

CONNECTING MODEL-BASED AND MODEL-FREE APPROACHES TO LINEAR LEAST SQUARES REGRESSION

Lutz Dümbgen¹

Department of Mathematics and Statistics, University of Bern, Bern, Switzerland

Laurie Davies

University of Duisburg-Essen, Essen, Germany

SUMMARY

In a regression setting with a response vector and given regressor vectors, a typical question is to what extent the response is related to these regressors, specifically, how well it can be approximated by a linear combination of the latter. Classical methods for this question are based on statistical models for the conditional distribution of the response, given the regressors. In the present paper it is shown that various p-values resulting from this model-based approach have also a purely data-analytic, model-free interpretation. This finding is derived in a rather general context. In addition, we introduce equivalence regions, a reinterpretation of confidence regions in the model-free context.

Keywords: Gaussian covariate; Haar measure; Projection.

1. INTRODUCTION

Statistical inference with general linear models is a well-established and indispensable tool for data analysis. The standard output of statistical software for linear models includes least squares estimators of parameters and their standard errors as well as p-values for various linear hypotheses. While the latter are based on certain model assumptions, linear models can also be viewed as tools for purely exploratory data analysis. In such a model-free context, one might wonder whether the p-values for, say, the relevance of certain covariates are still meaningful. The surprising answer is yes, these p-values do have a very precise and new interpretation.

To formulate a first result, suppose we observe a response vector $y \in \mathbb{R}^n$ and $p < n$ linearly independent regressor vectors (regressors) $x_1, \dots, x_p \in \mathbb{R}^n$. For a given integer

¹ Corresponding Author. E-mail: lutz.duembgen@unibe.ch

$p_o \in \{0, \dots, p-1\}$, we would like to know whether the least squares approximation of $\mathbf{y}$ by a linear combination $\hat{\mathbf{y}}$ of $\mathbf{x}_1, \dots, \mathbf{x}_p$, that is,

$$\|\mathbf{y} - \hat{\mathbf{y}}\|^2 = \min_{\beta \in \mathbb{R}^p} \left\| \mathbf{y} - \sum_{j=1}^p \beta_j \mathbf{x}_j \right\|^2$$

with the standard Euclidean norm $\|\cdot\|$, is substantially better than the restricted least squares approximation of $\mathbf{y}$ by a linear combination $\hat{\mathbf{y}}_o$ of $\mathbf{x}_1, \dots, \mathbf{x}_{p_o}$ only, where $\hat{\mathbf{y}}_o := \mathbf{0}$ in case of $p_o = 0$. A classical, model-based answer is to assume that the $\mathbf{x}_j$ are fixed while

$$\mathbf{y} \sim N_n \left(\sum_{j=1}^p \beta_j \mathbf{x}_j, \sigma^2 \mathbf{I} \right), \quad (1)$$

with an unknown parameter vector $\beta \in \mathbb{R}^p$ and an unknown standard deviation $\sigma > 0$. Then an exact p-value of the null hypothesis that

$$\beta_j = 0 \text{ for } p_o < j \leq p$$

is given by

$$1 - F_{p-p_o, n-p} \left(\frac{\|\hat{\mathbf{y}} - \hat{\mathbf{y}}_o\|^2 / (p - p_o)}{\|\mathbf{y} - \hat{\mathbf{y}}\|^2 / (n - p)} \right), \quad (2)$$

where $F_{k,\ell}$ denotes the distribution function of Fisher's F distribution with k and ℓ degrees of freedom. Since $\|\mathbf{y} - \hat{\mathbf{y}}_o\|^2 = \|\mathbf{y} - \hat{\mathbf{y}}\|^2 + \|\hat{\mathbf{y}} - \hat{\mathbf{y}}_o\|^2$, one can deduce from well-known connections between chi-squared, gamma, F and beta distributions that the p-value (2) may be rewritten as

$$B_{(n-p)/2, (p-p_o)/2} \left(\frac{\|\mathbf{y} - \hat{\mathbf{y}}\|^2}{\|\mathbf{y} - \hat{\mathbf{y}}_o\|^2} \right), \quad (3)$$

where $B_{a,b}$ denotes the distribution function of the beta distribution $\text{Beta}(a, b)$ with parameters $a, b > 0$.

Now let us view all observation vectors $\mathbf{y}$ and $\mathbf{x}_1, \dots, \mathbf{x}_p$ as fixed. To measure to what extent $(\mathbf{x}_j)_{p_o < j \leq p}$ contributes substantially to the least squares fit $\hat{\mathbf{y}}$, let $\hat{\mathbf{y}}^*$ be the least squares fit of $\mathbf{y}$ after replacing $(\mathbf{x}_j)_{p_o < j \leq p}$ with a tuple $(\mathbf{x}_j^*)_{p_o < j \leq p}$ of independent random vectors $\mathbf{x}_j^* \sim N_n(\mathbf{0}, \mathbf{I})$. A precise measure of the relevance of $(\mathbf{x}_j)_{p_o < j \leq p}$ is given by the probability that $\|\mathbf{y} - \hat{\mathbf{y}}^*\|^2$ is smaller than $\|\mathbf{y} - \hat{\mathbf{y}}\|^2$. The smaller this probability, the higher is the relevance. Interestingly, it can be computed exactly and coincides with the p-value in Eq. (3).

LEMMA 1. *For arbitrary fixed, linearly independent vectors $\mathbf{y}, \mathbf{x}_1, \dots, \mathbf{x}_p \in \mathbb{R}^n$ and stochastically independent standard Gaussian random vectors $\mathbf{x}_j^* \in \mathbb{R}^n$, $p_o < j \leq p$,*

$$\frac{\|\mathbf{y} - \hat{\mathbf{y}}^*\|^2}{\|\mathbf{y} - \hat{\mathbf{y}}_o\|^2} \sim \text{Beta}((n-p)/2, (p-p_o)/2).$$

In particular,

$$\mathbb{P}(\|\mathbf{y} - \hat{\mathbf{y}}^*\|^2 \leq \|\mathbf{y} - \hat{\mathbf{y}}\|^2) = B_{(n-p)/2, (p-p_o)/2} \left(\frac{\|\mathbf{y} - \hat{\mathbf{y}}\|^2}{\|\mathbf{y} - \hat{\mathbf{y}}_o\|^2} \right).$$

This Lemma is essentially a variant of a classical result about the angle between a random linear subspace and a fixed vector, see for instance Theorem 1.1 of [Frankl and Maehara \(1990\)](#). A direct and self-contained proof will be given at the beginning of Section A. But Lemma 1 can be viewed as a special case of a more general connection between the model-based and model-free point of view which is elaborated in Section 2. In particular, the classical model-based p-values do not require a Gaussian distribution of $\mathbf{y}$, given $(\mathbf{x}_j)_{1 \leq j \leq p}$, and for the model-free interpretation, the random tuple $(\mathbf{x}_j^*)_{p_o < j \leq p}$ may have different distributions all of which lead to the p-value in Eq. (3). In Section 3 we discuss “equivalence” regions. In the model-based context, these are confidence regions for the unknown mean vector $\boldsymbol{\mu} = \mathbb{E}(\mathbf{y})$. Under the model-free point of view, the interpretation of these regions is somewhat different. To illustrate the concept, we describe relatively simple equivalence regions for a sparse signal vector.

Some final comments and an outlook to future work are given in Section 4. In particular, we explain how the considerations and results in the present paper are related to previous work about permutation tests in regression settings, a key reference being [Freedman and Lane \(1983\)](#) and a review of [Winkler et al. \(2014\)](#).

Technical details and proofs are deferred to Section A. Throughout this paper we use standard results from multivariate statistics and linear models as presented in standard textbooks, e.g. [Mardia et al. \(1979\)](#), [Eaton \(1983\)](#) and [Scheffé \(1959\)](#), without further reference.

2. THE F TEST AND OTHER METHODS REVISITED

We consider arbitrary vectors $\mathbf{y}$ and $\mathbf{x}_1, \dots, \mathbf{x}_p$ in $\mathbb{R}^n$. At first we discuss the question whether there is any association between $\mathbf{y}$ and $(\mathbf{x}_j)_{1 \leq j \leq p}$. In the introduction, this corresponds to $p_o = 0$. In Section 2.3 we return to situations in which the contribution of $p_o \in \{1, \dots, p-1\}$ regressors $\mathbf{x}_1, \dots, \mathbf{x}_{p_o}$ is not questioned. This includes linear regression models with an intercept, accommodated by the trivial regressor $\mathbf{x}_1 = (1)_{i=1}^n$.

