

On shrinkage estimation under divergence loss

Mohammad Arashi^{†*}, Hidekazu Tanaka[‡], Morteza Amini[§]

[†]*Department of Statistics, Faculty of Mathematical Sciences, Ferdowsi University of
Mashhad, Mashhad, Iran*

[‡]*Graduate School of Science, Osaka Prefecture University, Sakai, Japan*

[§]*Department of Statistics, School of Mathematics, Statistics, and Computer Science, College
of Science, University of Tehran, Tehran, Iran*

Email(s): arashi@um.ac.ir, tanaka@ms.osakafu-u.ac.jp, morteza.amini@ut.ac.ir

Abstract. In this paper, the superiority conditions for a general class of shrinkage estimators in the estimation problem of the normal mean are established under divergence loss. This approach is an extension of the work of Ghosh and Mergel (2009).

Keywords: Divergence loss; James-Stein estimate; Shrinkage estimate; Superharmonic function.

1 Introduction

In the line of seminal works of Stein (1956) and James and Stein (1961), there is growing interest in modifying and generalizing the latter work to bring a new estimator of shrinkage type in order to outperform the sample mean. Interesting studies may include in the couple of works done by Baranchik (1970), Efron and Morris (1973), Strawderman (1971), Faith (1978), Stein (1981), Brandwein and Strawderman (1980), Casella (1990), George (1991), Shao et al. (1994), Maruyama (2004), Srivastava and Kubokawa (2005), Ghosh et al. (2008), Wells and Zhou (2008), Ghosh and Mergel (2009) and Arashi and Tabatabay (2010) under different settings.

This work is arisen from the recent study due to Ghosh and Mergel (2009). They considerably investigated on the superiority conditions of Baranchik-type estimators over the sample mean in multivariate normal model, with divergence loss. In this paper, a minor extension is then carried out for another class of shrinkage estimators.

For the precise setup, first of all suppose that $\mathbf{X} \sim \mathcal{N}_p(\boldsymbol{\theta}, \sigma^2 I_p)$, where $\boldsymbol{\theta}$ and σ^2 are both unknown. Further, $S \sim (\sigma^2/(m+2))\chi_m^2$ is independent of $\mathbf{X}$. The aim of this work is to establish

*Corresponding author

Received: 01 November 2024 / Accepted: 15 March 2025

DOI: 10.22067/smeps.2025.90474.1034

conditions in which the general class of shrinkage estimators given by

$$\delta(\mathbf{X}) = \mathbf{X} + S\mathbf{g}(\mathbf{X}) \quad (1)$$

outperforms $\mathbf{X}$, where $\mathbf{g} : \mathbb{R}^p \rightarrow \mathbb{R}^p$ satisfies some regularity conditions which will be given later.

The outline of this paper is as follows: In Section 2, some preliminary results as well as some notations are given, while the main results are exhibited in Section 3. Some important remarks are also given in Section 4.

2 Preliminaries

Before revealing the main results, we express some useful notations. For any $\mathbf{x}, \mathbf{y} \in \mathbb{R}^p$, let

$$\mathbf{x} \cdot \mathbf{y} = \sum_{j=1}^p x_j y_j, \quad \|\mathbf{x}\|^2 = \sum_{i=1}^p x_i^2.$$

Definition 1. A function $h : \mathbb{R}^p \rightarrow \mathbb{R}$ is said to be almost differentiable if there exists a function $\nabla h : \mathbb{R}^p \rightarrow \mathbb{R}^p$ such that, for all $\mathbf{z} \in \mathbb{R}^p$,

$$h(\mathbf{x} + \mathbf{z}) - h(\mathbf{x}) = \int_0^1 \mathbf{z} \cdot \nabla h(\mathbf{x} + t\mathbf{z}) dt$$

for (Lebesgue measure) almost all $\mathbf{x} \in \mathbb{R}^p$. A function $\mathbf{g} : \mathbb{R}^p \rightarrow \mathbb{R}^p$ is almost differentiable if all its coordinate functions are almost differentiable. Essentially, ∇ is the vector differential operator of first partial derivatives with i^{th} coordinate $\nabla_i = \partial / \partial x_i$.

