

# DEBIASING CONVEX REGULARIZED ESTIMATORS AND INTERVAL ESTIMATION IN LINEAR MODELS

BY PIERRE C. BELLEC<sup>a</sup> AND CUN-HUI ZHANG<sup>b</sup>

Department of Statistics, Rutgers University, <sup>a</sup>[pierre.bellec@rutgers.edu](mailto:pierre.bellec@rutgers.edu), <sup>b</sup>[czhang@stat.rutgers.edu](mailto:czhang@stat.rutgers.edu)

New upper bounds are developed for the  $L_2$  distance between  $\xi/\text{Var}[\xi]^{1/2}$  and linear and quadratic functions of  $\mathbf{z} \sim N(\mathbf{0}, \mathbf{I}_n)$  for random variables of the form  $\xi = \mathbf{z}^\top f(\mathbf{z}) - \text{div } f(\mathbf{z})$ . The linear approximation yields a central limit theorem when the squared norm of  $f(\mathbf{z})$  dominates the squared Frobenius norm of  $\nabla f(\mathbf{z})$  in expectation.

Applications of this normal approximation are given for the asymptotic normality of debiased estimators in linear regression with correlated design and convex penalty in the regime  $p/n \leq \gamma$  for constant  $\gamma \in (0, \infty)$ . For the estimation of linear functions  $\langle \mathbf{a}_0, \boldsymbol{\beta} \rangle$  of the unknown coefficient vector  $\boldsymbol{\beta}$ , this analysis leads to asymptotic normality of the debiased estimate for most normalized directions  $\mathbf{a}_0$ , where “most” is quantified in a precise sense. This asymptotic normality holds for any convex penalty if  $\gamma < 1$  and for any strongly convex penalty if  $\gamma \geq 1$ . In particular, the penalty needs not be separable or permutation invariant. By allowing arbitrary regularizers, the results vastly broaden the scope of applicability of debiasing methodologies to obtain confidence intervals in high dimensions. In the absence of strong convexity for  $p > n$ , asymptotic normality of the debiased estimate is obtained for the Lasso and the group Lasso under additional conditions. For general convex penalties, our analysis also provides prediction and estimation error bounds of independent interest.

## 1. Introduction. Consider the linear model

$$(1.1) \quad \mathbf{y} = \mathbf{X}\boldsymbol{\beta} + \boldsymbol{\varepsilon}$$

with an unknown coefficient vector  $\boldsymbol{\beta} \in \mathbb{R}^p$ , a Gaussian noise vector  $\boldsymbol{\varepsilon} \sim N(\mathbf{0}, \sigma^2 \mathbf{I}_n)$  and a Gaussian design matrix  $\mathbf{X} \in \mathbb{R}^{n \times p}$  with i.i.d.  $N(\mathbf{0}, \boldsymbol{\Sigma})$  rows independent of  $\boldsymbol{\varepsilon}$ . We assume throughout the sequel that  $\boldsymbol{\Sigma}$  is invertible. The paper develops confidence intervals for  $\theta = \langle \mathbf{a}_0, \boldsymbol{\beta} \rangle$  from a given regularized initial estimator  $\widehat{\boldsymbol{\beta}} \in \mathbb{R}^p$ , using a technique referred to as *debiasing*: a correction to the initial estimate  $\langle \mathbf{a}_0, \widehat{\boldsymbol{\beta}} \rangle$  in the direction  $\mathbf{a}_0$  is constructed so that the “*debiased*” estimate can be used for inference about  $\theta = \langle \mathbf{a}_0, \boldsymbol{\beta} \rangle$ .

**1.1. Regularization induces bias.** If  $\mathbf{X}^\top \mathbf{X}$  is invertible, the unregulated least-squares estimate  $\widehat{\boldsymbol{\beta}}^{ls} = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y}$  is unbiased, that is,  $\mathbb{E}[\widehat{\boldsymbol{\beta}}^{ls} - \boldsymbol{\beta} | \mathbf{X}] = \mathbf{0}$ . On the other hand, if the square loss is regularized with an additive penalty,

$$(1.2) \quad \widehat{\boldsymbol{\beta}} = \arg \min_{\mathbf{b} \in \mathbb{R}^p} \|\mathbf{y} - \mathbf{X}\mathbf{b}\|^2 / (2n) + g(\mathbf{b})$$

for penalty functions commonly used in high-dimensional statistics such as  $g(\mathbf{b}) = \lambda \|\mathbf{b}\|_1$  for  $\lambda > 0$  (Lasso) or  $g(\mathbf{b}) = \mu \|\mathbf{b}\|_2^2$  for  $\mu > 0$  (ridge regression), then  $\widehat{\boldsymbol{\beta}}$  is biased.

For ridge regression  $\widehat{\boldsymbol{\beta}} = (\mathbf{X}^\top \mathbf{X} + n\mu \mathbf{I}_p)^{-1} \mathbf{X}^\top \mathbf{y}$ , this bias can be quantified explicitly when  $\boldsymbol{\Sigma} = \mathbf{I}_p$  as a shrinkage to the origin. Let  $\sum_{i=1}^r \mathbf{u}_i s_i \mathbf{v}_i^\top$  be the SVD of  $\mathbf{X}$  with  $s_i > 0$  and

Received July 2020; revised October 2022.

*MSC2020 subject classifications.* Primary 62H12, 62G15; secondary 62F35, 62J07.

*Key words and phrases.* Bias correction, central limit theorem, confidence intervals, convex regularization, Gaussian Poincaré inequality, high-dimensional linear models, Lasso, Stein’s formula, variance estimation.

$r = \min(n, p)$ . By rotational invariance,  $\mathbf{v}_i$  is independent of  $s_i$  and uniformly distributed in the unit sphere in  $\mathbb{R}^p$ . Thus, with  $G_\gamma$  being the Marchenko–Pastur law,

$$\mathbb{E}[\widehat{\boldsymbol{\beta}}] = \mathbb{E}\left[\sum_{i=1}^r \frac{s_i^2 \mathbf{v}_i \mathbf{v}_i^\top \boldsymbol{\beta}}{s_i^2 + n\mu}\right] = \mathbb{E}\left[\sum_{i=1}^r \frac{p^{-1} s_i^2}{s_i^2 + n\mu}\right] \boldsymbol{\beta} \approx \boldsymbol{\beta} \int \frac{(r/p)x}{x + (r/p)\mu} G_\gamma(dx) \quad \text{as } \frac{p}{n} \rightarrow \gamma.$$

The Lasso penalty  $g(\mathbf{b}) = \lambda \|\mathbf{b}\|_1$  also introduces bias. For example, for deterministic orthonormal designs, the Lasso estimator of the coefficient  $\beta_j$  is the soft-thresholding of  $N(\beta_j, \sigma^2/n)$ , which is again biased toward the origin. For Gaussian designs with  $\boldsymbol{\Sigma} = \mathbf{I}_p$  and in an average sense, the Lasso is approximately the soft-thresholding of  $N(\beta_j, \tau_*^2/n)$  with certain  $\tau_* \geq \sigma$  under proper conditions [1]. Thus, with  $s_1 = \#\{j : |\beta_j| > \lambda\}$ , the squared bias of the Lasso,  $\|\boldsymbol{\beta} - \mathbb{E}[\widehat{\boldsymbol{\beta}}]\|_2^2$ , is expected to have no smaller order than the lower bound  $s_1 \lambda^2$  for its  $\ell_2$  risk [3], Theorem 3.1. Alternative approaches were proposed to remove or reduce the bias of the Lasso for strong signals, for example, by using concave penalty functions (e.g., SCAD [24], MCP [48]) or iterated hard thresholding algorithms [13]. These approaches yield an error term of the order  $(\|\boldsymbol{\beta}\|_0 - s'_1) \lambda^2 + s'_1 \sigma^2/n$  where  $s'_1 = \{j = 1, \dots, p : |\beta_j| > c\lambda\}$  for some constant  $c > 0$  [25, 34], alleviating the bias of the Lasso for large coefficients at typical penalty levels  $\lambda > \sigma/n^{1/2}$ .

*Debiasing the Lasso, asymptotic normality and confidence intervals.* If the goal is the estimation of a single scalar parameter  $\theta = \langle \mathbf{a}_0, \boldsymbol{\beta} \rangle$  in a predetermined direction  $\mathbf{a}_0$  instead of the full vector  $\boldsymbol{\beta} \in \mathbb{R}^p$ , it is possible to correct the bias of the Lasso and to construct confidence intervals for  $\theta$ : there is already a vast literature on asymptotic normality of de-biased estimates in sparse linear regression for the Lasso [8, 9, 26–28, 33, 46, 51], among others. In this literature  $\mathbf{a}_0$  is usually the  $j$ th canonical basis vector and  $\beta_j$  the scalar parameter of interest. Given the Lasso  $\widehat{\boldsymbol{\beta}}$  as an initial estimator of  $\boldsymbol{\beta}$ , the idea is to add a debiasing term to achieve asymptotic normality, which then yields confidence intervals for  $\theta = \langle \mathbf{a}_0, \boldsymbol{\beta} \rangle$ . If  $s_0 = \|\boldsymbol{\beta}\|_0$  in (1.1), several debiased estimators have been proposed and their asymptotic normality hold under certain rate conditions on  $s_0, n, p$ . The earliest works on this topic [9, 26, 46, 51] provide asymptotic normality results in the regime  $s_0 \log(p)/\sqrt{n} \rightarrow 0$ . When  $s_0 \log(p)/\sqrt{n} \rightarrow 0$  indeed holds, the debiasing constructions in these papers are all first-order equivalent to each other, and under normalization  $\|\boldsymbol{\Sigma}^{-1/2} \mathbf{a}_0\|_2 = 1$  to

(1.3)

$$\widehat{\theta} = \underbrace{\langle \mathbf{a}_0, \widehat{\boldsymbol{\beta}} \rangle}_{\text{initial estimate}} + \underbrace{\|\mathbf{z}_0\|_2^{-2} \mathbf{z}_0^\top (\mathbf{y} - \mathbf{X} \widehat{\boldsymbol{\beta}})}_{\text{debiasing correction}},$$
$$\sqrt{n}(\widehat{\theta} - \theta) = \underbrace{\sqrt{n} \|\mathbf{z}_0\|_2^{-2} \mathbf{z}_0^\top \boldsymbol{\varepsilon}}_{\text{normal part}} + \underbrace{O_{\mathbb{P}}(R_n)}_{\text{remainder}},$$

where  $\mathbf{u}_0 = \boldsymbol{\Sigma}^{-1} \mathbf{a}_0 / \langle \mathbf{a}_0, \boldsymbol{\Sigma}^{-1} \mathbf{a}_0 \rangle$  and  $\mathbf{z}_0 = \mathbf{X} \mathbf{u}_0 \sim N(\mathbf{0}, \mathbf{I}_n)$ . While these works do not assume  $\boldsymbol{\Sigma}$  known and construct an estimated score vector  $\widehat{\mathbf{z}}$  for  $\mathbf{z}_0$ , the impact of using  $\widehat{\mathbf{z}}$  can be absorbed into the remainder in (1.3) with  $R_n = \sigma s_0 \log(p)/\sqrt{n}$ . The direction  $\mathbf{u}_0$  and the debiasing correction in (1.3) have a natural semiparametric interpretation [49]. Viewing  $\theta : \mathbb{R}^p \rightarrow \mathbb{R}$  as the function  $\theta(\boldsymbol{\beta}) = \langle \mathbf{a}_0, \boldsymbol{\beta} \rangle$ , the Fischer information for the estimation of  $\theta(\boldsymbol{\beta})$  in (1.1) is  $F_\theta = 1/(\sigma^2 \langle \mathbf{a}_0, \boldsymbol{\Sigma}^{-1} \mathbf{a}_0 \rangle)$ , and the direction  $\mathbf{u}_0$  above is the only  $\mathbf{u} \in \mathbb{R}^p$  with

(1.4)

$$\langle \nabla \theta(\widehat{\boldsymbol{\beta}}), \mathbf{u} \rangle = \langle \mathbf{a}_0, \mathbf{u} \rangle = 1$$

such that  $F_\theta$  is also the Fischer information in the one-dimensional submodel  $\{\widehat{\boldsymbol{\beta}} + t \mathbf{u}, t \in \mathbb{R}\}$ . For this reason, the line  $\{\widehat{\boldsymbol{\beta}} + t \mathbf{u}_0, t \in \mathbb{R}\}$  is referred to as the least-favorable one-dimensional submodel for the estimation of  $\theta$ . The normalization (1.4) ensures that  $\theta(\widehat{\boldsymbol{\beta}} + t \mathbf{u}) = \theta(\widehat{\boldsymbol{\beta}}) + t$  and  $\widehat{\theta} = \theta(\widehat{\boldsymbol{\beta}} + \widehat{t} \mathbf{u}_0)$  with  $\widehat{t} = \|\mathbf{z}_0\|_2^{-2} \mathbf{z}_0^\top (\mathbf{y} - \mathbf{X} \widehat{\boldsymbol{\beta}})$ , so that (1.3) replaces the initial  $\boldsymbol{\beta}$  with

its one-step correction  $\hat{\beta} + \hat{t}u_0$ , where  $\hat{t}$  maximizes the likelihood in the least-favorable sub-model. We refer to [10] for a systematic study of this semiparametric perspective.

If  $s_0 \log(p)/\sqrt{n} \rightarrow +\infty$  and  $\Sigma$  is unknown with bounded spectrum, the minimax estimation error of the form  $\sqrt{n}(\hat{\theta} - \theta)$  diverges for any estimator  $\hat{\theta}$  [16]. This rules out asymptotic normality results at the  $\sqrt{n}$  adjusted rate if  $s_0 \log(p)/\sqrt{n} \rightarrow +\infty$  and no further assumption is made on  $\Sigma$ . However, if  $z_0$  is known, (1.3) holds with  $R'_n = \sqrt{s_0 \log(p/s_0)/n}(1 + s_0/\sqrt{n})$ , providing asymptotic normality for sparsity levels  $s_0 \lesssim n^{2/3}$  up to logarithmic factors; cf. [8], Corollary 3.3. Similarly, [28], Theorem 3.8, provides (1.3) with  $a_0 = e_j \in \mathbb{R}^p$  a canonical basis vector and  $R'_n = \log(p)\sqrt{s_0/n} \max_j \|\Sigma^{-1}e_j\|_1$ . Already in the regime  $\sqrt{n} \lll s_0 \lll n^{2/3}$ , the arguments of [8, 28] differ significantly from the  $\ell_1$ - $\ell_\infty$  Hölder inequality argument of [9, 26, 46, 51]: while these earlier works prove asymptotic normality with a remainder term of order  $O_{\mathbb{P}}(s_0 \log(p)/\sqrt{n})$ , [8, 28] analyze explicitly the smaller order terms hidden in this  $O_{\mathbb{P}}(s_0 \log(p)/\sqrt{n})$  remainder.

For  $s_0 \ggg n^{2/3}$ , the debiasing correction in (1.3) needs to be modified:

$$(1.5) \quad \begin{aligned} \hat{\theta} &= \langle a_0, \hat{\beta} \rangle + (n - |\hat{S}|)^{-1} z_0^\top (y - X\hat{\beta}), \\ \sqrt{n}(\hat{\theta} - \theta) &= \sqrt{n} \|z_0\|_2^{-2} z_0^\top \varepsilon + O_{\mathbb{P}}(R'_n) \end{aligned}$$

with  $\hat{S} = \{j \in [p] : \hat{\beta}_j \neq 0\}$  and  $R'_n = \sigma(s_0 \log(p/s_0)/n)^{1/2}$ ; cf. [8], Theorem 3.1. For  $\|\Sigma^{-1/2}a_0\| = 1$ , the difference from (1.3) is the replacement of  $\|z_0\|_2^{-2} \approx n^{-1}$  in the debiasing correction with  $(n - |\hat{S}|)^{-1}$  to amplify it by a factor  $(1 - |\hat{S}|/n)^{-1}$ . This modification is required as soon as  $s_0 \ggg n^{2/3}$  up to logarithmic factors [8], Section 3. These asymptotic results for  $s_0 \ggg \sqrt{n}$  are amenable to the lack of knowledge of  $\Sigma$ : in this case, estimation of  $z_0$  is possible when  $\Sigma^{-1}a_0$  is sufficiently sparse; see [28] if the direction of interest  $a_0$  is canonical basis vector and [8], Section 2.2, for arbitrary direction  $a_0$ . These results [8, 28] for  $s_0 \ggg \sqrt{n}$  and correlated  $\Sigma$  are so far restricted to random Gaussian designs.

*Inflated asymptotic variance for nonvanishing prediction error.* In the results discussed so far for the Lasso,  $s_0 \log(p/s_0)/n \rightarrow 0$  or stronger conditions are required for asymptotic normality, and the asymptotic variance of  $\sqrt{n}(\hat{\theta} - \theta)$  is  $\sigma^2$ . The condition  $s_0 \log(p/s_0)/n \rightarrow 0$  implies the consistency of the Lasso in prediction and estimation thanks to error bounds of the form  $\|\Sigma^{1/2}(\hat{\beta} - \beta)\|_2^2 \lesssim s_0 \log(p/s_0)/n$  [5, 8, 38, 50]. It turns out that the asymptotic variance of  $\sqrt{n}(\hat{\theta} - \theta)$  is larger than  $\sigma^2$  if  $\|\Sigma^{1/2}(\hat{\beta} - \beta)\|_2^2$  does not vanish; this is the situation studied in the present work. The literature on asymptotic normality of debiased estimates in the regime

$$(1.6) \quad p/n \rightarrow \gamma \in (0, +\infty), \quad s_0/n \rightarrow \kappa \in (0, 1)$$

for constants  $\gamma, \kappa > 0$  is more scarce. In this regime where  $p, n$  and  $s_0$  are all of the same order, [27, 33] provide asymptotic normality results for the debiased Lasso (1.5) in the estimation of  $\beta_j$  (canonical  $a_0 = e_j$ ) in the isotropic Gaussian design. In these works, the asymptotic variance of  $\sqrt{n}(\hat{\theta} - \theta)$  equals a constant  $\tau_*^2$  satisfying the system of two nonlinear equations in [1] and [33], Proposition 3.1, Theorem 3.1. The constant  $\tau_*^2$  is related to the residual sum of squares [33], Corollary 4.1, and out-of-sample error [33], Theorem 3.2, as in

$$(1 - |\hat{S}|/n)^{-2} \|y - X\hat{\beta}\|_2^2/n \rightarrow^{\mathbb{P}} \tau_*^2, \quad \sigma^2 + \|\Sigma^{1/2}(\hat{\beta} - \beta)\|_2^2 \rightarrow^{\mathbb{P}} \tau_*^2,$$

where  $\rightarrow^{\mathbb{P}}$  denotes convergence in probability. These results for  $\Sigma = I_p$  highlight that the asymptotic variance is strictly larger than  $\sigma^2$  when  $p, n$  are of the same order as in (1.6). This phenomenon in the regime (1.6) is generic: for instance, the asymptotic variance is also larger than  $\sigma^2$  for all permutation-invariant penalty functions [17], Proposition 4.3.

In this regime where  $n$  and  $p$  are of the same order, [21, 23] proved asymptotic normality and characterized the variance for unregularized  $M$ -estimators. For  $M$ -estimators, a debiasing correction is unnecessary due to the absence of regularization, and a rotational invariance argument reduces the problem of correlated designs to a corresponding uncorrelated one [23], Lemma 1. However, this rotational invariance is lost in the presence of a penalty such as the  $\ell_1$ -norm. New techniques are called for to analyze the asymptotic behavior in the regime (1.6) and under correlated designs of estimators that are not rotational invariant. More recently, the approximate message Ppassing techniques used in [21, 27] were used to obtain similar results in logistic regression [40]; but again, these techniques cannot handle the Lasso penalty for correlated design. A more detailed comparison with these works is made in Section 3.8. To our knowledge, there is no previous asymptotic normality result for debiased estimates in the regime (1.6) for correlated designs in the presence of a penalty not depending on  $\Sigma$  (i.e., in situations where rotational invariance does not hold). A main goal of the paper is to fill this gap. Available techniques that tackle the regime (1.6) assume, in addition to uncorrelated design, that the penalty is invariant under permutations of the  $p$  coefficients [1, 15, 17, 33] and that the empirical distribution of the true  $\{\sqrt{n}\beta_j, j \leq p\}$  converges to some prior distribution. A second goal of the present paper is to show that asymptotic normality of debiased estimates can be obtained beyond the Lasso and beyond permutation-invariant penalty functions, without imposing the convergence of the empirical distribution of the normalized coefficients  $\{\sqrt{n}\beta_j, j \leq p\}$ .

1.2. *A general construction of debiased estimators.* This section describes a general approach to systematically construct de-biased estimates in the linear model (1.1) where  $X$  has i.i.d.  $N(\mathbf{0}, \Sigma)$  rows. Our goal is to construct confidence intervals for the one-dimensional parameter  $\theta = \langle \mathbf{a}_0, \boldsymbol{\beta} \rangle$ . Consider an initial estimator  $\hat{\boldsymbol{\beta}}$ , viewed as a function of  $(\mathbf{y}, X)$ , that is,  $\hat{\boldsymbol{\beta}} : \mathbb{R}^{n \times (1+p)} \rightarrow \mathbb{R}^p$  and assume that this function  $\hat{\boldsymbol{\beta}}$  is Fréchet<sup>1</sup> differentiable. For a given observed data  $(\mathbf{y}, X)$  from the linear model (1.1) and a  $\hat{\boldsymbol{\beta}}$  Fréchet differentiable at  $(\mathbf{y}, X)$ , there exist uniquely matrices  $\hat{H} \in \mathbb{R}^{n \times n}$  and  $\hat{G} \in \mathbb{R}^{n \times p}$  such that

$$\begin{aligned} X\hat{\boldsymbol{\beta}}(\mathbf{y} + \boldsymbol{\eta}, X) - X\hat{\boldsymbol{\beta}}(\mathbf{y}, X) &= \hat{H}^\top \boldsymbol{\eta} + o(\|\boldsymbol{\eta}\|), \\ \hat{\boldsymbol{\beta}}(\mathbf{y}, X + \boldsymbol{\eta} \mathbf{a}_0^\top) - \hat{\boldsymbol{\beta}}(\mathbf{y}, X) &= \hat{G}^\top \boldsymbol{\eta} + o(\|\boldsymbol{\eta}\|) \end{aligned} \tag{1.7}$$

for all  $\boldsymbol{\eta} \in \mathbb{R}^n$ . With  $X = (x_{ij})_{i \in [n], j \in [p]}$ , if the partial derivatives of  $\hat{\boldsymbol{\beta}}(\mathbf{y}, X)$  at the observed data  $(\mathbf{y}, X)$  are  $(\partial/\partial x_{ij})\hat{\boldsymbol{\beta}}(\mathbf{y}, X)$  and  $(\partial/\partial y_i)\hat{\boldsymbol{\beta}}(\mathbf{y}, X)$  then (1.7) implies  $\hat{H}^\top \mathbf{e}_i = X(\partial/\partial y_i)\hat{\boldsymbol{\beta}}(\mathbf{y}, X)$  and  $\hat{G}^\top \mathbf{e}_i = \sum_{j=1}^p \langle \mathbf{a}_0, \mathbf{e}_j \rangle (\partial/\partial x_{ij})\hat{\boldsymbol{\beta}}(\mathbf{y}, X)$  for canonical basis vectors  $\mathbf{e}_i \in \mathbb{R}^n$  and  $\mathbf{e}_j \in \mathbb{R}^p$ . The derivatives of  $\hat{\boldsymbol{\beta}}$  and the matrices  $\hat{H}$  and  $\hat{G}$  can be computed by only looking at the observed data  $(\mathbf{y}, X)$ , for instance by finite difference schemes.

Next, consider the function  $\phi$  defined as

$$\phi : \mathbb{R}^{n \times (1+p)} \rightarrow \mathbb{R}^n, \quad (\mathbf{y}, X) \mapsto \phi(\mathbf{y}, X) = X\hat{\boldsymbol{\beta}}(\mathbf{y}, X) - \mathbf{y}.$$

If  $\hat{\boldsymbol{\beta}}$  is differentiable at  $(\mathbf{y}, X)$ , then  $\phi$  is differentiable as well. By the product and chain rules,

$$\begin{aligned} \phi(\mathbf{y} + \boldsymbol{\eta}, X) - \phi(\mathbf{y}, X) &= [\hat{H} - I_n]^\top \boldsymbol{\eta} + o(\|\boldsymbol{\eta}\|), \\ \phi(\mathbf{y}, X + \boldsymbol{\eta} \mathbf{a}_0^\top) - \phi(\mathbf{y}, X) &= [\langle \mathbf{a}_0, \hat{\boldsymbol{\beta}} \rangle I_n + \hat{G} X^\top]^\top \boldsymbol{\eta} + o(\|\boldsymbol{\eta}\|). \end{aligned} \tag{1.8}$$

<sup>1</sup>Although the Fréchet derivative is the usual definition of derivative in finite dimension, we write Fréchet to emphasize that the derivative is linear. Linearity may fail for weaker notions such as Gateaux differentiability.

If the partial derivatives of  $\phi$  are  $(\partial/\partial y_i)\phi$  and  $(\partial/\partial x_{ij})\phi$ , the second line of the previous display is equivalently rewritten as

$$\sum_{j=1}^p \langle \mathbf{a}_0, \mathbf{e}_j \rangle (\partial/\partial x_{ij})\phi(\mathbf{y}, \mathbf{X}) = [\langle \mathbf{a}_0, \hat{\boldsymbol{\beta}} \rangle \mathbf{I}_n + \hat{\mathbf{G}} \mathbf{X}^\top]^\top \mathbf{e}_i$$

for each canonical basis vector  $\boldsymbol{\eta} = \mathbf{e}_i \in \mathbb{R}^n$ .

Observe that the arguments  $(\mathbf{y}, \mathbf{X})$  of  $\phi$  are centered and jointly normal random variables and their correlations are computed explicitly, for example,  $\mathbb{E}[x_{ij}y_l] = \mathbf{e}_j^\top \boldsymbol{\Sigma} \boldsymbol{\beta} \mathbf{I}_{\{i=l\}}$ , with basis vectors  $\mathbf{e}_j \in \mathbb{R}^p$ . One version of Stein's formula, also known as Gaussian integration by parts, is  $\mathbb{E}[Gh(Z_1, \dots, Z_q)] = \sum_{k=1}^q \mathbb{E}[GZ_k] \mathbb{E}[(\partial/\partial z_k)h(Z_1, \dots, Z_q)]$  provided that the function  $h(z_1, \dots, z_q)$  is differentiable and that  $G, Z_1, \dots, Z_q$  are centered jointly normal random variables, provided the existence of the expectations [41], Appendix A.4. We leverage this version of Stein's formula to obtain an unbiased estimating equation involving only one unknown parameter, the scalar  $\theta = \langle \mathbf{a}_0, \boldsymbol{\beta} \rangle$  of interest. For  $G_i = \mathbf{e}_i^\top \mathbf{X} \boldsymbol{\Sigma}^{-1} \mathbf{a}_0$ , we find  $\mathbb{E}[G_i y_k] = \mathbb{E}[G_i x_{kj}] = 0$  if  $i \neq k$  while  $\mathbb{E}[G_i x_{ij}] = \langle \mathbf{a}_0, \mathbf{e}_j \rangle$  and  $\mathbb{E}[G_i y_i] = \langle \mathbf{a}_0, \boldsymbol{\beta} \rangle$  so that by reading the partial derivatives in (1.8),

$$\begin{aligned} \mathbb{E}[G_i \phi_i(\mathbf{y}, \mathbf{X})] &= \langle \mathbf{a}_0, \boldsymbol{\beta} \rangle \mathbb{E}\left[\frac{\partial \phi_i}{\partial y_i}(\mathbf{y}, \mathbf{X})\right] + \sum_{j=1}^p \langle \mathbf{a}_0, \mathbf{e}_j \rangle \mathbb{E}\left[\frac{\partial \phi_i}{\partial x_{ij}}(\mathbf{y}, \mathbf{X})\right] \\ (1.9) \quad &= \mathbb{E}[\langle \mathbf{a}_0, \boldsymbol{\beta} \rangle (\hat{H}_{ii} - 1)] + \mathbb{E}[\langle \mathbf{a}_0, \hat{\boldsymbol{\beta}} \rangle + \mathbf{e}_i^\top \hat{\mathbf{G}} \mathbf{X}^\top \mathbf{e}_i]. \end{aligned}$$

Summing over  $i = 1, \dots, n$  and using  $\phi(\mathbf{y}, \mathbf{X}) = \mathbf{X} \hat{\boldsymbol{\beta}} - \mathbf{y}$ , we find that

$$\mathbb{E}[(\mathbf{X} \boldsymbol{\Sigma}^{-1} \mathbf{a}_0, \mathbf{X} \hat{\boldsymbol{\beta}} - \mathbf{y})] = \mathbb{E}[-\langle \mathbf{a}_0, \boldsymbol{\beta} \rangle \text{trace}[\mathbf{I}_n - \hat{\mathbf{H}}] + \langle \mathbf{a}_0, \hat{\boldsymbol{\beta}} \rangle n + \text{trace}[\mathbf{X}^\top \hat{\mathbf{G}}]].$$

To transform this equation into a form representative of the results of the paper, define the scalars  $\hat{\text{df}}$  and  $\hat{A}$  by

$$(1.10) \quad \hat{\text{df}} = \text{trace}[\hat{\mathbf{H}}], \quad \hat{A} = \text{trace}[\mathbf{X}^\top \hat{\mathbf{G}}] + \langle \mathbf{a}_0, \hat{\boldsymbol{\beta}} \rangle \hat{\text{df}}.$$

The notation  $\hat{\text{df}}$  underlines that  $\text{trace}[\hat{\mathbf{H}}]$  has the interpretation of degrees-of-freedom of the estimator  $\hat{\boldsymbol{\beta}}$  in Stein's Unbiased Risk Estimate (SURE) [37]: regarding  $\hat{\boldsymbol{\mu}} = \mathbf{X} \hat{\boldsymbol{\beta}}$  as an estimate of  $\boldsymbol{\mu} = \mathbf{X} \boldsymbol{\beta}$  in the Gaussian sequence model with observation  $\mathbf{y} = \boldsymbol{\mu} + \boldsymbol{\varepsilon}$ , the quantity  $\widehat{\text{SURE}} = \|\mathbf{y} - \hat{\boldsymbol{\mu}}\|^2 + 2\sigma^2 \hat{\text{df}} - \sigma^2 n$  is an unbiased estimate of the in-sample error  $\|\hat{\boldsymbol{\mu}} - \boldsymbol{\mu}\|^2 = \|\mathbf{X}(\hat{\boldsymbol{\beta}} - \boldsymbol{\beta})\|^2$ . With this notation, we obtain the *unbiased estimating equation*,

$$(1.11) \quad 0 = \mathbb{E}[(\mathbf{X} \boldsymbol{\Sigma}^{-1} \mathbf{a}_0, \mathbf{y} - \mathbf{X} \hat{\boldsymbol{\beta}}) + (n - \hat{\text{df}})(\langle \mathbf{a}_0, \hat{\boldsymbol{\beta}} \rangle - \theta) + \hat{A}],$$

where the only unobserved quantity inside the expectation is  $\theta = \langle \mathbf{a}_0, \boldsymbol{\beta} \rangle$ , the scalar parameter we wish to estimate. In the above application of Stein's formula,  $G_i = \mathbf{e}_i^\top \mathbf{X} \boldsymbol{\Sigma}^{-1} \mathbf{a}_0$  was chosen on purpose so that  $\boldsymbol{\beta}$  appears in (1.9) only through  $\langle \mathbf{a}_0, \boldsymbol{\beta} \rangle$  thanks to  $\mathbb{E}[G_i y_i] = \langle \mathbf{a}_0, \boldsymbol{\beta} \rangle$ . Note that replacing  $G_i$  in (1.9) by  $\mathbf{e}_i^\top \mathbf{X} \mathbf{u}$  for any  $\mathbf{u} \in \mathbb{R}^p$  not proportional to  $\boldsymbol{\Sigma}^{-1} \mathbf{a}_0$  brings a scalar projection of  $\boldsymbol{\beta}$  different from  $\langle \mathbf{a}_0, \boldsymbol{\beta} \rangle$ : this shows the unique role of the random vector  $\mathbf{X} \boldsymbol{\Sigma}^{-1} \mathbf{a}_0$  to derive an unbiased estimating equation for  $\theta = \langle \mathbf{a}_0, \boldsymbol{\beta} \rangle$ . It is notable that the direction  $\boldsymbol{\Sigma}^{-1} \mathbf{a}_0$  coincides with the least-favorable direction described around (1.4). Equation (1.11) is obtained for an arbitrary initial estimator  $\hat{\boldsymbol{\beta}}$  provided that its derivatives with respect to  $(\mathbf{y}, \mathbf{X})$  exist and the integrability conditions hold to ensure existence of the expectations involved. From (1.11), the method of moments suggests to estimate  $\theta$  with  $\hat{\theta} = \langle \mathbf{a}_0, \hat{\boldsymbol{\beta}} \rangle + (n - \hat{\text{df}})^{-1}(\langle \mathbf{X} \boldsymbol{\Sigma}^{-1} \mathbf{a}_0, \mathbf{y} - \mathbf{X} \hat{\boldsymbol{\beta}} \rangle + \hat{A})$ , which resembles (1.5) for the Lasso for  $\hat{\text{df}} = |\hat{S}|$  and  $\mathbf{X} \boldsymbol{\Sigma}^{-1} \mathbf{a}_0 = \mathbf{z}_0$  under the normalization  $\langle \mathbf{a}_0, \boldsymbol{\Sigma}^{-1} \mathbf{a}_0 \rangle = 1$ .

