

Article

Not peer-reviewed version

Intrinsically Optimal Equivariant Estimator for the Univariate Linear Normal Models

[Gloria García](#), [Marta Cubedo](#)^{*}, [Josep Maria Oller](#)

Posted Date: 13 October 2025

doi: 10.20944/preprints202510.0742.v1

Keywords: Equivariant estimator; Information metric; Riemannian risk; Intrinsic bias



Preprints.org is a free multidisciplinary platform providing preprint service that is dedicated to making early versions of research outputs permanently available and citable. Preprints posted at Preprints.org appear in Web of Science, Crossref, Google Scholar, Scilit, Europe PMC.

Copyright: This open access article is published under a Creative Commons CC BY 4.0 license, which permit the free download, distribution, and reuse, provided that the author and preprint are cited in any reuse.

Disclaimer/Publisher's Note: The statements, opinions, and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions, or products referred to in the content.

Article

Intrinsically Optimal Equivariant Estimator for the Univariate Linear Normal Models

Gloria García ¹, Marta Cubedo ^{2,*} and Josep M. Oller ²

¹ Department of Education and Professional Training, Gencat, 08021 Barcelona, Spain; ggarcia76@xtec.cat

² University of Barcelona, Avda. Diagonal 643, 08028 Barcelona, Spain; joller@ub.edu

* Correspondence: mcubedo@ub.edu

Abstract

In this paper we prove the existence and unicity of the minimum risk equivariant estimator (MIRE) for the univariate linear normal models, under the action of the subgroup of the affine group which leaves the column space of the design matrix invariant and using the framework of the intrinsic analysis of statistical estimation which uses the square of the Rao distance as the loss function, which expectation behaves very different to the corresponding of the square error loss when we work with small samples, and supplying also an intrinsic bias measure for any equivariant estimator. This estimator is studied and compared with the standard maximum likelihood estimator (MLE) in terms of its intrinsic risk and bias, showing explicitly its differences, apparent with small samples, fortunately small for large sample sizes. Moreover, an approximate and very simple estimator is also suggested, which reduces the majority of bad behaviour of MLE, of intrinsic risk and bias, for small samples.

Keywords: equivariant estimator; information metric; riemannian risk; intrinsic bias

PACS: 62F00, 60K35, 62A00

1. Introduction

The Fisher information matrix of a regular parametric family of probability distributions induces a natural Riemannian structure on the parameter space [1–3]. This structure provides the foundation for the intrinsic analysis of statistical estimation, see [4,5] and, beyond statistics, it has also been shown to reveal connections with physical laws [6].

In this intrinsic framework, estimators are assessed through tools that do not depend on any particular parametrization of the model. Consequently, non-intrinsic criteria such as the squared error loss should be replaced by intrinsic measures. A natural choice is the squared Riemannian (Rao) distance, which serves as an intrinsic loss function. The associated risk can behave very differently from that based on squared error loss, especially in small-sample regimes. Similarly, the classical definition of bias must be reformulated in terms of a vector field determined by the geometry of the model, whose squared norm provides a natural intrinsic bias measure.

The estimator itself should also be intrinsic, i.e., independent of parametrization. Consider, for instance, the exponential distribution under two parametrizations,

$$f(x, \theta) = \theta \exp(-\theta x) \mathbb{1}_{\mathbb{R}^+}(x) \quad \text{and} \quad f(x, \lambda) = \frac{1}{\lambda} \exp(-x/\lambda) \mathbb{1}_{\mathbb{R}^+}(x)$$

where $\mathbb{1}_{\mathbb{R}^+}(x)$ is the positive real numbers indicator. The UMVU estimator computed under each parametrization $\hat{\theta}$ and $\hat{\lambda}$ are not related by $\hat{\theta} = 1/\hat{\lambda}$. This apparent inconsistency arises because UMVU estimators rely on non-intrinsic notions such as unbiasedness and variance. As a result, the UMVU estimator is parametrization-dependent and cannot be regarded as intrinsically defined.

When the statistical model is invariant under the action of a transformation group on the sample space, it is natural to restrict attention to equivariant estimators, i.e., estimators U satisfying

$$U(g(x)) = \bar{g}(U(x)) \quad \text{for all } g \in \mathcal{G}_n$$

where $\mathcal{G}_n$ is the transformation group acting on the space of all samples of size n and $\bar{g} \in \bar{\mathcal{G}}$ is the induced transformation on the parameter space Θ . This restriction ensures logical consistency: only equivariant estimators remain coherent under data transformations. It is therefore a requirement that should be imposed prior to any attempt to minimize an intrinsic risk function.

As an illustration, consider the estimation of the unconstrained mean μ of a p -variate normal distribution. The MLE of μ (which really defines an intrinsic estimator) based on a sample of size n , is just the sample mean vector $\bar{X}_n$. However, for $p \geq 3$, the James–Stein estimator is known to have smaller quadratic risk [7]. In our view, this does not undermine the MLE. Although the James–Stein estimator enjoys a lower risk under the non-intrinsic quadratic loss, it lacks the essential property of equivariance, which is crucial from the intrinsic perspective.

In this work we focus on classical statistical estimation, as is standard in the absence of prior knowledge. Nonetheless, intrinsic Bayesian approaches based on non-informative priors are also possible [5,8], and represent an interesting avenue for future research. Here, we adopt the squared Rao distance as a natural loss function, since it is the conceptually simplest intrinsic analogue of the quadratic loss, despite potential computational challenges. For a broader background on statistical estimation and its intrinsic developments, see [9–12].

Specifically, we consider the univariate linear normal model with a fixed design matrix. First, we explicitly characterize the class of equivariant estimators under the action of a suitable subgroup of the affine group—namely, those affine transformations of the data that leave the column space of the design matrix unchanged. This subgroup is the largest one that preserves the model’s structure. Next, we prove the existence and uniqueness of the estimator within this class that minimizes intrinsic risk, i.e., the equivariant estimator that minimizes the mean squared Rao distance. We also derive an explicit expression for the intrinsic bias of any equivariant estimator.

Furthermore, we compare the intrinsic bias and risk of the proposed estimator with those of the MLE, highlighting their differences in small samples and showing how these differences diminish as the sample size grows. Finally, we propose a computable approximation of the optimal estimator, which mitigates most of the intrinsic bias and risk issues that affect the MLE in finite samples.

2. Equivariant estimators for linear models

Let us consider the univariate linear normal model,

$$\mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \mathbf{e} \quad (1)$$

where $\mathbf{y}$ is a $n \times 1$ random vector, $\mathbf{X}$ is a $n \times m$ matrix of known constants with $0 < \text{rank}(\mathbf{X}) = m \leq n$, $\boldsymbol{\beta}$ is a $m \times 1$ vector of unknown parameters to be estimated and $\mathbf{e}$ is the fluctuation or error of $\mathbf{y}$ about $\mathbf{X}\boldsymbol{\beta}$. We assume that the errors are unbiased, independent, with the same variance σ^2 and following a n -variate normal distribution, that is $\mathbf{e} \sim N_n(\mathbf{0}, \sigma^2 \mathbf{I}_n)$, where $\mathbf{I}_n$ is the $n \times n$ identity matrix. Therefore $\mathbf{y}$ distributes according to an element of the parametric family of probability distributions $\mathcal{P} = \{N_n(\mathbf{X}\boldsymbol{\beta}, \sigma^2 \mathbf{I}_n) \mid (\boldsymbol{\beta}, \sigma) \in \Theta\}$ with parameter space $\Theta = \mathbb{R}^m \times \mathbb{R}_+$, a $(m+1)$ -dimensional simply connected real manifold. Hereafter, we shall identify the elements of $\mathbb{R}^m$ with $m \times 1$ column vectors, when necessary.

