

## Bias Estimation of Linear Regression Model with Autoregressive Scheme using Simulation Study

SAJID ALI KHAN <sup>a</sup>, SAYYAD KHURSHID <sup>b</sup>, SHABNAM ARSHAD <sup>b</sup>, OWAIS MUSHTAQ <sup>b</sup>

<sup>a</sup> Green Hills Postgraduate College Rawalakot Poonch AJK

<sup>b</sup> Government Postgraduate College Boys Rawalakot AJK

• Received: 25 January 2021 • Accepted: 16 March 2021 • Published Online: 28 March 2021

### Abstract

In regression modeling, first-order auto correlated errors are often a problem, when the data also suffers from independent variables. Generalized Least Squares (GLS) estimation is no longer the best alternative to Ordinary Least Squares (OLS). The Monte Carlo simulation illustrates that regression estimation using data transformed according to the GLS method provides estimates of the regression coefficients which are superior to OLS estimates. In GLS, we observe that in sample size 200 and  $\sigma=3$  with correlation level 0.90 the bias of GLS  $\beta_0$  is  $-0.1737$ , which is less than all bias estimates, and in sample size 200 and  $\sigma = 1$  with correlation level 0.90 the bias of GLS  $\beta_0$  is  $8.6802$ , which is maximum in all levels. Similarly minimum and maximum bias values of OLS and GLS of  $\beta_1$  are  $-0.0816$ ,  $-7.6101$  and  $0.1371$ ,  $0.1383$  respectively. The average values of parameters of the OLS and GLS estimation with different size of sample and correlation levels are estimated. It is found that for large samples both methods give similar results but for small sample size GLS is best fitted as compared to OLS.

Keywords: OLS, GLS, Monte Carlo.

### 1. Introduction

We introduce about the linear regression estimation and the effect of serial correlation under first order autoregressive scheme by different scientists. The regression examination is a factual strategy broadly utilized in numerous fields, for example, financial aspects, innovation, sociologies and account. A direct relapse model is built to depict the connection between the needy variable and one or a few indicator factors. This regression model could be straightforward or different. Not with standing, in the direct relapse model, certain presumptions are made on how a dataset will be created by a hidden information producing measure.

\*Corresponding author: [sajid.ali680@gmail.com](mailto:sajid.ali680@gmail.com)

As indicated by these suspicions incorporate linearity, homoscedasticity, ordinariness and no autocorrelation between the blunder terms. Also, relapse model portrays the estimation of the needy variable as the amount of two sections, a deterministic part (illustrative factors) and the arbitrary part. The mistake term is basically an unsettling influence to a generally steady relationship and can catch the excess data in the reliant variable which couldn't be clarified by the autonomous factors.

Identifying with the supposition on the blunder term, on the off chance that the pre-sumption of no relationship in the mistake term is disregarded, at that point, the fundamental model would be ruined with the standard mistakes of the boundaries getting one-sided. Besides, if the mistakes are related, the least squares assessors are wasteful and the assessed fluctuations are not suitable. By definition, autocorrelation is the slack connection of a given arrangement with itself, slacked by various time units. Hence, while applying relapse models to financial or the board information within the sight of autocorrelation, the normal least squares assessment technique stops to give productive assessors and suitable fluctuations .

### **Objective of the Study**

The main objectives are:

- To apply the Monte Carlo method on linear regression model with First Order Auto-Regressive scheme.
- To compare the biases of OLS and GLS estimators in linear regression model with First Order Auto-Regressive scheme.

## **2. Literature review**

Yale and Forsythe (1976) in [1] assessed a basic straight relapse model and contrast this strategy and OLS through relative productivity estimations got from Monte Carlo tests and they likewise apply this technique to a genuine informational collection. This method is likewise done by Rivest (1994) [2] for different slanted disseminations. Additionally, to eliminate the high impact of peripheral perceptions, Hoo *et. al.*, (2002) [3] built up this methodology, examine the idea of powerful measurements and present the techniques for strong multivariate exception separating.

Cheung (2007)[4] proposes an altered least-squares relapse approach un-weighted least squares relapse with a Huber-White powerful standard blunder for assessment of danger contrasts. Four adaptations of the powerful standard mistake are thought of. The binomial, common least squares and adjusted least-squares assessors are thought about diagnostically in a basic circumstance of one introduction variable. Multivariable relapse examinations are recreated to exhibit the convenience of the methodology. For test sizes of roughly 200 or less, a little example adaptation of the hearty standard mistake is suggested. The strategy is represented with information from a patient review in which the binomial relapse neglects to meet in the investigations of four out of five result factors.

Kiviet (2011) in [5] concentrated in econometric hypothesis are enhanced by Monte Carlo reenactment examinations. These show the properties of elective derivation strategies when applied to tests drawn from generally completely engineered information creating measures. They ought to give data on how strategies, which might be sound asymptotically, act in limited examples and afterward reveal the impacts of, model attributes

too complex to even consider analyzing systematically. Additionally the understanding of applied examinations ought to frequently profit when enhanced by a devoted recreation study, in light of a plan motivated by the hypothesized real observational information producing measure, which would verge on bootstrapping.

Ayinde et al., (2012) in [6] saw the exhibitions of assessors of direct relapse model with auto corresponded mistake term have been credited to the nature and particular of the illustrative factors. The infringement of presumption of the autonomy of the informative factors isn't remarkable particularly in business, monetary and sociologies, prompting the improvement of numerous assessors. In addition, expectation is one of the principle embodiments of relapse examination.

In [7] Suvarna and Ismail (2016) utilized the direct factual model, scientists face assortment of issues because of non trial nature for example vulnerability about the idea of the mistake cycle, model misspecifications, subordinate regressors and so on The marvel of connected mistakes in direct relapse models including time arrangement information is called autocorrelation. Infringement of the supposition of free regressors prompts multicollinearity. Henceforth, Ordinary edge gauges are loose to be very useful if there should arise an occurrence of auto related relapse model with the multicollinearity issue.