Definition 2. A function $\mathbf{g} : \mathbb{R}^p \rightarrow \mathbb{R}^p$ is said to be homogeneous of degree -1 , if it satisfies

$$\mathbf{g}(\lambda \mathbf{x}) = \frac{1}{\lambda} \mathbf{g}(\mathbf{x})$$

for all real $\lambda \neq 0$ and for all $\mathbf{x} \in \mathbb{R}^p$.

As an example satisfying the regularity condition in Definition 2, consider the reciprocal function $\mathbf{g}(\mathbf{x}) = (g(x_1), \dots, g(x_p))'$ where $g(x) = x^{-1}$. Obviously, $g(\lambda x) = 1/(\lambda x) = \lambda^{-1}g(x)$.

In this paper, we employ the expectation of divergence loss, which includes Kullback-Leibler (KL for short) loss and Bhattacharyya-Hellinger (BH) loss, as a measurement. The divergence loss has been considered by many authors in other contexts. Among others, we refer to Amari (1982) and Cressie and Read (1984).

Definition 3. Suppose that $\mathcal{N}(\mathbf{x}|\boldsymbol{\theta}, \sigma^2 I_p)$ denotes the probability density function of $\mathcal{N}_p(\boldsymbol{\theta}, \sigma^2 I_p)$. For an estimator $\mathbf{a}$ of $\boldsymbol{\theta}$, the divergence loss is defined by

$$\begin{aligned} L_\beta(\boldsymbol{\theta}; \mathbf{a}) &= \frac{1}{\beta(1-\beta)} \left(1 - \int_{\mathbb{R}^p} \mathcal{N}^{1-\beta}(\mathbf{x}|\boldsymbol{\theta}, \sigma^2 I_p) \mathcal{N}^\beta(\mathbf{x}|\mathbf{a}, \sigma^2 I_p) d\mathbf{x} \right) \\ &= \frac{1}{b} \left(1 - \exp \left[-\frac{b}{2\sigma^2} \|\mathbf{a} - \boldsymbol{\theta}\|^2 \right] \right), \end{aligned} \quad (2)$$

where $b = \beta(1-\beta)$ and $\beta \in (0, 1)$.

The second equality in (2) is a consequence of Lemma 2.2 of [Ghosh et al. \(2008\)](#). The KL loss and the BH loss occur as special cases of the divergence loss when $\beta \rightarrow 0$ or $\beta \rightarrow 1$, and $\beta = 1/2$, respectively. Some graphical displays are presented in Figure 1, illustrating the relative behavior of the loss function for $\sigma = 0.5$. In this figure, β varies within the range $(0, 1)$, while $\|\mathbf{a} - \boldsymbol{\theta}\|^2$ changes from small to large values across different panes to demonstrate the effect of its magnitude. As shown, for moderate values of $\|\mathbf{a} - \boldsymbol{\theta}\|^2$, the shape of the loss function is clearly bath-shaped.

