

Research Paper

A hybrid estimator for the linear measurement error model: Computations and simulations

FATEMEH GHAPANI*

DEPARTMENT OF MATHEMATICS AND STATISTICS, SHOUSHTAR BRANCH, ISLAMIC AZAD
UNIVERSITY, SHOUSHTAR, IRAN

Received: August 21, 2024/ Revised: January 09, 2025/ Accepted: February 12, 2025

Abstract: In this paper, a hybrid estimator is introduced by combining the mixed estimator and the Kibria-Lukman estimator for linear measurement error models in the presence of multicollinearity. The asymptotic properties of the estimators are derived, and their performance is evaluated with respect to the mean square error matrix criterion. A simulation study is conducted to compare the performance of the proposed estimators with other available estimators under different sample sizes, variances, and degrees of multicollinearity. Additionally, a real-life dataset is analyzed to illustrate the findings of the paper.

Keywords: Mean squared error matrix; Modified ridge-type estimator; Multicollinearity; Stochastic linear restrictions.

Mathematics Subject Classification (2010): 62J05, 62J07.

1 Introduction

Consider the linear measurement error model

$$\begin{aligned} y &= Z\beta + \varepsilon, & \varepsilon &\sim N(0, \sigma^2 I_n), \\ X &= Z + U, & U &\sim MN(0, I_n \otimes \Sigma), \end{aligned} \quad (1)$$

where $y = (y_1, \dots, y_n)'$ is an $n \times 1$ vector of response variables, β is a $p \times 1$ vector of unknown parameters and Z is an $n \times p$ matrix of unobservable values of explanatory variables which can be observed through the matrix X with the measurement error U . Furthermore, it is assumed that ε and U are independent and Σ is a $p \times p$ matrix of

*Corresponding author: fatemeh.ghapani@iau.ac.ir

known values or estimable with nonnegative diagonal elements. The corrected score estimates (CSE) of β and σ^2 are given by $\hat{\beta} = (X'X - n\Sigma)^{-1}X'y = S^{-1}X'y$, where $S = X'X - n\Sigma$ and $\hat{\sigma}^2 = \frac{1}{n} \left[(y - X\hat{\beta})'(y - X\hat{\beta}) - n\hat{\beta}'\Sigma\hat{\beta} \right]$, respectively, (Nakamura, 1990). Whenever multicollinearity exists, the $\hat{\beta}$ suffers setbacks by yielding regression coefficients that produce incorrect signs, large standard errors, imprecise confidence intervals, and small t-ratios (Lukman et al., 2019). Some biased estimators have been developed to address the problem of multicollinearity such as the Stein estimator (Stein, 1956), principal components estimator (Massy, 1965), the ordinary ridge regression estimator by Hoerl and Kennard. (1970), and the modified ridge regression estimator (Swindel, 1976; Liu, 1993). Ridge-type estimators have been proposed as an alternative to the CSE. The Ridge estimator (RE) in the linear measurement error (LME) model is defined by Ghapani et al. (2015) as

$$\hat{\beta}_{RE} = (X'X - n\Sigma + kI_p)^{-1}X'y = S_k^{-1}X'y, \quad k > 0,$$

where $S_k = S + kI_p$. In fact, ridge estimator is obtained by solving the following extreme value problem

$$(y - X\beta)'(y - X\beta) - n\beta'\Sigma\beta + k(\beta'\beta - c),$$

in which, c is constant, k is Lagrange multiplier called the biasing parameter.

Another popular technique to overcome the collinearity problem is to consider parameter estimation in addition to the sample information, such as exact or stochastic linear restrictions on the unknown parameters (Rao et al., 2008). In LME models, Shalabh et al. (2007) obtained a consistent estimator of regression coefficients when prior information was available in the form of exact linear restrictions. In the context of multiple linear regression models, additional information in the form of a known covariance matrix of measurement errors associated with explanatory variables and a known matrix of reliability ratios of explanatory variables has been extensively utilized by Shalabh et al. (2009). Additionally, Shalabh et al. (2010) obtained consistent estimation of regression coefficients in an ultrastructural measurement error model using stochastic prior information. Ghapani et al. (2015) introduced a mixed ridge estimator for regression coefficients, and Ghapani and Babadi (2016) introduced a new ridge estimator when additional stochastic linear restrictions on the parameter vector are assumed to hold. Ehsanes Saleh and Shalabh (2020) proposed a class of Liu-type estimators by choosing five quasi-empirical Bayes estimators in the presence of measurement errors in the data. Ziae et al. (2021) presented a ridge estimator in replicated measurement error to overcome the multicollinearity problem.

The incorporation of prior information available in the form of stochastic linear restrictions may provide better estimators than those obtained without such information. In addition to model (1), assume that the vector of parameters β is subject to the following stochastic linear restrictions

$$r = R\beta + e, \quad e \sim N(0, \sigma^2 W), \quad (2)$$

where r is an $m \times 1$ random vector, R is a known $m \times p$ matrix of rank $m < p$, e is an $m \times 1$ vector of unobservable random errors and W is supposed to be known and

positive definite matrix. Also, it is assumed that ε is stochastically independent of e . With the use of restricted estimation method, the mixed estimator (ME) in LME models introduced by Ghapani et al. (2015) as

$$\hat{\beta}_{ME} = (X'X + R'W^{-1}R - n\Sigma)^{-1}(X'y + R'W^{-1}r) = S_r^{-1}(X'y + R'W^{-1}r),$$

where $S_r = S + R'W^{-1}R$. In recent years, biased estimation and estimation methods with prior information have often been combined to form a broader biased estimation. When prior information and sample information are not equally important, Ghapani et al. (2018) introduced a weighted mixed ridge estimator for regression coefficients in LME models. Additionally, in LME models, the mixed Liu estimator was obtained by Ghapani and Babadi. (2018), and a two-parameter estimator along with a two-parameter weighted mixed estimator were introduced by Ghapani and Babadi. (2022a). Wu (2021) proposed an alternative mixed Liu estimator in LME models with stochastic linear restrictions. The primary aim of this paper is to introduce a hybrid estimator in LME models by combining the Kibria-Lukman (KL) estimation procedure and the mixed estimation procedure.

