

A STATIONARY WEIBULL FRAILTY PROCESS

DEBASIS KUNDU *

Abstract

In this paper we have developed two correlated stationary processes X_n and Y_n , where $P(X_n = X_{n+1}) > 0$ and $P(Y_n = Y_{n+1}) > 0$. The model has been developed based on two real life data sets. Efficient inference procedures and goodness of fit test have been developed. The motivation of this work came when we were trying to analyze the gold price data $\{X_n\}$ of the Indian market and the exchange rate data $\{Y_n\}$ between Indian Rupees and US dollars during the same time period. Although, both of them are stationary processes, it is observed that there is a significant amount of time $X_n = X_{n+1}$ or $Y_n = Y_{n+1}$, and they cannot be ignored. Moreover, X_n and Y_n are dependent, and if we want to analyze the data during the same period, they should be analyzed together. In this paper first we have proposed a new stationary process based on Marshall-Olkin bivariate exponential distribution type construction combining with frailty, so that $P(X_n = X_{n+1}) > 0$. Different properties of the proposed bivariate stationary process have been established. The maximum likelihood estimators of the unknown parameters have been obtained. Goodness of fit test of the proposed model has been proposed. We have performed the analysis of gold price data and the exchange rate data using the proposed model and the results are quite satisfactory. Then we have extended the process to the bivariate case where $P(X_n = X_{n+1}) > 0$, $P(Y_n = Y_{n+1}) > 0$ and X_n, Y_n are dependent. We have demonstrated how it can be used in practice.

KEY WORDS AND PHRASES: Marshall-Olkin bivariate exponential distribution; Marshall-Olkin bivariate Weibull distribution; bivariate singular distribution; bivariate copula; positive dependence, maximum likelihood estimators; EM algorithm.

AMS SUBJECT CLASSIFICATIONS: 62F10, 62F03, 62H12.

*Department of Mathematics and Statistics, Indian Institute of Technology Kanpur, Pin 208016, India.
E-mail: kundu@iitk.ac.in, Phone no. 91-512-2597141, Fax no. 91-512-2597500.

1 MOTIVATION AND INTRODUCTION

The motivation of this work came when we were trying to analyze the gold price data $\{X_n\}$ of the Indian market, and the exchange rate data $\{Y_n\}$ of the Indian Rupees and US dollars during the same period. It has been observed that both $\{X_n\}$ and $\{Y_n\}$ are continuous random variables. They are stationary processes, but there are significant cases where $X_n = X_{n+1}$ or $Y_n = Y_{n+1}$ and they cannot be ignored. Further, X_n and Y_n are not independent. Hence, to analyze the two series separately leads to loss of information. If they are analyzed jointly then the information of one series can be used in the other. In this paper first we have developed a stationary process $\{X_n\}$ such that $P(X_n = X_{n+1}) > 0$. We have provided several properties and developed the classical inference procedure of the proposed process. Further, we have developed a bivariate stationary process $\{(X_n, Y_n) : n \geq 1\}$, so that both $\{X_n\}$ and $\{Y_n\}$ are continuous stationary processes, $P(X_n = X_{n+1}) > 0$, $P(Y_n = Y_{n+1}) > 0$, and X_n, Y_n are dependent. We have demonstrated how this process can be used in practice. It may be noted that although X_n 's are continuous random variables for all $n \geq 1$, since X_n and X_{n+1} are dependent, it is possible to generate $\{X_n; n \geq 1\}$, such that $P(X_n = X_{n+1}) > 0$, see for example Kundu [4]. The same is true for $\{Y_n; n \geq 1\}$ also.

It may be mentioned that there are several positive valued stationary processes available in the literature, for example the exponential process by Tavares [13], Weibull and gamma process by Sim [12], Pareto process by Yeh et al. [15], semi-Pareto process by Pillai [11], Weibull process by Jayakumar and Girish Babu [3] and see the references cited there in. Unfortunately in all these cases $P(X_n = X_{n+1}) = 0$. Arnold and Hallet [1] developed a Pareto process, but they did not provide any inference procedure of the process. Kundu [6] developed a proportional hazard class process which can take positive probability on $X_n = X_{n+1}$. But it is not immediate in both these cases how they can be generalized to bivariate or multivariate cases. It is an attempt towards that direction.

Part of the main idea of this process comes from the construction of the Marshall-Olkin bivariate exponential (MOBE) distribution, see for example Marshall and Olkin [9] or Kundu

[4] for this construction. Using this approach it is possible to generate a stationary process $\{X_n\}$, such that $P(X_n = X_{n+1}) > 0$, and this approach has been used by Kundu [6] to define a proportional hazard class process. But the same procedure cannot be extended to define a bivariate process. It may be recalled that Vaupel [14] introduced a mathematical model to introduce dependence between two variables through another random variable. It is known as frailty. In this paper we have incorporated frailty in the MOBE type distribution, to define a univariate stationary process $\{X_n\}$ and also a bivariate stationary process $\{X_n, Y_n\}$, so that X_n, Y_n can be dependent, and further $P(X_n = X_{n+1}) > 0$, $P(Y_n = Y_{n+1}) > 0$. We call the process $\{X_n\}$ as univariate Weibull-frailty process (UWFP) and the bivariate process $\{X_n, Y_n\}$ as the bivariate Weibull-frailty process (BWFP).

First we have provided different properties of the UWFP. Generation of the random sample from the UWFP is an important issue for simulation purposes, and we have addressed how to generate random samples from this process. It has been shown that (X_n, Y_n) has a copula which can be expressed in explicit forms. Hence, many dependence properties can be obtained from this copula function. The maximum likelihood estimators of the unknown parameters can not be obtained in explicit forms. It involves solving a four dimensional optimization problem. To avoid that we have used EM algorithm and profile maximization method. It involves solving two one dimensional optimization methods. The proposed EM algorithm can be used quite conveniently for implementation purposes. We have performed some simulation experiments to show how the EM algorithm works in practice and also we have analyzed two data sets namely the gold price data and exchange rate data to show the effectiveness of the proposed model. We have provided an empirical goodness of fit test procedure also. It shows that the proposed model can be used to analyze these two data sets.

Further, we have provided the construction of the BWFP and derive some of its properties. It has seven unknown parameters. In this case also the MLEs of unknown parameters cannot be obtained in explicit forms, and they can only be obtained by solving a seven dimensional optimization problem. We have used EM algorithm and profile maximization

method to compute the MLEs of the unknown parameters. The proposed EM algorithm involves solving three one dimensional optimization problems. Hence, the implementation of the proposed EM algorithm is quite convenient in practice. We have performed some simulation experiments to show how the EM algorithm works in practice for the bivariate case and also we have analyzed the two data sets jointly based on the proposed BWFP. The results are quite satisfactory.

The rest of the paper is organized as follows. In Section 2 we have provided some of the preliminaries and introduced the univariate and bivariate Weibull-frailty models. Based on Weibull-frailty distribution, UWFP process has been defined and their several properties have been provided in Section 3. The classical inference of the unknown parameters have been provided in Section 4. In Section 6 we have described an empirical goodness of fit test procedure. The analysis of the gold price data and exchange rate data have been presented in 7. In Section 8 we have introduced the BWFP and demonstrated how it can be used in practice. Finally in Section 9 we have concluded the paper.

2 NOTATIONS, PRELIMINARIES AND SOME NEW RESULTS

2.1 NOTATIONS AND PRELIMINARIES

We have used the following notations in this paper. A random variable X is said to have an Weibull distribution with the shape parameter α and the scale parameter λ , if it has the following probability density function (PDF).

$$f_{WE}(x; \alpha, \lambda) = \begin{cases} \alpha \lambda x^{\alpha-1} e^{-\lambda x^\alpha} & \text{if } x > 0 \\ 0 & \text{if } x \leq 0 \end{cases} \quad (1)$$

The corresponding cumulative distribution function (CDF), survival function (SF) and the hazard function (HF) for $x > 0$, respectively, will be as follows;

$$F_{WE}(x; \alpha, \lambda) = 1 - e^{-\lambda x^\alpha}, \quad S_{WE}(x; \alpha, \lambda) = e^{-\lambda x^\alpha}, \quad h_{WE}(x; \alpha, \lambda) = \alpha \lambda x^{\alpha-1}$$

From now on it will be denoted by WE(α, λ). A random variable X is said to have a gamma distribution with the shape parameter α and scale parameter λ , if it has the following

probability density function (PDF).

$$f_{GA}(x; \alpha, \lambda) = \begin{cases} \frac{\lambda^\alpha}{\Gamma(\alpha)} x^{\alpha-1} e^{-\lambda x} & \text{if } x > 0 \\ 0 & \text{if } x \leq 0. \end{cases} \quad (2)$$

From now on it will be denoted by $GA(\alpha, \lambda)$.

Marshall and Olkin [9] developed MOBE distribution. Following the same procedure several other bivariate distributions have been proposed in the literature. In particular Marshall-Olkin bivariate Weibull (MOBW) distribution has become very popular because of its wide scale applicability. Now we will briefly mention how the MOBW construction has been done and it will be used later. Let $U_1 \sim \text{WE}(\alpha, \lambda_1)$, $U_2 \sim \text{WE}(\alpha, \lambda_2)$ and $U_3 \sim \text{WE}(\alpha, \lambda_3)$ and they are independently distributed. Then (X, Y) , where $X = \min\{U_1, U_3\}$ and $Y = \min\{U_2, U_3\}$ is said to have MOBW distribution. The joint PDF of (X, Y) can be written as

$$f_{MOBW}(x, y) = \begin{cases} f_{WE}(x; \alpha, \lambda_1) f_{WE}(y; \alpha, \lambda_2 + \lambda_3) & \text{if } x < y \\ f_{WE}(x; \alpha, \lambda_1 + \lambda_3) f_{WE}(y; \alpha, \lambda_2) & \text{if } x > y \\ \frac{\lambda_3}{\lambda_1 + \lambda_2 + \lambda_3} f_{WE}(z; \alpha, \lambda_1 + \lambda_2 + \lambda_3) & \text{if } x = y = z. \end{cases} \quad (3)$$

It may be mentioned that although U_1, U_2 and U_3 are independently distributed, X and Y are dependent due to the presence of U_3 . The copula associated with the MOBW distribution is

$$\overline{C}(u, v) = \begin{cases} u^{\frac{\lambda_1}{\lambda_1 + \lambda_3}} v & \text{if } u^{\frac{1}{\lambda_1 + \lambda_3}} < v^{\frac{1}{\lambda_2 + \lambda_3}} \\ uv^{\frac{\lambda_2}{\lambda_1 + \lambda_3}} & \text{if } u^{\frac{1}{\lambda_1 + \lambda_3}} < v^{\frac{1}{\lambda_2 + \lambda_3}}. \end{cases}$$

It shows that the correlation between X and Y depends only of λ_1 , λ_2 and λ_3 and it is independent of α .

2.2 UNIVARIATE AND BIVARIATE WEIBULL-FRAILTY DISTRIBUTIONS

Definition: Suppose $U \sim \text{WE}(\alpha, \lambda)$, $V \sim \text{GA}(\beta, \beta)$ and they are independent. Then $X = \frac{U}{V^{1/\alpha}}$ is said to have univariate Weibull-frailty (UWF) distribution with parameters α , β and λ . It will be denoted by $\text{UWF}(\alpha, \beta, \lambda)$.