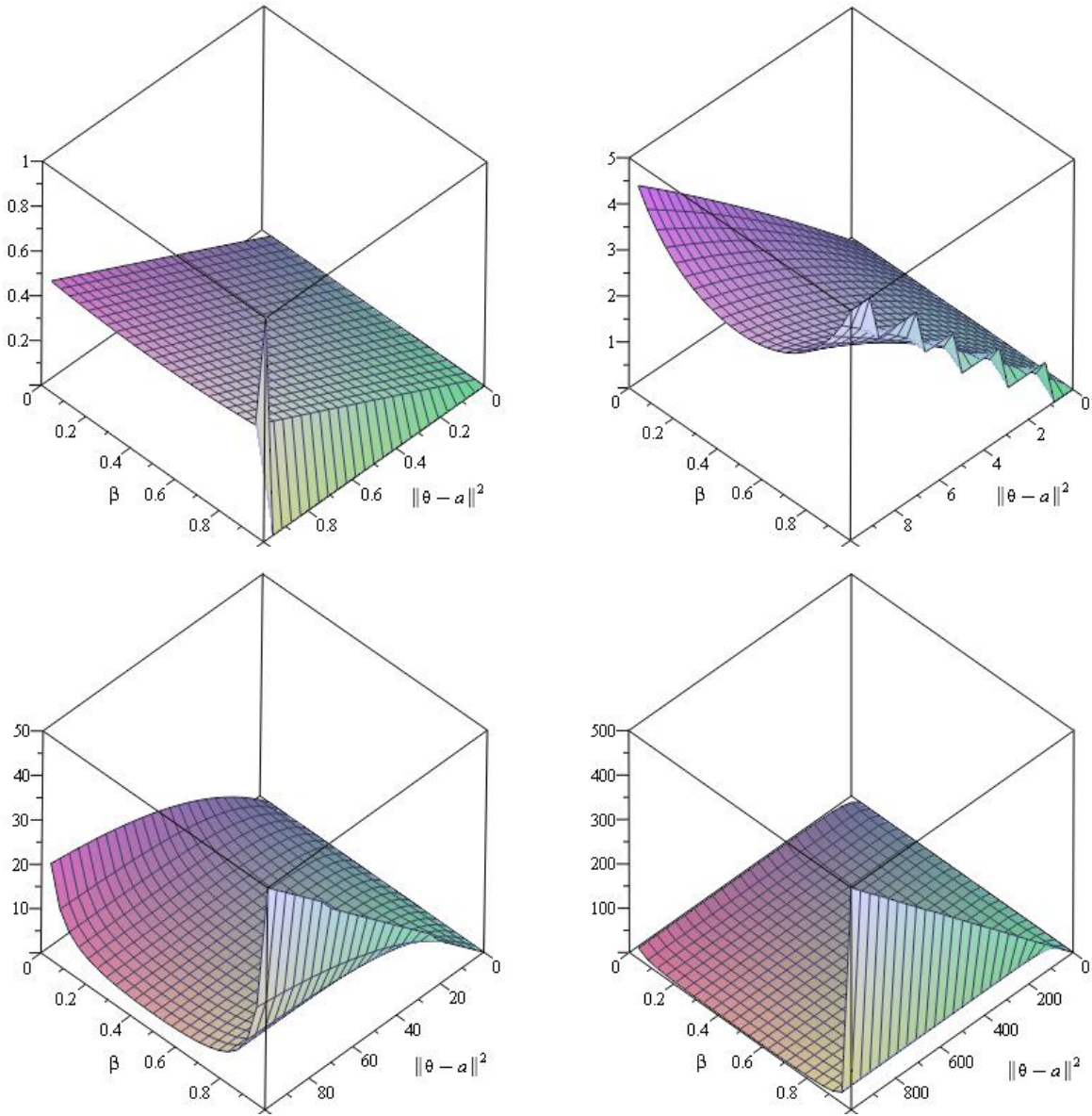


Figure 1: Behavior of the divergence loss relative to β and $\|\mathbf{a} - \boldsymbol{\theta}\|^2$.

To close this section, note that since $\|\mathbf{X} - \boldsymbol{\theta}\|^2 \sim \sigma^2 \chi_p^2$, under the loss, it can be directly concluded that the risk of $\mathbf{X}$ is given by

$$R_\beta(\boldsymbol{\theta}; \mathbf{X}) = \frac{1 - (1+b)^{-p/2}}{b}. \quad (3)$$

3 Main results

In this section, we first give sufficient conditions in which the estimator $\boldsymbol{\delta}(\mathbf{X})$ given by (1) dominates $\mathbf{X}$ for the case of $\Sigma = \sigma^2 I_p$ (Case I). Also, for the case of the general framework Σ (Case II), we give sufficient conditions so that an estimator dominates $\mathbf{X}$.

Case I: ($\Sigma = \sigma^2 I_p$, σ^2 is unknown)

Theorem 1. Assume $\mathbf{X} \sim \mathcal{N}_p(\boldsymbol{\theta}, \sigma^2 I_p)$ and $S \sim (\sigma^2/(m+2))\chi_m^2$ are independent, where $\boldsymbol{\theta}$ and σ^2 are both unknown. Further, assume that for a given estimator $\boldsymbol{\delta}(\mathbf{X})$ in (1), $\mathbf{g}$ is almost differentiable and homogenous function of degree -1 satisfying $E|(\partial/\partial W_i)g_i(\mathbf{W})| < \infty$ for all $\boldsymbol{\eta} \in \mathbb{R}^p$ and for $i = 1, \dots, p$, where $\mathbf{W} \sim \mathcal{N}_p(\boldsymbol{\eta}, I_p/(b+1))$ is independent of S . Then, the estimator $\boldsymbol{\delta}(\mathbf{X})$ has smaller risk than $\mathbf{X}$, under the divergence loss (2), provided that

$$\|\mathbf{g}(\mathbf{w})\|^2 + \frac{2}{b+1} \boldsymbol{\nabla} \cdot \mathbf{g}(\mathbf{w}) \leq 0$$

for all $\mathbf{w} \in \mathbb{R}^p$.