The paper is organized as follows: In Section 2, the KL and mixed Kibria-Lukman (MKL) estimators are defined for the vector of parameters in LME models. The asymptotic normal properties of estimators are presented in Section 3. The performance of the obtained estimators is examined with respect to mean square error matrix (MSEM) criterion in Section 4. The simulation results are discussed in Section 5 and a numerical example is provided in Section 6. Finally, some conclusions are presented in Section 7.

2 The proposed estimator

Following Kibria and Lukman (2020) consider the following equation in the LME model

$$\Phi_1 = (y - X\beta)'(y - X\beta) - n\beta'\Sigma\beta + k \left[(\beta + \hat{\beta})'(\beta + \hat{\beta}) - c \right].$$

Differentiating the function Φ_1 with respect to β and setting the result to zero, the KL estimator is given by

$$\hat{\beta}_{KL} = (X'X - n\Sigma + kI_p)^{-1}(X'y - k\hat{\beta}) = S_k^{-1}(X'y - k\hat{\beta}).$$

Letting $F_k = (S + kI_p)^{-1}(S - kI_p)$. So, $\hat{\beta}_{KL}$ can be expressed as

$$\hat{\beta}_{KL} = F_k S^{-1} X'y = F_k \hat{\beta}.$$

In LME model the mixed ridge estimator (MRE) is defined by Ghapani et al. (2015) as

$$\hat{\beta}_{MRE} = (X'X + R'W^{-1}R - n\Sigma + kI_p)^{-1}(X'y + R'W^{-1}r) = S_{rk}^{-1}(X'y + R'W^{-1}r),$$

where $S_{rk} = S_r + kI_p$. In order to incorporate the restrictions (2) in the KL estimator, the following equation is considered

$$\Phi_2 = (y - X\beta)'(y - X\beta) - n\beta'\Sigma\beta + k \left[(\beta + \hat{\beta}_{ME})'(\beta + \hat{\beta}_{ME}) - c \right] + (r - R\beta)'W^{-1}(r - R\beta).$$

By differentiating Φ_2 with respect to β and setting the result to zero, the following relationship is obtained

$$X'X\beta - X'y - n\Sigma\beta + k(\beta + \hat{\beta}_{ME}) + R'W^{-1}R\beta - R'W^{-1}r = 0.$$

The proposed estimator is denoted by $\hat{\beta}_{MKL}$ and is obtained as follows

$$\hat{\beta}_{MKL} = (X'X - n\Sigma + R'W^{-1}R + kI_p)^{-1}(X'y - k\hat{\beta}_{ME} + R'W^{-1}r).$$

$\hat{\beta}_{MKL}$ may be written as

$$\begin{aligned}\hat{\beta}_{MKL} &= (S + R'W^{-1}R + kI_p)^{-1}(X'y - k\hat{\beta}_{ME} + R'W^{-1}r) \\ &= S_r^{-1}(S_r + kI_p)^{-1}(S_r - kI_p)(X'y + R'W^{-1}r).\end{aligned}$$

By considering $F_{rk} = (S_r + kI_p)^{-1}(S_r - kI_p)$, $\hat{\beta}_{MKL}$ is written as follows

$$\hat{\beta}_{MKL} = F_{rk}S_r^{-1}(X'y + R'W^{-1}r) = F_{rk}\hat{\beta}_{ME}. \quad (3)$$

The following properties can be obtained

- When $k = 0$, then $\hat{\beta}_{MKL(k=0)} = S_r^{-1}(X'y + R'W^{-1}r) = \hat{\beta}_{ME}$, which is ME of β .
- When $R = 0$, then $\hat{\beta}_{MKL(R=0)} = S_k^{-1}(X'y - k\hat{\beta}) = \hat{\beta}_{KL}$, which is KL estimate of β .
- When $k = 0$ and $R = 0$ then $\hat{\beta}_{MKL(k=0, R=0)} = S^{-1}X'y = \hat{\beta}$, which is CSE of β .

3 Asymptotic properties

To investigate the asymptotic behavior of estimators, we utilize large sample asymptotic approximation theory to analyze their asymptotic distribution. It is assumed that all derivatives related to the log-likelihood exist and that the parameter β is identifiable. Additionally, as n approaches to infinity, the limits of $n^{-1}(Z'Z + R'W^{-1}R)$, $n^{-1}(Z'Z + R'W^{-1}R + kI_p)$ and $n^{-1}(Z'Z + R'W^{-1}R - kI_p)$ expressions are assumed to exist, and, E denotes the global expectation taken at the true value.

Theorem 3.1. *Under the above assumptions $\hat{\beta}_{MKL}$ is asymptotically normal with correspondent mean and covariance matrix as*

$$\begin{aligned}E(\hat{\beta}_{MKL}) &= G_{rk}\beta, \\ AVar(\hat{\beta}_{MKL}) &= A_r^{-1}(G_{rk}BG_{rk} + \sigma^2G_{rk}A_rG_{rk})A_r^{-1},\end{aligned}$$

where $B = (n\sigma^2 + \beta'Z'Z\beta)\Sigma$, $G_{rk} = (Z'Z + R'W^{-1}R + kI_p)^{-1}(Z'Z + R'W^{-1}R - kI_p)$ and $A_r = Z'Z + R'W^{-1}R$.

