

MULTIVARIATE GENERALIZED EXPONENTIAL DISTRIBUTION

JIANYONG MU AND YANLING WANG

COLLEGE OF MATHEMATICS AND INFORMATION SCIENCE, HENAN NORMAL
UNIVERSITY, HENAN, 453007, CHINA

E-MAILS: JIANYONGMU@163.COM, BIGDUCKWYL@163.COM

ABSTRACT. The paper mainly aims to extend the bivariate generalized exponential distribution into multivariate exponential distribution. It also provides the explicit forms of the joint cumulative distribution function and joint probability distribution function, and further discusses that the EM algorithm can be used to compute the maximum likelihood estimators of the unknown parameters.

AMS Classification: 62E15, 62F10, 62H10, 62H12

Keywords: Generalized exponential distribution, Joint probability density function, Maximum likelihood estimators, EM algorithm

1. INTRODUCTION

Generalized exponential(GE) distribution was introduced by Gupta and Kundu [3]. The GE distribution has lots of interesting properties and it can be used effectively to analyze several skewed life time data. The review article of Gupta and Kundu [2] have introduced a current account on GE distribution.

Recently, some works have been carried out on the extending of the GE to the multivariate set up. Sarhan and Balakrishman [5] have defined a new bivariate distribution using the GE distribution and exponential distribution and derived several interesting properties of this new distribution. Kundu and Gupta [4] have

provide a bivariate generalize exponential (BVGE) distribution so that the marginal distribution are GE distribution.

The paper mainly aims to extend the BVGE distribution into multivariate generalized exponential distribution (MVGE). The joint cumulative distribution function (CDF) and the joint probability density function (PDF) of the (MVGE) are provided in Section 2. The EM algorithm to compute the MLEs of the unknown parameters is provided in Section 3. An illustrative example for the EM algorithm to compute the MLEs of the unknown parameters is given in section 4.

2. MULTIVARIATE GENERALIZED EXPONENTIAL DISTRIBUTION

The univariate GE distribution has the following CDF and PDF respectively for $x > 0$.

$$F_{\text{GE}}(x; \alpha, \lambda) = (1 - e^{-\lambda x})^\alpha, \quad f_{\text{GE}}(x; \alpha, \lambda) = \alpha \lambda e^{-\lambda x} (1 - e^{-\lambda x})^{\alpha-1}.$$

Here $\alpha > 0$ and $\lambda > 0$ are the shape and scale parameters respectively. From now on a GE distribution with the shape parameter α and the scale parameter λ will be denoted by $\text{GE}(\alpha, \lambda)$.

Suppose $U_i \sim \text{GE}(\alpha_i, \lambda)$, $U \sim \text{GE}(\alpha, \lambda)$, $i = 1, 2, \dots, n$, and they are mutually independent. Here ' $\sim$ ' means follows or has the distribution. Now define $X_i = U_i \vee U$, $i = 1, 2, \dots, n$. Then we say that the multivariate vector $(X_1, X_2, \dots, X_n)$ has a multivariate generalized exponential distribution with the parameters $\alpha_1, \alpha_2, \dots, \alpha_n, \alpha$ and the scale parameter λ . It is denoted by $\text{MVGE}(n, \alpha_1, \alpha_2, \dots, \alpha_n, \alpha, \lambda)$.

Theorem 1. *If $(X_1, X_2, \dots, X_n) \sim \text{MVGE}(n, \alpha_1, \alpha_2, \dots, \alpha_n, \alpha, \lambda)$, then the joint CDF of $(X_1, X_2, \dots, X_n)$ is*

$$F_{\text{MVGE}}(x_1, x_2, \dots, x_n) = (1 - e^{-\lambda y})^\alpha \prod_{i=1}^n (1 - e^{-\lambda x_i})^{\alpha_i},$$

where $y = \min_i \{x_i\}$.