Proof. Let $\mathbf{Y} = \mathbf{X}/\sigma$, $\boldsymbol{\eta} = \boldsymbol{\theta}/\sigma$, $S^* = S/\sigma^2$. Then, by the homogeneity of $\mathbf{g}$ and the inequality $e^x - e^y \geq e^y(x - y)$ ($x, y \in \mathbb{R}$), the risk difference between $\mathbf{X}$ and $\boldsymbol{\delta}(\mathbf{X})$ is given by

$$\begin{aligned} R_\beta(\boldsymbol{\theta}; \mathbf{X}) - R_\beta(\boldsymbol{\theta}; \boldsymbol{\delta}(\mathbf{X})) &= \frac{1}{b} E \left[\exp \left(-\frac{b}{2\sigma^2} \|\mathbf{X} + S\mathbf{g}(\mathbf{X}) - \boldsymbol{\theta}\|^2 \right) - \exp \left(-\frac{b}{2\sigma^2} \|\mathbf{X} - \boldsymbol{\theta}\|^2 \right) \right] \\ &\geq -\frac{1}{2\sigma^2} E \left[\exp \left(-\frac{b}{2\sigma^2} \|\mathbf{X} - \boldsymbol{\theta}\|^2 \right) (S^2 \|\mathbf{g}(\mathbf{X})\|^2 + 2S(\mathbf{X} - \boldsymbol{\theta}) \cdot \mathbf{g}(\mathbf{X})) \right] \\ &= -\frac{1}{2} E \left[\exp \left(-\frac{b}{2} \|\mathbf{Y} - \boldsymbol{\eta}\|^2 \right) (S^{*2} \|\mathbf{g}(\mathbf{Y})\|^2 + 2S^*(\mathbf{Y} - \boldsymbol{\eta}) \cdot \mathbf{g}(\mathbf{Y})) \right]. \quad (4) \end{aligned}$$

Since $\mathbf{Y} \sim \mathcal{N}_p(\boldsymbol{\eta}, I_p)$ and $S^* \sim (m+2)^{-1} \chi_m^2$ are independent, and $E(S^*) = E(S^{*2}) = m/(m+2)$, the right hand side in (4) is rewritten by

$$\begin{aligned} &-\frac{m}{2(m+2)} E \left[\exp \left(-\frac{b}{2} \|\mathbf{Y} - \boldsymbol{\eta}\|^2 \right) (\|\mathbf{g}(\mathbf{Y})\|^2 + 2(\mathbf{Y} - \boldsymbol{\eta}) \cdot \mathbf{g}(\mathbf{Y})) \right] \\ &= -\frac{m}{2(m+2)} (b+1)^{-p/2} E [\|\mathbf{g}(\mathbf{W})\|^2 + 2\mathbf{g}(\mathbf{W}) \cdot (\mathbf{W} - \boldsymbol{\eta})]. \quad (5) \end{aligned}$$

By the Stein identity, we have

$$E[\mathbf{g}(\mathbf{W}) \cdot (\mathbf{W} - \boldsymbol{\eta})] = \frac{1}{b+1} E[\boldsymbol{\nabla} \cdot \mathbf{g}(\mathbf{W})]. \quad (6)$$