### 3. Material And Methods

#### 3.1. Monte Carlo Simulation

Estimation of parameters with different correlation levels in linear regression model and constructed Monte Carlo (MC) Simulation with different sizes of sample. A Simulation is the impersonation of the activity of a genuine cycle or framework over the long run. Regardless of whether done by hand or on a PC, reenactment includes the age of a counterfeit history of a framework and the perception of that fake history to draw derivations concerning the working qualities of the genuine framework .

This simulation study is performed for sample sizes  $n=50, 100, 200, 300$  and  $500$ . The values of the model parameters of betas (1, 1) and first order autoregressive coefficients ( $\rho = -0.9, -0.5, 0.5, 0.9$ ) with error variances (1, 3). All experimental results reported are based on 5000 replications of sample sizes. For data generation and analysis, R programming and Minitab have been used.

#### 3.2. Ordinary Least Squares

Ordinary Least Squares (OLS) regression is a measurable technique for investigation that appraises the connection between at least one free factors and a needy variable; the strategy gauges the relationship by limiting the amount of the squares in the contrast between the noticed and anticipated estimations of the reliant variable designed as a straight line . The simple linear regression model is:

$$Y = \beta_0 + \beta_1 X + e.$$

In matrix form

$$Y = X\beta + e.$$

Estimated model is

$$Y = X\beta + e.$$

$$Y - X\beta = e \quad Y = Y.$$

By minimizing the sum of squares of residuals that is

$$\sum e_i^2 = e^t e$$

$$e^t e = [Y - X\beta]^t [Y - X\beta]$$

$$e^t e = [Y^t - X^t \beta^t] [Y - X\beta]$$

$$e^t e = Y^t Y - Y^t X \beta - X^t \beta^t Y + X^t \beta^t X \beta$$

Since  $\beta^t X^t Y$  is scalar, therefore it is equal its transpose i.e.  $\beta^t X^t Y = \beta X Y^t$

$$e^t e = Y^t Y - \beta^t X^t Y - \beta^t X^t Y + X^t \beta^t X \beta$$

$$e^t e = Y^t Y - 2\beta^t X^t Y + X^t \beta \beta^t X = \beta^t.$$

Minimize with respect to  $\beta$  and equating zero.

$$\frac{\partial e^t e}{\partial \beta} = -2X^t Y + 2\beta X^t X$$

$$0 = -2X^t Y + 2\beta X^t X$$

$$2X^t Y = 2\beta X^t X$$

$$\frac{2X^t Y}{2X^t X} = \beta$$

$$\beta_{OLS} = (X^t X)^{-1} X^t Y$$

$$\text{Bias}(\beta_{OLS}) = \beta_{OLS} - \beta.$$

### 3.3. Generalized Least Squares

Generalized Least Squares (GLS) is a technique for fitting coefficients of illustrative factors that help to foresee the results of a reliant arbitrary variable. As its name recommends, GLS incorporates Ordinary Least Squares (OLS) as an uncommon case. GLS is likewise called "Aitken's assessor," after Aitken (1935) [8]. The simple linear regression model is:

$$Y = \beta_0 + \beta_1 X + e.$$

In matrix form

$$Y = X\beta + e.$$

Estimated model is

$$Y = X\beta + e.$$

$$Y = X\beta + eY = Y$$

with  $E(e) = 0$ , and  $\text{var}(e) = \sigma^2\Omega$ , where  $\Omega$  is a known  $n \times n$  matrix. Given that  $\sigma^2\Omega$  is a covariance matrix, we know that  $\Omega$  must be symmetric and non singular. Therefore we can define:

$$\Omega = K^t K = K K$$

with  $K$  called the square root of  $\Omega$ . Let we define:

$$\begin{cases} Y' = K^{-1}Y, \\ X' = K^{-1}X, \\ e' = K^{-1}e. \end{cases}$$

Note that:

$$E(e') = E(K^{-1}e) = K^{-1}E(e) = K^{-1}(0) = 0,$$

$$\text{Var}(e') = \text{Var}(K^{-1}e) = K^{-1}\text{Var}(e)K^{-1} = K^{-1}\sigma^2\Omega K^{-1} = \sigma^2 K^{-1}K K^{-1} = \sigma^2.$$

By minimizing the sum of squares of residuals that is  $\sum e_i^2 = e'^t e$

$$e'^t e = [Y' - X'\beta]^t [Y' - X'\beta],$$

$$e'^t e = [K^{-1}Y - K^{-1}X\beta]^t [K^{-1}Y - K^{-1}X\beta],$$

$$e'^t e = [K^{-1}(Y - X\beta)]^t [K^{-1}(Y - X\beta)],$$

$$e'^t e = (Y - X\beta)^t K^{-1} K^{-1} (Y - X\beta),$$

$$\begin{aligned}
e'^t e &= (Y - X\beta)^t \Omega^{-1} (Y - X\beta), \\
e'^t e &= (Y^t - X^t \beta^t) \Omega^{-1} (Y - X\beta), \\
e'^t e &= Y^t \Omega^{-1} Y - Y^t \Omega^{-1} X\beta - X^t \beta^t \Omega^{-1} Y + X^t \beta^t \Omega^{-1} X\beta, \\
e'^t e &= Y^t \Omega^{-1} Y - \beta^t X^t \Omega^{-1} Y - \beta^t X^t \Omega^{-1} Y + \beta^t X^t \Omega^{-1} X\beta, \\
e'^t e &= Y^t \Omega^{-1} Y - 2\beta^t X^t \Omega^{-1} Y + \beta^t X^t \Omega^{-1} X\beta.
\end{aligned}$$

Minimize with respect to  $\beta$  and equating zero.