It is useful at this point to specialize the above derivation to an estimator for which all derivatives can be computed explicitly. For Ridge regression with penalty  $g(\mathbf{b}) = \mu \|\mathbf{b}\|_2^2$  for

some  $\mu > 0$ ,  $\widehat{\boldsymbol{\beta}}(\mathbf{y}, \mathbf{X}) = (\mathbf{X}^\top \mathbf{X} + n\mu \mathbf{I}_p)^{-1} \mathbf{X}^\top \mathbf{y}$  and

$$(1.12) \quad \begin{aligned} \widehat{\mathbf{H}}^\top &= \mathbf{X}(\mathbf{X}^\top \mathbf{X} + n\mu \mathbf{I}_p)^{-1} \mathbf{X}^\top, \\ \widehat{\mathbf{G}}^\top &= (\mathbf{X}^\top \mathbf{X} + n\mu \mathbf{I}_p)^{-1} [\mathbf{a}_0(\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}})^\top - \mathbf{X}^\top \langle \mathbf{a}_0, \widehat{\boldsymbol{\beta}} \rangle]. \end{aligned}$$

Indeed, the derivatives of  $\widehat{\boldsymbol{\beta}}(\mathbf{y}, \mathbf{X})$  exist as it is the composition of elementary differentiable functions. Differentiation with respect to  $\mathbf{y}$  is straightforward as  $\widehat{\boldsymbol{\beta}}$  is linear in  $\mathbf{y}$ , while in order to compute  $\widehat{\mathbf{G}}$  we proceed by setting  $\mathbf{b}(t) = \widehat{\boldsymbol{\beta}}(\mathbf{y}, \mathbf{X}(t))$  with  $\mathbf{X}(t) = \mathbf{X} + t\eta\mathbf{a}_0^\top$ . Differentiation of the KKT conditions  $\mathbf{X}(t)^\top(\mathbf{y} - \mathbf{X}(t)\mathbf{b}(t)) = n\mu\mathbf{b}(t)$  at  $t = 0$  provides the directional derivative  $(d/dt)\mathbf{b}(t)|_{t=0} = \widehat{\mathbf{G}}^\top \eta$ . This gives (1.12). It follows from (1.12) that  $d\widehat{\mathbf{f}} = \text{trace}[\mathbf{X}(\mathbf{X}^\top \mathbf{X} + n\mu \mathbf{I}_p)^{-1} \mathbf{X}^\top]$  and  $\widehat{\mathbf{A}} = (\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}})^\top \mathbf{X}(\mathbf{X}^\top \mathbf{X} + n\mu \mathbf{I}_p)^{-1} \mathbf{a}_0$  for the quantities in (1.10) (for  $\widehat{\mathbf{A}}$ , note the fortuitous cancelation of the term  $\langle \mathbf{a}_0, \widehat{\boldsymbol{\beta}} \rangle d\widehat{\mathbf{f}}$ ). For the Lasso, similar differentiability formulae are derived in [8]. It is however, unclear how to obtain closed form formulae for the derivatives of  $\widehat{\boldsymbol{\beta}}$  for an arbitrary convex penalty  $g$  in (1.2).

We now set up some notation that will be useful for the rest of the paper, and derive again the unbiased estimating equation (1.11) using this new notation. Define

$$(1.13) \quad \mathbf{u}_0 = \boldsymbol{\Sigma}^{-1} \mathbf{a}_0 / \langle \mathbf{a}_0, \boldsymbol{\Sigma}^{-1} \mathbf{a}_0 \rangle, \quad \mathbf{z}_0 = \mathbf{X} \mathbf{u}_0, \quad \mathbf{Q}_0 = \mathbf{I}_{p \times p} - \mathbf{u}_0 \mathbf{a}_0^\top.$$

The normalizing constant in  $\mathbf{u}_0$  is such that  $\langle \mathbf{a}_0, \mathbf{u}_0 \rangle = 1$  holds so that the expression (1.13) for  $\mathbf{u}_0$  coincides with the direction of the least-favorable submodel discussed around (1.4). The vector  $\mathbf{z}_0$  is independent of  $\mathbf{X} \mathbf{Q}_0$  by construction as  $(\mathbf{z}_0, \mathbf{X} \mathbf{Q}_0)$  are jointly normal and uncorrelated. This follows by noting that  $\mathbf{X} \boldsymbol{\Sigma}^{-1/2}$  has i.i.d.  $N(0, 1)$  entries and

$$\mathbf{z}_0 = \mathbf{X} \boldsymbol{\Sigma}^{-1/2} \mathbf{v} / \|\boldsymbol{\Sigma}^{-1/2} \mathbf{a}_0\|, \quad \mathbf{X} \mathbf{Q}_0 = \mathbf{X} \boldsymbol{\Sigma}^{-1/2} (\mathbf{I}_p - \mathbf{v} \mathbf{v}^\top) \boldsymbol{\Sigma}^{1/2}$$

for the unit vector  $\mathbf{v} = \boldsymbol{\Sigma}^{-1/2} \mathbf{a}_0 / \|\boldsymbol{\Sigma}^{-1/2} \mathbf{a}_0\|$  as by construction of  $\mathbf{Q}_0$  matrix  $\mathbf{I}_p - \mathbf{v} \mathbf{v}^\top = \boldsymbol{\Sigma}^{1/2} \mathbf{Q}_0 \boldsymbol{\Sigma}^{-1/2}$  is the orthogonal projection onto  $\{\mathbf{v}\}^\perp$ . We summarize this as

$$(1.14) \quad \mathbf{X} = \mathbf{X} \mathbf{Q}_0 + \mathbf{z}_0 \mathbf{a}_0^\top \quad \text{with } \mathbf{z}_0 \sim N(\mathbf{0}, \|\boldsymbol{\Sigma}^{-1/2} \mathbf{a}_0\|^{-2} \mathbf{I}_n) \text{ independent of } \mathbf{X} \mathbf{Q}_0.$$

For brevity, we assume in the sequel and without loss of generality that the direction of interest  $\mathbf{a}_0$  is normalized such that

$$(1.15) \quad \|\boldsymbol{\Sigma}^{-1/2} \mathbf{a}_0\|^2 = \langle \mathbf{a}_0, \boldsymbol{\Sigma}^{-1} \mathbf{a}_0 \rangle = 1.$$

By definition of  $\mathbf{u}_0$  and  $\mathbf{z}_0$ , the normalization (1.15) gives  $\mathbf{z}_0 \sim N(\mathbf{0}, \mathbf{I}_n)$ .

Conditionally on  $(\mathbf{X} \mathbf{Q}_0, \boldsymbol{\varepsilon})$ , define the function  $f_{(\mathbf{X} \mathbf{Q}_0, \boldsymbol{\varepsilon})} : \mathbb{R}^n \rightarrow \mathbb{R}^n$  by

$$(1.16) \quad f_{(\mathbf{X} \mathbf{Q}_0, \boldsymbol{\varepsilon})}(\mathbf{z}_0) = \mathbf{X} \widehat{\boldsymbol{\beta}} - \mathbf{y}.$$

By (1.14) and the independence of  $\boldsymbol{\varepsilon}$  and  $\mathbf{X}$ , the conditional expectation given  $(\mathbf{X} \mathbf{Q}_0, \boldsymbol{\varepsilon})$  can be written as integrals against the Gaussian measure of  $\mathbf{z}_0$ , for example,

$$\mathbb{E}[\mathbf{z}_0^\top f_{(\mathbf{X} \mathbf{Q}_0, \boldsymbol{\varepsilon})}(\mathbf{z}_0) | (\mathbf{X} \mathbf{Q}_0, \boldsymbol{\varepsilon})] = \int (\mathbf{z}^\top f_{(\mathbf{X} \mathbf{Q}_0, \boldsymbol{\varepsilon})}(\mathbf{z})) e^{-\|\mathbf{z}\|_2^2/2} (\sqrt{2\pi})^{-n} d\mathbf{z}$$

since  $\mathbf{z}_0 \sim N(\mathbf{0}, \mathbf{I}_n)$ . As we argue conditionally on  $(\mathbf{X} \mathbf{Q}_0, \boldsymbol{\varepsilon})$ , we omit the dependence on  $(\mathbf{X} \mathbf{Q}_0, \boldsymbol{\varepsilon})$  and write simply  $f : \mathbb{R}^n \rightarrow \mathbb{R}^n$ . Since  $\mathbf{y} = \boldsymbol{\varepsilon} + \mathbf{X} \widehat{\boldsymbol{\beta}}$  and  $\mathbf{X} = \mathbf{X} \mathbf{Q}_0 + \mathbf{z}_0 \mathbf{a}_0^\top$ ,

$$f(\mathbf{z}_0) = \mathbf{X} \widehat{\boldsymbol{\beta}} (\boldsymbol{\varepsilon} + \mathbf{X} \mathbf{Q}_0 \boldsymbol{\beta} + \mathbf{z}_0 \mathbf{a}_0^\top \boldsymbol{\beta}, \mathbf{X} \mathbf{Q}_0 + \mathbf{z}_0 \mathbf{a}_0^\top) - \mathbf{X} \boldsymbol{\beta} - \boldsymbol{\varepsilon}.$$

The gradient  $\nabla f$  with respect to  $\mathbf{z}_0$ , holding  $(\mathbf{X} \mathbf{Q}_0, \boldsymbol{\varepsilon})$  fixed, can be computed by the product rule and the chain rule via (1.8):

$$(1.17) \quad \nabla f(\mathbf{z}_0)^\top = \mathbf{I}_n \langle \mathbf{a}_0, \widehat{\boldsymbol{\beta}} - \boldsymbol{\beta} \rangle + [\langle \mathbf{a}_0, \boldsymbol{\beta} \rangle \widehat{\mathbf{H}}^\top + \mathbf{X} \widehat{\mathbf{G}}^\top].$$

We adopt the usual convention that the gradient of a vector valued function is the transpose of its Jacobian. Computing the directional derivative of  $f$  in a direction  $\eta$  requires considering the difference of an expression at  $(\epsilon, XQ_0, z_0 + t\eta)$  minus the same expression at  $(\epsilon, XQ_0, z_0)$ , dividing by  $t$  and taking the limit as  $t \rightarrow 0$ ; this is equivalent to considering the difference of an expression at  $(\epsilon, X(t))$  with  $X(t) = X + t\eta a_0^\top$  minus the same expression at  $(\epsilon, X)$ , dividing by  $t$  and taking the limit as  $t \rightarrow 0$ .

Taking the trace of (1.17) and by definition of  $\widehat{df}$  and  $\widehat{A}$  in (1.10), the identity

$$(1.18) \quad \begin{aligned} -\xi_0 &\stackrel{\text{def}}{=} \text{div } f(z_0) - z_0^\top f(z_0) \\ &= (n - \widehat{df})(\langle a_0, \widehat{\beta} \rangle - \theta) + \langle z_0, y - X\widehat{\beta} \rangle + \widehat{A} \end{aligned}$$

holds where  $\text{div } f(z_0) = \text{trace}[\nabla f(z_0)]$ . Since  $\mathbb{E}[\text{div } f(z_0) - z_0^\top f(z_0) | (XQ_0, \epsilon)] = 0$  by Stein's formula [37], this provides the unbiased estimating equation (1.11). Reasoning conditionally on  $(\epsilon, XQ_0)$ , using Stein formulae with respect to  $z_0$  involving conditional expectations given  $(\epsilon, XQ_0)$  and gradients of the form  $\nabla f(z_0)$  holding  $(\epsilon, XQ_0)$  fixed will be a recurring theme throughout the paper. In this context, the function  $f$  itself depends on  $(\epsilon, XQ_0)$  as in (1.16), although the dependence on  $(\epsilon, XQ_0)$  is omitted for brevity.

In order to construct confidence intervals using the unbiased estimating equation (1.11), one may hope that the quantity (1.18) above is well behaved—ideally, approximately normal with mean zero and a variance that can be consistently estimated from the observed data. By the second-order Stein's formula in Proposition 2.1 below, which was already known to Stein [37], (8.6), in a different form, the conditional variance of (1.18) given  $(\epsilon, XQ_0)$  is

$$(1.19) \quad \begin{aligned} \text{Var}_0[\xi_0] &= \mathbb{E}_0[\|f(z_0)\|^2 + \text{trace}[\{\nabla f(z_0)\}^2]] \\ &= \mathbb{E}_0[V^*(\theta)] \quad \text{for } V^*(\theta) = \|y - X\widehat{\beta}\|^2 + \text{trace}[\{\nabla f(z_0)\}^2], \end{aligned}$$

where  $\mathbb{E}_0 = \mathbb{E}[\cdot | \epsilon, XQ_0]$  denotes the conditional expectation with respect to  $z_0$  given  $(\epsilon, XQ_0)$  and  $\text{Var}_0$  denotes the conditional variance given  $(\epsilon, XQ_0)$ . The gradient  $\nabla f(z_0)$  in (1.17) and the unbiased estimate  $V^*(\theta)$  of  $\text{Var}_0[\xi_0]$  only depend on the unknown parameter of interest  $\theta$  and observable quantities, and  $V^*(\theta)$  is quadratic in  $\theta$ .

Assume now we are in an ideal situation in the sense that both conditions below are satisfied: (i) The quantity (1.18) is approximately normally distributed conditionally on  $(\epsilon, XQ_0)$  and (ii)  $V^*(\theta)$  is a consistent estimator of (1.19), the conditional variance of the random variable (1.18). Then the set of  $\theta$  for which the inequality

$$(1.20) \quad [(n - \widehat{df})(\langle a_0, \widehat{\beta} \rangle - \theta) + \langle z_0, y - X\widehat{\beta} \rangle + \widehat{A}]^2 - V^*(\theta)z_{\alpha/2}^2 \leq 0$$

is satisfied is an  $(1 - \alpha)$ -confidence interval, where  $\mathbb{P}(|N(0, 1)| > z_{\alpha/2}) = 1 - \alpha$ . Solving the corresponding quadratic equality gives up to two solutions  $\Theta_1(z_{\alpha/2}) \leq \Theta_2(z_{\alpha/2})$  that are such that (1.20) holds with equality. These two solutions implicitly depend on the observables

$$\langle y - X\widehat{\beta}, z_0 \rangle, \quad \|y - X\widehat{\beta}\|^2, \quad \widehat{df}, \quad \widehat{A}, \quad a_0^\top \widehat{\beta}$$

and the derivatives of  $\widehat{\beta}$ . If the coefficient of  $\theta^2$  in the left-hand side of (1.20) is positive, (i.e., if the leading coefficient of (1.20), seen as a polynomial in  $\theta$  with data-driven coefficients, is positive), a  $(1 - \alpha)$  confidence interval for  $\theta = a_0^\top \beta$  is then given by

$$(1.21) \quad \widehat{CI} = [\Theta_1(z_{\alpha/2}), \Theta_2(z_{\alpha/2})].$$

We will show in the discussion surrounding (3.30) below that the dominant coefficient is positive and that the confidence interval is indeed of the above form if  $\widehat{\beta}$  is a convex penalized estimator. Although a variant of the above construction was briefly presented in [7], Section 6, (there the function  $z_0 \rightarrow XQ_0(\widehat{\beta} - \beta) - \epsilon$  is used), important questions remain unanswered to prove the validity of the general confidence interval in (1.21) and its applicability to commonly used regularized estimators.

1.3. *The rest of the paper is organized as follows.* Section 2 develops an  $L_2$  bound between  $\xi/\text{Var}[\xi]^{1/2}$  and  $N(0, 1)$  for random variables of the form  $\xi = \mathbf{z}^\top f(\mathbf{z}) - \text{div } f(\mathbf{z})$  where  $\mathbf{z} \sim N(\mathbf{0}, \mathbf{I}_n)$ . Section 3 uses this normal approximation to show the asymptotic normality of (1.18) and proves the consistency of the variance estimate  $V^*(\theta)$  in (1.19) in the regime where  $p$  and  $n$  are of the same order in the linear model (1.1) with correlated design. Section 4 provides closed-form formulas to apply the results in Section 3 to the Lasso, the group Lasso and twice continuously differentiable penalty functions. Section 7 contains the proofs of the results in Section 3. Appendix A provides a technical lemma on the integrability of smallest eigenvalue of Wishart matrices, Appendix B provides the proofs of the asymptotic normality results for the Lasso and group Lasso when  $p > n$  and Appendix C contains the proofs of the derivative formulae for the group Lasso.

1.4. *Notation.* For two reals  $\{a, b\}$ , let  $a \wedge b = \min\{a, b\}$ ,  $a \vee b = \max\{a, b\}$  and  $a_+ = a \vee 0$ . Let  $\mathbf{I}_d$  be the identity matrix of size  $d \times d$ , for example,  $d = n, p$ . For any  $p \geq 1$ , let  $[p]$  be the set  $\{1, \dots, p\}$ . Let  $\|\cdot\|$  be the Euclidean norm and  $\|\cdot\|_q$  the  $\ell_q$  norm of vectors for any  $q \geq 1$ , so that  $\|\cdot\| = \|\cdot\|_2$ . Let  $\|\cdot\|_{\text{op}}$  be the operator norm of matrices and  $\|\cdot\|_F$  the Frobenius norm. Let  $\phi_{\min}(\mathbf{S})$  be the smallest eigenvalue of a symmetric matrix  $\mathbf{S}$ . We use the notation  $\langle \cdot, \cdot \rangle$  for the canonical scalar product of vectors in  $\mathbb{R}^n$  or  $\mathbb{R}^p$ , that is,  $\langle \mathbf{a}, \mathbf{b} \rangle = \mathbf{a}^\top \mathbf{b}$  for two vectors  $\mathbf{a}, \mathbf{b}$  of the same dimension. For any event  $\Omega$ , denote by  $I_\Omega$  its indicator function. The unit sphere is  $S^{p-1} = \{\mathbf{x} \in \mathbb{R}^p : \|\mathbf{x}\| = 1\}$ . Convergence in distribution is denoted by  $\rightarrow^d$  and convergence in probability by  $\rightarrow^{\mathbb{P}}$ . Throughout the paper,  $C_0, C_1, \dots$  denote positive absolute constants,  $C_k(\gamma)$  positive constants depending on  $\gamma$  only and  $C_k(\gamma, \mu)$  on  $\{\gamma, \mu\}$  only.

For any vector  $\mathbf{v} = (v_1, \dots, v_p)^\top \in \mathbb{R}^p$  and set  $A \subset [p]$ , the vector  $\mathbf{v}_A \in \mathbb{R}^{|A|}$  is the restriction  $(v_j)_{j \in A}$ . For any  $n \times p$  matrix  $\mathbf{M}$  with columns  $(\mathbf{M}_1, \dots, \mathbf{M}_p)$  and any subset  $A \subset [p]$ , let  $\mathbf{M}_A = (\mathbf{M}_j, j \in A)$  be the matrix composed of columns of  $\mathbf{M}$  indexed by  $A$ . If  $\mathbf{M}$  is a symmetric matrix of size  $p \times p$  and  $A \subset [p]$ , then  $\mathbf{M}_{A,A}$  denotes the submatrix of  $\mathbf{M}$  with rows and columns in  $A$ , and  $\mathbf{M}_{A,A}^{-1}$  is the inverse of  $\mathbf{M}_{A,A}$ . For any square matrix  $\mathbf{M}$ , let  $\mathbf{M}^s = (\mathbf{M} + \mathbf{M}^\top)/2$  be its symmetrization giving the same quadratic form.

For a vector valued map  $h : \mathbb{R}^n \rightarrow \mathbb{R}^q$  with coordinates  $h_1, \dots, h_q : \mathbb{R}^n \rightarrow \mathbb{R}$ , the gradient  $\nabla h \in \mathbb{R}^{n \times q}$  is the matrix with columns  $\nabla h_1, \dots, \nabla h_q$ . Thus,  $\nabla h$  is the transpose of the Jacobian of  $h$  and  $h(\mathbf{x} + \boldsymbol{\eta}) = h(\mathbf{x}) + \nabla h(\mathbf{x})^\top \boldsymbol{\eta} + o(\|\boldsymbol{\eta}\|)$  if each coordinate  $h_i$  is Fréchet differentiable at  $\mathbf{x}$ . For deterministic matrices  $\mathbf{A} \in \mathbb{R}^{m \times q}$ ,  $\nabla(\mathbf{A}h) = (\nabla h)\mathbf{A}^\top \in \mathbb{R}^{n \times m}$ . For  $f$  in (2.1),  $\nabla f(\mathbf{x}) \in \mathbb{R}^{n \times n}$  and the divergence is  $\text{div } f(\mathbf{x}) = \text{trace}[\nabla f(\mathbf{x})]$ .

**2. Normal approximation in Stein's formula.** We develop in this section normal approximations for random variables of the form

$$(2.1) \quad \xi = \mathbf{z}^\top f(\mathbf{z}) - \text{div } f(\mathbf{z}),$$

for which Stein's formula [37] states  $\mathbb{E}[\xi] = 0$ , where  $\mathbf{z} \sim N(\mathbf{0}, \mathbf{I}_n)$  is standard normal and  $f : \mathbb{R}^n \rightarrow \mathbb{R}^n$ . We establish  $L_2$  bounds for the linear and quadratic approximations of  $\xi$  and construct consistent variance estimates in the related CLT.

Throughout this paper, the  $i$ th coordinate  $f_i$  of  $f$  is a function  $f_i : \mathbb{R}^n \rightarrow \mathbb{R}$  and its weak gradient is denoted by  $\nabla f_i$ . Similarly, the weak derivative of  $g(\mathbf{z})$  is denoted by  $\nabla g$ . We refer to [14], Section 1.5, for definitions of weak differentiability. For the application to asymptotic normality of debiased estimates in Section 3, the functions we will consider are locally Lipschitz. By Rademacher's theorem, locally Lipschitz functions are Fréchet differentiable almost everywhere, which is stronger than the existence of directional derivatives in all directions. In this case, the weak derivatives agree with the classical partial derivatives almost everywhere. As far as the application in Section 3 is concerned, the reader unfamiliar with

weak differentiability may consider the additional assumption that  $f$  is locally Lipschitz in the following results and replace weak derivatives with classical derivatives. The variance of (2.1) is given by the following proposition.

**PROPOSITION 2.1** (Second-order Stein formula, [37], equation (8.6), [7]). *Let  $\mathbf{z} \sim N(\mathbf{0}, \mathbf{I}_n)$  and  $f : \mathbb{R}^n \rightarrow \mathbb{R}^n$  be a function with each coordinate  $f_i$  being squared integrable and weakly differentiable with squared integrable gradient, that is,  $\mathbb{E}[f_i(\mathbf{z})^2] + \mathbb{E}[\|\nabla f_i(\mathbf{z})\|^2] < +\infty$ . Then*

$$(2.2) \quad \mathbb{E}[(\mathbf{z}^\top f(\mathbf{z}) - \operatorname{div} f(\mathbf{z}))^2] = \mathbb{E}[\|f(\mathbf{z})\|^2] + \mathbb{E}[\operatorname{trace}\{\nabla f(\mathbf{z})\}^2].$$

The above result, in the twice differentiable case, was known to Stein [37], equation (8.6). If  $f$  is twice differentiable, the result follows by a sequence of integration by parts. The differentiability requirement was relaxed to only once weakly differentiable  $f$  in [7] where statistical applications of this formula to such once differentiable  $f$  are discussed.

**2.1. Linear approximation.** The goal of the present section is to derive normal approximations and CLT for the random variable (2.1). The intuition is as follows. We are looking for linear approximation of the random variable (2.1), of the form  $\mathbf{z}^\top \boldsymbol{\mu} \sim N(0, \|\boldsymbol{\mu}\|^2)$  for some deterministic  $\boldsymbol{\mu} \in \mathbb{R}^n$ . We rewrite (2.1) as

$$(2.3) \quad \mathbf{z}^\top f(\mathbf{z}) - \operatorname{div} f(\mathbf{z}) = \underbrace{\mathbf{z}^\top \boldsymbol{\mu}}_{\text{linear part}} + \underbrace{\mathbf{z}^\top (f(\mathbf{z}) - \boldsymbol{\mu}) - \operatorname{div} f(\mathbf{z})}_{\text{remainder}}.$$

The remainder term above is mean zero with second moment equal to  $\mathbb{E}[\|f(\mathbf{z}) - \boldsymbol{\mu}\|^2] + \mathbb{E}[\operatorname{trace}\{\nabla f(\mathbf{z})\}^2]$  by Proposition 2.1. This second moment is minimized for  $\boldsymbol{\mu} = \mathbb{E}[f(\mathbf{z})]$ , hence  $\mathbf{z}^\top \mathbb{E}[f(\mathbf{z})]$  gives the best linear approximation of  $\xi$  in (2.1). The following result provides conditions on  $f$  under which the remainder term is negligible in (2.3).

**THEOREM 2.2.** *Let  $\mathbf{z} \sim N(\mathbf{0}, \mathbf{I}_n)$  and  $f$  be a function  $f : \mathbb{R}^n \rightarrow \mathbb{R}^n$ , with each coordinate  $f_i$  being squared integrable and weakly differentiable with squared integrable gradient, that is,  $\mathbb{E}[f_i(\mathbf{z})^2] + \mathbb{E}[\|\nabla f_i(\mathbf{z})\|^2] < +\infty$ . Then  $\xi = \mathbf{z}^\top f(\mathbf{z}) - \operatorname{div} f(\mathbf{z})$  satisfies*

$$(2.4) \quad \mathbb{E}[(\xi / \operatorname{Var}[\xi]^{1/2} - Z)^2] = \epsilon_1^2 + (1 - (1 - \epsilon_1^2)^{1/2})^2 = \epsilon_1^2 + c_1 \epsilon_1^4$$

with  $Z = \mathbf{z}^\top \mathbb{E}[f(\mathbf{z})] / \|\mathbb{E}[f(\mathbf{z})]\| \sim N(0, 1)$ , deterministic real  $1/4 \leq c_1 \leq 1$  and

$$(2.5) \quad \epsilon_1^2 \stackrel{\text{def}}{=} 1 - \frac{\|\mathbb{E}[f(\mathbf{z})]\|^2}{\operatorname{Var}[\xi]} \leq \bar{\epsilon}_1^2 \leq \frac{2\mathbb{E}[\|\nabla f(\mathbf{z})\|_F^2]}{\mathbb{E}[\|f(\mathbf{z})\|^2] + \mathbb{E}[\|\nabla f(\mathbf{z})\|_F^2]},$$

where  $\bar{\epsilon}_1^2 \stackrel{\text{def}}{=} 2\mathbb{E}[\|\{\nabla f(\mathbf{z})\}^s\|_F^2] / \{\|\mathbb{E}[f(\mathbf{z})]\|^2 + 2\mathbb{E}[\|\{\nabla f(\mathbf{z})\}^s\|_F^2]\}$ . Consequently,  $\sup_{t \in \mathbb{R}} |\mathbb{P}(\xi / \operatorname{Var}[\xi]^{1/2} \leq t) - \mathbb{P}(Z \leq t)| \leq C(\epsilon_1^2 + c_1 \epsilon_1^4)^{1/3}$  for  $C = 1 + (2\pi)^{-1/2}$ .

A direct consequence of Theorem 2.2 is  $\epsilon_1^2 \leq (2.4) \leq 2\epsilon_1^2 \leq 2\bar{\epsilon}_1^2$ . Inequality (2.4) provides an upper bound on the 2-Wasserstein distance between  $\xi / \operatorname{Var}[\xi]^{1/2}$  and  $Z \sim N(0, 1)$ . When  $\epsilon_1^2 \rightarrow 0$ , it gives a stronger  $L_2$  form of the CLT  $\xi / \operatorname{Var}[\xi]^{1/2} \xrightarrow{d} N(0, 1)$  in addition to the Kolmogorov distance bound in Theorem 2.2. The theorem follows from Proposition 2.1 and an application of the Gaussian Poincaré inequality.

**PROOF OF THEOREM 2.2.** Define  $Z = \mathbf{z}^\top \mathbb{E}[f(\mathbf{z})] / \|\mathbb{E}[f(\mathbf{z})]\|$  then  $Z \sim N(0, 1)$  and

$$\xi - \operatorname{Var}[\xi]^{1/2} Z = \mathbf{z}^\top g(\mathbf{z}) - \operatorname{div} g(\mathbf{z}),$$

where  $g(\mathbf{z}) = f(\mathbf{z}) - r\mathbb{E}f(\mathbf{z})$  and  $r = (\text{Var}[\xi]^{1/2}/\|\mathbb{E}f(\mathbf{z})\|)$ . By Proposition 2.1 applied to  $g$  and a bias-variance decomposition,

$$\begin{aligned} & \mathbb{E}[(\xi - \text{Var}[\xi]^{1/2}Z)^2] \\ &= \mathbb{E}\|f(\mathbf{z}) - r\mathbb{E}[f(\mathbf{z})]\|^2 + \mathbb{E}\text{trace}[\{\nabla f(\mathbf{z})\}^2] \\ &= \mathbb{E}\|f(\mathbf{z}) - \mathbb{E}[f(\mathbf{z})]\|^2 + \mathbb{E}\text{trace}[\{\nabla f(\mathbf{z})\}^2] + \{\text{Var}[\xi]^{1/2} - \|\mathbb{E}f(\mathbf{z})\|\}^2 \\ &= \text{Var}[\xi] - \|\mathbb{E}f(\mathbf{z})\|^2 + \{\text{Var}[\xi]^{1/2} - \|\mathbb{E}f(\mathbf{z})\|\}^2 \end{aligned}$$

thanks to  $(r - 1)\|\mathbb{E}f(\mathbf{z})\| = \text{Var}[\xi]^{1/2} - \|\mathbb{E}f(\mathbf{z})\|$ . Thus, (2.4) follows from the definition of  $\epsilon_1^2$  in (2.5). Moreover,  $\mathbb{E}[\|f(\mathbf{z}) - \mathbb{E}[f(\mathbf{z})]\|^2] \leq \mathbb{E}[\|\nabla f(\mathbf{z})\|_F^2]$  by the Gaussian Poincaré inequality and  $\|\mathbf{M}\|_F^2 + \text{trace}(\mathbf{M}^2) = 2\|\mathbf{M}^s\|_F^2$  for  $\mathbf{M} \in \mathbb{R}^{n \times n}$ . Hence, with  $a = \|\mathbb{E}f(\mathbf{z})\|^2$ ,  $b = \mathbb{E}\|f(\mathbf{z}) - \mathbb{E}f(\mathbf{z})\|^2$ ,  $c = \mathbb{E}\text{trace}[\{\nabla f(\mathbf{z})\}^2]$  and  $d = \mathbb{E}[\|\{\nabla f(\mathbf{z})\}^s\|_F^2]$  we have

$$\epsilon_1^2 = \frac{b + c}{a + b + c} \leq \frac{2d}{a + 2d} = \bar{\epsilon}_1^2 \leq \frac{2\mathbb{E}[\|\nabla f(\mathbf{z})\|_F^2]}{a + 2\mathbb{E}[\|\nabla f(\mathbf{z})\|_F^2]} \leq \frac{2\mathbb{E}[\|\nabla f(\mathbf{z})\|_F^2]}{\mathbb{E}[\|f(\mathbf{z})\|^2 + \|\nabla f(\mathbf{z})\|_F^2]}$$

thanks to another Gaussian Poincaré inequality for the last inequality. Finally,  $x^2/4 \leq (1 - \sqrt{1 - x})^2 \leq x^2$  holds for all  $x \in [0, 1]$ , which proves  $c_1 \in [1/4, 1]$ .

