



ARTICLE

Assessing the efficiency of the weighted mixed Liu estimator in linear models with measurement error

 Pragma Goyal,¹  Manoj Kumar Tiwari,^{*,2} and  Vikas Bist³

¹Department of Statistics, Panjab University, Chandigarh, India.

²Department of Statistics, College of Science, Sultan Qaboos University, Muscat, Oman.

³Department of Mathematics, Panjab University, Chandigarh, India.

*Corresponding author. Email: mantiwa@gmail.com

(Received: August 16, 2025; Revised: September 24, 2025; Accepted: November 14, 2025; Published: July 02, 2026)

Section Editor: Camilla Marques Barroso

Abstract

This article introduces the weighted mixed Liu estimator for linear measurement error models. Additionally, the analysis incorporates stochastic linear restrictions along with the presence of multicollinearity. The proposed estimator facilitates the allocation of varying weights to the auxiliary information in relation to sample information. The asymptotic properties of the resultant estimator are established, and the efficacy of various estimators is assessed utilizing the mean squared error matrix criterion. A simulation study and numerical example are provided to assess the theoretical results. The simulation results and numerical example indicate that the proposed estimator outperforms the existing estimators.

Keywords: Weighted Mixed Liu Estimator; Measurement Error; Stochastic Linear Restrictions; Mean Squared Error Matrix; Auxiliary Information.

1. Introduction

In a linear regression model, one of the fundamental assumptions is that all variables are free from any measurement error. However, researchers often face difficulties in addressing measurement errors that are present in observations. These errors deteriorate the optimal properties of estimators. One significant implication of measurement errors is that the ordinary least squares estimator (OLSE) of the regression coefficient becomes inconsistent. It is important to acknowledge that OLSE is considered to be the best linear unbiased estimator when there are no measurement errors in the data. Various methodologies have been developed to tackle the difficulties presented by measurement error (Fuller, 2009). Nakamura (1990) proposed a method to address the impact of

measurement errors on parameter estimation by correcting the score function. This methodology enables both inference and parameter estimation without requiring any additional assumptions.

Moreover, the general assumption is that all explanatory variables exhibit linear independence. However, if the explanatory variables deviate from this assumption, the issue of multicollinearity impacts the data significantly. Numerous techniques have been proposed to tackle this problem. Biased estimators are preferred over unbiased estimators as they have a lesser mean squared error (MSE) value. The ridge regression estimator proposed by Hoerl & Kennard (1970) is a notable example of biased estimator. To address the limitations of the ridge estimator, Kejian (1993) presented the Liu estimator. Kaçiranlar (2003) examined the performance of the Liu estimator in linear regression models under conditions where the homoscedasticity assumption is violated. Other researchers, such as (Özkale, 2009; Owolabi *et al.*, 2022; Ahmad & Aslam, 2022), and others have recommended the use of a two-parameter estimator for improved efficiency.

An alternative approach to overcome this issue is integrating parameter estimates with sample information, such as exact or stochastic restrictions on the unknown parameter vector, suggested by Rao (2008). Durbin (1953), Theil (1963) and Theil & Goldberger (1992) introduced the mixed estimator (ME) when additional stochastic linear restrictions are assumed to be held in an unknown parameter vector. This ME remains consistent even when multicollinearity is present. It satisfies the specified linear restrictions on regression coefficients. Additionally, it has lower variability compared to the OLSE. Sarkar (1992) proposed an estimator by incorporating the ideas of both the ridge estimator and ME. In addition, Özkale (2009) presented the stochastic restricted ridge regression estimator for linear models, while Li & Yang (2010) proposed the mixed ridge estimator. Many authors introduced different estimators by incorporating prior information for linear mixed models, such as Yang *et al.* (2014) and Kuran & Özkale (2016). In contrast, Yavarizadeh *et al.* (2022) proposed ridge estimation along with stochastic restriction using Nakamura's approach for linear mixed models (LMMs) with measurement error.

It is possible that the relative importance of auxiliary information and sample information may vary. Therefore, the approach of weighted mixed estimation was employed to fulfill this objective. Schaffrin & Toutenburg (1990) explored the field of weighted mixed regression. This led to the development of the weighted mixed estimator (WME). Li & Yang (2011) combined ridge estimator into the WME procedure and derived the weighted mixed ridge estimator (WMRE) in linear regression models. The WMRE for regression coefficients in linear models with measurement errors was proposed by Ghapani *et al.* (2018). The weighted ridge estimation was suggested in Yavarizadeh & Ahmed (2022) for LMMs with measurement error while considering stochastic linear mixed restrictions. Inspired by this, our main objective in this paper is to derive a novel estimator known as the weighted mixed Liu estimator (WMLE) for linear models with measurement error. The model integrates stochastic linear restrictions when multicollinearity is present. Furthermore, the asymptotic properties of the proposed estimator and its superiority over the MLE and WME in terms of the mean squared error matrix (MSEM) of estimators is theoretically analyzed.

The structure of the article is as follows: Section 2 provides a comprehensive explanation of linear measurement error models. It also introduces the WMLE within the framework of linear measurement error models that incorporate stochastic linear restrictions (2.1). In the subsection 2.3, we derive the asymptotic properties of the estimator, along with a comparison of WMLE to MLE and WME using the MSEM criterion to evaluate its performance (2.4). Section 3 presents the results and discussion of simulation findings (3.1) and numerical evaluation (3.2). Finally, Section 4 provides concluding remarks.

2. Materials and Methods

2.1 Model Description and Estimation

We consider the linear regression model with measurement error as

$$\begin{aligned} \mathbf{y} &= \mathbf{Z}\boldsymbol{\beta} + \boldsymbol{\epsilon}, \quad \boldsymbol{\epsilon} \sim (0, \sigma^2 \mathbf{I}), \\ \mathbf{X} &= \mathbf{Z} + \boldsymbol{\Delta}, \quad \boldsymbol{\Delta} \sim (0, \boldsymbol{\Sigma}) \end{aligned} \quad (1)$$

where $\mathbf{y}$ is an $(n \times 1)$ vector of response variables and $\mathbf{Z}$ is an $(n \times p)$ matrix of unobservable explanatory variables which can be observed through $\mathbf{X}$ with the measurement error. The measurement error matrix is defined as

$$\boldsymbol{\Delta} = (\delta_1, \delta_2, \dots, \delta_n)',$$

where each δ_i is an uncorrelated random vector with $E(\delta_i) = 0$ and $\text{Var}(\delta_i) = \boldsymbol{\Sigma}$. Here, $\boldsymbol{\Sigma}$ is a $(p \times p)$ known matrix with non-negative diagonal elements, and $\boldsymbol{\beta}$ is a $(p \times 1)$ vector of unknown parameters. It is assumed that $\boldsymbol{\epsilon}$ and $\boldsymbol{\Delta}$ are independent.

For model (1), the consistent estimator of $\boldsymbol{\beta}$ is given by

$$\hat{\boldsymbol{\beta}} = (\mathbf{X}'\mathbf{X} - n\boldsymbol{\Sigma})^{-1} \mathbf{X}'\mathbf{y}.$$

This estimator is obtained by minimizing the following expression with respect to $\boldsymbol{\beta}$:

$$S_1(\boldsymbol{\beta}) = (\mathbf{y} - \mathbf{X}\boldsymbol{\beta})'(\mathbf{y} - \mathbf{X}\boldsymbol{\beta}) - n \boldsymbol{\beta}'\boldsymbol{\Sigma}\boldsymbol{\beta}.$$

The consistent estimator of σ^2 is given by

$$\hat{\sigma}^2 = \frac{1}{n} (\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}})'(\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}).$$

Let us consider the auxiliary information about $\boldsymbol{\beta}$ be in the form of independent stochastic linear restrictions:

$$\mathbf{r} = \mathbf{R}\boldsymbol{\beta} + \mathbf{e}. \quad (2)$$

Here, $\mathbf{r}$ is a $(q \times 1)$ observable random vector, $\mathbf{R}$ is a $(q \times p)$ known matrix with rank $q < p$, and $\mathbf{e}$ is a $(q \times 1)$ error vector with $E(\mathbf{e}) = \mathbf{0}$ and $\text{Var}(\mathbf{e}) = \sigma^2 \mathbf{V}$. The matrix $\mathbf{V}$ is supposed to be known and positive definite. Further, $\mathbf{e}$ is assumed to be stochastically independent of $\boldsymbol{\epsilon}$ and $\boldsymbol{\Delta}$.