Proof. Since $E(X'X) = Z'Z + n\Sigma$, (Fung et al., 2003), then $X'X = Z'Z + n\Sigma + O_p(n^{\frac{1}{2}})$. Therefore, it can be

$$n^{-1}(X'X + R'W^{-1}R) = n^{-1}(Z'Z + R'W^{-1}R) + \Sigma + O_p(n^{-\frac{1}{2}}). \quad (4)$$

Additionally, $F_{rk} = G_{rk} + O_p(n^{-\frac{1}{2}})$. Therefore, it follows from (3) and (4) that

$$\sqrt{n}\hat{\beta}_{MKL} = \left[n^{-1}(Z'Z + R'W^{-1}R) + O_p(n^{-\frac{1}{2}}) \right]^{-1} n^{-\frac{1}{2}}$$

$$\begin{aligned} & \times \left[G_{rk}(X'y + R'W^{-1}r) + O_p(n^{\frac{1}{2}}) \right] \\ & = \left[I_p + O_p(n^{-\frac{1}{2}}) \right]^{-1} C^{-1} n^{-\frac{1}{2}} \left[G_{rk}(X'y + R'W^{-1}r) + O_p(n^{\frac{1}{2}}) \right]. \end{aligned}$$

From Taylor series expansion, $\left[I_p + O_p(n^{-\frac{1}{2}}) \right]^{-1} = \left[I_p + O_p(n^{-\frac{1}{2}}) \right]$ can be obtained and since the limit of $C = n^{-1}(Z'Z + R'W^{-1}R)$ exists, then the above Equation can be written as follows:

$$\sqrt{n}\hat{\beta}_{MKL} = C^{-1}\xi + O_p(n^{-\frac{1}{2}}), \quad (5)$$

where $\xi = n^{-\frac{1}{2}} [G_{rk}(X'y + R'W^{-1}r)]$ is asymptotically normal (Fung et al., 2003). Therefore, from $E [G_{rk}(X'y + R'W^{-1}r)] = G_{rk}A_r\beta$, the relation $E(\xi) = n^{-\frac{1}{2}}G_{rk}A_r\beta$ is obtained. Consequently, $\sqrt{n}(\hat{\beta}_{MKL} - G_{rk}\beta) = C^{-1}[\xi - E(\xi)] + O_p(n^{-\frac{1}{2}})$, which indicates that $\sqrt{n}(\hat{\beta}_{MKL} - G_{rk}\beta)$ has an asymptotic normal distribution with mean vector zero. Furthermore, it follows from (5) that $AVar(\sqrt{n}\hat{\beta}_{MKL}) = C^{-1}Var(\xi)C^{-1}$. The variance of ξ can be obtained from

$$\begin{aligned} Var(\xi) &= E_{y_r} [Var(\xi | y_r)] + Var_{y_r} [E(\xi | y_r)] \\ &= n^{-1}E_{y_r}(G_{rk}y'y\Sigma G_{rk}) + n^{-1}Var_{y_r} [G_{rk}(Z'y + R'W^{-1}r)], \end{aligned}$$

where E_{y_r} and Var_{y_r} denote the expectation and variance with respect to the random vector $y_r' = (y', r')$. Since $E(y'y) = n\sigma^2 + \beta'Z'Z\beta$ and $Var_{y_r} [G_{rk}(Z'y + R'W^{-1}r)] = \sigma^2G_{rk}A_rG_{rk}$, therefore, $Var(\xi) = n^{-1}(G_{rk}BG_{rk} + \sigma^2G_{rk}A_rG_{rk})$, and $AVar(\hat{\beta}_{MKL}) = A_r^{-1}(G_{rk}BG_{rk} + \sigma^2G_{rk}A_rG_{rk})A_r^{-1}$. $\square$

According to Theorem 3.1, the following results are obtained

- $\hat{\beta}$ has asymptotic normal distribution with $E(\hat{\beta}) = \beta$ and covariance matrix $AVar(\hat{\beta}) = A^{-1}(B + \sigma^2A)A^{-1}$, where $A = Z'Z$.
- $\hat{\beta}_{RE}$ has asymptotic normal distribution with $E(\hat{\beta}_{RE}) = A_k^{-1}A\beta$ and covariance matrix $AVar(\hat{\beta}_{RE}) = A_k^{-1}(B + \sigma^2A)A_k^{-1}$, where $A_k = Z'Z + kI_p$.
- $\hat{\beta}_{KL}$ has asymptotic normal distribution with $E(\hat{\beta}_{KL}) = G_k\beta$ and covariance matrix $AVar(\hat{\beta}_{KL}) = A^{-1}(G_kBG_k + \sigma^2G_kAG_k)A^{-1}$, where $G_k = (Z'Z + kI_p)^{-1}(Z'Z - kI_p)$.
- $\hat{\beta}_{ME}$ has asymptotic normal distribution with $E(\hat{\beta}_{ME}) = \beta$ and covariance matrix $AVar(\hat{\beta}_{ME}) = A_r^{-1}(B + \sigma^2A_r)A_r^{-1}$.
- $\hat{\beta}_{MRE}$ has asymptotic normal distribution with $E(\hat{\beta}_{MRE}) = A_{rk}^{-1}A_r\beta$ and covariance matrix $AVar(\hat{\beta}_{MRE}) = A_{rk}^{-1}(B + \sigma^2A_r)A_{rk}^{-1}$, where $A_{rk} = (Z'Z + R'W^{-1}R + kI_p)$.

4 Mean square error matrix comparisons

One of the criteria used to evaluate the performance of estimators is taken to be the MSEM criterion. Note that for any estimator $\hat{\beta}$ of β , the MSEM is defined as $MSEM(\hat{\beta}) = E(\hat{\beta} - \beta)(\hat{\beta} - \beta)' = Var(\hat{\beta}) + Bias(\hat{\beta})Bias(\hat{\beta})'$, where $Bias(\hat{\beta}) = E(\hat{\beta}) - \beta$ denotes the bias vector. Another criterion for evaluating an estimator is the MSE value which is defined as follows

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

Definition 4.1. If $\Delta = MSEM(\hat{\beta}_1) - MSEM(\hat{\beta}_2) \geq 0$ then $MSE(\hat{\beta}_1) - MSE(\hat{\beta}_2) \geq 0$. The reverse conclusion may not be correct. Thus, MSEM criterion is stronger than the MSE value (Rao et al., 2008). Therefore, comparing two estimators in the sense of MSEM criterion can be more suitable.