Substituting (5) and (6) in (4) concludes that

$$\begin{aligned} R_\beta(\boldsymbol{\theta}; \mathbf{X}) - R_\beta(\boldsymbol{\theta}; \boldsymbol{\delta}(\mathbf{X})) &\geq -\frac{m}{2(m+2)}(b+1)^{-p/2}E \left[\|\mathbf{g}(\mathbf{W})\|^2 + \frac{2}{b+1} \boldsymbol{\nabla} \cdot \mathbf{g}(\mathbf{W}) \right] \\ &\geq 0 \end{aligned} \quad (7)$$

for all $\boldsymbol{\theta} \in \mathbb{R}^p$ and for all $\sigma^2 > 0$. Here, we have to show that the inequality in (7) holds strictly for some $\boldsymbol{\theta}$ and σ^2 . Since $P\{S\mathbf{g}(\mathbf{X}) + 2(\mathbf{X} - \boldsymbol{\theta}) = \mathbf{0}\} = 0$, the equality in (4) holds only if $\mathbf{g}(\mathbf{X}) = \mathbf{0}$ (a.s.), that is, $\boldsymbol{\delta}(\mathbf{X}) = \mathbf{X}$ (a.s.). This completes the proof. $\square$

As a supplement, the earlier result can be proposed in a more general situation. In this regard, we have the following essential consequence.

Theorem 2. Suppose that $\mathbf{X} \sim \mathcal{N}_p(\boldsymbol{\theta}, \sigma^2 \mathbf{I}_p)$ and $S \sim (\sigma^2/(m+2))\chi_m^2$ are independent, where $\boldsymbol{\theta}$ and σ^2 are both unknown. Consider the class of shrinkage estimators

$$\boldsymbol{\delta}^*(\mathbf{X}) = \mathbf{X} + cS\mathbf{g}(\mathbf{X}),$$

where $\mathbf{g}$ is as in Theorem 1. Then, $\boldsymbol{\delta}^*(\mathbf{X})$ outperforms $\mathbf{X}$ under divergence loss provided that

(i) there exists a function $h(\cdot)$ such that, for all $\mathbf{x} \in \mathbb{R}^p$

$$\frac{1}{2}\|\mathbf{g}(\mathbf{x})\|^2 \leq h(\mathbf{x}) \leq -\frac{1}{1+b}\boldsymbol{\nabla} \cdot \mathbf{g}(\mathbf{x}),$$

(ii) $0 < c \leq 1$.

Proof. By the same way as the proof of Theorem 1, the risk difference between $\mathbf{X}$ and $\boldsymbol{\delta}^*(\mathbf{X})$ is given by

$$\begin{aligned} R_\beta(\boldsymbol{\theta}; \mathbf{X}) - R_\beta(\boldsymbol{\theta}; \boldsymbol{\delta}^*(\mathbf{X})) &\geq -\frac{m}{2(m+2)}(b+1)^{-p/2}E \left[c^2\|\mathbf{g}(\mathbf{W})\|^2 + \frac{2c}{b+1}\boldsymbol{\nabla} \cdot \mathbf{g}(\mathbf{W}) \right] \\ &\geq -\frac{m}{m+2}(b+1)^{-p/2}c(c-1)E[h(\mathbf{W})] \\ &\geq 0 \end{aligned} \quad (8)$$

for all $c \in (0, 1]$. By the same reason as Theorem 1, the equality in (8) holds for all $\boldsymbol{\theta}$ and for all σ^2 only if $\boldsymbol{\delta}(\mathbf{X}) = \mathbf{X}$ (a.s.). $\square$

Now, let $\mathcal{S}(p)$ be the set of all positive definite matrices of order p .

Case II: (Unknown $\Sigma \in \mathcal{S}(p)$)

In this case, first of all, let $n-1$ mutually independent random vectors $\mathbf{Z}_1, \dots, \mathbf{Z}_{n-1} \sim \mathcal{N}_p(\mathbf{0}, \Sigma)$ be independent of $\mathbf{Z}_n \sim \mathcal{N}_p(\boldsymbol{\theta}, \Sigma)$, in which $p \geq 3$. We deal with a more general type of improved shrinkage estimators, based on the sufficient statistic $(\mathbf{Z}_n, S)$, of the form

$$\boldsymbol{\gamma}(\mathbf{Z}_n, S) = \mathbf{Z}_n + \mathbf{F}(\mathbf{Z}_n, S), \quad (9)$$

where $S = \sum_{j=1}^{n-1} \mathbf{Z}_j \mathbf{Z}_j'$.