Definition: Suppose $U_1 \sim \text{WE}(\alpha, \lambda_1)$, $U_2 \sim \text{WE}(\alpha, \lambda_2)$, $U_3 \sim \text{WE}(\alpha, \lambda_3)$, $V \sim \text{GA}(\beta, \beta)$ and they are all independently distributed. Let us define

$$X = \min \left\{ \frac{U_1}{V^{1/\alpha}}, \frac{U_3}{V^{1/\alpha}} \right\} \quad \text{and} \quad Y = \min \left\{ \frac{U_2}{V^{1/\alpha}}, \frac{U_3}{V^{1/\alpha}} \right\}. \quad (4)$$

Then (X, Y) is said to have bivariate Weibull-frailty (BWF) distribution with parameters $\alpha, \beta, \lambda_1, \lambda_2, \lambda_3$, see for example Kundu [7]. It will be denoted by $\text{BWF}(\alpha, \beta, \lambda_1, \lambda_2, \lambda_3)$. It may be mentioned that in both the above definitions, it is important to take the scale and shape parameters of V to be same, otherwise, they are not identifiable. It is clear from the construction only.

The joint SF of (X, Y) for $x > 0$ and $y > 0$ becomes

$$\begin{aligned} S_{BWF}(x, y) &= P(X > x, Y > y) \\ &= \left[1 + \frac{\lambda_1}{\beta} x^\alpha + \frac{\lambda_2}{\beta} y^\alpha + \frac{\lambda_3}{\beta} \max\{x^\alpha, y^\alpha\} \right]^{-\beta} \\ &= \begin{cases} [1 + \theta_1 x^\alpha + (\theta_2 + \theta_3) y^\alpha]^{-\beta} & \text{if } x < y \\ [1 + (\theta_1 + \theta_3) x^\alpha + \theta_2 y^\alpha]^{-\beta} & \text{if } y \leq x. \end{cases} \end{aligned}$$

Here we have denoted $\theta_1 = \lambda_1/\beta$, $\theta_2 = \lambda_2/\beta$ and $\theta_3 = \lambda_3/\beta$. Hence, the marginal SFs of X and Y become

$$P(X > x) = [1 + (\theta_1 + \theta_3) x^\alpha]^{-\beta} \quad \text{and} \quad P(Y > y) = [1 + (\theta_2 + \theta_3) y^\alpha]^{-\beta}.$$

The PDFs of X and Y for $x > 0$ and $y > 0$ become

$$f_X(x) = \frac{\alpha\beta(\theta_1 + \theta_3)x^{\alpha-1}}{[1 + (\theta_1 + \theta_3)x^\alpha]^{\beta+1}} \quad \text{and} \quad f_Y(y) = \frac{\alpha\beta(\theta_2 + \theta_3)y^{\alpha-1}}{[1 + (\theta_2 + \theta_3)y^\alpha]^{\beta+1}},$$

respectively. Note that $X \sim \text{UWF}(\alpha, \beta, \theta_1 + \theta_3)$ and $Y \sim \text{UWF}(\alpha, \beta, \theta_2 + \theta_3)$.

Observe that $\{(X, Y)|V = v\} \sim \text{MOBW}(\alpha, \lambda_1 v, \lambda_2 v, \lambda_3 v)$. Hence, using (3) we can write the joint PDF of (X, Y) given $V = v$ as follows;

$$f_{(X,Y)|V=v}(x, y) = \begin{cases} f_{WE}(x; \alpha, \lambda_1 v) f_{WE}(y; \alpha, (\lambda_2 + \lambda_3) v) & \text{if } x < y \\ f_{WE}(x; \alpha, (\lambda_1 + \lambda_3) v) f_{WE}(y; \alpha, \lambda_2 v) & \text{if } y < x \\ \frac{\lambda_3}{\lambda_1 + \lambda_2 + \lambda_3} f_{WE}(z; \alpha, (\lambda_1 + \lambda_2 + \lambda_3) v) & \text{if } x = y = z. \end{cases} \quad (5)$$

Hence, the joint PDF of (X, Y) becomes

$$f_{BWF}(x, y) = \begin{cases} \frac{\alpha^2 \beta (\beta+1) \theta_1 \theta_2 x^{\alpha-1} y^{\alpha-1}}{(1 + \theta_1 x^\alpha + (\theta_2 + \theta_3) y^\alpha)^{\beta+2}} & \text{if } x < y \\ \frac{\alpha^2 \beta (\beta+1) \theta_1 \theta_2 x^{\alpha-1} y^{\alpha-1}}{(1 + (\theta_1 + \theta_3) x^\alpha + \theta_2 y^\alpha)^{\beta+2}} & \text{if } y < x \\ \frac{\alpha \theta_3 \beta z^{\alpha-1}}{(1 + (\theta_1 + \theta_2 + \theta_3) z^\alpha)^{\beta+1}} & \text{if } x = y = z. \end{cases} \quad (6)$$

The BWF distribution has an absolutely continuous part and a singular part. It is absolutely continuous on $0 < x \neq y < \infty$ and it has a singular part on $0 < x = y < \infty$. The joint PDF

(6) can be written as

$$f_{BWF}(x, y) = \frac{\theta_1 + \theta_2}{\theta_1 + \theta_2 + \theta_3} f_{ac}(x, y) + \frac{\theta_3}{\theta_1 + \theta_2 + \theta_3} f_{si}(x, y), \quad (7)$$

where

$$f_{ac}(x, y) = \frac{\theta_1 + \theta_2 + \theta_3}{\theta_1 + \theta_2} \begin{cases} \frac{\alpha^2 \beta (\beta+1) \theta_1 \theta_2 x^{\alpha-1} y^{\alpha-1}}{(1+\theta_1 x^\alpha + (\theta_2 + \theta_3) y^\alpha)^{\beta+2}} & \text{if } x < y \\ \frac{\alpha^2 \beta (\beta+1) \theta_1 \theta_2 x^{\alpha-1} y^{\alpha-1}}{(1+(\theta_1 + \theta_3) x^\alpha + \theta_2 y^\alpha)^{\beta+2}} & \text{if } y < x \end{cases}$$

and

$$f_{si}(x, y) = \begin{cases} \frac{\alpha \beta (\theta_1 + \theta_2 + \theta_3) z^{\alpha-1}}{(1+(\theta_1 + \theta_2 + \theta_3) z^\alpha)^{\beta+1}} & \text{if } x = y = z \\ 0 & \text{if } x \neq y. \end{cases}$$

It can be easily seen that

$$P(X < Y) = \frac{\theta_1}{\theta_1 + \theta_2 + \theta_3}, \quad P(Y < X) = \frac{\theta_2}{\theta_1 + \theta_2 + \theta_3}, \quad P(X = Y) = \frac{\theta_3}{\theta_1 + \theta_2 + \theta_3}.$$

The BWF has the following survival copula

$$\overline{C}(u, v) = \begin{cases} \left(\frac{\theta_1}{\theta_1 + \theta_3} (u^{-1/\beta} - 1) + v^{-1/\beta} \right)^{-\beta} & \text{if } (\theta_2 + \theta_3)(u^{-1/\beta} - 1) < (\theta_1 + \theta_3)(v^{-1/\beta} - 1) \\ \left(u^{-1/\beta} + \frac{\theta_2}{\theta_2 + \theta_3} (v^{-1/\beta} - 1) \right)^{-\beta} & \text{if } (\theta_2 + \theta_3)(u^{-1/\beta} - 1) > (\theta_1 + \theta_3)(v^{-1/\beta} - 1). \end{cases} \quad (8)$$

If $\theta_1 = \theta_2$ and if we denote $\delta = \frac{\theta_1}{\theta_1 + \theta_3} = \frac{\theta_2}{\theta_2 + \theta_3}$, then (8) becomes

$$\overline{C}(u, v) = \begin{cases} (\delta(u^{-1/\beta} - 1) + v^{-1/\beta})^{-\beta} & \text{if } u > v \\ (u^{-1/\beta} + \delta(v^{-1/\beta} - 1))^{-\beta} & \text{if } u < v. \end{cases} \quad (9)$$

It shows that the correlation between X and Y depends on λ_1 , λ_2 , λ_3 and β also. It is independent of α in this case also. Since it depends on β also, it brings more flexibility in the correlation between X and Y . Different dependency properties of BWF may be explored through this copula function.

3 UNIVARIATE WEIBULL-FRAILTY PROCESS

Let $\{U_n; n \geq 1\}$ be a sequence of independent and identically distributed (i.i.d.) $\text{WE}(\alpha, \lambda_1)$ random variables, $\{W_n; n \geq 1\}$ be a sequence of independent and identically distributed (i.i.d.) $\text{WE}(\alpha, \lambda_2)$ random variables and they are independent. Let V be a $\text{GA}(\beta, \beta)$ random variables and it is independent of $\{U_n; n \geq 1\}$ and $\{W_n; n \geq 1\}$. Now we define a discrete

state space continuous time stochastic process $\{X_n; n \geq 1\}$ as follows;

$$X_n = \min \left\{ \frac{U_n}{V^{1/\alpha}}, \frac{U_{n+1}}{V^{1/\alpha}}, \frac{W_n}{V^{1/\alpha}} \right\}; \quad n \geq 1. \quad (10)$$

We will call $\{X_n; n \geq 1\}$ as the univariate Weibull-frailty process (UWFP) with parameters $(\alpha, \lambda_1, \lambda_2, \beta)$. From now on, it will be denoted by $\text{UWFP}(\alpha, \lambda_1, \lambda_2, \beta)$. The generation of a random sample from the process $\{X_n\}$ is quite straightforward, it follows from the definition of the process (10). We have the following result.

Theorem 1: If $\{X_n\}$ is a $\text{UWFP}(\alpha, \lambda_1, \lambda_2, \beta)$, then

- (i) $\{X_n\}$ is a stationary process.
- (ii) $X_n \sim \text{UWF}(\alpha, \beta, 2\theta_1 + \theta_2)$, where $\theta_1 = \frac{\lambda_1}{\beta}$ and $\theta_2 = \frac{\lambda_2}{\beta}$

Proof: (i) follows from the definition only. The proof of (ii) follows from the definition of X_n as in (10) and by conditioning on V . ■

Theorem 2: Suppose $\{X_n\}$ is a $\text{UWFP}(\alpha, \lambda_1, \lambda_2, \beta)$.

- (i) The joint survival function of $X_1, \dots, X_n$ for $x_1 > 0, \dots, x_n > 0$, is

$$S_{X_1, \dots, X_n}(x_1, \dots, x_n) = \left(1 + \lambda_1 \left(x_1^\alpha + \sum_{i=1}^{n-1} z_i^\alpha + x_n^\alpha \right) + \lambda_2 \sum_{i=1}^n x_i^\alpha \right)^{-\beta},$$

here $z_i = \max\{x_i, x_{i+1}\}$, for $i = 1, \dots, n$.