Proof. Trivial and therefore it is omitted. □

Now let $S = \{i : x_i = \min_i \{x_i\}\}$, $y = \min_i \{x_i\}$ and $|\cdot|$ be the number of the elements in the set ' $\cdot$ ', then the joint PDF of $\text{MVGE}(n, \alpha_1, \alpha_2, \dots, \alpha_n, \alpha, \lambda)$ can be obtained.

Theorem 2. If $(X_1, X_2, \dots, X_n) \sim \text{MVGE}(n, \alpha_1, \alpha_2, \dots, \alpha_n, \alpha, \lambda)$, then the joint PDF of $(X_1, X_2, \dots, X_n)$ is

$$f_{\text{MVGE}}(x_1, x_2, \dots, x_n) = \begin{cases} \lambda^n e^{-\lambda \sum_{i=1}^n x_i} (\alpha + \alpha_k) (1 - e^{-\lambda x_k})^{\alpha + \alpha_k - 1} \prod_{i \notin S} \alpha_i (1 - e^{-\lambda x_i})^{\alpha_i - 1}, & \text{if } |S| = 1, \\ \lambda^{n+1-|S|} \alpha e^{-\lambda(y + \sum_{i \notin S} x_i)} (1 - e^{-\lambda y})^{\alpha - 1 + \sum_{i \in S} \alpha_i} \prod_{i \notin S} \alpha_i (1 - e^{-\lambda x_i})^{\alpha_i - 1}, & \text{if } |S| \geq 2, \end{cases}$$

where $x_k = \min_i \{x_i\}$.

Proof. When $|S| = 1$, PDF can be obtained simply by taking the mixed n th partial derivative of $F(x_1, x_2, \dots, x_n)$. When $|S| \geq 2$,

$$P\left(\left[\bigcap_{i \in S} \{X_i = y\}\right] \cap \left[\bigcap_{i \notin S} \{X_i \leq x_i\}\right] \middle| U = y\right) =$$

$$P\left(\left[\bigcap_{i \in S} \{U_i \leq y\}\right] \cap \left[\bigcap_{i \notin S} \{U_i \leq x_i\}\right]\right) =$$

$$\prod_{i \in S} F_{U_i}(y) \cdot \prod_{i \notin S} F_{U_i}(x_i).$$

$$f_{\text{MVGE}}(x_1, x_2, \dots, x_n | U = y) = \prod_{i \notin S} f_{U_i}(x_i) \prod_{i \in S} F_{U_i}(y).$$

Therefore,

$$f_{\text{MVGE}}(x_1, x_2, \dots, x_n) = f_{\text{MVGE}}(x_1, x_2, \dots, x_n | U = y) f_U(y) =$$

$$f_U(y) \prod_{i \notin S} f_{U_i}(x_i) \prod_{i \in S} F_{U_i}(y).$$

The result follows. $\square$

The following theorem provides the marginal and conditional results of the MVGE distribution.

Theorem 3. If $(X_1, X_2, \dots, X_n) \sim \text{MVGE}(n, \alpha_1, \alpha_2, \dots, \alpha_n, \alpha, \lambda)$, then

- (1) $X_i \sim \text{GE}(\alpha + \alpha_i, \lambda)$, $i = 1, 2, \dots, n$.
 (2) Suppose A is a nonempty subset of $\{1, 2, \dots, n\}$, then

$$P\left(\bigcap_{i \in A} \{X_i \leq x_i\}\right) = (1 - e^{-\lambda y})^\alpha \prod_{i \in A} (1 - e^{-\lambda x_i})^{\alpha_i},$$

where $y = \min_{i \in A} \{x_i\}$.

- (3) Suppose A, B are two disjoint nonempty subsets of $\{1, 2, \dots, n\}$, then the conditional distribution

$$P\left(\bigcap_{i \in A} \{X_i \leq x_i\} \middle| \bigcap_{j \in B} \{X_j \leq x_j\}\right) = \frac{(1 - e^{-\lambda y})^\alpha}{(1 - e^{-\lambda z})^\alpha} \prod_{i \in A} (1 - e^{-\lambda x_i})^{\alpha_i},$$

where $y = \min_{i \in A \cup B} \{x_i\}$ and $z = \min_{i \in B} \{x_i\}$.