$$\begin{aligned}
\frac{\partial e'^t e}{\partial \beta} &= -2X^t \Omega^{-1} Y + 2X^t \Omega^{-1} X\beta, \\
0 &= -2X^t \Omega^{-1} Y + 2\beta^t X^t \Omega^{-1} X, \\
2X^t \Omega^{-1} Y &= 2\beta^t X^t \Omega^{-1} X, \\
\beta &= \frac{X^t \Omega^{-1} Y}{X^t \Omega^{-1} X}, \\
\beta_{GLS} &= (X^t \Omega^{-1} X)^{-1} X^t \Omega^{-1} Y, \\
\text{Bias}(\beta_{GLS}) &= \beta_{GLS} - \beta.
\end{aligned}$$

### 3.4. First Order Auto-Regressive (AR) Scheme

First order autocorrelation results from correlation between the error terms of adjacent time periods. If first order autocorrelation is present, the error for one time period  $e_t$  is a function of the error of the previous time period  $e_{t-1}$  as follow:

$$e_t = \rho e_{t-1} + u_t,$$

where  $E(u_t) = 0$ ,  $E(u_t^2) = \sigma_u^2$  and  $\rho$  is the parameter depicting the functional relationship among observations of the error term,  $e_t$  and  $u_t$  is stochastic error term which is iid.

## 4. Results and Discussion

We estimate the average values of parameters with different correlation structures for linear regression technique. For the generation of data set we use Monte Carlo Simulation method with linear regression model of OLS and GLS with first order autoregressive scheme. The Monte Carlo Simulation results are given in Table 4.1.

**Fig. 4.1:** Bias OLS  $\beta_0$  when  $\sigma = 1$

In Fig. 4.1 the Bias of OLS ( $\beta_0$ ), we observe that in different sample sizes (50,100,200,300,500) and  $\sigma = 1$  with different correlation levels (-0.90,-0.50, 0.50, 0.90). In sample size 500 with correlation level 0.90 the bias of OLS ( $\beta_0$ ) is -0.0642 which is less than all others, and in sample size 200 with  $\sigma = 1$  the correlation level 0.90 have maximum value of bias OLS ( $\beta_0$ ) is 0.0917.

Table 1: Monte Carlo Simulation Estimation Varying Sample Sizes and Correlations

|                        | n=50                | n=100               | n=200               | n=300               | n=500               |
|------------------------|---------------------|---------------------|---------------------|---------------------|---------------------|
|                        | $\beta_0$ $\beta_1$ | $\beta_0$ $\beta_1$ | $\beta_0$ $\beta_1$ | $\beta_0$ $\beta_1$ | $\beta_0$ $\beta_1$ |
| <b>rho=-0.9, var=1</b> |                     |                     |                     |                     |                     |
| Bias OLS               | 0.0004 -0.0045      | 0.0006 -0.0197      | 0.0158 -0.0320      | 0.0098 0.0161       | -0.0023 -0.0042     |
| Bias GLS               | 0.0004 -0.0046      | 0.0006 -0.0199      | 0.0158 -0.0321      | 0.0099 0.0164       | -0.0233 -0.0042     |
| <b>rho=-0.5, var=1</b> |                     |                     |                     |                     |                     |
| Bias OLS               | 0.0006 -0.0091      | -0.0012 -0.0026     | 0.0200 -0.0199      | -0.0006 -0.0236     | -0.0090 -0.0045     |
| Bias GLS               | 0.0006 -0.0091      | -0.0012 -0.0026     | 0.0199 -0.0202      | -0.0006 -0.0234     | -0.0090 -0.0046     |
| <b>rho=0.5, var=1</b>  |                     |                     |                     |                     |                     |
| Bias OLS               | -0.0043 -0.0022     | 0.0059 -0.0019      | 0.0226 0.0083       | 0.0178 0.0001       | 0.0116 -0.0581      |
| Bias GLS               | -0.0044 -0.0011     | 0.0055 -0.0011      | 0.0218 0.0030       | 0.0179 -0.0001      | 0.0103 -0.0390      |
| <b>rho=0.9, var=1</b>  |                     |                     |                     |                     |                     |
| Bias OLS               | 0.0022 0.0035       | 0.0313 0.0133       | 0.0917 0.0233       | 0.0127 0.0125       | -0.064 -0.0105      |
| Bias GLS               | 0.0085 -0.0029      | 0.0100 -0.0042      | 8.6802 -7.6101      | 0.0039 -0.0531      | -0.1254 -0.0190     |
| <b>rho=-0.9, var=3</b> |                     |                     |                     |                     |                     |
| Bias OLS               | -0.0024 -0.0181     | -0.0028 0.0030      | -0.0218 -0.0633     | 0.0143 0.1373       | -0.0069 0.0682      |
| Bias GLS               | -0.0024 -0.0183     | -0.0028 0.0029      | -0.0218 -0.0640     | 0.0143 0.1383       | -0.0070 0.0687      |
| <b>rho=-0.5, var=3</b> |                     |                     |                     |                     |                     |
| Bias OLS               | 0.0029 -0.0126      | -0.0017 0.0077      | 0.0282 -0.0810      | 0.0374 0.0162       | 0.0128 0.0157       |
| Bias GLS               | 0.0029 -0.0128      | -0.0017 0.0077      | 0.0282 -0.0814      | 0.0374 0.0164       | 0.0128 0.0156       |
| <b>rho=0.5, var=3</b>  |                     |                     |                     |                     |                     |
| Bias OLS               | 0.0062 0.0093       | 0.0265 0.0148       | -0.0152 0.0078      | 0.0621 -0.0816      | -0.0747 -0.0314     |
| Bias GLS               | 0.0066 0.0047       | 0.0265 0.0121       | -0.0150 0.0174      | 0.0600 -0.0177      | -0.0732 -0.0148     |
| <b>rho=0.9, var=3</b>  |                     |                     |                     |                     |                     |
| Bias OLS               | 0.0315 -0.0089      | -0.0388 -0.0027     | -0.1625 0.0475      | 0.2006 -0.0251      | 0.5300 0.0794       |
| Bias GLS               | 0.3205 -0.0028      | -0.053 -0.0086      | -0.1737 -0.0123     | 0.1447 -0.0488      | 0.5223 0.0037       |