- (ii) $\{X_n | V = v\} \sim \text{WE}(\alpha, v(2\lambda_1 + \lambda_2))$.
- (iii) $\{(X_n, X_{n+1}) | V = v\} \sim \text{MOBW}(\alpha, v(\lambda_1 + \lambda_2), v(\lambda_1 + \lambda_2), v\lambda_1)$
- (iv) The joint survival function of X_n and X_{n+m} for $n \geq 1, m > 1$, given $V = v$, for $x_n > 0, x_{n+m} > 0, v > 0$, is

$$S_{X_n, X_{n+m} | V=v}(x, y) = P(X_n \geq x | V = v) P(X_{n+m} \geq y | V = v).$$

- (v) For any $x_1 > 0, \dots, x_n > 0, v > 0$,

$$P(X_n \geq x_n | X_{n-1} = x_{n-1}, \dots, x_1 = x_1, V = v) = P(X_n \geq x_n | X_{n-1} = x_{n-1}, V = v).$$

Proof: See in the Appendix. ■

Now we are going to derive the joint PDF of $X_1, \dots, X_{n+1}$, it will be helpful to develop the likelihood inference of the unknown parameter. The joint PDF of $\{(X_n, X_{n+1})|V = v\}$ is

$$f_{X_n, X_{n+1}|V=v}(x_n, x_{n+1}) = \begin{cases} f_{WE}(x_n; \alpha, (\lambda_1 + \lambda_2)v) f_{WE}(x_{n+1}; \alpha, (2\lambda_1 + \lambda_2)v) & \text{if } x_n < x_{n+1} \\ f_{WE}(x_n; \alpha, (2\lambda_1 + \lambda_2)v) f_{WE}(x_{n+1}; \alpha, (\lambda_1 + \lambda_2)v) & \text{if } x_n > x_{n+1} \\ \frac{\lambda_1}{3\lambda_1 + 2\lambda_2} f_{WE}(x_n; \alpha, (3\lambda_1 + 2\lambda_2)v) & \text{if } x_n = x_{n+1}. \end{cases} \quad (11)$$

Using (v) of Theorem 2, the joint PDF of $X_1, \dots, X_{n+1}$ given $V = v$, can be written as

$$\begin{aligned} f_{X_1, \dots, X_{n+1}|V=v}(x_1, \dots, x_{n+1}) &= f_{X_1|V=v}(x_1) f_{X_2|X_1=x_1, V=v}(x_2) \cdots f_{X_{n+1}|X_n=x_n, V=v}(x_{n+1}), \\ &= \frac{f_{X_1|V=v}(x_1) f_{X_1, X_2|V=v}(x_1, x_2) \cdots f_{X_n, X_{n+1}|V=v}(x_n, x_{n+1})}{\prod_{i=1}^n f_{X_i|V=v}(x_i)}, \\ &= \frac{\prod_{i=1}^n f_{X_i, X_{i+1}|V=v}(x_i, x_{i+1})}{\prod_{i=2}^n f_{X_i|V=v}(x_i)}, \end{aligned}$$

for $x_1 > 0, \dots, x_{n+1} > 0, v > 0$. Let us use the following notations: $I = \{1, \dots, n\}$, $I_1 = \{i : i \in I, x_i < x_{i+1}\}$, $I_2 = \{i : i \in I, x_i > x_{i+1}\}$, $I_3 = \{i : i \in I, x_i = x_{i+1}\}$ and n_1, n_2 and n_3 denote the number of points in I_1, I_2 and I_3 , respectively. Hence,

$$\begin{aligned} \prod_{i=1}^n f_{X_i, X_{i+1}|V=v}(x_i, x_{i+1}) &= (\lambda_1 + \lambda_2)^{n_1+n_2} (2\lambda_1 + \lambda_2)^{n_1+n_2} \lambda_1^{n_3} \alpha^{n+n_1+n_2} v^{n+n_1+n_2} \times \\ &\quad \left(\prod_{i \in I} x_i^{\alpha-1} \right) \left(\prod_{i \in I_1 \cup I_2} x_{i+1}^{\alpha-1} \right) e^{-v((\lambda_1 + \lambda_2)A'_1(\mathbf{x}; \alpha) + \lambda_1 A'_2(\mathbf{x}; \alpha))} \\ \prod_{i=2}^n f_{X_i|V=v}(x_i) &= \alpha^{n-1} v^{n-1} (2\lambda_1 + \lambda_2)^{n-1} \prod_{i=2}^n x_i^{\alpha-1} e^{-(2\lambda_1 + \lambda_2)v \sum_{i=2}^n x_i^\alpha}, \end{aligned} \quad (12)$$

here

$$A'_1(\mathbf{x}; \alpha) = \sum_{i=1}^n (x_i^\alpha + x_{i+1}^\alpha) \quad \text{and} \quad A'_2(\mathbf{x}; \alpha) = \sum_{i \in I_2 \cup I_3} x_i^\alpha + \sum_{i \in I_1} x_{i+1}^\alpha.$$

Hence,

$$\begin{aligned} f_{X_1, \dots, X_{n+1}|V=v}(x_1, \dots, x_{n+1}) &= \frac{\alpha^{n_1+n_2} (\lambda_1 + \lambda_2)^{n_1+n_2} \lambda_1^{n_3}}{(2\lambda_1 + \lambda_2)^{n_3}} \left(\prod_{i \in I_1 \cup I_2} x_{i+1}^{\alpha-1} \right) \times \\ &\quad v^{n_1+n_2} e^{-v((\lambda_1 + \lambda_2)A_1(\mathbf{x}; \alpha) + \lambda_1 A_2(\mathbf{x}; \alpha))}. \end{aligned} \quad (13)$$

$$A_1(\mathbf{x}; \alpha) = \sum_{i=1}^n x_{i+1}^\alpha + x_1^\alpha \quad \text{and} \quad A_2(\mathbf{x}; \alpha) = \sum_{i \in I_1} (x_{i+1}^\alpha - x_i^\alpha) + x_1^\alpha.$$

Therefore, the joint PDF of $X_1, \dots, X_n$ becomes

$$\begin{aligned}
f_{X_1, \dots, X_{n+1}}(x_1, \dots, x_{n+1}) &= \int_0^\infty f_{X_1, \dots, X_{n+1}|V=v}(x_1, \dots, x_n) f_{GA}(v; \beta, \beta) dv \\
&= \frac{\alpha^{n_1+n_2}(\lambda_1 + \lambda_2)^{n_1+n_2} \lambda_1^{n_3}}{(2\lambda_1 + \lambda_2)^{n_3}} \times \left(\prod_{i \in I_1 \cup I_2} x_{i+1}^{\alpha-1} \right) \times \frac{\Gamma(n_1 + n_2 + \beta)}{\Gamma(\beta)} \times \\
&\quad \frac{\beta^\beta}{(\beta + (\lambda_1 + \lambda_2)A_1(\mathbf{x}; \alpha) + \lambda_1 A_2(\mathbf{x}; \alpha))^{n_1+n_2+\beta}}. \tag{14}
\end{aligned}$$

The following results will be used in the implementation of the EM algorithm and they can be obtained after some calculations.

$$V|\{(X_1, \dots, X_{n+1})\} \sim \text{Gamma}(n_1 + n_2 + \beta, \beta + (\lambda_1 + \lambda_2)A_1(\mathbf{x}; \alpha) + \lambda_1 A_2(\mathbf{x}; \alpha)) \tag{15}$$

Hence,

$$E(V|\{(X_1, \dots, X_{n+1})\}) = \frac{n_1 + n_2 + \beta}{\beta + (\lambda_1 + \lambda_2)A_1(\mathbf{x}; \alpha) + \lambda_1 A_2(\mathbf{x}; \alpha)}. \tag{16}$$

Another important concept in any stationary time series is 'stopping time'. It has been discussed quite extensively in the time series literature. Christensen [2] discussed about optimal stopping time for autoregressive processes, Novikov and Shiryaev [10] discussed about the optimal stopping time for processes with independent increment, Kundu [6] discussed for the proportional hazard class processes. Let us define the stopping time N for the UWFP $\{X_n\}$, for a given L as follows;

$$\{N = n\} \Leftrightarrow \{X_1 > L, \dots, X_{n-1} > L, X_n < L\}.$$

Since, the event $\{N = n\}$ is independent of $X_{n+1}, X_{n+2}, \dots$, for all $n \geq 1$, therefore N is a stopping. Let us compute the probability mass function of N , for $n = 1, 2, \dots$ becomes

$$\begin{aligned}
P(N = n) &= P(X_1 > L, \dots, X_{n-1} > L, X_n < L) \\
&= P(X_1 > L, \dots, X_{n-1} > L) - P(X_1 > L, \dots, X_{n-1} > L, X_n > L) \\
&= (1 + n\lambda_1 L^\alpha + (n-1)\lambda_2 L^\alpha)^{-\beta} - (1 + (n+1)\lambda_1 L^\alpha + n\lambda_2 L^\alpha)^{-\beta}.
\end{aligned}$$

The k -th moment of N exists for $\beta > k + 1$ and it is

$$\begin{aligned}
E(N^k) &= \sum_{n=1}^{\infty} P(N = n) n^k \\
&= \sum_{n=1}^{\infty} n^k (1 + n\lambda_1 L^\alpha + (n-1)\lambda_2 L^\alpha)^{-\beta} - \sum_{n=1}^{\infty} n^k (1 + (n+1)\lambda_1 L^\alpha + n\lambda_2 L^\alpha)^{-\beta}.
\end{aligned}$$

Using the Hurwitz zeta function $\zeta(s, q) = \sum_{m=0}^{\infty} (m+q)^{-s}$, it can be seen that $E(N)$ exists for $\beta > 2$, $E(N^2)$ exists for $\beta > 3$ and for $t = L^\alpha$ they are as follows;

$$\begin{aligned} E(N) &= \frac{1}{(t(\lambda_1 + \lambda_2))^\beta} \zeta\left(\beta, 1 + \frac{1 - \lambda_2 t}{(\lambda_1 + \lambda_2)t}\right) \\ E(N^2) &= \frac{1}{(t(\lambda_1 + \lambda_2))^\beta} \left[2\zeta\left(\beta - 1, 1 + \frac{1 - \lambda_2 t}{(\lambda_1 + \lambda_2)t}\right) - \left(2\frac{1 - \lambda_2 t}{(\lambda_1 + \lambda_2)t} + 1\right) \zeta\left(\beta, 1 + \frac{1 - \lambda_2 t}{(\lambda_1 + \lambda_2)t}\right) \right]. \end{aligned}$$

4 INFERENCE

Let $\mathcal{D} = \{x_1, \dots, x_{n+1}\}$ be a sample from the UWFP($\alpha, \lambda_1, \lambda_2, \beta$). Based on $\mathcal{D}$ we would like to obtain the MLEs of the unknown parameters. The log-likelihood function of the observed sample $\mathcal{D}$ without the additive constant is

$$\begin{aligned} l(\alpha, \lambda_1, \lambda_2, \beta | \mathcal{D}) &= (n_1 + n_2) \ln \alpha + (n_1 + n_2) \ln(\lambda_1 + \lambda_2) + n_3 \ln \lambda_1 - n_3 \ln(2\lambda_1 + \lambda_2) + \\ &\quad \beta \ln \beta + \alpha \sum_{i \in I_1 \cup I_2} \ln x_{i+1} + \ln \Gamma(n_1 + n_2 + \beta) - \ln \Gamma(\beta) - \\ &\quad (n_1 + n_2 + \beta) \ln(\beta + (\lambda_1 + \lambda_2)A_1(\mathbf{x}; \alpha) + \lambda_1 A_2(\mathbf{x}; \alpha)) \end{aligned} \quad (17)$$

The maximization of (17) involves solving a four dimensional optimization problem. To avoid solving four dimensional optimization problem, we use EM algorithm combining with profile maximization approach. The basic idea is as follows. First we fixed β , and then we use EM algorithm to obtain the MLEs of α , λ_1 and λ_2 as a function of β . It involves solving a one-dimensional optimization problem at each M-step of the EM algorithm. Then using the profile likelihood method we obtain the MLE of β .