Denote by

$$O_E(n) = \{\mathbf{H} \in O(n) \mid \mathbf{y} \in E \Rightarrow \mathbf{H}\mathbf{y} \in E\}$$

where E is a subspace of $\mathbb{R}^n$ and $O(n)$ is the group of $n \times n$ orthogonal matrices with entries from $\mathbb{R}$. Define F as the subspace of $\mathbb{R}^n$ spanned by the columns of $\mathbf{X}$, that is $F = \text{Col}(\mathbf{X})$. Observe that $O_F(n) = O_{F^\perp}(n)$, $\mathbf{I}_n \in O_F(n)$, $\text{Col}(\mathbf{X}) = \text{Col}(\mathbf{H}\mathbf{X}) \quad \forall \mathbf{H} \in O_F(n)$ and if $\mathbf{H} \in O_F(n)$ then $\mathbf{H}^t \in O_F(n)$. Moreover, every $\mathbf{H} \in O_F(n)$ induces, in F and $F^\perp$, two isomorphisms preserving the Euclidean norm in each subspace.

$\mathcal{P}$ is invariant under the action of the subgroup of the affine group in $\mathbb{R}^n$ given by the family of transformations

$$g_{(a, H, c)}(\mathbf{y}) = a\mathbf{H}\mathbf{y} + \mathbf{c}, \mathbf{y} \in \mathbb{R}^n \quad (2)$$

where $a > 0, \mathbf{H} \in O_F(n), \mathbf{c} \in F$.

Observe that $\mathcal{G} = \{g_{(a, H, c)} \mid a > 0, \mathbf{H} \in O_F(n), \mathbf{c} \in F\}$ induces an action on the parameter space θ given by

$$\overline{g_{(a, H, c)}}(\boldsymbol{\beta}, \sigma) = (a(\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t \mathbf{H} \mathbf{X} \boldsymbol{\beta} + (\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t \mathbf{c}, a\sigma) \quad (3)$$

where this result is obtained taking into account that $\mathbf{X}(\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t$ is the projection matrix into F and thus, if $\mathbf{w} \in F$ there exists a unique $\boldsymbol{\eta}$ such that $\mathbf{w} = \mathbf{X}\boldsymbol{\eta}$ and $\mathbf{X}(\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t \mathbf{w} = \mathbf{w}$.

Since the family $\mathcal{P}$ is invariant under the action of $\mathcal{G}$, it is natural to restrict our attention to the class of equivariant estimators $\mathbf{U} = (\mathbf{U}_1, \mathbf{U}_2)$ of $(\boldsymbol{\beta}, \sigma)$ i.e. an estimator satisfying $\mathbf{U}(g_{(a, H, c)}(\mathbf{y})) = \overline{g_{(a, H, c)}}(\mathbf{U}(\mathbf{y}))$ for all $g_{(a, H, c)} \in \mathcal{G}$.

Proposition 1. Let $\mathbf{U}$ be an equivariant estimator of $(\boldsymbol{\beta}, \sigma)$. Then $\mathbf{U}$ belongs to the family $\{\mathbf{U}_\lambda, \lambda \in \mathbb{R}_+\}$ where,

$$\mathbf{U}_\lambda(\mathbf{y}) = \left((\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t \mathbf{y}, \lambda \left\{ \mathbf{y}^t (\mathbf{I}_n - \mathbf{X}(\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t) \mathbf{y} \right\}^{1/2} \right)$$

Proof. Let $\mathbf{U} = (\mathbf{U}_1, \mathbf{U}_2)$. The equivariance condition for $\mathbf{U}$ involves,

$$\begin{aligned} \mathbf{U}_1(a\mathbf{H}\mathbf{y} + \mathbf{c}) &= a(\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t \mathbf{H} \mathbf{X} \mathbf{U}_1(\mathbf{y}) + (\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t \mathbf{c} \\ \mathbf{U}_2(a\mathbf{H}\mathbf{y} + \mathbf{c}) &= a \mathbf{U}_2(\mathbf{y}) \end{aligned}$$

for any $a > 0, \mathbf{H} \in O_F(n), \mathbf{c} \in F$.

Any $\mathbf{y} \in \mathbb{R}^n$ can be written in a unique form as $\mathbf{y} = \mathbf{y}_F + \mathbf{y}_{F^\perp}$ where $\mathbf{y}_F \in F$ and $\mathbf{y}_{F^\perp} \in F^\perp$. Specifically

$$\mathbf{y}_F = \mathbf{X}(\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t \mathbf{y} \quad \text{and} \quad \mathbf{y}_{F^\perp} = (\mathbf{I}_n - \mathbf{X}(\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t) \mathbf{y} \quad (4)$$

If we choose $\mathbf{c} = -a\mathbf{H}\mathbf{y}_F$ in the previous expressions, we obtain

$$\mathbf{U}_1(a\mathbf{H}\mathbf{y}_{F^\perp}) = a(\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t \mathbf{H} \mathbf{X} \mathbf{U}_1(\mathbf{y}) - a(\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t \mathbf{H} \mathbf{y}_F \quad (5)$$

$$\mathbf{U}_2(a\mathbf{H}\mathbf{y}_{F^\perp}) = a \mathbf{U}_2(\mathbf{y}) \quad (6)$$

First we focus on $\mathbf{U}_1$. If we let $\mathbf{H}^* = (\mathbf{I}_n - 2\mathbf{X}(\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t)$, we have $\mathbf{H}^* \in O_F(n)$ with

$$\mathbf{H}^* \mathbf{v} = -\mathbf{v} \quad \forall \mathbf{v} \in F \quad \text{and} \quad \mathbf{H}^* \mathbf{v} = \mathbf{v} \quad \forall \mathbf{v} \in F^\perp$$

Now observe that (5) is satisfied for any $a > 0$ and $\mathbf{H}$ in $O_F(n)$, in particular for $\mathbf{I}_n$ and $\mathbf{H}^*$. Therefore

$$\begin{aligned} \mathbf{U}_1(a\mathbf{y}_{F^\perp}) &= a\mathbf{U}_1(\mathbf{y}) - a(\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t \mathbf{y}_F \\ \mathbf{U}_1(a\mathbf{H}^* \mathbf{y}_{F^\perp}) &= -a\mathbf{U}_1(\mathbf{y}) + a(\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t \mathbf{y}_F \end{aligned}$$

But $a\mathbf{H}^* \mathbf{y}_{F^\perp} = a\mathbf{y}_{F^\perp}$ which leads to

$$0 = \mathbf{U}_1(\mathbf{y}) - (\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t \mathbf{y}_F$$

From (4),

$$\mathbf{U}_1(\mathbf{y}) = (\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t \mathbf{y}, \quad \mathbf{y} \in \mathbb{R}^n$$

Next, we consider U_2 . Let $a = 1$ and $H = I_n$ in (6), it follows

$$U_2(\mathbf{y}_{F^\perp}) = U_2(\mathbf{y}) \quad (7)$$

Accordingly, it is enough to determine U_2 on $F^\perp$. Let us take a unit vector $\mathbf{z} \in F^\perp$ such that $U_2(\mathbf{z}) = \lambda > 0$. Then (6).

$$U_2(H\mathbf{z}) = U_2(\mathbf{z})$$

for any $H \in O_F(n)$. Observe that any arbitrary unit vector in $F^\perp$ can be written as $H\mathbf{z}$ for a proper $H \in O_F(n)$. Therefore, for any $\mathbf{y}_{F^\perp} \in F^\perp$, $\mathbf{y}_{F^\perp} \neq \mathbf{0}$, we have

$$U_2(\mathbf{y}_{F^\perp}) = \|\mathbf{y}_{F^\perp}\| U_2\left(\frac{\mathbf{y}_{F^\perp}}{\|\mathbf{y}_{F^\perp}\|}\right) = \|\mathbf{y}_{F^\perp}\| U_2(H\mathbf{z}) = \|\mathbf{y}_{F^\perp}\| U_2(\mathbf{z}) = \lambda \|\mathbf{y}_{F^\perp}\|$$

Observe also that if $\mathbf{y} = \mathbf{0}$ then $\mathbf{y}_{F^\perp} = \mathbf{0}$ and, choosing $\mathbf{c} = \mathbf{0}$, we have that $U_2(\mathbf{0}) = a U_2(\mathbf{0})$, $\forall a \in \mathbb{R}_+$. This implies $U_2(\mathbf{0}) = 0$.