Table 2: Bias OLS  $\beta_0$  when  $\sigma = 1$ 

| Sample Size | Rho     |         |         |         |
|-------------|---------|---------|---------|---------|
|             | -0.90   | -0.50   | 0.50    | 0.90    |
| 50          | 0.0004  | 0.0006  | -0.0043 | 0.0022  |
| 100         | 0.0006  | -0.0012 | 0.0059  | 0.0313  |
| 200         | 0.0158  | 0.0200  | 0.0226  | 0.0917  |
| 300         | 0.0098  | -0.0006 | 0.0178  | 0.0127  |
| 500         | -0.0023 | 0.0090  | 0.0116  | -0.0641 |

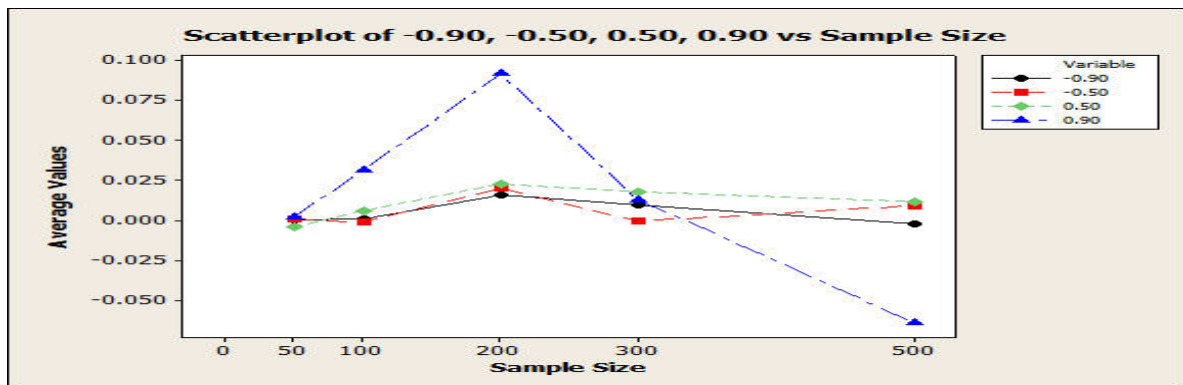
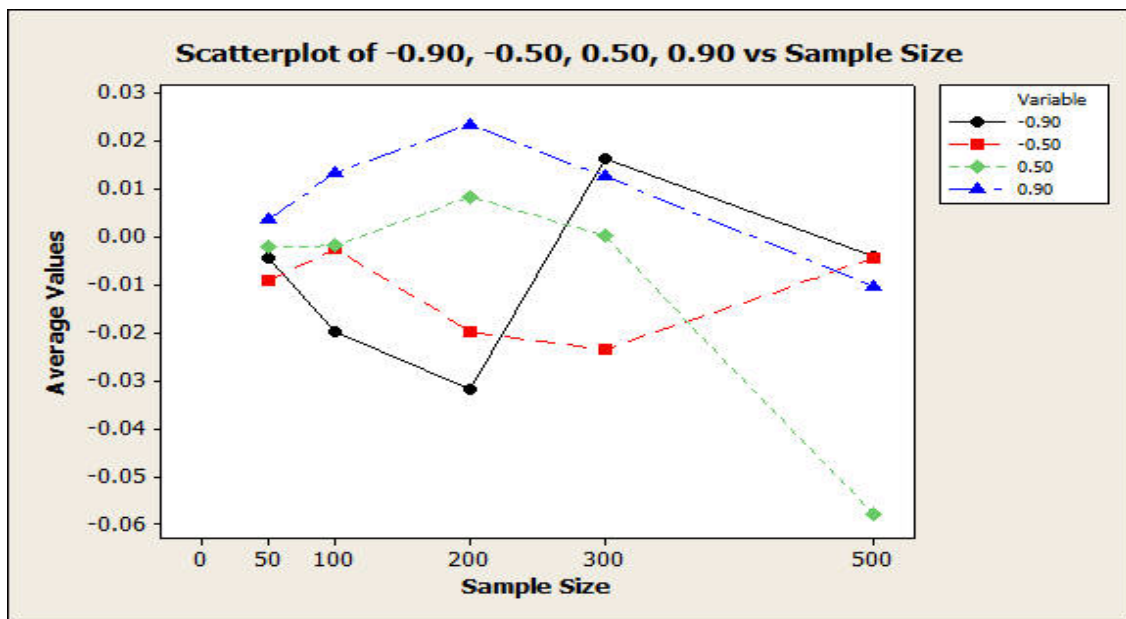


Table 3: Bias OLS  $\beta_1$  when  $\sigma = 1$ 

| Sample Size | Rho     |         |         |         |
|-------------|---------|---------|---------|---------|
|             | -0.90   | -0.50   | 0.50    | 0.90    |
| 50          | -0.0045 | -0.0091 | -0.0022 | 0.0035  |
| 100         | -0.0197 | -0.0026 | -0.0019 | 0.0133  |
| 200         | -0.0320 | -0.0199 | 0.0083  | 0.0233  |
| 300         | 0.0161  | -0.0236 | 0.0001  | 0.0125  |
| 500         | -0.0042 | -0.0045 | -0.0581 | -0.0105 |