To implement EM, it is assumed that the v is missing here. Hence, the complete observation is $\mathcal{D}_c = (\mathbf{x}, v)$. Based on the complete observation, for fixed β , the log-likelihood function without the additive constant becomes;

$$\begin{aligned} l_c(\alpha, \lambda_1, \lambda_2 | \beta, \mathcal{D}_c) &= (n_1 + n_2) \ln \alpha + (n_1 + n_2) \ln(\lambda_1 + \lambda_2) + n_3 \ln \lambda_1 - n_3 \ln(2\lambda_1 + \lambda_2) + \\ &\quad \alpha \sum_{i \in I_1 \cup I_2} \ln x_{i+1} - v((\lambda_1 + \lambda_2)A_1(\mathbf{x}; \alpha) + \lambda_1 A_2(\mathbf{x}; \alpha)). \end{aligned} \quad (18)$$

Suppose at the k -th stage the estimates of α , λ_1 and λ_2 are $\alpha^{(k)}$, $\lambda_1^{(k)}$ and $\lambda_2^{(k)}$, respectively. We further denote $a^{(k)} = E(V | \{(X_1, \dots, X_{n+1})\})$, where $a^{(k)}$ is obtained from (16) by replacing

α , λ_1 and λ_2 with $\alpha^{(k)}$, $\lambda_1^{(k)}$ and $\lambda_2^{(k)}$, respectively. At the k -th stage of the EM algorithm, the 'pseudo' log-likelihood function without the additive constant can be obtained as

$$\begin{aligned} l_s^{(k)}(\alpha, \lambda_1, \lambda_2 | \mathcal{D}) &= (n_1 + n_2) \ln \alpha + (n_1 + n_2) \ln(\lambda_1 + \lambda_2) + n_3 \ln \lambda_1 - n_3 \ln(2\lambda_1 + \lambda_2) + \\ &\quad \alpha \sum_{i \in I_1 \cup I_2} \ln x_{i+1} - a^{(k)}((\lambda_1 + \lambda_2)A_1(\mathbf{x}; \alpha) + \lambda_1 A_2(\mathbf{x}; \alpha)). \end{aligned} \quad (19)$$

At the M-step the $\alpha^{(k+1)}$, $\lambda_1^{(k+1)}$ and $\lambda_2^{(k+1)}$ can be obtained by maximizing (19) with respect to α , λ_1 and λ_2 , respectively. For fixed α , if $\lambda_1^{(k+1)}(\alpha)$ and $\lambda_2^{(k+1)}(\alpha)$ maximize (19), then

$$\begin{aligned} \lambda_1^{(k+1)}(\alpha) &= \frac{n_3}{a^{(k)}A_2(\mathbf{x}; \alpha)} - \frac{n_1 + n_2}{a^{(k)}(A_1(\mathbf{x}; \alpha) + A_2(\mathbf{x}; \alpha))} \quad \text{and} \\ \lambda_2^{(k+1)}(\alpha) &= \frac{2(n_1 + n_2)}{a^{(k)}(A_1(\mathbf{x}; \alpha) + A_2(\mathbf{x}; \alpha))} - \frac{n_3}{a^{(k)}A_2(\mathbf{x}; \alpha)}. \end{aligned}$$

Obtain $\alpha^{(k+1)}$ by the argument maximum of

$$\begin{aligned} g_1^{(k)}(\alpha) &= (n_1 + n_2) \ln \alpha + (n_1 + n_2) \ln(\lambda_1^{(k+1)}(\alpha) + \lambda_2^{(k+1)}(\alpha)) + n_3 \ln(\lambda_1^{(k+1)}(\alpha)) + \\ &\quad \alpha \sum_{i \in I_1 \cup I_2} \ln x_{i+1} - n_3 \ln(2\lambda_1^{(k+1)}(\alpha) + \lambda_2^{(k+1)}(\alpha)) - \\ &\quad a^{(k)}((\lambda_1^{(k+1)}(\alpha) + \lambda_2^{(k+1)}(\alpha))A_1(\mathbf{x}; \alpha) + \lambda_1^{(k+1)}(\alpha)A_2(\mathbf{x}; \alpha)). \end{aligned}$$

Once $\alpha^{(k+1)}$ is obtained $\lambda_1^{(k+1)}$ and $\lambda_2^{(k+1)}$ can be obtained as $\lambda_1^{(k+1)} = \lambda_1^{(k+1)}(\alpha^{(k+1)})$ and $\lambda_2^{(k+1)} = \lambda_2^{(k+1)}(\alpha^{(k+1)})$, respectively. Continue the process until the convergence is achieved. For fixed β , Let us denote these estimates as $\hat{\alpha}(\beta)$, $\hat{\lambda}_1(\beta)$ and $\hat{\lambda}_2(\beta)$. Now maximize the profile log-likelihood function $l(\hat{\alpha}(\beta), \hat{\lambda}_1(\beta), \hat{\lambda}_2(\beta), \beta | \mathcal{D})$ as defined in (17) to compute the MLE of β , say $\hat{\beta}$. Once $\hat{\beta}$ is obtained, the MLEs of α , λ_1 and λ_2 can be obtained as $\hat{\alpha}(\hat{\beta})$, $\hat{\lambda}_1(\hat{\beta})$ and $\hat{\lambda}_2(\hat{\beta})$, respectively.

5 SIMULATION RESULTS

We have performed some simulation experiments to see how the proposed method works in practice. We have considered four different sample sizes namely 25, 50, 75 and 100 and four different sets of parameter values; (i) $\alpha = 1.0$, $\lambda_1 = 0.1$, $\lambda_2 = 0.2$ and $\beta = 1.0$ (ii) $\alpha = 2.0$, $\lambda_1 = 0.1$, $\lambda_2 = 0.2$ and $\beta = 1.0$ (iii) $\alpha = 2.0$, $\lambda_1 = 0.1$, $\lambda_2 = 0.2$ and $\beta = 2.0$ (iv) α

$= 2.0$, $\lambda_1 = 0.5$, $\lambda_2 = 0.5$ and $\beta = 2.0$. In all the cases we have computed the maximum likelihood estimators of the unknown parameters based on the EM algorithm and report the average estimators and the corresponding mean squared errors based on 1000 replications. We stop the iterations whenever the relative difference of the log-likelihood function of the two consecutive iterations is less than 10^{-6} . The results are reported in Tables 1 to 4. In all the cases the EM algorithm stops in less than 15 iterations. In each case we have also computed the confidence intervals of the unknown parameters based on parametric bootstrap approach. It is based on 1000 replications. We have reported the average lengths and the coverage percentages. It is observed that as the sample size increases the biases, MSEs and the lengths of the confidence intervals decrease. In all the cases the coverage percentages are close to the nominal value.

Table 1: Average estimates and the associated mean squared errors. The average lengths of the parametric bootstrap confidence intervals (95%) and the coverage percentages are also reported in each box below the average estimates and MSEs. Here $\alpha = 1.0$, $\lambda_1 = 0.1$ and $\lambda_2 = 0.2$ and $\beta = 1.0$.

n	α	λ_1	λ_2	β
25	1.2513	0.1029	0.2132	1.2132
	(0.0391)	(0.00013)	(0.00042)	(0.2413)
	(0.0981)	(0.02235)	(0.06123)	(0.8845)
	(0.931)	(0.935)	(0.093)	(0.928)
50	1.2171	0.1025	0.2100	1.1987
	(0.0289)	(0.00009)	(0.00028)	(0.2014)
	(0.0912)	(0.01987)	(0.04234)	(0.7851)
	(0.938)	(0.941)	(0.942)	(0.939)
75	1.1913	0.1023	0.2089	1.1411
	(0.0161)	(0.00008)	(0.00021)	(0.1643)
	(0.0712)	(0.01156)	(0.03978)	(0.5867)
	(0.942)	(0.943)	(0.944)	(0.941)
100	1.0595	0.1020	0.2033	1.0312
	(0.0081)	(0.00002)	(0.00014)	(0.1024)
	(0.0514)	(0.00846)	(0.01657)	(0.3423)
	(0.948)	(0.949)	(0.951)	(0.946)

Table 2: Average estimates and the associated mean squared errors based on 1000 replications. The average lengths of the parametric bootstrap confidence intervals (95%) and the coverage percentages are also reported in each box below the average estimates and MSEs. Here $\alpha = 2.0$, $\lambda_1 = 0.1$, $\lambda_2 = 0.2$ and $\beta = 1.0$.

n	α	λ_1	λ_2	β
25	1.9082	0.1056	0.2239	1.3125
	(0.0469)	(0.00028)	(0.00173)	(0.3315)
	(0.0994)	(0.03365)	(0.09451)	(0.9564)
	(0.932)	(0.938)	(0.941)	(0.939)
50	1.9163	0.1031	0.2149	1.2567
	(0.0414)	(0.00017)	(0.00087)	(0.2867)
	(0.0687)	(0.01756)	(0.05378)	(0.5123)
	(0.938)	(0.943)	(0.948)	(0.944)
75	1.9884	0.1027	0.2123	1.1786
	(0.0311)	(0.00011)	(0.00063)	(0.1987)
	(0.0511)	(0.00954)	(0.03267)	(0.3128)
	(0.941)	(0.948)	(0.944)	(0.947)
100	1.9990	0.1021	0.2097	1.1156
	(0.0092)	(0.00003)	(0.00047)	(0.1423)
	(0.0237)	(0.00457)	(0.01318)	(0.1732)
	(0.951)	(0.949)	(0.952)	(0.951)

Table 3: Average estimates and the associated mean squared errors based on 1000 replications. The average lengths of the parametric bootstrap confidence intervals (95%) and the coverage percentages are also reported in each box below the average estimates and MSEs. Here $\alpha = 2.0$, $\lambda_1 = 0.1$, $\lambda_2 = 0.2$ and $\beta = 2.0$.

n	α	λ_1	λ_2	β
25	1.8864	0.0984	0.1934	2.3541
	(0.0467)	(0.00025)	(0.00033)	(0.4014)
	(0.0931)	(0.03267)	(0.03765)	(0.7834)
	(0.929)	(0.934)	(0.931)	(0.928)
50	1.8997	0.0990	0.1964	2.2987
	(0.0418)	(0.00016)	(0.00013)	(0.3656)
	(0.0714)	(0.01132)	(0.01023)	(0.5753)
	(0.938)	(0.941)	(0.939)	(0.944)
75	1.9845	0.0992	0.1972	2.2017
	(0.0382)	(0.00012)	(0.00008)	(0.2987)
	(0.0612)	(0.00545)	(0.00521)	(0.4154)
	(0.941)	(0.943)	(0.944)	(0.951)
100	1.9952	0.0996	0.1976	2.1765
	(0.0219)	(0.00001)	(0.00006)	(0.1656)
	(0.0398)	(0.00387)	(0.00267)	(0.2134)
	(0.949)	(0.953)	(0.949)	(0.948)

Table 4: Average estimates and the associated mean squared errors based on 1000 replications. The average lengths of the parametric bootstrap confidence intervals (95%) and the coverage percentages are also reported in each box below the average estimates and MSEs. Here $\alpha = 2.0$, $\lambda_1 = 0.5$, $\lambda_2 = 0.5$ and $\beta = 2.0$.