The asymptotic MSEM of the estimators $\hat{\beta}$, $\hat{\beta}_{KL}$, $\hat{\beta}_r$ and $\hat{\beta}_{MKL}$ obtained as follows

$$AMSEM(\hat{\beta}) = A^{-1}(B + \sigma^2 A)A^{-1},$$

$$AMSEM(\hat{\beta}_{RE}) = A_k^{-1}(B + \sigma^2 A)A_k^{-1} + bb', \quad (6)$$

$$AMSEM(\hat{\beta}_{KL}) = A^{-1}(G_k B G_k + \sigma^2 G_k A G_k)A^{-1} + b_1 b_1', \quad (7)$$

$$AMSEM(\hat{\beta}_{ME}) = A_r^{-1}(B + \sigma^2 A_r)A_r^{-1},$$

$$AMSEM(\hat{\beta}_{MRE}) = A_{rk}^{-1}(B + \sigma^2 A_r)A_{rk}^{-1} + b_2 b_2', \quad (8)$$

$$AMSEM(\hat{\beta}_{MKL}) = A_r^{-1}(G_{rk} B G_{rk} + \sigma^2 G_{rk} A_r G_{rk})A_r^{-1} + b_3 b_3', \quad (9)$$

where, $b = -kA_k^{-1}\beta$, $b_1 = -2kA_k^{-1}\beta$, $b_2 = -kA_{rk}^{-1}\beta$ and $b_3 = -2kA_{rk}^{-1}\beta$.

The following lemmas will be used to make some theoretical comparisons among estimators in the following section.

Lemma 4.2. Let M be a positive definite matrix, namely $M > 0$, be some vector, then $M - \alpha\alpha' \geq 0$ if and only if $\alpha'M^{-1}\alpha \leq 1$. (Farebrother, 1976)

Lemma 4.3. Let the two $n \times n$ matrices $M > 0$, $N \geq 0$, then $M > N$ if and only if $\lambda_{\max}(NM^{-1}) < 1$. (Rao et al., 2008)

Lemma 4.4. Let $\tilde{\beta}_j = A_j y$, $j = 1, 2$ be two competing homogeneous linear estimators of β . Suppose that $D = \text{Cov}(\tilde{\beta}_1) - \text{Cov}(\tilde{\beta}_2) > 0$, where $\text{Cov}(\tilde{\beta}_j)$, $j = 1, 2$ denotes the covariance matrix of $\tilde{\beta}_j$. Then $\Delta(\tilde{\beta}_1, \tilde{\beta}_2) = MSEM(\tilde{\beta}_1) - MSEM(\tilde{\beta}_2) \geq 0$ if and only if $d_2'(D + d_1 d_1')^{-1} d_2 \leq 1$, where $MSEM(\tilde{\beta}_j)$, d_j , $j = 1, 2$ denote the MSEM and bias vector of $\tilde{\beta}_j$, respectively. (Trenkler and Toutenburg, 1990)

4.1 Comparison between $\hat{\beta}_{KL}$ and $\hat{\beta}$

In order to compare $\hat{\beta}_{KL}$ and $\hat{\beta}$ in the MSEM criterion, the matrix difference between $AMSEM(\hat{\beta})$ and $AMSEM(\hat{\beta}_{KL})$ is considered as

$$\Delta_1 = AMSEM(\hat{\beta}) - AMSEM(\hat{\beta}_{KL}) = D_1 - b_1 b_1',$$

where $D_1 = 2kA_k^{-1}(A^{-1}B + BA^{-1} + 2\sigma^2 I_p)A_k^{-1}$. When $A_k > 0$ and $A > 0$ then $A_k^{-1} > 0$ and $A^{-1} > 0$. Furthermore $A^{-1}B$ and BA^{-1} have the same eigenvalues. The roots of the equation $|B - \lambda A| = 0$ are the eigenvalues of $A^{-1}B$. Can be written $0 = |B - \lambda A| = |A^{1/2}|^2 |A^{-1/2}BA^{-1/2} - \lambda I_p|$ and $\lambda_i(A^{-1}B) = \lambda_i(A^{-1/2}BA^{-1/2}) > 0$, since $A^{-1/2}BA^{-1/2} > 0$. Therefore $D_1 > 0$ is a p.d. matrix. Hence according to lemma 4.2 stated above, the following Theorem is concluded.

Theorem 4.5. The estimator $\hat{\beta}_{KL}$ is superior to the estimator $\hat{\beta}$ in the MSEM sense if and only if $b_1' D_1^{-1} b_1 \leq 1$.

4.2 Comparison between $\hat{\beta}_{KL}$ and $\hat{\beta}_{RE}$

From (6) and (7) we make

$$\Delta_2 = AMSEM(\hat{\beta}_{RE}) - AMSEM(\hat{\beta}_{KL}) = D_2 + bb' - b_1b'_1,$$

where $D_2 = AVar(\hat{\beta}_{RE}) - AVar(\hat{\beta}_{KL})$. It is obvious that $AVar(\hat{\beta}_{RE}) > 0$ and $AVar(\hat{\beta}_{KL}) > 0$. Therefore, when $\lambda_{\max} \left[AVar(\hat{\beta}_{KL})AVar(\hat{\beta}_{RE})^{-1} \right] < 1$, by applying lemma 4.3, $D_2 > 0$ is obtained. Thus, according to lemma 4.4, the following Theorem is obtained.