Concerning the regressors $\mathbf{x}_j$, suppose the raw data are given by a data matrix with n rows

$$[\mathbf{y}_i, \mathbf{w}_i^\top] = [\mathbf{y}_i, w_{i,1}, \dots, w_{i,d}], \quad 1 \leq i \leq n,$$

containing the values of a response and d covariates for each observation. If the covariates are numerical or 0-1-valued, the usual multiple linear regression model would consider the regressors $\mathbf{x}_1 := (1)_{i=1}^n$ and $\mathbf{x}_j := (w_{i,j-1})_{i=1}^n$, $2 \leq j \leq d+1$. More complex models would also include the $\binom{d}{2}$ interaction vectors $\mathbf{x}_{j(a,b)} := (w_{i,a} w_{i,b})_{i=1}^n$, $1 \leq a < b \leq d$. In general, with arbitrary types of covariates, one could think of $\mathbf{x}_j = (f_j(\mathbf{w}_i))_{i=1}^n$ with given basis functions $f_1, \dots, f_p$.

Let us introduce some notation. The unit sphere of $\mathbb{R}^n$ is denoted by $\mathbb{S}_n$, and $\mathbb{O}_n$ stands for the set of orthogonal matrices in $\mathbb{R}^{n \times n}$.

2.1. The model-based approach

We consider the regressors $x_1, \dots, x_p$ as fixed and y as a random vector. In settings with random regressors, the subsequent considerations concern the conditional distribution of y , given $x_1, \dots, x_p$. For simplicity we assume throughout that the distribution of y is continuous, i.e. $Pr(y = x) = 0$ for any $x \in \mathbb{R}^n$.

The null hypothesis of no relationship between y and the regressors $x_1, \dots, x_p$ can be specified by describing a distribution of y which does not depend on the latter vectors (or any other fixed regressors):

H_0 : The random vector y has a spherically symmetric distribution on $\mathbb{R}^n$.

That means, its length $\|y\|$ and direction $\|y\|^{-1}y$ are stochastically independent, where $\|y\|^{-1}y \sim \text{Unif}(\mathbb{S}_n)$, the uniform distribution on the unit sphere $\mathbb{S}_n$. It is well-known that this hypothesis H_0 encompasses the classical assumption that $y \sim N_n(0, \sigma^2 I)$ for some unknown $\sigma > 0$.

In order to derive p-values for H_0 , let $S(y) = S(y, x_1, \dots, x_p)$ be a test statistic such that high values indicate a potential violation of H_0 . Then, a p-value for H_0 is given by

$$\pi(y) := \mathbb{P}(S(\|y\|u) \geq S(y) \mid y), \quad (4)$$

where $u \sim \text{Unif}(\mathbb{S}_n)$ is independent from y . If $S(y)$ is scale-invariant in the sense that

$$S(cv) = S(v) \quad \text{for all } v \in \mathbb{R}^n \setminus \{0\} \text{ and } c > 0, \quad (5)$$

one can write

$$\pi(y) = 1 - F(S(y)), \quad (6)$$

with the distribution function F of $S(u)$,

$$F(x) := \mathbb{P}(S(u) \leq x).$$

Here one could also consider a random vector $z \sim N_n(0, I)$ instead of u .

EXAMPLE 2 (F TEST). If $x_1, \dots, x_p$ are linearly independent with $p < n$, and if $S(y)$ equals the F test statistic

$$S(y) := \frac{\|\hat{y}\|^2/p}{\|y - \hat{y}\|^2/(n-p)}, \quad (7)$$

then scale-invariance of the latter implies that the p-value $\pi(y)$ is given by the simplified formula in Eq. (6). Moreover, the distribution function F in Eq. (6) equals $F_{p, n-p}$, so $\pi(y)$ coincides with Eq. (2) in the special case of $p_0 = 0$. This follows from a standard

argument for linear models: let $\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_n$ be an orthonormal basis of $\mathbb{R}^n$ such that $\text{span}(\mathbf{x}_1, \dots, \mathbf{x}_p) = \text{span}(\mathbf{b}_1, \dots, \mathbf{b}_p)$. Then $\mathbf{z} \sim N_n(\mathbf{0}, \mathbf{I})$ has the same distribution as $\tilde{\mathbf{z}} = \sum_{i=1}^n z_i \mathbf{b}_i$, and

$$S(\tilde{\mathbf{z}}) = \frac{\sum_{j=1}^p z_j^2 / p}{\sum_{i=p+1}^n z_i^2 / (n-p)}$$

has distribution function $F_{p, n-p}$ by definition of Fisher's F distributions.

EXAMPLE 3 (MULTIPLE T TESTS). Suppose that the linear span $\mathbb{V}$ of the regressors $\mathbf{x}_1, \dots, \mathbf{x}_p$ satisfies $q := \dim(\mathbb{V}) < n$. Further let $\mathbb{A}$ be a subset of $\mathbb{V} \cap \mathbb{S}_n$. With the orthogonal projection $\hat{\mathbf{y}}$ of $\mathbf{y}$ onto $\mathbb{V}$, a possible test statistic is given by

$$S(\mathbf{y}) := \hat{\sigma}^{-1} \sup_{\mathbf{a} \in \mathbb{A}} |\mathbf{a}^\top \mathbf{y}| \quad \text{with} \quad \hat{\sigma} := (n-q)^{-1/2} \|\mathbf{y} - \hat{\mathbf{y}}\|. \quad (8)$$

Note that under H_0 , each term $\hat{\sigma}^{-1} \mathbf{a}^\top \mathbf{y}$ follows student's t distribution with $n-q$ degrees of freedom. This Example of $S(\cdot)$ is motivated by Tukey's studentized maximum modulus or studentized range test statistics (Miller, 1981).

EXAMPLE 4 (MULTIPLE F TESTS). Let p and $\mathbf{x}_1, \dots, \mathbf{x}_p$ be arbitrary, and let Λ be a family of subsets M of $\{1, \dots, p\}$ such that the vectors $\mathbf{x}_j$, $j \in M$, are linearly independent with $\#M < n$. With Π_M denoting the orthogonal projection from $\mathbb{R}^n$ onto $\text{span}(\mathbf{x}_j : j \in M)$, a possible test statistic is given by

$$S(\mathbf{y}) := \max_{M \in \Lambda} \frac{\|\Pi_M \mathbf{y}\|^2 / \#M}{\|\mathbf{y} - \Pi_M \mathbf{y}\|^2 / (n - \#M)}. \quad (9)$$

The idea behind this test statistic is that possibly $\mathbf{y} = \boldsymbol{\mu} + \boldsymbol{\varepsilon}$ with a random vector $\boldsymbol{\varepsilon}$ having spherically symmetric distribution and a fixed vector $\boldsymbol{\mu} \in \mathbb{R}^n$ such that

$$\|\Pi_M \boldsymbol{\mu}\|^2 \gg \|\boldsymbol{\mu} - \Pi_M \boldsymbol{\mu}\|^2$$

for some $M \in \Lambda$.

2.2. The model-free point of view

To elaborate on the connection between model-based and model-free approach, note first that the null hypothesis H_0 is equivalent to an orthogonal invariance property. With $\stackrel{d}{=}$ denoting equality in distribution, the alternative formulation reads as follows:

$$H'_0 : T\mathbf{y} \stackrel{d}{=} \mathbf{y} \text{ for any fixed } T \in \mathbb{O}_n.$$

Another equivalent formulation involves normalized Haar measure Haar_n on $\mathbb{O}_n$. This is the unique distribution of a random matrix $\mathbf{H} \in \mathbb{O}_n$ with left-invariant distribution in the sense that

$$\mathbf{TH} \stackrel{d}{=} \mathbf{H} \quad \text{for any fixed } \mathbf{T} \in \mathbb{O}_n.$$

For a thorough account of Haar measure we refer to [Eaton \(1989\)](#); in [Section A](#) we mention two explicit constructions of $\mathbf{H}$ and resulting properties. For the moment it suffices to know that also

$$\mathbf{H}^\top \stackrel{d}{=} \mathbf{H} \stackrel{d}{=} \mathbf{HT} \quad \text{for any fixed } \mathbf{T} \in \mathbb{O}_n.$$

Moreover, for any fixed unit vector $\mathbf{v} \in \mathbb{S}_n$, the random vector $\mathbf{H}\mathbf{v}$ is uniformly distributed on $\mathbb{S}_n$. Now the null hypothesis H'_0 may be reformulated as follows:

$$H''_0 : \text{If } \mathbf{H} \sim \text{Haar}_n \text{ is independent from } \mathbf{y}, \text{ then } \mathbf{H}\mathbf{y} \stackrel{d}{=} \mathbf{y}.$$

The equivalence of the null hypotheses H_0 , H'_0 and H''_0 is explained in [Section A](#).