In order to include the stochastic restrictions (2) when estimating parameters, we minimize

$$S_2(\boldsymbol{\beta}) = (\mathbf{y} - \mathbf{X}\boldsymbol{\beta})'(\mathbf{y} - \mathbf{X}\boldsymbol{\beta}) - n \boldsymbol{\beta}'\boldsymbol{\Sigma}\boldsymbol{\beta} + (\mathbf{r} - \mathbf{R}\boldsymbol{\beta})'\mathbf{V}^{-1}(\mathbf{r} - \mathbf{R}\boldsymbol{\beta}),$$

Differentiating $S_2(\boldsymbol{\beta})$ with respect to $\boldsymbol{\beta}$ and equating to zero, we obtain

$$\begin{aligned} \hat{\boldsymbol{\beta}}_{ME} &= (\mathbf{X}'\mathbf{X} - n\boldsymbol{\Sigma} + \mathbf{R}'\mathbf{V}^{-1}\mathbf{R})^{-1} (\mathbf{X}'\mathbf{y} + \mathbf{R}'\mathbf{V}^{-1}\mathbf{r}) \\ &= (\mathbf{S} + \mathbf{R}'\mathbf{V}^{-1}\mathbf{R})^{-1} (\mathbf{X}'\mathbf{y} + \mathbf{R}'\mathbf{V}^{-1}\mathbf{r}), \end{aligned}$$

where $\mathbf{S} = \mathbf{X}'\mathbf{X} - n\boldsymbol{\Sigma}$. Using Lemma 1 (presented in the Appendix), we can write

$$(\mathbf{S} + \mathbf{R}'\mathbf{V}^{-1}\mathbf{R})^{-1} = \mathbf{S}^{-1} - \mathbf{S}^{-1}\mathbf{R}'(\mathbf{V} + \mathbf{R}\mathbf{S}^{-1}\mathbf{R}')^{-1}\mathbf{R}\mathbf{S}^{-1}.$$

and,

$$(\mathbf{S}^{-1} - \mathbf{S}^{-1}\mathbf{R}'(\mathbf{V} + \mathbf{R}\mathbf{S}^{-1}\mathbf{R}')^{-1}\mathbf{R}\mathbf{S}^{-1})\mathbf{R}'\mathbf{V}^{-1}\mathbf{r} = \mathbf{S}^{-1}\mathbf{R}'(\mathbf{V} + \mathbf{R}\mathbf{S}^{-1}\mathbf{R}')^{-1}\mathbf{r}.$$

Therefore, $\hat{\boldsymbol{\beta}}_{ME}$ can be written as

$$\hat{\boldsymbol{\beta}}_{ME} = \hat{\boldsymbol{\beta}} + \mathbf{S}^{-1}\mathbf{R}'(\mathbf{V} + \mathbf{R}\mathbf{S}^{-1}\mathbf{R}')^{-1}(\mathbf{r} - \mathbf{R}\hat{\boldsymbol{\beta}}). \quad (3)$$

To mitigate the impact of collinearity, we opt for the Liu estimator of β instead of the ridge estimator. This choice is motivated by ability of the Liu estimator to address the limitations of the ridge estimator within the conditions of the measurement error model.

Therefore, minimizing $S_3(\beta) = (y - X\beta)'(y - X\beta) - n\beta'\Sigma\beta + (d\hat{\beta} - \beta)'(d\hat{\beta} - \beta)$ yields the following Liu estimator

$$\hat{\beta}_{LE}(d) = (S + I)^{-1}(X'y + d\hat{\beta}) = F_d\hat{\beta},$$

where $F_d = (S + I)^{-1}(S + dI)$.

Now, we are considering the Liu estimator under the stochastic linear restrictions in the measurement error model and substituting $\hat{\beta}$ with $\hat{\beta}_{LE}(d)$ in Eq.(3) to obtain mixed liu estimator,

$$\begin{aligned}\hat{\beta}_{MLE}(d) &= \hat{\beta}_{LE}(d) + S^{-1}R'(V + RS^{-1}R')^{-1}(r - R\hat{\beta}_{LE}(d)) \\ &= (S + R'V^{-1}R)^{-1}(F_dX'y + R'V^{-1}r).\end{aligned}$$

The following properties are obtained:

1. If we put $d = 1$ then F_d becomes identity matrix which implies $\hat{\beta}_{MLE}(d) = \hat{\beta}_{ME}$.
2. If stochastic linear restrictions are removed then $\hat{\beta}_{MLE}(d) = \hat{\beta}_{LE}$.

2.2 The Weighted Mixed Liu Estimator

Drawing from the recognition that the auxiliary and the sample information are not always equally significant, Schaffrin & Toutenburg (1990) introduced the concept of weighted mixed estimation. This approach involves the allocation of various weights to both the sample and the auxiliary information. Consequently, the WME of β for model (1) featuring stochastic linear restrictions (2) is achieved by minimizing

$$S_4(\beta) = (y - X\beta)'(y - X\beta) - n\beta'\Sigma\beta + w(r - R\beta)'V^{-1}(r - R\beta).$$

with respect to β . The weight, w ($0 \leq w \leq 1$), quantifies the relative significance of the auxiliary information compared to the sample information. The calculated weighted mixed estimator of β is given by

$$\hat{\beta}_{WME}(w) = (X'X + wR'V^{-1}R - n\Sigma)^{-1}(X'y + wR'V^{-1}r).$$

We may write the WME as

$$\hat{\beta}_{WME}(w) = \hat{\beta} + wS^{-1}R'(V + wRS^{-1}R')^{-1}(r - R\hat{\beta}). \quad (4)$$

Furthermore, a novel Liu type estimator is derived by merging the LE and WME entitled as WMLE. This will be achieved through the process of minimizing the following expression

$$S_5(\beta) = (y - X\beta)'(y - X\beta) - n\beta'\Sigma\beta + w(r - R\beta)'V^{-1}(r - R\beta) + (d\hat{\beta} - \beta)'(d\hat{\beta} - \beta)$$

with respect to β . Therefore, the WMLE is given by

$$\hat{\beta}_{WMLE}(w, d) = (S + wR'V^{-1}R)^{-1}(F_dX'y + wR'V^{-1}r). \quad (5)$$

Remark. The estimator $\hat{\beta}_{WMLE}$ simplifies to $\hat{\beta}_{LE}$, effectively ignoring the auxiliary information, on substituting $w = 0$. On the other hand, if we choose $w = 1$, then $\hat{\beta}_{WMLE}$ results into the mixed liu estimator $\hat{\beta}_{MLE}$. A value of w ranging from 0 to 1 signifies that the sample is given greater importance than the auxiliary information, while a value of w greater than 1 indicates the opposite scenario.

2.3 Asymptotic Properties of Estimators

It is challenging to obtain the precise distribution and finite sample properties of the estimators. This section will focus on the asymptotic properties of the estimators. It is assumed that as the sample size (n) approaches infinity, the limits of $n^{-1}(\mathbf{Z}'\mathbf{Z} + \mathbf{R}'\mathbf{V}^{-1}\mathbf{R})$ and $n^{-1}(\mathbf{Z}'\mathbf{Z} + w\mathbf{R}'\mathbf{V}^{-1}\mathbf{R})$ exist.

Theorem 1. The estimator $\hat{\beta}_{WMLE}$ has an asymptotic normal distribution with $E(\hat{\beta}_{WMLE}) = \beta + \mathbf{M}_w^{-1}(\mathbf{G}_d - \mathbf{I})\mathbf{Z}'\mathbf{Z}\beta$, and covariance matrix $AVar(\hat{\beta}_{WMLE}) = \mathbf{M}_w^{-1}[\mathbf{G}_d\mathbf{B}\mathbf{G}_d + \sigma^2(\mathbf{G}_d\mathbf{Z}'\mathbf{Z}\mathbf{G}_d + w^2\mathbf{R}'\mathbf{V}^{-1}\mathbf{R})]\mathbf{M}_w^{-1}$, where $\mathbf{M}_{(w,d)} = \mathbf{G}_d\mathbf{Z}'\mathbf{Z} + w\mathbf{R}'\mathbf{V}^{-1}\mathbf{R}$, $\mathbf{M}_w = \mathbf{M}_{(w,1)} = \mathbf{G}\mathbf{Z}'\mathbf{Z} + w\mathbf{R}'\mathbf{V}^{-1}\mathbf{R}$, and $\mathbf{B} = (n\sigma^2 + \beta'\mathbf{Z}'\mathbf{Z}\beta)\Sigma$.

