

SOME SHRINKAGE ESTIMATORS FOR ESTIMATING THE STANDARD DEVIATION AND ITS INVERSE FOR NORMAL PARENT

Housila P. Singh¹ and Vankim Chander²

ABSTRACT

Shrunk estimator has its importance in the small sample theory when appropriate prior information about the unknown parameter is supposed to be known. The present paper investigates classes of shrunk estimators for estimating standard deviation (σ) and its inverse (σ^{-1}) in the case of univariate normal parent. The need to study these parameters arises due to their importance to estimate different parametric functions such as Mean Deviation about Mean ($\sigma\sqrt{\frac{2}{\pi}}$), Process Capability Index (C_p), Standard Deviation (σ), Standard

Error of Mean $\left(\frac{\sigma}{\sqrt{n}}\right)$, Coefficient of Variation (C.V) with known mean etc.,

using the prior information or guessed value of σ . Simulation studies confirm the high efficiency of the developed classes of shrunk estimators when compared with their usual unbiased estimators and minimum mean squared error (MMSE) estimators.

Key words: Bias, Gauss-Laguerre integration method, Mean Squared Error, Normal parent, Percent Relative Efficiency (PRE), Prior information.

1. Introduction

The normal distribution has dominated statistical practice as well as theory. Any variable whose expression results from the additive contribution of many small effects will tend to be normally distributed. For measurements whose distributions are not normal, a simple transformation of the scale of measurement may induce approximate normality. The square root, and logarithm $\ln(x)$ are often

¹ hpsujn@rediffmail.com : School of Studies in Vikram University, Ujjain, MP, India 456010.

² vcupadhyay06@yahoo.co.in : SAHARA Arts and Management Academy, Lucknow, UP, India.

used as such transformations, so normal distribution is very important in statistical point of views.

1.1.The Model

Let $x_1, x_2, \dots, x_n$ be a random sample of size n , drawn from population $N(\mu, \sigma^2)$, the probability distribution function (pdf) of which is given by

$$f(x; \mu, \sigma^2) = \begin{cases} \frac{1}{\sigma\sqrt{2\pi}} \exp\left\{-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2\right\}, & -\infty < x < \infty, -\infty < \mu < \infty, \sigma > 0 \\ 0, & \text{otherwise.} \end{cases} \quad (1.1)$$

μ is the mean and σ is the standard deviation.

1.2.Classical Estimators

It is desired to estimate standard deviation (S.D) σ and its inverse σ^{-1} . The conventional unbiased estimators of σ and σ^{-1} are respectively given by

$$T_{(u,\sigma)} = c_1(n)s \quad (1.2)$$

and

$$T_{(u,\sigma^{-1})} = c_2(n)s^{-1}, \quad (1.3)$$

where

$$c_1(n) = \left(\frac{n-1}{2}\right)^{\frac{1}{2}} \frac{\Gamma\left(\frac{(n-1)}{2}\right)}{\Gamma\left(\frac{n}{2}\right)}, \quad (1.4)$$

$$c_2(n) = \left(\frac{n-1}{2}\right)^{-\frac{1}{2}} \frac{\Gamma\left(\frac{(n-1)}{2}\right)}{\Gamma\left(\frac{(n-2)}{2}\right)} \quad (1.5)$$

and

$$s^2 = \frac{1}{(n-1)} \sum_{i=1}^n (x_i - \bar{x})^2 \quad \text{with} \quad \bar{x} = \frac{1}{n} \sum_{i=1}^n x_i.$$

The variance of $T_{(u,\sigma)}$ and $T_{(u,\sigma^{-1})}$ are respectively given by

$$\text{Var}\{T_{(u,\sigma)}\} = v_1(n)\sigma^2 \quad (1.6)$$

and

$$\text{Var}\{T_{(u,\sigma^{-1})}\} = v_2(n)(1/\sigma^2) \quad (1.7)$$

where

$$v_1(n) = [(c_1(n))^2 - 1] \quad , \quad (1.8)$$

$$v_2(n) = [c_3(n) - 1] \quad (1.9)$$

and

$$c_3(n) = \frac{\Gamma\left(\frac{n-3}{2}\right)\Gamma\left(\frac{n-1}{2}\right)}{\Gamma^2\left(\frac{n-2}{2}\right)} \quad . \quad (1.10)$$

The minimum mean squared error (MMSE) estimator in the class $Q_b=bs$ (b , being suitably chosen constant such that mean squared error (MSE) of Q_b is minimum) is defined by

$$T_{(MMSE,\sigma)} = \frac{s}{c_1(n)} \quad , \quad (1.11)$$

with bias and MSE

$$\text{Bias}\{T_{(MMSE,\sigma)}\} = -\sigma \left\{ \frac{v_1(n)}{v_1(n)+1} \right\} \quad (1.12)$$

and

$$\text{MSE}\{T_{(MMSE,\sigma)}\} = \sigma^2 \left\{ \frac{v_1(n)}{v_1(n)+1} \right\} \quad (1.13)$$

respectively.

From (1.6) and (1.13), we have

$$\text{Var}\{T_{(u,\sigma)}\} - \text{MSE}\{T_{(MMSE,\sigma)}\} = \left\{ \frac{(v_1(n))^2}{v_1(n)+1} \right\} \sigma^2 > 0 \quad , \quad (1.14)$$

which gives the inequality

$$Var\{T_{(u,\sigma)}\} > MSE\{T_{(MMSE,\sigma)}\}. \quad (1.15)$$

Thus the MMSE estimator $T_{(MMSE,\sigma)}$ is more efficient than unbiased estimator $T_{(u,\sigma)}$.

Further the MMSE estimator in the class of estimators $d_{M^*} = \frac{M^*}{S}$ (M^* , being suitably chosen constant such that MSE of d_{M^*} is least) is given by

$$T_{(MMSE,\sigma^{-1})} = c_4(n) \left(\frac{1}{s} \right), \quad (1.16)$$

$$\text{where } c_4(n) = \sqrt{\left(\frac{2}{n-1} \right)} \frac{\Gamma\left(\frac{n-2}{2}\right)}{\Gamma\left(\frac{n-3}{2}\right)}.$$

The bias and MSE are respectively given by

$$Bias\{T_{(MMSE,\sigma^{-1})}\} = -\frac{1}{\sigma} \left\{ \frac{v_2(n)}{v_2(n) + 1} \right\} \quad (1.17)$$

and

$$MSE\{T_{(MMSE,\sigma^{-1})}\} = \frac{1}{\sigma^2} \left\{ \frac{v_2(n)}{v_2(n) + 1} \right\}. \quad (1.18)$$

From (1.10) and (1.18) we have

$$Var\{T_{(u,\sigma^{-1})}\} - MSE\{T_{(MMSE,\sigma^{-1})}\} = \frac{1}{\sigma^2} \left\{ \frac{\{v_2(n)\}^2}{v_2(n) + 1} \right\} > 0, \quad (1.19)$$

which yields the inequality

$$MSE\{T_{(MMSE,\sigma^{-1})}\} < Var\{T_{(u,\sigma^{-1})}\}. \quad (1.20)$$

Thus, the MMSE estimator $T_{(MMSE,\sigma^{-1})}$ is more efficient than unbiased estimator $T_{(u,\sigma^{-1})}$. Experimenter often possesses some knowledge of experimental conditions based on the acquaintance with the behaviour of the system under consideration or from past experience or from some extraneous source, and is thus in a position to give an adequate guess or an initial estimate, say θ_0 of the value of the unknown parameter θ . Thompson (1968 a, b) developed shrinkage estimator for the mean, Pandey and Singh (1977 a) studied the performance of the shrinkage estimator of the variance and many more authors Pandey (1979), Jani

(1991), Gangele and Singh (1994), Kourouklis (1994), Singh and Singh (1997), Singh et al. (1999, 2004), Singh and Saxena (2003 a, b, c, d) , Joshi (2005) , Singh and Chander (2007 a, b, c, d, e) have introduced shrinkage estimators for the parameters of normal/non-normal parents using prior knowledge of the unknown parameters.

In the present paper we have investigated some classes of shrinkage estimators for estimating standard deviation (σ) and its inverse (σ^{-1}) in the case of univariate normal parent. The problem of estimating these parameters arises due to their importance in estimating different parametric functions such as Mean Deviation about Mean $\left\{ \sigma \sqrt{\left(\frac{2}{\pi}\right)} \right\}$, Process Capability Index (C_p), Standard Deviation (σ), Standard Error of Mean $\left\{ \frac{\sigma}{\sqrt{n}} \right\}$, Coefficient of Variation (C.V) with known mean etc., using the prior information or guessed value of σ .

Simulation studies confirm the high efficiency of the developed classes of shrunk estimators when compared with their usual unbiased estimators and (MMSE) estimators.

2. The proposed class of shrinkage estimators for S.D. σ

Let σ_0 be the guessed value of σ available from the past experience. Then we define a class of estimators for σ as

$$\begin{aligned} T_{(s,\sigma)} &= \ddot{k} T_{(u,\sigma)} + (1 - \ddot{k}) \sigma_0 \\ &= \ddot{k} (T_{(u,\sigma)} - \sigma_0) + \sigma_0 \quad , \end{aligned} \quad (2.1)$$

where $\ddot{k}$ is a scalar such that $0 \leq \ddot{k} \leq 1$.

The bias and MSE of $T_{(s,\sigma)}$ are respectively given by

$$Bias\{T_{(s,\sigma)}\} = \{E(T_{(s,\sigma)}) - \sigma\} = \sigma(1 - \ddot{k})(\hat{\lambda} - 1) \quad (2.2)$$

and

$$\begin{aligned} MSE\{T_{(s,\sigma)}\} &= Var\{T_{(s,\sigma)}\} + \{Bias\{T_{(s,\sigma)}\}\}^2 \\ &= \ddot{k}^2 \sigma^2 [(c_1(n))^2 - 1] + \sigma^2 (1 - \ddot{k})^2 (\hat{\lambda} - 1)^2 \\ &= \sigma^2 [\ddot{k}^2 [(c_1(n))^2 - 1] + (1 - \ddot{k})^2 (\hat{\lambda} - 1)^2] \end{aligned}$$

$$= \sigma^2 [\ddot{k}^2 v_1(n) + (1 - \ddot{k})^2 (\hat{\lambda} - 1)^2] , \quad (2.3)$$

where $\hat{\lambda} = \frac{\sigma_0}{\sigma}$.