Proof. Trivial and therefore it is omitted. $\square$

3. MAXIMUM LIKELIHOOD ESTIMATION

When $|S| = 1$, the PDF of $(X_1, X_2, \dots, X_n)$ can be denoted by the convenient form:

$$f_{\text{MVGE}}(x_1, x_2, \dots, x_n) = \lambda^n e^{-\lambda \sum_{i=1}^n x_i} \prod_{i \in S} (\alpha + \alpha_i) (1 - e^{-\lambda y})^{\alpha + \alpha_i - 1} \prod_{i \notin S} \alpha_i (1 - e^{-\lambda x_i})^{\alpha_i - 1}.$$

Suppose $\{(X_{1j}, X_{2j}, \dots, X_{nj})\}_{j=1}^N$ is a random sample from $\text{MVGE}(n, \alpha_1, \alpha_2, \dots, \alpha_n, \alpha, \lambda)$. Let

$$S_j = \{i : X_{ij} = \min_i \{X_{ij}\} = Y_j\}, \quad s_j = |S_j|, \quad j = 1, 2, \dots, N,$$

$$I_1 = \{j : |S_j| = 1\}, \quad I_2 = \{j : |S_j| \geq 2\}, \quad n_1 = |I_1|, \quad n_2 = |I_2|.$$

The likelihood function can be written as

$$L(\alpha_1, \alpha_2, \dots, \alpha_n, \alpha, \lambda) = \prod_{j \in I_1} \left[\lambda^n e^{-\lambda \sum_{i=1}^n x_{ij}} \prod_{i \in S_j} (\alpha + \alpha_i) (1 - e^{-\lambda y_j})^{\alpha + \alpha_i - 1} \prod_{i \notin S_j} \alpha_i (1 - e^{-\lambda x_{ij}})^{\alpha_i - 1} \right]$$

$$\cdot \prod_{j \in I_2} \left[\lambda^{n+1-s_j} \alpha e^{-\lambda(y_j + \sum_{i \notin S_j} x_{ij})} (1 - e^{-\lambda y_j})^{\alpha-1 + \sum_{i \in S_j} \alpha_i} \prod_{i \notin S_j} \alpha_i (1 - e^{-\lambda x_{ij}})^{\alpha_i - 1} \right]$$

$$= \alpha^{n_2} \lambda^{nn_1 + (n+1)n_2 - \sum_{j \in I_2} s_j} \exp \left[-\lambda \left(\sum_{j \in I_1} \sum_{i=1}^n x_{ij} + \sum_{j \in I_2} \sum_{i \notin S_j} x_{ij} + \sum_{j \in I_2} y_j \right) \right]$$

$$\cdot \prod_{j=1}^N (1 - e^{-\lambda y_j})^{\alpha-1 + \sum_{i \in S_j} \alpha_i} \prod_{j=1}^N \prod_{i \notin S_j} \alpha_i (1 - e^{-\lambda x_{ij}})^{\alpha_i-1} \prod_{j \in I_1} \prod_{i \in S_j} (\alpha + \alpha_i).$$

Therefor the log-likelihood function is

$$\ln L = n_2 \ln \alpha + \left[nn_1 + (n+1)n_2 - \sum_{j \in I_2} s_j \right] \ln \lambda -$$

$$\lambda \left(\sum_{j \in I_1} \sum_{i=1}^n x_{ij} + \sum_{j \in I_2} \sum_{i \notin S_j} x_{ij} + \sum_{j \in I_2} y_j \right)$$

$$+ \sum_{j=1}^N \left(\alpha - 1 + \sum_{i \in S_j} \alpha_i \right) \ln(1 - e^{-\lambda y_j}) + \sum_{j=1}^N \sum_{i \notin S_j} [\ln \alpha_i + (\alpha_i - 1) \ln(1 - e^{-\lambda x_{ij}})]$$

$$+ \sum_{j \in I_1} \sum_{i \in S_j} \ln(\alpha + \alpha_i).$$

With regard to parameter estimation, we are going to carry out the discussion in two aspects, one is α is known, and the MLEs can be obtained by getting fixed point; the other is α is unknown, and the MLEs can be obtained by calculating EM.