For any  $\delta > 0$ , by Markov's inequality  $\mathbb{P}(\xi/\text{Var}[\xi]^{1/2} \leq t) - \mathbb{P}(Z \leq t) \leq \mathbb{P}(|\xi/\text{Var}[\xi]^{1/2} - Z| > \delta) + \mathbb{P}(Z \in [t, t + \delta]) \leq (\epsilon_1^2 + c_1\epsilon_1^4)/\delta^2 + \delta(2\pi)^{-1/2}$  since the standard normal pdf is uniformly bounded by  $(2\pi)^{-1/2}$ . Hence, with  $\delta = (\epsilon_1^2 + c_1\epsilon_1^4)^{1/3}$ , the above and a similar argument on  $[t - \delta, t]$  provide the Kolmogorov distance bound.  $\square$

Normal approximation results such as Theorem 2.2 are flexible tools as they let us derive asymptotic normality results by mechanically computing gradients: By Theorem 2.2, it suffices to show that the expectation of  $\|\nabla f(\mathbf{z})\|_F^2$  is negligible compared with that of  $\|f(\mathbf{z})\|^2$  to obtain  $\xi/\text{Var}[\xi]^{1/2} \xrightarrow{d} N(0, 1)$ . Normal approximations involving derivatives have been studied for random variables with the more general form  $W = g(\mathbf{z})$  for differentiable functions  $g: \mathbb{R}^n \rightarrow \mathbb{R}$ . The second-order Poincaré inequality of [19] bounds the total variation distance  $d_{\text{TV}}$  of  $g(\mathbf{z})$  to the Gaussian distribution using the first and second derivatives of  $g$ : [19], Theorem 2.2, specialized to  $W = g(\mathbf{z})$  with  $\mathbf{z} \sim N(\mathbf{0}, \mathbf{I}_n)$  states that

$$(2.6) \quad d_{\text{TV}}\{W, N(\mu_0, \sigma_0^2)\} \leq (2\sqrt{5}/\sigma_0^2)\mathbb{E}[\|\nabla g(\mathbf{z})\|^4]^{1/4}\mathbb{E}[\|\nabla^2 g(\mathbf{z})\|_{\text{op}}^4]^{1/4},$$

where  $W = g(\mathbf{z})$ ,  $\mathbf{z} \sim N(\mathbf{0}, \mathbf{I}_n)$ ,  $\mu_0 = \mathbb{E}[W]$  and  $\sigma_0^2 = \text{Var}[W]$ . Above,  $\nabla g$ ,  $\nabla^2 g$  denote the gradient and Hessian matrix of  $g$ . Inequality (2.6) provides a CLT for  $g(\mathbf{z})$  provided that the moments of the derivatives  $\mathbb{E}[\|\nabla g(\mathbf{z})\|^4]^{1/4}$  and  $\mathbb{E}[\|\nabla^2 g(\mathbf{z})\|_{\text{op}}^4]^{1/4}$  are negligible compared to the variance  $\sigma_0^2 = \text{Var}[g(\mathbf{z})]$ . Inequality (2.6) has been successfully applied to derive asymptotic normality of unregularized  $M$ -estimators when  $p/n \rightarrow \gamma < 1$  and the  $M$ -estimation loss is twice differentiable [31]. However, the (2.6) based approach is not applicable for regularized estimators such as the Lasso and group Lasso that are only once differentiable functions of  $(\mathbf{X}, \mathbf{y})$ . In fact, by Proposition 4.1 below, the Lasso is not twice differentiable as  $\text{trace}[(\partial/\partial \mathbf{y})\mathbf{X}\hat{\boldsymbol{\beta}}(\mathbf{y}, \mathbf{X})]$  is integer-valued. In Theorem 2.2, while  $\xi = \mathbf{z}^\top f(\mathbf{z}) - \text{div} f(\mathbf{z})$  already involves the derivatives of  $f$  through the divergence, the ratio  $\bar{\epsilon}_1^2$  that appears in the upper bound (2.5) only involves  $f$  and its gradient  $\nabla f$ ; the second derivatives of  $f$  need not exist. Section 3 uses Theorem 2.2 to provide asymptotic normality for debiasing estimators that are only once differentiable.

*Variance estimate.* It follows from Theorem 2.2 that random variables  $\xi$  of the form (2.1) are asymptotically normal under the condition  $1 - \|\mathbb{E}[f(\mathbf{z})]\|^2 / \text{Var}[\xi] \rightarrow 0$ , or under a somewhat stronger but more explicit condition  $\bar{\epsilon}_1^2 \rightarrow 0$  as in (2.5). The following theorem provides consistent estimates of  $\text{Var}[\xi]$ .

**THEOREM 2.3.** *Let  $f, \mathbf{z}, \xi, \epsilon_1^2$  and  $c_1 \in [1/4, 1]$  be as in Theorem 2.2. Then*

$$(2.7) \quad \begin{aligned} \mathbb{E}[(\|f(\mathbf{z})\| / \text{Var}[\xi]^{1/2} - 1)^2] &\leq \epsilon_1^2 - \mathbb{E}[\text{trace}(\{\nabla f(\mathbf{z})\}^2)] / \text{Var}[\xi] + c_1 \epsilon_1^4 \\ &\leq (1 - \epsilon_1^2) \bar{\epsilon}_1^2 / (2 - 2\bar{\epsilon}_1^2) + c_1 \epsilon_1^4 \end{aligned}$$

with  $\bar{\epsilon}_1^2 \stackrel{\text{def}}{=} 2\mathbb{E}[\|\nabla f(\mathbf{z})\|_F^2] / (\|\mathbb{E}[f(\mathbf{z})]\|^2 + 2\mathbb{E}[\|\nabla f(\mathbf{z})\|_F^2]) \geq \epsilon_1^2$ . Consequently,

$$(2.8) \quad \|f(\mathbf{z})\|^2 / \text{Var}[\xi] \rightarrow^{\mathbb{P}} 1 \quad \text{and} \quad \xi / \|f(\mathbf{z})\| \rightarrow^d N(0, 1)$$

when  $\epsilon_1^2 + \bar{\epsilon}_1^2 I\{\mathbb{E}[\text{trace}(\{\nabla f(\mathbf{z})\}^2)] < 0\} \rightarrow 0$ .

**PROOF OF THEOREM 2.3.** It follows from the Jensen inequality and (2.4) that

$$\begin{aligned} \mathbb{E}[(\|f(\mathbf{z})\| / \text{Var}[\xi]^{1/2} - 1)^2] &\leq \mathbb{E}[\|f(\mathbf{z})\|^2] / \text{Var}[\xi] + 1 - 2\|\bar{\mu}\| / \text{Var}[\xi]^{1/2} \\ &= \epsilon_1^2 - \mathbb{E}[\text{trace}(\{\nabla f(\mathbf{z})\}^2)] / \text{Var}[\xi] + c_1 \epsilon_1^4 \end{aligned}$$

with  $\bar{\mu} = \mathbb{E}[\nabla f(\mathbf{z})]$  due to  $\epsilon_1^2 = 1 - \|\bar{\mu}\|^2 / \text{Var}[\xi]$ . For the second inequality in (2.7),

$$\epsilon_1^2 - \mathbb{E}[\text{trace}(\{\nabla f(\mathbf{z})\}^2)] / \text{Var}[\xi] = \mathbb{E}[\|f(\mathbf{z}) - \bar{\mu}\|^2] / \text{Var}[\xi] \leq \mathbb{E}[\|\nabla f(\mathbf{z})\|_F^2] / \text{Var}[\xi],$$

thanks to the Gaussian Poincaré inequality, and  $(1 - \epsilon_1^2) \bar{\epsilon}_1^2 / (2 - 2\bar{\epsilon}_1^2)$  equals to the right-hand side above by the definition of  $\bar{\epsilon}_1^2$ .  $\square$

**2.2. Quadratic approximation.** The decomposition (2.3) is especially useful if the linear part  $\mathbf{z}^\top \bar{\mu}$  with  $\bar{\mu} = \mathbb{E}[f(\mathbf{z})]$  is a good approximation for  $\xi = \mathbf{z}^\top \nabla f(\mathbf{z}) - \text{div } f(\mathbf{z})$ . In some cases, for example, if  $f(\mathbf{z}) = \mathbf{A}\mathbf{z}$  for some square deterministic matrix  $\mathbf{A}$ , the decomposition (2.3) is uninformative. It is then natural to look for the best quadratic approximation of  $\xi$  in the sense of the  $L_2$  orthogonal projection to

$$\mathcal{H}_{1,2} = \{\xi_{\mu, \mathbf{A}} = \mu^\top \mathbf{z} + \mathbf{z}^\top \mathbf{A}\mathbf{z} - \text{trace}[\mathbf{A}] : \mu \in \mathbb{R}^n, \mathbf{A} \in \mathbb{R}^{n \times n}\} = \mathcal{H}_1 \oplus \mathcal{H}_2,$$

where  $\mathcal{H}_1 = \{\mu^\top \mathbf{z} : \mu \in \mathbb{R}^n\}$  and  $\mathcal{H}_2 = \{\mathbf{z}^\top \mathbf{A}\mathbf{z} - \text{trace}[\mathbf{A}] : \mathbf{A} \in \mathbb{R}^{n \times n}\}$  are  $L_2$  subspaces orthogonal to each other.

The calculation in (2.4) for  $\mathcal{H}_1$  is generic in the following sense. If  $\bar{\xi}$  is the  $L_2$  projection of a random variable  $\xi$  in  $L_2$ , then the sine of the  $L_2$ -angle between  $\bar{\xi}$  and  $\xi$  is  $\epsilon = (\mathbb{E}[(\xi - \bar{\xi})^2] / \mathbb{E}[\xi^2])^{1/2} = (1 - \mathbb{E}[\bar{\xi}^2] / \mathbb{E}[\xi^2])^{1/2}$  and

$$(2.9) \quad \mathbb{E}[(\xi / \text{Var}[\xi]^{1/2} - \bar{\xi} / \text{Var}[\bar{\xi}]^{1/2})^2] = 2(1 - \sqrt{1 - \epsilon^2}) = \epsilon^2 + c\epsilon^4$$

holds for some deterministic real  $1/4 \leq c \leq 1$ . Indeed, take  $\epsilon = \sin \alpha$  with  $\alpha$  being the  $L_2$ -angle between  $\xi$  and  $\bar{\xi}$ , so that (2.9) becomes  $(2 \sin(\alpha/2))^2 = 2(1 - \cos(\alpha)) = \epsilon^2 + c\epsilon^4$  as in the proof of Theorem 2.2.

The next result extends Theorem 2.2 to the  $L_2$  quadratic projections to  $\mathcal{H}_2$  and  $\mathcal{H}_{1,2}$ , and also gives Theorem 2.2 the interpretation as the  $L_2$  projection to  $\mathcal{H}_1$ .

**THEOREM 2.4.** *Let  $\mathbf{z} \sim N(\mathbf{0}, \mathbf{I}_n)$ ,  $f: \mathbb{R}^n \rightarrow \mathbb{R}^n$  satisfy the assumption of Theorem 2.2, and  $\xi = \mathbf{z}^\top f(\mathbf{z}) - \operatorname{div} f(\mathbf{z})$ . For  $\boldsymbol{\mu} \in \mathbb{R}^n$  and  $\mathbf{A} \in \mathbb{R}^{n \times n}$  let  $\xi_{\boldsymbol{\mu}, \mathbf{A}} = \mathbf{z}^\top (\boldsymbol{\mu} + \mathbf{A}\mathbf{z}) - \operatorname{trace} \mathbf{A}$ . Let  $\bar{\boldsymbol{\mu}} = \mathbb{E}[f(\mathbf{z})]$  and  $\bar{\mathbf{A}} = \mathbb{E}[\nabla f(\mathbf{z})]$ . Then  $\xi_{\bar{\boldsymbol{\mu}}, \bar{\mathbf{A}}}$  is the  $L_2$  projection of  $\xi$  to  $\mathcal{H}_{1,2}$  and*

$$(2.10) \quad \begin{aligned} \mathbb{E}[(\xi - \xi_{\bar{\boldsymbol{\mu}}, \bar{\mathbf{A}}})^2] &= \mathbb{E}[\|f(\mathbf{z}) - \bar{\boldsymbol{\mu}}\|^2 - \|\bar{\mathbf{A}}\|_F^2] + \mathbb{E} \operatorname{trace}[\{\nabla f(\mathbf{z}) - \bar{\mathbf{A}}\}^2] \\ &\leq 2\mathbb{E}[\|\nabla f(\mathbf{z}) - \bar{\mathbf{A}}\|_F^2]. \end{aligned}$$

Consequently,  $\xi_{\bar{\boldsymbol{\mu}}, \mathbf{0}}$  is the projection of  $\xi$  and  $\xi_{\bar{\boldsymbol{\mu}}, \bar{\mathbf{A}}}$  to  $\mathcal{H}_1$  with  $\mathbb{E}[(\xi_{\bar{\boldsymbol{\mu}}, \bar{\mathbf{A}}} - \xi_{\bar{\boldsymbol{\mu}}, \mathbf{0}})^2] = 2\|\bar{\mathbf{A}}\|_F^2$  and  $\xi_{\mathbf{0}, \bar{\mathbf{A}}}$  is the projection of  $\xi$  and  $\xi_{\bar{\boldsymbol{\mu}}, \bar{\mathbf{A}}}$  to  $\mathcal{H}_2$  with  $\mathbb{E}[(\xi_{\bar{\boldsymbol{\mu}}, \bar{\mathbf{A}}} - \xi_{\mathbf{0}, \bar{\mathbf{A}}})^2] = \|\bar{\boldsymbol{\mu}}\|^2$ .

For the projection  $\xi_{\bar{\boldsymbol{\mu}}, \bar{\mathbf{A}}}$  of  $\xi$  to  $\mathcal{H}_{1,2}$ ,  $\epsilon_{1,2}^2 \stackrel{\text{def}}{=} 1 - \mathbb{E}[\xi_{\bar{\boldsymbol{\mu}}, \bar{\mathbf{A}}}^2]/\mathbb{E}[\xi^2]$  satisfies  $\epsilon_{1,2}^2 \leq \bar{\epsilon}_{1,2}^2 \stackrel{\text{def}}{=} 2\mathbb{E}[\|\nabla f(\mathbf{z}) - \bar{\mathbf{A}}\|_F^2]/\{\|\bar{\boldsymbol{\mu}}\|^2 + 2\mathbb{E}[\|\nabla f(\mathbf{z})\|_F^2]\}$  and under the condition  $\epsilon_{1,2}^2 = o(1)$ ,

$$(2.11) \quad \|\bar{\mathbf{A}}^s\|_{\text{op}}^2/(\|\bar{\boldsymbol{\mu}}\|_2^2 + \|\bar{\mathbf{A}}^s\|_F^2) \rightarrow 0 \quad \Leftrightarrow \quad \xi/\operatorname{Var}[\xi]^{1/2} \rightarrow^d N(0, 1).$$

For the projection  $\xi_{\mathbf{0}, \bar{\mathbf{A}}}$  of  $\xi$  to  $\mathcal{H}_2$ ,  $\epsilon_2^2 = 1 - \mathbb{E}[\xi_{\mathbf{0}, \bar{\mathbf{A}}}^2]/\mathbb{E}[\xi^2]$  satisfies  $\epsilon_2^2 \leq \bar{\epsilon}_2^2 \stackrel{\text{def}}{=} \{\|\bar{\boldsymbol{\mu}}\|^2 + 2\mathbb{E}[\|\nabla f(\mathbf{z}) - \bar{\mathbf{A}}\|_F^2]\}/\{\|\bar{\boldsymbol{\mu}}\|^2 + 2\mathbb{E}[\|\nabla f(\mathbf{z})\|_F^2]\}$  and under the condition  $\epsilon_2^2 = o(1)$ ,

$$(2.12) \quad \|\bar{\mathbf{A}}^s\|_{\text{op}}^2/\|\bar{\mathbf{A}}^s\|_F^2 \rightarrow 0 \quad \Leftrightarrow \quad \xi/\operatorname{Var}[\xi]^{1/2} \rightarrow^d N(0, 1).$$

**PROOF OF THEOREM 2.4.** The function  $g(\mathbf{z}) = f(\mathbf{z}) - \boldsymbol{\mu} - \mathbf{A}^\top \mathbf{z}$  has gradient  $\nabla g = \nabla f - \mathbf{A}$ . Application of the second-order Stein's formula in Proposition 2.1 to  $g$  yields

$$\mathbb{E}[(\xi - \xi_{\boldsymbol{\mu}, \mathbf{A}})^2] = \mathbb{E}[\|f(\mathbf{z}) - \boldsymbol{\mu} - \mathbf{A}^\top \mathbf{z}\|^2] + \mathbb{E} \operatorname{trace}[\{\nabla f(\mathbf{z}) - \mathbf{A}\}^2] \stackrel{\text{def}}{=} I + II.$$

The first term is  $I = \mathbb{E}[\|f(\mathbf{z}) - \boldsymbol{\mu}\|^2] + \|\mathbf{A}\|_F^2 - 2\mathbb{E}[\mathbf{z}^\top \mathbf{A}(f(\mathbf{z}) - \boldsymbol{\mu})]$ . By Stein's formula and the linearity of the trace, we have

$$\begin{aligned} \|\mathbf{A}\|_F^2 - 2\mathbb{E}[\mathbf{z}^\top \mathbf{A}(f(\mathbf{z}) - \boldsymbol{\mu})] &= \|\mathbf{A}\|_F^2 - 2\mathbb{E} \operatorname{trace}(\nabla f(\mathbf{z}) \mathbf{A}^\top) \\ &= \|\mathbf{A}\|_F^2 - 2\operatorname{trace}[\mathbf{A}^\top \bar{\mathbf{A}}] \\ &= -\|\bar{\mathbf{A}}\|_F^2 + \|\mathbf{A} - \bar{\mathbf{A}}\|_F^2. \end{aligned}$$

We also have  $\mathbb{E}[\|\nabla f(\mathbf{z}) - \bar{\mathbf{A}}\|_F^2] = \mathbb{E}[\|\nabla f(\mathbf{z})\|_F^2] - \|\bar{\mathbf{A}}\|_F^2$  so that

$$I = \mathbb{E}[\|f(\mathbf{z}) - \boldsymbol{\mu}\|^2 - \|\nabla f(\mathbf{z})\|_F^2] + \mathbb{E}[\|\nabla f(\mathbf{z}) - \bar{\mathbf{A}}\|_F^2] + \|\mathbf{A} - \bar{\mathbf{A}}\|_F^2.$$

For the second term, using that  $\mathbb{E}[\nabla f(\mathbf{z}) - \bar{\mathbf{A}}] = 0$  we get

$$II = \mathbb{E} \operatorname{trace}[\{\nabla f(\mathbf{z}) - \mathbf{A}\}^2] = \mathbb{E} \operatorname{trace}[\{\nabla f(\mathbf{z}) - \bar{\mathbf{A}}\}^2] + \operatorname{trace}[(\bar{\mathbf{A}} - \mathbf{A})^2].$$

Due to  $\|\mathbf{M}\|_F^2 + \operatorname{trace}(\mathbf{M}^2) = 2\|\mathbf{M}^s\|_F^2$  for  $\mathbf{M} \in \mathbb{R}^{n \times n}$ , it follows that

$$\begin{aligned} \mathbb{E}[(\xi - \xi_{\boldsymbol{\mu}, \mathbf{A}})^2] &= \mathbb{E}[\|f(\mathbf{z}) - \bar{\boldsymbol{\mu}}\|^2 - \|\bar{\mathbf{A}}\|_F^2] + \mathbb{E} \operatorname{trace}[\{\nabla f(\mathbf{z}) - \bar{\mathbf{A}}\}^2] \\ &\quad + \|\boldsymbol{\mu} - \bar{\boldsymbol{\mu}}\|^2 + 2\|(\mathbf{A} - \bar{\mathbf{A}})^s\|_F^2 \end{aligned}$$

The optimality of  $\boldsymbol{\mu} = \bar{\boldsymbol{\mu}}$  and  $\mathbf{A} = \bar{\mathbf{A}}$  follows, so that  $\xi_{\bar{\boldsymbol{\mu}}, \bar{\mathbf{A}}}$  is the  $L_2$  projection of  $\xi$  to  $\mathcal{H}_{1,2}$ . Also, the first line above gives the formula of  $\mathbb{E}[(\xi - \xi_{\bar{\boldsymbol{\mu}}, \bar{\mathbf{A}}})^2]$  in (2.10), and the second line gives the formulas for the variances of  $\xi_{\bar{\boldsymbol{\mu}}, \mathbf{0}}$  and  $\xi_{\mathbf{0}, \bar{\mathbf{A}}}$ . The upper bound in (2.10) follows from  $\mathbb{E}[\|f(\mathbf{z}) - \bar{\boldsymbol{\mu}}\|^2 - \|\bar{\mathbf{A}}\|_F^2] \leq \mathbb{E}[\|\nabla f(\mathbf{z})\|_F^2] - \|\bar{\mathbf{A}}\|_F^2 = \mathbb{E}[\|\nabla f(\mathbf{z}) - \bar{\mathbf{A}}\|_F^2]$  thanks to the Gaussian Poincaré inequality. Inequality (2.10) is equivalent to

$$\mathbb{E}[\xi^2] = \mathbb{E}[\xi_{\bar{\boldsymbol{\mu}}, \bar{\mathbf{A}}}^2] + \mathbb{E}[(\xi - \xi_{\bar{\boldsymbol{\mu}}, \bar{\mathbf{A}}})^2] \leq \|\bar{\boldsymbol{\mu}}\|^2 + 2\|\bar{\mathbf{A}}^s\|_F^2 + 2\mathbb{E}[\|\nabla f(\mathbf{z}) - \bar{\mathbf{A}}\|_F^2],$$

which provides  $\epsilon_{1,2}^2 \leq \bar{\epsilon}_{1,2}^2$  and  $\epsilon_2^2 \leq \bar{\epsilon}_2^2$  by bounding from above the denominator in  $\epsilon_{1,2}^2 = 1 - \mathbb{E}[\xi_{\bar{\mu}, \bar{A}}^2] / \mathbb{E}[\xi^2]$  and  $\epsilon_2^2 = 1 - \mathbb{E}[\xi_{0, \bar{A}}^2] / \mathbb{E}[\xi^2]$ .

For (2.11), we write  $\xi_{\bar{\mu}, \bar{A}} = \sum_{j=1}^n \{a_j G_j + b_j (G_j^2 - 1)\}$  with i.i.d.  $G_j \sim N(0, 1)$ , where  $a_j = \mathbf{u}_j^\top \bar{\mu}$  and  $G_j = \mathbf{u}_j^\top \mathbf{z}$  with the eigenvalue decomposition  $\bar{\mathbf{A}}^s = \sum_{j=1}^n b_j \mathbf{u}_j \mathbf{u}_j^\top$ . Assume without loss of generality that  $\text{Var}(a_j G_j + b_j (G_j^2 - 1)) = a_j^2 + 2b_j^2$  is nonincreasing in  $j$  and that  $\text{Var}[\xi_{\bar{\mu}, \bar{A}}]$  satisfies  $\sum_{j=1}^n (a_j^2 + 2b_j^2) = 1$ . The condition on the left-hand side of (2.11) implies that the integer  $k_n \stackrel{\text{def}}{=} \lceil \|\bar{\mathbf{A}}^s\|_{\text{op}}^{-1} \rceil = \lceil (\max_j b_j)^{-1} \rceil$  satisfies  $k_n \rightarrow +\infty$  and  $\sum_{j=1}^{k_n} b_j^2 \ll 1 = \text{Var}[\xi_{\bar{\mu}, \bar{A}}]$ , so that

$$\xi_{\bar{\mu}, \bar{A}} = \sum_{j=1}^{k_n} a_j G_j + \sum_{j=k_n+1}^n \{a_j G_j + b_j (G_j^2 - 1)\} + o_{\mathbb{P}}(1).$$

Assuming that  $\sum_{j=1}^{k_n} a_j^2 \rightarrow c$  for some  $c \in [0, 1]$  by extracting a subsequence if necessary,  $k_n \rightarrow +\infty$  implies  $\max_{j > k_n} (a_j^2 + 2b_j^2) = a_{k_n+1}^2 + 2b_{k_n+1}^2 \rightarrow 0$  so that the second term above is independent of the first and approximately  $N(0, 1 - c)$  by the Lyapunov CLT when  $\sum_{j=1}^{k_n} a_j^2 \rightarrow c \leq 1$ . This proves that the LHS of (2.11) implies the RHS. Conversely, assume the asymptotic normality on the RHS so that  $\sum_{j=1}^n \{a_j G_j + b_j (G_j^2 - 1)\} \rightarrow N(0, 1)$ . Let  $W_j = a_j G_j + b_j (G_j^2 - 1)$  and  $j_n \leq n$ . As  $W_{j_n}$  is an independent component of the sum, for any  $(a_{j_n}, b_{j_n}) \rightarrow (a, b)$  along a subsequence with  $a^2 + b^2 > 0$ , we must have  $b = 0$  because  $W_{j_n} \xrightarrow{d} N(0, a^2 + 2b^2)$  by the Cramér–Lévy theorem and  $W_{j_n} \xrightarrow{d} aG + b(G^2 + 1)$  for some  $G \sim N(0, 1)$ . As  $j_n \leq n$  are arbitrary, this gives  $\|\bar{\mathbf{A}}^s\|_{\text{op}} = \max_{j=1, \dots, n} b_j^2 \rightarrow 0$ .  $\square$

*Variance estimate: Quadratic case.* Theorem 2.4 provides the quadratic normal approximation of  $\xi$  under the condition  $\bar{\epsilon}_{1,2}^2 \rightarrow 0$  with

$$(2.13) \quad \bar{\epsilon}_{1,2}^2 \geq \epsilon_{1,2}^2 = 1 - (\|\mathbb{E}[f(\mathbf{z})]\|^2 + 2\|\mathbb{E}[\nabla f(\mathbf{z})]\|_F^2) / \text{Var}[\xi],$$

where  $\bar{\epsilon}_{1,2}^2$  is defined using the upper bound  $\|\mathbb{E}[f(\mathbf{z})]\|^2 + 2\mathbb{E}[\|\nabla f(\mathbf{z})\|_F^2] \geq \text{Var}[\xi]$  established in (2.10) in the denominator on the right-hand side of (2.13).

**THEOREM 2.5.** *Let  $f, \mathbf{z}, \xi, \bar{\mu} = \mathbb{E}[f(\mathbf{z})]$  and  $\bar{\mathbf{A}} = \mathbb{E}[\nabla f(\mathbf{z})]$  be as in Theorem 2.4 and  $\widehat{\text{Var}}[\xi] = \|f(\mathbf{z})\|^2 + \text{trace}[\{\nabla f(\mathbf{z})\}^2]$ . Then*

$$(2.14) \quad \mathbb{E}[\widehat{\text{Var}}[\xi] / \text{Var}[\xi] - 1] \leq 2\bar{\epsilon}_{1,2}^2 + 2\bar{\epsilon}_{1,2} C_0 + C_0 \|\bar{\mathbf{A}}\|_{\text{op}} / \text{Var}[\xi]^{1/2}$$

with  $\bar{\epsilon}_{1,2} \stackrel{\text{def}}{=} \{2\mathbb{E}[\|\nabla f(\mathbf{z}) - \bar{\mathbf{A}}\|_F^2] / \text{Var}[\xi]\}^{1/2}$  and  $C_0 \stackrel{\text{def}}{=} \{(\|\bar{\mu}\|^2 + 2\|\bar{\mathbf{A}}\|_F^2) / \text{Var}[\xi]\}^{1/2}$ . Consequently, under the conditions  $\{\mathbb{E}[\|\nabla f(\mathbf{z}) - \bar{\mathbf{A}}\|_F^2] + \|\bar{\mathbf{A}}\|_{\text{op}}\} / (\|\bar{\mu}\|^2 + \|\bar{\mathbf{A}}^s\|_F^2) = o(1)$  and  $\|\bar{\mathbf{A}}\|_F^2 / (\|\bar{\mu}\|^2 + \|\bar{\mathbf{A}}^s\|_F^2) = O(1)$ ,

$$(2.15) \quad \widehat{\text{Var}}[\xi] / \text{Var}[\xi] \rightarrow^{\mathbb{P}} 1 \quad \text{and} \quad (\widehat{\text{Var}}[\xi])^{-1/2} \xi \rightarrow^d N(0, 1).$$

It follows from the second-order Stein formula in Proposition 2.1 that  $\widehat{\text{Var}}[\xi]$  is an unbiased estimator of  $\text{Var}[\xi]$ . Moreover, when  $\mathbb{E}[\|\nabla f(\mathbf{z}) - \bar{\mathbf{A}}\|_F^2]$ ,  $\|\bar{\mathbf{A}}\|_F$  and  $\|\bar{\mathbf{A}}\|_{\text{op}}$  are all equivalent to their symmetric counterparts,  $\mathbb{E}[\|\nabla f(\mathbf{z}) - \bar{\mathbf{A}}\|_F^2] \asymp \mathbb{E}[\|\{\nabla f(\mathbf{z}) - \bar{\mathbf{A}}\}^s\|_F^2]$ ,  $\|\bar{\mathbf{A}}\|_F \asymp \|\bar{\mathbf{A}}^s\|_F$  and  $\|\bar{\mathbf{A}}\|_{\text{op}} \asymp \|\bar{\mathbf{A}}^s\|_{\text{op}}$ , the condition for (2.15) holds if and only if  $\epsilon_{1,2}^2 + \|\bar{\mathbf{A}}^s\|_{\text{op}}^2 / (\|\bar{\mu}\|^2 + \|\bar{\mathbf{A}}^s\|_F^2) = o(1)$  for the quantities in (2.13) and (2.11). The proof is given in Appendix D.

**3. Debiasing general convex regularizers.** Our main application of the normal approximation in Theorem 2.2 concerns debiasing regularized estimators of the form

$$(3.1) \quad \hat{\boldsymbol{\beta}} = \arg \min_{\mathbf{b} \in \mathbb{R}^p} \{ \|\mathbf{y} - \mathbf{X}\mathbf{b}\|^2 / (2n) + g(\mathbf{b}) \}$$

for convex  $g : \mathbb{R}^p \rightarrow \mathbb{R}$  in the linear model (1.1). Throughout, let  $\mathbf{h} = \hat{\boldsymbol{\beta}} - \boldsymbol{\beta}$  be the error vector,  $\mathbf{a}_0 \in \mathbb{R}^p$  be a direction of interest,  $\theta = \langle \mathbf{a}_0, \boldsymbol{\beta} \rangle$  be the target of statistical inference, and  $\mathbf{u}_0, \mathbf{z}_0$  and  $\mathbf{Q}_0$  be as in (1.13) so that (1.14) holds.

**3.1. Assumption.** We say that  $g$  is  $\mu$ -strongly convex with respect to the norm  $\mathbf{b} \rightarrow \|\boldsymbol{\Sigma}^{1/2}\mathbf{b}\|$  if its symmetric Bregman divergence is bounded from below as

$$(3.2) \quad (\tilde{\mathbf{b}} - \mathbf{b})^\top ((\partial g)(\tilde{\mathbf{b}}) - (\partial g)(\mathbf{b})) \geq \mu \|\boldsymbol{\Sigma}^{1/2}(\tilde{\mathbf{b}} - \mathbf{b})\|^2$$

for some  $\mu > 0$ . Here, the interpretation of (3.2) is its validity for all choices in the sub-differential  $(\partial g)(\tilde{\mathbf{b}})$  and  $(\partial g)(\mathbf{b})$ . Condition (3.2) holds for any convex  $g$  for  $\mu = 0$ . If  $g$  is twice differentiable, (3.2) holds if and only if  $\mu\boldsymbol{\Sigma}$  is a lower bound for the Hessian of  $g$ . However, (3.2) may also hold for nondifferentiable  $g$ , for example, the elastic-net penalty with  $\boldsymbol{\Sigma} = \mathbf{I}_p$ . Our results require the following assumption.