From (1.6) and (2.3), we have

$$\begin{aligned} Var\{T_{(u,\sigma)}\} - MSE\{T_{(s,\sigma)}\} &= \sigma^2 [v_1(n) - \ddot{k}^2 v_1(n) - (1 - \ddot{k})^2 (\hat{\lambda} - 1)^2] \\ &= \sigma^2 [v_1(n)(1 - \ddot{k}^2) - (1 - \ddot{k})^2 (\hat{\lambda} - 1)^2] , \end{aligned}$$

which is positive if

$$[v_1(n)(1 - \ddot{k}^2) - (1 - \ddot{k})^2 (\hat{\lambda} - 1)^2] > 0 ,$$

$$\text{i.e. if } [v_1(n)(1 - \ddot{k}) - (1 - \ddot{k})(\hat{\lambda} - 1)^2] > 0 , (1 - \ddot{k}) > 0$$

$$\text{i.e. if } \ddot{k} [v_1(n) + (\hat{\lambda} - 1)^2] > [(\hat{\lambda} - 1)^2 - v_1(n)]$$

$$\text{i.e. if } \ddot{k} > \left\{ \frac{(\hat{\lambda} - 1)^2 - v_1(n)}{(\hat{\lambda} - 1)^2 + v_1(n)} \right\}$$

$$\text{i.e. } \left\{ \frac{(\hat{\lambda} - 1)^2 - v_1(n)}{(\hat{\lambda} - 1)^2 + v_1(n)} \right\} < \ddot{k} \leq 1 . \quad (2.4)$$

Further, from (1.13) and (2.3), we have

$$MSE\{T_{(MMSE,\sigma)}\} - MSE\{T_{(s,\sigma)}\} = \sigma^2 \left\{ \frac{v_1(n)}{1 + v_1(n)} \right\}$$

$$- \ddot{k}^2 v_1(n) - (1 - \ddot{k})^2 (\hat{\lambda} - 1)^2 > 0$$

$$\text{if } \ddot{k}^2 \left\{ (1 - \hat{\lambda}^2) + v_1(n) \right\} - 2\ddot{k}(\hat{\lambda} - 1)^2 + (\hat{\lambda} - 1)^2 - \left\{ \frac{v_1(n)}{1 + v_1(n)} \right\} > 0$$

$$\text{i.e. } (a - b) < \ddot{k} < (a + b) \quad (2.5)$$

where

$$a = \frac{(\hat{\lambda} - 1)^2}{v_1(n) + (\hat{\lambda} - 1)^2},$$

$$b = \frac{v_1(n)\sqrt{\hat{\lambda}(2 - \hat{\lambda})}}{\sqrt{\{1 + v_1(n)\}}\{v_1(n) + (\hat{\lambda} - 1)^2\}}, 0 < \hat{\lambda} < 2.$$

2.1. Numerical Illustrations

To see the performance of the proposed estimator $T_{(s,\sigma)}$ over the usual unbiased estimator $T_{(u,\sigma)}$ and MMSE estimator $T_{(MMSE,\sigma)}$, we have computed the percentage relative efficiencies (PRE's) of $T_{(s,\sigma)}$ with respect to $T_{(u,\sigma)}$ and $T_{(MMSE,\sigma)}$ using the formulae:

$$PRE\{T_{(s,\sigma)}, T_{(u,\sigma)}\} = \frac{MSE\{T_{(u,\sigma)}\}}{MSE\{T_{(s,\sigma)}\}} * 100 = \frac{Var\{T_{(u,\sigma)}\}}{MSE\{T_{(s,\sigma)}\}} * 100$$

$$= \frac{1}{\ddot{k}^2 + \{v_1(n)\}^{-1}(1 - \ddot{k})^2(\hat{\lambda} - 1)^2} * 100 \quad (2.6)$$

and

$$PRE\{T_{(s,\sigma)}, T_{(MMSE,\sigma)}\} = \frac{MSE\{T_{(MMSE,\sigma)}\}}{MSE\{T_{(s,\sigma)}\}} * 100$$

$$= \frac{v_1(n)}{\{v_1(n)\ddot{k}^2 + (1 - \ddot{k})^2(\hat{\lambda} - 1)^2\}\{1 + v_1(n)\}} * 100, \quad (2.7)$$

for different values of $\ddot{k}$, $\hat{\lambda}$, n and findings are displayed in Table 2.2. It is observed from the Table 2.2 that:

- the proposed estimator $T_{(s,\sigma)}$ is better than the usual unbiased estimator $T_{(u,\sigma)}$ and the MMSE estimator $T_{(MMSE,\sigma)}$ with substantial gain in efficiency when $\hat{\lambda} \in [0.70, 1.30]$ and $\ddot{k} \in [0.25, 0.75]$.
- larger gain in efficiency is observed when the sample size n is small.

- (c) for fixed $\ddot{k}$ and n , the value of PRE decreases as $\hat{\lambda}$ goes away from unity.
- (d) for fixed $\ddot{k}$ and $\hat{\lambda}$, the value of PRE decreases as n increases.
- (e) the gain in efficiency by using $T_{(s,\sigma)}$ over MMSE estimator $T_{(MMSE,\sigma)}$ is fewer than by using $T_{(s,\sigma)}$ over the unbiased estimator $T_{(u,\sigma)}$.

Thus, the proposed estimator is more beneficial when σ_0 , the guessed value is in the vicinity of the true value of σ and when the sample size n is small. In practice, when the observations are expensive, such small sample sizes may be all that are available.

We have also computed the range of $\ddot{k}$ (for different values of n and $\hat{\lambda}$) in which the suggested estimator $T_{(s,\sigma)}$ is better than $T_{(u,\sigma)}$ and $T_{(MMSE,\sigma)}$ using the expressions (2.4) and (2.5) respectively and the results are tabled in Table 2.1. It is observed from Table 2.1 that the effective range of $\ddot{k}$ become shorter as sample size n goes up for fixed $\hat{\lambda}$, also for fixed n the effective range becomes shorter as $\hat{\lambda}$ goes away from unity.

Table 2.1 Range of $\ddot{k}$ for which the proposed estimator $T_{(s,\sigma)}$ is better than $T_{(u,\sigma)}$ and MMSE estimator $T_{(MMSE,\sigma)}$

$\hat{\lambda} \downarrow n \rightarrow$	3	5	7	9	11
0.3	(0.28, 1.0) (0.42, 0.87)	(0.58, 1.0) (0.65, 0.93)	(0.70, 1.00) (0.75, 0.95)	(0.77, 1.00) (0.80, 0.96)	(0.81, 1.00) (0.84, 0.97)
0.5	(0.00, 1.00) (0.08, 0.88)	(0.31, 1.00) (0.37, 0.94)	(0.49, 1.00) (0.53, 0.96)	(0.59, 1.00) (0.62, 0.97)	(0.66, 1.00) (0.69, 0.97)
0.7	(0.00, 1.00) (0.00, 0.88)	(0.00, 1.00) (0.00, 0.94)	(0.02, 1.00) (0.06, 0.96)	(0.17, 1.00) (0.20, 0.97)	(0.27, 1.00) (0.30, 0.97)
0.8	(0.00, 1.00) (0.00, 0.89)	(0.00, 1.00) (0.00, 0.94)	(0.00, 1.00) (0.00, 0.96)	(0.00, 1.00) (0.00, 0.97)	(0.00, 1.00) (0.00, 0.98)
0.9	(0.00, 1.00) (0.00, 0.89)	(0.00, 1.00) (0.00, 0.94)	(0.00, 1.00) (0.00, 0.96)	(0.00, 1.00) (0.00, 1.00)	(0.00, 1.00) (0.00, 0.98)
1	(0.00, 1.00) (0.00, 0.89)	(0.00, 1.00) (0.00, 0.94)	(0.00, 1.00) (0.00, 0.96)	(0.00, 1.00) (0.00, 1.00)	(0.00, 1.00) (0.00, 0.98)
1.1	(0.00, 1.00) (0.00, 0.89)	(0.00, 1.00) (0.00, 0.94)	(0.00, 1.00) (0.00, 0.96)	(0.00, 1.00) (0.00, 1.00)	(0.00, 1.00) (0.00, 0.98)
1.2	(0.00, 1.00) (0.00, 0.89)	(0.00, 1.00) (0.00, 0.94)	(0.00, 1.00) (0.00, 0.96)	(0.00, 1.00) (0.00, 1.00)	(0.00, 1.00) (0.00, 0.98)
1.3	(0.00, 1.00) (0.00, 0.88)	(0.00, 1.00) (0.00, 0.94)	(0.02, 1.00) (0.06, 0.96)	(0.17, 1.00) (0.20, 0.97)	(0.27, 1.00) (0.30, 0.97)
1.5	(0.00, 1.00) (0.08, 0.88)	(0.31, 1.00) (0.37, 0.94)	(0.49, 1.00) (0.53, 0.96)	(0.59, 1.00) (0.62, 0.97)	(0.66, 1.00) (0.69, 0.97)
1.7	(0.28, 1.00) (0.42, 0.87)	(0.58, 1.00) (0.65, 0.93)	(0.70, 1.00) (0.75, 0.95)	(0.77, 1.00) (0.80, 0.96)	(0.81, 1.00) (0.84, 0.97)
1.9	(0.50, 1.00) (0.65, 0.85)	(0.72, 1.00) (0.80, 0.92)	(0.81, 1.00) (0.86, 0.94)	(0.85, 1.00) (0.90, 0.96)	(0.88, 1.00) (0.92, 0.97)
2	(0.57, 1.00) n.e.	(0.77, 1.00) n.e.	(0.84, 1.00) n.e.	(0.88, 1.00) n.e.	(0.90, 1.00) n.e.

Note : a) The **bold** figures in parentheses show the range of $\ddot{k}$ for which the proposed estimator $T_{(s,\sigma)}$ is better than MMSE estimator $T_{(MMSE,\sigma)}$. b) **n.e** stands for does not exit.