Theorem 4.6. *When $\lambda_{\max} \left[AVar(\hat{\beta}_{KL})AVar(\hat{\beta}_{RE})^{-1} \right] < 1$, the $\hat{\beta}_{KL}$ is superior to the estimator $\hat{\beta}_{RE}$ in the MSEM sense, if and only if $b'_1(D_2 + bb')^{-1}b_1 \leq 1$.*

4.3 Comparison between $\hat{\beta}_{MKL}$ and $\hat{\beta}_{ME}$

In order to compare MSEM values between $\hat{\beta}_{MKL}$ and $\hat{\beta}_{ME}$, the following difference is considered

$$\Delta_3 = AMSEM(\hat{\beta}_{ME}) - AMSEM(\hat{\beta}_{MKL}) = D_3 - b_3b'_3,$$

where $D_3 = AVar(\hat{\beta}_{ME}) - AVar(\hat{\beta}_{MKL}) = 2kA_{rk}^{-1}(A_r^{-1}B + BA_r^{-1} + 2\sigma^2I_p)A_{rk}^{-1}$. This implies that $D_3 > 0$ is clearly a p.d. matrix. Hence according to lemma 4.2, it is concluded the following Theorem.

Theorem 4.7. *The necessary and sufficient conditions for $\hat{\beta}_{MKL}$ to be superior to $\hat{\beta}_{ME}$ under MSEM sense as $b'_3D_3^{-1}b_3 \leq 1$.*

4.4 Comparison between $\hat{\beta}_{MKL}$ and $\hat{\beta}_{MRE}$

From (8) and (9), it can be written

$$\Delta_4 = AMSEM(\hat{\beta}_{MRE}) - AMSEM(\hat{\beta}_{MKL}) = D_4 + b_2b'_2 - b_3b'_3,$$

where $D_4 = AVar(\hat{\beta}_{MRE}) - AVar(\hat{\beta}_{MKL})$. When $\lambda_{\max} \left[AVar(\hat{\beta}_{MKL})AVar(\hat{\beta}_{MRE})^{-1} \right] < 1$, by applying lemma 4.3, it follows that $D_4 > 0$. Furthermore, according to lemma 4.4, the findings are summarized in the following Theorem.

Theorem 4.8. *The necessary and sufficient conditions for $\hat{\beta}_{MKL}$ to be superior to $\hat{\beta}_{ME}$ under MSEM sense as $b'_3(D_4 + b_2b'_2)^{-1}b_3 \leq 1$.*

4.5 Comparison between $\hat{\beta}_{MKL}$ and $\hat{\beta}_{KL}$

In order to compare $\hat{\beta}_{MKL}$ with $\hat{\beta}_{KL}$ in the MSEM sense, the following difference is considered

$$\Delta_5 = AMSEM(\hat{\beta}_{KL}) - AMSEM(\hat{\beta}_{MKL}) = D_5 + b_1b'_1 - b_3b'_3,$$

where $D_5 = AVar(\hat{\beta}_{KL}) - AVar(\hat{\beta}_{MKL})$. When $\lambda_{\max} \left[AVar(\hat{\beta}_{MKL})AVar(\hat{\beta}_{KL})^{-1} \right] < 1$, according to lemma 4.3, $D_5 > 0$ is obtained. So, using lemma 4.4 the following Theorem is conclude.

Theorem 4.9. *When $\lambda_{\max} \left(AVar(\hat{\beta}_{MKL})AVar(\hat{\beta}_{KL})^{-1} \right) < 1$, the $\hat{\beta}_{MKL}$ is superior to the estimator $\hat{\beta}_{KL}$ in the MSEM sense, if and only if $b'_3(D_5 + b_1b'_1)^{-1}b_3 \leq 1$.*

From the above theorems, it can be concluded that the $\hat{\beta}_{MKL}$ can perform better than the $\hat{\beta}_{KL}$ under certain conditions.

4.6 Determination of parameter k

An important issue in the study of ridge regression is determining an appropriate parameter k . In this section, the selection of parameter k is discussed based on the MSEM criterion. In this method, the ridge parameter k is derived such that the matrix Δ_1 is positive definite. The matrix Δ_1 can be expressed as follows

$$\Delta_1 = 2kA_k^{-1}(A^{-1}B + BA^{-1} + 2\sigma^2I_p - 2k\beta'\beta)A_k^{-1},$$

when $A_k^{-1} > 0$, $A^{-1}B > 0$ and $BA^{-1} > 0$, then, Δ_1 is a p.d. matrix if $2\sigma^2I_p - 2k\beta'\beta$ is p.s.d matrix. Thus a sufficient condition is $k < \frac{\sigma^2}{\beta'\beta}$. For practical purpose, σ^2 and β are replaced with the suitable estimates. Then, the estimated value of k is obtained as $\hat{k} < \frac{\hat{\sigma}^2}{\hat{\beta}'\hat{\beta}}$.

5 Simulation study

In this section, to further investigate the behavior of the estimators, a Monte Carlo simulation study was conducted to compare the performances of mentioned estimators. For this purpose, The j th data set was generated from the simulated data as follows

$$\begin{aligned} y_j &= Z\beta + \varepsilon_j, \\ X_j &= Z + U_j, \quad j = 1, \dots, 2000, \\ r_j &= R\beta + e_j, \end{aligned}$$

where, $y_j = (y_{1j}, \dots, y_{nj})'$, $Z = (z^{(1)}, z^{(2)}, z^{(3)})$, $z^{(l)} = (z_{1l}, \dots, z_{nl})'$, $l = 1, 2, 3$ and $\varepsilon_j \sim N(0, \sigma^2 I_n)$ is rewritten in accordance with y_j . Furthermore, $r_j = (r_{1j}, \dots, r_{mj})'$, $R = (R^{(1)}, R^{(2)}, R^{(3)})$, $R^{(l)} = (R_{1l}, \dots, R_{ml})'$, $l = 1, 2, 3$ and $e_j \sim N(0, \sigma^2 I_m)$ is rewritten in accordance with r_l . Also, it is assumed that $R_{sl} \sim N(0, 1)$, $s = 1, \dots, m$, $l = 1, 2, 3$, $\sigma^2 = 0.5$, $\sigma^2 = 1$, $U_l \sim MN(0, I_n \otimes \Sigma)$ and $\Sigma = \text{diag}(0.05, 0.05, 0.05)$. To achieve different degrees of collinearity, following McDonald and Galarneau. (1975), the explanatory variables are generated using the following