From now on suppose that the test statistic $S(\mathbf{y}, \mathbf{x}_1, \dots, \mathbf{x}_p)$ is orthogonally invariant in the sense that

$$S(\mathbf{T}\mathbf{y}, \mathbf{T}\mathbf{x}_1, \dots, \mathbf{T}\mathbf{x}_p) = S(\mathbf{y}, \mathbf{x}_1, \dots, \mathbf{x}_p) \quad \text{for all } \mathbf{y} \in \mathbb{R}^n \text{ and } \mathbf{T} \in \mathbb{O}_n. \quad (10)$$

Since $\mathbf{x} \mapsto \mathbf{T}\mathbf{x}$ preserves inner products, a sufficient condition for orthogonal invariance of the test statistic S is that $S(\mathbf{y}, \mathbf{x}_1, \dots, \mathbf{x}_p)$ depends only on the inner products $\mathbf{y}^\top \mathbf{y}$, $\mathbf{y}^\top \mathbf{x}_j$ and $\mathbf{x}_j^\top \mathbf{x}_k$, $1 \leq j \leq k \leq p$. Then, the p-value in [Eq. \(4\)](#) may be rewritten as follows:

$$\begin{aligned} \pi(\mathbf{y}) &= \mathbb{P}(S(\|\mathbf{y}\| \mathbf{u}, \mathbf{x}_1, \dots, \mathbf{x}_p) \geq S(\mathbf{y}) | \mathbf{y}) \\ &= \mathbb{P}(S(\mathbf{H}\mathbf{y}, \mathbf{x}_1, \dots, \mathbf{x}_p) \geq S(\mathbf{y}) | \mathbf{y}) \\ &= \mathbb{P}(S(\mathbf{y}, \mathbf{H}^\top \mathbf{x}_1, \dots, \mathbf{H}^\top \mathbf{x}_p) \geq S(\mathbf{y}) | \mathbf{y}) \\ &= \mathbb{P}(S(\mathbf{y}, \mathbf{H}\mathbf{x}_1, \dots, \mathbf{H}\mathbf{x}_p) \geq S(\mathbf{y}, \mathbf{x}_1, \dots, \mathbf{x}_p) | \mathbf{y}), \end{aligned}$$

where $\mathbf{H} \sim \text{Haar}_n$ is independent from $\mathbf{y}$. If we adopt the model-free point of view and consider all vectors $\mathbf{y}$ and $\mathbf{x}_1, \dots, \mathbf{x}_p$ as fixed, we may write

$$\pi(\mathbf{y}) = \mathbb{P}(S(\mathbf{y}, \mathbf{H}\mathbf{x}_1, \dots, \mathbf{H}\mathbf{x}_p) \geq S(\mathbf{y}, \mathbf{x}_1, \dots, \mathbf{x}_p)).$$

Thus, $\pi(\mathbf{y})$ measures the strength of the apparent association between $\mathbf{y}$ and the regressor tuple $(\mathbf{x}_1, \dots, \mathbf{x}_p)$, as quantified by the test statistic $S(\mathbf{y}, \mathbf{x}_1, \dots, \mathbf{x}_p)$, by comparing the latter value with $S(\mathbf{y}, \mathbf{H}\mathbf{x}_1, \dots, \mathbf{H}\mathbf{x}_p)$. That means, the regressor tuple $(\mathbf{x}_1, \dots, \mathbf{x}_p)$ undergoes a random orthogonal transformation, and there is certainly no “true association” between $\mathbf{y}$ and $(\mathbf{H}\mathbf{x}_1, \dots, \mathbf{H}\mathbf{x}_p)$. To make the latter point rigorous, note that if

$H, J \sim \text{Haar}_n$ and $u \sim \text{Unif}(\mathbb{S}_n)$ are independent (while y is fixed), then $J^\top H \sim \text{Haar}_n$ too, whence

$$\begin{aligned} S(y, Hx_1, \dots, Hx_p) &\stackrel{d}{=} S(y, J^\top Hx_1, \dots, J^\top Hx_p) \\ &= S(Jy, Hx_1, \dots, Hx_p) \\ &\stackrel{d}{=} S(\|y\|u, Hx_1, \dots, Hx_p). \end{aligned}$$

Finally, recall that the method in the introduction with $p_o = 0$ amounts to replacing the fixed regressors $x_1, \dots, x_p$ with independent random vectors $x_1^*, \dots, x_p^* \sim N_n(0, I)$. But in connection with the F test statistic, this has the same effect as replacing the former with $(Hx_1, \dots, Hx_p)$. Indeed, in case of linearly independent vectors $x_1, \dots, x_p$, the value of $S(y)$ in Eq. (7) depends only on y and the p -dimensional linear space $\text{span}(x_1, \dots, x_p)$. Moreover, the distributions of $\text{span}(Hx_1, \dots, Hx_p)$ and of $\text{span}(x_1^*, \dots, x_p^*)$ coincide, see Section A.

2.3. Composite null models

Quite often, the potential influence of some regressors $x_1, \dots, x_{p_o}$ with $1 \leq p_o < p$ is out of question or not of primary interest, and the main question is whether the regressors $x_{p_o+1}, \dots, x_p$ are really relevant for the approximation of y . One possibility to deal with that is to “residualize” the response y and the regressors $x_{p_o+1}, \dots, x_p$, that is, to project them onto the orthogonal complement of $\text{span}(x_1, \dots, x_{p_o})$. In case of $p_o = 1$ and $x_1 = (1)_{i=1}^n$, this boils down to centering y and $x_2, \dots, x_p$.

More formally, assuming without loss of generality that $x_1, \dots, x_{p_o}$ are linearly independent, let $b_1, b_2, \dots, b_n$ be an orthonormal basis of $\mathbb{R}^n$ such that $b_1, \dots, b_{p_o}$ is a basis of $\text{span}(x_1, \dots, x_{p_o})$. With $B = [b_{p_o+1}, \dots, b_n] \in \mathbb{R}^{n \times (n-p_o)}$, the model Equation (1) implies that

$$B^\top y \sim N_{n-p_o} \left(\sum_{j=p_o+1}^p \beta_j B^\top x_j, \sigma^2 I \right).$$

Generally, the previous model-based and model-free considerations can be applied to $B^\top y \in \mathbb{R}^{n-p_o}$ in place of y .

Applying the model-based or model-free approach with the F test statistic in Eq. (7) to the transformed observations $B^\top y$ and $B^\top x_{p_o+1}, \dots, B^\top x_p$ yields the p-value of Eq. (3) in the introduction, see also the first part of the proof of Lemma 1.