Table 2.2 PRE of $T_{(s,\sigma)}$ with respect to unbiased estimator $T_{(u,\sigma)}$ and MMSE estimator $T_{(MMSE,\sigma)}$, for different values of $n = 2, 3, 5, 7, 9$; $\ddot{k} = 0.25$ (0.25) 0.75 and $\hat{\lambda}$

$\ddot{k}$	$\hat{\lambda}$ $\downarrow n \rightarrow$	PRE of $T_{(s,\sigma)}$ w. r. t $T_{(u,\sigma)}$					PRE of $T_{(s,\sigma)}$ w. r. t $T_{(MMSE,\sigma)}$				
		2	3	5	7	9	2	3	5	7	9
0.25	0.3	183.36	93.35	46.42	30.78	23	116.73	73.32	41.02	28.33	21.61
	0.5	323.76	173.26	88.52	59.23	44.47	206.11	136.08	78.21	54.52	41.78
	0.7	661.41	403.59	223.87	154.37	117.71	421.07	316.98	197.8	142.08	110.6
	0.8	981.18	690.39	428.72	309.97	242/55	624.64	542.23	378.8	285.28	227.89
	0.9	1382.08	1203.57	950.67	784.12	666.9	879.86	945.28	839.99	721.7	626.59
	1	1600	1600	1600	1600	1600	1018.59	1256.64	1413.72	1472.62	1503.3
	1.1	1382.08	1203.57	950.67	784.12	666.9	879.86	945.28	839.99	721.7	626.59
	1.2	981.18	690.39	428.72	309.97	242/55	624.64	542.23	378.8	285.28	227.89
	1.3	661.41	403.59	223.87	154.37	117.71	421.07	316.98	197.8	142.08	110.6
	1.5	323.76	173.26	88.52	59.23	44.47	206.11	136.08	78.21	54.52	41.78
	1.7	183.36	93.35	46.42	30.78	23	116.73	73.32	41.02	28.33	21.61
	Range of λ	[0.06,1.94]	[0.33,1.67]	[0.54,1.46]	[0.63,1.37]	[0.68,1.32]	[0.24,1.76]	[0.41,1.59]	[0.57,1.43]	[0.64,1.36]	[0.69,1.31]
0.5	0.3	215.23	143.2	84.77	60.02	46.42	137.02	112.47	74.9	55.24	43.61
	0.5	278.17	208.89	138.06	102.82	81.86	177.09	164.06	121.99	94.64	76.91
	0.7	345.52	300.89	237.67	196.03	166.72	219.97	236.32	210	180.42	156.65
	0.8	373.8	348.92	306.85	273.52	246.63	237.97	274.04	271.13	251.74	231.73
	0.9	393.11	385.88	371.78	358.55	346.18	250.26	303.07	328.5	330	325.26
	1	400	400	400	400	400	254.65	314.16	353.43	368.16	375.83
	1.1	393.11	385.88	371.78	358.55	346.18	250.26	303.07	328.5	330	325.26
	1.2	373.8	348.92	306.85	273.52	246.63	237.97	274.04	271.13	251.74	231.73
	1.3	345.52	300.89	237.67	196.03	166.72	219.97	236.32	210	180.42	156.65
	1.5	278.17	208.89	138.06	102.82	81.86	177.09	164.06	121.99	94.64	76.91
	1.7	215.23	143.2	84.77	60.02	46.42	137.02	112.47	74.9	55.24	43.61
	Range of λ	[0,2.3]	[0.1,0.9]	[0.38,1.62]	[0.5,1.5]	[0.57,1.43]	[0.07,1.93]	[0.24,1.76]	[0.43,1.57]	[0.52,1.48]	[0.58,1.42]
0.75	0.3	162.3	148.24	125.8	109.1	96.28	103.32	116.43	111.15	100.42	90.46
	0.5	169.53	161.37	146.83	134.56	124.16	107.92	126.74	129.73	123.85	116.66
	0.7	174.72	171.5	165.24	159.35	153.86	111.23	134.7	146	146.67	144.56
	0.8	176.4	174.93	172	169.1	166.29	112.3	137.39	151.95	155.63	156.24
	0.9	177.43	177.06	176.29	175.52	174.76	112.96	139.06	155.77	161.55	164.2
	1	177.78	177.78	177.78	177.78	177.78	113.18	139.63	157.08	163.62	167.03
	1.1	177.43	177.06	176.29	175.52	174.76	112.96	139.06	155.77	161.55	164.2
	1.2	176.4	174.93	172	169.1	166.29	112.3	137.39	151.95	155.63	156.24
	1.3	174.72	171.5	165.24	159.35	153.86	111.23	134.7	146	146.67	144.56
	1.5	169.53	161.37	146.83	134.56	124.16	107.92	126.74	129.73	123.85	116.66
	1.7	162.3	148.24	125.8	109.1	96.28	103.32	116.43	111.15	100.42	90.46
	Range of λ	[0.2,99]	[0.2,3]	[0.04,1.96]	[0.24,1.76]	[0.34,1.66]	[0.18,1.82]	[0.18,1.82]	[0.18,1.82]	[0.3,1.7]	[0.38,1.62]

3. Some modified shrinkage estimators for S.D. σ

The $MSE \{T_{(s,\sigma)}\}$ given by equation (2.3) is minimised for

$$\begin{aligned} k_{opt}^{(a)} &= \frac{(1 - \hat{\lambda})^2}{(1 - \hat{\lambda})^2 + v_1(n)} \\ &= \frac{(\sigma - \sigma_0)^2}{(\sigma - \sigma_0)^2 + \sigma^2 v_1(n)} \end{aligned} \quad (3.1)$$

Now, replacing σ and σ^2 by their unbiased estimator $T_{(u,\sigma)}$ and $T_{(u,\sigma^2)} = s^2$ respectively in (3.1), we get a consistent estimate of $k_{opt}^{(a)}$ as

$$\hat{k}_{opt}^{(1)} = \frac{(T_{(u,\sigma)} - \sigma_0)^2}{(T_{(u,\sigma)} - \sigma_0)^2 + v_1(n)s^2} . \quad (3.2)$$

Replacing σ by its unbiased estimator $T_{(u,\sigma)}$ in (3.1), we get

$$\hat{k}_{opt}^{(2)} = \frac{(T_{(u,\sigma)} - \sigma_0)^2}{(T_{(u,\sigma)} - \sigma_0)^2 + v_1(n)\{1 + v_1(n)\}s^2} , \quad (3.3)$$

and replacing σ by its unbiased estimator $T_{(u,\sigma)}$ and σ^2 by its MMSE estimator

$T_{(MMSE,\sigma^2)} = \left(\frac{n-1}{n+1}\right)s^2$ in (3.1), we get

$$\hat{k}_{opt}^{(3)} = \frac{(T_{(u,\sigma)} - \sigma_0)^2}{(T_{(u,\sigma)} - \sigma_0)^2 + v_1(n)\left(\frac{n-1}{n+1}\right)s^2} . \quad (3.4)$$

Many more consistent estimators of $k_{opt}^{(a)}$ can be obtained. However, a more flexible estimator of $k_{opt}^{(a)}$ can be obtained as :

$$\hat{k}_{opt}^{(\alpha)} = \frac{(T_{(u,\sigma)} - \sigma_0)^2}{(T_{(u,\sigma)} - \sigma_0)^2 + \alpha_1 s^2}, \quad (3.5)$$

where $\alpha_1 (\geq 0)$ is a suitably chosen constant.

Thus, the resulting modified shrinkage estimators of S.D. σ are given by

$$T_{(s,\sigma)}^{(1)} = \sigma_0 + \frac{(T_{(u,\sigma)} - \sigma_0)^3}{(T_{(u,\sigma)} - \sigma_0)^2 + v_1(n)s^2} \quad (3.6)$$

$$T_{(s,\sigma)}^{(2)} = \sigma_0 + \frac{(T_{(u,\sigma)} - \sigma_0)^3}{(T_{(u,\sigma)} - \sigma_0)^2 + v_1(n)(1 + v_1(n))s^2}, \quad (3.7)$$

$$T_{(s,\sigma)}^{(3)} = \sigma_0 + \frac{(T_{(u,\sigma)} - \sigma_0)^3}{(T_{(u,\sigma)} - \sigma_0)^2 + \left(\frac{n-1}{n+1}\right)v_1(n)s^2}, \quad (3.8)$$

and

$$T_{(s,\sigma)}^{(4)} = \sigma_0 + \frac{(T_{(u,\sigma)} - \sigma_0)^3}{(T_{(u,\sigma)} - \sigma_0)^2 + \alpha_1 s^2}. \quad (3.9)$$

The biases and MSEs of the estimators $T_{(s,\sigma)}^{(i)}$, $i = 1$ to 4 are respectively given by

$$Bias\{T_{(s,\sigma)}^{(i)}\} = (\sigma_0 - \sigma) + \int_0^\infty \left[\frac{(T_{(u,\sigma)} - \sigma_0)^3}{\{(T_{(u,\sigma)} - \sigma_0)^2 + j s^2\}} \right] f(s) ds \quad (3.10)$$

and

$$MSE\{T_{(s,\sigma)}^{(i)}\} = (\sigma_0 - \sigma)^2 + 2(\sigma_0 - \sigma) \int_0^\infty \left[\frac{(T_{(u,\sigma)} - \sigma_0)^3}{\{(T_{(u,\sigma)} - \sigma_0)^2 + j s^2\}} \right] f(s) ds$$

$$+ \int_0^{\infty} \left[\frac{(T_{(u,\sigma)} - \sigma_0)^6}{\left\{ (T_{(u,\sigma)} - \sigma_0)^2 + j s^2 \right\}^3} \right] f(s) ds, \quad (3.11)$$

where $j = v_1(n)$, $v_1(n)\{1 + v_1(n)\}$, $v_1(n)\left(\frac{n-1}{n+1}\right)$ and $\alpha_1(\geq 0)$.