$$z_{il} = (1 - \rho^2)^{\frac{1}{2}} w_{il} + \rho w_{i,p+1}, \quad i = 1, \dots, n, \quad l = 1, \dots, p,$$

where w_{ij} are independent standard normal random numbers and ρ is specified so that the correlation between any two explanatory variables is given by ρ^2 . The values of ρ are chosen to be 0.75, 0.80, 0.85, 0.90 and 0.95 to show the weakly to severely collinearity between the explanatory variables, respectively. The data are standardized to obtain $Z'Z$ in correlation form. In this experiment, $n = 30, 60$ or 100 , $p = 3$ and $m = 1$, were choose, respectively. For each set of explanatory variables the coefficient

vector corresponding to the largest eigenvalue $Z'Z$ was considered. First the estimators $\hat{\beta}$, $\hat{\sigma}^2$ and $\hat{Z}$ are calculated. An estimate of Z can be derived as, $\hat{Z} = X + \hat{\sigma}_v^{-2} \hat{v} \hat{\beta}' \Sigma$ (Ghapani et al., 2015), where $\hat{v}_i = y_i - x_i' \hat{\beta}$ and $\hat{\sigma}_v^2 = \hat{\sigma}^2 + \hat{\beta}' \Sigma \hat{\beta}$. The replication of this simulation study is 2000 times for different sample sizes and σ^2 . The simulation study was carried out using the R software. For each replicate, the estimated mean square error MSE values of the estimators by using the equation below

$$\text{MSE}(\tilde{\beta}) = \frac{1}{2000} \sum_{j=1}^{2000} \sum_{l=1}^3 (\tilde{\beta}_{lj} - \beta_l)^2,$$

where $\tilde{\beta}_{lj}$ denotes the estimate of the l th parameter in j th replication and β is the true parameter values. The results are shown in Table 1.

Table 1: Estimated MSE values of the proposed estimators.

n	ρ	σ^2	$\hat{\beta}$	$\hat{\beta}_{RE}$	$\hat{\beta}_{KL}$	$\hat{\beta}_{ME}$	$\hat{\beta}_{MRE}$	$\hat{\beta}_{MKL}$
30	0.75	5	0.7832	0.5723	0.4552	0.7784	0.5696	0.4536
		10	1.5369	1.0100	0.8143	1.5276	1.0058	0.8115
	0.80	5	0.9286	0.6590	0.5170	0.9214	0.6552	0.5147
		10	1.8132	1.1609	0.9393	1.8001	1.1548	0.9352
	0.85	5	1.1858	0.8185	0.6325	1.1733	0.8122	0.6287
		10	2.2976	1.4421	1.1637	2.2740	1.4319	1.1566
	0.90	5	1.7162	1.1457	0.8875	1.6885	1.1320	0.8791
		10	3.2845	2.0258	1.6754	3.2328	2.0035	1.6885
	0.95	5	3.4317	2.3445	1.8941	3.3164	2.2809	1.8486
		10	6.3886	4.1813	3.5798	6.1815	4.0737	3.4936
60	0.75	5	0.4980	0.3888	0.3142	0.4965	0.3878	0.3136
		10	0.9724	0.6765	0.5260	0.9694	0.6750	0.5251
	0.80	5	0.6066	0.4606	0.3646	0.6043	0.4592	0.3637
		10	1.1771	0.7969	0.6141	1.1728	0.7946	0.6128
	0.85	5	0.7902	0.5811	0.4507	0.7864	0.5788	0.4493
		10	1.5200	1.0013	0.7708	1.5127	0.9976	0.7686
	0.90	5	1.1573	0.8251	0.6332	1.1491	0.8203	0.6303
		10	2.1954	1.4237	1.1133	2.1803	1.4161	1.1085
	0.95	5	2.2214	1.6158	1.2932	2.1921	1.5977	1.2808
		10	4.1096	2.8309	2.3447	2.0569	2.8009	2.3228
100	0.75	5	0.2701	0.2319	0.2012	0.2693	0.2313	0.2008
		10	0.5292	0.4136	0.3356	0.5277	0.4127	0.3351
	0.80	5	0.3288	0.2749	0.2330	0.3276	0.2741	0.2325
		10	0.6406	0.4839	0.3841	0.6383	0.4827	0.3834
	0.85	5	0.4291	0.3483	0.2876	0.4270	0.3469	0.2867
		10	0.8293	0.6061	0.4716	0.8254	0.6041	0.4705
	0.90	5	0.6325	0.4958	0.3980	0.6280	0.4930	0.3963
		10	1.2061	0.8551	0.6569	1.1978	0.8508	0.6545
	0.95	5	1.2083	0.9207	0.7352	1.1932	0.9113	0.7292
		10	2.2495	1.5927	1.2680	2.2223	1.5778	1.2582

Based on Table 1 the following conclusions are drawn:

- The estimated MSE values of all estimates increase with the increase of the levels of collinearity and σ^2 .
- The estimated MSE values of the estimators decrease whenever n increase.
- $\hat{\beta}_{RE}$ and $\hat{\beta}_{KL}$ estimator uniformly dominate the $\hat{\beta}$ estimator.
- In all cases, the new estimator MKL has smaller estimated MSE values than the other

existing estimators. Therefore, we can conclude that the proposed estimator performs better than the other estimators based on estimated MSE values. Additionally, the proposed estimator can also perform better than the other estimators under MSEM criterion in cases where the conditions of the theorems mentioned in the previous section are met.