Case 1. α is known.

In this case for fixed λ , The MLE of $\alpha_1, \alpha_2, \dots, \alpha_n$, say $\hat{\alpha}_i(\lambda)$, $i = 1, 2, \dots, n$ can be obtained as the solutions of the following equations:

$$\frac{N - r_i}{\alpha_i} + \frac{r_{i1}}{\alpha + \alpha_i} = - \sum_{j: S_j \ni i} \ln(1 - e^{-\lambda y_j}) - \sum_{j: S_j \not\ni i} \ln(1 - e^{-\lambda x_{ij}}), \quad i = 1, 2, \dots, n,$$

where $r_{i1} = |\{j : j \in I_1, i \in S_j\}|$, $r_{i2} = |\{j : j \in I_2, i \in S_j\}|$, $r_i = r_{i1} + r_{i2}$. It is not difficult to show that

$$(1) \quad \hat{\alpha}_i(\lambda) = \frac{N - k_i \alpha - r_{i2} + \sqrt{(N - k_i \alpha - r_{i2})^2 + 4k_i(N - r_i)\alpha}}{2k_i}, \quad i = 1, 2, \dots, n,$$

where

$$k_i = - \left(\sum_{j:S_j \ni i} \ln(1 - e^{-\lambda y_j}) + \sum_{j:S_j i} \ln(1 - e^{-\lambda x_{ij}}) \right).$$

Once $\hat{\alpha}_i(\lambda)$ ($i = 1, 2, \dots, n$) are obtained, the MLE of λ can be obtained by maximizing the profile log-likelihood of λ . It can be obtained as a solution of the following fixed point type equation:

$$(2) \quad g(\lambda) = \lambda,$$

where

$$\begin{aligned} g(\lambda) = & \left[nn_1 + (n+1)n_2 - \sum_{j \in I_2} s_j \right] \left[\sum_{j \in I_1} \sum_{i=1}^n x_{ij} + \sum_{j \in I_2} \sum_{i \notin S_j} x_{ij} + \right. \\ & \sum_{j \in I_2} y_j - \sum_{j=1}^N \left(\alpha - 1 + \sum_{i \in S_j} \hat{\alpha}_i(\lambda) \right) \frac{y_j e^{-\lambda y_j}}{(1 - e^{-\lambda y_j})} - \\ & \left. \sum_{j=1}^N \sum_{i \notin S_j} (\hat{\alpha}_i(\lambda) - 1) \frac{x_{ij} e^{-\lambda x_{ij}}}{(1 - e^{-\lambda x_{ij}})} \right]^{-1}. \end{aligned}$$

Simple iterative procedure as follows can be use to compute the MLEs. We start the initial guess of λ as $\lambda^{(0)}$. Obtain $\hat{\alpha}_i(\lambda^{(0)})$ ($i = 1, 2, \dots, n$) from (1). Compute $\lambda^{(1)} = g(\lambda^{(0)})$ using (2). Replace $\lambda^{(0)}$ by $\lambda^{(1)}$ and repeat the process. The process continues until $|\lambda^{(i)} - \lambda^{(i+1)}| < \varepsilon$, where ε is some pre-assigned tolerance level.

Case 2. α is also unknown.