3. EQUIVALENCE REGIONS

3.1. General considerations

Model-based approach. Let $\mathbb{M} \subset \mathbb{R}^n$ be a given set. We assume that $\mathbf{y}$ is a random vector such that

$$\mathbf{y} \stackrel{d}{=} \boldsymbol{\mu} + \boldsymbol{\varepsilon}, \quad (11)$$

with an unknown fixed parameter vector $\boldsymbol{\mu} \in \mathbb{M}$ and a random vector $\boldsymbol{\varepsilon}$ with spherically invariant distribution on $\mathbb{R}^n$, where $\mathbb{P}(\boldsymbol{\varepsilon} = \mathbf{0}) = 0$. Now, let $S(\mathbf{y}) = S(\mathbf{y}, \mathbf{x}_1, \dots, \mathbf{x}_p)$ be a test statistic which is scale-invariant in $\mathbf{y} \neq \mathbf{0}$, and let $S(\mathbf{0}) = 0$. For a given (small) number $\alpha \in (0, 1)$ we define the equivalence region

$$C_\alpha(\mathbf{y}) = C_\alpha(\mathbf{y}, \mathbf{x}_1, \dots, \mathbf{x}_p) := \{\mathbf{m} \in \mathbb{M} : S(\mathbf{y} - \mathbf{m}) \leq \kappa_\alpha\},$$

where κ_α is the $(1 - \alpha)$ -quantile of the distribution of $S(\mathbf{u}) \stackrel{d}{=} S(\mathbf{z})$ with random vectors $\mathbf{u} \sim \text{Unif}(\mathbb{S}_n)$ and $\mathbf{z} \sim N_n(\mathbf{0}, \mathbf{I})$. This defines an $(1 - \alpha)$ confidence region for $\boldsymbol{\mu}$ in the sense that in case of Eq. (11),

$$\mathbb{P}(C_\alpha(\mathbf{y}) \ni \boldsymbol{\mu}) = \mathbb{P}(S(\mathbf{z}) \leq \kappa_\alpha) \geq 1 - \alpha.$$

Model-free interpretation. We consider $\mathbf{y}$ as fixed and assume in addition that $S(\mathbf{y}, \mathbf{x}_1, \dots, \mathbf{x}_p)$ is orthogonally invariant. Then, the equivalence region $C_\alpha(\mathbf{y}) = C_\alpha(\mathbf{y}, \mathbf{x}_1, \dots, \mathbf{x}_p)$ consists of all vectors $\mathbf{m} \in \mathbb{M}$ such that the association between $\mathbf{y} - \mathbf{m}$ and the tuple $(\mathbf{x}_1, \dots, \mathbf{x}_p)$ is not substantially stronger than the association between $\mathbf{y} - \mathbf{m}$ and the randomly rotated tuple $(H\mathbf{x}_1, \dots, H\mathbf{x}_p)$, where $H \sim \text{Haar}_n$. Precisely, the value $S(\mathbf{y} - \mathbf{m}, \mathbf{x}_1, \dots, \mathbf{x}_p)$, our measure of association, is not larger than the $(1 - \alpha)$ -quantile of $S(\mathbf{y} - \mathbf{m}, H\mathbf{x}_1, \dots, H\mathbf{x}_p)$.

Example 2 (continued) Let $\mathbb{M} = \text{span}(\mathbf{x}_1, \dots, \mathbf{x}_p) = X\mathbb{R}^p$ with the matrix $X = [\mathbf{x}_1, \dots, \mathbf{x}_p] \in \mathbb{R}^{n \times p}$. Then, the equivalence region equals

$$C_\alpha(\mathbf{y}) = \{X\boldsymbol{\eta} : \boldsymbol{\eta} \in \mathbb{R}^p, \|\hat{\mathbf{y}} - X\boldsymbol{\eta}\|^2 \leq pF_{p, n-p}^{-1}(1 - \alpha)\hat{\sigma}^2(\mathbf{y})\},$$

where $\hat{\sigma}^2(\mathbf{y}) = \|\mathbf{y} - \hat{\mathbf{y}}\|^2 / (n - p)$. The corresponding set $\tilde{C}_\alpha(\mathbf{y}) = \{\boldsymbol{\eta} \in \mathbb{R}^p : X\boldsymbol{\eta} \in C_\alpha(\mathbf{y})\}$ is Scheffé's well-known confidence ellipsoid for the unknown parameter $\boldsymbol{\beta} \in \mathbb{R}^p$ such that $\boldsymbol{\mu} = X\boldsymbol{\beta}$.

3.2. Inference on a sparse signal

Suppose that $p = n$, and that the vectors $\mathbf{x}_1, \dots, \mathbf{x}_n$ are linearly independent. In that case, $X := [\mathbf{x}_1, \dots, \mathbf{x}_n]$ is nonsingular, and assuming Eq. (11), the least squares estimator of $\boldsymbol{\beta} := X^{-1}\boldsymbol{\mu}$ is given by $\hat{\boldsymbol{\beta}}(\mathbf{y}) := X^{-1}\mathbf{y}$. Writing $X^{-1} = [\mathbf{a}_1, \dots, \mathbf{a}_n]^\top$, the Gauss–Markov

estimator of β_i equals $\hat{\beta}_i(\mathbf{y}) = \mathbf{a}_i^\top \mathbf{y}$. In case of $\mathbf{y} \sim N_n(\boldsymbol{\mu}, \sigma^2 \mathbf{I})$, it has distribution $N(\beta_i, \|\mathbf{a}_i\|^2 \sigma^2) = N(\beta_i, ((\mathbf{X}^\top \mathbf{X})^{-1})_{ii} \sigma^2)$.

Suppose that $\boldsymbol{\mu} = \mathbf{X}\boldsymbol{\beta}$ is sparse in the sense that

$$\|\boldsymbol{\beta}\|_0 := \#\{i \leq n : \beta_i \neq 0\}$$

is relatively small compared with n . Then a possible test statistic for the null hypothesis “ $\boldsymbol{\mu} = \mathbf{0}$ ” is given by

$$S_\ell(\mathbf{y}) = S_\ell(\mathbf{y}, \mathbf{x}_1, \dots, \mathbf{x}_n) := G_1(\mathbf{y})/G_\ell(\mathbf{y})$$

for some integer $\ell \in \{2, \dots, n\}$, where $G_1(\mathbf{y}) \geq G_2(\mathbf{y}) \geq \dots \geq G_n(\mathbf{y})$ are the values $\|\mathbf{a}_1\|^{-1}|\mathbf{a}_1^\top \mathbf{y}|, \|\mathbf{a}_2\|^{-1}|\mathbf{a}_2^\top \mathbf{y}|, \dots, \|\mathbf{a}_n\|^{-1}|\mathbf{a}_n^\top \mathbf{y}|$ in descending order. The idea behind this choice is that $G_\ell(\mathbf{y})$ depends mostly on the noise vector $\boldsymbol{\varepsilon}$ if $\|\boldsymbol{\beta}\|_0 < \ell$, while $G_1(\mathbf{y})$ is mainly driven by $\boldsymbol{\mu}$ in case of $\|\boldsymbol{\beta}\| \gg 0$. Scale-invariance of the test statistic $S_\ell(\mathbf{y}, \mathbf{x}_1, \dots, \mathbf{x}_n)$ in $\mathbf{y}$ is obvious. Since $\mathbf{X}^{-1} = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top$, one may write $\|\mathbf{a}_i\|^{-2} = (\mathbf{X}^\top \mathbf{X})_{ii}$ and $\mathbf{a}_i^\top \mathbf{y} = ((\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y})_i$, whence $S_\ell(\dots)$ is also orthogonally invariant.

Denoting the $(1-\alpha)$ -quantile of $S_\ell(\mathbf{z})$, $\mathbf{z} \sim N_n(\mathbf{0}, \mathbf{I})$, with $\kappa_{\ell, \alpha}$ and setting $\mathbb{M} = \mathbb{R}^n$, we obtain the equivalence region

$$C_{\ell, \alpha}(\mathbf{y}) := \{\mathbf{m} \in \mathbb{R}^n : S_\ell(\mathbf{y} - \mathbf{m}) \leq \kappa_{\ell, \alpha}\}.$$

This region is rather useless per se. But if we restrict our attention to vectors $\mathbf{m} = \mathbf{X}\mathbf{b}$, where $\|\mathbf{b}\|_0$ is bounded by a given number, we end up with the equivalence regions

$$\tilde{C}_{k, \ell, \alpha}(\mathbf{y}) := \{\mathbf{b} \in \mathbb{R}^n : \|\mathbf{b}\|_0 \leq k, S_\ell(\mathbf{y} - \mathbf{X}\mathbf{b}) \leq \kappa_{\ell, \alpha}\}, \quad 1 \leq k < n,$$

which are potentially useful in case of $k < \ell$. In this manuscript we only prove a first result about these equivalence regions.

LEMMA 5. *Let*

$$k_{\ell, \alpha}(\mathbf{y}) := \ell - 1 - \max\{j - i : 1 \leq i \leq j \leq \ell, G_i(\mathbf{y})/G_j(\mathbf{y}) \leq \kappa_{\ell, \alpha}\}.$$

Then, $\tilde{C}_{\alpha, k, \ell}(\mathbf{y}) \neq \emptyset$ if and only if $k \geq k_{\ell, \alpha}(\mathbf{y})$.