n	α	λ_1	λ_2	β
25	2.1656	0.4660	0.4663	2.3314
	(0.0558)	(0.0035)	(0.0035)	(0.3878)
	(0.0992)	(0.1160)	(0.1221)	(0.9984)
	(0.939)	(0.941)	(0.942)	(0.941)
50	2.1561	0.4776	0.4776	2.3011
	(0.0514)	(0.0018)	(0.0019)	(0.3564)
	(0.0910)	(0.0698)	(0.0712)	(0.8922)
	(0.944)	(0.942)	(0.944)	(0.945)
75	2.1357	0.4877	0.4824	2.1887
	(0.0417)	(0.0014)	(0.0014)	(0.2765)
	(0.0723)	(0.0445)	(0.0412)	(0.6734)
	(0.948)	(0.947)	(0.951)	(0.949)
100	2.1096	0.4850	0.4848	2.1541
	(0.0225)	(0.0011)	(0.0012)	(0.1553)
	(0.0423)	(0.0298)	(0.0310)	(0.3798)
	(0.951)	(0.952)	(0.949)	(0.951)

6 GOODNESS OF FIT FOR THE UNIVARIATE PROCESS

We have proposed the following test to verify whether the proposed model can be used for a given data set or not. The test was originally suggested by Kundu [5] and it is based on order statistics approach. The problem can be mathematically formulated as follows. Suppose $\{x_1, \dots, x_n\}$ is sample from a stationary sequence $\{X_1, \dots, X_n\}$, we want to test the following hypothesis:

$$H_0 : \{X_1, \dots, X_n\} \sim \text{UWFP}(\alpha, \beta, \lambda_1, \lambda_2).$$

Let us denote $X_{1:n} < \dots < X_{n:n}$ as the ordered $\{X_1, \dots, X_n\}$ and the corresponding sample version of $\{x_1, \dots, x_n\}$ is $\{x_{1:n} < \dots < x_{n:n}\}$. Further, let us denote $a_1 = E_{H_0}(X_{1:n}) \dots, a_n = E_{H_0}(X_{n:n})$ as their ordered expected values under H_0 . Clearly, $a_1, \dots, a_n$ depend on α, β, λ_1 and λ_2 , but we do not make it explicit for brevity. We use the following statistic for the goodness of fit test:

$$W_n = \max_{1 \leq i \leq n} |X_{i:n} - a_i|.$$

It is expected that W_n should be small under H_0 , and large otherwise. We use the following test criterion for a given level of significance $0 < \gamma < 1$

$$\text{Reject } H_0 \text{ if } W_n > c_n(\gamma),$$

where $c_n(\gamma)$ is such that

$$P_{H_0}(W_n > c_n(\gamma)) = \gamma.$$

Clearly, $c_n(\gamma)$ also depends on the true parameter values, but we do not make it explicit here also. It is difficult to compute $c_n(\gamma)$ even asymptotically. Hence, we suggest to use parametric bootstrap technique to approximate $c_n(\gamma)$ from the given sample $\{x_1, \dots, x_n\}$.

Algorithm:

Step 1: Obtain $\hat{\alpha}, \hat{\beta}, \hat{\lambda}_1$ and $\hat{\lambda}_2$, the MLEs of α, β, λ_1 and λ_2 , respectively, based on the sample $\{x_1, \dots, x_n\}$.

- Step 2: Generate a sample of size n from a UWFP($\hat{\alpha}, \hat{\beta}, \hat{\lambda}_1, \hat{\lambda}_2$) and order them. Let us denote them as $(x_1^1, \dots, x_{n:n}^1)$. Repeat this procedure B times, and obtain $(x_{1:n}^b, \dots, x_{n:n}^b); b = 1, \dots, B$.

- Step 3: Obtain estimates of $a_1, \dots, a_n$ as

$$\hat{a}_i = \frac{1}{B} \sum_{b=1}^B x_{i:n}^b; \quad i = 1, \dots, n.$$

- Step 4: Compute

$$w^b = \max_{1 \leq i \leq n} \{|x_{i:n}^b - \hat{a}_i|\}; \quad b = 1, \dots, B.$$

- Step 5: Order $\{w_1, \dots, w_B\}$ as $w^{(1)} < \dots < w^{(B)}$, then $\hat{c}_n(\gamma) = w^{[100(1-\gamma)]}$ is an estimate of $c_n(\gamma)$.

Hence, if $w_n = \max_{1 \leq i \leq n} |x_{i:n} - \hat{a}_i| > \hat{c}_n(\gamma)$, we reject the null hypothesis with $\gamma\%$ level of confidence, otherwise we accept the null hypothesis.

7 DATA ANALYSIS

In this section we have provided the analyses of two data sets, namely the gold price data set in the Indian market and the Indian Rupees and US Dollars exchange rate data set. Both the data are obtained on the same period of 60 days starting from March 01, 2023.

7.1 GOLD PRICE DATA SET:

The gold price data has been plotted in Figure 1. Some of the basic statistics are as follows: The minimum and maximum values are 4230.02 and 4515.21, respectively. The 25th percentile, median and 75th percentile values are 4433.57, 4486.16 and 4515.21, respectively. The augmented Dickey-Fuller test suggests that it is from a stationary process.

The auto-correlation function (ACF) and partial auto-correlation function (PACF) have been plotted in Figure 2.

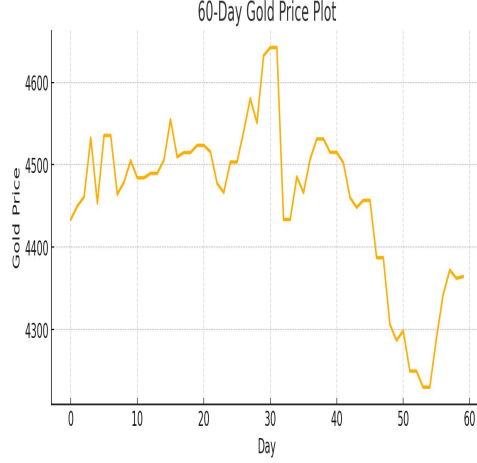


Figure 1: (a) The plot of the gold price data.

Now to compute the MLEs of the unknown parameters, we have subtracted 4225 from all the data points and we have used the proposed EM algorithm as it has been proposed in this paper. We have obtained the MLEs of the unknown parameters, and at the last stage of the EM algorithm based on the method proposed by Louis [8] we have obtained the 95% confidence intervals of the unknown parameters. The MLEs and the corresponding confidence intervals are reported in brackets: $\hat{\alpha} = 2.1518(0.7545)$, $\hat{\beta} = 1.6956(0.4562)$, $\hat{\lambda}_1 = 1.6376 \times 10^{-2}(0.9167 \times 10^{-3})$, $\hat{\lambda}_2 = 3.0394 \times 10^{-2}(1.1165 \times 10^{-3})$, The corresponding log-likelihood value becomes -324.7195. We have performed the goodness of fit test as proposed in Section 6. Based on 500 bootstrap replications, the p -value becomes 0.344, hence we can say that the null hypothesis that the data has been obtained from a UWFP cannot be rejected.

7.2 INDIAN RUPEES AND US DOLLARS EXCHANGE RATE DATA

In Figure 3 we have plotted the exchange rate data. Some of the basic statistics of the exchange rate data are as follows: The minimum, maximum and range are 81.81, 82.77 and 0.96, respectively. The 25th percentile, median and the 75th percentile are 82.1910, 82.0750 and 82.4025, respectively. In this case also the ADF test suggests that it is obtained from a stationary process. The ACF and PACF plots of the exchange rate data are provided in Figure 4 Before computing the MLEs of the unknown parameters using the EM algorithm

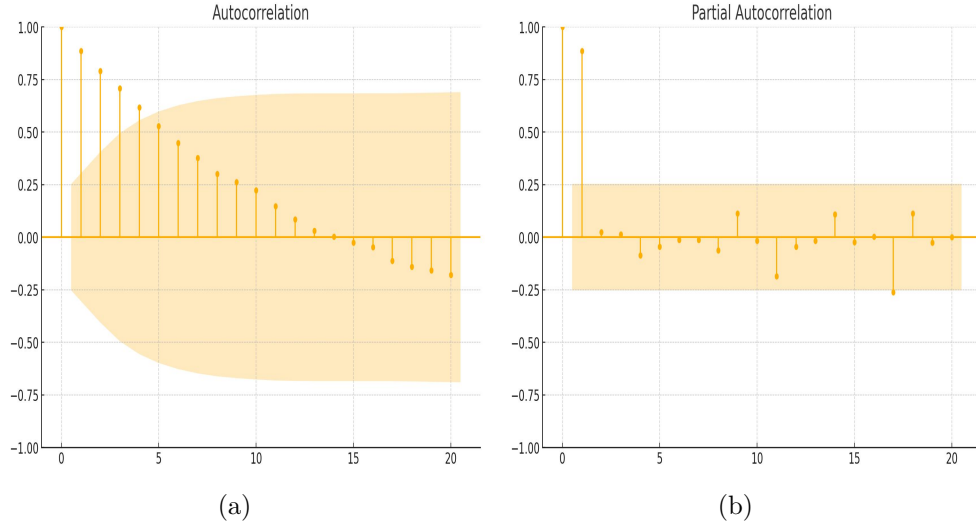


Figure 2: (a) The autocorrelation function and (b) the partial autocorrelation function of the gold price data.

proposed here, we have subtracted 82.5 from each data set. The MLEs and the associated 95% confidence intervals (reported in brackets) are as follows: $\hat{\alpha} = 2.7800(0.6754)$, $\hat{\beta} = 1.7564(0.4321)$, $\hat{\lambda}_1 = 0.2075(0.0134)$, $\hat{\lambda}_2 = 1.2274(0.4564)$. The corresponding log-likelihood value becomes -37.3628. In this case also we have performed the goodness of fit test. Based on 500 bootstrap replications, the p -value becomes 0.216. Hence in this case also the null hypothesis that the data has been obtained from a UWFP cannot be rejected.

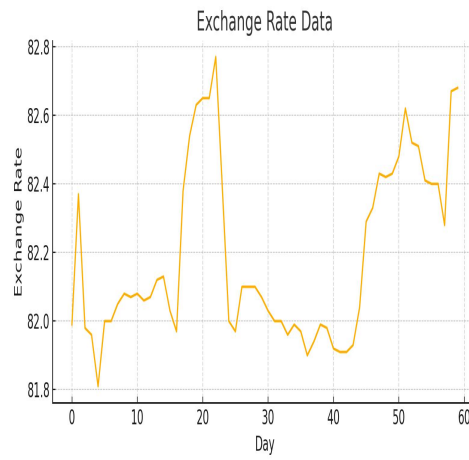


Figure 3: (a) The plot of the exchange rate data.

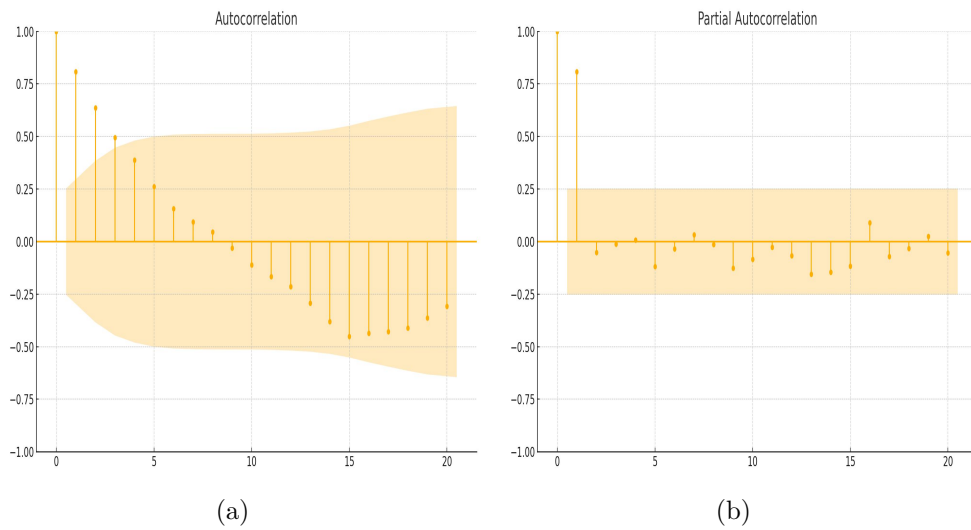


Figure 4: (a) The autocorrelation function and (b) the partial autocorrelation function of the exchange rate data.