Proof. Since $E(\mathbf{X}'\mathbf{X}) = \mathbf{Z}'\mathbf{Z} + n\Sigma$, (Fung et al., 2003), we have

$$\mathbf{X}'\mathbf{X} = \mathbf{Z}'\mathbf{Z} + n\Sigma + O_p(n^{1/2}),$$

then, we can write

$$\hat{\beta}_{WMLE} = (\mathbf{X}'\mathbf{X} - n\Sigma + w\mathbf{R}'\mathbf{V}^{-1}\mathbf{R})^{-1}[(\mathbf{G}_d + O_p(n^{-1/2}))\mathbf{X}'\mathbf{y} + w\mathbf{R}'\mathbf{V}^{-1}\mathbf{r}], \quad (6)$$

it follows from (5) and (6) that

$$\begin{aligned} \sqrt{n}\hat{\beta}_{WMLE} &= n[(\mathbf{Z}'\mathbf{Z} + w\mathbf{R}'\mathbf{V}^{-1}\mathbf{R}) + O_p(n^{1/2})]^{-1}n^{-1/2}[\mathbf{G}_d\mathbf{X}'\mathbf{y} + w\mathbf{R}'\mathbf{V}^{-1}\mathbf{r} + O_p(n^{-1/2})] \\ &= [\mathbf{I}_p + O_p(n^{-1/2})]^{-1}[n^{-1}(\mathbf{Z}'\mathbf{Z} + w\mathbf{R}'\mathbf{V}^{-1}\mathbf{R})]^{-1}n^{-1/2}(\mathbf{G}_d\mathbf{X}'\mathbf{y} + w\mathbf{R}'\mathbf{V}^{-1}\mathbf{r} + O_p(n^{-1/2})) \\ &= [\mathbf{I}_p + O_p(n^{-1/2})][n^{-1}(\mathbf{Z}'\mathbf{Z} + w\mathbf{R}'\mathbf{V}^{-1}\mathbf{R})]^{-1}[n^{-1/2}(\mathbf{G}_d\mathbf{X}'\mathbf{y} + w\mathbf{R}'\mathbf{V}^{-1}\mathbf{r} + O_p(n^{-1/2}))] \end{aligned} \quad (7)$$

Moreover, since the limit of $\mathbf{C} = n^{-1}(\mathbf{Z}'\mathbf{Z} + w\mathbf{R}'\mathbf{V}^{-1}\mathbf{R})$ exists, then (7) can be written as

$$\sqrt{n}\hat{\beta}_{WMLE} = \mathbf{C}^{-1}\zeta + O_p(n^{-1/2}), \quad (8)$$

where $\zeta = n^{-1/2}(\mathbf{G}_d\mathbf{X}'\mathbf{y} + w\mathbf{R}'\mathbf{V}^{-1}\mathbf{r})$ is asymptotically normal (Fung et al., 2003). So it follows from $E(\mathbf{G}_d\mathbf{X}'\mathbf{y} + w\mathbf{R}'\mathbf{V}^{-1}\mathbf{r}) = (\mathbf{G}_d\mathbf{Z}'\mathbf{Z} + w\mathbf{R}'\mathbf{V}^{-1}\mathbf{R})\beta$ that

$$E(\zeta) = n^{-1/2}\mathbf{M}_{w,d}\beta.$$

Consequently, we have

$$\sqrt{n}[\hat{\beta}_{WMLE}(w, d) - \mathbf{M}_w^{-1}\mathbf{M}_{(w,d)}\beta] = \mathbf{C}^{-1}(\zeta - E(\zeta)) + O_p(n^{-1/2})$$

which indicate that $\sqrt{n}[\hat{\beta}_{WMLE}(w, d) - \mathbf{M}_w^{-1}\mathbf{M}_{(w,d)}\beta]$ is asymptotically normal with mean zero. Furthermore, from (8) we have

$$AVar[\sqrt{n}\hat{\beta}_{WMLE}(w, d)] = \mathbf{C}^{-1}\zeta\mathbf{C}^{-1}.$$

The variance of ζ can be obtained from

$$\begin{aligned} Var(\zeta) &= E_{\mathbf{y}^*}[Var(\zeta|\mathbf{y}^*)] + Var_{\mathbf{y}^*}[E(\zeta|\mathbf{y}^*)] \\ &= n^{-1}E_{\mathbf{y}^*}(\mathbf{G}_d\mathbf{y}'\mathbf{y}\Sigma\mathbf{G}_d) + n^{-1}Var_{\mathbf{y}^*}(\mathbf{G}_d\mathbf{Z}'\mathbf{y} + w\mathbf{R}'\mathbf{V}^{-1}\mathbf{r}) \end{aligned}$$

where $E_{\mathbf{y}^*}$ and $Var_{\mathbf{y}^*}$ denote the expectation and variance with respect to the random vector $\mathbf{y}^{*'} = (\mathbf{y}', \mathbf{r}')$.

Since $E(\mathbf{y}'\mathbf{y}) = n\sigma^2 + \boldsymbol{\beta}'\mathbf{Z}'\mathbf{Z}\boldsymbol{\beta}$ and $\text{Var}_{\mathbf{y}^*}(\mathbf{G}_d\mathbf{Z}'\mathbf{y} + w\mathbf{R}'\mathbf{V}^{-1}\mathbf{r}) = \sigma^2(\mathbf{G}_d\mathbf{Z}'\mathbf{Z}\mathbf{G}_d + w^2\mathbf{R}'\mathbf{V}^{-1}\mathbf{R})$, therefore, we have

$$\text{Var}(\zeta) = n^{-1}(\mathbf{G}_d\mathbf{B}\mathbf{G}_d + \sigma^2(\mathbf{G}_d\mathbf{Z}'\mathbf{Z}\mathbf{G}_d + w^2\mathbf{R}'\mathbf{V}^{-1}\mathbf{R})).$$

Thus,

$$\text{AVar}[\hat{\boldsymbol{\beta}}_{WMLE}(w, d)] = \mathbf{M}_w^{-1}[\mathbf{G}_d\mathbf{B}\mathbf{G}_d + \sigma^2(\mathbf{G}_d\mathbf{Z}'\mathbf{Z}\mathbf{G}_d + w^2\mathbf{R}'\mathbf{V}^{-1}\mathbf{R})]\mathbf{M}_w^{-1}.$$

□

Corollary 1. The estimator $\hat{\boldsymbol{\beta}}_{MLE}(d)$ has an asymptotic normal distribution with $E[\hat{\boldsymbol{\beta}}_{MLE}(d)] = \mathbf{M}^{-1}\mathbf{M}_d\boldsymbol{\beta}$, and covariance matrix $\text{AVar}[\hat{\boldsymbol{\beta}}_{MLE}(d)] = \mathbf{M}^{-1}[\mathbf{G}_d\mathbf{B}\mathbf{G}_d + \sigma^2(\mathbf{G}_d\mathbf{Z}'\mathbf{Z}\mathbf{G}_d + \mathbf{R}'\mathbf{V}^{-1}\mathbf{R})]\mathbf{M}^{-1}$.

Corollary 2. The estimator $\hat{\boldsymbol{\beta}}_{WME}(w)$ has an asymptotic normal distribution with $E[\hat{\boldsymbol{\beta}}_{WME}(w)] = \mathbf{M}_w^{-1}\mathbf{M}_w\boldsymbol{\beta} = \boldsymbol{\beta}$, and covariance matrix $\text{AVar}[\hat{\boldsymbol{\beta}}_{WME}(w)] = \mathbf{M}_w^{-1}[\mathbf{B} + \sigma^2(\mathbf{Z}'\mathbf{Z} + w^2\mathbf{R}'\mathbf{V}^{-1}\mathbf{R})]\mathbf{M}_w^{-1}$.

Corollary 3. The estimator $\hat{\boldsymbol{\beta}}_{ME}(d)$ has an asymptotic normal distribution with $E[\hat{\boldsymbol{\beta}}_{ME}(d)] = \mathbf{M}^{-1}\mathbf{M}\boldsymbol{\beta} = \boldsymbol{\beta}$, and covariance matrix $\text{AVar}[\hat{\boldsymbol{\beta}}_{ME}(d)] = \mathbf{M}^{-1}(\mathbf{B} + \sigma^2\mathbf{M})\mathbf{M}^{-1}$.