In the model-based context, $k_{\ell, \alpha}$ is a lower $(1-\alpha)$ -confidence bound for $\|\boldsymbol{\beta}\|_0$, and $\tilde{C}_{k, \ell, \alpha}$ is a $(1-\alpha)$ -confidence region for $\boldsymbol{\beta}$ under the additional assumption that $\|\boldsymbol{\beta}\|_0 \leq k$.

3.3. Special case: the Gaussian sequence model

Suppose that $\mathbf{x}_1, \dots, \mathbf{x}_n$ is the standard basis of $\mathbb{R}^n$. Then $(G_i(\mathbf{z}))_{i=1}^n$ has the same distribution as $(\tilde{\Phi}^{-1}(U_{(n+1-i)}))_{i=1}^n$, where $U_{(1)} < \dots < U_{(n)}$ are the order statistics of independent random variables $U_1, \dots, U_n \sim \text{Unif}[0, 1]$, and $\tilde{\Phi}$ is the distribution function of $|z_1|$,

that is, $\tilde{\Phi}(r) = 2\Phi(r) - 1$ for $r \geq 0$. Since $U_{(n+1-\ell)} = 1 - \ell/n + O_p(n^{-1/2})$ as $n \rightarrow \infty$, a reasonable choice for ℓ seems to be $\ell \approx (1 - \tilde{\Phi}(1))n = 2\Phi(-1)n \approx 0.32n$. Then $\kappa_{\ell,\alpha}$ is approximately the $(1 - \alpha)$ -quantile of $G_1(\mathbf{z}) = \max_{1 \leq i \leq n} |z_i| = \sqrt{2 \log n} (1 + o_p(1))$.

Specifically, let $n = 100$ and $\ell = 32$. Numerical computations (outlined in Section A) yield $\kappa_{\ell,0.01} = 4.083$.

Now we consider the model-based setting with $\mathbf{y} \sim N_n(\boldsymbol{\mu}, \mathbf{I})$ and two different choices for $\boldsymbol{\mu}$. The proof of Lemma 5 includes the construction of a particular vector $\hat{\boldsymbol{\mu}}(\mathbf{y}) = \hat{\boldsymbol{\mu}}_{\ell,\alpha}(\mathbf{y})$ in $C_{\ell,\alpha}(\mathbf{y})$ such that $\|\hat{\boldsymbol{\mu}}\|_0 = k_{\ell,\alpha}(\mathbf{y})$. Now we investigated the distribution of the latter number, of the set $\{i \leq n : \hat{\mu}_i \neq 0\}$ and of the cosine of the angle between $\boldsymbol{\mu}$ and $\hat{\boldsymbol{\mu}}$.

In the first scenario, we considered the sparse vector $\boldsymbol{\mu} = (10, -6, 3, 0, \dots, 0)^\top$. In 100'000 Monte Carlo simulations we estimated the joint distribution of

$$\text{TP}(\mathbf{y}) := \{i \in \{1, 2, 3\} : \hat{\mu}_i(\mathbf{y}) \neq 0\} \quad \text{and} \quad \text{FP}(\mathbf{y}) := \#\{i > 3 : \hat{\mu}_i(\mathbf{y}) \neq 0\},$$

the set of nonzero components of $\boldsymbol{\mu}$ which were detected correctly (“true positives”) and the number of “false positives”, respectively. Table 1 contains estimated probabilities rounded to four digits. It turned out that $\text{FP}(\mathbf{y}) \geq 2$ with estimated probability less than $10^{-4}/2$, and $\text{FP}(\mathbf{y}) \geq 1$ with estimated probability 0.0058. The first component of $\boldsymbol{\mu}$ was always identified correctly as non-zero, whereas the second and third component stayed sometimes undetected.

In the second scenario, we considered the vector $\boldsymbol{\mu} = ((-1)^{i-1} 2^{-(i-1)/2})_{i=1}^n$. Although $\|\boldsymbol{\mu}\|_0 = n$, the first four to six components of $\boldsymbol{\mu}$ contain the main signal, because $\|\boldsymbol{\mu}\|^{-2} \sum_{i>4} \mu_i^2 = 0.0625$ and $\|\boldsymbol{\mu}\|^{-2} \sum_{i>6} \mu_i^2 < 0.0157$. Figure 1 shows the (estimated) distribution of $k_{\ell,\alpha}(\mathbf{y})$. Recall that the latter is a lower $(1 - \alpha)$ -confidence bound for $\|\boldsymbol{\beta}\|_0$.

Finally, Figure 2 show boxplots of the distribution of

$$\cos \angle(\boldsymbol{\mu}, \mathbf{y}) \quad \text{and} \quad \cos \angle(\boldsymbol{\mu}, \hat{\boldsymbol{\mu}})$$

for both scenarios. These plots illustrate that the sparse estimator $\hat{\boldsymbol{\mu}}$ captures $\boldsymbol{\mu}$ better than the raw data $\mathbf{y}$.

TABLE 1
Joint distribution of $\text{TP}(\mathbf{y})$ and $\text{FP}(\mathbf{y})$ in 1st scenario (in percent).

$\text{TP}(\mathbf{y}) =$	$\{1, 2, 3\}$	$\{1, 2\}$	$\{1, 3\}$	$\{1\}$	Otherwise	Total
$\text{FP}(\mathbf{y}) = 0$	11.13	82.70	0.35	5.24	0.00	99.42
$\text{FP}(\mathbf{y}) = 1$	0.12	0.42	0.00	0.04	0.00	0.58
Total	11.25	83.12	0.35	5.28	0.00	100.00

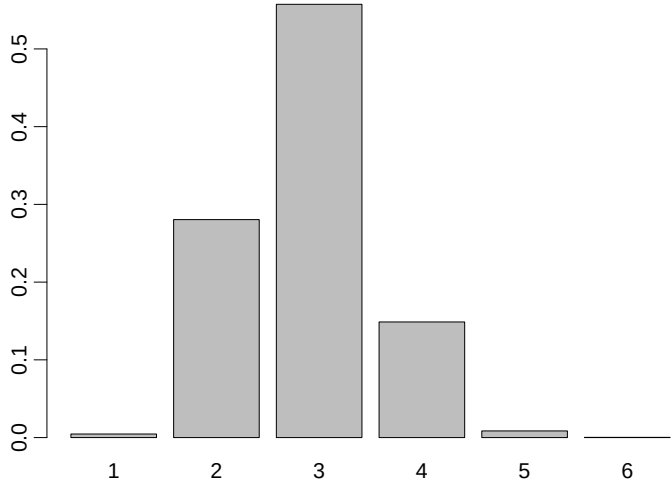


Figure 1 – Distribution of $k_{\ell,\alpha}(y)$ in 2nd scenario.

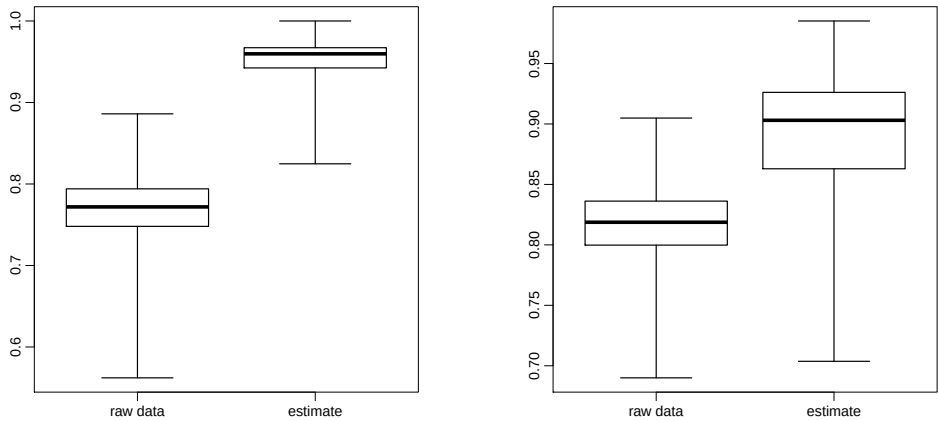


Figure 2 – Distribution of $\cos \angle(\mu, y)$ and $\cos \angle(\mu, \hat{\mu})$ in 1st scenario (left) and 2nd scenario (right).