More details can be seen in Figure 1, where all simulation results are summarized in box plots. These results are for the case $n = 100$, $\sigma^2 = 10$ and $\rho = 0.90$. The horizontal lines represent the true parameter values. An estimator is considered more accurate if its median (center line of the box) is close to the horizontal line. A shorter box indicates more consistent estimates.

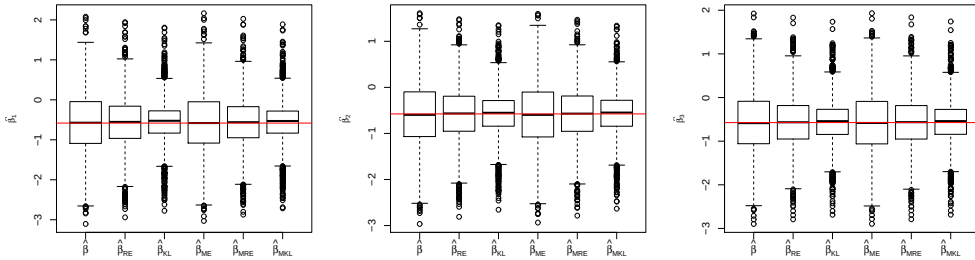


Figure 1: Boxplot presentations of the estimated parameters for $n = 100$, $\sigma^2 = 10$ and $\rho = 0.90$ from 2000 simulated data sets.

Fewer and less extreme outliers suggest better robustness. Based on the plot, $\hat{\beta}_{KL}$ and $\hat{\beta}_{MKL}$ performs the best, as their median is closest to the horizontal line, they have a relatively small spread, and fewer extreme outliers compared to the others. Furthermore, based on the boxplot when there are stochastic linear restrictions on the parameter vector, the MKL estimator has better performance than the other estimators.

6 Example: Health club data

To illustrate theoretical results, the health club data provided in Chatterjee and Hadi (1988) are considered. The data set is derived from the health records of 30 employees who were regular members of a company's health club. The data include four regression variables: weight in pounds (X_1), resting pulse rate per minute (X_2), arm and leg strength (number of pounds an employee was able to lift) (X_3) and time (in seconds) in a $\frac{1}{4}$ -mile mal run (X_4). The response (y) is the weight in pounds. The authors admitted the problem of measurement errors in the data due to rounding off the actual observational values to the nearest integer unit. Then it may be reasoned that the measurement errors are uniformly distributed over the interval $(-0.5, 0.5)$ with $\Sigma = \text{diag}(\frac{1}{12}, \frac{1}{12}, \frac{1}{12}, \frac{1}{12})$ as covariance matrix of errors. This data set has also been studied by Zhao and Lee (1994), Ghapani and Babadi. (2022a) and Ghapani and Babadi. (2022b). For this data set, the scaled condition number is $\sqrt{\frac{\lambda_{\max}}{\lambda_{\min}}} = 49.42$, where $\lambda_{\max}$

and $\lambda_{\min}$ are the largest and smallest eigenvalues of $\hat{Z}'\hat{Z}$ which indicates moderate to strong correlation in the data set. Consider the 23th element of y vector and 23th row of X matrix as r and R , respectively. Because of having the most time, the 23th element of y vector is arbitrary chosen. But, the other elements of vector and corresponding elements of the X matrix can also be taken as r and R , respectively. Then we fit a linear measurement error models with stochastic linear restrictions and obtain the estimates of parameters.

To obtain the estimated MSE values replace the unknown model parameters with the appropriate estimators in the corresponding theoretical expressions. The estimated MSE values and standard error of mentioned estimators are presented in Table 2. It can be seen from Table 2 that the proposed estimator performed the best in the sense of smaller MSE. Also, for more convenience, plot of the estimated MSE values of all estimators versus k given in interval $[0, 50]$ with increment 0.1 is given in Figure 2. The MSE values of $\hat{\beta}_{RE}$, $\hat{\beta}_{KL}$, $\hat{\beta}_{MRE}$ and $\hat{\beta}_{MKL}$ decreases when k increases.

Table 2: Estimated MSE values and standard error (Sd) of the of the mentioned estimators.

	$\hat{\beta}$	$\hat{\beta}_{RE}$	$\hat{\beta}_{KL}$	$\hat{\beta}_{ME}$	$\hat{\beta}_{MRE}$	$\hat{\beta}_{MKL}$
MSE values	1.2686	1.2028	1.1712	1.1377	1.0847	1.0619
Sd	2.0983	2.0321	1.9662	2.0402	1.9801	1.9203

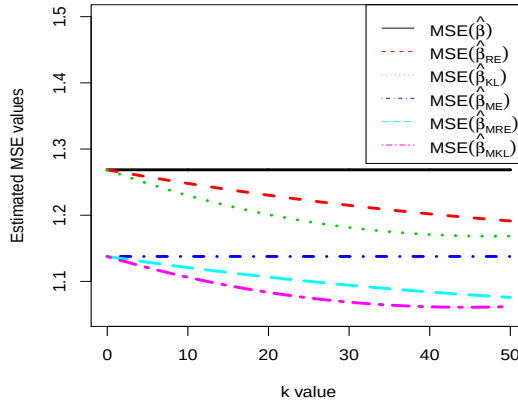


Figure 2: Estimated MSE values of the estimators.

Furthermore, from the Figure 2, it is obvious that the mentioned estimators can perform better than the $\hat{\beta}$ in MSEM criterion.