In this case we suggest EM algorithm to compute the MLEs of the unknown parameters. We treat this as a missing value problem as follows. Assume that for a random vector $(X_1, X_2, \dots, X_n)$, there is an associate random variable Δ , $\Delta = 1$ or 0 , if $U < \min_i \{U_i\}$ or $U > \min_i \{U_i\}$. Therefor, if $|S| \geq 2$, $\Delta = 0$, but if $|S| = 1$, then Δ is missing.

Now we provide the 'E'-step and 'M'-step of the EM algorithm. In the 'E'-step, we treat the observations as the complete observations if $|S| \geq 2$. If $|S| = 1$, we treat it as a missing observation, and form the 'pseudo observation' as $(x_1, x_2, \dots, x_n, \Delta)$. $\Delta = 1$ in the conditional probability $u_k = P\{U < U_k | X_1 = x_1, X_2 = x_2, \dots, X_n = x_n\}$ or 0 in the probability $1 - u_k$, where k is the only element in the S .

When $|S| = 1$, the conditional probability u_k can be obtained as

$$\begin{aligned} u_k &= P\{U < U_k | X_1 = x_1, X_2 = x_2, \dots, X_n = x_n\} \\ &= \frac{\lambda \alpha_k (1 - e^{-\lambda x_k})^{\alpha_k - 1} (1 - e^{-\lambda x_k})^\alpha \prod_{i \notin S} \lambda \alpha_i (1 - e^{-\lambda x_i})^{\alpha_i - 1}}{\lambda^n e^{-\lambda \sum_{i=1}^n x_i} (\alpha + \alpha_k) (1 - e^{-\lambda x_k})^{\alpha + \alpha_k - 1} \prod_{i \notin S} \alpha_i (1 - e^{-\lambda x_i})^{\alpha_i - 1}} = \frac{\alpha_k}{\alpha_k + \alpha}. \end{aligned}$$

For the sample $\{(x_{1j}, x_{2j}, \dots, x_{nj})\}_{j=1}^N$ from $MVGE(n, \alpha_1, \alpha_2, \dots, \alpha_n, \alpha, \lambda)$, the log-likelihood function of the 'pseudo data' can be written as

$$\begin{aligned} l_{\text{pseudo}} = & n_2 \ln \alpha + \left[nn_1 + (n+1)n_2 - \sum_{j \in I_2} s_j \right] \ln \lambda - \\ & \lambda \left(\sum_{j \in I_1} \sum_{i=1}^n x_{ij} + \sum_{j \in I_2} \sum_{i \notin S_j} x_{ij} + \sum_{j \in I_2} y_j \right) \\ & + \sum_{j=1}^N \sum_{i \notin S_j} [\ln \alpha_i + (\alpha_i - 1) \ln(1 - e^{-\lambda x_{ij}})] + \sum_{j \in I_2} \left(\alpha - 1 + \sum_{i \in S_j} \alpha_i \right) \ln(1 - e^{-\lambda y_j}) \\ & + \sum_{j \in I_1} \left[\left(\sum_{i \in S_j} u_i \right) \left(\sum_{i \in S_j} \ln \alpha_i + \sum_{i \in S_j} (\alpha_i + \alpha - 1) \ln(1 - e^{-\lambda y_j}) \right) \right. \\ & \left. + \left(1 - \sum_{i \in S_j} u_i \right) \left(\sum_{i \in S_j} \ln \alpha + \sum_{i \in S_j} (\alpha_i + \alpha - 1) \ln(1 - e^{-\lambda y_j}) \right) \right], \end{aligned}$$

where

$$(3) \quad u_i = \frac{\alpha_i}{\alpha_i + \alpha}, \quad i = 1, 2, \dots, n.$$