Now, transform the above expression in the form of $\hat{\lambda}$, for this we are using transformation $y = \frac{(n-1)s^2}{2\sigma^2}$, we have

$$\begin{aligned} & MSE \left\{ T_{(s,\sigma)}^{(i)} \right\} = \\ & \sigma^2 \left\{ (\hat{\lambda} - 1)^2 + \frac{2(\hat{\lambda} - 1)}{\Gamma\left(\frac{n-1}{2}\right)} \int_0^{\infty} \left[\frac{\left\{ C_5(n)y^{\frac{1}{2}} - \hat{\lambda} \right\}^3}{\left\{ \left(C_5(n)y^{\frac{1}{2}} - \hat{\lambda} \right)^2 + \left(\frac{2}{n-1} \right) j y \right\}} \right] e^{-y} y^{\left(\frac{n-3}{2}\right)} dy \right. \\ & \left. + \frac{1}{\Gamma\left(\frac{n-1}{2}\right)} \int_0^{\infty} \left[\frac{\left\{ C_5(n)y^{\frac{1}{2}} - \hat{\lambda} \right\}^6}{\left\{ \left(C_5(n)y^{\frac{1}{2}} - \hat{\lambda} \right)^2 + \left(\frac{2}{n-1} \right) j y \right\}^2} \right] e^{-y} y^{\left(\frac{n-3}{2}\right)} dy \right\} \quad (3.12) \end{aligned}$$

where

$$C_5(n) = \frac{\Gamma\left(\frac{n-1}{2}\right)}{\Gamma\left(\frac{n}{2}\right)}.$$

The PRE of $T_{(s,\sigma)}^{(i)}$ w.r.t. unbiased estimator $T_{(u,\sigma)}$ and MMSE estimator $T_{(MMSE,\sigma)}$ for different values of $i = 1, 2, 3$ and $j = v_1(n)$, $v_1(n)\{1 + v_1(n)\}$ and $v_1(n)\left(\frac{n-1}{n+1}\right)$ are given in the Tables 3.1, 3.2

and 3.3 respectively. It is observed that the proposed modified estimator $T_{(s,\sigma)}^{(i)}$ is better than $T_{(u,\sigma)}$ and $T_{(MMSE,\sigma)}$ when $\hat{\lambda} \in [0.7, 1.3]$ and $2 \leq n \leq 9$. For fixed $\hat{\lambda}$, the $PRE\{T_{(s,\sigma)}^{(1)}, T_{(u,\sigma)}\}$ and $PRE\{T_{(s,\sigma)}^{(1)}, T_{(MMSE,\sigma)}\}$ decrease as n increases and for fixed n these values decrease as $\hat{\lambda}$ goes away from unity where PRE's attain its maximum. Similar trend is found for $T_{(s,\sigma)}^{(2)}$ and $T_{(s,\sigma)}^{(3)}$ also. We also observe that the gain in efficiency w.r.t. $T_{(MMSE,\sigma)}$ is lesser than $T_{(u,\sigma)}$. Tables 3.1, 3.2 and 3.3 show that the estimator $T_{(s,\sigma)}^{(2)}$ is the best (in the sense of having smallest MSE) among $T_{(s,\sigma)}^{(1)}$, $T_{(s,\sigma)}^{(2)}$ and $T_{(s,\sigma)}^{(3)}$ followed by $T_{(s,\sigma)}^{(3)}$. Thus, the modified estimator $T_{(s,\sigma)}^{(2)}$ is to be preferred in practice when $\hat{\lambda}$ moves in the vicinity of unity and larger efficiency is observed when sample size n is small.

All the integration and gamma values are calculated using Gauss-Laguerre integration formula (10 point) with the help of computer programs developed in C++ platform.

Table 3.1 $PRE\{T_{(s,\sigma)}^{(1)}, T_{(u,\sigma)}\}$ and $PRE\{T_{(s,\sigma)}^{(1)}, T_{(MMSE,\sigma)}\}$

$\hat{\lambda} \downarrow n \rightarrow$	2	3	5	7	9
0.4	221.09 140.75	136.16 106.94	116.19 102.66	99.79 92.77	97.34 91.45
0.5	251.83 160.32	166.21 130.54	117.54 103.86	99.96 92.01	98.49 92.53
0.7	376.11 239.44	263.07 206.62	131.36 116.06	151.07 139.04	119.63 112.4
0.8	489.72 311.76	275.15 216.11	163.78 144.72	180.54 166.17	136.59 128.34
0.9	605.7 385.6	272.57 214.08	199.68 176.43	195.89 180.29	185.67 174.45
1	650.39 414.05	263.12 206.66	211.21 186.62	204 187.76	211.98 199.17
1.1	613.81 390.76	244.21 191.81	203.74 180.02	176.27 162.24	196.91 185.01
1.2	542.48 345.35	213.59 167.76	180.2 159.22	141.54 130.27	156.23 146.79
1.3	464.72 295.85	178.39 140.11	145.56 128.61	121.04 111.41	111.02 104.31
1.5	321.31 204.55	127.85 100.41	94.51 83.51	89.13 82.04	77.01 72.36
1.7	212.99 135.6	104.99 82.46	77.25 68.26	71.32 65.65	70.49 66.23
Range of $\hat{\lambda}$	[0 , 2.1] [0.1 , 1.85]	[0.2 , 1.7] [0.3 , 1.5]	[0.25 , 1.42] [0.37 , 1.36]	[0.57 , 1.42] [0.6 , 1.36]	[0.55 , 1.32] [0.6 , 1.31]

Table 3.2 $PRE\{T_{(s,\sigma)}^{(2)}, T_{(u,\sigma)}\}$ and $PRE\{T_{(s,\sigma)}^{(2)}, T_{(MMSE,\sigma)}\}$

$\hat{\lambda} \downarrow \quad n \rightarrow$	2	3	5	7	9
0.4	244.45 155.62	140.73 110.53	116 102.5	97.76 92.25	96.85 91
0.5	292.32 186.1	175.48 137.83	118.1 104.35	99.78 91.84	98.21 92.27
0.7	508.15 323.5	294.55 231.34	135.09 119.36	152.63 140.48	120.29 113.02
0.8	729.82 464.62	318.93 250.49	170.12 150.31	185.01 170.28	138.98 130.58
0.9	984.12 626.51	320.8 251.95	208.46 184.19	203.88 187.65	190.35 178.84
1	1062.06 676.13	308.45 242.26	220.57 194.89	213.09 196.12	217.32 204.19
1.1	913.93 581.83	280.45 220.26	211.49 186.86	182.15 167.65	200.31 188.21
1.2	715.82 455.7	238.31 187.17	185.43 163.84	143.74 132.3	157.83 148.29
1.3	554.45 352.97	193.51 151.98	148.68 131.37	121.03 111.4	111.71 104.96
1.5	341.86 217.63	131.48 103.26	94.86 83.82	88.44 81.4	76.34 71.72
1.7	217.6 138.53	103.1 80.97	75.86 67.02	70.39 64.79	69.5 65.3
Range of $\hat{\lambda}$	[0 , 2] [0.06 , 1.85]	[0.08 , 1.72] [0.3 , 1.52]	[0.1 , 1.47] [0.38 , 1.51]	[0.51 , 1.42] [0.58 , 1.35]	[0.52 , 1.33] [0.59 , 1.31]

Note: **Bold** figures denote the PREs/ranges w.r.t. MMSE estimator

Table 3.3 $PRE \{T_{(s,\sigma)}^{(3)}, T_{(u,\sigma)}\}$ and $PRE \{T_{(s,\sigma)}^{(3)}, T_{(MMSE,\sigma)}\}$

$\hat{\lambda} \downarrow n \rightarrow$	2	3	5	7	9
0.4	164.72 104.86	122.18 95.96	114.92 101.54	102.04 93.91	98.65 92.69
0.5	176.47 112.35	139.68 109.71	114.61 101.27	100.19 92.22	99.09 93.1
0.7	209.87 133.61	191.84 150.67	120.46 106.44	145.15 133.59	117.24 110.16
0.8	240.44 153.07	191.66 150.53	145.36 128.44	166.05 152.83	128.78 120.99
0.9	271.08 172.58	188.67 148.18	174.48 154.17	172.53 158.79	170.61 160.3
1	286.91 182.65	185.34 145.57	184.12 162.69	177.95 163.78	194.78 183.01
1.1	291.1 185.32	179.22 140.76	180.25 159.27	158.82 146.17	185.33 174.13
1.2	287.18 182.83	165.85 130.26	163.51 144.48	134.87 124.13	150.46 141.37
1.3	275.62 175.46	146.84 115.33	135.35 119.59	121.15 111.5	108.45 101.9
1.5	233.62 148.73	120.73 94.82	93.61 82.71	91.07 83.82	79.49 74.68
1.7	179.12 114.03	111.31 87.43	82.05 72.5	74.46 68.53	73.91 69.44
Range of $\hat{\lambda}$	[0.001, 2] [0.35, 1.77]	[0.01, 1.89] [0.5, 1.40]	[0.1, 1.44] [0.58, 1.37]	[0.45, 1.43] [0.6, 1.35]	[0.53, 1.30] [0.6, 1.3]

Note: **Bold** figures denote the PREs/ranges w.r.t. MMSE estimator

Expression (3.1) shows that the optimum value of $\ddot{k}$ depends on the unknown parameter σ . Since σ_0 is the guessed or prior value of σ , one may replace σ by $\dot{\alpha}\sigma_0$, where $\dot{\alpha}$ is a positive constant. Thus, putting $\sigma = \dot{\alpha}\sigma_0$ ($\dot{\alpha} > 0$), in (3.1), we get

$$k_{opt}^{(b)} = \frac{(\dot{\alpha} - 1)^2}{(\dot{\alpha} - 1)^2 + v_1(n)\dot{\alpha}^2}.$$

Thus, the resulting estimator of σ is

$$T_{(s,\sigma)}^{(*,opt)} = \frac{(1 - \dot{\alpha})^2 T_{(u,\sigma)} + v_1(n)\sigma_0\dot{\alpha}^2}{(1 - \dot{\alpha})^2 + v_1(n)\dot{\alpha}^2}, \quad (3.13)$$

The bias of $T_{(s,\sigma)}^{(*,opt)}$ is given by

$$\begin{aligned} Bias\{T_{(s,\sigma)}^{(*,opt)}\} &= E(T_{(s,\sigma)}^{(*,opt)}) - \sigma \\ &= -\frac{(1-\hat{\lambda})v_1(n)\sigma\dot{\alpha}^2}{(1-\dot{\alpha})^2 + v_1(n)\dot{\alpha}^2} \end{aligned}$$

The MSE of $T_{(s,\sigma)}^{(*,opt)}$ is given by

$$\begin{aligned} MSE\{T_{(s,\sigma)}^{(*,opt)}\} &= Var\{T_{(s,\sigma)}^{(*,opt)}\} + [Bias\{T_{(s,\sigma)}^{(*,opt)}\}]^2 \\ &= \frac{(1-\dot{\alpha})^4 v_1(n)\sigma^2}{\{(1-\dot{\alpha})^2 + v_1(n)\dot{\alpha}^2\}^2} + \frac{(1-\hat{\lambda})^2 (v_1(n))^2 \dot{\alpha}^4 \sigma^2}{\{(1-\dot{\alpha})^2 + v_1(n)\dot{\alpha}^2\}^2} \\ &= \frac{\sigma^2 v_1(n) \{(1-\dot{\alpha})^4 + (1-\hat{\lambda})^2 v_1(n)\dot{\alpha}^4\}}{\{(1-\dot{\alpha})^2 + v_1(n)\dot{\alpha}^2\}^2}, \end{aligned}$$

which is smaller than the variance of the conventional unbiased estimator $T_{(u,\sigma)}$ if

$$\frac{\{(1-\dot{\alpha})^4 + (1-\hat{\lambda})^2 v_1(n)\dot{\alpha}^4\}}{\{(1-\dot{\alpha})^2 + v_1(n)\dot{\alpha}^2\}^2} < 1$$

$$\text{i.e. if } \dot{\alpha} \leq \left[1 + \sqrt{\frac{1}{2} \{(1-\hat{\lambda})^2 - v_1(n)\}} \right]^{-1}$$

which shows that $\dot{\alpha}$ should be between zero and one (i.e. $0 < \dot{\alpha} \leq 1$) for gain in efficiency.