2.4 Comparison among biased estimators

When an estimator demonstrates bias, evaluating the performance of an estimator through the MSEM criterion becomes a more suitable approach. In this section, we compare the WMLE to the MLE and WME. The MSEM of the estimator $\hat{\boldsymbol{\beta}}$ in relation to $\boldsymbol{\beta}$ is defined in the following manner:

$$\begin{aligned} \text{MSEM}(\hat{\boldsymbol{\beta}}) &= E[(\hat{\boldsymbol{\beta}} - \boldsymbol{\beta})'(\hat{\boldsymbol{\beta}} - \boldsymbol{\beta})] \\ &= \text{Var}(\hat{\boldsymbol{\beta}}) + \text{Bias}(\hat{\boldsymbol{\beta}})' \text{Bias}(\hat{\boldsymbol{\beta}}) \end{aligned}$$

where $\text{Bias}(\hat{\boldsymbol{\beta}}) = E(\hat{\boldsymbol{\beta}}) - \boldsymbol{\beta}$. There is one more criterion to check the superiority of an estimator over others, that is, using scalar mean squared error (SMSE), which is given by:

$$\text{SMSE}(\hat{\boldsymbol{\beta}}) = \text{tr}[\text{MSEM}(\hat{\boldsymbol{\beta}})] = \text{tr}[\text{Var}(\hat{\boldsymbol{\beta}})] + \text{Bias}(\hat{\boldsymbol{\beta}})' \text{Bias}(\hat{\boldsymbol{\beta}})$$

Let us suppose, we have considered two estimators $\hat{\boldsymbol{\beta}}_1$ and $\hat{\boldsymbol{\beta}}_2$, then the estimator $\hat{\boldsymbol{\beta}}_2$ is superior to $\hat{\boldsymbol{\beta}}_1$ in the MSEM sense if and only if

$$\Delta(\hat{\boldsymbol{\beta}}_1, \hat{\boldsymbol{\beta}}_2) = \text{MSEM}(\hat{\boldsymbol{\beta}}_1) - \text{MSEM}(\hat{\boldsymbol{\beta}}_2) \geq 0$$

If $\Delta > 0$, then $\hat{\boldsymbol{\beta}}_2$ is said to be strongly superior to $\hat{\boldsymbol{\beta}}_1$.

Remark. If $\Delta = \text{MSEM}(\hat{\boldsymbol{\beta}}_1) - \text{MSEM}(\hat{\boldsymbol{\beta}}_2) \geq 0$, then $\text{SMSE}(\hat{\boldsymbol{\beta}}_1) - \text{SMSE}(\hat{\boldsymbol{\beta}}_2) \geq 0$. The inverse reference may not be correct. Thus, MSEM criterion is more robust than the SMSE value (Rao & Toutenburg, 1999). Thus, it may be more appropriate to compare two estimators using the MSEM criterion.

Now, we can write the asymptotic mean squared error matrix (AMSEM) of the estimators $\hat{\boldsymbol{\beta}}_{WMLE}$, $\hat{\boldsymbol{\beta}}_{WME}$ and $\hat{\boldsymbol{\beta}}_{MLE}$ as follows:

$$\text{AMSEM}[\hat{\boldsymbol{\beta}}_{WMLE}(w, d)] = \mathbf{M}_w^{-1}[\mathbf{G}_d\mathbf{B}\mathbf{G}_d + \sigma^2(\mathbf{G}_d\mathbf{Z}'\mathbf{Z}\mathbf{G}_d + w^2\mathbf{R}'\mathbf{V}^{-1}\mathbf{R})]\mathbf{M}_w^{-1} + \mathbf{b}\mathbf{b}'$$

$$AMSEM[\hat{\beta}_{WME}(w)] = \mathbf{M}_w^{-1} [\mathbf{B} + \sigma^2(\mathbf{Z}'\mathbf{Z} + w^2\mathbf{R}'\mathbf{V}^{-1}\mathbf{R})] \mathbf{M}_w^{-1}$$

$$AMSEM[\hat{\beta}_{MLE}(d)] = \mathbf{M}^{-1} [\mathbf{G}_d\mathbf{B}\mathbf{G}_d + \sigma^2(\mathbf{G}_d\mathbf{Z}'\mathbf{Z}\mathbf{G}_d + \mathbf{R}'\mathbf{V}^{-1}\mathbf{R})] \mathbf{M}^{-1} + \mathbf{b}_1\mathbf{b}'_1,$$

where $\mathbf{b}_1 = \mathbf{M}^{-1}(\mathbf{G}_d - \mathbf{I})\mathbf{Z}'\mathbf{Z}\boldsymbol{\beta}$ and $\mathbf{b} = \mathbf{M}_w^{-1}(\mathbf{G}_d - \mathbf{I})\mathbf{Z}'\mathbf{Z}\boldsymbol{\beta}$.

Theorem 2. $\hat{\beta}_{WME}(w, d)$ is superior to the estimator $\hat{\beta}_{WME}(w)$ in MSEM sense, if and only if $\mathbf{b}'\mathbf{H}_1^{-1}\mathbf{b} \leq 1$.

Proof. In order to check the superiority between $\hat{\beta}_{WME}(w, d)$ and $\hat{\beta}_{WME}(w)$, we consider the AMSEM difference:

$$\begin{aligned} D_1 &= AMSEM[\hat{\beta}_{WME}(w)] - AMSEM[\hat{\beta}_{WME}(w, d)] \\ &= \mathbf{M}_w^{-1} [\mathbf{B} + \sigma^2(\mathbf{Z}'\mathbf{Z} + w^2\mathbf{R}'\mathbf{V}^{-1}\mathbf{R})] \mathbf{M}_w^{-1} - \\ &\quad \mathbf{M}_w^{-1} [\mathbf{G}_d\mathbf{B}\mathbf{G}_d + \sigma^2(\mathbf{G}_d\mathbf{Z}'\mathbf{Z}\mathbf{G}_d + w^2\mathbf{R}'\mathbf{V}^{-1}\mathbf{R})] \mathbf{M}_w^{-1} - \mathbf{b}\mathbf{b}' \\ &= \mathbf{M}_w^{-1} [\mathbf{B} - \mathbf{G}_d\mathbf{B}\mathbf{G}_d + \sigma^2(\mathbf{Z}'\mathbf{Z} - \mathbf{G}_d\mathbf{Z}'\mathbf{Z}\mathbf{G}_d)] \mathbf{M}_w^{-1} - \mathbf{b}\mathbf{b}' \\ &= \mathbf{H}_1 - \mathbf{b}\mathbf{b}', \end{aligned} \quad (9)$$

where $\mathbf{H}_1 = \mathbf{M}_w^{-1} [\mathbf{B} - \mathbf{G}_d\mathbf{B}\mathbf{G}_d + \sigma^2(\mathbf{Z}'\mathbf{Z} - \mathbf{G}_d\mathbf{Z}'\mathbf{Z}\mathbf{G}_d)] \mathbf{M}_w^{-1}$.

Now, we have to check if $\mathbf{H}_1$ is positive definite (p.d.) matrix. For $\mathbf{A} = \mathbf{Z}'\mathbf{Z} \geq 0$, there exist some orthogonal matrix $\mathbf{Q}$, such that $\mathbf{A} = \mathbf{Q}'\boldsymbol{\Lambda}\mathbf{Q}$, where $\boldsymbol{\Lambda} = \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_p)$, $\lambda_i \geq 0$. Therefore, we can compute that $\mathbf{Z}'\mathbf{Z} - \mathbf{G}_d\mathbf{Z}'\mathbf{Z}\mathbf{G}_d = \mathbf{Q}'\boldsymbol{\tau}\mathbf{Q} = \mathbf{Q}'\text{diag}(\gamma_1, \dots, \gamma_p)\mathbf{Q}$, where $\boldsymbol{\tau} = (\boldsymbol{\Lambda} - (\boldsymbol{\Lambda} + \mathbf{I})^{-1}(\boldsymbol{\Lambda} + d\mathbf{I}))\boldsymbol{\Lambda}(\boldsymbol{\Lambda} + \mathbf{I})^{-1}(\boldsymbol{\Lambda} + d\mathbf{I})$, for $d \geq 0$, $\lambda_i \geq 0$, so $\gamma_i \geq 0$ which means $\mathbf{Z}'\mathbf{Z} - \mathbf{G}_d\mathbf{Z}'\mathbf{Z}\mathbf{G}_d \geq 0$. Furthermore, $\mathbf{B} - \mathbf{G}_d\mathbf{B}\mathbf{G}_d \geq 0$ and $\mathbf{M}_w^{-1} \geq 0$. This implies that, the matrix $\mathbf{H}_1$ is also a p.d. matrix. Hence, using Lemma 2, we have $\mathbf{H}_1 - \mathbf{b}\mathbf{b}' \geq 0$ if and only if $\mathbf{b}'\mathbf{H}_1^{-1}\mathbf{b} \leq 1$.