Definition 4. Suppose that $\mathcal{N}(\mathbf{x}|\boldsymbol{\theta}, \Sigma)$ denotes the probability density function of $\mathcal{N}_p(\boldsymbol{\theta}, \Sigma)$. For an estimator $\mathbf{a}$ of $\boldsymbol{\theta}$, the generalized divergence loss (GDL) function is defined by

$$\begin{aligned} L_{\beta}^{\Sigma}(\boldsymbol{\theta}; \mathbf{a}) &= \frac{1}{\beta(1-\beta)} \left(1 - \int_{\mathbb{R}^p} \mathcal{N}^{1-\beta}(\mathbf{x}|\boldsymbol{\theta}, n^{-1}\Sigma) \mathcal{N}^{\beta}(\mathbf{x}|\mathbf{a}, n^{-1}\Sigma) d\mathbf{x} \right) \\ &= \frac{1}{b} \left(1 - \exp \left[-\frac{nb}{2} \|\mathbf{a} - \boldsymbol{\theta}\|_{\Sigma}^2 \right] \right), \end{aligned} \quad (10)$$

where $\|\mathbf{a} - \boldsymbol{\theta}\|_{\Sigma}^2 = (\mathbf{a} - \boldsymbol{\theta})' \Sigma^{-1} (\mathbf{a} - \boldsymbol{\theta})$.

Theorem 3. Suppose that

$$E [\mathbf{F}'(\mathbf{U}, S) \mathbf{F}(\mathbf{U}, S)] < \infty$$

for all $\boldsymbol{\theta} \in \mathbb{R}^p$ and for all $\Sigma \in \mathcal{S}(p)$, where $\mathbf{U} \sim \mathcal{N}_p(\boldsymbol{\theta}, \Sigma/(nb+1))$ is independent of S . Then, the estimator $\gamma(\mathbf{Z}_n, S)$ given by (9) has smaller risk than $\mathbf{Z}_n$, under GDL function given by (10) provided that

$$2\nabla_{\mathbf{u}} \cdot \mathbf{F}(\mathbf{u}, S) + (n-p+2)(nb+1)\mathbf{F}'(\mathbf{u}, S)S^{-1}\mathbf{F}(\mathbf{u}, S) \leq 0,$$

where $\nabla_{\mathbf{u}}$ states getting derivative with respect to $\mathbf{u}$.

Proof. By the same way as the proof of Theorem 1, the risk difference between $\mathbf{Z}_n$ and $\gamma(\mathbf{Z}_n, S)$ is given by

$$R_{\beta}(\boldsymbol{\theta}; \mathbf{Z}_n) - R_{\beta}(\boldsymbol{\theta}; \gamma(\mathbf{Z}_n, S)) \geq -\frac{n}{2} E \left[(\|\mathbf{F}(\mathbf{Z}_n, S)\|_{\Sigma}^2 + 2\mathbf{F}'(\mathbf{Z}_n, S)\Sigma^{-1}(\mathbf{Z}_n - \boldsymbol{\theta})) e^{(-\frac{nb}{2}\|\mathbf{Z}_n - \boldsymbol{\theta}\|_{\Sigma}^2)} \right]. \quad (11)$$

Let $\mathbf{V}_i = \mathbf{Z}_i/\sqrt{nb+1}$ for $i = 1, \dots, n-1$, $T = \sum_{j=1}^{n-1} \mathbf{V}_j \mathbf{V}_j'$, $\mathbf{G}(\mathbf{u}, T) = \mathbf{F}(\mathbf{u}, (nb+1)T)$ and $\Lambda = \Sigma/(nb+1)$. Then, the right hand side in (11) is rewritten by

$$-\frac{n}{2}(nb+1)^{-1-p/2} E [\mathbf{G}'(\mathbf{U}, T)\Lambda^{-1}\mathbf{G}(\mathbf{U}, T) + 2\mathbf{G}'(\mathbf{U}, T)\Lambda^{-1}(\mathbf{U} - \boldsymbol{\theta})].$$

By applying Lemma 1 of Fourdrinier et al. (2003), we see that

$$\begin{aligned} &E [\mathbf{G}'(\mathbf{U}, T)\Lambda^{-1}\mathbf{G}(\mathbf{U}, T) + 2\mathbf{G}'(\mathbf{U}, T)\Lambda^{-1}(\mathbf{U} - \boldsymbol{\theta})] \\ &\leq E [2\nabla_{\mathbf{U}} \cdot \mathbf{G}(\mathbf{U}, T) + (n-p+2)\mathbf{G}(\mathbf{U}, T)'T^{-1}\mathbf{G}(\mathbf{U}, T)]. \end{aligned} \quad (12)$$