Now the 'M' step involves the maximization of the l_{pseudo} with respect to $\alpha_1, \alpha_2, \dots, \alpha_n, \alpha$ and λ at each stem. For fixed λ , the maximization of l_{pseudo} occurs at

$$(4) \quad \hat{\alpha}_i(\lambda) = - \frac{N - r_i + u_i r_{i1}}{\sum_{j: S_j \ni i} \ln(1 - e^{-\lambda y_j}) + \sum_{j: S_j \ni i} \ln(1 - e^{-\lambda x_{ij}})}, \quad i = 1, 2, \dots, n,$$

$$(5) \quad \hat{\alpha}(\lambda) = - \frac{n_2 + \sum_{i=1}^n r_{i1}(1 - u_i)}{\sum_{j=1}^N \ln(1 - e^{-\lambda y_j})}.$$

The MLE of λ can be obtained as a solution of the following fixed point type equation:

$$(6) \quad g(\lambda) = \lambda,$$

where

$$\begin{aligned} g(\lambda) = & \left[nn_1 + (n+1)n_2 - \sum_{j \in I_2} s_j \right] \left[\sum_{j \in I_1} \sum_{i=1}^n x_{ij} + \sum_{j \in I_2} \sum_{i \notin S_j} x_{ij} + \right. \\ & \left. \sum_{j \in I_2} y_j - \sum_{j=1}^N \left(\hat{\alpha}(\lambda) - 1 + \sum_{i \in S_j} \hat{\alpha}_i(\lambda) \right) \frac{y_j e^{-\lambda y_j}}{(1 - e^{-\lambda y_j})} \right] \end{aligned}$$

$$\sum_{j=1}^N \sum_{i \notin S_j} (\hat{\alpha}_i(\lambda) - 1) \frac{x_{ij} e^{-\lambda x_{ij}}}{(1 - e^{-\lambda x_{ij}})]^{-1}.$$

Now we describe how to compute $(t + 1)$ -th step from the t -th step in the EM algorithm.

- Step 1: Suppose at the t -th step the estimates of α_i, α and λ are $\alpha_i^{(t)}, \alpha^{(t)}$ and $\lambda^{(t)}$ respectively.
- Step 2: Compute u_i using $\alpha_i^{(t)}$ and $\alpha^{(t)}$ by (3).
- Step 3: Find $\lambda^{(t+1)}$ by solving (6).
- Step 4: Once $\lambda^{(t+1)}$ is obtained, compute $\alpha_i^{(t+1)} = \hat{\alpha}_i(\lambda^{(t+1)})$ and $\alpha^{(t+1)} = \hat{\alpha}(\lambda^{(t+1)})$ by (4) and (5).

In order to reduce the workload of finding fixed point in equation (6), we can adopt the improved method as follows by replacing Step 3 with $\lambda^{(t+1)} = g(\lambda^{(t)})$.

- Step 1: Suppose at the t -th step the estimates of α_i, α and λ are $\alpha_i^{(t)}, \alpha^{(t)}$ and $\lambda^{(t)}$ respectively.
- Step 2: Compute u_i using $\alpha_i^{(t)}$ and $\alpha^{(t)}$ by (3).
- Step 3: Compute $\alpha_i^{(t+1)} = \hat{\alpha}_i(\lambda^{(t)})$ and $\alpha^{(t+1)} = \hat{\alpha}(\lambda^{(t)})$ by (4) and (5).
- Step 4: Compute $\lambda^{(t+1)} = g(\lambda^{(t)})$ using $\alpha_i^{(t+1)}$ and $\alpha^{(t+1)}$ by (6).

Stop iterating if some convergence criterion is satisfied.

4. ILLUSTRATIVE EXAMPLE

In order to illustrate the EM type algorithm derived above, we use a simulated data set from a MVGE distribution. Consider the data of Table 1 generated from MVGE(4, 1, 2, 4, 5, 3, 3). That is to say $n = 4, \alpha_1 = 1, \alpha_2 = 2, \alpha_3 = 4, \alpha_4 = 5, \alpha = 3$ and $\lambda = 3$.