4. FINAL COMMENTS AND OUTLOOK

Let us relate the considerations in the present paper to research about permutation tests in general regression models. We start from the model-based approach with $y = X\beta + \varepsilon$ with fixed $X = [x_1, \dots, x_p] \in \mathbb{R}^{n \times p}$, $\beta \in \mathbb{R}^p$ and a random error $\varepsilon \in \mathbb{R}^n$. When testing the null hypothesis that $(\beta_j)_{p_0 < j \leq p} = 0$, the goal is to relax the assumption of a spherically symmetric distribution of ε to assuming exchangeability only. That is, for any fixed permutation τ of $\{1, \dots, n\}$, the distribution of $\tau(\varepsilon)$ coincides with the distribution of ε , where $\tau(v) := (v_{\tau(i)})_{i=1}^n$ for $v \in \mathbb{R}^n$. Indeed, if $p_0 = 0$ or $p_0 = 1$ and $x_1 = (1)_{i=1}^n$, the null hypothesis can be tested exactly with a permutation test in which the original data (y, X) are compared with $(\tau(y), X)$ for all (or many randomly chosen) permutations τ . In other cases, however, it is not obvious how to perform a valid permutation test. [Kennedy \(1995\)](#) describes potential pitfalls, and [Winkler et al. \(2014\)](#) provide a good overview of proposed permutation schemes. Apart from the aforementioned simple setting, all these proposals are justified by asymptotic considerations only. A particularly interesting reference is [Freedman and Lane \(1983\)](#), because they motivated their permutation test also by the wish to interpret the resulting p-values in a model-free way. Their approach works as follows: Let $\hat{\varepsilon}_o := y - \hat{y}_o$ be the residual vector under the null model. Then they propose to compare the F test statistic for the original observations (y, X) with the F test statistic applied to $(\hat{y}_o + \tau(\hat{\varepsilon}_o), X)$. The collection of data $(\hat{y}_o + \tau(\hat{\varepsilon}_o), X)$, where τ is an arbitrary permutation of $\{1, \dots, n\}$, serves as a data-generated reference set of data to be compared with the original (y, X) . If the latter sticks out in terms of the F test statistic, this is viewed as model-free evidence that the regressors x_j , $p_0 < j \leq p$, are relevant. [Freedman and Lane \(1983\)](#) show that under certain regularity assumptions, the resulting p-value is asymptotically equivalent to the classical p-value (2). Thus our Lemma 1 may be viewed as a computationally simple and exact alternative to the paradigm of [Freedman and Lane \(1983\)](#).

Sections 3.2 and 3.3 describe a model-free approach to variable selection. Since it is computationally intensive, at least for non-orthogonal regressors $x_1, \dots, x_n$, [Davies and Dümbgen \(2024\)](#) develop an alternative simpler procedure for model-free variable selection. The price to be paid for the simplification is that the selection procedure is possibly conservative in the sense of selecting too few variables. The idea of Gaussian covariate vectors might also be applicable for order selection in autoregressive time series models.

ACKNOWLEDGEMENTS

The authors are grateful to a reviewer and the editor for constructive comments. Part of this work was supported by Swiss National Science Foundation.

APPENDIX

A. TECHNICAL DETAILS AND PROOFS

A.1. Proof of Lemma 1

Note that $\mathbf{y}$ is the sum of the three orthogonal vectors $\hat{\mathbf{y}}_o$, $\hat{\mathbf{y}}^* - \hat{\mathbf{y}}_o$ and $\mathbf{y} - \hat{\mathbf{y}}^*$. Thus,

$$\frac{\|\mathbf{y} - \hat{\mathbf{y}}^*\|^2}{\|\mathbf{y} - \hat{\mathbf{y}}_o\|^2} = \frac{\|\mathbf{y} - \hat{\mathbf{y}}_o\|^2 - \|\mathbf{y} - \hat{\mathbf{y}}^*\|^2}{\|\mathbf{y} - \hat{\mathbf{y}}_o\|^2} = 1 - \frac{\|\hat{\mathbf{y}}^* - \hat{\mathbf{y}}_o\|^2}{\|\mathbf{y} - \hat{\mathbf{y}}_o\|^2},$$

whence the claim is equivalent to

$$\frac{\|\hat{\mathbf{y}}^* - \hat{\mathbf{y}}_o\|^2}{\|\mathbf{y} - \hat{\mathbf{y}}^*\|^2} \sim \text{Beta}((p - p_o)/2, (n - p)/2).$$

Let us first recall two well-known facts about a standard Gaussian random vector $\mathbf{z} = (z_i)_{i=1}^n$ in $\mathbb{R}^n$:

(F1) For any fixed matrix $\mathbf{B} \in \mathbb{O}_n$, the random vector $\mathbf{B}\mathbf{z}$ is standard Gaussian too. Equivalently, for any orthonormal basis $\mathbf{b}_1, \dots, \mathbf{b}_n$ of $\mathbb{R}^n$, the random vector $\sum_{i=1}^n z_i \mathbf{b}_i$ is standard Gaussian.

(F2) For any $k \in \{1, \dots, n-1\}$, the random variable $\sum_{j=1}^k z_j^2 / \sum_{i=1}^n z_i^2$ follows the beta distribution with parameters $k/2$ and $(n-k)/2$.

Step 1: Reduction to the case of $p_o = 0$. Suppose that $p_o \geq 1$. Let Π_o be the orthogonal projection from $\mathbb{R}^n$ onto $\text{span}(\mathbf{x}_1, \dots, \mathbf{x}_{p_o})$, so $\hat{\mathbf{y}}_o = \Pi_o \mathbf{y}$. The vector $\hat{\mathbf{y}}^* - \hat{\mathbf{y}}_o$ is the orthogonal projection of $\mathbf{y} - \Pi_o \mathbf{y}$ onto the linear span of the vectors $\mathbf{x}_j^* - \Pi_o \mathbf{x}_j^*$, $p_o < j \leq p$. All these vectors lie in the $(n - p_o)$ -dimensional linear space $\{\mathbf{x}_1, \dots, \mathbf{x}_{p_o}\}^\perp$, and the independent random vectors $\mathbf{x}_j^* - \Pi_o \mathbf{x}_j^*$, $p_o < j \leq p$, follow a standard Gaussian distribution on that space. The latter claim follows from (F1) when choosing an orthonormal basis of $\mathbb{R}^n$ such that $n - p_o$ of these basis vectors span $\{\mathbf{x}_1, \dots, \mathbf{x}_{p_o}\}^\perp$. Hence, we may assume without loss of generality that $p_o = 0$ and $\hat{\mathbf{y}}_o = \mathbf{0}$.

Step 2: The case of $p_o = 0$. We argue similarly as [Frankl and Machara \(1990\)](#). With the random matrix $\mathbf{X} := [\mathbf{x}_1^*, \dots, \mathbf{x}_p^*] \in \mathbb{R}^{n \times p}$, the orthogonal projection $\hat{\mathbf{y}}^*$ is given by $\hat{\mathbf{y}}^* = \mathbf{X}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y}$, and

$$\|\hat{\mathbf{y}}^*\|^2 = \mathbf{y}^\top \mathbf{X}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y}.$$

The distribution of this random variable depends only on $\|\mathbf{y}\|^2$. Indeed, let $\mathbf{B} \in \mathbb{O}_n$ such that $\mathbf{B}\mathbf{y} = \|\mathbf{y}\| \mathbf{e}$ with $\mathbf{e} = (1, 0, \dots, 0)^\top$. Since $\tilde{\mathbf{X}} := \mathbf{B}\mathbf{X}$ has the same distribution as $\mathbf{X}$, we may conclude that

$$\|\hat{\mathbf{y}}^*\|^2 = \|\mathbf{y}\|^2 \mathbf{e}^\top \tilde{\mathbf{X}}(\tilde{\mathbf{X}}^\top \tilde{\mathbf{X}})^{-1} \tilde{\mathbf{X}}^\top \mathbf{e} \stackrel{d}{=} \|\mathbf{y}\|^2 \mathbf{e}^\top \mathbf{X}(\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{e}.$$