Finally, from (7) and (4) we have

$$U_2(\mathbf{y}) = \lambda \left\{ \mathbf{y}^t \left(I_n - \mathbf{X}(\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t \right) \mathbf{y} \right\}^{1/2}$$

□

Observe that the standard maximum likelihood estimator, MLE, for the present model is an equivariant estimator, with $\lambda = \frac{1}{\sqrt{n}}$.

3. Minimum Riemannian risk estimators

In the framework of intrinsic analysis, where the loss function is the square of the Rao distance, the Riemannian distance induced by the information metric in the parameter space Θ , once the class of the equivariant estimators has been determined a natural question arises: which is the equivariant estimator that minimizes the risk?

First of all, we summarize the basic geometric results corresponding to the model (1) which are going to be used hereafter. We are going to use a *standardized* version of the information metric, given by the usual *information metric* corresponding to this linear model divided by a constant factor n , i.e. the number of rows of matrix $\mathbf{X}$. This metric is given by

$$ds^2 = \frac{1}{n \sigma^2} \left(d\boldsymbol{\beta}^t \mathbf{X}^t \mathbf{X} d\boldsymbol{\beta} + 2n (d\sigma)^2 \right), \quad (8)$$

which is, up to a linear coordinate change, the *Poincaré hyperbolic metric* of the upper half space $\mathbb{R}^m \times \mathbb{R}^+$, see [13]. The Riemannian curvature $\kappa = -\frac{1}{2}$ is constant and negative and the unique geodesic, parameterized by the arc-length, which connects two points $\theta_1 = (\boldsymbol{\beta}_1, \sigma_1)$ and $\theta_2 = (\boldsymbol{\beta}_2, \sigma_2) \in \Theta$, when $\boldsymbol{\beta}_1 \neq \boldsymbol{\beta}_2$, is given by:

$$\begin{aligned} \boldsymbol{\beta}(s) &= 2n K^{-2} \tanh\left(\frac{s}{\sqrt{2}} + \epsilon\right) (\mathbf{X}^t \mathbf{X})^{-1/2} \mathbf{C} + \mathbf{D} \\ \sigma(s) &= K^{-1} \operatorname{sech}\left(\frac{s}{\sqrt{2}} + \epsilon\right) \end{aligned} \quad (9)$$

where s is the arc-length, $\mathbf{C}$ and $\mathbf{D}$ are $m \times 1$ vectors whose components, and also ϵ , are convenient real integration constants, such that $\boldsymbol{\beta}(0) = \boldsymbol{\beta}_1$, $\boldsymbol{\beta}(\rho_{12}) = \boldsymbol{\beta}_2$, $\sigma(0) = \sigma_1$, $\sigma(\rho_{12}) = \sigma_2$ being ρ_{12} the

Riemannian distance between θ_1 and θ_2 . Finally, K is given by $K^2 = 2n \mathbf{C}^t \mathbf{C}$. When $\beta_1 = \beta_2$, the geodesic is given by

$$\beta(s) = \mathbf{D} \quad \sigma(s) = B e^{\pm \frac{s}{\sqrt{2}}} \quad (10)$$

where B is a positive integration constant.

The Rao distance ρ between the points θ_1 and θ_2 is

$$\rho_{12} \equiv \rho(\theta_1, \theta_2) = \sqrt{2} \ln \left(\frac{1 + \delta(\theta_1, \theta_2)}{1 - \delta(\theta_1, \theta_2)} \right) = 2\sqrt{2} \operatorname{arctanh}(\delta(\theta_1, \theta_2)) \quad (11)$$

where

$$\delta(\theta_1, \theta_2) = \left(\frac{d_M^2(\beta_1, \beta_2) + 2 \frac{(\sigma_1 - \sigma_2)^2}{\sigma_1 \sigma_2}}{d_M^2(\beta_1, \beta_2) + 2 \frac{(\sigma_1 + \sigma_2)^2}{\sigma_1 \sigma_2}} \right)^{1/2} \quad (12)$$

and

$$d_M^2(\theta_1, \theta_2) = \frac{1}{n \sigma_1 \sigma_2} (\beta_1 - \beta_2)^t \mathbf{X}^t \mathbf{X} (\beta_1 - \beta_2) \quad (13)$$

or, equivalently,

$$\rho(\theta_1, \theta_2) = \sqrt{2} \operatorname{arccosh} \left(\frac{1}{4} d_M^2(\theta_1, \theta_2) + \frac{\sigma_1^2 + \sigma_2^2}{2 \sigma_1 \sigma_2} \right) \quad (14)$$

Let $\exp_{\theta_1}^{-1}(\theta_2)$ be the inverse of the exponential map corresponding to Levi-Civita connection and $W^1, \dots, W^m, W^{m+1}$ its components corresponding to the basis field $(\partial/\partial\beta^1)_{\theta_1}, \dots, (\partial/\partial\beta^m)_{\theta_1}, (\partial/\partial\sigma)_{\theta_1}$. Then, we have

$$\begin{aligned} W^i &= \frac{\rho_{12}/\sqrt{2}}{\sinh(\rho_{12}/\sqrt{2})} (\beta_2^i - \beta_1^i) \frac{\sigma_1}{\sigma_2}, \quad i = 1, \dots, m \\ W^{m+1} &= \frac{\rho_{12}/\sqrt{2}}{\sinh(\rho_{12}/\sqrt{2})} \left(\cosh(\rho_{12}/\sqrt{2}) - \frac{\sigma_1}{\sigma_2} \right) \sigma_1 \end{aligned} \quad (15)$$

It is well known that the Riemannian distance induced by the information metric is invariant under equivariant estimator transformations. We shall supply a direct and alternative proof for the linear model setting.

Proposition 2. *The Rao distance ρ given by (11) is invariant under the action of the induced group by $\mathcal{G}$ on the parameter space, $\overline{\mathcal{Q}}$. In other words*

$$\rho(\overline{g(a, H, c)} \theta_1, \overline{g(a, H, c)} \theta_2) = \rho(\theta_1, \theta_2)$$

Proof: Observe that

$$\mathbf{H}\mathbf{X}(\beta_1 - \beta_2) \in F, \forall \mathbf{H} \in O_F(n)$$

and taking into account that $\mathbf{X}(\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t$ is the projection matrix into F , we have

$$\mathbf{X}(\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t \mathbf{H}\mathbf{X}(\beta_1 - \beta_2) = \mathbf{H}\mathbf{X}(\beta_1 - \beta_2)$$

Therefore

$$(\beta_1 - \beta_2)^t \mathbf{X}^t \mathbf{H}^t \mathbf{X}(\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t \mathbf{H}\mathbf{X}(\beta_1 - \beta_2) = (\beta_1 - \beta_2)^t \mathbf{X}^t \mathbf{X}(\beta_1 - \beta_2)$$

and the invariance of δ and ρ trivially follows. $\square$

Proposition 3. $\bar{\mathcal{G}}$ acts transitively on Θ .