The initial values were set arbitrarily equal to 1 for all the parameters. The algorithm stopped iterating when $|\alpha^{(t+1)} - \alpha^t| + |\lambda^{(t+1)} - \lambda^t| + \sum_{i=1}^n |\alpha_i^{(t+1)} - \alpha_i^{(t)}|$ became smaller than $\varepsilon = 10^{-20}$. The results of the algorithm are described in the Table 2. The history of the parameter values with respect to the iterations can be seen in Fig 1 and the log-likelihood in Fig 2. The algorithm converged quite quickly, after 27 iterations, and likelihood function increased slightly after the 7th iteration.

Several other initial values were used without change in the final solution. It is quite interesting that even if one starts from a point far away from the ML estimates, like all initial values were set equal to 100 or 0.001, the algorithm converged after

TABLE 1. Generated data from MVGE(4, 1, 2, 4, 5, 3, 3)

X_{1j}	X_{2j}	X_{3j}	X_{4j}	X_{1j}	X_{2j}	X_{3j}	X_{4j}
0.99944	0.47636	0.71409	0.66901	0.71381	0.71381	0.71381	0.71381
0.73972	0.68781	0.57526	0.39909	0.40368	0.40368	0.68684	0.97716
0.57422	0.37965	0.72303	1.0294	0.26589	0.53703	0.37587	0.57905
0.84955	0.65345	0.41808	0.60043	0.50979	0.58124	0.57916	0.74492
0.91354	1.0532	0.53702	1.2687	0.52697	0.28602	0.71	0.28602
0.6097	0.6097	0.99598	0.6097	0.46255	0.46255	1.137	0.46255
0.22635	0.22635	0.36721	0.7812	0.99515	1.407	1.9987	1.024
0.13678	0.19679	0.14439	0.95652	0.51786	0.51786	0.51786	0.82409
0.49158	1.1204	0.58332	0.61123	0.90827	1.3015	0.90827	0.90827
1.0307	1.0307	1.0307	1.0307	0.49335	0.49335	0.56141	1.4303
1.2659	1.2659	1.2659	1.2659	0.49056	0.49056	1.0499	0.80781
0.59419	0.41158	0.83274	0.61983	0.48514	0.48514	0.71445	0.80901
0.26119	0.26119	0.36274	0.86958	0.4545	0.47347	0.57101	0.4545
1.0388	1.0388	1.0388	1.0388	1.1278	1.1278	1.1278	1.1278
0.32219	0.87412	1.0832	0.77066	0.38307	0.12114	0.18947	0.79022
0.66597	0.98912	1.0127	0.82568	0.36704	0.36704	0.36704	0.42129
1.4984	1.4984	1.4984	1.4984	0.2947	0.15208	0.50265	0.81072
0.68739	0.68739	0.85648	0.68739	0.72097	0.72097	0.72097	0.72097
1.0321	1.0321	1.0321	1.0321	0.67068	0.67068	0.67068	0.86155
0.26364	0.36653	0.81361	0.79972	1.7222	1.7222	1.7222	1.7222
0.52796	1.275	0.6329	1.2268	0.5728	0.5728	0.84123	1.266
0.91378	1.5292	0.91378	0.91378	0.44501	0.44501	1.1133	0.45791
0.69513	0.65221	0.46867	0.46867	0.5447	1.0199	0.39466	0.46379
1.0902	1.0902	1.0902	1.0902	0.64386	0.64386	0.75035	0.64386
0.77571	0.77571	0.77571	2.0186	0.77356	0.77356	0.77356	2.2579

TABLE 2. Results of the algorithm to the data of Table 1

Parameters	α_1	α_2	α_3	α_4	α	λ
Exact values	1	2	4	5	3	3
Estimates	1.3119	1.7867	3.9829	5.8183	3.2544	3.2108