Since $T = S/(nb+1)$, the right hand side in (12) is expressed as

$$E [2\nabla_{\mathbf{U}} \cdot \mathbf{F}(\mathbf{U}, S) + (n-p+2)(nb+1)\mathbf{F}'(\mathbf{U}, S)S^{-1}\mathbf{F}(\mathbf{U}, S)]. \quad (13)$$

From (11) to (13), we have

$$\begin{aligned} R_{\beta}(\boldsymbol{\theta}; \mathbf{Z}_n) - R_{\beta}(\boldsymbol{\theta}; \gamma(\mathbf{Z}_n, S)) &\geq -\frac{n}{2}(nb+1)^{-1-p/2} \\ &\quad \times E [2\nabla_{\mathbf{U}} \cdot \mathbf{F}(\mathbf{U}, S) + (n-p+2)(nb+1)\mathbf{F}'(\mathbf{U}, S)S^{-1}\mathbf{F}(\mathbf{U}, S)] \\ &\geq 0 \end{aligned} \quad (14)$$

for all $\boldsymbol{\theta} \in \mathbb{R}^p$ and $\Sigma \in \mathcal{S}(p)$. By the same reason as Theorem 1, the equality in (14) holds only if $\mathbf{F}(\mathbf{Z}_n, S) = \mathbf{0}$ (a.s.). $\square$

4 Conclusions

In this paper, we establish the conditions under which a general class of shrinkage estimators outperforms the consistent estimator of a multivariate normal population mean when using a divergence loss function. The results obtained apply to both known and unknown covariance scenarios. Furthermore, our findings extend the earlier work of [Ghosh and Mergel \(2009\)](#) to a broader class of dominant estimators. The proofs of the proposed theorems are presented in a more straightforward manner than in the previous reference. This method can also be applied to exponential-type loss functions, such as LINEX, and reflected normal losses. Importantly, the conditions for superiority are robust concerning the squared error loss function, which further validates our derivations, as the divergence loss encompasses the square error loss as a special case.

4.1 Easy understanding

In conclusion, we provide a simple approach for making decisions regarding the superiority conditions based on divergence loss. As previously mentioned, the square error loss is a specific case of divergence loss. It is important to note that the graphs of both losses are parabolic. Additionally, the graph of the risk of $\mathbf{X}$ in Equation (3) also follows a parabolic form (see Figure 2). From a graphical perspective, for any estimator of $\boldsymbol{\theta}$ to be superior to $\mathbf{X}$, it must have a risk that is lower than that of the parabolic shape. Therefore, we can conclude that determining the superiority conditions for the proposed shrinkage estimators can be achieved by examining them under the square error loss. Interestingly, the conditions derived in this study align exactly with those found in the work of [Ghosh and Mergel \(2009\)](#), as they are the same under square error loss when $b \rightarrow 0$.

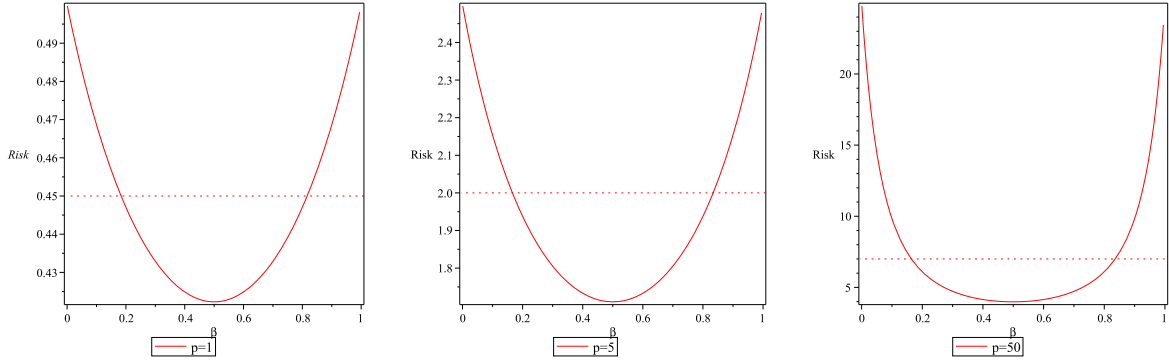


Figure 2: Risk of $\mathbf{X}$ under the divergence loss.