But then we may replace $\mathbf{y}$ with $\|\mathbf{y}\| \|\mathbf{z}\|^{-1} \mathbf{z}$ with a standard Gaussian random vector $\mathbf{z}$ which is independent from X . Conditional on X , this vector $\mathbf{z}$ has the same distribution as $\sum_{i=1}^n z_i \mathbf{b}_i$ with an orthonormal basis $\mathbf{b}_1, \dots, \mathbf{b}_n$ of $\mathbb{R}^n$ such that $\text{span}(\mathbf{x}_1^*, \dots, \mathbf{x}_p^*) = \text{span}(\mathbf{b}_1, \dots, \mathbf{b}_p)$, and the orthogonal projection of $\mathbf{z}$ onto the latter space equals $\sum_{j=1}^p z_j \mathbf{b}_j$. Consequently, conditional on X ,

$$\frac{\|\hat{\mathbf{y}}^*\|^2}{\|\mathbf{y}\|^2} \stackrel{d}{=} \left\| \sum_{i=1}^n z_i \mathbf{b}_i \right\|^{-2} \left\| \sum_{j=1}^p z_j \mathbf{b}_j \right\|^2 = \sum_{j=1}^p z_j^2 / \sum_{i=1}^n z_i^2$$

follows a beta distribution with parameters $p/2$ and $(n-p)/2$. Since this does not depend on X , we may conclude that the ratio $\|\hat{\mathbf{y}}^*\|^2 / \|\mathbf{y}\|^2$ has distribution $\text{Beta}(p/2, (n-p)/2)$. $\square$

A.2. Haar measure on $\mathbb{O}_n$ in a nutshell

Recall that Haar_n is the unique probability measure on $\mathbb{O}_n$ such that a random matrix $H \sim \text{Haar}_n$ satisfies $TH \stackrel{d}{=} H$ for any fixed $T \in \mathbb{O}_n$ (left-invariance). For the reader's convenience, we explain a few well-known basic facts here.

Existence. To construct such a random matrix explicitly, let $Z \in \mathbb{R}^{n \times n}$ be a random matrix such that $TZ \stackrel{d}{=} Z$ for any fixed $T \in \mathbb{O}_n$, and $\text{rank}(Z) = n$ almost surely. This is true, for instance, if the n^2 entries of Z are independent with distribution $N(0, 1)$. Then one can easily verify that $H := Z(Z^\top Z)^{-1/2}$ is a random orthogonal matrix with left-invariant distribution. Another possible construction would be to apply Gram-Schmidt orthogonalization to the columns $\mathbf{z}_1, \mathbf{z}_2, \dots, \mathbf{z}_n$ of Z . This leads to a random matrix $H \in \mathbb{O}_n$ with left-invariant distribution, because the coefficients for the Gram-Schmidt procedure involve only the inner products $\mathbf{z}_i^\top \mathbf{z}_j$, and these remain unchanged if Z is replaced with TZ for some $T \in \mathbb{O}_n$.

Inversion-invariance and uniqueness. Let J, H be independent random matrices with left-invariant distribution on $\mathbb{O}_n$. By conditioning on J one can show that $J^\top H \stackrel{d}{=} H$, and conditioning on H reveals that $J^\top H = (H^\top J)^\top \stackrel{d}{=} J^\top$. Hence, $J^\top \stackrel{d}{=} H$. In the special case that J is an independent copy of H , this shows that

$$H^\top \stackrel{d}{=} H.$$

And then, with the original J , we see that $J = (J^\top)^\top \stackrel{d}{=} H^\top \stackrel{d}{=} H$, i.e.

$$J \stackrel{d}{=} H.$$

Right-invariance. If $\mathbf{H} \sim \text{Haar}_n$, then its distribution is right-invariant in the sense that $\mathbf{HT} \stackrel{d}{=} \mathbf{H}$ for any fixed $\mathbf{T} \in \mathbb{O}_n$. This follows easily from left-invariance and inversion-invariance.

Random linear subspaces of $\mathbb{R}^n$. The second construction of $\mathbf{H} = [\mathbf{h}_1, \dots, \mathbf{h}_n]$ via Gram-Schmidt may be applied to a random matrix $\mathbf{Z}$ with independent columns $\mathbf{z}_1, \mathbf{z}_2, \dots, \mathbf{z}_n \sim N_n(\mathbf{0}, \mathbf{I})$. This shows that the first column of $\mathbf{H}$ has distribution $\text{Unif}(\mathbb{S}_n)$. Combined with right-invariance, applied to permutation matrices, this shows that any column of $\mathbf{H}$ has distribution $\text{Unif}(\mathbb{S}_n)$. Furthermore, for $k \in \{1, 2, \dots, n\}$ and arbitrary linearly independent vectors $\mathbf{x}_1, \dots, \mathbf{x}_k \in \mathbb{R}^n$,

$$\text{span}(\mathbf{z}_1, \dots, \mathbf{z}_k) \stackrel{d}{=} \text{span}(\mathbf{H}_1, \dots, \mathbf{H}_k) \stackrel{d}{=} \text{span}(\mathbf{H}\mathbf{x}_1, \dots, \mathbf{H}\mathbf{x}_k).$$

The first equality follows from the construction of $\mathbf{H}$. For the second one, let $\mathbf{B} \in \mathbb{R}^{k \times k}$ be a nonsingular matrix such that the columns of $[\mathbf{t}_1, \dots, \mathbf{t}_k] = [\mathbf{x}_1, \dots, \mathbf{x}_k]\mathbf{B}$ are orthonormal. Extending them to an orthonormal basis $\mathbf{t}_1, \mathbf{t}_2, \dots, \mathbf{t}_n$ of $\mathbb{R}^n$,

$$\text{span}(\mathbf{H}\mathbf{x}_1, \dots, \mathbf{H}\mathbf{x}_k) = \mathbf{H}[\mathbf{x}_1, \dots, \mathbf{x}_k]\mathbb{R}^k = \mathbf{H}[\mathbf{t}_1, \dots, \mathbf{t}_k]\mathbb{R}^k,$$

and this is the linear span of the first k columns of $\mathbf{H}[\mathbf{t}_1, \dots, \mathbf{t}_n]$. But by right-invariance, the latter matrix has the same distribution as $\mathbf{H}$, because $[\mathbf{t}_1, \dots, \mathbf{t}_n] \in \mathbb{O}_n$.

A.3. Equivalence of the null hypotheses H_0 , H'_0 and H''_0

Recall that we consider a random vector $\mathbf{y}$ with continuous distribution on $\mathbb{R}^d$. Suppose first that H_0 is true. Then $\mathbf{y}$ has the same distribution as $\|\mathbf{y}\| \|\mathbf{z}\|^{-1} \mathbf{z}$, where $\mathbf{y}$ and $\mathbf{z} \sim N_n(\mathbf{0}, \mathbf{I})$ are independent. For any fixed $\mathbf{T} \in \mathbb{O}_n$, orthogonal invariance of the standard Gaussian distribution implies that $\mathbf{T}\mathbf{z} \stackrel{d}{=} \mathbf{z}$, so

$$\mathbf{T}\mathbf{y} \stackrel{d}{=} \|\mathbf{y}\| \|\mathbf{z}\|^{-1} \mathbf{T}\mathbf{z} = \|\mathbf{y}\| \|\mathbf{T}\mathbf{z}\|^{-1} \mathbf{T}\mathbf{z} \stackrel{d}{=} \|\mathbf{y}\| \|\mathbf{z}\|^{-1} \mathbf{z} \stackrel{d}{=} \mathbf{y}.$$

Consequently, H'_0 is true as well.

Now suppose that H'_0 is true. If $\mathbf{H} \sim \text{Haar}_n$ and $\mathbf{y}$ are independent, then for any measurable set $B \subset \mathbb{R}^n$,

$$\mathbb{P}(\mathbf{H}\mathbf{y} \in B) = \mathbb{E} \mathbb{P}(\mathbf{H}\mathbf{y} \in B | \mathbf{H}) = \mathbb{E} \mathbb{P}(\mathbf{y} \in B | \mathbf{H}) = \mathbb{P}(\mathbf{y} \in B),$$

where $\mathbb{P}(\mathbf{H}\mathbf{y} \in B | \mathbf{H}) = \mathbb{P}(\mathbf{y} \in B | \mathbf{H})$ by H'_0 .

Finally, suppose that H''_0 is true. Then for any measurable set $B \subset \mathbb{R}^n$,

$$\begin{aligned} \mathbb{P}(\mathbf{y} \in B) &= \mathbb{P}(\mathbf{H}\mathbf{y} \in B) = \mathbb{E} \mathbb{P}(\mathbf{H}\mathbf{y} \in B | \mathbf{y}) \\ &= \mathbb{E} \mathbb{P}(\|\mathbf{y}\| \mathbf{u} \in B | \mathbf{y}) = \mathbb{P}(\|\mathbf{y}\| \mathbf{u} \in B), \end{aligned}$$

where $\mathbf{u} \sim \text{Unif}(\mathbb{S}_n)$ and $\mathbf{y}$ are independent. Thus H_0 is satisfied as well.