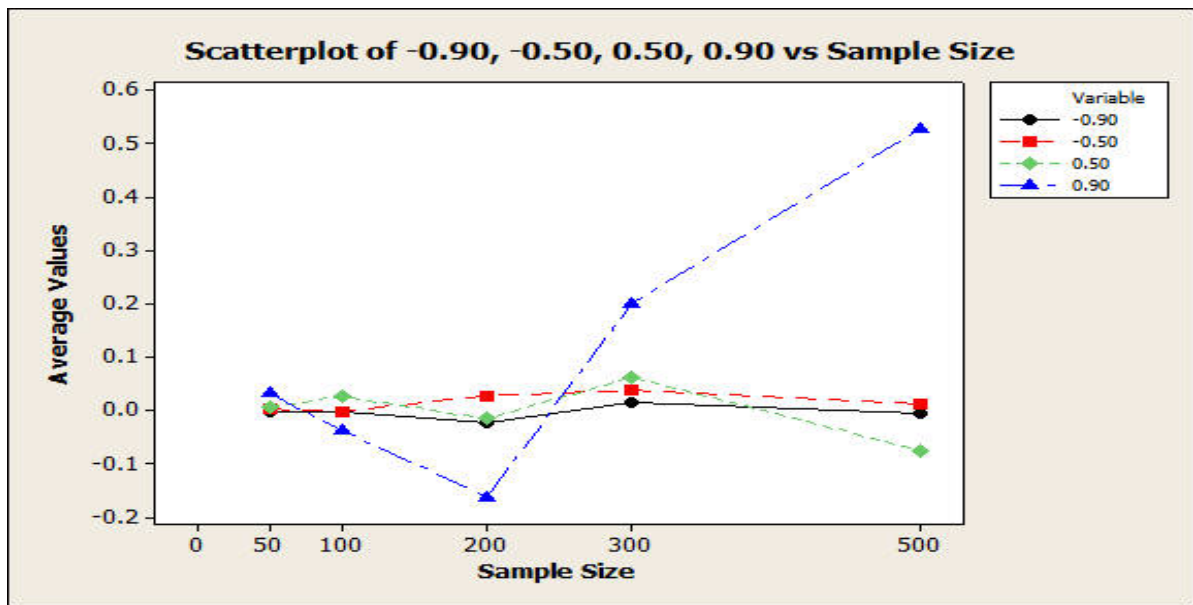
Fig. 4.2: Bias OLS  $\beta_1$  when  $\sigma = 1$ 

In Fig. 4.2 the bias of OLS ( $\beta_1$ ), we observe that in different sample sizes (50,100,200,300,500) and  $\sigma = 1$  with different correlation levels (-0.90,-0.50, 0.50, 0.90). In sample size 500 with correlation level 0.50 the bias of OLS ( $\beta_1$ ) is -0.0581 which is less than all others, and in sample size 200 with  $\sigma = 1$  the correlation level 0.90 have maximum value of bias OLS ( $\beta_1$ ) is 0.0233.



Table 4: Bias OLS  $\beta_0$  when  $\sigma = 3$ 

| Sample Size | Rho     |         |         |         |
|-------------|---------|---------|---------|---------|
|             | -0.90   | -0.50   | 0.50    | 0.90    |
| 50          | -0.0024 | 0.0029  | 0.0062  | 0.0315  |
| 100         | -0.0028 | -0.0017 | 0.0265  | -0.0388 |
| 200         | -0.0218 | 0.0282  | -0.0152 | -0.1625 |
| 300         | 0.0143  | 0.0374  | 0.0621  | 0.2006  |
| 500         | -0.0069 | 0.0128  | -0.0747 | 0.5300  |

Fig. 4.3: Bias OLS  $\beta_0$  when  $\sigma = 3$ 

In Fig. 4.3 the bias of OLS ( $\beta_0$ ), we observe that in different sample sizes (50,100,200,300,500) and  $\sigma = 3$  with different correlation levels (-0.90,-0.50, 0.50, 0.90). In sample size 200 with correlation level 0.90 the bias of OLS ( $\beta_0$ ) is -0.1625, which is less than all others, and in sample size 500 with  $\sigma = 3$ , the correlation level 0.90 have maximum value of bias OLS ( $\beta_0$ ) is 0.5300.

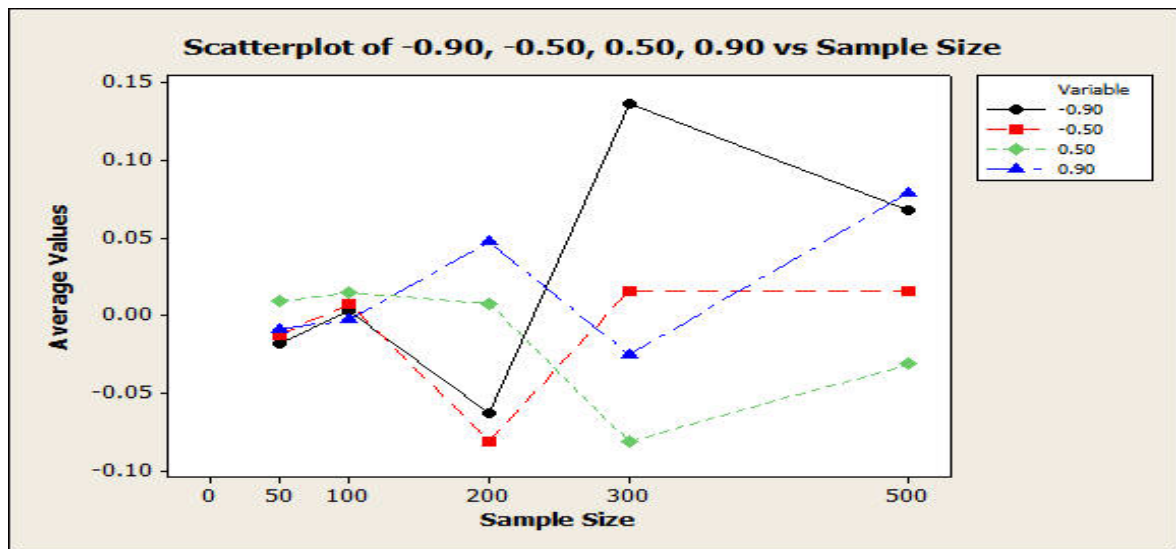
Fig. 4.4: Bias OLS  $\beta_1$  when  $\sigma = 3$ 

In Fig. 4.4 the bias of OLS ( $\beta_1$ ), we observe that in different sizes (50,100,200,300,500) and  $\sigma = 3$  with different correlation levels (-0.90,-0.50, 0.50, 0.90). In sample size 300 with correlation level 0.50 the bias of OLS ( $\beta_1$ ) is -0.0816, which is less than all others, and in sample size 300 with  $\sigma = 3$  the correlation level -0.90 have maximum value of bias OLS ( $\beta_1$ ) is 0.1373.