Proof: The transitivity follows observing that a is an arbitrary positive real number and $\mathbf{X}(\mathbf{X}^t\mathbf{X})^{-1}\mathbf{X}^t$ is the projection matrix into F with rank $\mathbf{X} = m$. $\square$

Since ρ , and thus ρ^2 , is invariant under the action of $\bar{\mathcal{G}}$ and $\bar{\mathcal{G}}$ acts transitively on Θ , the distribution of $\rho^2(\mathbf{U}_\lambda(\mathbf{y}), \theta)$ does not depend on θ , and therefore, the risk of any equivariant estimator remains constant and independent of the target parameter provided that this risk is finite. More precisely, observe that if we let

$$\mathbf{z} = \frac{1}{\sigma} (\mathbf{y}_F - \mathbf{X}\boldsymbol{\beta}), \quad \mathcal{V} = \|\mathbf{z}\| \quad \text{and} \quad \mathcal{W} = \frac{1}{\sigma} \|\mathbf{y}_{F^\perp}\| \quad (16)$$

from (1) and (4) we clearly have that $\mathbf{z} \sim N_n(\mathbf{0}, \mathbf{X}(\mathbf{X}^t\mathbf{X})^{-1}\mathbf{X}^t)$ with a rank m idempotent covariance matrix, and $\mathcal{V}^2$ and $\mathcal{W}^2$ are independent random variables following a chi-square distribution with m and $(n - m)$ degrees of freedom, equal to the dimensions of F and $F^\perp$ since $\mathcal{V}^2$ and $\mathcal{W}^2$ are quadratic forms based on the projection matrices on these subspaces of $\mathbb{R}^n$ and $\mathbf{y}_F$ (or $\mathbf{z}$) and $\mathbf{y}_{F^\perp}$ are independent random vectors. Therefore, since $\mathbf{X}\mathbf{U}_{\lambda 1}(\mathbf{y}) = \mathbf{y}_F$ and $\mathbf{U}_{\lambda 2}(\mathbf{y}) = \lambda \|\mathbf{y}_{F^\perp}\|$, we have that

$$d_M^2(\mathbf{U}_\lambda(\mathbf{y}), \theta) = \frac{1}{n \mathbf{U}_{\lambda 2}(\mathbf{y}) \sigma} (\mathbf{U}_{\lambda 1}(\mathbf{y}) - \boldsymbol{\beta})^t \mathbf{X}^t \mathbf{X} (\mathbf{U}_{\lambda 1}(\mathbf{y}) - \boldsymbol{\beta}) = \frac{\mathcal{V}^2}{n \lambda \mathcal{W}} \quad (17)$$

$$\delta(\mathbf{U}_\lambda(\mathbf{y}), \theta) = \left(\frac{\frac{\mathcal{V}^2}{n \lambda \mathcal{W}} + 2 \frac{(\lambda \mathcal{W} - 1)^2}{\lambda \mathcal{W}}}{\frac{\mathcal{V}^2}{n \lambda \mathcal{W}} + 2 \frac{(\lambda \mathcal{W} + 1)^2}{\lambda \mathcal{W}}} \right)^{1/2} \quad (18)$$

and

$$\rho(\mathbf{U}_\lambda(\mathbf{y}), \theta) = 2\sqrt{2} \operatorname{arctanh} \left(\left(\frac{\mathcal{V}^2 + 2n(\lambda \mathcal{W} - 1)^2}{\mathcal{V}^2 + 2n(\lambda \mathcal{W} + 1)^2} \right)^{1/2} \right) \quad (19)$$

or

$$\rho(\mathbf{U}_\lambda(\mathbf{y}), \theta) = \sqrt{2} \operatorname{arccosh} \left(\frac{1}{4} \frac{\mathcal{V}^2}{n \lambda \mathcal{W}} + \frac{1}{2} \left(\lambda \mathcal{W} + \frac{1}{\lambda \mathcal{W}} \right) \right) \quad (20)$$

which have a distribution which depend only on $\mathcal{V}^2$ and $\mathcal{W}^2$, independent random variables with fixed distribution, whatever the value of θ .

Since the risk of any equivariant estimator remains constant on the parameter space, it's enough to examine it at one point, for instance at the point $(\mathbf{0}, 1) \in \Theta$. Let us denote the expectation with respect to the n -variate linear normal model $N_n(\mathbf{X}\boldsymbol{\beta}, \sigma^2 \mathbf{I}_n)$ by $E_{(\boldsymbol{\beta}, \sigma)}$ and by E the $E_{(\mathbf{0}, 1)}$. We can prove the following propositions.

Proposition 4. If $n \geq m + 1$ we have:

$$E_{(\boldsymbol{\beta}, \sigma)} \left(\rho^2(\mathbf{U}_\lambda(\mathbf{y}), (\boldsymbol{\beta}, \sigma)) \right) = E \left(\rho^2(\mathbf{U}_\lambda(\mathbf{y}), (\mathbf{0}, 1)) \right) < \infty, \quad \forall \lambda > 0$$

for any $(\boldsymbol{\beta}, \sigma) \in \Theta$.

Proof: From (14) and (13), since

$$1 + \frac{1}{2!} \frac{\rho^2}{2} + \frac{1}{4!} \frac{\rho^4}{4} \leq \cosh\left(\frac{\rho}{2}\right) = \frac{1}{4} d_M^2(\boldsymbol{\theta}_1, \boldsymbol{\theta}_2) + \frac{\sigma_1^2 + \sigma_2^2}{2 \sigma_1 \sigma_2} \quad (21)$$

we have

$$0 \leq \rho^2(\theta_1, \theta_2) \leq 12 \left(\sqrt{1 + \frac{d_M^2(\theta_1, \theta_2)}{6} + \frac{1}{3} \frac{(\sigma_1 - \sigma_2)^2}{\sigma_1 \sigma_2}} - 1 \right) \quad (22)$$

developing the square of the difference and taking into account that the standard Euclidean norm of a vector is less or equal to the absolute value of the sum of its components, we obtain

$$0 \leq \rho^2(\theta_1, \theta_2) \leq 2\sqrt{6} d_M(\theta_1, \theta_2) + 4\sqrt{3} \left(\sqrt{\frac{\sigma_1}{\sigma_2}} + \sqrt{\frac{\sigma_2}{\sigma_1}} \right) + 4\sqrt{3} - 12 \quad (23)$$

Notice that both bounds (22) and (23) are invariant under the action of the induced group on the parameter space.

As we mentioned before, from [14], it is enough to prove that the risk is finite at $(0, 1) \in \Theta$. Taking into account (16) it follows, from (23), that

$$0 \leq \rho^2(\mathbf{U}_\lambda(\mathbf{y}), (0, 1)) \leq 2\sqrt{6} \sqrt{\frac{\mathcal{V}^2}{\lambda \mathcal{W} n}} + 4\sqrt{3} \left(\sqrt{\lambda \mathcal{W}} + \frac{1}{\sqrt{\lambda \mathcal{W}}} \right) + 4\sqrt{3} - 12 \quad (24)$$

Observe that if Q has a chi-square distribution with k degrees of freedom

$$E(Q^\alpha) = 2^\alpha \frac{\Gamma(\frac{k}{2} + \alpha)}{\Gamma(\frac{k}{2})} \quad \text{provided that} \quad \frac{k}{2} + \alpha > 0$$