8 BIVARIATE WEIBULL-FRAILTY PROCESS

8.1 MODEL FORMULATION AND MLES

In this section we will define the bivariate Weibull-frailty process (BWFP) process extending the UWFP process to the bivariate case. First consider the sequences of random variables $\{U_n; n \geq 1\}$, $\{W_n; n \geq 1\}$ and V same as defined before. Let us consider two more sequences of random variables $\{S_n; n \geq 1\}$ and $\{T_n; n \geq 1\}$ where S_n 's are i.i.d., $S_n \sim \text{WE}(\delta, \gamma_1)$, T_n 's are i.i.d., $T_n \sim \text{WE}(\delta, \gamma_2)$, further $\{S_n; n \geq 1\}$ and $\{T_n; n \geq 1\}$ are independently distributed. Now we define the bivariate process $\{(X_n, Y_n); n \geq 1\}$ as follows;

$$X_n = \min \left\{ \frac{U_n}{V^{1/\alpha}}, \frac{U_{n+1}}{V^{1/\alpha}}, \frac{W_n}{V^{1/\alpha}} \right\}; \quad \text{and} \quad Y_n = \min \left\{ \frac{S_n}{V^{1/\delta}}, \frac{S_{n+1}}{V^{1/\delta}}, \frac{T_n}{V^{1/\delta}} \right\}; \quad n \geq 1. \quad (20)$$

We call this process as the BWFP with parameters $(\alpha, \delta, \lambda_1, \lambda_2, \gamma_1, \gamma_2, \beta)$. It will be denoted by $\text{BWFP}(\alpha, \delta, \lambda_1, \lambda_2, \gamma_1, \gamma_2, \beta)$. The generation from a BWFP is also quite straight forward, and it follows from the definition itself.

Theorem 3: If $\{X_n, Y_n\}$ is a $\text{BWFP}(\alpha, \delta, \lambda_1, \lambda_2, \gamma_1, \gamma_2, \beta)$, then

- (i) $\{X_n, Y_n\}$ is a stationary process.
- (ii) $\{X_n\}$ is a $\text{UWFP}(\alpha, \lambda_1, \lambda_2, \beta)$
- (iii) $\{Y_n\}$ is a $\text{UWFP}(\delta, \gamma_1, \gamma_2, \beta)$
- (iv) If $\alpha = \delta$ and $M_n = \min\{X_n, Y_n\}$, then $\{M_n\}$ is a $\text{UWFP}(\alpha, \lambda_1 + \gamma_1, \lambda_2 + \gamma_2, \beta)$
- (v) (X_n, Y_n) is absolutely continuous and the joint survival function of (X_n, Y_n) for $x > 0$ and $y > 0$, is

$$S_{X_n, Y_n}(x, y) = P(X_n > x, Y_n > y) = (1 + (2\theta_1 + \theta_2)x^\alpha + (2\xi_1 + \xi_2)y^\delta)^{-\beta},$$

where $\theta_1 = \frac{\lambda_1}{\beta}$, $\theta_2 = \frac{\lambda_2}{\beta}$, $\xi_1 = \frac{\gamma_1}{\beta}$ and $\xi_2 = \frac{\gamma_2}{\beta}$.

- (vi) If we use the following standard notations: $Z_n = (X_n, Y_n)^\top$, $z_n = (x_n, y_n)^\top$ and $Z_n > (=)z_n$ means $X_n > (=)x_n$ and $Y_n > (=)y_n$, then for $z_n > 0$, $z_{n+m} > 0$, and $v > 0$,

$$S_{Z_n, Z_{n+m}|V=v}(z_n, z_{n+m}) = P(Z_n > z_n | V = v)P(Z_{n+m} > z_{n+m} | V = v)$$

(vii) For $z_1 > 0, \dots, z_n > 0, z_{n+m} > 0$ and $v > 0$,

$$P(Z_n > z_n | Z_{n-1} = z_{n-1}, \dots, Z_1 = z_1, V = v) = P(Z_n \geq z_n | Z_{n-1} = z_{n-1}, V = v).$$

Proof: See in the Appendix. ■

It may be mentioned that the construction (20) is a very natural construction to bring dependence between X_n and Y_n . It is interesting to see that the random vector (X_n, Y_n) which is obtained using the construction (20) has the following survival copula

$$\overline{C}(u, v) = \frac{1}{(1 + u^{-1/\beta} + v^{-1/\beta})^\beta}.$$

It is the famous Clayton copula. Several properties of this copula is well known in the literature. Let us denote the bivariate vector $Z_n = (X_n, Y_n)^\top$, for $n = 1, 2, \dots$. First, we would like to obtain the joint PDF of $Z_1, \dots, Z_{n+1}$ and it will be useful to develop the likelihood inference of the unknown parameters. The joint PDF of $\{(Z_n, Z_{n+1}) | V = v\}$ is

$$f_{Z_n, Z_{n+1} | V=v}(z_n, z_{n+1}) = f_{X_n, X_{n+1} | V=v}(x_n, x_{n+1}) f_{Y_n, Y_{n+1} | V=v}(y_n, y_{n+1}) \quad (21)$$

If $\{X_n\}$ is a UWFP($\alpha, \lambda_1, \lambda_2, \beta$), then $f_{X_n, X_{n+1} | V=v}(x_n, x_{n+1})$ is provided in (11), similarly, when $\{Y_n\}$ is a UWFP($\delta, \gamma_1, \gamma_2, \beta$), then $f_{Y_n, Y_{n+1} | V=v}(y_n, y_{n+1})$ can be obtained. Now using (vii) of Theorem 3,

$$\begin{aligned} f_{Z_1, \dots, Z_{n+1} | V=v}(z_1, \dots, z_{n+1}) &= f_{Z_1 | V=v}(z_1) f_{Z_2 | Z_1=z_1, V=v}(z_2) \dots f_{Z_{n+1} | Z_1=z_1, \dots, Z_n=z_n, V=v}(z_{n+1}) \\ &= f_{Z_1 | V=v}(z_1) f_{Z_2 | Z_1=z_1, V=v}(z_2) \dots f_{Z_{n+1} | Z_n=z_n, V=v}(z_{n+1}) \\ &= \frac{\prod_{i=1}^n f_{Z_i, Z_{i+1} | V=v}(z_i, z_{i+1})}{\prod_{i=2}^n f_{Z_i | V=v}(z_i)} \\ &= \left\{ \frac{\prod_{i=1}^n f_{X_i, X_{i+1} | V=v}(x_i, x_{i+1})}{\prod_{i=2}^n f_{X_i | V=v}(x_i)} \right\} \left\{ \frac{\prod_{i=1}^n f_{Y_i, Y_{i+1} | V=v}(y_i, y_{i+1})}{\prod_{i=2}^n f_{Y_i | V=v}(y_i)} \right\}. \end{aligned}$$

Let us introduce the following notations: $J_1 = \{i : i \in I, y_i < y_{i+1}\}$, $J_2 = \{i : i \in I, y_i > y_{i+1}\}$, $J_3 = \{i : i \in I, y_i = y_{i+1}\}$, m_1, m_2 and m_3 denote the number of elements in J_1 , J_2 and J_3 , respectively. Here I is same as defined before. Similarly, we define $\mathbf{y} = (y_1, \dots, y_{n+1})$,

$$B_1(\mathbf{y}; \delta) = \sum_{i=1}^n y_{i+1}^\delta + y_1^\delta \quad \text{and} \quad B_2(\mathbf{y}; \delta) = \sum_{i \in J_1} (y_{i+1}^\delta - y_i^\delta) + y_1^\delta.$$

The joint PDF of $Z_1, \dots, Z_{n+1}$ becomes

$$\begin{aligned}
f_{Z_1, \dots, Z_{n+1}}(z_1, \dots, z_{n+1}) &= \int_0^\infty f_{Z_1, \dots, Z_{n+1}|V=v}(z_1, \dots, z_{n+1}) f_{GA}(v; \beta, \beta) dv \\
&= \frac{\alpha^{n_1+n_2} (\lambda_1 + \lambda_2)^{n_1+n_2} \lambda_1^{n_3}}{(2\lambda_1 + \lambda_2)^{n_3}} \times \frac{\delta^{m_1+m_2} (\gamma_1 + \gamma_2)^{m_1+m_2} \gamma_1^{m_3}}{(2\gamma_1 + \gamma_2)^{m_3}} \times \left(\prod_{i \in I_1 \cup I_2} x_{i+1}^{\alpha-1} \right) \\
&\quad \times \left(\prod_{i \in J_1 \cup J_2} y_{i+1}^{\delta-1} \right) \times \frac{\Gamma(n_1 + n_2 + m_1 + m_2 + \beta)}{\Gamma(\beta)} \times g_1(\boldsymbol{\Theta}; \mathbf{x}, \mathbf{y})
\end{aligned}$$

where $\boldsymbol{\Theta} = (\alpha, \delta, \lambda_1, \lambda_2, \gamma_1, \gamma_2, \beta)$, and if

$$C(\alpha, \delta, \lambda_1, \lambda_2, \gamma_1, \gamma_2) = (\lambda_1 + \lambda_2)A_1(\mathbf{x}; \alpha) + \lambda_1 A_2(\mathbf{x}; \alpha) + (\gamma_1 + \gamma_2)B_1(\mathbf{y}; \delta) + \gamma_1 B_2(\mathbf{y}; \delta),$$

then

$$g_1(\boldsymbol{\Theta}; \mathbf{x}, \mathbf{y}) = \frac{\beta^\beta}{(\beta + C(\alpha, \delta, \lambda_1, \lambda_2, \gamma_1, \gamma_2))^{n_1+n_2+m_1+m_2+\beta}}.$$

Therefore, the log-likelihood function without the additive constant becomes

$$\begin{aligned}
l(\boldsymbol{\Theta}|\mathbf{x}, \mathbf{y}) &= (n_1 + n_2) \ln \alpha + (n_1 + n_2) \ln(\lambda_1 + \lambda_2) + n_3 \ln \lambda_1 - n_3 \ln(2\lambda_1 + \lambda_2) + \alpha \sum_{i \in I_1 \cup I_2} x_{i+1} + \\
&\quad (m_1 + m_2) \ln \delta + (m_1 + m_2) \ln(\gamma_1 + \gamma_2) + m_3 \ln \gamma_1 - m_3 \ln(2\gamma_1 + \gamma_2) + \delta \sum_{i \in I_1 \cup I_2} y_{i+1} + \\
&\quad \ln \Gamma(n_1 + n_2 + m_1 + m_2 + \beta) - \ln \Gamma(\beta) + \ln g_1(\boldsymbol{\Theta}; \mathbf{x}, \mathbf{y}).
\end{aligned} \tag{22}$$