FIGURE 1. The history of the parameter values with respect to the iterations

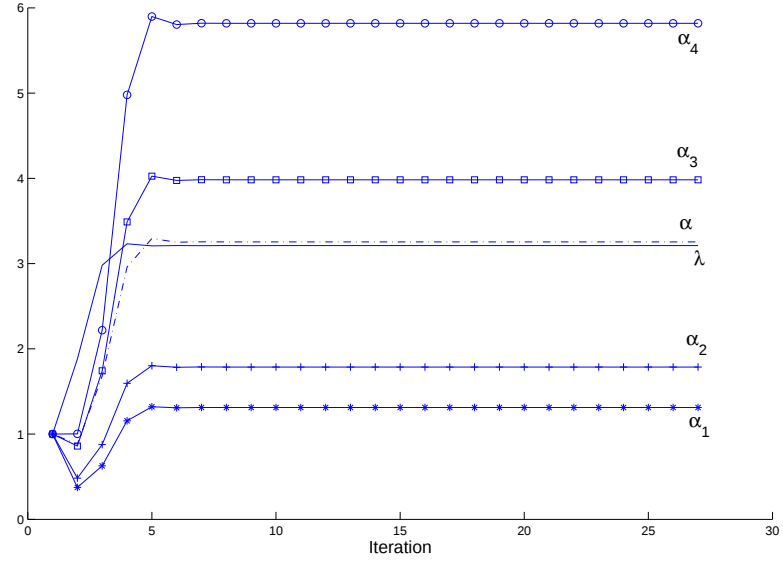
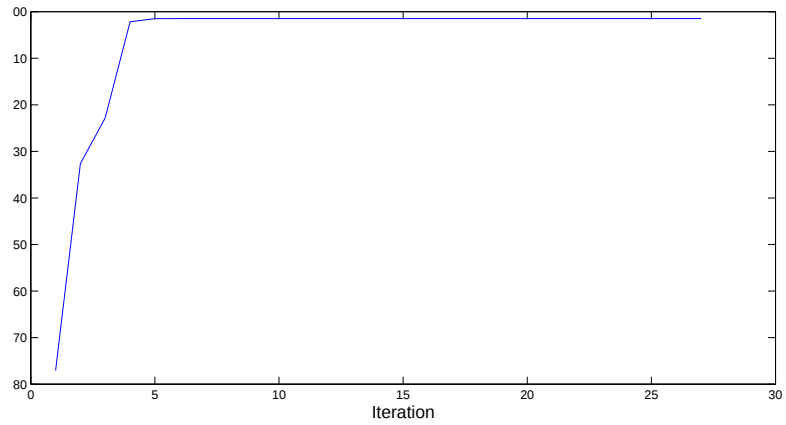


FIGURE 2. The history of the log-likelihood values with respect to the iterations



27 or 55 iteration (with the same stopping criterion). This is an indication that the choice of initial values is not so important.

5. CONCLUDING REMARKS

In this paper, we have propose the multivariate generalized exponential distribution function whose marginals are generalized exponential distributions. The explicit forms of the joint distribution function and the joint density function have been given, so that it can be used in practice. An EM type algorithm for the MLEs of the unknown parameters of the MVBE distribution was described.

REFERENCES

- [1] B. Bemis, L.J. Bain and J.J. Higgins, Estimation and hypothesis testing for the parameters of bivariate exponential distribution, *Journal of the American Statistical Association* 67 (1972) 927-929.
- [2] R.D. Gupta, Generalized exponential distributions: Existing results and some recent developments, *Journal of Statistical Planning and Inference* 137 (2007) 3525-3536.
- [3] R.D. Gupta and D. Kundu, Generalized exponential distributions, *Australian and New Zealan Journal of Statistics* 41 (1999) 173-188.
- [4] D. Kundu and R.D. Gupta, Bivariate generalized exponential distribution, *Journal of Multivariate Anlysis* 100 (2009) 581-593.
- [5] A. Sarhan and N. Balakrishnan, A new class of bivariate distribution and its mixture, *Journal of the Multivariate Analysis* 98 (2007) 1508-1527.