4. The suggested class of shrinkage estimators for inverse of standard deviation σ^{-1}

Let us define a class of estimator for σ^{-1} as

$$\begin{aligned} T_{(s,\sigma^{-1})} &= w T_{(u,\sigma^{-1})} + (1-w) \sigma_0^{-1} \\ &= w (T_{(u,\sigma^{-1})} - \sigma_0^{-1}) + \sigma_0^{-1}, \end{aligned} \quad (4.1)$$

where w is a scalar such that $0 \leq w \leq 1$.

The bias and MSE of $T_{(s,\sigma^{-1})}$ are respectively given by

$$\begin{aligned} \text{Bias} \{T_{(s,\sigma^{-1})}\} &= E(T_{(s,\sigma^{-1})}) - \sigma^{-1} \\ &= \sigma^{-1} (1-w)(\lambda^* - 1) \end{aligned} \quad (4.2)$$

and

$$\begin{aligned} \text{MSE} \{T_{(s,\sigma^{-1})}\} &= \text{Var} \{T_{(s,\sigma^{-1})}\} + [\text{Bias} \{T_{(s,\sigma^{-1})}\}]^2 \\ &= \sigma^{-2} [w^2 v_2(n) + (1-w)^2 (\lambda^* - 1)^2], \end{aligned} \quad (4.3)$$

where $\lambda^* = \lambda^{-1}$.

From (1.7) and (4.3), we have

$$\text{Var} \{T_{(s,\sigma^{-1})}\} - \text{MSE} \{T_{(s,\sigma^{-1})}\} = \sigma^{-2} [v_2(n) (1-w^2) - (1-w)^2 (\lambda^* - 1)^2],$$

which is positive if

$$\begin{aligned} [v_2(n) (1-w^2) - (1-w)^2 (\lambda^* - 1)^2] &> 0, \\ \text{i.e. if } w &> \left[\frac{\{(\lambda^* - 1)^2 - v_2(n)\}}{\{(\lambda^* - 1)^2 + v_2(n)\}} \right] \\ \text{i.e. } \left[\frac{\{(\lambda^* - 1)^2 - v_2(n)\}}{\{(\lambda^* - 1)^2 + v_2(n)\}} \right] &< w \leq 1. \end{aligned} \quad (4.4)$$

Further, from (1.18) and (4.3), we have

$$MSE \left\{ T_{(MMSE, \sigma^{-1})} \right\} - MSE \left\{ T_{(s, \sigma^{-1})} \right\} > 0 ,$$

$$\text{if } (c-d) < w < (c+d) , \quad (4.5)$$

where

$$c = \frac{(\lambda^* - 1)^2}{\{v_2(n) + (\lambda^* - 1)^2\}} ,$$

$$d = \frac{v_2(n) \sqrt{\lambda^* (2 - \lambda^*)}}{\sqrt{\{1 + v_2(n)\} \{v_2(n) + (\lambda^* - 1)^2\}}} , 0 < \lambda^* < 2 .$$

4.1. PRE and Empirical Study

To see the performance of the proposed estimator $T_{(s, \sigma^{-1})}$ over the usual unbiased estimator $T_{(u, \sigma^{-1})}$ and MMSE estimator $T_{(MMSE, \sigma^{-1})}$, we have computed the percentage relative efficiencies (PRE's) of $T_{(s, \sigma^{-1})}$ with respect to $T_{(u, \sigma^{-1})}$ and $T_{(MMSE, \sigma^{-1})}$ using the formulae:

$$PRE \left\{ T_{(s, \sigma^{-1})}, T_{(u, \sigma^{-1})} \right\} = \frac{MSE \left\{ T_{(u, \sigma^{-1})} \right\}}{MSE \left\{ T_{(s, \sigma^{-1})} \right\}} * 100 = \frac{Var \left\{ T_{(u, \sigma^{-1})} \right\}}{MSE \left\{ T_{(s, \sigma^{-1})} \right\}} * 100$$

$$= \frac{1}{\left[w^2 + \{v_2(n)\}^{-1} (1-w)^2 \{\lambda^* - 1\}^2 \right]} \quad (4.6)$$

and

$$PRE \left\{ T_{(s, \sigma^{-1})}, T_{(MMSE, \sigma^{-1})} \right\} = \frac{MSE \left\{ T_{(MMSE, \sigma^{-1})} \right\}}{MSE \left\{ T_{(s, \sigma^{-1})} \right\}} * 100$$

$$= \frac{v_2(n)}{\left[v_2(n) w^2 + (1-w)^2 \{\lambda^* - 1\}^2 \right] \{1 + v_2(n)\}} * 100 , \quad (4.7)$$

for different w, λ^*, n and findings are given in Table 4.2. It is observed that:

- the proposed estimator $T_{(s, \sigma^{-1})}$ is better than the usual unbiased estimator $T_{(u, \sigma^{-1})}$ and MMSE estimator $T_{(MMSE, \sigma^{-1})}$ with substantial gain in efficiency when $\lambda^* \in [0.70, 1.30] \Rightarrow \lambda \in [0.77, 1.43]$ and $w \in [0.25, 0.75]$.
- larger gain in efficiency is observed when sample size n is small.
- for fixed w and λ^* , the value of PRE decreases as n increases.
- for fixed w and n , the value of PRE decreases as λ^* goes away from unity.
- the gain in efficiency by using $T_{(s, \sigma^{-1})}$ over MMSE estimator $T_{(MMSE, \sigma^{-1})}$ is fewer than by using $T_{(s, \sigma^{-1})}$ over unbiased estimator $T_{(u, \sigma^{-1})}$.

Thus, the proposed estimator is more beneficial when σ_0 , the guessed value is close to the true value of σ and when the sample size n is small.

The ranges of w (for different values of n and λ^*) in which $T_{(s, \sigma^{-1})}$ is better than $T_{(u, \sigma^{-1})}$ and $T_{(MMSE, \sigma^{-1})}$ are given in the Table 4.1 (using equations (4.4) and (4.5)). It is observed that the effective range of w becomes shorter as sample size n goes up for fixed λ^* , also for fixed n the effective range becomes shorter as λ^* goes away from unity.

Table 4.1 Range of w for which the proposed estimator $T_{(s,\sigma^{-1})}$ is better than $T_{(u,\sigma^{-1})}$ and MMSE estimator $T_{(MMSE,\sigma^{-1})}$.

$\lambda^* \downarrow n \rightarrow$	7	9	11	13	15
0.4	(0.46,1.00) (0.53, 0.93)	(0.61,1.00) (0.66, 0.95)	(0.70,1.00) (0.73, 0.97)	(0.75,1.00) (0.78, 0.97)	(0.79,1.00) (0.81, 0.98)
0.5	(0.31,1.00) (0.37, 0.94)	(0.49,1.00) (0.53, 0.96)	(0.59,1.00) (0.62, 0.97)	(0.66,1.00) (0.69, 0.97)	(0.71,1.00) (0.73, 0.98)
0.7	(0.00,1.00) (0.00, 0.94)	(0.02,1.00) (0.06, 0.96)	(0.17,1.00) (0.20, 0.97)	(0.27,1.00) (0.30, 0.97)	(0.36,1.00) (0.38, 0.98)
0.8	(0.00,1.00) (0.00, 0.94)	(0.00,1.00) (0.00, 0.96)	(0.00,1.00) (0.00, 0.97)	(0.00,1.00) (0.00, 0.98)	(0.00,1.00) (0.00, 0.98)
0.9	(0.00,1.00) (0.00, 0.94)	(0.00,1.00) (0.00, 0.96)	(0.00,1.00) (0.00, 1.00)	(0.00,1.00) (0.00, 0.98)	(0.00,1.00) (0.00, 0.98)
1	(0.00,1.00) (0.00, 0.94)	(0.00,1.00) (0.00, 0.96)	(0.00,1.00) (0.00, 1.00)	(0.00,1.00) (0.00, 0.98)	(0.00,1.00) (0.00, 0.98)
1.1	(0.00,1.00) (0.00, 0.94)	(0.00,1.00) (0.00, 0.96)	(0.00,1.00) (0.00, 1.00)	(0.00,1.00) (0.00, 0.98)	(0.00,1.00) (0.00, 0.98)
1.2	(0.00,1.00) (0.00, 0.94)	(0.00,1.00) (0.00, 0.96)	(0.00,1.00) (0.00, 1.00)	(0.00,1.00) (0.00, 0.98)	(0.00,1.00) (0.00, 0.98)
1.3	(0.00,1.00) (0.00, 0.94)	(0.02,1.00) (0.06, 0.96)	(0.17,1.00) (0.20, 0.97)	(0.27,1.00) (0.30, 0.97)	(0.36,1.00) (0.38, 0.98)
1.5	(0.31,1.00) (0.37, 0.94)	(0.49,1.00) (0.53, 0.96)	(0.59,1.00) (0.62, 0.97)	(0.66,1.00) (0.69, 0.97)	(0.71,1.00) (0.73, 0.98)
1.7	(0.58,1.00) (0.65,0.93)	(0.70,1.00) (0.75, 0.95)	(0.77,1.00) (0.80, 0.96)	(0.81,1.00) (0.84,0.97)	(0.84,1.00) (0.86,0.98)
1.9	(0.72,1.00) (0.80,0.92)	(0.81,1.00) (0.86, 0.94)	(0.85,1.00) (0.90, 0.96)	(0.88,1.00) (0.92, 0.97)	(0.90,1.00) (0.93, 0.97)
2	(0.77,1.00) n.e.	(0.84,1.00) n.e.	(0.88,1.00) n.e.	(0.90,1.00) n.e.	(0.92,1.00) n.e.