A.4. Proof of Lemma 5

We first construct a particular point $\hat{\beta} = \hat{\beta}_{\ell, \alpha}(\mathbf{y}) \in \mathbb{R}^n$ such that $S_\ell(\mathbf{y} - X\hat{\beta}) \leq \kappa_{\ell, \alpha}$ and $\|\hat{\beta}\|_0 = k_{\ell, \alpha}(\mathbf{y})$. To this end let $(\sigma(1), \sigma(2), \dots, \sigma(n))$ be a permutation of $(1, 2, \dots, n)$ such that $G_i(\mathbf{y}) = \|\mathbf{a}_{\sigma(i)}\|^{-1} |\mathbf{a}_{\sigma(i)}^\top \mathbf{y}|$ for $1 \leq i \leq n$. Let $k_{\ell, \alpha}(\mathbf{y}) = \ell - 1 - j_o + i_o$ with indices $1 \leq i_o \leq j_o \leq \ell$ such that $|G_{i_o}(\mathbf{y})|/|G_{j_o}(\mathbf{y})| \leq \kappa_{\ell, \alpha}$. Now we set

$$\hat{\beta}_{\sigma(i)} := \begin{cases} 0 & \text{if } i_o \leq i \leq j_o \text{ or } i > \ell, \\ \text{sign}(\mathbf{a}_{\sigma(i)}^\top \mathbf{y}) (|\mathbf{a}_{\sigma(i)}^\top \mathbf{y}| - \|\mathbf{a}_{\sigma(i)}\| G_{i_o}(\mathbf{y})) & \text{if } 1 \leq i < i_o, \\ \text{sign}(\mathbf{a}_{\sigma(i)}^\top \mathbf{y}) (|\mathbf{a}_{\sigma(i)}^\top \mathbf{y}| - \|\mathbf{a}_{\sigma(i)}\| G_{j_o}(\mathbf{y})) & \text{if } j_o < i \leq \ell. \end{cases}$$

Obviously, this defines a vector $\hat{\beta} \in \mathbb{R}^n$ such that $\|\hat{\beta}\|_0 = k_{\ell, \alpha}(\mathbf{y})$. Note also that the numbers $\|\mathbf{a}_{\sigma(i)}\|^{-1} |\mathbf{a}_{\sigma(i)}^\top (\mathbf{y} - X\hat{\beta})|$ are nonincreasing in $i \in \{1, \dots, n\}$ with

$$\|\mathbf{a}_{\sigma(1)}\|^{-1} |\mathbf{a}_{\sigma(1)}^\top (\mathbf{y} - X\hat{\beta})| = G_{i_o}(\mathbf{y}), \quad \|\mathbf{a}_{\sigma(\ell)}\|^{-1} |\mathbf{a}_{\sigma(\ell)}^\top (\mathbf{y} - X\hat{\beta})| = G_{j_o}(\mathbf{y}).$$

Consequently, $\hat{\beta} \in \tilde{C}_{k_{\ell, \alpha}(\mathbf{y})}(\mathbf{y})$ for $k \geq k_{\ell, \alpha}(\mathbf{y})$.

On the other hand, let $k_{\ell, \alpha}(\mathbf{y}) > 0$. We have to show that $S_\ell(\mathbf{y} - X\mathbf{b}) > \kappa_{\ell, \alpha}(\mathbf{y})$ for any $\mathbf{b} \in \mathbb{R}^n$ with $k := \|\mathbf{b}\|_0 < k_{\ell, \alpha}(\mathbf{y})$. Then there exist indices $1 \leq i(1) < i(2) < \dots < i(n-k) \leq n$ and $1 \leq j(1) < j(2) < \dots < j(n-k) \leq n$ such that $G_{i(s)}(\mathbf{y}) = G_{j(s)}(\mathbf{y} - X\mathbf{b})$ for $1 \leq s \leq n-k$. But at least $n-k - (n-\ell) = \ell - k$ of the indices $j(1), \dots, j(n-k)$ are contained in $\{1, \dots, \ell\}$. Consequently, if $S_\ell(\mathbf{y} - X\mathbf{b}) \leq \kappa_{\ell, \alpha}$, then

$$\begin{aligned} \kappa_{\ell, \alpha} &\geq G_1(\mathbf{y} - X\mathbf{b})/G_\ell(\mathbf{y} - X\mathbf{b}) \\ &\geq G_{j(1)}(\mathbf{y} - X\mathbf{b})/G_{j(\ell-k)}(\mathbf{y} - X\mathbf{b}) \\ &= G_{i(1)}(\mathbf{y})/G_{i(\ell-k)}(\mathbf{y}), \end{aligned}$$

whence $k_{\ell, \alpha}(\mathbf{y}) \leq \ell - 1 - (i(\ell-k) - i(1)) \leq \ell - 1 - (\ell - k - 1) = k$, a contradiction to $k < k_{\ell, \alpha}(\mathbf{y})$. $\square$

A.5. Technical details for Section 3.3

For any $x \geq 1$,

$$\begin{aligned} \mathbb{P}(S_\ell(\mathbf{z}) \leq x) &= \mathbb{P}(\tilde{\Phi}^{-1}(U_n) \leq x \tilde{\Phi}^{-1}(U_{(n+1-\ell)})) \\ &= \mathbb{E} \mathbb{P}(U_n \leq \tilde{\Phi}[x \tilde{\Phi}^{-1}(U_{(n+1-\ell)})] \mid U_{(n+1-\ell)}) \\ &= \mathbb{E} \left(\left(\frac{\tilde{\Phi}[x \tilde{\Phi}^{-1}(U_{(n+1-\ell)})] - U_{(n+1-\ell)}}{1 - U_{(n+1-\ell)}} \right)^{\ell-1} \right), \end{aligned}$$

where the last step follows from the fact that conditionally on $U_{(n+1-\ell)}$, the random variable $U_{(n)}$ has the same distribution as the maximum of $\ell - 1$ independent random variables with uniform distribution on $[U_{(n+1-\ell)}, 1]$. Since $U_{(n+1-\ell)}$ has distribution $\text{Beta}(n + 1 - \ell, \ell)$, the latter expectation may be expressed as

$$\int_0^1 \left(\frac{\tilde{\Phi}[\chi \tilde{\Phi}^{-1}(B^{-1}(u))] - B^{-1}(u)}{1 - B^{-1}(u)} \right)^{\ell-1} du,$$

where B^{-1} is the quantile function of $\text{Beta}(n + 1 - \ell, \ell)$. This integral can be computed numerically, and the quantile $\chi_{\ell, \alpha}$ can be found via bisection.

REFERENCES

- L. DAVIES, L. DÜMBGEN (2024). *Gaussian covariates and linear regression*. Preprint in preparation.
- M. L. EATON (1983). *Multivariate Statistics: a Vector Space Approach*. Wiley Series in Probability and Statistics, Wiley and Sons, Chichester.
- M. L. EATON (1989). *Group invariance applications in statistics*. NSF-CBMS Regional Conference Series in Probability and Statistics, 1. Institute of Mathematical Statistics, Hayward, CA; American Statistical Association, Alexandria, VA.
- P. FRANKL, H. MAEHARA (1990). *Some geometric applications of the beta distribution*. Annals of the Institute of Statistical Mathematics, 42, no. 3, pp. 463–474.
- D. FREEDMAN, D. LANE (1983). *A nonstochastic interpretation of reported significance levels*. Journal of Business & Economic Statistics, 1, no. 4, pp. 292–298.
- F. E. KENNEDY (1995). *Randomization tests in econometrics*. Journal of Business & Economic Statistics, 13, no. 1, pp. 85–94.
- K. MARDIA, J. KENT, J. BIBBY (1979). *Multivariate Analysis*. Academic Press, London.
- R. G. MILLER, JR. (1981). *Simultaneous statistical inference, Second ed.* Springer-Verlag, New York.
- H. SCHEFFÉ (1959). *Analysis of Variance*. John Wiley and Sons, New York.
- A. M. WINKLER, G. R. RIDGWAY, M. A. WEBSTER, S. M. SMITH, T. E. NICHOLS (2014). *Permutation inference for the general linear model*. NeuroImage, 92, pp. 381–397.