Fig. 4.5: Bias GLS  $\beta_0$  when  $\sigma = 1$

Table 5: Bias OLS  $\beta_1$  when  $\sigma = 3$ 

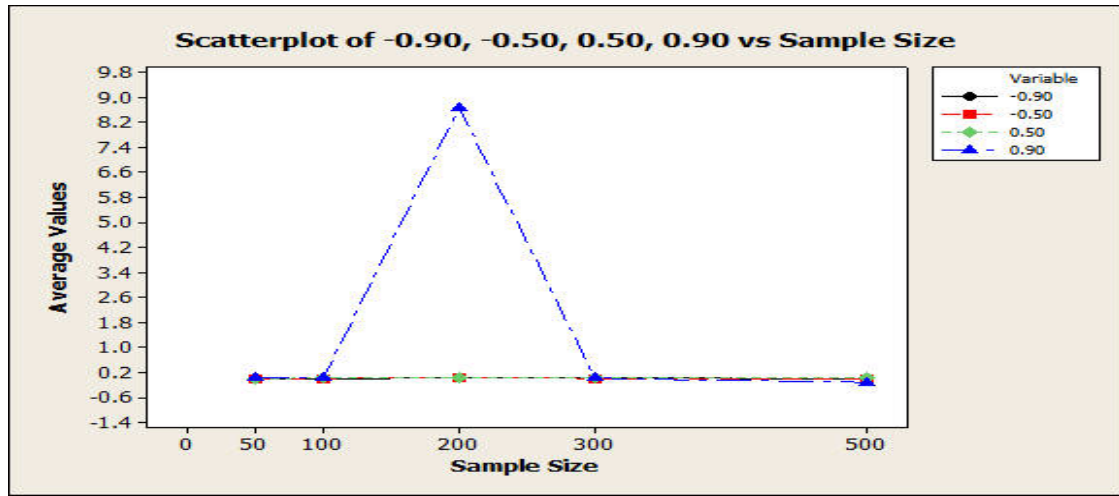
| Sample Size | Rho     |         |         |         |
|-------------|---------|---------|---------|---------|
|             | -0.90   | -0.50   | 0.50    | 0.90    |
| 50          | -0.0181 | -0.0126 | 0.0093  | -0.0089 |
| 100         | 0.0030  | 0.0077  | 0.0148  | -0.0027 |
| 200         | -0.0633 | -0.0810 | 0.0078  | 0.0475  |
| 300         | 0.1373  | 0.0162  | -0.0816 | -0.0251 |
| 500         | 0.0682  | 0.0157  | -0.0314 | 0.0794  |

Table 6: Bias GLS  $\beta_0$  when  $\sigma = 1$ 

| Sample Size | Rho     |         |         |         |
|-------------|---------|---------|---------|---------|
|             | -0.90   | -0.50   | 0.50    | 0.90    |
| 50          | 0.0004  | 0.0006  | -0.0044 | 0.0085  |
| 100         | 0.0006  | -0.0012 | 0.0055  | 0.0100  |
| 200         | 0.0158  | 0.0199  | 0.0218  | 8.6802  |
| 300         | 0.0099  | -0.0006 | 0.0179  | 0.0039  |
| 500         | -0.0233 | -0.0090 | 0.0103  | -0.1254 |

In Fig. 4.5 the bias of GLS ( $\beta_0$ ), we observe that in different sample sizes (50,100,200,300,500) and  $\sigma = 1$  with different correlation levels (-0.90,-0.50, 0.50, 0.90). In sample size 500 with correlation level 0.90 the bias of GLS ( $\beta_0$ ) is -0.1254, which is less than all others, and in sample size 200 with  $\sigma = 1$  the correlation level 0.90 have maximum value of bias GLS ( $\beta_0$ ) is 8.6802.

Fig. 4.6: Bias GLS  $\beta_1$  when  $\sigma = 1$

Table 7: Bias GLS  $\beta_1$  when  $\sigma = 1$ 

| Sample Size | Rho     |         |         |         |
|-------------|---------|---------|---------|---------|
|             | -0.90   | -0.50   | 0.50    | 0.90    |
| 50          | -0.0046 | -0.0091 | -0.0011 | -0.0029 |
| 100         | -0.0199 | -0.0026 | -0.0011 | -0.0042 |
| 200         | -0.0321 | -0.0202 | 0.0030  | -7.6101 |
| 300         | 0.0164  | -0.0234 | -0.0001 | -0.0531 |
| 500         | -0.0042 | -0.0046 | -0.0390 | -0.0190 |