Therefore, since $\mathcal{V}^2$ and $\mathcal{W}^2$ are independent random variables following a central chi-square distribution with m and $(n - m)$ degrees of freedom, we have

$$E(\mathcal{V}) = \sqrt{2} \frac{\Gamma(\frac{m+1}{2})}{\Gamma(\frac{m}{2})}, \quad E(\sqrt{\mathcal{W}}) = \sqrt[4]{2} \frac{\Gamma(\frac{2(n-m)+1}{4})}{\Gamma(\frac{n-m}{2})} \quad \text{provided that} \quad n - m > -1/2$$

and

$$E\left(\frac{1}{\sqrt{\mathcal{W}}}\right) = \frac{1}{\sqrt[4]{2}} \frac{\Gamma(\frac{2(n-m)-1}{4})}{\Gamma(\frac{n-m}{2})} \quad \text{provided that} \quad n - m > 1/2$$

Then, taking the average, it follows that

$$\begin{aligned} E\left(\rho^2(\mathbf{U}_\lambda(\mathbf{y}), (0, 1))\right) &< \frac{2\sqrt{6}}{\sqrt{\lambda} n} E(\mathcal{V}) E\left(\frac{1}{\sqrt{\mathcal{W}}}\right) + \\ &+ 4\sqrt{3} \left(\sqrt{\lambda} E(\sqrt{\mathcal{W}}) + \frac{1}{\sqrt{\lambda}} E\left(\frac{1}{\sqrt{\mathcal{W}}}\right) \right) + 4\sqrt{3} - 12 \\ &< +\infty, \quad \forall \lambda > 0 \text{ and } n - m > 1/2 \end{aligned}$$

Since n and m are positive integers being $n > m$, we conclude that the risk is finite if $n \geq m + 1$. $\square$

Proposition 4 is a sufficient condition for the existence of the Riemannian risk of the equivariant estimator $\mathbf{U}_\lambda$, thus

$$\Phi(\lambda) = E\left(\rho^2(\mathbf{U}_\lambda(\mathbf{y}), (0, 1))\right), \quad \lambda > 0, \quad (25)$$

is well defined for $n \geq m + 1$, which we shall assume hereafter.

Proposition 5. *There exists a unique minimizer $\lambda_{nm} > 0$ of the Riemannian risk given by Φ .*

Proof:

Let us consider the Riemannian risk at $\lambda = e^{\frac{s}{2\sqrt{2}}}$ as a function of s , that is

$$F(s) = \Phi(e^{\frac{s}{2\sqrt{2}}}), \quad s \in \mathbb{R}$$

The particular selection of λ , from which F follows, relies on the Riemannian structure of Θ induced by the information metric. The Riemannian curvature is constant and equal to $-\frac{1}{2}$ and taking into account (10) we have that

$$\gamma(s) = \left((\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t \mathbf{y}, e^{\frac{s}{\sqrt{2}}} \left\{ \mathbf{y}^t \left(\mathbf{I}_n - \mathbf{X}(\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t \right) \mathbf{y} \right\}^{1/2} \right)$$

is a geodesic in Θ ; precisely a geodesic parameterized by the arc-length, see [13] for further details.

Then, following [15], the real valued function

$$s \mapsto \rho^2(\gamma(s), (\mathbf{0}, 1))$$

is strictly convex. Since almost surely convexity of a stochastic process carries over the mean of a process, the map F is strictly convex as well.

On the other hand, from Fatou's Lemma $F(s) \rightarrow +\infty$ as $s \rightarrow -\infty$ or $s \rightarrow \infty$. This, together with the strict convexity of F yield the existence of a unique minimizer s^* of the function F , which depends on n and m .

Finally, since the map $s \mapsto e^{\frac{s}{2\sqrt{2}}}$ is a strictly monotonous function; must exist a unique λ_{nm} , namely $\lambda_{nm} = e^{\frac{s^*}{2\sqrt{2}}}$, such that $\Phi(\lambda_{nm}) = \min_{\lambda > 0} \Phi(\lambda)$. $\square$

In fact, this result guarantees the unicity of the MIRE, although a numerical analysis is required to obtain it explicitly (see next section). It could be useful to develop a simple approximate estimator, that shall be referred hereafter a -MIRE, obtained, luckily, minimizing a convenient upper bound of $\Phi(\lambda)$. Since

$$1 + \frac{1}{2!} \frac{\rho^2}{2} \leq \cosh\left(\frac{\rho}{2}\right) = \frac{1}{4} d_M^2(\theta_1, \theta_2) + \frac{\sigma_1^2 + \sigma_2^2}{2 \sigma_1 \sigma_2} \quad (26)$$

we shall have

$$0 \leq \rho^2(\mathbf{u}_\lambda(\mathbf{y}), (\mathbf{0}, 1)) \leq \frac{1}{n \lambda} \frac{\mathcal{V}^2}{\mathcal{W}} + 2 \left(\sqrt{\lambda \mathcal{W}} - \frac{1}{\sqrt{\lambda \mathcal{W}}} \right)^2 \quad (27)$$

and therefore

$$0 \leq \Phi(\lambda) \leq H(\lambda) = \frac{1}{\lambda} \frac{\Gamma(\frac{n-m}{2} - \frac{1}{2})}{\sqrt{2} \Gamma(\frac{n-m}{2})} \left(\frac{m}{n} + 2 \right) + 2\sqrt{2} \frac{\Gamma(\frac{n-m}{2} + \frac{1}{2})}{\Gamma(\frac{n-m}{2})} \quad (28)$$

the upper bound $H(\lambda)$ it is clearly a convex functions with an absolute minimum attained when λ satisfy

$$\widetilde{\lambda}_{nm} = \left(\sqrt{1 + \frac{m}{2n}} \right) \frac{1}{\sqrt{n-m-1}} \quad (29)$$

Furthermore, given an arbitrary m , we have

$$\lim_{n \rightarrow \infty} \left(\widetilde{\lambda}_{nm} \right)^2 n = 1 \quad (30)$$

and, therefore, a -MIRE is very close to MLE for large values of n . Observe also that it is possible to compute a -MIRE for $n > m + 1$, a condition which is slightly stronger that the result required for the existence of MIRE in proposition (4).

A further aspect is the intrinsic bias of the equivariant estimators. In fact connections between minimum risk, bias and invariance have been established, see [14]. Since the action of the group G is

not commutative, we cannot guarantee the unbiasedness of the MIRE and an additional analysis must be performed. First of all we are going to compute the vector bias, see [4], a quantitative measure of the bias which is compatible with Lehmann results.

Let $A_\theta(\lambda) = \exp_\theta^{-1}(\mathbf{U}_\lambda(\mathbf{y}))$ and $A_\theta^1(\lambda), \dots, A_\theta^m(\lambda), A_\theta^{m+1}(\lambda)$ be the components of $A_\theta(\lambda)$ corresponding to the basis field $(\partial/\partial\beta^1)_\theta, \dots, (\partial/\partial\beta^m)_\theta, (\partial/\partial\sigma)_\theta$. With matrix notation, $A_{\theta 1}(\lambda) = (A_\theta^1(\lambda), \dots, A_\theta^m(\lambda))^t$. Furthermore, let us define $h(x) \equiv x/\sinh x$ for $x \neq 0$ and $h(0) \equiv 0$; taking into account (16), (19) and from (11) and (15) we have

$$\begin{aligned} \mathbf{X}A_{\theta 1}(\lambda) &= f_\lambda(\mathcal{V}, \mathcal{W}) \mathbf{z} \frac{\sigma}{\lambda \mathcal{W}} \\ A_\theta^{m+1}(\lambda) &= f_\lambda(\mathcal{V}, \mathcal{W}) g_\lambda(\mathcal{V}, \mathcal{W}) \sigma \end{aligned} \quad (31)$$

where

$$f_\lambda(\mathcal{V}, \mathcal{W}) \equiv h(\rho(\mathbf{U}_\lambda(\mathbf{y}), \theta)/\sqrt{2})$$

and

$$g_\lambda(\mathcal{V}, \mathcal{W}) \equiv \cosh(\rho(\mathbf{U}_\lambda(\mathbf{y}), \theta)/\sqrt{2}) - \frac{1}{\lambda \mathcal{W}}$$