Note : a) The **bold** figures in parentheses show the range of w for which the proposed estimator $T_{(u,\sigma^{-1})}$ is better than MMSE estimator $T_{(MMSE,\sigma^{-1})}$. b) **n.e** stands for does not exit.

Table 4.2 PRE of $T_{(s,\sigma^{-1})}$ with respect to unbiased estimator $T_{(u,\sigma^{-1})}$ and MMSE estimator $T_{(MMSE,\sigma^{-1})}$, for different values of $n = 2,3,5,7,9$; $w = 0.25$ (0.25) 0.75 and λ^*

w	λ^* ↓ n→	PRE of $T_{(s,\sigma)}$ w. r. t. $T_{(u,\sigma^{-1})}$					PRE of $T_{(s,\sigma)}$ w. r. t. $T_{(MMSE,\sigma^{-1})}$				
		7	9	11	13	15	7	9	11	13	15
0.25	0.3	46.42	30.78	23	18.36	15.27	41.02	28.33	21.61	17.46	14.65
	0.5	88.52	59.23	44.47	35.59	29.66	78.21	54.52	41.78	33.86	28.45
	0.7	223.87	154.37	117.71	95.1	79.76	197.8	142.08	110.6	90.47	76.51
	0.8	428.72	309.96	242.55	199.17	168.93	378.8	285.28	227.89	189.47	162.05
	0.9	950.67	784.12	666.9	580.06	513.18	839.99	721.7	626.59	551.81	492.26
	1	1600	1600	1600	1600	1600	1413.72	1472.62	1503.3	1522.09	1534.78
	1.1	950.67	784.12	666.9	580.06	513.18	839.99	721.7	626.59	551.81	492.26
	1.2	428.72	309.96	242.55	199.17	168.93	378.8	285.28	227.89	189.47	162.05
	1.3	223.87	154.37	117.71	95.1	79.76	197.8	142.08	110.6	90.47	76.51
	1.5	88.52	59.23	44.47	35.59	29.66	78.21	54.52	41.78	33.86	28.45
	1.7	46.42	30.78	23	18.36	15.27	41.02	28.33	21.61	17.46	14.65
	Range of λ^*	[0.54,1.46]	[0.63,1.37]	[0.68,1.32]	[0.71,1.29]	[0.74,1.26]	[0.57,1.43]	[0.64,1.36]	[0.69,1.31]	[0.72,1.28]	[0.75,1.25]
0.5	0.3	84.77	60.02	46.42	37.83	31.92	74.9	55.24	43.61	35.99	30.62
	0.5	138.06	102.82	81.86	67.98	58.12	121.99	94.64	76.91	64.67	55.75
	0.7	237.67	196.03	166.72	145.01	128.3	210	180.42	156.65	137.95	123.07
	0.8	306.85	273.52	246.63	224.53	206.05	271.13	251.74	231.73	213.6	197.65
	0.9	371.78	358.55	346.18	334.62	323.81	328.5	330	325.26	318.33	310.61
	1	400	400	400	400	400	353.43	368.16	375.83	380.52	383.69
	1.1	371.78	358.55	346.18	334.62	323.81	328.5	330	325.26	318.33	310.61
	1.2	306.85	273.52	246.63	224.53	206.05	271.13	251.74	231.73	213.6	197.65
	1.3	237.67	196.03	166.72	145.01	128.3	210	180.42	156.65	137.95	123.07
	1.5	138.06	102.82	81.86	67.98	58.12	121.99	94.64	76.91	64.67	55.75
	1.7	84.77	60.02	46.42	37.83	31.92	74.9	55.24	43.61	35.99	30.62
	Range of λ^*	[0.38,1.62]	[0.5,1.5]	[0.57,1.43]	[0.61,1.39]	[0.65,1.35]	[0.43,1.57]	[0.52,1.48]	[0.58,1.42]	[0.63,1.37]	[0.66,1.34]
0.75	0.3	125.8	109.1	96.28	86.15	77.93	111.15	100.42	90.46	81.95	74.76
	0.5	146.83	134.56	124.16	115.24	107.51	129.73	123.85	116.66	109.63	103.12
	0.7	165.24	159.35	153.86	148.72	143.91	146	146.67	144.56	141.48	138.05
	0.8	171.98	169.09	166.29	163.57	160.95	151.95	155.63	156.24	155.61	154.38
	0.9	176.29	175.52	174.76	174	173.25	155.77	161.55	164.2	165.53	166.19
	1	177.78	177.78	177.78	177.78	177.78	157.08	163.62	167.03	169.12	170.53
	1.1	176.29	175.52	174.76	174	173.25	155.77	161.55	164.2	165.53	166.19
	1.2	171.98	169.09	166.29	163.57	160.95	151.95	155.63	156.24	155.61	154.38
	1.3	165.24	159.35	153.86	148.72	143.91	146	146.67	144.56	141.48	138.05
	1.5	146.83	134.56	124.16	115.24	107.51	129.73	123.85	116.66	109.63	103.12
	1.7	125.8	109.1	96.28	86.15	77.93	111.15	100.42	90.46	81.95	74.76
	Range of λ^*	[0.06,1.94]	[0.24,1.76]	[0.33,1.67]	[0.42,1.58]	[0.46,1.54]	[0.18,1.2]	[0.3,1.7]	[0.38,1.62]	[0.44,1.56]	[0.49,1.51]

5. Improved classes of estimators for estimating the inverse of standard deviation (i.e. σ^{-1})

Minimizing the $MSE \left\{ T_{(s, \sigma^{-1})} \right\}$, given by equation (4.3) with respect to w we get the optimum value of w as

$$\begin{aligned} w_{opt}^{(c)} &= \frac{(1 - \lambda^*)^2}{(1 - \lambda^*)^2 + v_2(n)} \\ &= \frac{(\sigma^{-1} - \sigma_0^{-1})^2}{(\sigma^{-1} - \sigma_0^{-1})^2 + \sigma^{-2} v_2(n)}. \end{aligned} \quad (5.1)$$

Now, replacing σ^{-1} and σ^{-2} by their unbiased estimator $T_{(u, \sigma^{-1})}$ and $T_{(u, \sigma^{-2})} = \left(\frac{n-3}{n-1} \right) s^2$ respectively in (5.1), we get a consistent estimate of $w_{opt}^{(c)}$ as:

$$\hat{w}_{opt}^{(c1)} = \frac{(T_{(u, \sigma^{-1})} - \sigma_0^{-1})^2}{(T_{(u, \sigma^{-1})} - \sigma_0^{-1})^2 + \left\{ \frac{n-3}{n-1} \right\} v_2(n) s^{-2}}. \quad (5.2)$$

Replacing σ^{-1} by its unbiased estimator $T_{(u, \sigma^{-1})}$ in (5.1), we get

$$\hat{w}_{opt}^{(c2)} = \frac{(T_{(u, \sigma^{-1})} - \sigma_0^{-1})^2}{(T_{(u, \sigma^{-1})} - \sigma_0^{-1})^2 + v_2(n) \{c_2(n)\}^2 s^{-2}}, \quad (5.3)$$

and replacing σ^{-1} by its unbiased estimator $T_{(u, \sigma^{-1})}$ and σ^{-2} by its MMSE estimator $T_{(MMSE, \sigma^{-2})} = \left(\frac{n-5}{n-1} \right) s^{-2}$ in (5.1), we get

$$\hat{w}_{opt}^{(c3)} = \frac{(T_{(u, \sigma^{-1})} - \sigma_0^{-1})^2}{(T_{(u, \sigma^{-1})} - \sigma_0^{-1})^2 + v_2(n) \left\{ \frac{n-5}{n-1} \right\} s^{-2}}. \quad (5.4)$$

Many more consistent estimators of $w_{opt}^{(c)}$ can be obtained. However, a more flexible estimator of $w_{opt}^{(c)}$ can be obtained as :

$$\hat{w}_{opt}^{(\beta)} = \frac{\left(T_{(u, \sigma^{-1})} - \sigma_0^{-1}\right)^2}{\left(T_{(u, \sigma^{-1})} - \sigma_0^{-1}\right)^2 + \beta_1 s^{-2}}, \quad (5.5)$$

where $\beta_1 (\geq 0)$ is a suitably chosen constant.

