for the generalized pareto distribution. *Extremes* **7**, 237–251 (2004)
- [5] Hazra, A.: Minimum density power divergence estimation for the generalized exponential distribution. *Communications in Statistics-Theory and Methods* **54**(4), 1050–1070 (2025)
- [6] Cressie, N., Read, T.R.: Multinomial goodness-of-fit tests. *Journal of the Royal Statistical Society Series B: Statistical Methodology* **46**(3), 440–464 (1984)
- [7] Kullback, S., Leibler, R.A.: On information and sufficiency. *The Annals of Mathematical Statistics* **22**(1), 79–86 (1951)
- [8] Basu, A., Lindsay, B.G.: Minimum disparity estimation for continuous models: efficiency, distributions and robustness. *Annals of the Institute of Statistical Mathematics* **46**, 683–705 (1994)

- [9] Beran, R.: Minimum hellinger distance estimates for parametric models. *The annals of Statistics*, 445–463 (1977)
- [10] Beran, R.: An efficient and robust adaptive estimator of location. *The Annals of Statistics*, 292–313 (1978)
- [11] Fujisawa, H., Eguchi, S.: Robust estimation in the normal mixture model. *Journal of Statistical Planning and Inference* **136**(11), 3989–4011 (2006)
- [12] Strupczewski, W., Kochanek, K., Weglarczyk, S., Singh, V.: On robustness of large quantile estimates of log-gumbel and log-logistic distributions to largest element of the observation series: Monte carlo results vs. first order approximation. *Stochastic Environmental Research and Risk Assessment* **19**, 280–291 (2005)
- [13] Strupczewski, W., Kochanek, K., Weglarczyk, S., Singh, V.: On robustness of large quantile estimates to largest elements of the observation series. *Hydrological Processes: An International Journal* **21**(10), 1328–1344 (2007)
- [14] Huber, P.J.: *Robust Statistical Procedures*. SIAM, ??? (1996)
- [15] Hampel, F.R., Ronchetti, E.M., Rousseeuw, P.J., Stahel, W.A.: *Robust Statistics: the Approach Based on Influence Functions*. John Wiley & Sons, ??? (2011)
- [16] Hampel, F.R.: The influence curve and its role in robust estimation. *Journal of the American Statistical Association* **69**(346), 383–393 (1974)
- [17] Kallianpur, G., Rao, C.R.: On fisher’s lower bound to asymptotic variance of a consistent estimate. *Sankhyā: The Indian Journal of Statistics (1933-1960)* **15**(4), 331–342 (1955)
- [18] Basu, A., Shioya, H., Park, C.: *Statistical Inference: the Minimum Distance Approach*. CRC press, ??? (2011)
- [19] Kundu, D., Gupta, R.D.: A class of absolutely continuous bivariate distributions. *Statistical Methodology* **7**(4), 464–477 (2010)
- [20] McGilchrist, C., Aisbett, C.: Regression with frailty in survival analysis. *Biometrics*, 461–466 (1991)
- [21] Rousseeuw, P.J., Ruts, I., Tukey, J.W.: The bagplot: a bivariate boxplot. *The American Statistician* **53**(4), 382–387 (1999)
- [22] Kumar, S., Kundu, D., Mitra, S.: Estimation of five parameter bivariate modified weibull singular distribution. *Communications in Statistics - Simulation and Computation*, 1–24 (2024) <https://doi.org/10.1080/03610918.2024.2434709>
- [23] Efron, B.: *The Jackknife, the Bootstrap and Other Resampling Plans*. SIAM, ??? (1982)

- [24] Alam, K., Abernathy, R., Williams, C.L.: Multivariate goodness-of-fit tests based on statistically equivalent blocks. *Communications in Statistics-Theory and Methods* **22**(6), 1515–1533 (1993)
- [25] Alam, K., Williams, C.L.: A multivariate goodness-of-fit test for stochastically ordered distributions. *Biometrical journal* **37**(8), 945–956 (1995)
- [26] Chernick, M.R.: *Bootstrap Methods: A Guide for Practitioners and Researchers*. John Wiley & Sons, ??? (2011)