Let $\mathbf{B}_\theta(\lambda) = E_\theta(\exp_\theta^{-1}(\mathbf{U}_\lambda(\mathbf{y})))$ be the intrinsic bias vector corresponding to an equivariant estimator $\mathbf{U}_\lambda(\mathbf{y}) = (\mathbf{U}_{\lambda 1}(\mathbf{y}), \mathbf{U}_{\lambda 2}(\mathbf{y}))$ evaluated at the point $\theta = (\beta, \sigma)$ and let $B_\theta^1(\lambda), \dots, B_\theta^m(\lambda), B_\theta^{m+1}(\lambda)$ be their components. In matrix notation, $\mathbf{B}_{\theta 1}(\lambda) = (B_\theta^1(\lambda), \dots, B_\theta^m(\lambda))^t$. We have

Proposition 6. *If $n \geq m + 1$, the bias vector is finite and*

$$\begin{aligned} \mathbf{B}_{\theta 1}(\lambda) &= \mathbf{0} \\ B_\theta^{m+1}(\lambda) &= E(f_\lambda(\mathcal{V}, \mathcal{W}) g_\lambda(\mathcal{V}, \mathcal{W})) \sigma \end{aligned} \quad (32)$$

where $\mathcal{V}^2$ and $\mathcal{W}^2$ are independent random variables following a chi-square distribution with m and $n - m$ degrees of freedom respectively.

Moreover, the square of the norm of the bias vector is constant and given by

$$\|\mathbf{B}_\theta(\lambda)\|_\theta^2 = 2 (E(f_\lambda(\mathcal{V}, \mathcal{W}) g_\lambda(\mathcal{V}, \mathcal{W})))^2 \quad (33)$$

Proof: Observe that if $n \geq m + 1$ we have

$$\begin{aligned} \|\mathbf{B}_\theta(\lambda)\|_\theta^2 &\leq E_\theta(\|\exp_\theta^{-1}(\mathbf{U}_\lambda(\mathbf{y}))\|_\theta^2) \leq E_\theta(\|\exp_\theta^{-1}(\mathbf{U}_\lambda(\mathbf{y}))\|_\theta^2) \\ &= E_\theta(\rho^2(\mathbf{U}_\lambda(\mathbf{y}), \theta)) < \infty \end{aligned}$$

where $\|\cdot\|_\theta$ denotes the Riemannian norm at the tangent space at θ .

On the other hand taking into account (31) and defining $\mathbf{z}$ as in (16) observe that $\mathcal{V} = \|\mathbf{z}\| = \|\mathbf{z}\|$, $\mathbf{z}$ is independent of $\mathcal{W}$ and $\mathbf{z}$ and $-\mathbf{z}$ has the same distribution. Then we have

$$\begin{aligned} \mathbf{X}B_{\theta 1}(\lambda) &= E_\theta\left(f_\lambda(\mathcal{V}, \mathcal{W}) \mathbf{z} \frac{\sigma}{\lambda \mathcal{W}}\right) \\ &= E_\theta\left(f_\lambda(\mathcal{V}, \mathcal{W}) (-\mathbf{z}) \frac{\sigma}{\lambda \mathcal{W}}\right) = -E_\theta\left(f_\lambda(\mathcal{V}, \mathcal{W}) \mathbf{z} \frac{\sigma}{\lambda \mathcal{W}}\right) = \mathbf{0} \end{aligned}$$

and since $\text{rank}(\mathbf{X}) = m$ it follows that $\mathbf{B}_{\theta 1}(\lambda) = \mathbf{0}$.

$B_\theta^{m+1}(\lambda)$ is obtained directly from (31). The distribution of $\mathbf{z}$ and $\mathcal{W}^2$ follow from basic properties of multivariate normal distribution. Finally, the norm of the bias vector field follows from (32) and (8). $\square$

We may remark, finally, that the norm of the bias vector field of any equivariant estimator, $\|\mathbf{B}_\theta(\lambda)\|_\theta$, is invariant under the action of the induced group, $\mathcal{G}$, on the parameter space and since this group acts transitively on Θ , this quantity must be constant, which is clear from (33).

4. Numerical Evaluation

In this section we are going to compare, numerically, the MIRE estimator with the standard MLE. Observe that both estimators only differ in the estimation of the parameter σ (or σ^2). Precisely, the MIRE and the MLE of σ^2 are, respectively,

$$\hat{\sigma}^2 = \lambda_{nm}^2 \mathbf{y}^t (\mathbf{I}_n - \mathbf{X}(\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t) \mathbf{y} \quad \hat{\sigma}^{*2} = \frac{1}{n} \mathbf{y}^t (\mathbf{I}_n - \mathbf{X}(\mathbf{X}^t \mathbf{X})^{-1} \mathbf{X}^t) \mathbf{y} \quad (34)$$

namely, they only differ by the factor $\xi_{nm} = \lambda_{nm}^2 n$, since $\hat{\sigma}^2 = \xi_{nm} \hat{\sigma}^{*2}$.

In order to compare MLE with the MIRE we have computed the factor ξ_{nm} , the intrinsic risk and the square of the norm of the bias vector for each estimators. All the computations have been performed using Mathematica 10.2.

Moreover, we have suggested a rather simple approximation for ξ_{nm} , or λ_{nm} , which allow us to approximate MIRE estimator through (29), i.e.

$$\widetilde{\xi_{nm}} = \left(1 + \frac{m}{2n}\right) \frac{n}{n - m - 1} \quad (35)$$

The corresponding estimator, that shall be referred hereafter as *a*-MIRE, has been also compared with MLE and MIRE in terms of intrinsic risk.

The results are summarized in the following figures which summarize the tables given in the Appendix, see [16] for a convincing argument for the use of plots over tables. In Figure 2 numerical results of ξ_{nm} (left) and $\Phi(\lambda_{nm}) = \Phi(\sqrt{\xi_{nm}/n})$ (right) are displayed graphically. Observe that, for m fixed, as n increases $\Phi(\sqrt{\xi_{nm}/n})$ goes to zero and ξ_{nm} goes to one. The exact numerical values are given in the Appendix, Table 1.

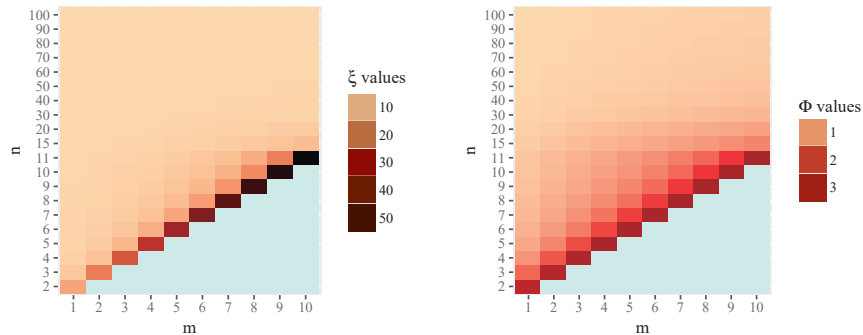


Figure 1. Numerical results for ξ_{nm} (left) and $\Phi(\lambda_{nm})$ (right), with $\xi_{nm} = \lambda_{nm}^2 n$.