Thus, we may infer that the suggested $\hat{\beta}_{WME}(w, d)$ estimator has the potential to outperform the $\hat{\beta}_{WME}(w)$ estimator under specific conditions.

It is seen that the matrix $\mathbf{B} - \mathbf{G}_d\mathbf{B}\mathbf{G}_d$ is a positive definite. Therefore, by applying Lemma 2, we may conclude that $\mathbf{D}_1$ is p.d. matrix if $\sigma^2(\mathbf{Z}'\mathbf{Z} - \mathbf{G}_d\mathbf{Z}'\mathbf{Z}\mathbf{G}_d - (\mathbf{G}_d - \mathbf{I})\mathbf{Z}'\mathbf{Z}\boldsymbol{\beta}\boldsymbol{\beta}'(\mathbf{G}_d - \mathbf{I}))$ is a positive semi-definite matrix. By defining $\boldsymbol{\alpha} = \mathbf{Q}'\boldsymbol{\beta}$, we may express $\boldsymbol{\beta} = \mathbf{Q}\boldsymbol{\alpha}$. For this situation, we may write the expression $\sigma^2(\mathbf{Z}'\mathbf{Z} - \mathbf{G}_d\mathbf{Z}'\mathbf{Z}\mathbf{G}_d - (\mathbf{G}_d - \mathbf{I})\mathbf{Z}'\mathbf{Z}\boldsymbol{\beta}\boldsymbol{\beta}'(\mathbf{G}_d - \mathbf{I})) = \mathbf{Q}\boldsymbol{\tau}\mathbf{Q}' = \mathbf{Q}\text{diag}(\tau_1, \dots, \tau_p)\mathbf{Q}'$, where $\tau_i = \frac{1-d}{(\lambda_i+1)^2} [(\sigma^2\lambda_i + \lambda_i^2\alpha_i^2)d - \lambda_i^2\alpha_i^2 + \sigma^2\lambda_i(2\lambda_i + 1)]$. For $d \leq 1$ and λ_i , we have $\tau \geq 0$ if and only if $d \geq \frac{\lambda_i\alpha_i^2 - \sigma^2(2\lambda_i+1)}{\sigma^2 + \lambda_i\alpha_i^2}$.

Thus, we can select a biasing parameter d within the range of $(\frac{\lambda_i\alpha_i^2 - \sigma^2(2\lambda_i+1)}{\sigma^2 + \lambda_i\alpha_i^2}, 1)$. $\square$

Theorem 3. If $\mathbf{b}'(\mathbf{H}_2 + \mathbf{b}_1\mathbf{b}'_1)\mathbf{b} < 1$, $\hat{\beta}_{WME}(w, d)$ is superior to the estimator $\hat{\beta}_{MLE}(d)$ in MSEM sense, provided $\lambda_{\max}\{AVar(\hat{\beta}_{WME}(w, d))(AVar(\hat{\beta}_{MLE}(d)))^{-1}\} < 1$.

Proof. To check the superiority between $\hat{\beta}_{WME}(w, d)$ and $\hat{\beta}_{MLE}(d)$, we consider the following difference:

$$\begin{aligned} D_2 &= AMSEM[\hat{\beta}_{MLE}(d)] - AMSEM[\hat{\beta}_{WME}(w, d)] \\ &= \mathbf{H}_2 + \mathbf{b}_1\mathbf{b}'_1 - \mathbf{b}\mathbf{b}', \end{aligned}$$

where $\mathbf{H}_2 = AVar(\hat{\beta}_{MLE}(d)) - AVar(\hat{\beta}_{WME}(w, d))$, $\mathbf{b}_1 = \mathbf{M}^{-1}(\mathbf{G}_d - \mathbf{I})\mathbf{Z}'\mathbf{Z}\boldsymbol{\beta}$ and $\mathbf{b} = \mathbf{M}_w^{-1}(\mathbf{G}_d - \mathbf{I})\mathbf{Z}'\mathbf{Z}\boldsymbol{\beta}$.

As it is evident that $AVar(\hat{\beta}_{MLE}(d)) > 0$ and $AVar(\hat{\beta}_{WME}(w, d)) > 0$, it follows that when

$\lambda_{\max}\{AVar(\hat{\beta}_{WME}(w, d))[AVar(\hat{\beta}_{MLE}(d))]^{-1}\} < 1$, we obtain $\mathbf{H}_2 > 0$ (Rao & Toutenburg, 1999).

In addition, $\mathbf{D}_2 \geq 0$ if $\mathbf{b}'(\mathbf{H}_2 + \mathbf{b}_1\mathbf{b}'_1)\mathbf{b} < 1$ (Trenkler & Toutenburg, 1990). $\square$

3. Results and Discussion

3.1 Simulation Study

In this section, we do a simulation study to evaluate the effectiveness of the specified estimators using the Mean Squared Error (MSE) criterion across various parameters employed in the simulation experiment. In order to create varying levels of collinearity, as described in McDonald & Galarneau (1975), the explanatory variables are constructed using the following equation

$$z_{ij} = (1 - \rho^2)^{\frac{1}{2}} v_{ij} + \rho v_{i,p+1}, \quad i = 1, 2, \dots, n, j = 1, 2, \dots, p.$$

where v_{ij} are independent random numbers drawn from standard normal distribution, and ρ denotes the correlation between explanatory variables. Three distinct correlation sets are examined based on the values $\rho = 0.80, 0.90$, and 0.95 . Furthermore, the independent variables are standardized to generate the correlation matrix $\mathbf{Z}'\mathbf{Z}$. In this study, we utilize the values of $p = 8$ and $q = 4$ when the sample size n takes on the values of either 25 or 200. The t^{th} set of simulated data is generated in the following manner:

$$\begin{aligned} \mathbf{y}_t &= \mathbf{Z}\boldsymbol{\beta} + \boldsymbol{\epsilon}_t, \\ \mathbf{X}_t &= \mathbf{Z} + \boldsymbol{\Delta}_t, \quad t = 1, 2, \dots, 5000, \\ \mathbf{r}_t &= \mathbf{R}\boldsymbol{\beta} + \mathbf{e}_t \end{aligned}$$

where $\mathbf{y}_t = (y_{1t}, \dots, y_{nt})'$, $\mathbf{Z} = (\mathbf{z}^{(1)}, \dots, \mathbf{z}^{(j)})'$, $\mathbf{z}^{(j)} = (z_{1j}, \dots, z_{nj})'$, $j = 1, \dots, p$ and $\boldsymbol{\epsilon}_t$ are independent normal $(0, \sigma^2 \mathbf{I}_n)$ pseudo-random numbers. Furthermore, $\mathbf{r}_t = (r_{1t}, \dots, r_{qt})'$, $\mathbf{R} = (\mathbf{R}^{(1)}, \dots, \mathbf{R}^{(p)})'$, $\mathbf{R}^{(j)} = (R_{1j}, \dots, R_{qj})'$, $j = 1, \dots, p$, and $\mathbf{e}_t \sim N(0, \sigma^2 \mathbf{I}_q)$. Also, we assume that $R_{sj} \sim N(0, 1)$, $s = 1, \dots, q, j = 1, \dots, p$, $\sigma^2 = 0.16$ and $\sigma^2 = 0.49$, $\boldsymbol{\Delta} \sim MN(0, \mathbf{I}_n \otimes \boldsymbol{\Sigma})$ and $\boldsymbol{\Sigma} = (0.2\mathbf{I}_p, 0.6\mathbf{I}_p)$, where $\mathbf{I}_p$ is the identity matrix of order p . The weight of auxiliary information, denoted as w , is selected as 0.4, 0.6, 0.8, and 0.1.

In this study, we utilize the vector of regression coefficients that corresponds to the largest eigenvalue of $\mathbf{Z}'\mathbf{Z}$ for each set of explanatory variables. Initially, we computed estimators $\hat{\boldsymbol{\beta}}$ and $\hat{\mathbf{Z}}$. The formula for estimating $\mathbf{Z}$ is $\hat{\mathbf{Z}} = \mathbf{X} + \hat{\sigma}_{\hat{\mathbf{f}}}^{-2} + \hat{\mathbf{f}}\hat{\boldsymbol{\beta}}'\boldsymbol{\Sigma}$ (Rasekh, 2006), where $\hat{f}_i = y_i - \mathbf{x}_i'\hat{\boldsymbol{\beta}}$ and $\hat{\sigma}_{\hat{\mathbf{f}}}^2 = \hat{\sigma}^2 + \hat{\boldsymbol{\beta}}'\boldsymbol{\Sigma}\hat{\boldsymbol{\beta}}$. The variables are generated for each choice of ρ , w , and σ^2 .