Hence, it is clear that the for the bivariate case the MLEs of the unknown parameters can be obtained by solving a seven dimensional optimization problem. In this case also we use EM algorithm and profile likelihood method to reduce the seven dimensional optimization problem. The idea is same as before. First it is assumed that β is fixed, and for fixed β we obtain the MLEs of the other parameters based on the EM algorithm. It is observed that the EM algorithm involves in solving two one dimensional optimization problem at each M-step. Hence, it reduces the computational burden significantly. To implement the EM algorithm, it is assumed that the observation v is missing. We need the following quantities for implementing the EM algorithm.

$$V|\{(Z_1, \dots, Z_{n+1})\} \sim \text{Gamma}(n_1 + n_2 + \beta, C(\alpha, \delta, \lambda_1, \lambda_2, \gamma_1, \gamma_2))$$

Hence,

$$E(V|\{(Z_1, \dots, Z_{n+1})\}) = \frac{n_1 + n_2 + \beta}{\beta + C(\alpha, \delta, \lambda_1, \lambda_2, \gamma_1, \gamma_2)}. \tag{23}$$

Here also v is missing. Hence, the complete observation is $\mathcal{D}_{cb} = (\mathbf{x}, \mathbf{y}, v)$. Based on the complete observation, for fixed β , the log-likelihood function without the additive constant becomes;

$$\begin{aligned}
l_{cb}(\boldsymbol{\Theta}_1 | \mathcal{D}_{cb}, \beta) &= (n_1 + n_2) \ln \alpha + (n_1 + n_2) \ln(\lambda_1 + \lambda_2) + n_3 \ln \lambda_1 - n_3 \ln(2\lambda_1 + \lambda_2) + \\
&\quad \alpha \sum_{i \in I_1 \cup I_2} \ln x_{i+1} - v((\lambda_1 + \lambda_2)A_1(\mathbf{x}; \alpha) + \lambda_1 A_2(\mathbf{x}; \alpha)) + \\
&\quad (m_1 + m_2) \ln \delta + (m_1 + m_2) \ln(\gamma_1 + \gamma_2) + m_3 \ln \gamma_1 - m_3 \ln(2\gamma_1 + \gamma_2) + \\
&\quad \delta \sum_{i \in J_1 \cup J_2} \ln y_{i+1} - v((\gamma_1 + \gamma_2)B_1(\mathbf{y}; \delta) + \gamma_1 B_2(\mathbf{y}; \delta)), \tag{24}
\end{aligned}$$

where $\boldsymbol{\Theta}_1 = (\alpha, \lambda_1, \lambda_2, \delta, \gamma_1, \gamma_2)$. Let us use the following notations: suppose at the k -th stage the estimates of $\alpha, \delta, \lambda_1, \lambda_2, \gamma_1$ and γ_2 are $\alpha^{(k)}, \delta^{(k)}, \lambda_1^{(k)}, \lambda_2^{(k)}, \gamma_1^{(k)}$ and $\gamma_2^{(k)}$, respectively. We further denote $c^{(k)} = E(V | \{(Z_1, \dots, Z_{n+1})\})$, where $c^{(k)}$ is obtained from (23) by replacing $\alpha, \delta, \lambda_1, \lambda_2, \gamma_1$ and γ_2 with $\alpha^{(k)}, \delta^{(k)}, \lambda_1^{(k)}, \lambda_2^{(k)}, \gamma_1^{(k)}$ and $\gamma_2^{(k)}$, respectively. At the k -th stage of the EM algorithm, the 'pseudo' log-likelihood function without the additive constant can be obtained as

$$\begin{aligned}
l_{sb}^{(k)}(\boldsymbol{\Theta}_1 | \mathcal{D}, \beta) &= (n_1 + n_2) \ln \alpha + (n_1 + n_2) \ln(\lambda_1 + \lambda_2) + n_3 \ln \lambda_1 - n_3 \ln(2\lambda_1 + \lambda_2) + \\
&\quad \alpha \sum_{i \in I_1 \cup I_2} \ln x_{i+1} - c^{(k)}((\lambda_1 + \lambda_2)A_1(\mathbf{x}; \alpha) + \lambda_1 A_2(\mathbf{x}; \alpha)) + \\
&\quad (m_1 + m_2) \ln \delta + (m_1 + m_2) \ln(\gamma_1 + \gamma_2) + m_3 \ln \gamma_1 - m_3 \ln(2\gamma_1 + \gamma_2) + \\
&\quad \delta \sum_{i \in J_1 \cup J_2} \ln y_{i+1} - c^{(k)}((\gamma_1 + \gamma_2)B_1(\mathbf{y}; \delta) + \gamma_1 B_2(\mathbf{y}; \delta)). \tag{25}
\end{aligned}$$

Observe that in $l_{sb}^{(k)}(\boldsymbol{\Theta}_1 | \mathcal{D})$, $(\alpha, \lambda_1, \lambda_2)$ has been separated from $(\delta, \gamma_1, \gamma_2)$. Hence, at the M-step the $\alpha^{(k+1)}, \lambda_1^{(k+1)}, \lambda_2^{(k+1)}$, can be obtained by maximizing

$$\begin{aligned}
l_{sb,1}^{(k)}(\alpha, \lambda_1, \lambda_2 | \mathcal{D}, \beta) &= (n_1 + n_2) \ln \alpha + (n_1 + n_2) \ln(\lambda_1 + \lambda_2) + n_3 \ln \lambda_1 - n_3 \ln(2\lambda_1 + \lambda_2) + \\
&\quad \alpha \sum_{i \in I_1 \cup I_2} \ln x_{i+1} - c^{(k)}((\lambda_1 + \lambda_2)A_1(\mathbf{x}; \alpha) + \lambda_1 A_2(\mathbf{x}; \alpha)), \tag{26}
\end{aligned}$$

with respect to α, λ_1 and λ_2 . Similarly, $\delta^{(k+1)}, \gamma_1^{(k+1)}$ and $\gamma_2^{(k+1)}$ can be obtained by maximizing

$$\begin{aligned}
l_{sb,2}^{(k)}(\boldsymbol{\Theta}_1 | \mathcal{D}, \beta) &= (m_1 + m_2) \ln \delta + (m_1 + m_2) \ln(\gamma_1 + \gamma_2) + m_3 \ln \gamma_1 - m_3 \ln(2\gamma_1 + \gamma_2) + \\
&\quad \delta \sum_{i \in J_1 \cup J_2} \ln y_{i+1} - c^{(k)}((\gamma_1 + \gamma_2)B_1(\mathbf{y}; \delta) + \gamma_1 B_2(\mathbf{y}; \delta)), \tag{27}
\end{aligned}$$

with respect to δ , γ_1 and γ_2 .

For fixed α , if $\widehat{\lambda}_1^{(k+1)}(\alpha)$ and $\widehat{\lambda}_2^{(k+1)}(\alpha)$ maximize (26), then

$$\begin{aligned}\widehat{\lambda}_1^{(k+1)}(\alpha) &= \frac{n_3}{c^{(k)}A_2(\mathbf{x}; \alpha)} - \frac{n_1 + n_2}{c^{(k)}(A_1(\mathbf{x}; \alpha) + A_2(\mathbf{x}; \alpha))} \\ \widehat{\lambda}_2^{(k+1)}(\alpha) &= \frac{2(n_1 + n_2)}{c^{(k)}(A_1(\mathbf{x}; \alpha) + A_2(\mathbf{x}; \alpha))} - \frac{n_3}{c^{(k)}A_2(\mathbf{x}; \alpha)}.\end{aligned}$$

Obtain $\alpha^{(k+1)}$ by the argument maximum of

$$\begin{aligned}g_7^{(k)}(\alpha) &= (n_1 + n_2) \ln \alpha + (n_1 + n_2) \ln(\lambda_1^{(k+1)}(\alpha) + \lambda_2^{(k+1)}(\alpha)) + n_3 \ln(\lambda_1^{(k+1)}(\alpha)) + \\ &\alpha \sum_{i \in I_1 \cup I_2} \ln x_{i+1} - n_3 \ln(2\lambda_1^{(k+1)}(\alpha) + \lambda_2^{(k+1)}(\alpha)) - \\ &a^{(k)}((\lambda_1^{(k+1)}(\alpha) + \lambda_2^{(k+1)}(\alpha))A_1(\mathbf{x}; \alpha) + \lambda_1^{(k+1)}(\alpha)A_2(\mathbf{x}; \alpha)).\end{aligned}$$

Once $\alpha^{(k+1)}$ is obtained $\lambda_1^{(k+1)}$ and $\lambda_2^{(k+1)}$ can be obtained as $\lambda_1^{(k+1)} = \lambda_1^{(k+1)}(\alpha^{(k+1)})$ and $\lambda_2^{(k+1)} = \lambda_2^{(k+1)}(\alpha^{(k+1)})$, respectively.

Similarly, for fixed δ , if $\widehat{\gamma}_1^{(k+1)}(\delta)$ and $\widehat{\gamma}_2^{(k+1)}(\delta)$ maximize (27), then

$$\begin{aligned}\widehat{\gamma}_1^{(k+1)}(\delta) &= \frac{m_3}{c^{(k)}B_2(\mathbf{y}; \delta)} - \frac{m_1 + m_2}{c^{(k)}(B_1(\mathbf{y}; \delta) + B_2(\mathbf{y}; \delta))} \\ \widehat{\gamma}_2^{(k+1)}(\delta) &= \frac{2(m_1 + m_2)}{c^{(k)}(B_1(\mathbf{y}; \delta) + B_2(\mathbf{y}; \delta))} - \frac{m_3}{c^{(k)}B_2(\mathbf{y}; \delta)}.\end{aligned}$$

Obtain $\delta^{(k+1)}$ by the argument maximum of

$$\begin{aligned}g_8^{(k)}(\delta) &= (m_1 + m_2) \ln \delta + (m_1 + m_2) \ln(\gamma_1^{(k+1)}(\delta) + \gamma_2^{(k+1)}(\delta)) + m_3 \ln(\gamma_1^{(k+1)}(\delta)) + \\ &\delta \sum_{i \in J_1 \cup J_2} \ln y_{i+1} - m_3 \ln(2\gamma_1^{(k+1)}(\delta) + \gamma_2^{(k+1)}(\delta)) - \\ &c^{(k)}((\gamma_1^{(k+1)}(\delta) + \gamma_2^{(k+1)}(\delta))B_1(\mathbf{y}; \delta) + \gamma_1^{(k+1)}(\delta)B_2(\mathbf{y}; \delta)).\end{aligned}$$

Once $\delta^{(k+1)}$ is obtained $\gamma_1^{(k+1)}$ and $\gamma_2^{(k+1)}$ can be obtained as $\gamma_1^{(k+1)} = \gamma_1^{(k+1)}(\delta^{(k+1)})$ and $\gamma_2^{(k+1)} = \gamma_2^{(k+1)}(\delta^{(k+1)})$, respectively.

Continue the process until the convergence takes place, and let us denote the estimates as $\widehat{\Theta}_1(\beta) = (\widehat{\alpha}(\beta), \widehat{\lambda}_1(\beta), \widehat{\lambda}_2(\beta), \widehat{\delta}(\beta), \widehat{\gamma}_1(\beta), \widehat{\gamma}_2(\beta))$. Finally obtain the MLE of β , $\widehat{\beta}$, by maximizing the profile log-likelihood function $l(\widehat{\Theta}_1(\beta), \beta | \mathbf{x}, \mathbf{y})$ as in (22). Once $\widehat{\beta}$ is obtained the MLE of Θ_1 can be obtained as $\widehat{\Theta}_1(\widehat{\beta})$.