**ASSUMPTION 3.1.** (i) Let  $\gamma > 0$ ,  $\mu \in [0, \frac{1}{2}]$  be constants such that  $\mu + (1 - \gamma)_+ > 0$ , that is, either  $\mu > 0$  or  $\gamma < 1$  must hold. Consider a sequence of regression problems (1.1) with  $n, p \rightarrow +\infty$  and  $p/n \leq \gamma$ . The penalty  $g : \mathbb{R}^p \rightarrow \mathbb{R}$  in (3.1) is convex and (3.2) holds. The rows of  $\mathbf{X}$  are i.i.d.  $N(\mathbf{0}, \boldsymbol{\Sigma})$  with invertible  $\boldsymbol{\Sigma}$  and the noise  $\boldsymbol{\varepsilon} \sim N(\mathbf{0}, \sigma^2 \mathbf{I}_n)$  is independent of  $\mathbf{X}$ . (ii)  $\mathbf{a}_0 \in \mathbb{R}^p$  is a sequence of vectors normalized with  $\|\boldsymbol{\Sigma}^{-1/2}\mathbf{a}_0\| = 1$ .

Note that if (3.2) holds for  $\mu \geq 0$  it also holds for  $\mu' = \min(\frac{1}{2}, \mu)$  and we may thus assume  $\mu \in [0, \frac{1}{2}]$  without loss of generality. Strongly convex objective functions admit unique minimizers. Since  $\gamma < 1$  implies  $\mathbb{P}(\phi_{\min}(\boldsymbol{\Sigma}^{-1/2}\mathbf{X}^\top \mathbf{X} \boldsymbol{\Sigma}^{-1/2}) > 0) = 1$  (cf. Appendix A) and the objective function of the optimization problem (3.1) is  $(\phi_{\min}(\boldsymbol{\Sigma}^{-1/2}\mathbf{X}^\top \mathbf{X} \boldsymbol{\Sigma}^{-1/2}/n) + \mu)$ -strongly convex, Assumption 3.1 grants almost surely the existence and uniqueness of the minimizer (3.1).

**3.2. Gradient with respect to  $\mathbf{y}$  and effective degrees-of-freedom.** Consider a penalized estimator (3.1) viewed as a function  $\hat{\boldsymbol{\beta}} = \hat{\boldsymbol{\beta}}(\mathbf{y}, \mathbf{X})$ . For every  $\mathbf{X} \in \mathbb{R}^{n \times p}$ , the map  $\mathbf{y} \mapsto \mathbf{X}\hat{\boldsymbol{\beta}}(\mathbf{y}, \mathbf{X})$  is 1-Lipschitz (cf. Proposition 7.3). By Rademacher's theorem, for almost every  $\mathbf{y}$  there exists a unique matrix  $\hat{\mathbf{H}} \in \mathbb{R}^{n \times n}$  such that

$$(3.3) \quad \mathbf{X}\hat{\boldsymbol{\beta}}(\mathbf{y} + \boldsymbol{\eta}, \mathbf{X}) = \mathbf{X}\hat{\boldsymbol{\beta}}(\mathbf{y}, \mathbf{X}) + \hat{\mathbf{H}}^\top \boldsymbol{\eta} + o(\|\boldsymbol{\eta}\|),$$

as in (1.7), that is,  $\hat{\mathbf{H}}$  is the gradient of the map  $\mathbf{y} \mapsto \mathbf{X}\hat{\boldsymbol{\beta}}(\mathbf{y}, \mathbf{X})$ . Furthermore,  $\hat{\mathbf{H}}$  is symmetric with eigenvalues in  $[0, 1]$ ; see Proposition 7.3 for the existence of  $\hat{\mathbf{H}}$  and its properties. While existence of  $\hat{\mathbf{H}}$  was assumed in (1.7) in the Introduction, for penalized estimators (3.1) the matrix  $\hat{\mathbf{H}}$  provably exists for almost every  $\mathbf{y}$  by Proposition 7.3.

Table 1 provides closed-form expressions of  $\hat{\mathbf{H}}$  for specific penalty functions  $g$ . The proofs of these closed-form expressions will be given in Section 4. An advantage of defining  $\hat{\mathbf{H}}$  as the Fréchet derivative of the Lipschitz map  $\mathbf{y} \mapsto \mathbf{X}\hat{\boldsymbol{\beta}}(\mathbf{y}, \mathbf{X})$  is that this definition applies to any convex penalty  $g$ , even though for arbitrary penalty  $g$  we are unable to provide a closed-form expression for  $\hat{\mathbf{H}}$ . Finally, define the effective degrees-of-freedom  $\hat{\text{df}}$  of  $\hat{\boldsymbol{\beta}}$  by

$$(3.4) \quad \hat{\text{df}} = \text{trace}[\hat{\mathbf{H}}]$$

as in (1.10). Because  $\hat{\mathbf{H}}$  is symmetric with eigenvalues in  $[0, 1]$  (cf. Proposition 7.3),  $0 \leq \hat{\text{df}} \leq n$  holds almost surely. The matrix  $\hat{\mathbf{H}}$  and the scalar  $\hat{\text{df}}$  play a major role in our analysis.

TABLE 1

Closed-form expressions  $\hat{\mathbf{H}}$  from equation (3.3) for specific convex penalty functions  $g : \mathbb{R}^p \rightarrow \mathbb{R}$ . For the Lasso and elastic-net,  $\hat{\mathbf{S}} = \{j \in [p] : \hat{\beta}_j \neq 0\}$ . For the group Lasso,  $\hat{\mathbf{S}}$  and  $\mathbf{M}$  are given in Section 4.3

| Penalty                                                                                                              | $\hat{\mathbf{H}} \in \mathbb{R}^{n \times n}$                                                                                                                                   | Justification                   |
|----------------------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------|
| $g(\mathbf{b}) = \lambda \ \mathbf{b}\ _1$ (Lasso)                                                                   | $\mathbf{X}_{\hat{\mathbf{S}}}(\mathbf{X}_{\hat{\mathbf{S}}}^\top \mathbf{X}_{\hat{\mathbf{S}}})^{-1} \mathbf{X}_{\hat{\mathbf{S}}}^\top$                                        | [44], Proposition 4.1           |
| $g(\mathbf{b}) = \mu \ \mathbf{b}\ _2^2$ (Ridge)                                                                     | $\mathbf{X}(\mathbf{X}^\top \mathbf{X} + n\mu \mathbf{I}_p)^{-1} \mathbf{X}^\top$                                                                                                | (1.12), Section 4.1             |
| $g(\mathbf{b}) = \lambda \ \mathbf{b}\ _1 + \mu \ \mathbf{b}\ _2^2$ (Elastic-Net)                                    | $\mathbf{X}_{\hat{\mathbf{S}}}(\mathbf{X}_{\hat{\mathbf{S}}}^\top \mathbf{X}_{\hat{\mathbf{S}}} + n\mu \mathbf{I}_{ \hat{\mathbf{S}} })^{-1} \mathbf{X}_{\hat{\mathbf{S}}}^\top$ | [44], (28), [7], Section 3.5.3, |
| $g(\mathbf{b}) = \ \mathbf{b}\ _{\text{GL}} = \sum_{k=1}^K \lambda_k \ \mathbf{b}_{G_k}\ _2$<br>(group Lasso (3.33)) | $\mathbf{X}_{\hat{\mathbf{S}}}(\mathbf{X}_{\hat{\mathbf{S}}}^\top \mathbf{X}_{\hat{\mathbf{S}}} + \mathbf{M})^{-1} \mathbf{X}_{\hat{\mathbf{S}}}^\top$                           | [45], Proposition 4.2           |
| $g(\mathbf{b})$ twice continuously differentiable                                                                    | $\mathbf{X}(\mathbf{X}^\top \mathbf{X} + n\nabla^2 g(\hat{\boldsymbol{\beta}}))^{-1} \mathbf{X}^\top$                                                                            | Section 4.1                     |
| $g(\mathbf{b})$ arbitrary convex function                                                                            | symmetric with eigenvalues in $[0, 1]$                                                                                                                                           | Proposition 7.3                 |

3.3. *Approximation for  $\xi_0 = \mathbf{z}_0^\top f(\mathbf{z}_0) - \text{div } f(\mathbf{z}_0)$  and the debiased vector  $\hat{\boldsymbol{\beta}}^{(\text{de-bias})}$ .* Consider, for a fixed value of  $(\mathbf{X} \mathbf{Q}_0, \boldsymbol{\varepsilon})$  the function  $f_{(\mathbf{X} \mathbf{Q}_0, \boldsymbol{\varepsilon})} : \mathbb{R}^n \rightarrow \mathbb{R}^n$  given by

$$(3.5) \quad f_{(\mathbf{X} \mathbf{Q}_0, \boldsymbol{\varepsilon})}(\mathbf{z}_0) = f(\mathbf{z}_0) = \mathbf{X} \hat{\boldsymbol{\beta}} - \mathbf{y}.$$

For brevity, we will often omit the dependence on  $(\mathbf{X} \mathbf{Q}_0, \boldsymbol{\varepsilon})$  of  $f$  as discussed after (1.16). The Fréchet gradient  $\nabla f(\mathbf{z}_0)$ , where it exists, is uniquely defined by

$$(3.6) \quad f_{(\mathbf{X} \mathbf{Q}_0, \boldsymbol{\varepsilon})}(\mathbf{z}_0 + \boldsymbol{\eta}) - f_{(\mathbf{X} \mathbf{Q}_0, \boldsymbol{\varepsilon})}(\mathbf{z}_0) = [\nabla f(\mathbf{z}_0)]^\top \boldsymbol{\eta} + o(\|\boldsymbol{\eta}\|)$$

and the divergence by  $\text{div } f(\mathbf{z}_0) = \text{trace}[\nabla f(\mathbf{z}_0)]$ . If  $\tilde{\boldsymbol{\beta}} = \arg \min_{\mathbf{b} \in \mathbb{R}^p} (\|\boldsymbol{\varepsilon} - \tilde{\mathbf{X}}(\mathbf{b} - \boldsymbol{\beta})\|^2 / (2n) + g(\mathbf{b}))$  with  $\tilde{\mathbf{X}} = \mathbf{X} + \boldsymbol{\eta} \mathbf{a}_0^\top$ , then (3.6) is equivalent to

$$(3.7) \quad (\tilde{\mathbf{X}}(\tilde{\boldsymbol{\beta}} - \boldsymbol{\beta}) - \boldsymbol{\varepsilon}) - (\mathbf{X} \hat{\boldsymbol{\beta}} - \mathbf{y}) = [\nabla f(\mathbf{z}_0)]^\top \boldsymbol{\eta} + o(\|\boldsymbol{\eta}\|).$$

By Stein's formula, we have conditionally on  $(\mathbf{X} \mathbf{Q}_0, \boldsymbol{\varepsilon})$  that almost surely

$$(3.8) \quad \mathbb{E}[\xi_0 | (\mathbf{X} \mathbf{Q}_0, \boldsymbol{\varepsilon})] = 0 \quad \text{for } \xi_0 = \mathbf{z}_0^\top f(\mathbf{z}_0) - \text{div } f(\mathbf{z}_0).$$

As in (1.18) for the general case discussed in the Introduction, (3.8) gives an unbiased estimating equation for  $\theta = \langle \mathbf{a}_0, \boldsymbol{\beta} \rangle$ . The next lemma provides an expression for  $\nabla f(\mathbf{z}_0)$ .

LEMMA 3.1. *Let Assumption 3.1(i) be fulfilled,  $\mathbf{a}_0 \in \mathbb{R}^p$  and  $\hat{\mathbf{H}}$  be as in (3.3). Then*

$$(3.9) \quad \nabla f(\mathbf{z}_0)^\top = (\mathbf{I}_n - \hat{\mathbf{H}})^\top \langle \mathbf{a}_0, \mathbf{h} \rangle + \mathbf{w}_0(\mathbf{y} - \mathbf{X} \hat{\boldsymbol{\beta}})^\top$$

*satisfies (3.6) for some random  $\mathbf{w}_0 \in \mathbb{R}^n$  almost surely. If additionally  $\|\boldsymbol{\Sigma}^{-1/2} \mathbf{a}_0\| = 1$ , then*

$$(3.10) \quad \|\mathbf{w}_0\|^2 \leq n^{-1} \min\{(4\mu)^{-1}, \phi_{\min}(\boldsymbol{\Sigma}^{-1/2} \mathbf{X}^\top \mathbf{X} \boldsymbol{\Sigma}^{-1/2} / n)^{-1}\}.$$

Lemma 3.1 is proved in Section 7.1. Although we do not use this fact in any results, we mention here in passing that vector  $\mathbf{w}_0$  in (3.9) is linear in  $\mathbf{a}_0$  in the sense that  $\mathbf{w}_0$  can be chosen of the form  $\mathbf{W} \boldsymbol{\Sigma}^{-1/2} \mathbf{a}_0$  for some matrix  $\mathbf{W} \in \mathbb{R}^{n \times p}$ . Indeed, the proof of Lemma 3.1 shows that the map  $(\boldsymbol{\varepsilon}, \mathbf{X}) \mapsto \mathbf{X}(\hat{\boldsymbol{\beta}} - \boldsymbol{\beta}) - \boldsymbol{\varepsilon}$  is Fréchet differentiable at almost every point by Rademacher's theorem. At such a point, with  $\mathbf{a}_1, \mathbf{a}_2 \in \mathbb{R}^p$ ,  $t \in \mathbb{R}$ , the linear combination  $\mathbf{a}_3 = \mathbf{a}_1 + t\mathbf{a}_2$  and the perturbed design matrix  $\tilde{\mathbf{X}} = \mathbf{X} + \boldsymbol{\eta}(\mathbf{a}_1 + t\mathbf{a}_2)^\top$ , linearity of the Fréchet derivative implies that

$$\begin{aligned} & (\tilde{\mathbf{X}}(\tilde{\boldsymbol{\beta}} - \boldsymbol{\beta}) - \boldsymbol{\varepsilon}) - (\mathbf{X} \hat{\boldsymbol{\beta}} - \mathbf{y}) \\ &= (\langle \mathbf{a}_1 + t\mathbf{a}_2, \mathbf{h} \rangle (\mathbf{I}_n - \hat{\mathbf{H}})^\top + (\mathbf{w}_1 + t\mathbf{w}_2)(\mathbf{y} - \mathbf{X} \hat{\boldsymbol{\beta}})^\top) \boldsymbol{\eta} + o(\|\boldsymbol{\eta}\|), \end{aligned}$$

TABLE 2  
*Closed-form expressions for  $\mathbf{w}_0 \in \mathbb{R}^n$  in Lemma 3.1 for specific convex penalties  $g : \mathbb{R}^p \rightarrow \mathbb{R}$*

| Penalty                                                                                                              | Vector $\mathbf{w}_0 \in \mathbb{R}^n$ in Lemma 3.1                                                                                                                      | Justification        |
|----------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----------------------|
| $g(\mathbf{b}) = \lambda \ \mathbf{b}\ _1$ (Lasso)                                                                   | $\mathbf{X}_{\widehat{\mathcal{S}}}(\mathbf{X}_{\widehat{\mathcal{S}}}^\top \mathbf{X}_{\widehat{\mathcal{S}}})^{-1}(\mathbf{a}_0)_{\widehat{\mathcal{S}}}$              | [8], Proposition 4.1 |
| $g(\mathbf{b}) = \mu \ \mathbf{b}\ _2^2$ (Ridge)                                                                     | $\mathbf{X}(\mathbf{X}^\top \mathbf{X} + n\mu \mathbf{I}_p)^{-1} \mathbf{a}_0$                                                                                           | Section 4.1          |
| $g(\mathbf{b}) = \ \mathbf{b}\ _{\text{GL}} = \sum_{k=1}^K \lambda_k \ \mathbf{b}_{G_k}\ _2$<br>(group Lasso (3.33)) | $\mathbf{X}_{\widehat{\mathcal{S}}}(\mathbf{X}_{\widehat{\mathcal{S}}}^\top \mathbf{X}_{\widehat{\mathcal{S}}} + \mathbf{M})^{-1}(\mathbf{a}_0)_{\widehat{\mathcal{S}}}$ | Proposition 4.2      |
| $g(\mathbf{b})$ twice continuously differentiable                                                                    | $\mathbf{X}(\mathbf{X}^\top \mathbf{X} + n\nabla^2 g(\widehat{\boldsymbol{\beta}}))^{-1} \mathbf{a}_0$                                                                   | Section 4.1          |

where  $\mathbf{w}_1$  and  $\mathbf{w}_2$  denote the  $\mathbf{w}_0$  from (3.9) for  $\mathbf{a}_0 = \mathbf{a}_1$  and  $\mathbf{a}_0 = \mathbf{a}_2$ , respectively. Hence, with  $\mathbf{w}_3 = \mathbf{w}_1 + t\mathbf{w}_2$ , (3.9) holds for  $(\mathbf{a}_0, \mathbf{w}_0) = (\mathbf{a}_3, \mathbf{w}_3)$ . This proves that  $\mathbf{w}_0$  is linear in  $\mathbf{a}_0$ , that is, it is of the form  $\mathbf{w}_0 = \mathbf{W}\boldsymbol{\Sigma}^{-1/2}\mathbf{a}_0$  for some matrix  $\mathbf{W} \in \mathbb{R}^{n \times p}$ . One way to construct  $\mathbf{W}$  explicitly is the following: define the  $j$ th column of  $\mathbf{W}$  as the vector  $\mathbf{w}_0$  corresponding to  $\mathbf{a}_0 = \boldsymbol{\Sigma}^{1/2}\mathbf{e}_j$  where  $\mathbf{e}_j$  is the  $j$ th canonical basis vector. The linearity proved above for any linear combination  $\mathbf{a}_3$  then implies that (3.9) holds for  $\mathbf{w}_0 = \mathbf{W}\boldsymbol{\Sigma}^{-1/2}\mathbf{a}_0$  for any  $\mathbf{a}_0 \in \mathbb{R}^p$ . Inequality (3.10) provides an upper bound on the operator norm of  $\mathbf{W}$ . Linearity of  $\mathbf{w}_0$  with respect to  $\mathbf{a}_0$  and explicit matrices  $\mathbf{W}$  can be seen for some penalty functions in Table 2. Such Fréchet differentiability with respect to  $\mathbf{X}$  is used in [4, 6] to develop estimates of  $\|\boldsymbol{\Sigma}^{1/2}(\widehat{\boldsymbol{\beta}} - \boldsymbol{\beta})\|^2$  in linear models with Gaussian covariates similar to the present paper.

By taking the trace of (3.9), we obtain almost surely under Assumption 3.1,

$$\begin{aligned} -\xi_0 &= \operatorname{div} f(\mathbf{z}_0) - \langle \mathbf{z}_0, f(\mathbf{z}_0) \rangle \\ (3.11) \qquad &= (n - \widehat{\text{df}})(\langle \mathbf{a}_0, \widehat{\boldsymbol{\beta}} \rangle - \theta) + \langle \mathbf{z}_0 + \mathbf{w}_0, \mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}} \rangle \\ &= (n - \widehat{\text{df}})(\widehat{\theta} - \theta) \end{aligned}$$

for

$$(3.12) \qquad \widehat{\theta} \stackrel{\text{def}}{=} \langle \mathbf{a}_0, \widehat{\boldsymbol{\beta}} \rangle + (n - \widehat{\text{df}})^{-1} \langle \mathbf{z}_0 + \mathbf{w}_0, \mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}} \rangle.$$

In the present context of the regularized estimator  $\widehat{\boldsymbol{\beta}}$  in (3.1) with its effective degrees-of-freedom  $\widehat{\text{df}}$  defined in (3.4) and under Assumption 3.1, the quantities  $\xi_0, \widehat{\text{df}}, \widehat{\theta}$  in the previous display coincide with the random variables with the same name in (1.18). By (3.8), equality  $\mathbb{E}[\xi_0] = 0$  holds so that

$$\begin{aligned} 0 &= \mathbb{E}[-\xi_0] \\ (3.13) \qquad &= \mathbb{E}[(n - \widehat{\text{df}})(\widehat{\theta} - \theta)] \\ &= \mathbb{E}[(n - \widehat{\text{df}})(\langle \mathbf{a}_0, \widehat{\boldsymbol{\beta}} \rangle - \theta) + \langle \mathbf{z}_0 + \mathbf{w}_0, \mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}} \rangle] \end{aligned}$$

by taking expectations in (3.11). This provides a first evidence that the correction  $(n - \widehat{\text{df}})^{-1} \langle \mathbf{z}_0 + \mathbf{w}_0, \mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}} \rangle$  indeed removes the bias, at least after multiplication of  $(\widehat{\theta} - \theta)$  by  $(n - \widehat{\text{df}})$ . Since  $\mathbf{z}_0 = \mathbf{X}\boldsymbol{\Sigma}^{-1}\mathbf{a}_0$  under the normalization (1.15), the unbiased estimating equation (3.13) is the specialization of (1.11) from the Introduction to the penalized estimator (3.1), for which we have the gradient expression (3.9) in terms of  $\mathbf{w}_0$  and  $\widehat{\mathbf{H}}$ . If  $\widehat{\boldsymbol{\beta}}$  is given by (3.1), by identifying the terms in (1.18) and (3.11) we see that the random variable  $\widehat{A}$  defined in (1.10) of the Introduction is here  $\widehat{A} = \langle \mathbf{w}_0, \mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}} \rangle$  with  $\mathbf{w}_0$  given by Lemma 3.1. The following lemma shows that this term is negligible.

LEMMA 3.2. *Under Assumption 3.1, there exists  $\Omega_n$  with  $\mathbb{P}(\Omega_n^c) \leq C_0(\gamma, \mu)n^{-1/2}$  and*

$$(3.14) \qquad \mathbb{E}[I_{\Omega_n} \langle \mathbf{w}_0, \mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}} \rangle^2 / \operatorname{Var}_0[\xi_0]] \leq C_1(\gamma, \mu)n^{-1}.$$

The proof is given in Section 7.5. Since  $\mathbb{P}(\Omega_n) \rightarrow 1$ , inequality (3.14) implies  $\langle \mathbf{w}_0, \mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}} \rangle^2 / \text{Var}_0[\xi_0] \rightarrow^{\mathbb{P}} 0$ . This motivates the definition

$$(3.15) \quad \hat{\boldsymbol{\beta}}^{(\text{de-bias})} \stackrel{\text{def}}{=} \hat{\boldsymbol{\beta}} + (n - \widehat{\text{df}})^{-1} \boldsymbol{\Sigma}^{-1} \mathbf{X}^\top (\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}).$$

The debiased estimate  $\langle \mathbf{a}_0, \hat{\boldsymbol{\beta}}^{(\text{de-bias})} \rangle$  in direction  $\mathbf{a}_0$  is obtained from  $\hat{\theta}$  in (3.12) by dropping the smaller order term  $(n - \widehat{\text{df}})^{-1} \langle \mathbf{w}_0, \mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}} \rangle$ . By Slutsky's theorem, (3.14) implies

$$(3.16) \quad \frac{\xi_0}{\text{Var}_0[\xi_0]^{1/2}} \rightarrow^d F \quad \text{if and only if} \quad \frac{(n - \widehat{\text{df}}) \langle \mathbf{a}_0, \hat{\boldsymbol{\beta}}^{(\text{de-bias})} - \boldsymbol{\beta} \rangle}{\text{Var}_0[\xi_0]^{1/2}} \rightarrow^d F$$

for any candidate limiting distribution  $F$ . As  $\mathbb{E}[\xi_0] = 0$ , this suggests that the simpler correction in (3.15) also corrects the bias of  $\hat{\boldsymbol{\beta}}$ . By Prohorov's theorem, there exists a subsequence and limiting distribution  $F$  such that (3.16) holds in this subsequence. While  $F$  is mean zero as  $\xi_0 / \text{Var}_0[\xi_0]$  has mean zero and variance one,  $F$  has variance at most one by Fatou's lemma. However, the normality of  $F$  is unclear at this point.

To obtain more precise information on the limiting distribution and the deviations of  $\xi_0$ , the next subsections build estimate of its variance and derive asymptotic normality results by showing that  $F = N(0, 1)$  for most directions  $\mathbf{a}_0$ . The next result provides a loose data-driven upper bound on the error  $\langle \mathbf{a}_0, \hat{\boldsymbol{\beta}}^{(\text{de-bias})} \rangle - \theta$ .

**THEOREM 3.3.** *Under Assumption 3.1, there exists  $\Omega_n$  with  $\mathbb{P}(\Omega_n^c) \leq C_0(\gamma, \mu)n^{-1/2}$  and*

$$(3.17) \quad \mathbb{E}[I_{\Omega_n} (n - \widehat{\text{df}})^2 \langle \mathbf{a}_0, \hat{\boldsymbol{\beta}}^{(\text{de-bias})} - \boldsymbol{\beta} \rangle^2 / \|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2] \leq C_2(\gamma, \mu).$$

Furthermore,  $|\langle \mathbf{a}_0, \hat{\boldsymbol{\beta}}^{(\text{de-bias})} - \boldsymbol{\beta} \rangle| = O_{\mathbb{P}}(1) \|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\| / (n - \widehat{\text{df}}) = O_{\mathbb{P}}(1) \|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\| / n$ .

Theorem 3.3 is proved in Section 7.6. If  $\|\mathbf{X}(\hat{\boldsymbol{\beta}} - \boldsymbol{\beta})\|^2 / n = O_{\mathbb{P}}(\sigma^2)$ , then  $\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\| / n = O_{\mathbb{P}}(1)\sigma / \sqrt{n}$  is of the same order as the width of confidence intervals based on the least-squares estimator as  $n \rightarrow +\infty$  while  $p$  remains fixed. Theorem 3.3 shows that under this mild additional assumption on the prediction error  $\|\mathbf{X}(\hat{\boldsymbol{\beta}} - \boldsymbol{\beta})\|^2 / n$ , the second term in (3.15) indeed corrects the bias, achieving  $\langle \mathbf{a}_0, \hat{\boldsymbol{\beta}}^{(\text{de-bias})} - \boldsymbol{\beta} \rangle = O_{\mathbb{P}}(1)\sigma / \sqrt{n}$ .

**3.4. Variance estimates.** By Proposition 2.1, the conditional variance  $\text{Var}_0[\xi_0]$  can be written as  $\text{Var}_0[\xi_0] = \mathbb{E}_0[V^*(\theta)]$  for

$$(3.18) \quad V^*(\theta) \stackrel{\text{def}}{=} \|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2 + \text{trace}[\{\nabla f(\mathbf{z}_0)\}^2].$$

We allow the variance estimate to depend on the unknown  $\theta = \langle \mathbf{a}_0, \boldsymbol{\beta} \rangle$  as the resulting pivotal quantity,  $-V^*(\theta)^{-1/2}\xi_0 = V^*(\theta)^{-1/2}(n - \widehat{\text{df}})(\hat{\theta} - \theta)$  via (3.11), would depend on  $\theta$  anyway. While  $V^*(\theta)$  itself can be used to estimate  $\text{Var}_0[\xi_0]$ , its sign is unclear. The following simplified version of it, obtained by removing the smaller order terms in  $V^*(\theta)$ ,

$$(3.19) \quad \begin{aligned} \widehat{V}(\theta) &\stackrel{\text{def}}{=} \|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2 + \text{trace}[(\widehat{\mathbf{H}} - \mathbf{I}_n)^2] \langle \mathbf{a}_0, \hat{\boldsymbol{\beta}} - \boldsymbol{\beta} \rangle^2 \\ &= \|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2 + \|\widehat{\mathbf{H}} - \mathbf{I}_n\|_F^2 \langle \mathbf{a}_0, \mathbf{h} \rangle^2, \end{aligned}$$

is nonnegative. This follows from Proposition 7.3 since  $\mathbf{I}_n - \widehat{\mathbf{H}}$  is almost surely positive semidefinite. Lemma 3.4 below shows that the relative bias  $\mathbb{E}_0[\widehat{V}(\theta)] / \text{Var}_0[\xi_0] - 1$  converges to 0 in probability, that is,  $\widehat{V}(\theta)$  is asymptotically unbiased for  $\text{Var}_0[\xi_0]$ .

**LEMMA 3.4.** *Under Assumption 3.1, there exists  $\Omega_n$  with  $\mathbb{P}(\Omega_n^c) \leq C_0(\gamma, \mu)n^{-1/2}$  and*

$$(3.20) \quad \mathbb{E}\left[I_{\Omega_n} \left| \frac{\mathbb{E}_0[\widehat{V}(\theta)]}{\text{Var}_0[\xi_0]} - 1 \right| \right] \leq \mathbb{E}\left[I_{\Omega_n} \frac{|\mathbb{E}_0[\widehat{V}(\theta) - V^*(\theta)]|}{\text{Var}_0[\xi_0]} \right] \leq \frac{C_3(\gamma, \mu)}{n}.$$

An alternative variance estimate that does not depend on the unknown parameter  $\theta$  is given by replacing  $\theta$  in  $\widehat{V}(\theta)$  by the point estimate  $\langle \mathbf{a}_0, \widehat{\boldsymbol{\beta}}^{(\text{de-bias})} \rangle$  with  $\widehat{\boldsymbol{\beta}}^{(\text{de-bias})}$  in (3.15):

$$(3.21) \quad \check{V}(\mathbf{a}_0) = \widehat{V}(\langle \mathbf{a}_0, \widehat{\boldsymbol{\beta}}^{(\text{de-bias})} \rangle) = \|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2 + \|\mathbf{I}_n - \widehat{\mathbf{H}}\|_F^2 \frac{\langle \mathbf{z}_0, \mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}} \rangle^2}{(n - \widehat{\text{df}})^2}.$$

The next lemma provides  $\check{V}(\mathbf{a}_0)/\widehat{V}(\theta) \rightarrow^{\mathbb{P}} 1$  and that  $\check{V}(\mathbf{a}_0)$  is also asymptotically unbiased in the sense  $\mathbb{E}_0[\check{V}(\mathbf{a}_0)]/\text{Var}_0[\xi_0] \rightarrow^{\mathbb{P}} 1$ . Lemmas 3.4 and 3.5 are proved in Section 7.5.

LEMMA 3.5. *Under Assumption 3.1, there exists  $\Omega_n$  with  $\mathbb{P}(\Omega_n^c) \leq C_0(\gamma, \mu)n^{-1/2}$  and*

$$(3.22) \quad \max \left\{ \mathbb{E} \left[ I_{\Omega_n} \left| \frac{\check{V}(\mathbf{a}_0)^{1/2}}{\widehat{V}(\theta)^{1/2}} - 1 \right|^2 \right], \mathbb{E} \left[ I_{\Omega_n} \left| \frac{\mathbb{E}_0[\check{V}(\mathbf{a}_0)]^{1/2}}{\mathbb{E}_0[\widehat{V}(\theta)]^{1/2}} - 1 \right|^2 \right] \right\} \leq \frac{C_4(\gamma, \mu)}{n}.$$

3.5. *Asymptotic normality of debiased estimates.* Throughout this section,  $\Phi(t) = \mathbb{P}(N(0, 1) \leq t)$  denotes the standard normal cdf. For a given penalty function  $g : \mathbb{R}^p \rightarrow \mathbb{R}$ , we define the deterministic oracle  $\boldsymbol{\beta}^*$  and its associated noiseless prediction risk  $R_*$  by

$$(3.23) \quad \begin{aligned} \boldsymbol{\beta}^* &\stackrel{\text{def}}{=} \arg \min_{\mathbf{b} \in \mathbb{R}^p} \|\boldsymbol{\Sigma}^{1/2}(\mathbf{b} - \mathbf{b})\|^2/2 + g(\mathbf{b}), \\ \mathbf{h}^* &\stackrel{\text{def}}{=} \boldsymbol{\beta}^* - \boldsymbol{\beta}, \\ R_* &\stackrel{\text{def}}{=} \sigma^2 + \|\boldsymbol{\Sigma}^{1/2}\mathbf{h}^*\|^2. \end{aligned}$$

Our first result provides asymptotic normality of the debiased estimate when the error  $\langle \mathbf{a}_0, \widehat{\boldsymbol{\beta}} - \boldsymbol{\beta} \rangle$  of  $\widehat{\boldsymbol{\beta}}$  in direction  $\mathbf{a}_0$  is negligible compared to  $R_*$ .