The experiment was replicated 5000 times by generating new error terms. The performance of the estimator was assessed by evaluating the estimated MSE values. The MSE is calculated using the formula

$$MSE(\tilde{\boldsymbol{\beta}}) = \frac{1}{5000} \sum_{t=1}^{5000} \sum_{j=1}^p (\tilde{\beta}_{jt} - \beta_j)^2$$

where $\tilde{\beta}_{jt}$ represents the estimate of the j^{th} parameter in the t^{th} replication and $\boldsymbol{\beta}$ is the true parameter value. The findings are succinctly presented in Tables 1–4.

Table 1. Estimated MSE values of the estimators when sample size is 25 and $\Sigma = 0.6I_8$

$\Sigma = 0.6I_8$		n=25			
		w=1	w=0.8	w=0.6	w=0.4
$\rho=0.8 \sigma^2=0.16$					
	MLE	1.020734	1.020364	1.02119	1.020618
	WME	1.025478	1.025784	1.027013	1.027363
	WMLE	1.014748	1.01633	1.019111	1.020618
$\rho=0.8 \sigma^2=0.49$					
	MLE	1.037282	1.038861	1.037804	1.038702
	WME	1.032081	1.034031	1.033846	1.035312
	WMLE	1.031205	1.034695	1.03572	1.038702
$\rho=0.9 \sigma^2=0.16$					
	MLE	1.029828	1.029452	1.03006	1.029708
	WME	1.053256	1.051083	1.049872	1.048835
	WMLE	1.023669	1.025576	1.02798	1.029708
$\rho=0.9 \sigma^2=0.49$					
	MLE	1.059256	1.060249	1.055705	1.058207
	WME	1.060256	1.069384	1.062781	1.06431
	WMLE	1.052928	1.055975	1.053682	1.058207
$\rho=0.95 \sigma^2=0.16$					
	MLE	1.049655	1.049427	1.049729	1.050097
	WME	1.094101	1.089614	1.086041	1.084113
	WMLE	1.043815	1.045587	1.047676	1.050097
$\rho=0.95 \sigma^2=0.49$					
	MLE	0.9762794	1.084912	1.083249	1.083386
	WME	0.9331777	1.10991	1.104796	1.1014
	WMLE	0.9179566	1.080801	1.081163	1.083386

Table 2. Estimated MSE values of the estimators when sample size is 200 and $\Sigma = 0.6I_8$

$\Sigma = 0.6I_8$		n=200			
		w=1	w=0.8	w=0.6	w=0.4
$\rho=0.8 \sigma^2=0.16$					
	MLE	1.029384	1.028822	1.028522	1.028821
	WME	1.030451	1.030289	1.030209	1.030714
	WMLE	1.027627	1.027688	1.027946	1.028821
$\rho=0.8 \sigma^2=0.49$					
	MLE	1.031704	1.031972	1.031421	1.031283
	WME	1.031216	1.031631	1.031356	1.031459
	WMLE	1.029938	1.030784	1.030833	1.031283
$\rho=0.9 \sigma^2=0.16$					
	MLE	1.017865	1.017845	1.01814	1.017889
	WME	1.020625	1.020563	1.020892	1.020703
	WMLE	1.016261	1.016768	1.017602	1.017889
$\rho=0.9 \sigma^2=0.49$					
	MLE	1.02137	1.021896	1.021518	1.022098
	WME	1.021816	1.022498	1.022137	1.022781
	WMLE	1.019692	1.020812	1.02099	1.022098
$\rho=0.95 \sigma^2=0.16$					
	MLE	1.02425	1.024047	1.023843	1.024476
	WME	1.028844	1.02847	1.028076	1.028521
	WMLE	1.023043	1.022721	1.023362	1.024476
$\rho=0.95 \sigma^2=0.49$					
	MLE	1.029202	1.030348	1.02919	1.029239
	WME	1.031268	1.032202	1.030779	1.030689
	WMLE	1.027703	1.029349	1.028665	1.030689

Table 3. Estimated MSE values of the estimators when sample size is 25 and $\Sigma = 0.2I_8$

$\Sigma = 0.2I_8$		n=25			
		w=1	w=0.8	w=0.6	w=0.4
$\rho=0.8 \sigma^2=0.16$					
	MLE	1.009667	1.052974	1.002869	1.003332
	WME	0.969992	0.967226	0.981287	1.003922
	WMLE	0.957275	1.063893	0.990630	1.003332
$\rho=0.8 \sigma^2=0.49$					
	MLE	1.024978	1.060059	1.011679	1.10259
	WME	0.978816	0.970484	0.985234	1.09922
	WMLE	0.961892	1.035163	0.998072	1.10259
$\rho=0.9 \sigma^2=0.16$					
	MLE	0.973106	0.970578	0.971714	0.973432
	WME	0.937451	0.946752	0.950754	0.974498
	WMLE	0.927060	0.936791	0.959611	0.973432
$\rho=0.9 \sigma^2=0.49$					
	MLE	1.012021	0.984547	0.984538	0.979070
	WME	0.957628	0.941881	0.955906	0.977257
	WMLE	0.932255	0.957983	0.970575	0.979070
$\rho=0.95 \sigma^2=0.16$					
	MLE	0.962802	0.956534	0.958363	1.005703
	WME	0.957628	0.932685	0.937208	1.006905
	WMLE	0.932255	0.921860	0.945659	1.005703
$\rho=0.95 \sigma^2=0.49$					
	MLE	0.976279	0.970949	0.964862	1.12496
	WME	0.933177	0.943583	1.126072	0.940604
	WMLE	0.917956	0.928321	0.952165	1.12496

Table 4. Estimated MSE values of the estimators when sample size is 200 and $\Sigma = 0.2I_8$

$\Sigma = 0.2I_8$		n=200			
		w=1	w=0.8	w=0.6	w=0.4
$\rho=0.8 \sigma^2=0.16$					
	MLE	0.956257	0.955708	0.956033	0.955753
	WME	0.950697	0.952208	0.954353	0.956138
	WMLE	0.950445	0.951902	0.954112	0.955753
$\rho=0.8 \sigma^2=0.49$					
	MLE	0.95744	0.957703	0.957790	1.10259
	WME	0.951337	0.953399	0.955338	0.957394
	WMLE	0.951642	0.953843	0.955844	0.957684
$\rho=0.9 \sigma^2=0.16$					
	MLE	0.929153	0.929328	0.929179	0.929268
	WME	0.924235	0.925992	0.927621	0.929632
	WMLE	0.923944	0.929268	0.927433	0.925770
$\rho=0.9 \sigma^2=0.49$					
	MLE	0.931195	0.930594	0.931145	0.930944
	WME	0.925921	0.9270708	0.928951	0.930720
	WMLE	0.925455	0.926725	0.929319	0.930944
$\rho=0.95 \sigma^2=0.16$					
	MLE	0.915426	0.916126	0.915756	0.915732
	WME	0.910945	0.912865	0.914311	0.916094
	WMLE	0.910500	0.912728	0.914047	0.914160
$\rho=0.95 \sigma^2=0.49$					
	MLE	0.917433	0.917348	0.917668	0.917724
	WME	0.912399	0.913721	0.915628	0.917511
	WMLE	0.912171	0.913930	0.915943	0.917724

It is observed in the Table 1 that as σ^2 increases from 0.16 to 0.49, MSE increases in all the six cases. But it is not same when $\rho = 0.95$ and $w = 0.4$. When correlation coefficient increases, there is increase in the values of MSE in all the considered cases. In case of $\sigma^2 = 0.16$, $\rho = 0.90$ and $\rho = 0.95$, MSE of $\hat{\beta}_{WME}(w)$ increases with increase in weight and decreases in other cases.

For Table 2, it is noticed that value of MSE increases when σ^2 increases from 0.16 to 0.49. The MSE value drops almost in every instance when we change the value of ρ from 0.80 to 0.90 or 0.80 to 0.95. An observation is made that the MSE values of $\hat{\beta}_{WME}(w)$ do not follow a fixed pattern when $\rho = 0.95$. In some cases, they are growing, while in other cases, they are shrinking. The $\hat{\beta}_{WMLE}(w, d)$ values, on the other hand, drop in every instance with different weights.