Acknowledgments

We would like to express our gratitude to the two anonymous reviewers for their constructive comments, which prompted us to include many additional details in the paper and significantly

improve its presentation. Mohammad Arashi's work is based on the research supported in part by the Iran National Science Foundation (INSF) grant No. 4015320.

References

- Amari, S. (1982). Differential geometry of curved exponential families-curvatures and information loss. *The Annals of Statistics*, **10**, 357–387.
- Arashi, M. and Tabatabaey, S. M. M. (2010). A note on classical Stein-type estimators in elliptically contoured models. *Journal of Statistical Planning and Inference*, **140**, 1206–1213.
- Baranchik, A. J. (1970). A family of minimax estimators of the mean of a multivariate normal distribution. *The Annals of Mathematical Statistics*, **41**, 642–645.
- Brandwein, A. C. and Strawderman, W. E. (1980). Minimax estimators of location parameters for spherically symmetric distributions with concave loss. *The Annals of Statistics*, **8**, 279–284.
- Casella, G. (1990). Estimators with nondecreasing risk: Application of a chi-square identity. *Statistics and Probability Letters*, **10**, 107–109.
- Cressie, N. and Read, T. R. C. (1984). Multinomial goodness-of-fit tests. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, **46**, 440–464.
- Efron, B. and Morris, C. (1973). Stein's estimation rule and its competitors-An empirical Bayes approach. *Journal of the American Statistical Association*, **68**, 117–130.
- Faith, R. E. (1978). Minimax Bayes estimators of a multivariate normal mean. *Journal of Multivariate Analysis*, **8**, 372–379.
- Fourdrinier, D., Strawderman, W. E. and Wells, M. T. (2003). Robust shrinkage estimation for elliptically symmetric distributions with unknown covariance matrix. *Journal of Multivariate Analysis*, **85**, 24–39.
- George, E. I. (1991). Shrinkage domination in a multivariate common mean problem. *The Annals of Statistics*, **19**, 952–960.
- Ghosh, M., Mergel, V. and Datta, G. S. (2008). Estimation, prediction and the Stein phenomenon under divergence loss, *Journal of Multivariate Analysis*, **99**, 1941–1961.
- Ghosh, M. and Mergel, V. (2009). On the Stein phenomenon under divergence loss and an unknown variance-covariance matrix. *Journal of Multivariate Analysis*, **100**, 2331–2336.
- James, W. and Stein, C. (1961). Estimation of quadratic loss. In *Proceedings of the Fourth Berkeley Symposium on Mathematical Statistics and Probability*, **1**, 361–379.
- Maruyama, Y. (2004). Stein's idea and minimax admissible estimation of a multivariate normal mean, *Journal of Multivariate Analysis*, **88**, 320–334.

- Shao, P. Yi-Shi and Strawderman, W. E. (1994). Improving on the James-Stein positive-part estimator. *The Annals of Statistics*, **22**, 1517–1538.
- Srivastava, M. S. and Kubokawa, T. (2005). Minimax multivariate empirical Bayes estimators under multicollinearity. *Journal of Multivariate Analysis*, **93**, 394–416.
- Stein, C. (1981). Estimation of the mean of a multivariate normal distribution. *The Annals of Statistics*, **9**, 1135–1151.
- Stein, C. (1956). Inadmissibility of the usual estimator for the mean of a multivariate distribution. In *Proceedings of the Third Berkeley Symposium on Mathematical Statistics and Probability*, **1**, 197–206.
- Strawderman, W. E. (1971). Proper Bayes minimax estimators of the multivariate normal mean. *The Annals of Mathematical Statistics*, **42**, 385–388.
- Wells, M. T. and Zhou, G. (2008). Generalized Bayes minimax estimators of the mean of multivariate normal distribution with unknown variance. *Journal of Multivariate Analysis*, **99**, 2208–2220.