THEOREM 3.6. *Let Assumption 3.1 be fulfilled. Let  $\widehat{\boldsymbol{\beta}}^{(\text{de-bias})}$  be as in (3.15). Then, for any  $\mathbf{a}_0$  with  $\|\boldsymbol{\Sigma}^{-1/2}\mathbf{a}_0\| = 1$  such that  $\langle \mathbf{a}_0, \mathbf{h} \rangle^2/R_* \rightarrow^{\mathbb{P}} 0$ ,*

$$\sup_{t \in \mathbb{R}} \left[ \left| \mathbb{P} \left( \frac{\xi_0}{V_0^{1/2}} \leq t \right) - \Phi(t) \right| + \left| \mathbb{P} \left( \frac{\langle \mathbf{a}_0, \widehat{\boldsymbol{\beta}}^{(\text{de-bias})} \rangle - \theta}{V_0^{1/2}/(n - \widehat{\text{df}})} \leq t \right) - \Phi(t) \right| \right] \rightarrow 0,$$

where  $V_0$  denotes any of the four quantities:  $\text{Var}_0[\xi_0]$ ,  $\|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2$ ,  $\widehat{V}(\theta)$  or  $\check{V}(\mathbf{a}_0)$ .

Theorem 3.6 is proved in Section 7.6. The theorem, as well as its variants below, are obtained by applying Theorem 2.2 conditionally on  $(\boldsymbol{\varepsilon}, \mathbf{X}\mathbf{Q}_0)$  to the function  $f(\mathbf{z}_0)$  in (3.5). This argument relies on the normality of  $\mathbf{z}_0$  conditionally on  $(\boldsymbol{\varepsilon}, \mathbf{X}\mathbf{Q}_0)$ , and thus the Gaussian design assumption. Here is an outline. Define

$$(3.24) \quad \delta_1^2(\mathbf{a}_0) \stackrel{\text{def}}{=} \mathbb{E}_0[\|\nabla f(\mathbf{z}_0)\|_F^2] / \{\mathbb{E}_0[\|f(\mathbf{z}_0)\|^2] + \mathbb{E}_0[\|\nabla f(\mathbf{z}_0)\|_F^2]\},$$

where  $\mathbb{E}_0$  and  $\text{Var}_0$  are the conditional expectation and conditional variance given  $(\mathbf{X}\mathbf{Q}_0, \boldsymbol{\varepsilon})$ . It is sufficient to show that  $\delta_1^2(\mathbf{a}_0) \rightarrow^{\mathbb{P}} 0$  in order to prove asymptotic normality of  $\xi_0/\text{Var}_0[\xi_0]^{1/2}$  by Theorem 2.2 and of  $\xi_0/\|f(\mathbf{z}_0)\|$  by Theorem 2.3 since  $\delta_1^2 = \delta_1^2(\mathbf{a}_0)$  satisfies  $2\delta_1^2 \geq \max(\epsilon_1^2, \bar{\epsilon}_1^2)$  for the  $\epsilon_1^2, \bar{\epsilon}_1^2$  in Theorems 2.2 and 2.3. The proof makes rigorous the following informal bound:

$$\delta_1^2(\mathbf{a}_0) = \frac{\mathbb{E}_0[\|\nabla f(\mathbf{z}_0)\|_F^2]}{\mathbb{E}_0[\|f(\mathbf{z}_0)\|^2] + \mathbb{E}_0[\|\nabla f(\mathbf{z}_0)\|_F^2]} \lesssim \frac{\mathbb{E}_0[\|\mathbf{I}_n - \widehat{\mathbf{H}}\|_F^2 \langle \mathbf{a}_0, \mathbf{h} \rangle^2]}{C_*^2(\gamma, \mu)nR_*} + O_{\mathbb{P}}(n^{-1/2})$$

for some constant  $C_*(\gamma, \mu)$ , by establishing a lower bound on  $\|f(\mathbf{z}_0)\|^2 = \|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2$  for the denominator (Lemmas 7.4, 7.6 and 7.7), and by showing that the rank one term  $\mathbf{w}_0(\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}})^\top$

in (3.9) is negligible in the numerator. Finally,  $\|\mathbf{I}_n - \widehat{\mathbf{H}}\|_F^2 \leq n$  always holds by Proposition 7.3 and  $\langle \mathbf{a}_0, \mathbf{h} \rangle^2 / R_*$  is shown to be uniformly integrable, so that the assumption  $\langle \mathbf{a}_0, \mathbf{h} \rangle^2 / R_* \rightarrow_{\mathbb{P}} 0$  grants  $\mathbb{E}[\delta_1^2(\mathbf{a}_0)] \rightarrow 0$ . The next two results identify directions  $\mathbf{a}_0$  such that  $\langle \mathbf{a}_0, \mathbf{h} \rangle^2 / R_* \rightarrow_{\mathbb{P}} 0$  holds.

**THEOREM 3.7.** *There exists an absolute constant  $C^* > 0$  such that the following holds. Let Assumption 3.1 be fulfilled,  $\widehat{\boldsymbol{\beta}}^{(\text{de-bias})}$  be as in (3.15). Then for any increasing sequence  $a_p \rightarrow +\infty$  (e.g.,  $a_p = \log \log p$ ), the subset*

$$(3.25) \quad \overline{S} = \{\mathbf{v} \in S^{p-1} : \mathbb{E}[(\boldsymbol{\Sigma}^{1/2} \mathbf{v}, \mathbf{h})^2 / \|\boldsymbol{\Sigma}^{1/2} \mathbf{h}\|^2] \leq C^*/a_p\}$$

*of the unit sphere  $S^{p-1}$  in  $\mathbb{R}^p$  has relative volume  $|\overline{S}|/|S^{p-1}| \geq 1 - 2e^{-p/a_p}$  and*

$$(3.26) \quad \sup_{\mathbf{a}_0 \in \boldsymbol{\Sigma}^{1/2} \overline{S}} \sup_{t \in \mathbb{R}} \left[ \left| \mathbb{P} \left( \frac{\xi_0}{V_0^{1/2}} \leq t \right) - \Phi(t) \right| + \left| \mathbb{P} \left( \frac{\langle \mathbf{a}_0, \widehat{\boldsymbol{\beta}}^{(\text{de-bias})} - \boldsymbol{\beta} \rangle}{V_0^{1/2}/(n - \widehat{\text{df}})} \leq t \right) - \Phi(t) \right| \right] \rightarrow 0$$

*where  $V_0$  denotes any of the four quantities:  $\text{Var}_0[\xi_0]$ ,  $\|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2$ ,  $\widehat{V}(\theta)$  or  $\check{V}(\mathbf{a}_0)$ . Furthermore, with  $\mathbf{e}_j \in \mathbb{R}^p$  the  $j$ th canonical basis vector and  $\phi_{\text{cond}}(\boldsymbol{\Sigma}) = \|\boldsymbol{\Sigma}\|_{\text{op}} \|\boldsymbol{\Sigma}^{-1}\|_{\text{op}}$ , the asymptotic normality in (3.26) uniformly holds over at least  $(p - \phi_{\text{cond}}(\boldsymbol{\Sigma})a_p/C^*)$  canonical directions in the sense that  $J_p = \{j \in [p] : \mathbf{e}_j / \|\boldsymbol{\Sigma}^{-1/2} \mathbf{e}_j\| \in \boldsymbol{\Sigma}^{1/2} \overline{S}\}$  has cardinality  $|J_p| \geq p - \phi_{\text{cond}}(\boldsymbol{\Sigma})a_p/C^*$ .*

Theorem 3.7 is proved in Section 7.6. For a given sequence of directions  $\mathbf{a}_0 \in \boldsymbol{\Sigma}^{1/2} S^{p-1}$ , if  $b_p = \mathbb{E}[\langle \mathbf{a}_0, \mathbf{h} \rangle^2 / \|\boldsymbol{\Sigma}^{1/2} \mathbf{h}\|^2] \rightarrow 0$  then it follows by choosing  $a_p = C^*/b_p$  that  $\mathbf{a}_0 \in \boldsymbol{\Sigma}^{1/2} \overline{S}$  for the  $\overline{S}$  in (3.25) so that (3.26) implies that asymptotic normality holds for this sequence of  $\mathbf{a}_0$ . In other words, asymptotic normality holds for all  $\mathbf{a}_0$  such that  $\mathbb{E}[\langle \mathbf{a}_0, \mathbf{h} \rangle^2 / \|\boldsymbol{\Sigma}^{1/2} \mathbf{h}\|^2] \rightarrow 0$ . Thus, a sequence of directions  $\mathbf{a}_0$  for which asymptotic normality does not follow from Theorem 3.7 is a sequence such that  $\mathbb{E}[\langle \mathbf{a}_0, \mathbf{h} \rangle^2 / \|\boldsymbol{\Sigma}^{1/2} \mathbf{h}\|^2]$  does not vanish, that is, the squared error  $\langle \mathbf{a}_0, \mathbf{h} \rangle^2$  in direction  $\mathbf{a}_0$  carries a constant fraction of the full prediction error  $\|\boldsymbol{\Sigma}^{1/2} \mathbf{h}\|^2$ . Such direction  $\mathbf{a}_0$  must thus be very special, which is embodied by the exponentially small bound on the relative volume  $|\overline{S} \setminus S^{p-1}|/|S^{p-1}|$ .

**THEOREM 3.8.** *Under Assumption 3.1, there exists  $\Omega_n$  with  $\mathbb{P}(\Omega_n^c) \leq C_0(\gamma, \mu)n^{-1/2}$  and*

$$(3.27) \quad \mathbb{E}[I_{\Omega_n}(\langle \mathbf{a}_0, \widehat{\boldsymbol{\beta}} - \boldsymbol{\beta} \rangle + (n - \widehat{\text{df}})^{-1} \mathbf{z}_0^\top (\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}))^2] \leq R_* C_5(\gamma, \mu)/n$$

*If additionally  $g$  is a seminorm, then  $|\mathbf{z}_0^\top (\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}})|/n = |\mathbf{a}_0^\top \boldsymbol{\Sigma}^{-1} \mathbf{X}^\top (\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}})|/n \leq g(\boldsymbol{\Sigma}^{-1} \mathbf{a}_0)$  always holds by properties of the subdifferential of a norm. Consequently, if  $g(\boldsymbol{\Sigma}^{-1} \mathbf{a}_0)^2 / R_* \rightarrow 0$  then  $\langle \mathbf{a}_0, \mathbf{h} \rangle^2 / R_* \rightarrow_{\mathbb{P}} 0$  and the conclusions of Theorem 3.6 hold.*

Theorem 3.8 is proved in Section 7.6. The first part of the theorem says that the estimation error  $\langle \mathbf{a}_0, \mathbf{h} \rangle$  is essentially  $-\langle \mathbf{z}_0, \mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}} \rangle / (n - \widehat{\text{df}})$  up to an error term of order  $R_* O_{\mathbb{P}}(n^{-1/2})$ , so that  $\langle \mathbf{a}_0, \mathbf{h} \rangle^2 / R_* \rightarrow_{\mathbb{P}} 0$  if and only if  $(n - \widehat{\text{df}})^{-2} \langle \mathbf{z}_0, \mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}} \rangle^2 / R_* \rightarrow_{\mathbb{P}} 0$ . Combined with the fact that  $(1 - \widehat{\text{df}}/n)$  is bounded away from 0 by Lemma 7.4, this implies that  $\langle \mathbf{a}_0, \mathbf{h} \rangle^2 / R_* \rightarrow_{\mathbb{P}} 0$  if and only if  $n^{-2} \langle \mathbf{z}_0, \mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}} \rangle^2 / R_* \rightarrow_{\mathbb{P}} 0$ . The last part of the theorem relies on the property  $g(\mathbf{u}) \leq \sup_{\mathbf{s} \in \partial g(\mathbf{u})} \mathbf{u}^\top \mathbf{s}$  for any norm  $g$  where  $\partial g(\mathbf{u})$  is the subdifferential of  $g$  at  $\mathbf{u}$ . This property also holds if  $g$  is a seminorm; however, it is unclear how to extend the last part of the above theorem if  $g$  is not a seminorm.

For the Lasso, the penalty function is  $g(\mathbf{b}) = \lambda \|\mathbf{b}\|_1$  and condition  $g(\boldsymbol{\Sigma}^{-1} \mathbf{a}_0)^2 / R_* \rightarrow 0$  becomes  $\lambda^2 \|\boldsymbol{\Sigma}^{-1} \mathbf{a}_0\|_1^2 / R_* \rightarrow 0$ . Typically, the tuning parameter  $\lambda$  is chosen as  $\lambda \propto \sigma(2 \log(p/k)/n)^{1/2}$  with  $k = 1$  [11], among others, or  $k = s_0$  [3, 5, 25, 30, 39], where  $s_0 =$

$\|\boldsymbol{\beta}\|_0$ . For such choices, the condition  $\lambda^2 \|\boldsymbol{\Sigma}^{-1} \mathbf{a}_0\|_1^2 / R_* \rightarrow 0$  can be written as  $\|\boldsymbol{\Sigma}^{-1} \mathbf{a}_0\|_1 = (R_*/\sigma^2)^{1/2} o(\sqrt{n/\log(p/k)})$  and since  $R_* \geq \sigma^2$ , a sufficient condition is  $\|\boldsymbol{\Sigma}^{-1} \mathbf{a}_0\|_1 = o(\sqrt{n/\log(p/k)})$ . If  $\mathbf{a}_0 = \mathbf{e}_j$  is a vector of the canonical basis, the normalization (1.15) gives  $(\boldsymbol{\Sigma}^{-1})_{jj} = 1$  and  $\|\boldsymbol{\Sigma}^{-1} \mathbf{e}_j\|_1$  is the  $\ell_1$  norm of the  $j$ th column of  $\boldsymbol{\Sigma}^{-1}$ . The condition  $\|\boldsymbol{\Sigma}^{-1} \mathbf{a}_0\|_1 = o(\sqrt{n/\log(p/k)})$  allows, for instance, the  $j$ th column of  $\boldsymbol{\Sigma}^{-1}$  to have  $o(\sqrt{n/\log(p/k)})$  constant entries. This assumption is weaker than that of some previous studies; for instance, [28] requires  $\|\boldsymbol{\Sigma}^{-1} \mathbf{a}_0\|_1 = O(1)$  for  $\mathbf{a}_0 = \mathbf{e}_j$ . The following example illustrates the benefit of picking a proper penalty level  $\lambda$ .

EXAMPLE 1. Let  $p/n \rightarrow \gamma < 1$  and  $g(\mathbf{b}) = \lambda \|\mathbf{b}\|_1$ .

(i) For  $\lambda = 0$ , the Lasso and debiased Lasso are both identical to the least squares estimator and the debiasing correction proportional to  $\mathbf{z}_0^\top (\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}})$  is 0 since  $\mathbf{X}^\top (\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}) = 0$ , so that  $\hat{\theta} - \theta = \langle \mathbf{a}_0, \mathbf{h} \rangle = \mathbf{a}_0^\top (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \boldsymbol{\varepsilon}$  in (3.12),  $\widehat{\text{df}} = p$ ,  $\hat{V}(\theta) \approx \|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2 \sim \sigma^2 \chi_{n-p}^2$  and  $\sqrt{n}(\hat{\theta} - \theta) \xrightarrow{d} N(0, \sigma^2/(1 - \gamma))$ .

(ii) Suppose  $|\hat{S}|/n + \|\mathbf{X}\mathbf{h}\|^2/n + \|\boldsymbol{\Sigma}^{1/2} \mathbf{h}\|^2 = o_{\mathbb{P}}(1)$  for suitable  $\lambda \geq \sigma \sqrt{2 \log(p/s_0)/n}$  as in [8, 50]. Then  $\hat{V}(\theta) = (1 + o_{\mathbb{P}}(1))n\sigma^2$  and  $\sqrt{n}(\hat{\theta} - \theta) \xrightarrow{d} N(0, \sigma^2)$ .

3.6. *Confidence intervals.* Theorems 3.6 to 3.8 are valid for any choice of the variance estimate among  $\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2$ ,  $\hat{V}(\theta)$  in (3.19) and  $\check{V}(\mathbf{a}_0)$  in (3.22) for directions  $\mathbf{a}_0$  such that  $\langle \mathbf{a}_0, \mathbf{h} \rangle / R_* \rightarrow^{\mathbb{P}} 0$  holds. For such direction  $\mathbf{a}_0$ , the choice  $\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2$  leads to the narrowest confidence interval for  $\theta$ , namely

$$(3.28) \quad \mathbb{P}(\theta \in [\langle \mathbf{a}_0, \hat{\boldsymbol{\beta}}^{(\text{de-bias})} \rangle \pm z_{\alpha/2}(n - \widehat{\text{df}})^{-1} \|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|]) \rightarrow 1 - \alpha,$$

where  $[u \pm v]$  denotes the interval  $[u - v, u + v]$ ,  $\mathbb{P}(|N(0, 1)| > z_{\alpha/2}) = \alpha$  and  $\hat{\boldsymbol{\beta}}^{(\text{de-bias})}$  is the debiased estimator in (3.15). The choice  $\check{V}(\mathbf{a}_0)$  leads to intervals with larger multiplicative coefficient for  $z_{\alpha/2}$ , namely

$$(3.29) \quad \theta \in \left[ \langle \mathbf{a}_0, \hat{\boldsymbol{\beta}}^{(\text{de-bias})} \rangle \pm z_{\alpha/2} \left( \frac{\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2}{(n - \widehat{\text{df}})^2} + \frac{\|\mathbf{I}_n - \hat{\mathbf{H}}\|_F^2 \langle \mathbf{z}_0, \mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}} \rangle^2}{(n - \widehat{\text{df}})^4} \right)^{1/2} \right]$$

has probability converging to  $1 - \alpha$  for directions  $\mathbf{a}_0$  satisfying any of the above theorems. For such directions, the choice  $\hat{V}(\theta)$  justifies confidence intervals of the form (1.21) as

$$(3.30) \quad ((n - \widehat{\text{df}})(\langle \mathbf{a}_0, \hat{\boldsymbol{\beta}} \rangle - \theta) + \langle \mathbf{z}_0, \mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}} \rangle)^2 - \hat{V}(\theta) z_{\alpha/2}^2 \leq 0$$

holds with probability converging to  $1 - \alpha$ . Given the expression for  $\hat{V}(\theta)$  in (3.19), the left-hand side of (3.30) is a quadratic polynomial in  $\theta$  with dominant coefficient  $(n - \widehat{\text{df}})^2 - z_{\alpha/2}^2 \|\mathbf{I}_n - \hat{\mathbf{H}}\|_F^2$ . Since  $\|\mathbf{I}_n - \hat{\mathbf{H}}\|_F^2 \leq n - \widehat{\text{df}}$  almost surely by properties of  $\hat{\mathbf{H}}$  in Proposition 7.3 and  $n - \widehat{\text{df}} \geq C_*(\gamma, \mu)n$  for some constant  $C_*(\gamma, \mu)$  with probability approaching one by Lemma 7.4, in the same event the dominant coefficient is positive. The intersection of events (3.30) and  $\{n - \widehat{\text{df}} \geq C_*(\gamma, \mu)n\}$  has probability converging to  $1 - \alpha$  and in this event,  $\theta \in [\Theta_1(z_{\alpha/2}), \Theta_2(z_{\alpha/2})]$  where  $\Theta_1(z_{\alpha/2}), \Theta_2(z_{\alpha/2})$  are the two real roots of the left-hand side of (3.30) as a quadratic function of  $\theta$ .

3.7. *Variance spike.* One can pick any choice among the three variance estimates in Theorem 3.6 because it assumes  $\langle \mathbf{a}_0, \mathbf{h} \rangle / R_* \rightarrow^{\mathbb{P}} 0$  and this limit in probability implies both  $\hat{V}(\theta)/\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2 \rightarrow^{\mathbb{P}} 1$  and  $\check{V}(\mathbf{a}_0)/\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2 \rightarrow^{\mathbb{P}} 1$ . These limits in probability to 1 are made rigorous by (3.22) and by the lower bound  $\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2 \geq R_* n (C_*^2(\gamma, \mu) - O_{\mathbb{P}}(n^{-1/2}))$  obtained from Lemmas 7.4, 7.6 and 7.7 as explained in (7.38) of the proof.

The reason that the estimates  $\widehat{V}(\theta)$  and  $\check{V}(\mathbf{a}_0)$  were introduced is that the simpler estimate  $\|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2$  is not asymptotically unbiased for  $\text{Var}_0[\xi_0]$  for directions  $\mathbf{a}_0$  such that  $\langle \mathbf{a}_0, \mathbf{h} \rangle^2 / R_*$  does not converge to 0 in probability: While the relative bias of  $\widehat{V}(\theta)$  and  $\check{V}(\mathbf{a}_0)$  provably converges to 0 in Lemma 3.4 and (3.22) for all directions  $\mathbf{a}_0$ , the same cannot be said for the simpler estimate  $\|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2$ .

**THEOREM 3.9.** *Let Assumption 3.1 be fulfilled. Then the following are equivalent:*

- (i)  $\|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2 / \text{Var}_0[\xi_0] \rightarrow^{\mathbb{P}} 1$ ,
- (ii)  $\mathbb{E}_0[\|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2] / \text{Var}_0[\xi_0] \rightarrow^{\mathbb{P}} 1$ ,
- (iii)  $\langle \mathbf{a}_0, \mathbf{h} \rangle^2 / R_* \rightarrow^{\mathbb{P}} 0$ ,
- (iv)  $\langle \mathbf{a}_0, \mathbf{h} \rangle^2 n / \|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2 \rightarrow^{\mathbb{P}} 0$ ,
- (v)  $\langle \mathbf{z}_0, \mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}} \rangle^2 / (n \|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2) \rightarrow^{\mathbb{P}} 0$ ,
- (vi)  $\widehat{V}(\theta) / \|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2 \rightarrow^{\mathbb{P}} 1$ ,
- (vii)  $\check{V}(\mathbf{a}_0) / \|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2 \rightarrow^{\mathbb{P}} 1$ .

Theorem 3.9 is proved in Section 7.5. It shows that for the directions  $\mathbf{a}_0$  such that  $\langle \mathbf{a}_0, \mathbf{h} \rangle^2 n / \|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2 \rightarrow^{\mathbb{P}} 0$  does not hold, for example, directions such that the error  $\langle \mathbf{a}_0, \mathbf{h} \rangle^2$  is of the same order as the average squared residual  $\|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2 / n$  (see Lemma 7.2 in Section 7.2), the simpler estimate  $\|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2$  fails to account for a nonnegligible part of the variance  $\text{Var}_0[\xi_0]$  by item (i) above. The goal of the estimates  $\widehat{V}(\theta)$  and  $\check{V}(\mathbf{a}_0)$  is to repair this as  $\mathbb{E}_0[\widehat{V}(\theta)] / \text{Var}_0[\xi_0] \rightarrow^{\mathbb{P}} 1$  and  $\mathbb{E}_0[\check{V}(\mathbf{a}_0)] / \text{Var}_0[\xi_0] \rightarrow^{\mathbb{P}} 1$  hold for all directions by Lemmas 3.4 and 3.5, even for directions  $\mathbf{a}_0$  such that (i)–(vii) above fail. Note that the quantity  $\langle \mathbf{z}_0, \mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}} \rangle^2 / (n \|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2)$  in item (v) is observable (i.e., does not depend on  $\boldsymbol{\beta}$ ), so that (i)–(vii) are expected to hold when this quantity is sufficiently small.

For directions  $\mathbf{a}_0$  such that  $\langle \mathbf{a}_0, \mathbf{h} \rangle^2 n / \|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2 \rightarrow^{\mathbb{P}} 0$  does not hold, we expect a variance spike, that is, an extra additive term in the variance estimate equal to  $\|\mathbf{I}_n - \widehat{\mathbf{H}}\|_F^2 \langle \mathbf{a}_0, \mathbf{h} \rangle^2$  in  $\widehat{V}(\theta)$  and to  $\|\mathbf{I}_n - \widehat{\mathbf{H}}\|_F^2 \langle \mathbf{z}_0, \mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}} \rangle^2 / (n - \text{df})^2$  in  $\check{V}(\mathbf{a}_0)$ . The confidence interval (3.28) that does not take into account this variance spike is expected to be too narrow and to suffer from undercoverage for directions  $\mathbf{a}_0$  with large  $\langle \mathbf{a}_0, \mathbf{h} \rangle^2 n / \|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2$ . The wider confidence interval (3.29) is expected to repair this, although for directions  $\mathbf{a}_0$  such that  $\langle \mathbf{a}_0, \mathbf{h} \rangle^2 n / \|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2 \rightarrow^{\mathbb{P}} 0$  does not hold the current theory does not prove whether the asymptotic distribution is normal. The theoretical evidence that the variance spike occurs is grounded in the relative asymptotic unbiasedness of  $\widehat{V}(\theta)$  in (3.20) and of  $\check{V}(\mathbf{a}_0)$  in (3.22), combined with the negative result for the simpler variance estimate  $\|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2$  via Theorem 3.9 as discussed above. Figure 1 in Section 5 illustrates the variance spike on simulations for the Lasso and direction  $\mathbf{a}_0$  proportional to the first canonical basis vector.

The second term in the variance estimates (3.19) and (3.21) is necessary for certain  $\mathbf{a}_0$  for the estimate  $\widehat{\boldsymbol{\beta}} = \mathbf{0}$ , which corresponds to (1.2) with penalty  $g(\mathbf{0}) = 0$  and  $g(\mathbf{b}) = +\infty$  for  $\mathbf{b} \neq \mathbf{0}$ . For  $\boldsymbol{\Sigma} = \mathbf{I}_p$ ,  $\mathbf{a}_0 = \boldsymbol{\beta} / \|\boldsymbol{\beta}\|$  and  $V_* = 2\|\boldsymbol{\beta}\|^2 + \sigma^2$ ,

$$(3.31) \quad \frac{-\xi_0}{\widehat{V}(\theta)} = \frac{-n\langle \mathbf{a}_0, \boldsymbol{\beta} \rangle + \mathbf{z}_0^\top \mathbf{y}}{(\|\mathbf{y}\|^2 + n\langle \mathbf{a}_0, \boldsymbol{\beta} \rangle^2)^{1/2}} = \frac{(nV_*)^{-1/2} \sum_{i=1}^n (-\|\boldsymbol{\beta}\| + z_{0i}\epsilon_i + \|\boldsymbol{\beta}\|z_{0i}^2)}{(\|\mathbf{y}\|^2/n + \|\boldsymbol{\beta}\|^2)^{1/2} / V_*^{1/2}}.$$

By the CLT, the numerator of the rightmost quantity converges to  $N(0, 1)$  and in the denominator,  $(\|\mathbf{y}\|^2/n + \|\boldsymbol{\beta}\|^2) / V_* \rightarrow^{\mathbb{P}} 1$  by the weak law of large numbers, so that (3.31) converges in law to  $N(0, 1)$  by Slutsky's theorem. On the other hand, if the variance estimate  $\|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2$  is used instead of  $\widehat{V}(\theta)$  in the denominator, the CLT  $-\xi_0 / \|\mathbf{y}\| \rightarrow^d N(0, V_* / (\|\boldsymbol{\beta}\|^2 + \sigma^2))$  still holds but the asymptotic variance is  $1 + (1 + \sigma^2 / \|\boldsymbol{\beta}\|^2)^{-1} > 1$ .

3.8. *Relaxing strong convexity when  $p > n$ .* The previous theorems are valid under Assumption 3.1: Either  $\gamma < 1$  and  $g$  is an arbitrary convex function, or  $\gamma \geq 1$  and  $g$  is required to be strongly convex with parameter  $\mu$ . If  $p/n > 1$  and the penalty  $g$  is not strongly convex, the techniques of the present paper still provide asymptotic normality results under additional assumptions as shown in the following result.

Consider either the Lasso

$$(3.32) \quad \hat{\boldsymbol{\beta}} = \arg \min_{\mathbf{b} \in \mathbb{R}^p} \{ \|\mathbf{y} - \mathbf{X}\mathbf{b}\|^2 / (2n) + \lambda \|\mathbf{b}\|_1 \}$$

for some  $\lambda > 0$  or the group Lasso norm  $\|\cdot\|_{\text{GL}}$  and group Lasso  $\hat{\boldsymbol{\beta}}$  defined as

$$(3.33) \quad \hat{\boldsymbol{\beta}} = \arg \min_{\mathbf{b} \in \mathbb{R}^p} \{ \|\mathbf{y} - \mathbf{X}\mathbf{b}\|^2 / (2n) + \|\mathbf{b}\|_{\text{GL}} \}, \quad \|\mathbf{b}\|_{\text{GL}} = \sum_{k=1}^K \lambda_k \|\mathbf{b}_{G_k}\|_2,$$

where  $(G_1, \dots, G_K)$  is a partition of  $\{1, \dots, p\}$  into  $K$  nonoverlapping groups and  $\lambda_1, \dots, \lambda_K > 0$  are tuning parameters. If each  $G_k$  is a singleton and  $\lambda_k = \lambda > 0$  for all  $k = 1, \dots, p$ , then (3.33) reduces to the Lasso (3.32).

**THEOREM 3.10.** *Let  $\gamma \geq 1$ ,  $\kappa \in (0, 1)$  be constants independent of  $\{n, p\}$ . Consider a sequence of regression problems with  $p/n \leq \gamma$  and invertible  $\boldsymbol{\Sigma}$ . Assume that the group Lasso estimator  $\hat{\boldsymbol{\beta}}$  in (3.33) satisfies*

$$(3.34) \quad \mathbb{P}(\|\hat{\boldsymbol{\beta}}\|_0 \leq \kappa n / 2) \rightarrow 1.$$

*If  $\mathbf{a}_0$  is such that  $\|\boldsymbol{\Sigma}^{-1/2} \mathbf{a}_0\| = 1$  and  $\langle \mathbf{a}_0, \hat{\boldsymbol{\beta}} - \boldsymbol{\beta} \rangle^2 / R_* \xrightarrow{\mathbb{P}} 0$  for the  $R_*$  in (3.23), then*

$$(3.35) \quad \sup_{t \in \mathbb{R}} |\mathbb{P}(\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^{-1} ((n - \text{df}) \langle \mathbf{a}_0, \hat{\boldsymbol{\beta}} - \boldsymbol{\beta} \rangle + \mathbf{z}_0^\top (\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}})) \leq t) - \Phi(t)| \rightarrow 0.$$

*Furthermore, for any  $a_p$  with  $a_p \rightarrow \infty$  and  $\bar{S}$  in (3.25), the relative volume bound given after (3.25) holds, and the asymptotic normality (3.35) holds uniformly over all  $\mathbf{a}_0 \in \boldsymbol{\Sigma}^{1/2} \bar{S}$  and uniformly over at least  $(p - \phi_{\text{cond}}(\boldsymbol{\Sigma}) a_p / C^*)$  canonical directions in the sense that  $J_p = \{j \in [p] : \mathbf{e}_j / \|\boldsymbol{\Sigma}^{-1/2} \mathbf{e}_j\| \in \boldsymbol{\Sigma}^{1/2} \bar{S}\}$  has cardinality  $|J_p| \geq p - \phi_{\text{cond}}(\boldsymbol{\Sigma}) a_p / C^*$ .*

Theorem 3.10 is proved in Appendix B. The strong convexity requirement in Assumption 3.1 is relaxed and replaced by assuming the high-probability bound  $\|\hat{\boldsymbol{\beta}}\|_0 \leq \kappa n / 2$  on the number of nonzero coordinates. Surprisingly, no conditions are required on the true regression vector  $\boldsymbol{\beta}$  or on the tuning parameters, although these quantities affect whether  $\mathbb{P}(\|\hat{\boldsymbol{\beta}}\|_0 \leq \kappa n / 2) \rightarrow 1$  is satisfied. Figure 2 in Section 6 illustrates Theorem 3.10 on simulated data.