Thus, the resulting modified shrinkage estimators of the inverse of S.D. σ^{-1} are given by

$$T_{(s, \sigma^{-1})}^{(1)} = \sigma_0^{-1} + \frac{\left(T_{(u, \sigma^{-1})} - \sigma_0^{-1}\right)^3}{\left(T_{(u, \sigma^{-1})} - \sigma_0^{-1}\right)^2 + \left\{\frac{n-3}{n-1}\right\} v_2(n) s^{-2}}, \quad (5.6)$$

$$T_{(s, \sigma^{-1})}^{(2)} = \sigma_0^{-1} + \frac{\left(T_{(u, \sigma^{-1})} - \sigma_0^{-1}\right)^3}{\left(T_{(u, \sigma^{-1})} - \sigma_0^{-1}\right)^2 + v_2(n) \{c_2(n)\}^2 s^{-2}}, \quad (5.7)$$

$$T_{(s, \sigma^{-1})}^{(3)} = \sigma_0^{-1} + \frac{\left(T_{(u, \sigma^{-1})} - \sigma_0^{-1}\right)^3}{\left(T_{(u, \sigma^{-1})} - \sigma_0^{-1}\right)^2 + \left\{\frac{n-5}{n-1}\right\} v_2(n) s^{-2}} \quad (5.8)$$

and

$$T_{(s, \sigma^{-1})}^{(4)} = \sigma_0^{-1} + \frac{\left(T_{(u, \sigma^{-1})} - \sigma_0^{-1}\right)^3}{\left\{\left(T_{(u, \sigma^{-1})} - \sigma_0^{-1}\right)^2 + \beta_1 s^{-2}\right\}}. \quad (5.9)$$

The biases and MSE's of the estimators $T_{(s, \sigma^{-1})}^{(i)}$, $i = 1$ to 4 are respectively given by

$$Bias\{T_{(s, \sigma^{-1})}^{(i)}\} = (\sigma_0^{-1} - \sigma^{-1}) + \int_0^\infty \left[\frac{\left(T_{(u, \sigma^{-1})} - \sigma_0^{-1}\right)^3}{\left\{\left(T_{(u, \sigma^{-1})} - \sigma_0^{-1}\right)^2 + j' s^{-2}\right\}} \right] f(s) ds \quad (5.10)$$

and

$$\begin{aligned}
MSE \left\{ T_{(s, \sigma^{-1})}^{(i)} \right\} &= \left(\sigma_0^{-1} - \sigma^{-1} \right)^2 + \\
&2 \left(\sigma_0^{-1} - \sigma^{-1} \right) \int_0^\infty \left[\frac{\left(T_{(u, \sigma^{-1})} - \sigma_0^{-1} \right)^3}{\left\{ \left(T_{(u, \sigma^{-1})} - \sigma_0^{-1} \right)^2 + j' s^{-2} \right\}} \right] f(s) ds \\
&+ \int_0^\infty \left[\frac{\left(T_{(u, \sigma^{-1})} - \sigma_0^{-1} \right)^6}{\left\{ \left(T_{(u, \sigma^{-1})} - \sigma_0^{-1} \right)^2 + j' s^{-2} \right\}^2} \right] f(s) ds
\end{aligned} \tag{5.11}$$

where $j' = v_2(n) \left(\frac{n-3}{n-1} \right)$, $v_2(n) \{c_2(n)\}^2$, $v_2(n) \left(\frac{n-5}{n-1} \right)$, $(\beta_1 \geq 0)$.

Now, transform the above expression in the form of λ^* , for this we are using transformation $y = \frac{(n-1)s^2}{2\sigma^2}$, we have $MSE \left\{ T_{(s, \sigma^{-1})}^{(i)} \right\} =$

$$\begin{aligned}
&\sigma^{-2} \left\{ \left(\lambda^* - 1 \right)^2 + \frac{2 \left(\lambda^* - 1 \right)}{\Gamma \left(\frac{n-1}{2} \right)} \int_0^\infty \left[\frac{\left(C_6(n) y^{-\frac{1}{2}} - \lambda^* \right)^3}{\left\{ \left(C_6(n) y^{-\frac{1}{2}} - \lambda^* \right)^2 + \left(\frac{n-1}{2} \right) j' (y)^{-1} \right\}} \right] e^{-y} y^{\left(\frac{n-3}{2} \right)} dy \right. \\
&+ \left. \frac{1}{\Gamma \left(\frac{n-1}{2} \right)} \int_0^\infty \left[\frac{\left(C_6(n) y^{-\frac{1}{2}} - \lambda^* \right)^6}{\left\{ \left(C_6(n) y^{-\frac{1}{2}} - \lambda^* \right)^2 + \left(\frac{2}{n-1} \right) j' (y)^{-1} \right\}^2} \right] e^{-y} y^{\left(\frac{n-3}{2} \right)} dy \right\}
\end{aligned} \tag{5.12}$$

where $C_6(n) = \frac{\Gamma \left(\frac{n-1}{2} \right)}{\Gamma \left(\frac{n-2}{2} \right)}$.

The PRE of $T_{(s,\sigma^{-1})}^{(i)}$ w.r.t unbiased estimator $T_{(u,\sigma^{-1})}$ and MMSE estimator $T_{(MMSE,\sigma^{-1})}$ for different values $j' = v_2(n) \left(\frac{n-3}{n-1} \right)$, $v_2(n) \{c_2(n)\}^2$ and $v_2(n) \left(\frac{n-5}{n-1} \right)$ respectively for $i=1,2,3$ are given in the Tables 5.1, 5.2 and 5.3. It is observed that the proposed modified estimators $T_{(s,\sigma^{-1})}^{(i)}$ are better than $T_{(u,\sigma^{-1})}$ and $T_{(MMSE,\sigma^{-1})}$ when $\lambda^* \in [0.7, 1.3] \Rightarrow \lambda \in [0.77, 1.43]$ and $7 \leq n \leq 15$. For fixed λ^* , the $PRE \{T_{(s,\sigma^{-1})}^{(i)}, T_{(u,\sigma^{-1})}\}$ and $PRE \{T_{(s,\sigma^{-1})}^{(i)}, T_{(MMSE,\sigma^{-1})}\}$ decrease as n increases and for fixed n these values decreases as λ^* goes away from unity where PRE's attain its maximum. We also observe that the gain in efficiency w.r.t. $T_{(MMSE,\sigma^{-1})}$ is lesser than $T_{(u,\sigma^{-1})}$. Tables 5.1, 5.2 and 5.3 show that the estimator $T_{(s,\sigma^{-1})}^{(2)}$ is the best (in the sense of having smallest MSE) among $T_{(s,\sigma^{-1})}^{(1)}, T_{(s,\sigma^{-1})}^{(2)}$ and $T_{(s,\sigma^{-1})}^{(3)}$ followed by $T_{(s,\sigma^{-1})}^{(1)}$. Thus, the modified estimator $T_{(s,\sigma^{-1})}^{(2)}$ is to be preferred in practice when λ^* moves toward unity and larger efficiency is observed when sample size n is small.

Expression (5.1) shows that the optimum value of w depends on the unknown parameter σ^{-1} . Since σ_0^{-1} is the guessed or prior value of σ^{-1} , one may replace σ^{-1} by $\ddot{\beta}\sigma_0^{-1}$, where $\ddot{\beta}$ is a positive constant. Thus, putting $\sigma^{-1} = \ddot{\beta}\sigma_0^{-1}$ ($\ddot{\beta} > 0$), in (5.1), we get

$$k_{opt}^{(d)} = \frac{(\ddot{\beta} - 1)^2}{(\ddot{\beta} - 1)^2 + v_2(n)\ddot{\beta}^2}.$$

Thus, the resulting estimator of σ^{-1} is

$$T_{(s,\sigma^{-1})}^{(*,opt)} = \frac{(1 - \ddot{\beta})^2 T_{(u,\sigma^{-1})} + v_2(n)\sigma_0^{-1}\ddot{\beta}^2}{((1 - \ddot{\beta})^2 + v_2(n)\ddot{\beta}^2)}. \quad (5.13)$$

The bias of $T_{(s,\sigma^{-1})}^{(*,opt)}$ is given by

$$\text{Bias} \left\{ T_{(s, \sigma^{-1})}^{(*, opt)} \right\} = - \frac{(1 - \lambda^*) v_2(n) \sigma^{-1} \ddot{\beta}^2}{\left\{ (1 - \ddot{\beta})^2 + v_2(n) \ddot{\beta}^2 \right\}} .$$

The MSE of $T_{(s, \sigma)}^{(*, opt)}$ is given by

$$\text{MSE} \left\{ T_{(s, \sigma)}^{(*, opt)} \right\} = \frac{\sigma^{-2} v_2(n) \left\{ (1 - \ddot{\beta})^4 + (1 - \lambda^*)^2 v_2(n) \ddot{\beta}^4 \right\}}{\left\{ (1 - \ddot{\beta})^2 + v_2(n) \ddot{\beta}^2 \right\}^2} ,$$

which is smaller than the variance of the conventional unbiased estimator $T_{(u, \sigma^{-1})}$ if

$$\frac{\left\{ (1 - \ddot{\beta})^4 + (1 - \lambda^*)^2 v_2(n) \ddot{\beta}^4 \right\}}{\left\{ (1 - \ddot{\beta})^2 + v_2(n) \ddot{\beta}^2 \right\}^2} < 1$$

$$\text{i.e. if } \ddot{\beta} \leq \left\{ 1 + \sqrt{\frac{1}{2} \left\{ (1 - \lambda^*)^2 - v_2(n) \right\}} \right\}^{-1} , \quad (5.14)$$

which shows that $\ddot{\beta}$ should be between zero and one (i.e. $0 < \ddot{\beta} \leq 1$) for gain in efficiency.

Table 5.1 $PRE \{T_{(s,\sigma^{-1})}^{(1)}, T_{(u,\sigma^{-1})}\}$ and $PRE \{T_{(s,\sigma^{-1})}^{(1)}, T_{(MMSE,\sigma^{-1})}\}$

$\lambda^* \downarrow$ $n \rightarrow$	7	9	11	13	15
0.3	99.4 92.24	98.41 94.26	97.72 94.63	97.26 95.38	99.91 95.83
0.5	103.92 91.82	97.64 89.87	94.37 88.67	93.43 88.88	93.08 89.28
0.7	136.25 120.38	121.72 112.03	109.18 102.58	103.69 98.65	96.33 92.4
0.8	170.34 150.5	160.61 147.82	140.94 132.42	132.5 126.05	131.37 126.02
0.9	211.78 187.12	196.21 180.59	203.14 190.87	176.64 168.04	180.62 173.26
1	225.12 198.91	216.16 198.96	228.07 214.28	216.34 205.81	206.5 198.08
1.1	214.58 189.6	191.31 176.08	206.74 194.24	194.1 184.65	164.49 157.78
1.2	194.02 171.43	156.21 143.77	152.7 143.47	167.13 158.99	129.54 124.26
1.3	163.63 144.58	139.9 128.76	109.48 102.86	127.92 121.69	118.46 113.63
1.5	104.83 92.63	121.61 111.92	88.44 83.1	77.53 73.76	85.6 82.12
1.7	82.11 72.55	91.4 84.12	88.81 83.44	74.19 70.57	71.38 68.47
Range of λ^*	[0.45 , 1.51] [0.6 , 1.46]	[0.58 , 1.57] [0.65 , 1.56]	[0.65 , 1.34] [0.69 , 1.31]	[0.69 , 1.37] [0.71 , 1.35]	[0.73 , 1.39] [0.74 , 1.39]