8.2 GOODNESS OF FIT TEST FOR THE BIVARIATE PROCESS

In this case we have proposed the following test similar to the univariate process. The problem can be stated as follows. Suppose $\{(x_k, y_k); k = 1, \dots, n\}$ is a sample from a stationary sequence $\{(X_k, Y_k); k = 1, \dots, n\}$, we want to test the following hypothesis:

$$H_0 : \{(X_1, Y_1), \dots, (X_n, Y_n)\} \sim BWFP(\alpha, \lambda_1, \lambda_2, \delta, \gamma_1, \gamma_2, \beta). \quad (28)$$

We use the same notation as before, i.e. $Y_{1:n} < \dots < Y_{n:n}$ as the ordered $\{Y_1, \dots, Y_n\}$, similarly, the corresponding sample version as the small letters. Further, we denote $b_1 = E_{H_0}(Y_{1:n}), \dots, b_n = E_{H_0}(Y_{n:n})$. Let us denote

$$W_n = \max_{1 \leq i \leq n} |X_{i:n} - a_i| \quad \text{and} \quad S_n = \max_{1 \leq i \leq n} |Y_{i:n} - b_i|,$$

then we use the statistic

$$T_n = \max\{W_n, S_n\}$$

for the goodness of test H_0 as in (28). We use the following test criterion for a given level of significance $0 < \gamma < 1$,

$$\text{Reject } H_0 \text{ if } T_n < d_n(\gamma),$$

where $d_n(\gamma)$ is such that

$$P_{H_0}(T_n > d_n(\gamma)) = \gamma.$$

We can compute the approximate value of $d_n(\gamma)$ as in Section 6.

8.3 SIMULATION RESULTS AND DATA ANALYSIS

We have performed some simulation experiments to see how the proposed EM algorithm works in practice. We have considered the following parameter set: $\alpha = 1.5$, $\lambda_1 = 0.1$, $\lambda_2 = 0.2$, $\delta = 1.5$, $\gamma_1 = 0.1$, $\gamma_2 = 0.2$, and $\beta = 1.0$. We have replicated the process 1000 times and have reported the average estimates and the MSEs for different sample sizes in Table 5. We have also reported the length of the 95% parametric bootstrap confidence interval of each parameter based on 1000 replications and also the associated coverage percentages

below the MSEs. In this case also it has been observed that as the sample size increases the biases, MSEs and the lengths of the confidence intervals decrease. In all the cases the coverage percentages are close to the nominal value.

Table 5: Average estimates and the associated mean squared errors are reported. The average lengths of the parametric bootstrap confidence intervals (95%) and the coverage percentages are also reported in each box below the average estimates and MSEs. Here $\alpha = \delta = 1.0$, $\lambda_1 = \gamma_1 = 0.1$, $\lambda_2 = \gamma_2 = 0.2$ and $\beta = 1.0$.

n	α	λ_1	λ_2	δ	γ_1	γ_2	β
25	1.2499	0.1029	0.2132	1.2557	0.1041	0.2016	1.2118
	(0.3686)	(0.00011)	(0.00038)	(0.3876)	(0.00012)	(0.00041)	(0.2413)
	(0.8912)	(0.00543)	(0.01988)	(0.7842)	(0.00435)	(0.00351)	(0.4321)
	(0.932)	(0.934)	(0.941)	(0.942)	(0.939)	(0.933)	(0.940)
50	1.2227	0.1016	0.2111	1.2267	0.1098	0.2006	1.2087
	(0.2765)	(0.00009)	(0.00026)	(0.2667)	(0.00008)	(0.00026)	(0.2014)
	(0.6865)	(0.00234)	(0.01456)	(0.4768)	(0.00256)	(0.00256)	(0.2678)
	(0.941)	(0.941)	(0.940)	(0.943)	(0.939)	(0.943)	(0.944)
75	1.1945	0.1011	0.2065	1.1969	0.1017	0.2001	1.1668
	(0.1586)	(0.00007)	(0.00018)	(0.1459)	(0.00006)	(0.00020)	(0.1321)
	(0.3676)	(0.00192)	(0.01148)	(0.3564)	(0.00198)	(0.00143)	(0.1994)
	(0.943)	(0.945)	(0.949)	(0.948)	(0.946)	(0.947)	(0.951)
100	1.1231	0.1001	0.2021	1.0312	0.1013	0.2000	1.0123
	(0.0781)	(0.00002)	(0.00011)	(0.0668)	(0.00002)	(0.00013)	(0.1011)
	(0.1478)	(0.00101)	(0.01111)	(0.1567)	(0.00076)	(0.00078)	(0.1012)
	(0.951)	(0.951)	(0.952)	(0.949)	(0.949)	(0.950)	(0.950)

Now we have fitted BWFP to the bivariate data sets. We have used the EM algorithm to compute the MLEs of the unknown parameters. The MLEs and the associated 95% confidence intervals (reported within brackets) are as follows: $\hat{\alpha} = 2.1989(0.7723)$, $\hat{\delta} = 2.1121(0.6856)$, $\hat{\beta} = 1.2998(0.2897)$, $\hat{\lambda}_1 = 1.6024 \times 10^{-2}(0.4565 \times 10^{-3})$, $\hat{\lambda}_2 = 1.7511 \times 10^{-2}(0.9786 \times 10^{-3})$, $\hat{\gamma}_1 = 0.1801(0.0116)$ and $\hat{\gamma}_2 = 1.5992(0.3978)$. The corresponding log-likelihood values is -364.3698. We have performed the proposed goodness of fit test, and based on 500 bootstrap samples we have obtained the p -value of the proposed test as 0.156. Hence, we cannot reject the hypothesis (28). Now the question comes whether we should use one bivariate series or two separate univariate series. If we perform the likelihood ratio test, then based on the log-likelihood values, the p -value becomes greater than 0.25. Hence, it

indicates that we should use one bivariate series rather than two univariate series to analyze these data sets.

9 CONCLUSIONS

In this paper first we have proposed a stationary process where $P(X_n = X_{n+1}) > 0$. It has been obtained based on frailty and Marshall-Olkin bivariate Weibull distribution. We have derived different properties of the proposed process and also developed a very efficient EM algorithm to compute the MLEs of the unknown parameters. We have also developed a goodness of fit test which can be implemented very conveniently in practice. We have analyzed two data sets where there are significant amount $X_n = X_{n+1}$ using UWFP and the results are quite satisfactory. Further we have proposed BWFP where the marginals are UWFPs and they are correlated. It has famous Clayton copula structure. In this case both the data sets can be analyzed simultaneously. We have derived several properties and also a very efficient EM algorithm to reduce computational burden mainly to compute the MLEs of the unknown parameters. It is observed that the proposed method can be used in practice quite conveniently.

ACKNOWLEDGEMENTS:

The author would like to thank two unknown reviewers and the Associate Editor for their constructive comments which have helped to improve the paper significantly.

APPENDIX A

Proof of Theorem 2: To prove (i), observe that

$$\begin{aligned}
S_{X_1, \dots, X_n}(x_1, \dots, x_n) &= P(X_1 \geq x_1, \dots, X_n \geq x_n) \\
&= P\left(\bigcap_{i=1}^n \{U_i \geq V^{1/\alpha} x_i, U_{i+1} \geq V^{1/\alpha} x_i, W_i \geq V^{1/\alpha} x_i\}\right) \\
&= \int_0^\infty P\left(\bigcap_{i=1}^n \{U_i \geq v^{1/\alpha} x_i, U_{i+1} \geq v^{1/\alpha} x_i, W_i \geq v^{1/\alpha} x_i\} \middle| V = v\right) f_{GA}(v; \beta, \beta) dv \\
&= \int_0^\infty e^{-v((\lambda_1(x_1^\alpha + \sum_{i=1}^{n-1} z_i^\alpha + x_n^\alpha) + \lambda_2 \sum_{i=1}^n x_i^\alpha))} f_{GA}(v; \beta, \beta) dv \\
&= \left(1 + \lambda_1 \left(x_1^\alpha + \sum_{i=1}^{n-1} z_i^\alpha + x_n^\alpha\right) + \lambda_2 \sum_{i=1}^n x_i^\alpha\right)^{-\beta}.
\end{aligned}$$

The proof of (ii), (iii) and (iv) follow from the definition of X_n as defined in (10). Using (iv), (v) can be easily obtained as $X_n|V$ is independent of $X_{n-k}|V$, for any $k > 1$. ■

Proof of Theorem 3: The proof of (i), (ii) and (iii) follow from the definition. To prove (iv), let us define $R_n = \min\{U_n, S_n\}$ and $Q_n = \min\{W_n, T_n\}$. Observe that

$$M_n = \min\{X_n, Y_n\} = \min\left\{\frac{R_n}{V^{1/\alpha}}, \frac{R_{n+1}}{V^{1/\alpha}}, \frac{Q_n}{V^{1/\alpha}}\right\}.$$

Hence, the result follows. The proof of (v) can be obtained by conditioning on V . The proof of (vi) follows from the definition of X_n and Y_n (20). Finally, the proof of (vii) can be obtained using (vi). ■

Declaration of Funding: No funding was received to prepare the manuscript.

References

- [1] Arnold, B. C. and Hallet, T. A characterization of the pareto process among stationary processes of the form $x_n = c \min(x_{n-1}, y_n)$ with application to inference and simulation. *Statistics and Probability Letters*, 8:377 – 380, 1989.
- [2] Christensen, S. Phase-type distributions and optimal stopping for autoregressive processes. *Journal of Applied Probability*, 49:22 – 39, 2012.

- [3] Jayakumar, K. and Babu, M. G. Some generalizations of weibull distribution and related processes. *Journal of Statistical Theory and Applications*, 14:425 – 434, 2015.
- [4] Kundu, D. Bivariate distributions with singular components. *Springer Handbook of Engineering Statistics*, Editor: Hoang Pham, pages 733–761, 2023.
- [5] Kundu, D. A stationary weibull-process and its applications. *Journal of Applied Statistics*, 50:2681–2700, 2023.
- [6] Kundu, D. A stationary proportional hazard class process and its applications. *Methodology and Computing in Applied Probability*, 26:Article No. 45, 2024.
- [7] Kundu, D. On bivariate weibull frailty model. *Special Proceedings of the Annual International Conference of the Society of Statistics, Computer and Applications*, pages 13–30, 2025.
- [8] Louis, T. Finding the observed information matrix when using the em algorithm. *Journal of the Royal Statistical Society, Ser B*, 44:226–233, 1982.
- [9] Marshall, A. and Olkin, I. A multivariate exponential distribution. *Journal of American Statistical Association*, 62:30–44, 1967.
- [10] Navikov, A. and Shiryaev, A. On solution of the optimal stopping problem for processes with independent increments. *Stochastics*, 79:393–406, 2007.
- [11] Pillai, R. Semi-pareto processes. *Journal of Applied Probability*, 28:461–465, 1991.
- [12] Sim, C. Simulation of weibull and gamma autoregressive stationary processes. *Communication in Statistics - Simulation and Computation*, 15:1141 – 1146, 1980.
- [13] Tavares, L. V. An exponential markovian stationary process. *Journal of Applied Probability*, 17:1117–1120, 1980.
- [14] Vaupel, J. W., Manton, K., and Stallard, E. The impact of heterogeneity in individual frailty on the dynamics of mortality. *Demography*, 16:439 – 454, 1979.

- [15] Yeh, H., Arnold, B., and Robertson, C. Pareto process. *Journal of Applied Probability*, 25:291 – 301, 1988.