*Comparison with existing works on the Lasso.* The Lasso is largely the most studied initial estimator in previous literature on debiasing and asymptotic normality, so it provides a level playing field to compare our method with existing results. In the approximate message passing (AMP) literature, which includes most existing works in the  $n/p \rightarrow \gamma$  regime, for example, [21, 23, 27] or more recently [33, 40, 42, 43], it is assumed that  $\boldsymbol{\Sigma} = \mathbf{I}_p$  and that the empirical distribution  $G_{n,p}(t) = p^{-1} \sum_{j=1}^p I\{\sqrt{n} \beta_j \leq t\}$  converges in distribution and in the second moment to some “prior”  $G$  as  $n, p \rightarrow +\infty$ . Assume these conditions and consider the  $j$ th component  $\hat{\beta}_j^{(\text{de-bias})}$  of  $\hat{\boldsymbol{\beta}}^{(\text{de-bias})}$  in (3.15) for  $\boldsymbol{\Sigma} = \mathbf{I}_p$ , that is,

$$\hat{\beta}_j^{(\text{de-bias})} = \hat{\beta}_j + \langle \mathbf{X} \mathbf{e}_j, \mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}} \rangle / (n - \text{df}).$$

Then the Lasso has the interpretation as its soft thresholded debiased version,

$$\widehat{\beta}_j = \eta(\widehat{\beta}_j^{(\text{de-bias})}; \lambda/(1 - \widehat{\text{df}}/n)) \quad \text{where } \eta(u; t) = \text{sgn}(u)(|u| - t)_+$$

and the main thrust of the AMP theory is that the joint empirical distribution of the debiased errors and the true coefficients,

$$H_{n,p}(u, t) = p^{-1} \sum_{j=1}^p I\{\sqrt{n}\widehat{\beta}_j^{(\text{de-bias})} - \sqrt{n}\beta_j \leq u, \sqrt{n}\beta_j \leq t\},$$

converges in distribution and the second moment to the limit  $H$  with independent  $N(0, \tau_0)$  and  $G$  components, where  $\tau_0$  is characterized by a system of nonlinear equations with 2 or 3 unknowns. These nonlinear equations depend on the loss (here, the  $\ell_2$  loss), the penalty (here, the  $\ell_1$ -norm), the distribution of the noise, as well as the prior distribution that governs the empirical distribution of the coefficients of  $\beta$ . We note that these works typically assume that  $X$  has  $N(0, 1/n)$  entries, so that their coefficient vector is equivalent to our  $\sqrt{n}\beta$ . For instance, [33] characterizes the limit of the empirical distribution of  $(\sqrt{n}\widehat{\beta}^{(\text{de-bias})}, \sqrt{n}\beta)$  in terms of two parameters,  $\{\tau_*(\lambda), \kappa_*(\lambda)\}$ , that are defined as solutions of the nonlinear equations in [33], Proposition 3.1; see also [17], Proposition 4.3, for similar results applicable to permutation-invariant penalty. This approach presents some drawbacks: For instance, it requires the convergence of the empirical distribution  $G_{n,p}$  to a limit (which can be viewed as a prior), it yields the limiting distribution for the joint empirical distribution  $H_{n,p}$  of the estimation errors and the unknown coefficients but not for a fixed coordinate.

The above Theorem 3.10 for the Lasso differs from this previous literature in major ways. First, it provides a limiting distribution for the de-biased version of  $\langle \mathbf{a}_0, \widehat{\beta} \rangle$  for a single, fixed direction  $\mathbf{a}_0$ : Theorem 3.10 does not involve the empirical distributions of  $\sqrt{n}\beta$ ,  $\sqrt{n}\widehat{\beta}$  or its debiased version. This contrasts with previous literature on the  $n/p \rightarrow \gamma$  regime where the confidence interval guarantee holds on average over the coefficients  $\{1, \dots, p\}$  [21, 23, 27, 40]. This improvement is important in practice: if the practitioner is interested in the effect of a specific effect  $j_0 \in \{1, \dots, p\}$ , it is important to construct confidence intervals with strict type I error control for  $\beta_{j_0}$ , as opposed to a controlled type I error that only holds on average over all coefficients. Another feature of the results in this paper is that there is no need to assume a prior on the coefficients of  $\beta$  in the limit.

Surprisingly, Theorems 3.6, 3.7 and 3.10 and their proofs completely bypass solving the nonlinear equations that appear in the aforementioned works as the nonlinearity is directly treated here with the normal approximation in Theorem 2.2. Asymptotic normality in Theorems 3.6, 3.7 and 3.10 is obtained for a fixed direction  $\mathbf{a}_0$  (or a fixed coordinate  $j \in \{1, \dots, p\}$  when  $\mathbf{a}_0 = \mathbf{e}_j$ ), and the correlations in  $\Sigma$  are handled with a direct approach. Since the first version of this paper was made public, extensions of works cited in the two previous paragraphs were developed [18, 32] to obtain, for  $\Sigma \neq \mathbf{I}_p$ , asymptotic normality results in an averaged sense over  $\{1, \dots, p\}$ . It is unclear at this point whether these methods can yield asymptotic normality for a fixed coordinate instead of in an averaged sense.

**4. Examples.** We now present three penalty functions for which closed-form expressions for  $\widehat{H}$  and  $\mathbf{w}_0$  are available. In this section, when computing gradients with respect to  $\mathbf{z}_0$  in order to find closed-form expressions for  $\mathbf{w}_0$  in (3.9), we consider  $(X\mathbf{Q}_0, \boldsymbol{\varepsilon})$  fixed as in (3.7). Explicitly,  $\nabla \widehat{\beta}(\mathbf{z}_0)^\top$  is uniquely defined as

$$(4.1) \quad \widetilde{\beta} - \widehat{\beta} = [\nabla \widehat{\beta}(\mathbf{z}_0)]^\top \boldsymbol{\eta} + o(\|\boldsymbol{\eta}\|),$$

where  $\widehat{\beta} = \arg \min_{\mathbf{b} \in \mathbb{R}^p} \|X(\mathbf{b} - \beta) - \boldsymbol{\varepsilon}\|^2/(2n) + g(\mathbf{b})$  and  $\widetilde{\beta} = \arg \min_{\mathbf{b} \in \mathbb{R}^p} \|(X + \boldsymbol{\eta}\mathbf{a}_0^\top)(\mathbf{b} - \beta) - \boldsymbol{\varepsilon}\|^2/(2n) + g(\mathbf{b})$ . When computing gradients with respect to  $\mathbf{y}$  in order to find closed-form expressions for  $\widehat{H}$  in (3.3), we view  $\beta(\mathbf{y}, X) = \arg \min_{\mathbf{b} \in \mathbb{R}^p} \|X\mathbf{b} - \mathbf{y}\|^2/(2n) + g(\mathbf{b})$  as

a function of  $(\mathbf{y}, \mathbf{X})$  and if  $\mathbf{y} \mapsto \widehat{\boldsymbol{\beta}}(\mathbf{y}, \mathbf{X})$  is differentiable at  $\mathbf{y}$  for a fixed  $\mathbf{X}$  then

$$(4.2) \quad \widehat{\boldsymbol{\beta}}(\tilde{\mathbf{y}}, \mathbf{X}) - \widehat{\boldsymbol{\beta}}(\mathbf{y}, \mathbf{X}) = [(\partial/\partial \mathbf{y})\widehat{\boldsymbol{\beta}}(\mathbf{y}, \mathbf{X})](\tilde{\mathbf{y}} - \mathbf{y}) + o(\|\tilde{\mathbf{y}} - \mathbf{y}\|),$$

where  $(\partial/\partial \mathbf{y})\widehat{\boldsymbol{\beta}}(\mathbf{y}, \mathbf{X}) \in \mathbb{R}^{p \times n}$  is the Jacobian. Once the Jacobian  $(\partial/\partial \mathbf{y})\widehat{\boldsymbol{\beta}}(\mathbf{y}, \mathbf{X})$  is computed,  $\widehat{\mathbf{H}}$  in (3.3) is given by  $\widehat{\mathbf{H}}^\top = \mathbf{X}(\partial/\partial \mathbf{y})\widehat{\boldsymbol{\beta}}(\mathbf{y}, \mathbf{X})$ . We use the Jacobian notation  $(\partial/\partial \mathbf{y})\widehat{\boldsymbol{\beta}}(\mathbf{y}, \mathbf{X})$  when computing the derivatives with respect to  $\mathbf{y}$  to avoid confusion with the gradient  $\nabla \boldsymbol{\beta}(\mathbf{z}_0)$  in (4.1).

**4.1. Twice continuously differentiable penalty.** The simplest example for which closed-form expressions for  $\widehat{\mathbf{H}}$ ,  $\widehat{\mathbf{d}}\mathbf{f}$ ,  $\mathbf{w}_0$  can be obtained is that of twice continuously differentiable and strongly convex penalty  $g$ . If  $g$  is strongly convex, Lemma 7.1 proves that the Fréchet derivative of  $\mathbf{h} = \widehat{\boldsymbol{\beta}} - \boldsymbol{\beta}$  with respect to  $(\boldsymbol{\varepsilon}, \mathbf{X})$  exist for almost every  $(\boldsymbol{\varepsilon}, \mathbf{X})$  by Rademacher's theorem. At a point  $(\boldsymbol{\varepsilon}, \mathbf{X})$  where the derivative exist, we obtain a closed-form expression for the gradient (3.9) as follows. The KKT conditions of the optimization problem (3.1) read  $\mathbf{X}^\top(\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}) = \mathbf{X}^\top(\boldsymbol{\varepsilon} - \mathbf{X}\mathbf{h}) = n\nabla g(\widehat{\boldsymbol{\beta}})$ . Differentiation with respect to  $\mathbf{z}_0$  for a fixed  $(\boldsymbol{\varepsilon}, \mathbf{X}, \mathbf{Q}_0)$  as in (4.1) gives

$$\{\mathbf{X}^\top \mathbf{X} + n[\nabla^2 g(\widehat{\boldsymbol{\beta}})]\}(\nabla \widehat{\boldsymbol{\beta}}(\mathbf{z}_0))^\top = \mathbf{a}_0(\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}})^\top - \mathbf{X}^\top \langle \mathbf{a}_0, \mathbf{h} \rangle.$$

By the product rule, this provides the derivative of  $f(\mathbf{z}_0) = \mathbf{X}\mathbf{h} - \boldsymbol{\varepsilon}$ , namely

$$(4.3) \quad \nabla f(\mathbf{z}_0)^\top = \mathbf{X}(\mathbf{X}^\top \mathbf{X} + n\nabla g(\widehat{\boldsymbol{\beta}}))^{-1}[\mathbf{a}_0(\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}})^\top - \langle \mathbf{a}_0, \mathbf{h} \rangle \mathbf{X}^\top] + \mathbf{I}_n \langle \mathbf{a}_0, \mathbf{h} \rangle.$$

Regarding  $\widehat{\mathbf{H}}$  involving differentiation with respect to  $\mathbf{y}$ , the Lipschitz condition of the map  $\mathbf{y} \mapsto \widehat{\boldsymbol{\beta}}$  for strongly convex  $g$  follows from (7.19) in the proof of Proposition 7.3. Hence, the Jacobian in (4.2) exists almost everywhere, and differentiation of the KKT conditions for fixed  $\mathbf{X}$  gives  $(\mathbf{X}^\top \mathbf{X} + n\nabla^2 g(\widehat{\boldsymbol{\beta}}))(\partial/\partial \mathbf{y})\widehat{\boldsymbol{\beta}}(\mathbf{y}, \mathbf{X}) = \mathbf{X}^\top$  so that

$$\widehat{\mathbf{H}} = (\mathbf{X}(\partial/\partial \mathbf{y})\widehat{\boldsymbol{\beta}}(\mathbf{y}, \mathbf{X}))^\top = \mathbf{X}(\mathbf{X}^\top \mathbf{X} + n\nabla^2 g(\widehat{\boldsymbol{\beta}}))^{-1} \mathbf{X}^\top.$$

Identity (4.3) combined with this expression for  $\widehat{\mathbf{H}}$  provides (3.9) with

$$\mathbf{w}_0 = \mathbf{X}\{\mathbf{X}^\top \mathbf{X} + n\nabla g(\widehat{\boldsymbol{\beta}})\}^{-1} \mathbf{a}_0.$$

**4.2. Lasso.** Consider the Lasso  $\widehat{\boldsymbol{\beta}}$  in (3.32). For  $(\boldsymbol{\varepsilon}, \mathbf{X})$  with continuous distribution such as Gaussian under consideration here, almost surely  $\widehat{\boldsymbol{\beta}}$  is unique and

$$(4.4) \quad \mathbf{X}_{\widehat{\mathcal{S}}}^\top(\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}})/n = \lambda \operatorname{sgn}(\widehat{\boldsymbol{\beta}}_{\widehat{\mathcal{S}}}), \quad \|\mathbf{X}_{\widehat{\mathcal{S}}^c}^\top(\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}})/n\|_\infty < \lambda, \quad \operatorname{rank}(\mathbf{X}_{\widehat{\mathcal{S}}}) = |\widehat{\mathcal{S}}|,$$

for the Lasso as in [44, 48] and [7], Proposition 3.9, so that the Jacobian of the mapping  $(\mathbf{z}_0, \boldsymbol{\varepsilon}, \mathbf{X}, \mathbf{Q}_0) \rightarrow \mathbf{X}\widehat{\boldsymbol{\beta}}$  with respect to  $\mathbf{z}_0$  and  $\boldsymbol{\varepsilon}$  can be computed directly by differentiating the KKT condition as in [7, 8, 44]. The following proposition provides closed-form expressions for the gradients for the Lasso estimator, which are valid almost surely and require no assumption on the sparsity of  $\boldsymbol{\beta}$  or the penalty level.

**PROPOSITION 4.1.** *Let  $\lambda > 0$  and consider the Lasso  $\widehat{\boldsymbol{\beta}}$  in (3.32). Let  $\widehat{\mathcal{S}} = \{j \in [p] : \widehat{\beta}_j \neq 0\}$ . For almost every  $(\boldsymbol{\varepsilon}, \mathbf{X}) \in \mathbb{R}^{n \times (p+1)}$ , there exists a neighborhood of  $(\boldsymbol{\varepsilon}, \mathbf{X})$  in which the map  $(\boldsymbol{\varepsilon}, \mathbf{X}) \mapsto \widehat{\mathcal{S}}$  is constant,  $|\widehat{\mathcal{S}}| \leq n$ ,  $\mathbf{X}_{\widehat{\mathcal{S}}}^\top \mathbf{X}_{\widehat{\mathcal{S}}}$  is invertible and the map  $(\boldsymbol{\varepsilon}, \mathbf{X}) \mapsto \widehat{\boldsymbol{\beta}}$  is Lipschitz. In this neighborhood, almost surely  $[\nabla \widehat{\boldsymbol{\beta}}(\mathbf{z}_0)]_{\widehat{\mathcal{S}}^c}^\top = \mathbf{0} \in \mathbb{R}^{n \times |\widehat{\mathcal{S}}^c|}$ ,*

$$[\nabla \widehat{\boldsymbol{\beta}}(\mathbf{z}_0)]_{\widehat{\mathcal{S}}}^\top = (\mathbf{X}_{\widehat{\mathcal{S}}}^\top \mathbf{X}_{\widehat{\mathcal{S}}})^{-1}((\mathbf{a}_0)_{\widehat{\mathcal{S}}}^\top (\mathbf{X}\widehat{\boldsymbol{\beta}} - \mathbf{y})^\top - \mathbf{X}_{\widehat{\mathcal{S}}}^\top \langle \mathbf{a}_0, \mathbf{h} \rangle) \in \mathbb{R}^{n \times |\widehat{\mathcal{S}}|},$$

$$\widehat{\mathbf{H}} = \mathbf{X}_{\widehat{\mathcal{S}}}(\mathbf{X}_{\widehat{\mathcal{S}}}^\top \mathbf{X}_{\widehat{\mathcal{S}}})^{-1} \mathbf{X}_{\widehat{\mathcal{S}}}^\top, \widehat{\mathbf{d}}\mathbf{f} = |\widehat{\mathcal{S}}| \text{ and (3.9) holds with } \mathbf{w}_0 = \mathbf{X}_{\widehat{\mathcal{S}}}(\mathbf{X}_{\widehat{\mathcal{S}}}^\top \mathbf{X}_{\widehat{\mathcal{S}}})^{-1}(\mathbf{a}_0)_{\widehat{\mathcal{S}}}.$$

PROOF OF PROPOSITION 4.1. Proposition 3.9 in [7] proves, for almost every  $(\boldsymbol{\varepsilon}, X)$ , the uniqueness of  $\hat{\boldsymbol{\beta}}$  and (4.4). Let  $(\bar{\boldsymbol{\varepsilon}}, \bar{X}) \in \mathbb{R}^{n \times (p+1)}$  be a point at which (4.4) holds. Let  $\bar{y} = \bar{X}\bar{\boldsymbol{\beta}} + \bar{\boldsymbol{\varepsilon}}$ ,  $\bar{S} = \text{supp}(\hat{\boldsymbol{\beta}}(\bar{y}, \bar{X}))$  and  $s = \bar{X}^\top(\bar{y} - \bar{X}\hat{\boldsymbol{\beta}}(\bar{y}, \bar{X})) / (n\lambda)$ . At  $(y, X) = (\bar{y}, \bar{X})$ , the unique solution of (4.4) is given by the analytic expression

$$\hat{\boldsymbol{\beta}}_{\bar{S}} = (X_{\bar{S}}^\top X_{\bar{S}})^{-1} (X_{\bar{S}}^\top y - n\lambda s_{\bar{S}}), \quad \hat{\boldsymbol{\beta}}_{\bar{S}^c} = \mathbf{0}_{\bar{S}^c}.$$

Moreover, the above expression gives the unique solution of (4.4) in an open neighborhood of  $(\bar{\boldsymbol{\varepsilon}}, \bar{X})$  in  $\mathbb{R}^{n \times (p+1)}$  in which  $\hat{S} = \bar{S}$ ,  $\text{sgn}(\hat{\boldsymbol{\beta}}_{\hat{S}}) = s_{\bar{S}}$  and  $\text{rank}(X_{\hat{S}}) = |\bar{S}|$  are constants. Differentiating this expression immediately yields the formulas for  $\hat{H}$ ,  $\hat{\mathbf{d}}\mathbf{f}$  and  $[\nabla \hat{\boldsymbol{\beta}}(z_0)^\top]_{\hat{S}^c}$ . For  $[\nabla \hat{\boldsymbol{\beta}}(z_0)^\top]_{\hat{S}}$ , differentiating both sides of  $X_{\bar{S}}^\top (X_{\bar{S}}(\hat{\boldsymbol{\beta}} - \boldsymbol{\beta}) - \boldsymbol{\varepsilon}) = -n\lambda s_{\bar{S}}$  yields

$$(a_0)_{\bar{S}}^\top (X_{\bar{S}}\hat{\boldsymbol{\beta}} - y)^\top + X_{\bar{S}}^\top a_0^\top h + X_{\bar{S}}^\top X_{\bar{S}} [\nabla \hat{\boldsymbol{\beta}}(z_0)^\top]_{\bar{S}} = \mathbf{0}$$

due to  $X = XQ_0 + z_0 a_0^\top$ . Finally, the formula for  $w_0$  follows from  $(\partial/\partial z_0)(X\hat{\boldsymbol{\beta}} - y) = X(\partial/\partial z_0)\hat{\boldsymbol{\beta}} + I_n a_0^\top h$  and simple algebra.  $\square$

4.3. *Group Lasso.* Consider a partition  $(G_1, \dots, G_K)$  of  $\{1, \dots, p\}$  and the group Lasso estimator in (3.33). Let  $\hat{B} = \{k \in [K] : \|\hat{\boldsymbol{\beta}}_{G_k}\| \neq 0\}$  be the set of active groups and  $\hat{S} = \bigcup_{k \in \hat{B}} G_k$  the union of all active groups. Define the block diagonal matrix  $M = \text{diag}((M_{G_k, G_k})_{k \in \hat{B}}) \in \mathbb{R}^{|\hat{S}| \times |\hat{S}|}$  by

$$(4.5) \quad M_{G_k, G_k} = n\lambda_k \|\hat{\boldsymbol{\beta}}_{G_k}\|^{-1} (I_{G_k} - \|\hat{\boldsymbol{\beta}}_{G_k}\|^{-2} \hat{\boldsymbol{\beta}}_{G_k} \hat{\boldsymbol{\beta}}_{G_k}^\top), \quad M \in \mathbb{R}^{|\hat{S}| \times |\hat{S}|}.$$

The following proposition provides closed-form expressions for the gradients for the group Lasso estimator and related quantities  $\hat{H}$  and  $w_0$  in terms of  $\hat{S}$  and  $M$ . Its proof is given in Appendix C. Note that the formula for  $\hat{H}$  was known [45].

PROPOSITION 4.2. *The following holds for almost every  $(\bar{y}, \bar{X}) \in \mathbb{R}^{n \times (1+p)}$ . The set  $\bar{B} = \{k \in [K] : \|\bar{\boldsymbol{\beta}}_{G_k}\| > 0\}$  of active groups is the same for all minimizers  $\bar{\boldsymbol{\beta}}$  of (3.33) at  $(\bar{y}, \bar{X})$  and  $\hat{B} = \bar{B}$  for all  $(y, X)$  in a sufficiently small neighborhood of  $(\bar{y}, \bar{X})$ . If additionally  $\bar{X}_{\bar{S}}^\top \bar{X}_{\bar{S}}$  is invertible where  $\bar{S} = \bigcup_{k \in \bar{B}} G_k$  then the map  $(y, X) \mapsto \hat{\boldsymbol{\beta}}$  is Lipschitz in a sufficiently small neighborhood of  $(\bar{y}, \bar{X})$ . In this neighborhood, we have*

$$[\nabla \hat{\boldsymbol{\beta}}(z_0)]_{\hat{S}^c} = 0, \quad [\nabla \hat{\boldsymbol{\beta}}(z_0)]_{\hat{S}}^\top = (X_{\hat{S}}^\top X_{\hat{S}} + M)^{-1} [(a_0)_{\hat{S}}^\top (y - X\hat{\boldsymbol{\beta}})^\top - \langle a_0, h \rangle X_{\hat{S}}^\top],$$

$$\hat{H} = X_{\hat{S}} (X_{\hat{S}}^\top X_{\hat{S}} + M)^{-1} X_{\hat{S}}^\top \text{ and (3.9) holds with } w_0 = X_{\hat{S}} (X_{\hat{S}}^\top X_{\hat{S}} + M)^{-1} (a_0)_{\hat{S}}.$$

**5. Simulations: Lasso and variance spike.** Figure 1 illustrates the variance spike phenomenon of Section 3.7 for the Lasso and  $a_0$  proportional to the first canonical basis vector. The data is generated as follows:  $(s, n, p, \sigma^2) = (200, 750, 1000, 1.0)$ , coefficient vector  $\boldsymbol{\beta}$  is  $s$ -sparse with  $\beta_1 = 20$ ,  $\beta_j = \pm 1$  for  $j = 2, \dots, s$  (independent random signs),  $\beta_j = 0$  for  $j > s$ ; inverse covariance matrix  $\Sigma^{-1} = I_p + 0.9s^{-1/2}(\mathbf{e}_1 \text{sgn}(\boldsymbol{\beta})^\top + \text{sgn}(\boldsymbol{\beta})\mathbf{e}_1^\top)$ , direction  $a_0 = \mathbf{e}_1 / (\mathbf{e}_1^\top \Sigma^{-1} \mathbf{e}_1)^{1/2}$  for  $\mathbf{e}_1 \in \mathbb{R}^p$  the first canonical basis vector. 512 repetitions were used and  $(\Sigma, \boldsymbol{\beta})$  are the same across these repetitions. We see that  $V_0 = \|y - X\hat{\boldsymbol{\beta}}\|^2$  yields an empirical standard deviation (std) substantially larger than 1 (second column), whereas using  $V_0 = \hat{V}(\theta)$  repairs this with an std close to 1 (third column). This choice of  $(\Sigma, \boldsymbol{\beta})$  is a minor modification of [8], Example 2.1 and Figure 2: it is constructed so that  $\langle a_0, \hat{\boldsymbol{\beta}} - \boldsymbol{\beta} \rangle^2$  captures a substantial fraction of the full prediction error  $\|\Sigma^{1/2}(\hat{\boldsymbol{\beta}} - \boldsymbol{\beta})\|^2$ .

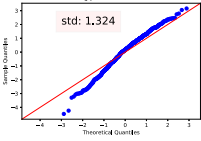
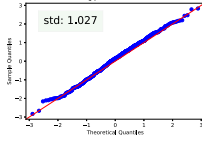
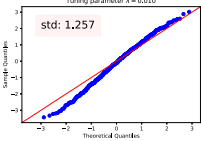
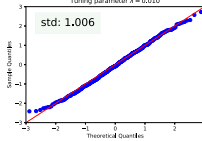
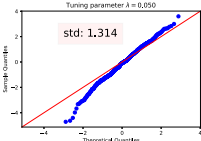
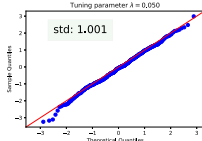
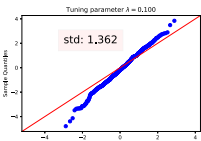
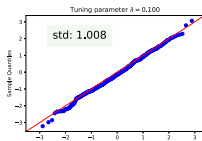
| $\lambda$ | $\frac{(n-\widehat{\text{df}})\langle \mathbf{a}_0, \widehat{\boldsymbol{\beta}}^{(\text{de-bias})} - \boldsymbol{\beta} \rangle}{\ \mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\ }$ | $\frac{(n-\widehat{\text{df}})\langle \mathbf{a}_0, \widehat{\boldsymbol{\beta}}^{(\text{de-bias})} - \boldsymbol{\beta} \rangle}{\widehat{V}(\theta)^{1/2}}$ | $\ \boldsymbol{\Sigma}^{1/2}(\widehat{\boldsymbol{\beta}} - \boldsymbol{\beta})\ ^2$ | $\langle \mathbf{a}_0, \widehat{\boldsymbol{\beta}} - \boldsymbol{\beta} \rangle^2$ |
|-----------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------|
| 0.005     |                                                                                                            |                                                                              | $3.23 \pm 0.45$                                                                      | $0.32 \pm 0.11$                                                                     |
| 0.01      |                                                                                                            |                                                                              | $2.62 \pm 0.41$                                                                      | $0.40 \pm 0.12$                                                                     |
| 0.05      |                                                                                                            |                                                                              | $4.39 \pm 0.86$                                                                      | $1.81 \pm 0.39$                                                                     |
| 0.1       |                                                                                                            |                                                                              | $11.70 \pm 2.40$                                                                     | $5.6 \pm 1.14$                                                                      |

FIG. 1. Standard normal QQ-plots of  $\langle \mathbf{a}_0, \widehat{\boldsymbol{\beta}}^{(\text{de-bias})} - \boldsymbol{\beta} \rangle / V_0^{1/2}$  with  $V_0 = \|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2$  (second column) and with  $V_0 = \widehat{V}(\theta) = \|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2 + (n - \widehat{\text{df}})\langle \mathbf{a}_0, \widehat{\boldsymbol{\beta}} - \boldsymbol{\beta} \rangle^2$  (third column), prediction error  $\|\boldsymbol{\Sigma}^{1/2}(\widehat{\boldsymbol{\beta}} - \boldsymbol{\beta})\|^2$  (fourth column) and squared estimation error in direction  $\mathbf{a}_0$  (fifth column) for the Lasso (3.32) for each  $\lambda \in \{0.005, 0.01, 0.05, 0.1\}$  with the data-generating process described in Section 5.

**6. Simulations: Group Lasso.** Figure 2 illustrates Theorem 3.10 for the group Lasso (3.33) with standard normal QQ-plots of  $\langle \mathbf{a}_0, \widehat{\boldsymbol{\beta}}^{(\text{de-bias})} - \boldsymbol{\beta} \rangle (n - \widehat{\text{df}}) / \|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|$  for  $(n, p, \sigma^2) = (600, 900, 2)$  and the group Lasso (3.33) with 30 nonoverlapping groups each of size 30, where all  $\lambda_k$  in (3.33) are equal to a single parameter  $\lambda$ . The unknown coefficient vector  $\boldsymbol{\beta}$  is the same across all 256 repetitions and has 240 nonzero coefficients, all equal to 1 and belonging to 8 groups (so that the group sparsity of  $\boldsymbol{\beta}$  is 8, and within these 8 groups all coefficients are equal to 1). The design covariance  $\boldsymbol{\Sigma}$  is generated once as  $\boldsymbol{\Sigma} = \mathbf{W} / (5p)$  where  $\mathbf{W}$  has Wishart distribution with covariance  $\mathbf{I}_p$  and  $5p$  degrees-of-freedom. This choice of  $(\boldsymbol{\beta}, \boldsymbol{\Sigma})$  is the same across all 256 repetitions. The direction of interest is  $\mathbf{a}_0 = \mathbf{e}_1 / \|\boldsymbol{\Sigma}^{-1/2} \mathbf{e}_1\|$  where  $\mathbf{e}_1 \in \mathbb{R}^p$  is the first canonical basis vector. The first 8 plots above are standard normal QQ-plots across 256 repetitions for 8 different choices of  $\lambda$ . The ninth plot shows, for each  $\lambda$ , boxplots of  $\widehat{\tau}^2 = (1 - \widehat{\text{df}}/n)^{-2} \|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2 / n$  across the 256 repetitions. This  $\widehat{\tau}$  is proportional to the length of the corresponding confidence interval (3.28) so that the smallest confidence interval (3.28) is achieved for  $\lambda = 0.138$ .

**7. Proof of the main results in Section 3.** In order to prove Theorems 3.6 and 3.7, we apply the bound on the normal approximation in Theorem 2.2. We recall here some notation used throughout the proof. Let  $\widehat{\boldsymbol{\beta}}$  be the estimator (3.1),  $\widehat{\mathbf{H}}$  the gradient of  $\mathbf{y} \rightarrow \mathbf{X}\widehat{\boldsymbol{\beta}}$  as in (1.8),  $\mathbf{a}_0 \in \mathbb{R}^p$  with  $\|\boldsymbol{\Sigma}^{-1/2} \mathbf{a}_0\| = 1$ ,  $\mathbf{z}_0$  and  $\mathbf{Q}_0$  as defined in (1.13),

$$(7.1) \quad \theta = \langle \mathbf{a}_0, \boldsymbol{\beta} \rangle, \quad f(\mathbf{z}_0) = \mathbf{X}\widehat{\boldsymbol{\beta}} - \mathbf{y}, \quad \xi_0 = \mathbf{z}_0^\top f(\mathbf{z}_0) - \text{div } f(\mathbf{z}_0).$$

Vector  $\mathbf{w}_0 \in \mathbb{R}^n$  is given by Lemma 3.1. The oracle  $\boldsymbol{\beta}^*$  and its associated noiseless prediction risk  $R_*$  are given by (3.23). Throughout,  $\mathbb{E}_0$  denotes the conditional expectation given  $(\boldsymbol{\varepsilon}, \mathbf{X}\mathbf{Q}_0)$  and  $\text{Var}_0$  the conditional variance given  $(\boldsymbol{\varepsilon}, \mathbf{X}\mathbf{Q}_0)$ .

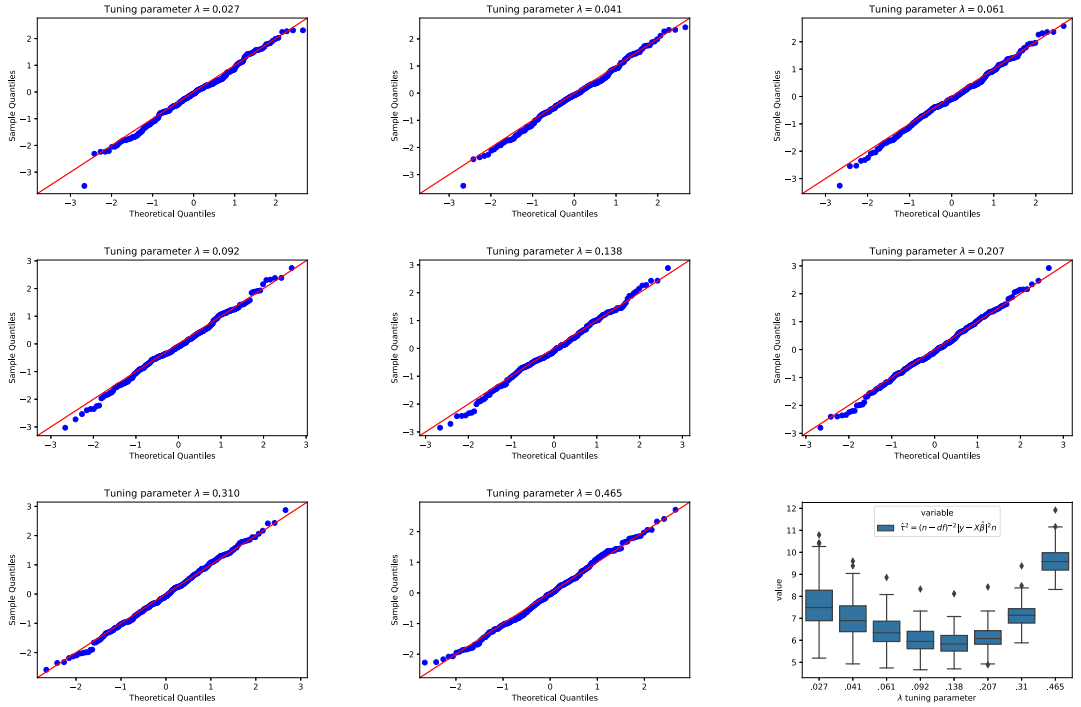


FIG. 2. Standard normal QQ-plots of  $(a_0, \hat{\beta}^{(\text{de-bias})} - \beta)(n - \text{df}) / \|y - X\hat{\beta}\|$  for the group Lasso (3.33). The data-generating process is described in Section 6.