Table 5.2 $PRE \{T_{(s,\sigma^{-1})}^{(2)}, T_{(u,\sigma^{-1})}\}$ and $PRE \{T_{(s,\sigma^{-1})}^{(2)}, T_{(MMSE,\sigma^{-1})}\}$

$\lambda^* \downarrow \quad n \rightarrow$	7	9	11	13	15
0.3	99.68 91.61	95.8 93.7	95.28 94.22	99.93 95.06	99.65 95.59
0.5	104 91.89	97.12 89.39	93.8 88.13	92.91 88.38	92.62 88.85
0.7	139.96 123.67	123.04 113.25	109.56 102.94	103.76 98.71	96.33 92.4
0.8	178.8 157.98	163.98 150.93	143.28 134.62	133.55 127.04	132.03 126.65
0.9	225.48 199.23	204.02 187.78	208.44 195.85	180.73 171.93	183.19 175.73
1	240.15 212.19	227.22 209.13	235.05 220.85	222.48 211.65	211.42 202.81
1.1	226.45 200.08	199.72 183.82	212.19 199.36	197.5 187.89	167.43 160.61
1.2	202.1 178.57	160.24 147.49	155.62 146.21	168.22 160.03	130.2 124.89
1.3	168.75 149.1	141.03 129.8	110.68 103.99	128.28 122.03	118.09 113.28
1.5	106.78 94.35	120.63 111.02	87.75 82.45	77.32 73.56	85.29 81.82
1.7	81.9 72.36	90.85 83.62	87.56 82.27	73.37 69.79	70.85 67.96
Range of λ^*	[0.41 , 1.53] [0.59 , 1.46]	[0.58 , 1.60] [0.65 , 1.57]	[0.65 , 1.34] [0.69 , 1.31]	[0.69 , 1.37] [0.71 , 1.36]	[0.72 , 1.41] [0.74 , 1.39]

Note: **Bold** figures denote the PREs/ranges w.r.t. MMSE estimator

Table 5.3 $PRE \{T_{(s,\sigma^{-1})}^{(3)}, T_{(u,\sigma^{-1})}\}$ and $PRE \{T_{(s,\sigma^{-1})}^{(3)}, T_{(MMSE,\sigma^{-1})}\}$

$\lambda^* \downarrow$ $n \rightarrow$	7	9	11	13	15
0.3	99.1 92.41	99.34 95.57	98.13 95.83	97.10 96.37	96.95 96.64
0.5	102.16 90.27	99.25 91.34	96.41 90.58	95.37 90.73	94.8 90.93
0.7	119.86 105.9	115.4 106.21	107.49 100.99	103.22 98.2	96.25 92.33
0.8	134.79 119.09	146.49 134.83	130.97 123.06	128.21 121.96	128.5 123.26
0.9	158.34 139.91	166.19 152.96	181.79 170.81	160.51 152.69	170.33 163.38
1	167.83 148.29	175.13 161.19	200.82 188.68	192.64 183.26	187.55 179.91
1.1	166.88 147.45	159.09 146.42	184.95 173.77	180.27 171.49	152.89 146.66
1.2	158.82 140.33	139.72 128.6	140.63 132.13	162.01 154.12	126.81 121.64
1.3	140.34 124	134.25 123.56	104.46 98.15	126.02 119.88	119.8 114.92
1.5	95.95 84.78	124.58 114.66	91.46 85.93	78.49 74.67	86.76 83.22
1.7	84.4 74.57	93.08 85.67	93.97 88.29	77.74 73.96	73.65 70.65
Range of λ^*	[0.44 , 1.47] [0.65 , 1.40]	[0.6 , 1.65] [0.67 , 1.60]	[0.64 , 1.32] [0.70 , 1.25]	[0.69 , 1.37] [0.71 , 1.35]	[0.72 , 1.42] [0.75 , 1.40]

Note: **Bold figures** denote the PREs/ranges w.r.t. MMSE estimator

6. Simulation study

In this section we have carried out a simulation study in order to see the performance of various proposed estimators $T_{(s,\theta)}^{(i)}$ ($i = 1, 2, 3$) relative to unbiased and MMSE.

For simulation study a random sample of 10 observations is generated from a normal population with mean $\mu = 20$ and standard deviation $\sigma = 5$ (see, Table 6.1, using SPSS), if the prior estimate of standard deviation σ is available from the past behaviour of the population as $\sigma_0 = 4$, i.e. $\lambda = 0.8$ and $\lambda^* = 1.25$.

Table 6.1 Simulated data from $N(20, 25)$

22.13	16.03	15.85	17.71	25.37
24.73	23.68	11.93	25.91	22.93

The different parametric functions are estimated using their usual estimators and developed modified estimators $T_{(s,\theta)}^{(1)}$, $T_{(s,\theta)}^{(2)}$ and $T_{(s,\theta)}^{(3)}$, where θ may be σ and σ^{-1} and gain in efficiency with respect to conventional unbiased estimator, MMSE estimator and themselves (i.e. $T_{(s,\theta)}^{(1)}$, $T_{(s,\theta)}^{(2)}$ and $T_{(s,\theta)}^{(3)}$) are given in Table 6.2.

It is observed from Table 6.2 that the proposed classes of estimators are more efficient than the usual unbiased and MMSE estimators with considerable gain in efficiency.

Table 6.2 Simulation Results

Parametric Function	True Value	Estimator used	Estimate	RV or RMSE	$T_{(s,\theta)}^{(1)}$	PRE of	
						$T_{(s,\theta)}^{(2)}$	$T_{(s,\theta)}^{(3)}$
Standard Deviation (σ)	5	Unbiased	4.9925	0.057	139.71	141.44	132.87
		MMSE	4.7232	0.0539	132.11	133.75	125.64
		$T_{(s,\theta)}^{(1)}$	4.4197	0.0408	100	101.24	95.105
		$T_{(s,\theta)}^{(2)}$	4.4063	0.0403	98.775	100	93.939
		$T_{(s,\theta)}^{(3)}$	4.4489	0.0429	105.15	106.45	100
Mean Deviation about Mean	3.9894	Unbiased	3.983	0.057	139.71	141.44	132.87
		MMSE	3.769	0.0539	132.11	133.75	125.64
		$T_{(s,\theta)}^{(1)}$	3.526	0.0408	100	101.24	95.105
		$T_{(s,\theta)}^{(2)}$	3.516	0.0403	98.775	100	93.939
		$T_{(s,\theta)}^{(3)}$	3.550	0.0429	105.15	106.45	100
**Process Capability Index	0.1	Unbiased	0.094	0.07379	126.79	129	117.88
		MMSE	0.088	0.06872	118.08	120.14	109.78
		$T_{(s,\theta)}^{(1)}$	0.107	0.0582	100	101.75	92.971
		$T_{(s,\theta)}^{(2)}$	0.107	0.0572	98.28	100	91.374
		$T_{(s,\theta)}^{(3)}$	0.104	0.0626	107.56	109.44	100
Coefficient of Variation	0.25	Unbiased	0.2496	0.057	139.71	141.44	132.87
		MMSE	0.2362	0.0539	132.11	133.75	125.64
		$T_{(s,\theta)}^{(1)}$	0.2198	0.0408	100	101.24	95.105
		$T_{(s,\theta)}^{(2)}$	0.2203	0.0403	98.775	100	93.939
		$T_{(s,\theta)}^{(3)}$	0.2224	0.0429	105.15	106.45	100

Note: ** Specification limits are set as 20 ± 1.5 so that $USL=21.5$ and $LSL=18.5$. θ can be σ or σ^{-1} as the case.

REFERENCES

- JANI, P. N. (1991). "A Class of Shrinkage Estimators for the Scale Parameter of the Experimental Distribution", IEEE Trans. Rel. 40(1), p. 68–70.
- KOUROUKLIS, S. (1994). "Estimation in the Two Parameters of the Exponential Distribution with Prior Information, IEEE. Trans. Rel. 43(3), 446–450.
- PANDEY, B. N. (1979). "On Shrinkage Estimation of Normal Population Variance", Communications in Statistics – Theory and Methods, 8, 359–365.
- PANDEY, B. N. AND SINGH. J. (1977). "Estimation of the Variance of Normal Population Using Prior Information", Journal of the Indian Statistical Association, 15, 141–150.
- SAXENA, S. AND SINGH, H. P. (2004). "Estimating Various Measures in Normal Population Through a Single Class of Estimators". Journal of the Korean Statistical Society, 33:3, pp 323–337.
- SINGH, H. P. AND CHANDER, V. (2008 a). "Some Classes of Shrinkage Estimators for Estimating the Standard Deviation Towards an Interval of Normal Distribution", Model Assisted Statistics and Application, 3, 1, 71–85.
- SINGH, H. P. AND CHANDER, V. (2008 b). "A General Procedure for Estimating Various Measures of Normal Distribution Using Prior Knowledge of Exponential Distribution", Statistics in Transition, New Series 9, (1), 139–158.
- SINGH, H. P. AND CHANDER, V. (2008 c). "Estimating the Variance of an Exponential Distribution in the Presence of Large True Observations", Austrian Journal of Statistics, 32(2), 207–216.
- SINGH, H. P. AND CHANDER, V. (2008 d). "Some Classes of Shrinkage Estimators for Estimating the Scale Parameter Towards an Interval of Exponential Distribution", Journal of Probability and Statistical Science, 6(1), 69–84.
- SINGH, H. P. AND CHANDER, V. (2008 e). "Estimation of Scale Parameters Towards an Interval of Exponential Distribution", Bulletin of Statistics and Economics (BSE), 2, A08,65–71.
- SINGH, H. P. AND SAXENA, S. (2003). "An Improved Class of Shrinkage Estimators for the Variance of a Normal Population", Statistics in Transition, 6, 119–129.
- SINGH, H. P. AND SINGH, R. (1997). "A Class of Shrinkage Estimators for the Variance of a Normal Population", Microelectronics and Reliability, 37, 863–867.

- SINGH, H. P. , SHUKLA, S. K. AND KATYAR, N. P. (1999). "Estimation of Standard Deviation in Normal Distribution with Prior Information", Proceedings of the National Academic Sciences India", 69, 183–189.
- THOMPSON, J. R. (1968). "Some Shrinkage Techniques for Estimating the Mean", Journal of the American Statistical Association, 63, 113–122.