Table 3 and Table 4 depicts that MSE values of $\hat{\beta}_{WME}(w)$ and $\hat{\beta}_{WMLE}(w, d)$ decreases with increase in weight essentially in every instance. There is a drop in MSE for all of the scenarios that were considered when the correlation coefficient increases. Additionally, MSE values exhibits negative correlation with sample size (n), indicating that as n increases, the MSE value of all three estimators decreases.

In the process of comparing Table 1 with Table 2 under various circumstances, we have observed that the MSE reduces for different weights as the sample size grows. The exceptions to this rule are when $\rho = 0.8$ and $\sigma^2 = 0.16$, as well as when $\rho = 0.95$ and $\sigma^2 = 0.49$. It is generally true that the MSE increases in each table when the weights are decreased. There is a reduction in MSE when the value of Σ goes from $0.6I_8$ to $0.2I_8$.

The lowest MSE value is observed when $\rho = 0.95$, $\sigma^2 = 0.49$, $\Sigma = 0.2I_8$ and $n = 200$ while the highest value occurs when $\rho = 0.90$, $\sigma^2 = 0.49$, $\Sigma = 0.6I_8$ and $n = 25$. Overall, the proposed estimator $\hat{\beta}_{WMLE}(w, d)$ is shown to be superior to both the estimators $\hat{\beta}_{MLE}(d)$ and $\hat{\beta}_{WME}(w)$. By reorganizing and clarifying these observations, the analysis becomes more comprehensible, effectively highlighting the trends and relationships within the data.

3.2 Health Club Data

We use the health club data provided in Chatterjee & Hadi (2009) to demonstrate our theoretical results. The data set is derived from the health records of 30 employees of a company’s health club. The variables that have been measured are:
Y = time (in seconds) in a one-mile run,
 X_1 = weight in pounds,
 X_2 = resting pulse rate per minute,
 X_3 = arm and leg strength (number of pounds an employee was able to lift),
 X_4 = time (in seconds) in a 1/4-mile ma1 run.

It was recognized that the observed values were rounded to the nearest whole number, suggesting the presence of measurement error in the dataset. Therefore, it may be inferred that the measurement errors conform to a uniform distribution spanning from -0.5 to 0.5 and exhibit a co-variance matrix of errors $\Sigma = \text{diag}(\frac{1}{12}, \frac{1}{12}, \frac{1}{12}, \frac{1}{12})$. Zhao *et al.* (1994) and Ghapani & Babadi (2022) have also utilized this dataset. In this dataset, the response variable is represented by a vector **y** with dimensions of 30×1 , and the regression matrix **X** has dimensions of 30×4 . The eigen values of $X'X$ are obtained as $\lambda_1 = 2422304.5099$, $\lambda_2 = 10444.4942$, $\lambda_3 = 5075.3011$, and $\lambda_4 = 991.6948$. To evaluate the extent of multicollinearity, we compute the condition number using the following formula:

$$N = \sqrt{\frac{\lambda_1}{\lambda_4}} = 49.42.$$

The condition number value suggests a moderate to strong association within the data set. In accordance with the information provided in Ghapani & Babadi (2022), we have chosen the identical components of **y** and **X** as **r** and **R**, respectively. Additionally, we have the option to select the remaining components of the **y** vector and the corresponding components of the **X** matrix as **r** and **R**, respectively. Next, we employ a linear measurement error model with stochastic linear restrictions to get the estimates of parameters. We substitute the unknown model parameters with the relevant estimators in the corresponding theoretical expression to obtain the estimated MSE values and estimated eigenvalue of **D**. Table 5 reveals a moderate correlation among the variables, suggesting the existence of multicollinearity. There is a strong link between variable X_1 and other variables.

Table 5. Correlations among variables of health club data

	<i>y</i>	<i>X</i> ₁	<i>X</i> ₂	<i>X</i> ₃	<i>X</i> ₄
<i>y</i>	1	0.7989	0.5011	0.4446	0.8480
<i>X</i> ₁	0.7989	1	0.4199	0.7365	0.6431
<i>X</i> ₂	0.5011	0.4199	1	0.0603	0.5388
<i>X</i> ₃	0.4446	0.7365	0.0603	1	0.4001
<i>X</i> ₄	0.8480	0.6431	0.5388	0.4001	1

The estimated MSE values of the specified estimators for various values of weights are presented in Table 6. The estimator $\hat{\beta}_{WMLE}(w, d)$ has a lower estimated MSE value compared to the estimators $\hat{\beta}_{MLE}(d)$ and $\hat{\beta}_{WME}(w)$. Furthermore, it is apparent that as the values of w increase, the estimated MSE values of $\hat{\beta}_{WME}(w)$ and $\hat{\beta}_{WMLE}(w, d)$ decrease.

Table 6. Estimated MSE values of different estimators

	WME	MLE	WMLE
$w = 0.4$	18232.54	18154.35	18132.66
$w = 0.6$	18024.32	18154.35	17913.45
$w = 0.8$	17845.16	18154.35	17468.23
$w = 1.0$	17453.71	18154.35	173162.03

Table 7. Estimated eigenvalues of D

	λ_1	λ_2	λ_3	λ_4
$w = 0.4$	21.6832	0.00119	0.0000156	0.0000000155
$w = 0.6$	20.84236	0.0023425	0.0005691	0.00000428
$w = 0.8$	17.5415	0.0004689	0.00007834	0.00000192
$w = 1.0$	15.1238	0.000123	0.00009235	0.00000843

Table 8. Values of $b'H_1^{-1}b$ for different weights

w	0.4	0.6	0.8	1
$b'H_1^{-1}b$	0.0002145	0.0001768	0.00008636	0.0003940

Table 9. Values of $b'(H_2 + b_1b_1)b$ for different weights

w	0.4	0.6	0.8	1
$b'(H_2 + b_1b_1)b$	0.00005771	0.000096532	0.000038541	0.000026424

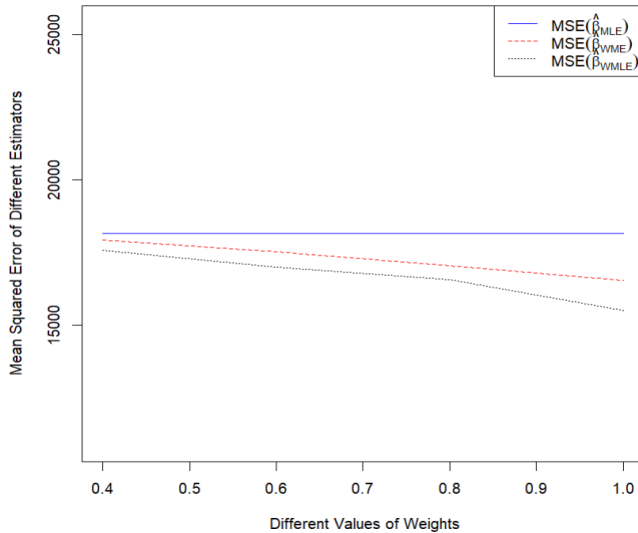


Figure 1. Estimated MSE values of the Estimators.

Table 7 shows that all the estimated eigenvalues of $\mathbf{D}$ are consistently positive for various values of w . The statement suggests that the difference $AMSEM[\hat{\beta}_{WME}(w)] - AMSEM[\hat{\beta}_{WMLE}(w, d)]$ is positive definite. Moreover, in all instances, we have the inequality $\mathbf{b}'\mathbf{H}_1^{-1}\mathbf{b} \leq 1$ as depicted in Table 8. This result demonstrates that the estimator $\hat{\beta}_{WMLE}(w, d)$ is superior to the estimated coefficient $\hat{\beta}_{WME}(w)$ in terms of $MSEM$ sense, which aligns with the theoretical result stated in Theorem 2. In Table 9, it is observed that the criterion $\mathbf{b}'(\mathbf{H}_2 + \mathbf{b}_1\mathbf{b}_1')\mathbf{b} < 1$ is satisfied for various weights. This implies that the estimated coefficient $\hat{\beta}_{WMLE}(w, d)$ is superior to the estimated coefficient $\hat{\beta}_{MLE}(w)$ in the sense of the $MSEM$ which is in accordance with the theoretical result given in Theorem 3. For the purpose of graphically evaluating the superiority, a graph is constructed between the estimated MSE values and various weights. In addition, we have observed that the MSE value of WMLE is reduced in comparison to the values of WME and MLE. Therefore, it can be demonstrated graphically as well that the proposed estimator is superior to the other two estimators mentioned.