**7.1. Lipschitzness of regularized least-squares and existence of  $w_0$ .** By Rademacher's theorem, a Lipschitz function  $U \rightarrow \mathbb{R}$  for some open set  $U \subset \mathbb{R}^q$  is Fréchet differentiable almost everywhere in  $U$ . The following lemma is the device that verifies the Lipschitz condition for the mappings  $(\epsilon, X) \mapsto \hat{\beta}$  and  $(\epsilon, X) \mapsto Xh - \epsilon$  in certain open set  $U$ , and consequently, differentiability almost everywhere in  $U$ .

**LEMMA 7.1.** Let  $\beta \in \mathbb{R}^p$ ,  $X$  and  $\tilde{X}$  be two design matrices of size  $n \times p$ , and  $\epsilon$  and  $\tilde{\epsilon}$  two noise vectors in  $\mathbb{R}^n$ . Let  $g : \mathbb{R}^p \rightarrow \mathbb{R}$  be convex such that minimizers

$$\hat{\beta} \in \arg \min_{b \in \mathbb{R}^p} \left\{ \frac{\|\epsilon + X(\beta - b)\|^2}{2n} + g(b) \right\}, \quad \tilde{\beta} \in \arg \min_{b \in \mathbb{R}^p} \left\{ \frac{\|\tilde{\epsilon} + \tilde{X}(\beta - b)\|^2}{2n} + g(b) \right\}$$

exist. Let  $h = \hat{\beta} - \beta$ ,  $f = Xh - \epsilon$ ,  $\tilde{h} = \tilde{\beta} - \beta$ ,  $\tilde{f} = \tilde{X}\tilde{h} - \tilde{\epsilon}$ . Let also  $D_g(\tilde{\beta}, \hat{\beta}) = (\tilde{\beta} - \hat{\beta})^\top \{(\partial g)(\tilde{\beta}) - (\partial g)(\hat{\beta})\}$  where  $(\partial g)(\tilde{\beta}) = n^{-1} \tilde{X}^\top (\tilde{\epsilon} - \tilde{X}\tilde{h})$  is the subdifferential at  $\tilde{\beta}$  given by the optimality condition of the above optimization problem and similarly for  $(\partial g)(\hat{\beta})$ , with  $D_g(\hat{\beta}, \tilde{\beta}) \geq 0$  by the monotonicity of the subdifferential. Then

$$\begin{aligned} (7.2) \quad & nD_g(\tilde{\beta}, \hat{\beta}) + \|f - \tilde{f}\|^2 \\ &= (\tilde{h} - h)^\top (X - \tilde{X})^\top f + (\tilde{\epsilon} - \epsilon + (X - \tilde{X})h)^\top (f - \tilde{f}) \\ &= \text{trace}[(X - \tilde{X})^\top (f\tilde{h}^\top - \tilde{f}h^\top)] + (\tilde{\epsilon} - \epsilon)^\top (f - \tilde{f}). \end{aligned}$$

If  $g$  is coercive (i.e.,  $g(x) \rightarrow +\infty$  as  $\|x\| \rightarrow +\infty$ ) then the map  $(\epsilon, X) \mapsto \epsilon - Xh$  is Holder continuous with coefficient 1/2 on every compact. We also have

$$\begin{aligned} (7.3) \quad & nD_g(\tilde{\beta}, \hat{\beta}) + \|X(\hat{\beta} - \tilde{\beta})\|^2/2 + \|\tilde{X}(\hat{\beta} - \tilde{\beta})\|^2/2 \\ &= (\hat{\beta} - \tilde{\beta})^\top (X^\top \epsilon - \tilde{X}^\top \tilde{\epsilon}) + (\hat{\beta} - \tilde{\beta})^\top (X^\top X - \tilde{X}^\top \tilde{X})(h + \tilde{h})/2 \end{aligned}$$

$$(7.4) \quad \leq \|\hat{\beta} - \tilde{\beta}\| \|\epsilon - \tilde{\epsilon}\| \|X + \tilde{X}\|_{\text{op}}/2 \\ + \|\hat{\beta} - \tilde{\beta}\| \|X - \tilde{X}\|_{\text{op}} [(\|\epsilon + \tilde{\epsilon}\|/2 + (\|X\|_{\text{op}} + \|\tilde{X}\|_{\text{op}}))(\|\tilde{h}\| + \|h\|)/2].$$

If either  $g$  is strongly convex or if there exists a constant  $\bar{\kappa} > 0$  and a bounded neighborhood  $\mathcal{N}$  of  $(\bar{\epsilon}, \bar{X})$  such that  $\|\hat{\beta} - \tilde{\beta}\| \bar{\kappa} \leq \|X(\hat{\beta} - \tilde{\beta})\|$  for all  $\{(\epsilon, X), (\tilde{\epsilon}, \tilde{X})\} \subset \mathcal{N}$  then the map  $(\epsilon, X) \mapsto \hat{\beta}$  is Lipschitz in  $\mathcal{N}$ .

PROOF OF LEMMA 7.1. The KKT conditions for  $\hat{\beta}$  and  $\tilde{\beta}$  provide

$$n(\hat{\beta} - \tilde{\beta})^\top (\partial g)(\hat{\beta}) = (\tilde{h} - h)^\top X^\top f, \\ n(\tilde{\beta} - \hat{\beta})^\top (\partial g)(\tilde{\beta}) = (h - \tilde{h})^\top \tilde{X}^\top \tilde{f}.$$

Summing and adding  $\|f - \tilde{f}\|^2 = (\tilde{\epsilon} - \epsilon)^\top (f - \tilde{f}) + (Xh - \tilde{X}\tilde{h})^\top (f - \tilde{f})$  on both sides,

$$nD_g(\hat{\beta}, \tilde{\beta}) + \|f - \tilde{f}\|^2 = \tilde{h}^\top (X - \tilde{X})^\top f + h^\top (\tilde{X} - X)^\top \tilde{f} + (\tilde{\epsilon} - \epsilon)^\top (f - \tilde{f})$$

so that (7.2) holds.

By optimality of  $\hat{\beta}$ ,  $\|f\|^2/(2n) + g(\hat{\beta}) \leq \|X\hat{\beta} + \epsilon\|^2/(2n)$ . If  $g$  is coercive, this implies that for every compact  $K \subset \mathbb{R}^{n \times (1+p)}$ ,  $\|f\| + \|h\|$  and  $\|\tilde{f}\| + \|\tilde{h}\|$  are bounded by a constant depending only on  $g, \beta, n, K$  if  $\{(\epsilon, X), (\tilde{\epsilon}, \tilde{X})\} \subset K$ . In this case, (7.2) implies that  $\|\tilde{f} - f\|^2 \leq (\|\tilde{X} - X\|_F + \|\epsilon - \tilde{\epsilon}\|)C(g, \beta, n, K)$  for some other constant depending on  $g, \beta, n, K$  only. This implies Holder continuity of  $(\epsilon, X) \mapsto \epsilon - Xh$  with Holder coefficient 1/2 on every compact.

For (7.3) and (7.4), the KKT conditions for  $\hat{\beta}$  yield

$$n(\hat{\beta} - \tilde{\beta})^\top (\partial g)(\hat{\beta}) + \|X(\hat{\beta} - \tilde{\beta})\|^2/2 = (\hat{\beta} - \tilde{\beta})^\top X^\top (\epsilon - X(h + \tilde{h})/2).$$

Summing the above and its  $\tilde{\beta}$  counterpart yields the equality (7.3). Writing  $X^\top \epsilon - \tilde{X}^\top \tilde{\epsilon} = (\tilde{X} + X)^\top (\epsilon - \tilde{\epsilon})/2 + (X - \tilde{X})^\top (\tilde{\epsilon} + \epsilon)/2$  and similarly  $X^\top X - \tilde{X}^\top \tilde{X} = (X + \tilde{X})^\top (X - \tilde{X})/2 + (X - \tilde{X})^\top (X + \tilde{X})/2$ , inequality (7.4) follows. To prove the Lipschitz condition in  $\mathcal{N}$ , we note that for a fixed value of  $(\epsilon, X, h)$ , the right-hand side of (7.4) is linear in  $\|\tilde{h}\|$  while the left-hand side is quadratic in  $\|\tilde{h}\|$  thanks to either strong convexity of  $g$  or the assumption on  $\bar{\kappa}$ . This implies that  $\|\tilde{h}\|$  is bounded uniformly for all  $(\tilde{\epsilon}, \tilde{X})$  in  $\mathcal{N}$ . Since  $\epsilon, X, \tilde{\epsilon}, \tilde{X}, \|\tilde{h}\|, \|h\|$  are all bounded in  $\mathcal{N}$ , (7.4) divided by  $\|\hat{\beta} - \tilde{\beta}\|$  provides the desired Lipschitz property.  $\square$

LEMMA 3.1. Let Assumption 3.1(i) be fulfilled,  $\mathbf{a}_0 \in \mathbb{R}^p$  and  $\hat{H}$  be as in (3.3). Then

$$(3.9) \quad \nabla f(\mathbf{z}_0)^\top = (\mathbf{I}_n - \hat{H})^\top \langle \mathbf{a}_0, h \rangle + \mathbf{w}_0(\mathbf{y} - X\hat{\beta})^\top$$

satisfies (3.6) for some random  $\mathbf{w}_0 \in \mathbb{R}^n$  almost surely. If additionally  $\|\Sigma^{-1/2}\mathbf{a}_0\| = 1$ , then

$$(3.10) \quad \|\mathbf{w}_0\|^2 \leq n^{-1} \min\{(4\mu)^{-1}, \phi_{\min}(\Sigma^{-1/2}X^\top X\Sigma^{-1/2}/n)^{-1}\}.$$

PROOF OF LEMMA 3.1. Under Assumption 3.1, Lemma 7.1 implies that the map  $(\epsilon, X) \mapsto f = Xh - \epsilon$  is Lipschitz in an open neighborhood of almost every point, and thus  $\hat{H}$  and  $\nabla f(\mathbf{z}_0)$  are defined as Fréchet derivatives almost surely in (3.3) and (3.6), respectively. To prove (3.9), that is, that the range of  $\nabla f(\mathbf{z}_0) - \langle \mathbf{a}_0, h \rangle (\mathbf{I}_n - \hat{H})$  is the linear span of  $f$ , we study the directional derivative in a direction  $\eta \in \mathbb{R}^n$ . For two pairs  $(\epsilon, X)$  and  $(\tilde{\epsilon}, \tilde{X})$  with  $\tilde{X} = X + t\eta\mathbf{a}_0^\top = X\mathbf{Q}_0 + (t\eta + \mathbf{z}_0)\mathbf{a}_0^\top$  and  $\tilde{\epsilon} = \epsilon + t\eta\langle \mathbf{a}_0, h \rangle$ , consider the solutions  $\hat{\beta}$  and  $\tilde{\beta}$  defined in Lemma 7.1 and  $\phi_t = \tilde{X}(\tilde{\beta} - \beta) - \tilde{\epsilon}$  with  $\phi_0 = X(\hat{\beta} - \beta) - \epsilon = f$ . When the map  $(\epsilon, X) \mapsto f$  is Fréchet differentiable at  $(\epsilon, X)$ ,

$$(7.5) \quad \lim_{t \rightarrow 0+} (\phi_t - \phi_0)/t = (\nabla f(\mathbf{z}_0) - \langle \mathbf{a}_0, h \rangle (\mathbf{I}_n - \hat{H}))^\top \eta$$

by the chain rule and the linearity of the Fréchet derivative, noting that  $(\partial/\partial \boldsymbol{\varepsilon})(\boldsymbol{\varepsilon} - \mathbf{X}\mathbf{h}) = \mathbf{I}_n - \tilde{\mathbf{H}}$ . For this specific choice of  $(\tilde{\boldsymbol{\varepsilon}}, \tilde{\mathbf{X}})$ , we have

$$(7.6) \quad (\tilde{\mathbf{X}} - \mathbf{X})\mathbf{h} + \boldsymbol{\varepsilon} - \tilde{\boldsymbol{\varepsilon}} = 0.$$

It follows that the second term in the first line of (7.2) is zero, so that (7.2) gives

$$\mu n \|\boldsymbol{\Sigma}^{1/2}(\hat{\boldsymbol{\beta}} - \tilde{\boldsymbol{\beta}})\|^2 + \|\boldsymbol{\phi}_0 - \boldsymbol{\phi}_t\|^2 \leq |\langle \mathbf{a}_0, \mathbf{h} - \tilde{\mathbf{h}} \rangle t \boldsymbol{\eta}^\top \mathbf{f}|$$

due to  $\tilde{\mathbf{X}} - \mathbf{X} = t\boldsymbol{\eta}\mathbf{a}_0^\top$ . Consequently,  $\boldsymbol{\phi}_t - \boldsymbol{\phi}_0 = 0$  when  $\boldsymbol{\eta}^\top \mathbf{f} = 0$ . This and (7.5) give (3.9). Moreover, for  $\mathbf{f} \neq \mathbf{0}$ ,  $\mathbf{w}_0 = \lim_{t \rightarrow 0}(\boldsymbol{\phi}_t - \boldsymbol{\phi}_0)/t$  for  $\boldsymbol{\eta} = -\mathbf{f}/\|\mathbf{f}\|^2$ , so that (3.10) is an upper bound for  $\lim_{t \rightarrow 0+} \|\boldsymbol{\phi}_t - \boldsymbol{\phi}_0\|/t$  in the case of  $\|\boldsymbol{\Sigma}^{-1/2}\mathbf{a}_0\| = 1 = -\boldsymbol{\eta}^\top \mathbf{f}$  where

$$\mu n \|\boldsymbol{\Sigma}^{1/2}(\hat{\boldsymbol{\beta}} - \tilde{\boldsymbol{\beta}})\|^2 + \|\boldsymbol{\phi}_0 - \boldsymbol{\phi}_t\|^2 \leq |t| \|\boldsymbol{\Sigma}^{1/2}(\mathbf{h} - \tilde{\mathbf{h}})\|$$

by the previous display. For  $\mu > 0$ , the above inequality gives  $\|\boldsymbol{\phi}_0 - \boldsymbol{\phi}_t\|^2 \leq t^2(4\mu n)^{-1}$  using  $uv \leq u^2/4 + v^2$ . For  $\mu = 0$ ,

$$\phi_{\min}(\tilde{\mathbf{W}}) \|\boldsymbol{\Sigma}^{1/2}(\mathbf{h} - \tilde{\mathbf{h}})\|^2 \leq \|\tilde{\mathbf{X}}^\top(\mathbf{h} - \tilde{\mathbf{h}})\|^2 = \|\boldsymbol{\phi}_t - \boldsymbol{\phi}_0\|^2$$

with  $\phi_{\min}(\tilde{\mathbf{W}})$  being the smallest eigenvalue of  $\tilde{\mathbf{W}} = \boldsymbol{\Sigma}^{-1/2}\tilde{\mathbf{X}}^\top\tilde{\mathbf{X}}\boldsymbol{\Sigma}^{-1/2}$ . Hence, (3.10) holds in either cases.  $\square$

**7.2. Loss equivalence to oracle estimators.** To apply Theorem 2.2 with respect to  $\mathbf{z}_0$  to  $\mathbf{f}$  in (7.1), we will need to control expectations involving  $\|\mathbf{w}_0\|$ ,  $\langle \mathbf{a}_0, \mathbf{h} \rangle$ ,  $\|\mathbf{X}\mathbf{h}\|$  and  $\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|$ . To this end, define the random variables  $F_+$  and  $F$  by

$$(7.7) \quad F_+ \stackrel{\text{def}}{=} (\|\mathbf{g}\|^2/n) \vee (\|\boldsymbol{\varepsilon}\|^2/(\sigma^2 n)) \vee (\|\boldsymbol{\varepsilon} - \mathbf{X}\mathbf{h}^*\|^2/(nR_*)) \vee 1$$

with  $\mathbf{g} = \mathbf{X}\mathbf{h}^*/\|\boldsymbol{\Sigma}^{1/2}\mathbf{h}^*\|$  and the  $\mathbf{h}^*$  and  $R_*$  in (3.23), and

$$(7.8) \quad F \stackrel{\text{def}}{=} 2/[1 \wedge \max\{\mu, \phi_{\min}(\boldsymbol{\Sigma}^{-1/2}(\mathbf{X}^\top \mathbf{X}/n)\boldsymbol{\Sigma}^{-1/2})\}].$$

We note that the three random vectors  $\boldsymbol{\varepsilon}/\sigma$ ,  $\mathbf{g}$  and  $(\boldsymbol{\varepsilon} - \mathbf{X}\mathbf{h}^*)/R_*$  have  $N(\mathbf{0}, \mathbf{I}_n)$  distribution, so that  $F_+$  is of the form  $F_+ = \max_{i=1,2,3} W_i/n$  where each  $W_i$  has the  $\chi_n^2$  distribution. Thus, by Proposition A.1 and properties of the  $\chi_n^2$  distribution,

$$(7.9) \quad \mathbb{E}[F_+^4] \vee \mathbb{E}[F_+^{10}] \leq C(\gamma, \mu), \quad \mathbb{E}[(F_+ - 1)^2] \leq 3 \text{Var}[\chi_n^2]/n^2 = 6/n.$$

It follows from (1.15), Lemma 7.2 below and (3.10) that almost surely

$$(7.10) \quad \langle \mathbf{a}_0, \mathbf{h} \rangle^2 \leq \|\boldsymbol{\Sigma}^{1/2}\mathbf{h}\|^2 \vee (\|\mathbf{X}\mathbf{h}\|^2/n) \leq F_+ F^2 R_*, \quad \|\mathbf{w}_0\|^2 \leq F/(2n),$$

$$(7.11) \quad \|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2/n \leq 2F_+ + 2F_+ F^2 R_* \leq 4F_+ F^2 R_*,$$

for the  $\mathbf{w}_0$  in Lemma 3.1. The moment inequalities in (7.9) and the almost sure bounds (7.10)–(7.11) allow us to control expectations involving  $\|\mathbf{w}_0\|$ ,  $\langle \mathbf{a}_0, \mathbf{h} \rangle$ ,  $\|\mathbf{X}\mathbf{h}\|$  and  $\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|$  throughout the proofs. The following lemma provides the first inequality in (7.10).

**LEMMA 7.2 (Deterministic lemma).** *Consider the linear model (1.1) and a convex penalty  $g(\cdot)$ . Let  $\hat{\boldsymbol{\beta}}$  in (3.1) and  $\boldsymbol{\beta}^*$ ,  $\mathbf{h}^*$ ,  $R_*$  be defined in (3.23). Suppose the penalty satisfies  $\mathbf{u}^\top \{(\partial g)(\mathbf{u} + \boldsymbol{\beta}^*) - (\partial g)(\boldsymbol{\beta}^*)\} \geq \mu \|\boldsymbol{\Sigma}^{1/2}\mathbf{u}\|^2 \forall \mathbf{u} \in \mathbb{R}^p$  with  $\mu \in [0, 1/2]$ . Let  $F_+$  be defined in (7.7) and  $F$  any random variable satisfying*

$$(7.12) \quad 1 \leq F/2 \quad \text{and either} \quad \|\boldsymbol{\Sigma}^{1/2}\mathbf{h}\|^2/(n\|\mathbf{X}\mathbf{h}\|^2) \leq F/2 \quad \text{or} \quad \mu^{-1} = F/2,$$

for instance (7.8). Then

$$(7.13) \quad \|\boldsymbol{\Sigma}^{1/2}\mathbf{h}\|^2 \leq F^2 \max(\bar{\sigma}^2, \|\boldsymbol{\Sigma}^{1/2}\mathbf{h}^*\|^2) \leq F_+ F^2 R_*,$$

$$(7.14) \quad \|\mathbf{X}\mathbf{h}\|^2/n \leq \max\{2F\bar{\sigma}^2, \bar{\sigma}^2 + F^2\|\boldsymbol{\Sigma}^{1/2}\mathbf{h}^*\|^2\} \leq F_+ F^2 R_*,$$

where  $\bar{\sigma}^2 = F_+ \sigma^2 + (F_+ - 1)\|\boldsymbol{\Sigma}^{1/2}\mathbf{h}^*\|^2 = (F_+ - 1)R_* + \sigma^2$ .

PROOF OF LEMMA 7.2. The KKT conditions for  $\widehat{\boldsymbol{\beta}}$ , that is,  $n\partial g(\widehat{\boldsymbol{\beta}}) = \mathbf{X}^\top(\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}})$ , yield

$$\begin{aligned} 2(\widehat{\boldsymbol{\beta}} - \boldsymbol{\beta}^*)^\top (\partial g)(\widehat{\boldsymbol{\beta}}) &= 2(\widehat{\boldsymbol{\beta}} - \boldsymbol{\beta}^*)^\top \mathbf{X}^\top(\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}})/n \\ &= (\|\mathbf{y} - \mathbf{X}\boldsymbol{\beta}^*\|^2 - \|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2 - \|\mathbf{X}(\widehat{\boldsymbol{\beta}} - \boldsymbol{\beta}^*)\|^2)/n \\ &= (\|\mathbf{X}\mathbf{h}^*\|^2 - \|\mathbf{X}\mathbf{h}\|^2 - \|\mathbf{X}(\widehat{\boldsymbol{\beta}} - \boldsymbol{\beta}^*)\|^2 + 2\boldsymbol{\varepsilon}^\top \mathbf{X}(\mathbf{h} - \mathbf{h}^*))/n \\ &\leq (\|\mathbf{X}\mathbf{h}^*\|^2 - \|\mathbf{X}\mathbf{h}\|^2 + \|\boldsymbol{\varepsilon}\|^2)/n. \end{aligned}$$

Similarly, the KKT conditions  $-\boldsymbol{\Sigma}\mathbf{h}^* = \partial g(\boldsymbol{\beta}^*)$  for  $\boldsymbol{\beta}^*$  yield

$$(7.15) \quad 2(\boldsymbol{\beta}^* - \widehat{\boldsymbol{\beta}})^\top (\partial g)(\boldsymbol{\beta}^*) + \|\boldsymbol{\Sigma}^{1/2}(\widehat{\boldsymbol{\beta}} - \boldsymbol{\beta}^*)\|^2 \leq \|\boldsymbol{\Sigma}^{1/2}\mathbf{h}\|^2 - \|\boldsymbol{\Sigma}^{1/2}\mathbf{h}^*\|^2.$$

Summing the two above displays yields

$$\begin{aligned} (7.16) \quad &(1 + 2\mu)\|\boldsymbol{\Sigma}^{1/2}(\widehat{\boldsymbol{\beta}} - \boldsymbol{\beta}^*)\|^2 \\ &\leq (F_+ - 1)\|\boldsymbol{\Sigma}^{1/2}\mathbf{h}^*\|^2 + \|\boldsymbol{\Sigma}^{1/2}\mathbf{h}\|^2 - \|\mathbf{X}\mathbf{h}\|^2/n + F_+\sigma^2 \\ &= \bar{\sigma}^2 + \|\boldsymbol{\Sigma}^{1/2}\mathbf{h}\|^2 - \|\mathbf{X}\mathbf{h}\|^2/n. \end{aligned}$$

For  $\|\boldsymbol{\Sigma}^{1/2}\mathbf{h}\| \geq F\|\boldsymbol{\Sigma}^{1/2}\mathbf{h}^*\|$ , by the triangle inequality

$$(7.17) \quad \|\boldsymbol{\Sigma}^{1/2}\mathbf{h}\|^2(1 - 1/F)^2 \leq \|\boldsymbol{\Sigma}^{1/2}(\widehat{\boldsymbol{\beta}} - \boldsymbol{\beta}^*)\|^2$$

provides a lower bound on the left-hand side of (7.16) so that

$$\begin{aligned} (7.18) \quad \bar{\sigma}^2 &\geq \begin{cases} \{(1 - 1/F)^2 + 2/F - 1\}\|\boldsymbol{\Sigma}^{1/2}\mathbf{h}\|^2 & \text{if } \|\mathbf{X}\mathbf{h}\|^2/n \geq (2/F)\|\boldsymbol{\Sigma}^{1/2}\mathbf{h}\|^2, \\ \{(1 - 1/F)^2(1 + 2\mu) - 1\}\|\boldsymbol{\Sigma}^{1/2}\mathbf{h}\|^2 & \text{if } \|\mathbf{X}\mathbf{h}\|^2/n < (2/F)\|\boldsymbol{\Sigma}^{1/2}\mathbf{h}\|^2, \end{cases} \\ &\geq F^{-2}\|\boldsymbol{\Sigma}^{1/2}\mathbf{h}\|^2 \end{aligned}$$

due to  $F = 2/\mu \geq 4$  in the second case. This gives (7.13). For  $\|\boldsymbol{\Sigma}^{1/2}\mathbf{h}\| \geq F\|\boldsymbol{\Sigma}^{1/2}\mathbf{h}^*\|$  by (7.16), (7.17) and (7.18) we have

$$\|\mathbf{X}\mathbf{h}\|^2/n \leq \bar{\sigma}^2 + \|\boldsymbol{\Sigma}^{1/2}\mathbf{h}\|^2\{1 - (1 - 1/F)^2\} \leq \bar{\sigma}^2 + F^2\bar{\sigma}^2\{1 - (1 - 1/F)^2\} = 2F\bar{\sigma}^2,$$

and for  $\|\boldsymbol{\Sigma}^{1/2}\mathbf{h}\| < F\|\boldsymbol{\Sigma}^{1/2}\mathbf{h}^*\|$  we have  $\|\mathbf{X}\mathbf{h}\|^2/n \leq \bar{\sigma}^2 + F^2\|\boldsymbol{\Sigma}^{1/2}\mathbf{h}^*\|^2$ , and thus (7.14) holds.  $\square$

### 7.3. Existence and properties of $\widehat{\mathbf{H}}$ and $\widehat{\text{df}}$ .

PROPOSITION 7.3. Let  $\mathbf{X} \in \mathbb{R}^{n \times p}$  be any fixed design matrix, and  $\widehat{\boldsymbol{\beta}}(\mathbf{y}) = \arg \min_{\mathbf{b} \in \mathbb{R}^p} \{\|\mathbf{y} - \mathbf{X}\mathbf{b}\|^2/(2n) + g(\mathbf{b})\}$ . Then the following statements hold:

(i)  $\|\mathbf{X}(\widehat{\boldsymbol{\beta}}(\mathbf{y}) - \widehat{\boldsymbol{\beta}}(\tilde{\mathbf{y}}))\| \leq \|\mathbf{y} - \tilde{\mathbf{y}}\|$  for all  $\mathbf{y}, \tilde{\mathbf{y}} \in \mathbb{R}^n$ , that is,  $\mathbf{y} \mapsto \mathbf{X}\widehat{\boldsymbol{\beta}}(\mathbf{y})$  is 1-Lipschitz. Its gradient  $\widehat{\mathbf{H}}$  exists almost everywhere by Rademacher's theorem, that is, for almost every  $\mathbf{y}$  there exists  $\widehat{\mathbf{H}} \in \mathbb{R}^{n \times n}$  with  $\|\widehat{\mathbf{H}}\|_{\text{op}} \leq 1$  such that  $\mathbf{X}\widehat{\boldsymbol{\beta}}(\tilde{\mathbf{y}}) = \mathbf{X}\widehat{\boldsymbol{\beta}}(\mathbf{y}) + \widehat{\mathbf{H}}^\top \boldsymbol{\eta} + o(\|\boldsymbol{\eta}\|)$ .

(ii) For almost every  $\mathbf{y}$ , matrix  $\widehat{\mathbf{H}}$  is symmetric with eigenvalues in  $[0, 1]$ . Consequently, with  $\widehat{\text{df}} = \text{trace}(\widehat{\mathbf{H}})$  as degrees-of-freedom,  $(n - \widehat{\text{df}})(1 - \widehat{\text{df}}/n) \leq \|\mathbf{I}_n - \widehat{\mathbf{H}}\|_F^2 \leq n - \widehat{\text{df}}$ .

PROOF. A proof of (i) is given in [2]. For completeness, the argument is the following: by (7.3) with  $\tilde{\mathbf{X}} = \mathbf{X}$ ,  $\tilde{\boldsymbol{\varepsilon}} = \tilde{\mathbf{y}} - \mathbf{X}\boldsymbol{\beta}$  and  $\boldsymbol{\varepsilon} = \mathbf{y} - \mathbf{X}\boldsymbol{\beta}$  we have

$$(7.19) \quad nD_g(\widehat{\boldsymbol{\beta}}(\tilde{\mathbf{y}}), \widehat{\boldsymbol{\beta}}(\mathbf{y})) + \|\mathbf{X}\widehat{\boldsymbol{\beta}}(\tilde{\mathbf{y}}) - \mathbf{X}\widehat{\boldsymbol{\beta}}(\mathbf{y})\|^2 \leq (\mathbf{y} - \tilde{\mathbf{y}})^\top \mathbf{X}(\widehat{\boldsymbol{\beta}}(\mathbf{y}) - \widehat{\boldsymbol{\beta}}(\tilde{\mathbf{y}})).$$

Using  $D_g(\widehat{\boldsymbol{\beta}}(\tilde{\mathbf{y}}), \widehat{\boldsymbol{\beta}}(\mathbf{y})) \geq 0$  by monotonicity of the subdifferential and the Cauchy–Schwarz inequality yields the desired Lipschitz property. For (ii), define

$$\begin{aligned} u(\mathbf{y}) &= (\|\mathbf{y}\|^2 - \|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}(\mathbf{y})\|^2)/2 - ng(\widehat{\boldsymbol{\beta}}(\mathbf{y})) \\ &= \sup_{\mathbf{b} \in \mathbb{R}^p} \{\mathbf{y}^\top \mathbf{X}\mathbf{b} - \|\mathbf{X}\mathbf{b}\|^2/2 - ng(\mathbf{b})\}. \end{aligned}$$

The function  $u : \mathbb{R}^n \rightarrow \mathbb{R}$  is convex in  $\mathbf{y}$  as a supremum of affine functions, and  $\mathbf{X}\widehat{\boldsymbol{\beta}}(\mathbf{y})$  is a subgradient of  $u$  at  $\mathbf{y}$ . Alexandrov's theorem as stated in [35], Theorem D.2.1, states that any convex  $u : \mathbb{R}^n \rightarrow \mathbb{R}$  is twice differentiable at  $\mathbf{y}$  for almost every  $\mathbf{y}$  in the following sense:  $u$  is Fréchet differentiable at  $\mathbf{y}$  with gradient  $\nabla u(\mathbf{y})$  and there exists a symmetric positive semidefinite matrix  $\mathbf{S}$  such that for every  $\varepsilon > 0$  there exists  $\delta > 0$  such that for all  $\tilde{\mathbf{y}} \in \mathbb{R}^n$ ,

$$\|\tilde{\mathbf{y}} - \mathbf{y}\| \leq \delta \quad \text{implies} \quad \sup_{\mathbf{v} \in \partial u(\tilde{\mathbf{y}})} \|\mathbf{v} - \nabla u(\mathbf{y}) - \mathbf{S}(\tilde{\mathbf{y}} - \mathbf{y})\| \leq \varepsilon \|\tilde{\mathbf{y}} - \mathbf{y}\|.$$

By (i) and the definition of  $\widehat{\mathbf{H}}$ , for almost every  $\mathbf{y}$  it holds that  $\mathbf{X}\widehat{\boldsymbol{\beta}}(\tilde{\mathbf{y}}) = \mathbf{X}\widehat{\boldsymbol{\beta}}(\mathbf{y}) + \widehat{\mathbf{H}}^\top(\tilde{\mathbf{y}} - \mathbf{y}) + o(\|\tilde{\mathbf{y}} - \mathbf{y}\|)$ . Combining these two results and taking  $\mathbf{v} = \mathbf{X}\widehat{\boldsymbol{\beta}}(\tilde{\mathbf{y}})$ , we get that  $\mathbf{S} = \widehat{\mathbf{H}}$  for almost every  $\mathbf{y}$ .  $\square$

LEMMA 7.4. *Let Assumption 3.1 be fulfilled with  $n \geq 2$ . Then there exists an event  $\Omega_0$  independent of  $(\mathbf{z}_0, \boldsymbol{\varepsilon})$  such that*

$$\Omega_0 \subset \{n - \widehat{\text{df}} \geq \|\mathbf{I}_n - \widehat{\mathbf{H}}\|_F^2 \geq C_*(\gamma, \mu)n\} \quad \text{with} \quad \begin{cases} \mathbb{P}(\Omega_0^c) = 0 & \text{if } \gamma < 1, \\ \mathbb{P}(\Omega_0^c) \leq e^{-n/2} & \text{if } \gamma \geq 1, \end{cases}$$

where  $C_*(\gamma, \mu) \in (0, 1)$  depends on  $\{\gamma, \mu\}$  only.