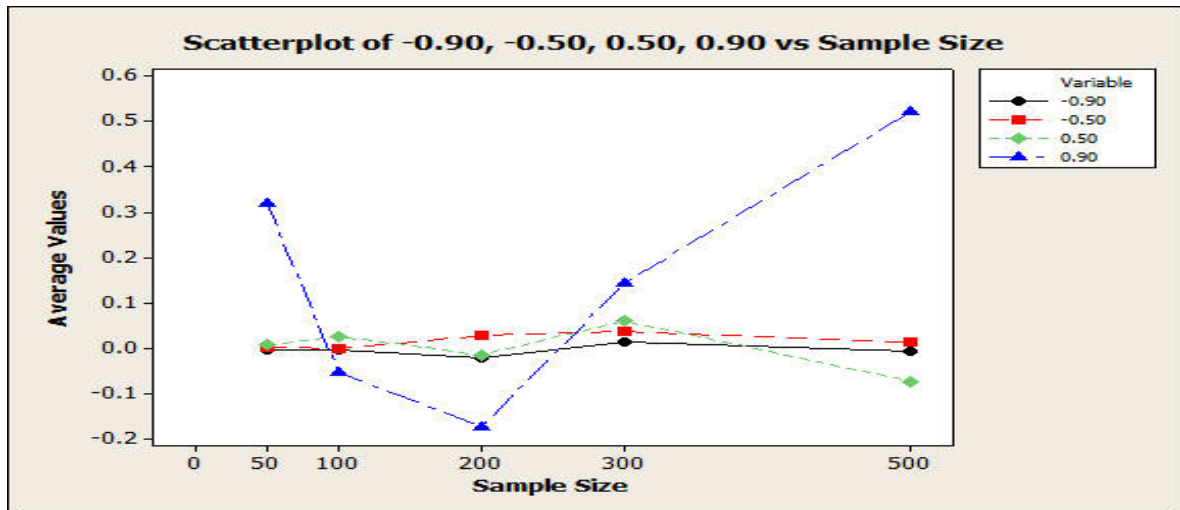
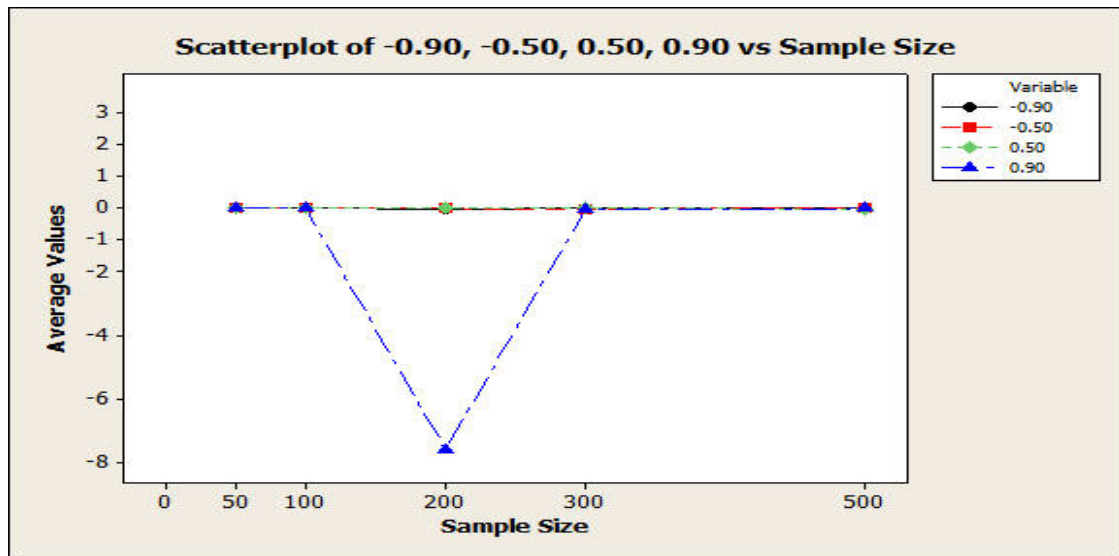
In Fig. 4.6 the bias of GLS ( $\beta_1$ ), we observe that in different sample sizes (50,100,200,300,500) and  $\sigma = 1$  with different correlation levels (-0.9,-0.50, 0.50, 0.90). In sample size 200 with correlation level 0.90 the bias of GLS ( $\beta_1$ ) is -7.6101, which is less than all others, and in sample size 300 with  $\sigma = 1$  the correlation level -0.90 have maximum value of bias GLS ( $\beta_1$ ) is 0.0164.

Table 8: Bias GLS  $\beta_0$  when  $\sigma = 3$ 

| Sample Size | Rho     |         |         |         |
|-------------|---------|---------|---------|---------|
|             | -0.90   | -0.50   | 0.50    | 0.90    |
| 50          | -0.0024 | 0.0029  | 0.0066  | 0.3205  |
| 100         | -0.0028 | -0.0017 | 0.0265  | -0.0534 |
| 200         | -0.0218 | 0.0282  | -0.0150 | -0.1737 |
| 300         | 0.0143  | 0.0374  | 0.0600  | 0.1447  |
| 500         | -0.0070 | 0.0128  | -0.0732 | 0.5223  |

**Fig. 4.7:** Bias GLS  $\beta_0$  when  $\sigma = 3$ 

In Fig. 4.7 the bias of GLS ( $\beta_0$ ), we observe that in different sample sizes (50,100,200,300,500) and  $\sigma = 3$  with different correlation (-0.90,-0.50, 0.50, 0.90). In sample size 200 with



correlation level 0.90 the bias of GLS ( $\beta_0$ ) is -0.1737, which is less than all others, and in sample size 500 with  $\sigma = 3$  the correlation level 0.90 have maximum value of bias GLS ( $\beta_0$ ) is 0.5223.

Table 9: Bias GLS  $\beta_1$  when  $\sigma = 3$

| Sample Size | Rho     |         |         |         |
|-------------|---------|---------|---------|---------|
|             | -0.90   | -0.50   | 0.50    | -0.50   |
| 50          | -0.0183 | -0.0128 | 0.0047  | -0.0028 |
| 100         | 0.0029  | 0.0077  | 0.0121  | -0.0086 |
| 200         | -0.0640 | -0.0814 | 0.0174  | -0.0123 |
| 300         | 0.1383  | 0.0164  | -0.0177 | -0.0488 |
| 500         | 0.0687  | 0.0156  | -0.0148 | 0.0037  |