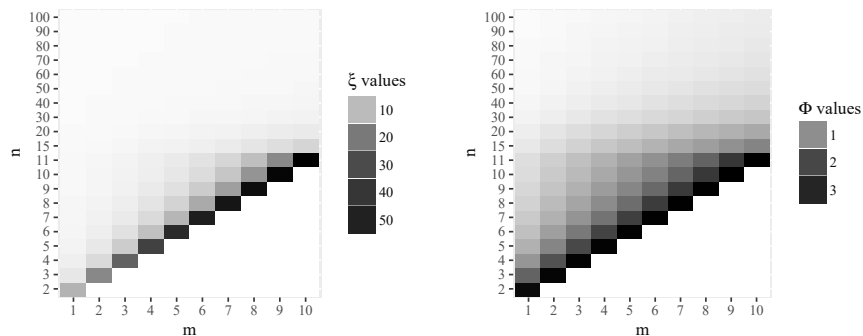


Figure 2. Numerical results for ξ_{nm} (left) and $\Phi(\lambda_{nm})$ (right), with $\xi_{nm} = \lambda_{nm}^2 n$.

Figure 4 shows graphically the percentage of intrinsic risk increment for MLE (left) and a -MIRE (right), that is

$$100 \frac{\left(\Phi\left(\frac{1}{\sqrt{n}}\right) - \Phi(\lambda_{nm})\right)}{\Phi(\lambda_{nm})} \quad \text{and} \quad 100 \frac{\left(\Phi(\widetilde{\lambda}_{nm}) - \Phi(\lambda_{nm})\right)}{\Phi(\lambda_{nm})} \quad (36)$$

respectively, for some values of n and m . The exact numerical results are given in the Appendix, Table 2.

Observe that if we approximate the MIRE by the MLE for a certain value of m , the relative difference of risk decreases as n increases. Notice that these relative differences of risks are rather moderate or small (about 10-15%) only if $n > 10m$. Let us remark also the non appropriate behavior of the MLE for small values of $n - m$. This regards to the intrinsic risk increment: when we use MLE instead of MIRE this increment oscillates between 80,5% and 251,1% when $n - m = 1$ with m from 1 to 10. On the other hand, the behavior of the a -MIRE is reasonably good, with its intrinsic risk very similar to the risk of the MIRE estimator: here the percentage on risk increment is less than 1% for all studied cases and also this percentage decreases as n increases. In fact this percentage is lower than 1% when $n - m \geq 4$, which indicates the extraordinary degree of approximation, being therefore the a -MIRE a reasonable and useful approximation of MIRE. In particular as an example, for a two way analysis of variance, with a and b levels for factors A and B respectively, with a single replicate for treatment we shall have

$$\tilde{\lambda} = \left(\sqrt{1 + \frac{a+b-1}{2ab}} \right) \frac{1}{\sqrt{ab-a-b}} \quad (37)$$

with $a > 1$ and $b > 1$, while the corresponding quantity for MLE is $\lambda^* = \frac{1}{\sqrt{ab}}$.

In particular if $a = 3$ and $b = 5$ we shall have

$$\tilde{\lambda} = \left(\sqrt{1 + \frac{7}{30}} \right) \frac{1}{\sqrt{7}} = 0,419750$$

and $\lambda^* = \frac{1}{\sqrt{15}} = 0,258199$, which is sensibly different. At this moment, it could be useful to recall that the quadratic loss and the squared of the Riemannian distance behaves very different, see [5].

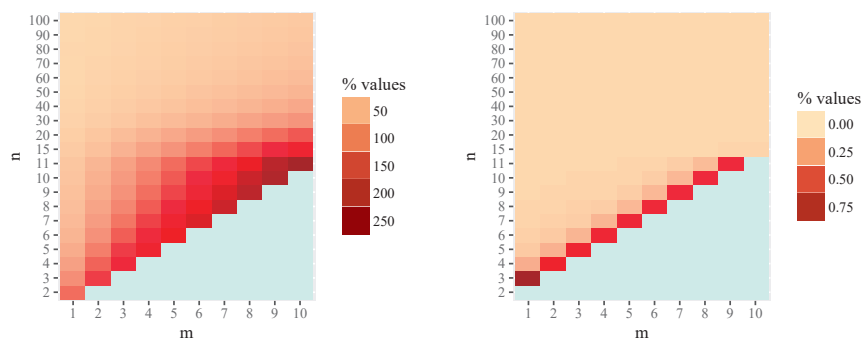


Figure 3. Percentage of intrinsic risk increment for MLE (left) and a -MIRE (right).

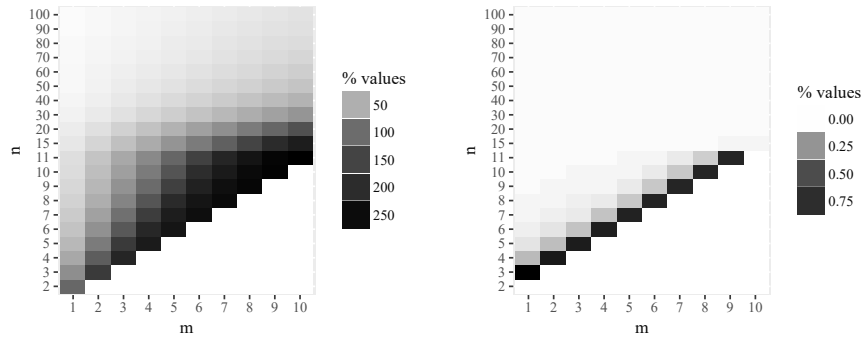


Figure 4. Percentage of intrinsic risk increment for MLE (left) and a -MIRE (right).

Figure 6 displays graphically the numerical results of the $m + 1$ component of the bias vector divided by σ , i.e. the unique non zero physical component of this vector field, for MIRE (left) and MLE (right), $\frac{1}{\sigma} B_{\theta}^{m+1}(\lambda_{nm})$ and $\frac{1}{\sqrt{n}\sigma} B_{\theta}^{m+1}(\frac{1}{\sqrt{n}})$ respectively, for some values of n and m . Observe the sign differences, meaning that MIRE overestimates, in average, σ while MLE underestimates this quantity, in average as well. This follows from the equation of the geodesics of the present model (9). The exact numerical values are given in the Appendix, Table 3.

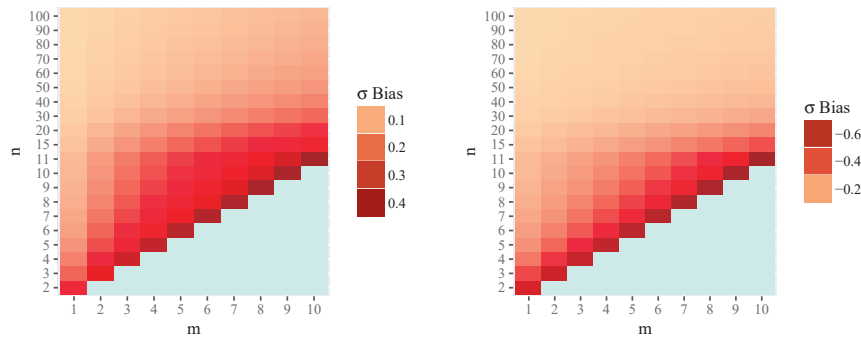


Figure 5. Numerical results for the unique non-null physical component of bias vector field $\frac{1}{\sigma} B_{\theta}^{m+1}(\lambda)$, for MIRE (left, with $\lambda = \lambda_{nm}$) and MLE (right, with $\lambda = 1/\sqrt{n}$).