7 Summary and concluding remarks

To address the problem of multicollinearity, this paper proposes the estimators $\hat{\beta}_{KL}$ and $\hat{\beta}_{MKL}$ for the vector of parameters in a linear measurement error model, building on

the research of Kibria and Lukman (2020). The performance of the estimators has been evaluated using the asymptotic MSEM criterion. A simulation study and a numerical example were conducted to compare the performance of the estimators and support the theoretical analysis. The simulation results indicate that as the sample size increases, the MSE values of the estimators decrease, even in the presence of high correlation among regressors. Overall, with increasing levels of collinearity and σ^2 , the estimated MSE values of the different estimators also increase. From the numerical example results, it appears that when k increases, the estimated MSE values of $\hat{\beta}_{RE}$, $\hat{\beta}_{KL}$ and $\hat{\beta}_{MKL}$ decrease. Additionally, the simulation and numerical example results indicate that the proposed estimators provide better outcomes than the other estimators under certain conditions.

Acknowledgement

The author would like to express gratitude to the editor and the anonymous referees for their valuable comments and suggestions, which have contributed to a significant improvement in the presentation of this paper.

References

- Chatterjee, S. and Hadi, A.S. (1988). *Sensitivity Analysis in Linear Regression*. New York: Wiley.
- Ehsanes Saleh, A.M. and Shalabh, (2020). On Liu-type biased estimators in measurement error models. *A Journal of Theoretical and Applied Statistics*, **54**(6):1171–1213.
- Farebrother, R.W. (1976). Further results on the mean square error of ridge regression. *Journal of the Royal Statistical Society. Series B (Methodological)*, **38**(3):248–250.
- Fung, W.K., Zhong, X.P. and Wei, B.C. (2003). On estimation and influence diagnostics in linear mixed measurement error models. *American Journal of Mathematical and Management Sciences*, **23**(1-2):37–59.
- Ghapani, F., Rasekh, A.R., Akhoond, M.R. and Babadi, B. (2015). Detection of outliers and influential observations in linear ridge measurement error models with stochastic linear restrictions. *Journal of Sciences Islamic Republic of Iran*, **26**(4):355–366.
- Ghapani, F. and Babadi, B. (2016). A new ridge estimator in linear measurement error model with stochastic linear restrictions. *Journal of the Iranian Statistical Society*, **15**(2):87–103.
- Ghapani, F. and Babadi, B. (2018). Mixed Liu estimator in linear measurement error models. *Communications in Statistics-Theory and Methods*, **47**(7):1561–1570.
- Ghapani, F., Rasekh, A.R. and Babadi, B. (2018). The weighted ridge estimator in stochastic restricted linear measurement error models. *Statistical Papers*, **59**(2):709–723.

- Ghapani, F. and Babadi, B. (2022a). Two parameter weighted mixed estimator in linear measurement error models. *Communications in Statistics-Simulation and Computation.*, **51**(12):6936–6946.
- Ghapani, F. and Babadi, B. (2022b). Diagnostic measures for the restricted Liu estimator in linear measurement error models. *Journal of Statistical Computation and Simulation*, **92**(11):2257–2272.
- Hoerl, A.E. and Kennard, R.W. (1970). Ridge regression: biased estimation for non-orthogonal problems. *Technometrics*, **12**(1):69–82.
- Kibria, B.M.G. and Lukman, A.F. (2020). A new ridge-type estimator for the linear regression model. Simulations and applications. *Scientific Scientifica*, **2020**(1):9758378.
- Liu, K. (1993). A new class of biased estimate in linear regression. *Communications in Statistics-Theory and Methods*, **22**(2):393–402.
- Lukman, A.F., Ayinde, K., Binuomote, S. and Onate, A.C. (2019). Modified ridge-type estimator to combat multicollinearity. *Journal of Chemometrics*, **33**(5):e3125.
- Massy, W.F. (1965). Principal components regression in exploratory statistical research. *Journal of the American Statistical Association*, **60**:234–256.
- McDonald, C. and Galarneau, D.A. (1975). A Monte Carlo evaluation of some Ridge-type estimators. *Journal of the American Statistical Association*, **70**(350):407–416.
- Nakamura, T. (1990). Corrected score function for errors-in-variables models: Methodology and application to generalized linear models. *Biometrika*, **77**(1):127–137.
- Rao, C.R., Toutenburg, H., Shalabh and Heumann, C. (2008). *Linear Models and Generalizations*. Berlin: Springer.
- Shalabh., Sh., Gaurav, G. and Misra, N. (2007). Restricted regression estimation in measurement error models. *Computational Statistics Data Analysis*, **52**(2):1149–1166.
- Shalabh., Sh., Gaurav, G. and Misra, N. (2009). Use of prior information in the consistent estimation of regression coefficients in measurement error models. *Journal of Multivariate Analysis*, **100**(7):1498–1520.
- Shalabh, Sh., Gaurav, G. and Misra, N. (2010). Consistent estimation of regression coefficients in ultrastructural measurement error model using stochastic prior information. *Statistical Papers*, **51**(3):717–748.
- Stein, C. (1956). Inadmissibility of the usual estimator for mean of multivariate normal distribution. *Proceedings of the Third Berkeley Symposium on Mathematical Statistics and Probability*, **3**(1):197–206.
- Swindel, B.F. (1976). Good ridge estimators based on prior information. *Communications in Statistics-Theory and Methods*, **5**(11):1065–1075.

- Trenkler, G., Toutenburg, H. (1990). Mean squared error matrix comparisons between biased estimators: an overview of recent results. *Statistical Papers*, **31**:165–179.
- Wu, J. (2021). The mixed Liu estimator in stochastic restricted linear measurement error mode. *Journal of Mathematics*, **2021**(1):6624923.
- Zhao, Y. and Lee, A.H. (1994). Influence diagnostics for generalized linear measurement error models. *Biometrics*, **50**(4):1117–1128.
- Ziae, A.R., Zare, K. and Sheikhi, R. (2021). Performance of Ridge regression approach in linear measurement error models with replicated data. *Journal of Mathematical Extension*, **15**(4):1–23.