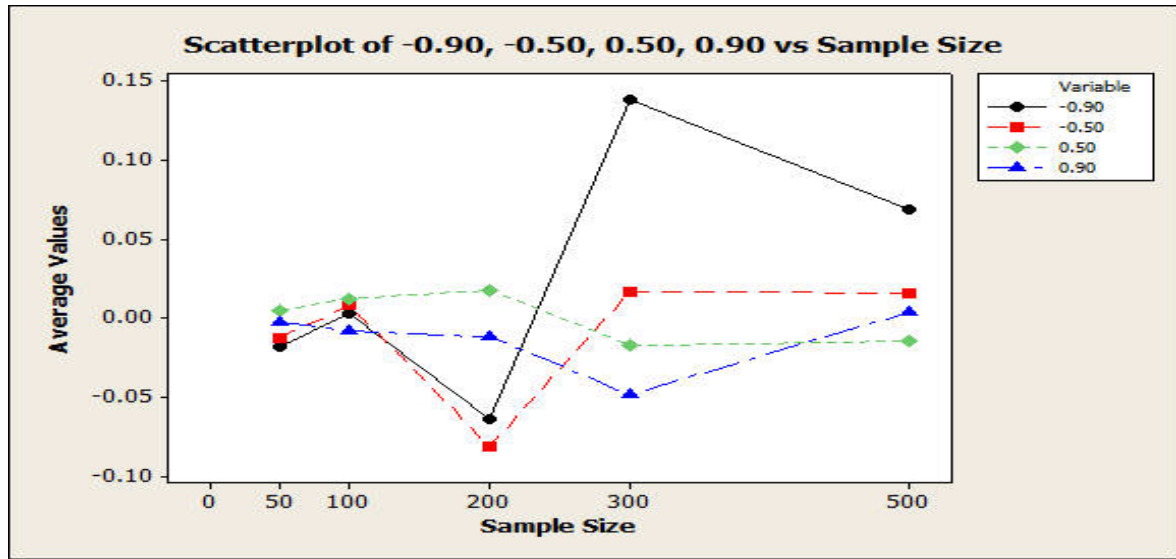


Fig. 4.8: Bias GLS  $\beta_1$  when  $\sigma = 3$

In Fig. 4.8 the bias of GLS ( $\beta_1$ ), we observe that in different sample sizes (50,100,200,300,500) and  $\sigma = 3$  with different correlation levels (-0.9,-0.50, 0.50, 0.90). In sample size 200 with correlation level -0.50 the bias of GLS ( $\beta_1$ ) is -0.0814, which is less than all others, and in sample size 300 with  $\sigma = 3$  the correlation level -0.90 have maximum value of bias GLS ( $\beta_1$ ) is 0.1383.

## 5. conclusion

We have discussed the Ordinary Least Squares and Generalized Least Squares techniques and estimate with First Order Auto-Regressive (AR1) scheme from different correlation levels by using simple linear regression model. For this purpose, we use simulation of Monte Carlo study and the experiment is repeated 5000 times and performed for different sample sizes.

The average values of parameters of the Ordinary Least Squares and Generalized Least Squares estimation with different size of sample and correlation levels are estimated. When the bias values of Ordinary Least Squares and Generalized Least Squares is not normal with haphazard manner of average values.

Comparing the bias of OLS and GLS of ( $\beta_0$ ), we observe that in different sample sizes (50,100,200,300,500) and  $\sigma = 1,3$  with different correlation levels (-0.90,-0.50, 0.50, 0.90). In sample size 200 with correlation level 0.90 the bias of OLS ( $\beta_0$ ) is -0.1625, which is less than all other estimates, and in sample size 500 with the correlation level 0.90 have maximum value of bias OLS ( $\beta_0$ ) is 0.5300.

In GLS, we observe that in sample size 200 and  $\sigma = 3$  with correlation level 0.90 the bias of GLS ( $\beta_0$ ) is -0.1737, which is less than all bias estimates, and in sample size 200 and  $\sigma = 1$  with correlation level 0.90 the bias of GLS ( $\beta_0$ ) is 8.6802, which is maximum in all levels. Similarly minimum and maximum bias values of OLS and GLS of ( $\beta_1$ ) are -0.0816, -7.6101 and 0.1371, 0.1383 respectively. These result shows GLS is best in different sample sizes and correlation situations.

## References

- [1] Yale C and Forsythe AB (1976). *Winsorized Regression*. Technometrics, **18**, 291-300.
- [2] Rivest L (1994). *Statistical Properties of Winsorized Means for Skewed Distributions*. Biometrika, **81** (2), 373-383. 42-48.<https://doi.org/10.1093/biomet/81.2.373>
- [3] Hoo K, Tvarlapati K, Piovoso M and Hajare R (2002). *A Method of Robust Multivariate Outlier Replacement*. Computers and Chemical Engineering, **26** (1), 17-39. 42-48.[https://doi.org/10.1016/S0098-1354\(01\)00734-7](https://doi.org/10.1016/S0098-1354(01)00734-7)
- [4] Cheung YB (2007). *A Modified Least-Squares Regression Approach to the Estimation of Risk Differences*. American Journal of Epidemiology **166** (11), 1337-1344. 42-48.<https://doi.org/10.1093/aje/kwm223>
- [5] Kiviet JF (2011). *Monte Carlo Simulation for Econometricians*. Foundations and Trends in Econometrics, **5** (1-2), 1-81. 42-48.<http://dx.doi.org/10.1561/08000000011>
- [6] Ayinde K, Adedayo DA and Adepoju AA (2012). *Estimators of Linear Regression Model with Autocorrelated Error Terms and Prediction using Correlated Uniform Regressors*. International Journal of Engineering Science and Technology, **4** (11), 4629-4638.
- [7] Suvarna M and Ismail B (2016). *Estimation of Linear Regression Model with Correlated Regressors in the Presence of Autocorrelation*. International Journal of Statistics and Applications, **6** (2), 35-39.
- [8] Aitken AC (1935). *On Least Squares and Linear Combinations of Observations*. Proceedings of the Royal Statistical Society **55**, 42-48.<https://doi.org/10.1017/S0370164600014346>