4. Conclusions

This article proposes the weighted mixed Liu estimator, denoted as $\hat{\beta}_{WMLE}(w, d)$, for estimating the vector of parameters in a linear measurement error model. The assumption is made that the parameter vector is subject to stochastic linear restrictions. We conducted a thorough analysis of the performance of the WMLE compared to the WME and MLE. Our analysis showed that the WMLE performs better than the WME and MLE based on the $MSEM$ criterion, but only under specified conditions. In conclusion, a simulation study and a numerical example were performed to demonstrate the efficacy of the suggested estimator. The use of graphical representation is also mentioned in order to provide more clarification regarding superiority of the estimators. Moreover, the numerical dataset revealed that when the value of w increases, the MSE values of the WMLE and WME decrease. This implies that acquiring additional dependent prior knowledge enables a more accurate estimation of the parameters.

For future work, one can consider weighted mixed Liu estimator for linear mixed measurement error model.

Acknowledgments

The authors wish to thank the editor and anonymous referees for valuable suggestions which improved the quality of the presentation.

Conflicts of Interest

The authors declare no conflict of interest.

Author Contributions

Conceptualization: TIWARI, M. K. **Methodology:** GOYAL, P. **Formal analysis:** GOYAL, P. **Software analysis:** GOYAL, P. **Validation:** BIST, V.; TIWARI, M. K. **Investigation:** GOYAL, P.; BIST, V. **Data curation:** GOYAL, P. **Writing- original draft:** GOYAL, P. **Writing- review:** Tiwari, M. K.; BIST V. **Writing- editing:** GOYAL, P. **Supervision:** TIWARI, M. K.

References

1. Ahmad, S. & Aslam, M. Another proposal about the new two-parameter estimator for linear regression model with correlated regressors. *Communications in Statistics-Simulation and Computation* **51**, 3054–3072. doi:10.1080/03610918.2019.1705975 (2022).
2. Chatterjee, S. & Hadi, A. S. *Sensitivity analysis in linear regression* doi:10.1002/9780470316764 (John Wiley & Sons, 2009).
3. Durbin, J. A note on regression when there is extraneous information about one of the coefficients. *Journal of the American Statistical Association* **48**, 799–808. doi:10.1080/01621459.1953.10501201 (1953).
4. Farebrother, R. in *Statisticians of the Centuries* 101–104 (Springer, 2001). doi:10.1007/978-1-4613-0179-0_20.
5. Fuller, W. A. *Measurement Error Models* doi:10.1002/9780470316665 (John Wiley & Sons, 2009).
6. Fung, W.-K., Zhong, X.-P. & Wei, B.-C. On estimation and influence diagnostics in linear mixed measurement error models. *American Journal of Mathematical and Management Sciences* **23**, 37–59. doi:10.1080/01966324.2003.10737603 (2003).
7. Ghapani, F., Rasekh, A. & Babadi, B. The weighted ridge estimator in stochastic restricted linear measurement error models. *Statistical Papers* **59**, 709–723. doi:10.1007/s00362-016-0786-3 (2018).
8. Ghapani, F. & Babadi, B. Two parameter weighted mixed estimator in linear measurement error models. *Communications in Statistics-Simulation and Computation* **51**, 6936–6946. doi:10.1080/03610918.2020.1825736 (2022).
9. Hoerl, A. E. & Kennard, R. W. Ridge Regression: Biased Estimation for Nonorthogonal Problems. *Technometrics* **12**, 55–67. doi:10.1080/00401706.1970.10488634 (1970).
10. Kaçiranlar, S. Liu estimator in the general linear regression model. *Journal of Applied Statistical Science* **13**, 229–234 (2003).
11. Kejian, L. A new class of biased estimate in linear regression. *Communications in Statistics-Theory and Methods* **22**, 393–402. doi:10.1080/03610929308831027 (1993).

12. Kuran, Ö. & Özkale, M. R. Gilmour's approach to mixed and stochastic restricted ridge predictions in linear mixed models. *Linear Algebra and Its Applications* **508**, 22–47. doi:10.1016/j.laa.2016.06.040 (2016).
13. Li, Y. & Yang, H. A new ridge-type estimator in stochastic restricted linear regression. *Statistics* **45**, 123–130. doi:10.1080/02331880903573153 (2011).
14. Li, Y. & Yang, H. A new stochastic mixed ridge estimator in linear regression model. *Statistical Papers* **51**, 315–323. doi:10.1007/s00362-008-0169-5 (2010).
15. McDonald, G. C. & Galarneau, D. I. A Monte Carlo evaluation of some ridge-type estimators. *Journal of the American Statistical Association* **70**, 407–416. doi:10.1080/01621459.1975.10479882 (1975).
16. Nakamura, T. Corrected score function for errors-in-variables models: Methodology and application to generalized linear models. *Biometrika* **77**, 127–137. doi:10.1093/biomet/77.1.127 (1990).
17. Owolabi, A. T., Ayinde, K., Idowu, J. I., Oladapo, O. J. & Lukman, A. F. A New Two-Parameter Estimator in the Linear Regression Model with Correlated Regressors. *Journal of Statistics Applications & Probability* **11**, 499–512. doi:10.18576/jsap/110211 (2022).
18. Özkale, M. A stochastic restricted ridge regression estimator. *Journal of Multivariate Analysis* **100**, 1706–1716. doi:10.1016/j.jmva.2009.02.005 (2009).
19. Rao, C. R. *Linear models and generalizations* doi:10.1007/978-3-540-74227-2 (Springer, 2008).
20. Rao, C. R. & Toutenburg, H. in *Linear Models: Least Squares and Alternatives* 97–136 (Springer, 1999). doi:10.1007/0-387-22752-0_3.
21. Rasekh, A. Local influence in measurement error models with ridge estimate. *Computational Statistics & Data Analysis* **50**, 2822–2834. doi:10.1016/j.csda.2005.04.022 (2006).
22. Sarkar, N. A new estimator combining the ridge regression and the restricted least squares methods of estimation. *Communications in Statistics-Theory and Methods* **21**, 1987–2000. doi:10.1080/03610929208830893 (1992).
23. Schaffrin, B. & Toutenburg, H. Weighted mixed regression. *Zeitschrift für Angewandte Mathematik und Mechanik* **70**. Referenceid:1309446, T735–T738 (1990).
24. Theil, H. On the use of incomplete prior information in regression analysis. *Journal of the American Statistical Association* **58**, 401–414. doi:10.1080/01621459.1963.10500854 (1963).
25. Theil, H. & Goldberger, A. S. in *Henri Theil's Contributions to Economics and Econometrics: Econometric Theory and Methodology* 317–332 (Springer, 1992). doi:10.1007/978-94-011-2546-8_18.
26. Trenkler, G. & Toutenburg, H. Mean squared error matrix comparisons between biased estimators—An overview of recent results. *Statistical Papers* **31**, 165–179. doi:10.1007/BF02924687 (1990).
27. Yang, H., Ye, H. & Xue, K. A further study of predictions in linear mixed models. *Communications in Statistics-Theory and Methods* **43**, 4241–4252. doi:10.1080/03610926.2012.725497 (2014).
28. Yavarizadeh, B. & Ahmed, S. E. The weighted ridge estimation for linear mixed models with measurement error under stochastic linear mixed restrictions. *Communications in Statistics-Theory and Methods* **51**, 1605–1621. doi:10.1080/03610926.2021.1927097 (2022).
29. Yavarizadeh, B., Rasekh, A., Ahmed, S. E. & Babadi, B. Ridge estimation in linear mixed measurement error models with stochastic linear mixed restrictions. *Communications in Statistics-Simulation and Computation* **51**, 3037–3053. doi:10.1080/03610918.2019.1705974 (2022).
30. Zhao, Y., Lee, A. H. & Van Hui, Y. Influence diagnostics for generalized linear measurement error models. *Biometrics*, 1117–1128. doi:10.2307/2533448 (1994).

Appendix

Lemma 1. Assume that square matrices A and C are not singular and B and D are matrices with proper orders. Then $(A + BCD)^{-1} = A^{-1} - A^{-1}B(C^{-1} + DA^{-1}B)^{-1}DA^{-1}$ Rao & Toutenburg (1999).

Lemma 2. Farebrother (2001): Let M be a positive definite matrix, namely $M \geq 0$, α be some vector, then $M - \alpha\alpha' \geq 0$ if and only if $\alpha'M^{-1}\alpha \leq 1$.