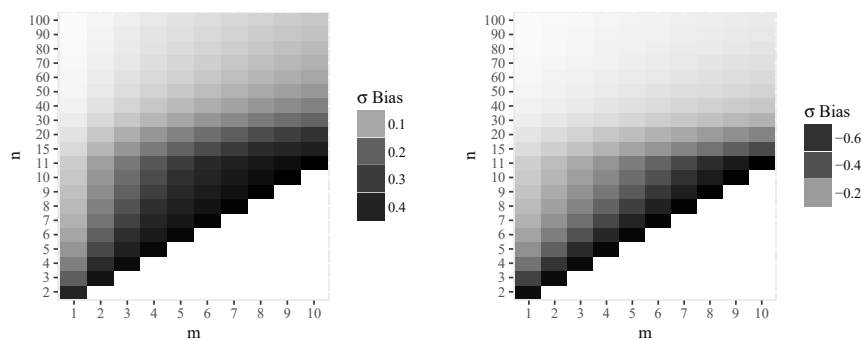


Figure 6. Numerical results for the unique non-null physical component of bias vector field $\frac{1}{\sigma} B_{\theta}^{m+1}(\lambda)$, for MIRE (left, with $\lambda = \lambda_{nm}$, positive values) and MLE (right, with $\lambda = 1/\sqrt{n}$, negative values).

Figure 8 shows graphically the numerical results of the percentage of intrinsic risk due to bias for MIRE and MLE, that is

$$100 \frac{\|B\|_{\theta}^2(\lambda_{nm})}{\Phi(\lambda_{nm})} \quad \text{and} \quad 100 \frac{\|B\|_{\theta}^2(\frac{1}{\sqrt{n}})}{\Phi(\lambda_{nm})} \quad (38)$$

respectively, for some values of n and m . Observe that the bias is moderate, with respect the intrinsic risk, in both estimators. The bias of MIRE estimator is smaller than the bias of MLE for small values of m and the opposite for large values. The exact numerical results are given in the Appendix, Table 4.

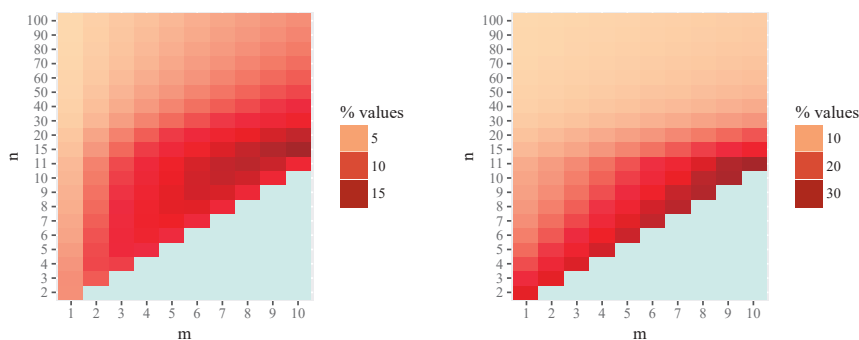


Figure 7. Percentage of intrinsic risk due to bias for MIRE (left) and MLE (right), referred to the risk of MIRE.

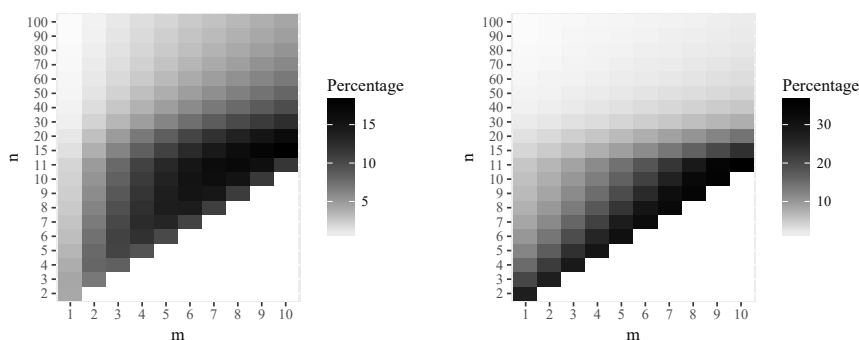


Figure 8. Percentage of intrinsic risk due to bias for MIRE (left) and MLE (right), referred to the risk of MIRE.

Acknowledgements: We have to thank the referees and the editor comments and suggestions for the improvement of this paper.

5. Appendix

[illegible]

Table 2: Percentage of intrinsic risk increment for MLE (above) and aMBRE (below).

n	1	2	3	4	5	6	7	8	9	10
2	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01
3	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01
4	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01
5	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01
6	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01
7	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01
8	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01
9	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01
10	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01

Table 3: Numerical results for the unique non-null physical component of base vector field $\hat{f}^{(n)}(x)$, for aMBRE (above, with $\lambda \rightarrow \lambda_{n-1}$) and MLE (below, with $\lambda \rightarrow 1/\sqrt{n}$).

n	1	2	3	4	5	6	7	8	9	10
2	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01
3	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01
4	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01
5	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01
6	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01
7	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01
8	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01
9	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01
10	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01

Table 4: Percentage of intrinsic risk due to bias for aMBRE (above) and MLE (below), referred to the risk of aMBRE.

n	1	2	3	4	5	6	7	8	9	10
2	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01
3	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01
4	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01
5	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01
6	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01
7	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01
8	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01
9	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01
10	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01	0.01

References

1. Rao, C. Information and accuracy attainable in the estimation of statistical parameters. *Bull. Calcutta Math. Soc.* **1945**, *37*, 81–91.
2. Burbea, J.; Rao, C. Entropy differential metric, distance and divergence measures in probability spaces: a unified approach. *J. Multivar. Anal.* **1982**, *12*, 575–596.
3. Burbea, J. Informative geometry of probability spaces. *Expo. Math.* **1986**, *4*, 347–378.
4. Oller, J.; Corcuera, J. Intrinsic Analysis of the Statistical Estimation. *Ann. Stat.* **1995**, *23*, 1562–1581.
5. García, G.; Oller, J. What does intrinsic mean in statistical estimation? *Sort* **2006**, *2*, 125–146.
6. Bernal-Casas, D.; Oller, J. Variational Information Principles to Unveil Physical Laws. *Mathematics* **2024**, *12*, 3941. <https://doi.org/10.3390/math12243941>.
7. Muirhead, R. *Aspects of Multivariate Statistical Theory*; John Wiley & Sons Inc.: New York, NY, USA, 1982.
8. Bernardo, J.; Juárez, M. Intrinsic estimation. In *Bayesian Statistics 7*; Bernardo, J.; Bayarri, M.; Berger, J.; Dawid, A.; Hackerman, W.; Smith, A.; West, M., Eds.; Oxford University Press: Berlin, Germany, 2003; pp. 465–476.

9. Lehmann, E.; Casella, G. *Theory of Point Estimation*; Springer Science & Business Media: New York, NY, USA, 2006.
10. ichi Amari, S. *Information Geometry and Its Applications*; Springer: Tokyo, 2016.
11. Ay, N.; Jost, J.; Lê, H.V.; Schwachhöfer, L. *Information Geometry*; Springer: Cham, 2017.
12. Nielsen, F. An Elementary Introduction to Information Geometry. *Entropy* **2020**, *22*, 1100. <https://doi.org/10.3390/e22101100>.
13. Burbea, J.; Oller, J. The information metric for univariate linear elliptic models. *Stat. Decis.* **1988**, *6*, 209–221.
14. Lehmann, E. A general concept of unbiasedness. *Ann. Math. Stat.* **1951**, *22*, 587–592.
15. Karcher, H. Riemannian Center of Mass and Mollifier Smoothing. *Commun. Pure Appl. Math.* **1977**, *30*, 509–541.
16. Gelman, A.; Pasarica, C.; Dodhia, R. Let's practice what we preach: turning tables into graphs. *Am. Stat.* **2002**, *56*, 121–130.

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.