PROOF OF LEMMA 7.4. If  $\gamma < 1$ , the choice  $C_*(\gamma, \mu) = (1 - \gamma)$  works with probability one because  $\text{rank}(\widehat{\mathbf{H}}) \leq \text{rank}(\mathbf{X}) \leq p \leq \gamma n$  and  $\|\widehat{\mathbf{H}}\|_{\text{op}} \leq 1$ .

If  $\gamma \geq 1$ , then we have  $\mu > 0$  in Assumption 3.1. Let  $\Omega_0 = \{\|\mathbf{X}\mathbf{Q}_0\boldsymbol{\Sigma}^{-1/2}\|_{\text{op}} \leq \sqrt{p} + 2\sqrt{n}\}$ . By [20], Theorem II.13,  $\mathbb{P}(\Omega_0) \geq 1 - e^{-n/2}$  due to  $\mathbf{X}\mathbf{Q}_0\boldsymbol{\Sigma}^{-1/2} = \mathbf{X}\boldsymbol{\Sigma}^{-1/2}(\mathbf{I}_p - (\boldsymbol{\Sigma}^{-1/2}\mathbf{a}_0)(\boldsymbol{\Sigma}^{-1/2}\mathbf{a}_0)^\top)$  with  $\|\boldsymbol{\Sigma}^{-1/2}\mathbf{a}_0\| = 1$ . Next, we hold  $\mathbf{X}$  fixed and study the derivatives of  $\mathbf{X}\widehat{\boldsymbol{\beta}}$  with respect to  $\mathbf{y}$ . Let  $\widehat{\boldsymbol{\beta}}(\mathbf{y})$  be as in Proposition 7.3. Let  $\mathbf{P} = \mathbf{I}_n - \mathbf{z}_0\mathbf{z}_0^\top/\|\mathbf{z}_0\|^2$  be the projection onto  $\{\mathbf{z}_0\}^\perp$  so that  $\mathbf{P}\mathbf{X} = \mathbf{P}\mathbf{X}\mathbf{Q}_0$ . Let  $\mathbf{y}, \tilde{\mathbf{y}}$  be such that  $\mathbf{z}_0^\top(\mathbf{y} - \tilde{\mathbf{y}}) = 0$ , or equivalently  $\mathbf{P}(\mathbf{y} - \tilde{\mathbf{y}}) = \mathbf{y} - \tilde{\mathbf{y}}$ . By (7.19) and (3.2),

$$\begin{aligned} n\mu\|\boldsymbol{\Sigma}^{1/2}(\widehat{\boldsymbol{\beta}}(\mathbf{y}) - \widehat{\boldsymbol{\beta}}(\tilde{\mathbf{y}}))\|^2 + \|\mathbf{X}(\widehat{\boldsymbol{\beta}}(\mathbf{y}) - \widehat{\boldsymbol{\beta}}(\tilde{\mathbf{y}}))\|^2 &\leq (\mathbf{y} - \tilde{\mathbf{y}})^\top \mathbf{X}(\widehat{\boldsymbol{\beta}}(\mathbf{y}) - \widehat{\boldsymbol{\beta}}(\tilde{\mathbf{y}})) \\ &= (\mathbf{y} - \tilde{\mathbf{y}})^\top \mathbf{P}\mathbf{X}(\widehat{\boldsymbol{\beta}}(\mathbf{y}) - \widehat{\boldsymbol{\beta}}(\tilde{\mathbf{y}})). \end{aligned}$$

On  $\Omega_0$ ,  $\mu(\sqrt{\gamma} + 2)^{-2}\|\mathbf{X}\mathbf{Q}_0(\widehat{\boldsymbol{\beta}}(\mathbf{y}) - \widehat{\boldsymbol{\beta}}(\tilde{\mathbf{y}}))\|^2 \leq n\mu\|\boldsymbol{\Sigma}^{1/2}(\widehat{\boldsymbol{\beta}}(\mathbf{y}) - \widehat{\boldsymbol{\beta}}(\tilde{\mathbf{y}}))\|^2$ . Combined with the above display, this implies  $(1 + \mu(\sqrt{\gamma} + 2)^{-2})\|\mathbf{P}\mathbf{X}(\widehat{\boldsymbol{\beta}}(\mathbf{y}) - \widehat{\boldsymbol{\beta}}(\tilde{\mathbf{y}}))\| \leq \|\mathbf{P}(\mathbf{y} - \tilde{\mathbf{y}})\|$ . With  $\tilde{\mathbf{y}} = \mathbf{y} + \mathbf{P}\boldsymbol{\eta}$  and by definition of  $\widehat{\mathbf{H}}$ , we have  $L^{-1}\|\mathbf{P}\widehat{\mathbf{H}}\mathbf{P}\boldsymbol{\eta}\| \leq \|\mathbf{P}\boldsymbol{\eta}\| + o(\|\boldsymbol{\eta}\|)$  for  $L = (1 + \mu(\sqrt{\gamma} + 2)^{-2})^{-1}$ , hence  $\|\mathbf{P}\widehat{\mathbf{H}}\mathbf{P}\|_{\text{op}} \leq L$ . Since  $\text{rank}(\mathbf{P}) = n - 1$ , by Cauchy's interlacing theorem  $\mathbf{I}_n - \widehat{\mathbf{H}}$  has at least  $n - 1$  eigenvalues no smaller than  $1 - L > 0$ . Finally, since  $\mathbf{I}_n - \widehat{\mathbf{H}}$  is symmetric with eigenvalues in  $[0, 1]$  by Proposition 7.3,

$$n - \widehat{\text{df}} = \text{trace}[\mathbf{I}_n - \widehat{\mathbf{H}}] \geq \|\mathbf{I}_n - \widehat{\mathbf{H}}\|_F^2 \geq (n - 1)(1 - L)^2 \geq nC_*$$

with  $C_* = (1 - L)^2/2$  thanks to  $n \geq 2$ .  $\square$

7.4. *Lower bound on  $\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2/n$ .* The following lemmas are useful to bound from below the denominator in (3.24).

LEMMA 7.5. *Let Assumption 3.1 be fulfilled. Then  $\mathbb{E}[\xi_0^2] \leq C_6(\gamma, \mu)nR_*$  and*

$$(7.20) \quad \mathbb{E}[(1 - \widehat{\mathbf{d}}\mathbf{f}/n)\langle \mathbf{a}_0, \mathbf{h} \rangle + n^{-1}\langle \mathbf{z}_0, \mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}} \rangle]^2/R_* \leq C_7(\gamma, \mu)n^{-1},$$

$$(7.21) \quad \mathbb{E}[I_{\Omega_0}(\langle \mathbf{a}_0, \mathbf{h} \rangle + (n - \widehat{\mathbf{d}}\mathbf{f})^{-1}\langle \mathbf{z}_0, \mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}} \rangle)^2]/R_* \leq C_8(\gamma, \mu)n^{-1},$$

where  $\Omega_0$  is the event from Lemma 7.4.

PROOF OF LEMMA 7.5. By (3.9) combined with (3.10), (7.11) and (7.9),

$$(7.22) \quad \mathbb{E}[\langle \mathbf{w}_0, \mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}} \rangle^2] \leq \mathbb{E}[\|\mathbf{w}_0\|^2\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2] \leq \mathbb{E}[(F/(2n))4nF_+F^2] \leq C_9(\gamma, \mu)R_*$$

Similarly, by definition of  $V^*(\theta)$  in (3.18),  $\mathbb{E}[\xi_0^2] = \mathbb{E}[V^*(\theta)] = \mathbb{E}[\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2 + \|\nabla f(\mathbf{z}_0)\|_F^2]$ . Using (7.10)–(7.11), we have  $\mathbb{E}[\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2] \leq 4nR_*\mathbb{E}[F_+F^2]$  and

$$(7.23) \quad \begin{aligned} \mathbb{E}[\|\nabla f(\mathbf{z}_0)\|_F^2] &\leq \mathbb{E}[2\|\mathbf{I}_n - \hat{\mathbf{H}}\|_F^2\langle \mathbf{a}_0, \mathbf{h} \rangle^2 + 2\langle \mathbf{w}_0, \mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}} \rangle^2] \\ &\leq n\mathbb{E}[2\langle \mathbf{a}_0, \mathbf{h} \rangle^2] + C_{10}(\gamma, \mu)R_* \\ &\leq nR_*C_{11}(\gamma, \mu) \end{aligned}$$

thanks to  $\nabla f(\mathbf{z}_0)$  in (3.9), and  $(a + b)^2 \leq 2(a^2 + b^2)$  for the first inequality,  $\|\mathbf{I}_n - \hat{\mathbf{H}}\|_F^2 \leq n$  by Proposition 7.3 and (7.22) for the second inequality, and (7.10)–(7.9) for the third inequality. This provides  $\mathbb{E}[\xi_0^2] \leq C_{12}(\gamma, \mu)nR_*$ . Next, (7.20) holds due the bound (7.22) and the relationship in (3.11) between  $\xi_0$ ,  $\langle \mathbf{w}_0, \mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}} \rangle$  and the integrand in left-hand side of (7.20). Then (7.21) follows from (7.20) and  $I_{\Omega_0}(1 - \widehat{\mathbf{d}}\mathbf{f}/n)^{-2} \leq C_*(\gamma, \mu)^{-2}$  by Lemma 7.4.  $\square$

LEMMA 7.6. *Let  $\hat{\boldsymbol{\beta}}$  be as in (3.1) for convex  $g$  and let  $\boldsymbol{\beta}^*$ ,  $\mathbf{h}^*$ ,  $R_*$  be as in (3.23). Then*

$$(7.24) \quad (1 - \widehat{\mathbf{d}}\mathbf{f}/n)^2/8 \leq \|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2/(nR_*) + \Delta_n^a + \Delta_n^b + \Delta_n^c$$

$$(7.25) \quad \leq V^*(\theta)/(nR_*) + \Delta_n^d + \Delta_n^a + \Delta_n^b + \Delta_n^c$$

where  $V^*(\theta)$  is defined in (3.18) and  $\Delta_n^a, \dots, \Delta_n^d$  are nonnegative terms defined as

$$(7.26) \quad \Delta_n^a \stackrel{\text{def}}{=} \sigma^2 |(1 - \widehat{\mathbf{d}}\mathbf{f}/n) - \boldsymbol{\varepsilon}^\top (\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}})/(n\sigma^2)|^2/R_*,$$

$$(7.27) \quad \Delta_n^b \stackrel{\text{def}}{=} (F_+ - 1)_+ \|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2/(nR_*),$$

$$(7.28) \quad \Delta_n^c \stackrel{\text{def}}{=} |(1 - \widehat{\mathbf{d}}\mathbf{f}/n)\langle \mathbf{a}_*, \mathbf{h} \rangle - \mathbf{g}^\top (\mathbf{X}\hat{\boldsymbol{\beta}} - \mathbf{y})/n|^2/R_*$$

$$\Delta_n^d \stackrel{\text{def}}{=} n^{-1} |2\mathbf{w}_0^\top (\mathbf{I}_n - \hat{\mathbf{H}})(\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}})\langle \mathbf{a}_0, \mathbf{h} \rangle|/R_*.$$

where  $\mathbf{g} = \mathbf{X}\mathbf{h}^*/\|\boldsymbol{\Sigma}^{1/2}\mathbf{h}^*\|$  and  $\mathbf{a}_* = \boldsymbol{\Sigma}\mathbf{h}^*/\|\boldsymbol{\Sigma}^{1/2}\mathbf{h}^*\|$ .

PROOF OF LEMMA 7.6. By the triangle inequality and definitions of  $\Delta_n^a$ ,  $\Delta_n^c$ ,

$$(7.29) \quad (1 - \widehat{\mathbf{d}}\mathbf{f}/n)\sigma \leq (\boldsymbol{\varepsilon}/\sigma)^\top (\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}})/n + (\Delta_n^a R_*)^{1/2},$$

$$(7.30) \quad (1 - \widehat{\mathbf{d}}\mathbf{f}/n)\langle \mathbf{a}_*, \mathbf{h} \rangle \leq (\mathbf{g}^\top (\mathbf{X}\hat{\boldsymbol{\beta}} - \mathbf{y}))/n + (\Delta_n^c R_*)^{1/2},$$

$$(1 - \widehat{\mathbf{d}}\mathbf{f}/n)(\sigma^2 + \mathbf{h}^\top \boldsymbol{\Sigma} \mathbf{h}^*) \leq (\boldsymbol{\varepsilon} - \mathbf{X}\mathbf{h}^*)^\top (\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}})/n + (\Delta_n^a)^{1/2} R_* + (\Delta_n^c)^{1/2} R_*,$$

where the last line follows from the weighted sum  $\sigma(7.29) + \|\Sigma^{1/2}\mathbf{h}^*\|(7.30)$  and using  $\sigma \vee \|\Sigma^{1/2}\mathbf{h}^*\| \leq R_*^{1/2}$  for the last two terms. By the KKT conditions of  $\hat{\beta}$  and  $\beta^*$ ,

$$\begin{aligned}(\beta^* - \hat{\beta})^\top \partial g(\beta^*) &= (\mathbf{h} - \mathbf{h}^*)^\top \Sigma \mathbf{h}^*, \\(\hat{\beta} - \beta^*)^\top \partial g(\hat{\beta}) &= (\hat{\beta} - \beta^*)^\top X^\top (\mathbf{y} - X\hat{\beta})/n.\end{aligned}$$

Summing these equalities and using the monotonicity of the subdifferential yields

$$\begin{aligned}(7.31) \quad & \|\Sigma^{1/2}\mathbf{h}^*\|^2 + \|\mathbf{y} - X\hat{\beta}\|^2/n \\& \leq \mathbf{h}^\top \Sigma \mathbf{h}^* + (\hat{\beta} - \beta^*)^\top X^\top (\mathbf{y} - X\hat{\beta})/n + \|\mathbf{y} - X\hat{\beta}\|^2/n \\& = \mathbf{h}^\top \Sigma \mathbf{h}^* + (\mathbf{y} - X\beta^*)^\top (\mathbf{y} - X\hat{\beta})/n.\end{aligned}$$

Combining (7.31) multiplied by  $1 - \hat{\mathbf{d}}\mathbf{f}/n$  with the line after (7.30) gives

$$\begin{aligned}& (1 - \hat{\mathbf{d}}\mathbf{f}/n)(R_* + \|\mathbf{y} - X\hat{\beta}\|^2/n) \\& \leq (2 - \hat{\mathbf{d}}\mathbf{f}/n)(\mathbf{y} - X\beta^*)^\top (\mathbf{y} - X\hat{\beta})/n + (\Delta_n^a)^{1/2}R_* + (\Delta_n^c)^{1/2}R_* \\& \leq 2\|\mathbf{y} - X\beta^*\|\|\mathbf{y} - X\hat{\beta}\|/n + 2(\max\{\Delta_n^a, \Delta_n^c\})^{1/2}R_*\end{aligned}$$

using the Cauchy–Schwarz inequality and  $(2 - \hat{\mathbf{d}}\mathbf{f}/n) \leq 2$  for the last inequality. Using  $(2a + 2b)^2 \leq 8(a^2 + b^2)$  for the right-hand side with  $\|\mathbf{y} - X\beta^*\|^2\|\mathbf{y} - X\hat{\beta}\|^2/(n^2R_*^2) \leq \|\mathbf{y} - X\hat{\beta}\|^2/(nR_*) + \Delta_n^b$  completes the proof of (7.24). The second inequality, (7.25), then follows from (3.9), (3.18) and

$$\begin{aligned}& \text{trace}[(\mathbf{I}_n - \hat{\mathbf{H}})\langle \mathbf{a}_0, \mathbf{h} \rangle + \mathbf{w}_0(\mathbf{y} - X\hat{\beta})]^2 \\& = \|\mathbf{I}_n - \hat{\mathbf{H}}\|_F^2 \langle \mathbf{a}_0, \mathbf{h} \rangle^2 + (\mathbf{w}_0^\top (\mathbf{y} - X\hat{\beta}))^2 + 2\mathbf{w}_0^\top (\mathbf{I}_n - \hat{\mathbf{H}})(\mathbf{y} - X\hat{\beta})\langle \mathbf{a}_0, \mathbf{h} \rangle,\end{aligned}$$

which implies  $V^*(\theta) - \|\mathbf{y} - X\hat{\beta}\|^2 \geq -n\Delta_n^d R_*$ .  $\square$

**LEMMA 7.7.** Define  $\Delta_n \stackrel{\text{def}}{=} \Delta_n^a + \Delta_n^b + \Delta_n^c + \Delta_n^d$  where  $\Delta_n^a, \dots, \Delta_n^d$  are defined in Lemma 7.6. Under Assumption 3.1, we have  $\mathbb{E}[\Delta_n] \leq C(\gamma, \mu)n^{-1/2}$ .

**PROOF OF LEMMA 7.7.** We bound each of  $\Delta_n^a, \Delta_n^b, \Delta_n^c, \Delta_n^d$  separately. We have  $\Delta_n^b \leq (F_+ - 1)4F_+F^2$  by (7.11) so that  $\mathbb{E}[\Delta_n^b] \leq C_{13}(\gamma, \mu)n^{-1/2}$  by virtue of (7.9). For  $\Delta_n^a$ , we have  $\Delta_n^a = n^{-2}\sigma^{-2}|\sigma^2(n - \hat{\mathbf{d}}\mathbf{f}) - \boldsymbol{\varepsilon}^\top (\mathbf{y} - X\hat{\beta})|^2/R_*$ . By the second-order Stein formula (Proposition 2.1) with respect to  $\boldsymbol{\varepsilon}$  conditionally on  $X$ ,

$$\mathbb{E}[\Delta_n^a] = n^{-2}\mathbb{E}[\|\mathbf{y} - X\hat{\beta}\|^2/R_* + \sigma^2 \text{trace}(\{\mathbf{I}_n - \hat{\mathbf{H}}\}^2)/R_*] \leq n^{-1}\mathbb{E}[4F_+F^2 + 1],$$

where we used  $\text{trace}(\{\mathbf{I}_n - \hat{\mathbf{H}}\}^2) \leq n$  from Proposition 7.3 and (7.11) for the inequality. Thanks to (7.9), this shows that  $\mathbb{E}[\Delta_n^a] \leq n^{-1}C_{14}(\gamma, \mu)$ . Similarly, for  $\Delta_n^d$  in (7.25),  $\Delta_n^d \leq 2n^{-1}\|\mathbf{w}_0\|\|\mathbf{y} - X\hat{\beta}\|\|\langle \mathbf{a}_0, \mathbf{h} \rangle\|/R_* \leq n^{-1}2(F/2)^{1/2}2F_+F^2$ , hence  $\mathbb{E}[\Delta_n^d] \leq n^{-1}C_{15}(\gamma, \mu)$  by (7.11) and (7.9). For  $\Delta_n^c$ , we have  $\mathbf{g} = \mathbf{z}_0$  for  $\mathbf{a}_0 = \mathbf{a}_*$  so that  $\mathbb{E}[\Delta_n^c] \leq C_{16}(\gamma, \mu)n^{-1}$  by (7.20).  $\square$

**7.5. Event  $\Omega_n$ .** With  $\Omega_0, C_*(\gamma, \mu)$  in Lemma 7.4 and  $\Delta_n$  in Lemma 7.7, let

$$(7.32) \quad \Omega_n = \Omega_0 \cap \{\mathbb{E}_0[\Delta_n] \vee \Delta_n \leq C_*^2(\gamma, \mu)/16\}.$$

By the union bound, Markov's inequality and the bound on  $\mathbb{E}[\Delta_n]$  in Lemma 7.7,

$$(7.33) \quad \mathbb{P}(\Omega_n^c) \leq \mathbb{P}(\Omega_0^c) + C_{17}(\gamma, \mu)n^{-1/2} \leq C_0n^{-1/2}$$

thanks to  $\mathbb{P}(\Omega_0^c) \leq e^{-n/2}$  in Lemma 7.4. By (7.24),

$$(7.34) \quad \Omega_n \subset \{\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2 \geq R_* n C_*^2(\gamma, \mu)/16\}.$$

Since  $\Omega_0^c$  is independent of  $\mathbf{z}_0$ , taking the condition expectation  $\mathbb{E}_0$  of (7.25) in  $\Omega_0$  gives

$$(7.35) \quad \Omega_n \subset \{\text{Var}_0[\xi_0] \wedge \mathbb{E}_0[\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2] \geq R_* n C_*^2(\gamma, \mu)/16\}.$$

Proofs of Lemmas 3.2, 3.4 and 3.5 and Theorem 3.9.

LEMMA 3.2. *Under Assumption 3.1, there exists  $\Omega_n$  with  $\mathbb{P}(\Omega_n^c) \leq C_0(\gamma, \mu)n^{-1/2}$  and*

$$(3.14) \quad \mathbb{E}[I_{\Omega_n} \langle \mathbf{w}_0, \mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}} \rangle^2 / \text{Var}_0[\xi_0]] \leq C_{18}(\gamma, \mu)n^{-1}.$$

PROOF OF LEMMA 3.2. With  $\Omega_n$  in (7.32), (3.14) follows from (7.35) and (7.22).  $\square$

LEMMA 3.4. *Under Assumption 3.1, there exists  $\Omega_n$  with  $\mathbb{P}(\Omega_n^c) \leq C_0(\gamma, \mu)n^{-1/2}$  and*

$$(3.20) \quad \mathbb{E}\left[I_{\Omega_n} \left| \frac{\mathbb{E}_0[\hat{V}(\theta)]}{\text{Var}_0[\xi_0]} - 1 \right| \right] \leq \mathbb{E}\left[I_{\Omega_n} \frac{\mathbb{E}_0[|\hat{V}(\theta) - V^*(\theta)|]}{\text{Var}_0[\xi_0]} \right] \leq \frac{C_{19}(\gamma, \mu)}{n}.$$

PROOF OF LEMMA 3.4. Let  $\Omega_n$  be as in (7.32). The first inequality in (3.20) follows from the triangle inequality. By (7.35), we have  $\mathbb{E}[I_{\Omega_n} \mathbb{E}_0[|\hat{V}(\theta) - V^*(\theta)|] / \text{Var}_0[\xi_0]] \leq \mathbb{E}[|\hat{V}(\theta) - V^*(\theta)|]16/(nR_*C_*^2(\gamma, \mu))$ . With  $V^*(\theta)$ ,  $\hat{V}(\theta)$  in (3.18)–(3.19) and  $\nabla f(\mathbf{z}_0)^\top$  in (3.9),

$$V^*(\theta) - \hat{V}(\theta) = \langle \mathbf{w}_0, \mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}} \rangle^2 + 2\mathbf{w}_0^\top (\mathbf{I}_n - \hat{\mathbf{H}})(\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}) \langle \mathbf{a}_0, \mathbf{h} \rangle.$$

Using  $\|\mathbf{I}_n - \hat{\mathbf{H}}\|_{\text{op}} \leq 1$  from Proposition 7.3 and (7.10)–(7.11), we find by the Cauchy–Schwarz inequality  $|V^*(\theta) - \hat{V}(\theta)| \leq (2 + 2\sqrt{2})R_*F_+F^3$ . The proof of (3.20) is complete by virtue of Hölder’s inequality and the moment bounds (7.9).  $\square$

LEMMA 3.5. *Under Assumption 3.1, there exists  $\Omega_n$  with  $\mathbb{P}(\Omega_n^c) \leq C_0(\gamma, \mu)n^{-1/2}$  and*

$$(3.22) \quad \max\left\{\mathbb{E}\left[I_{\Omega_n} \left| \frac{\check{V}(\mathbf{a}_0)^{1/2}}{\hat{V}(\theta)^{1/2}} - 1 \right|^2\right], \mathbb{E}\left[I_{\Omega_n} \left| \frac{\mathbb{E}_0[\check{V}(\mathbf{a}_0)]^{1/2}}{\mathbb{E}_0[\hat{V}(\theta)]^{1/2}} - 1 \right|^2\right]\right\} \leq \frac{C_{20}(\gamma, \mu)}{n}.$$

PROOF OF LEMMA 3.5. By the triangle inequality for the Euclidean norm in  $\mathbb{R}^2$ ,

$$(7.36) \quad |\check{V}(\mathbf{a}_0)^{1/2} - \hat{V}(\theta)^{1/2}| \leq \|\mathbf{I}_n - \hat{\mathbf{H}}\|_F |\langle \mathbf{a}_0, \mathbf{h} \rangle + (n - \widehat{\text{df}})^{-1} \langle \mathbf{z}_0, \mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}} \rangle|.$$

Let  $\Omega_n$  be as in (7.32). Using  $\hat{V}(\theta) \geq \|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2$ , the lower bound (7.34) and  $\|\mathbf{I}_n - \hat{\mathbf{H}}\|_F^2 \leq n$  by Proposition 7.3,

$$\begin{aligned} & \mathbb{E}[I_{\Omega_n} |\check{V}(\mathbf{a}_0)^{1/2} / \hat{V}(\theta)^{1/2} - 1|^2] \\ & \leq \mathbb{E}[I_{\Omega_n} (\langle \mathbf{a}_0, \mathbf{h} \rangle + (n - \widehat{\text{df}})^{-1} \langle \mathbf{z}_0, \mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}} \rangle)^2] 16/(R_*C_*^2(\gamma, \mu)) \end{aligned}$$

so that  $\Omega_n \subset \Omega_0$  and (7.21) completes the proof for the first term in the maximum in (3.22). For the second term in the maximum, by the triangle inequality for the norm  $\mathbb{E}_0[(\cdot)^2]^{1/2}$ , we have  $|\mathbb{E}_0[\check{V}(\mathbf{a}_0)]^{1/2} - \mathbb{E}_0[\hat{V}(\theta)]^{1/2}| \leq \mathbb{E}_0[|\check{V}(\mathbf{a}_0)^{1/2} - \hat{V}(\theta)^{1/2}|^2]^{1/2}$ . The proof is completed by using again (7.36), the lower bound (7.35) on  $\mathbb{E}_0[\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2]$  in  $\Omega_n$  and the same argument as for the first term in the maximum.  $\square$

THEOREM 3.9. *Let Assumption 3.1 be fulfilled. Then the following are equivalent:*

- (i)  $\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2 / \text{Var}_0[\xi_0] \xrightarrow{\mathbb{P}} 1$ ,

- (ii)  $\mathbb{E}_0[\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2]/\text{Var}_0[\xi_0] \xrightarrow{\mathbb{P}} 1$ ,
- (iii)  $\langle \mathbf{a}_0, \mathbf{h} \rangle^2/R_* \xrightarrow{\mathbb{P}} 0$ ,
- (iv)  $\langle \mathbf{a}_0, \mathbf{h} \rangle^2 n / \|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2 \xrightarrow{\mathbb{P}} 0$ ,
- (v)  $\langle \mathbf{z}_0, \mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}} \rangle^2 / (n \|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2) \xrightarrow{\mathbb{P}} 0$ ,
- (vi)  $\hat{V}(\theta) / \|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2 \xrightarrow{\mathbb{P}} 1$ ,
- (vii)  $\check{V}(\mathbf{a}_0) / \|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2 \xrightarrow{\mathbb{P}} 1$ .

PROOF OF THEOREM 3.9. (v)  $\Leftrightarrow$  (iv) is due to  $C_*(\gamma, \mu)n \leq \|\mathbf{I}_n - \hat{\mathbf{H}}\|_F^2 \leq n$  in  $\Omega_n$  by Lemma 7.4 and Proposition 7.3 combined with (3.17).

(iv)  $\Leftrightarrow$  (vi) follows from  $C_*(\gamma, \mu)n \leq \|\mathbf{I}_n - \hat{\mathbf{H}}\|_F^2 \leq n$  in  $\Omega_n$ .

(vi)  $\Leftrightarrow$  (vii) is proved in Lemma 3.5.

(iii)  $\Rightarrow$  (i), (iii)  $\Rightarrow$  (vi) and (iii)  $\Rightarrow$  (vii) are shown in the proof of Theorem 3.6.

(iv)  $\Rightarrow$  (iii) follows from  $\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2/(nR_*) = O_{\mathbb{P}}(1)$  by (7.11) and (7.9).

(iii)  $\Rightarrow \delta_1^2 \xrightarrow{\mathbb{P}} 0$  was shown in the proof of Theorem 3.6, and  $\delta_1^2 \xrightarrow{\mathbb{P}} 0$  implies  $\mathbb{E}[\|\nabla f(\mathbf{z}_0)\|_F^2]/\mathbb{E}_0[\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2] \xrightarrow{\mathbb{P}} 0$  and  $\mathbb{E}[\text{trace}[(\nabla f(\mathbf{z}_0))^2]]/\mathbb{E}_0[\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2] \xrightarrow{\mathbb{P}} 0$  so that (iii)  $\Rightarrow$  (ii) holds.

By Lemma 3.4, (ii) implies that  $\mathbb{E}_0[\hat{V}(\theta)]/\mathbb{E}_0[\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2] - 1 = \mathbb{E}_0[\|\mathbf{I}_n - \hat{\mathbf{H}}\|_F^2 \langle \mathbf{a}_0, \mathbf{h} \rangle^2]/\mathbb{E}_0[\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2]$  converges to 0 in probability. Since  $\mathbb{E}_0[\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2]/(nR_*) = O_{\mathbb{P}}(1)$  by (7.11) and (7.9), combined with  $\|\mathbf{I}_n - \hat{\mathbf{H}}\|_F^2 \geq C_*^2(\gamma, \mu)n$  in  $\Omega_0$  by Lemma 7.4, this implies  $I_{\Omega_0}\mathbb{E}_0[\langle \mathbf{a}_0, \mathbf{h} \rangle^2]/R_* = \mathbb{E}_0[I_{\Omega_0}\langle \mathbf{a}_0, \mathbf{h} \rangle^2]/R_* \xrightarrow{\mathbb{P}} 0$  as  $\Omega_0$  is independent of  $\mathbf{z}_0$ . Thus, (ii) implies (iii) by Markov's inequality with respect to  $\mathbb{E}_0$ .

Finally, to show (i)  $\Leftrightarrow$  (ii), we have by the Gaussian Poincaré inequality

$$\text{Var}_0[\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2] \leq \mathbb{E}_0[\|\nabla f(\mathbf{z}_0)(\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}})\|^2] \leq \mathbb{E}_0[\|\nabla f(\mathbf{z}_0)\|_{\text{op}}^2 \|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2].$$

With  $\nabla f(\mathbf{z}_0)$  in (3.9) and the bounds (7.10)–(7.11), we have  $\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2 \leq 4nF_+F^2R_*$  and  $\|\nabla f(\mathbf{z}_0)\|_{\text{op}}^2 \leq 2(2F_+F^3 + F_+F^2)R_*$  thanks to  $\|\mathbf{I}_n - \hat{\mathbf{H}}\|_{\text{op}} \leq 1$  by Proposition 7.3. Combined with the lower bound (7.35) on  $\text{Var}_0[\xi_0]$  and the moment upper bounds (7.9), we obtain  $\mathbb{E}[I_{\Omega_n} \text{Var}_0[\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2]/\text{Var}_0[\xi_0]^2] \leq C_{21}(\gamma, \mu)n^{-1/2}$ , which gives (vii)  $\Leftrightarrow$  (i).  $\square$

### 7.6. Proofs of Theorem 3.3 and asymptotic normality results.

THEOREM 3.3. Under Assumption 3.1, there exists  $\Omega_n$  with  $\mathbb{P}(\Omega_n^c) \leq C_0(\gamma, \mu)n^{-1/2}$  and

$$(3.17) \quad \mathbb{E}[I_{\Omega_n}(n - \widehat{\text{df}})^2 \langle \mathbf{a}_0, \hat{\boldsymbol{\beta}}^{(\text{de-bias})} - \boldsymbol{\beta} \rangle^2 / \|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2] \leq C_{22}(\gamma, \mu).$$

Furthermore,  $|\langle \mathbf{a}_0, \hat{\boldsymbol{\beta}}^{(\text{de-bias})} - \boldsymbol{\beta} \rangle| = O_{\mathbb{P}}(1)\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|/(n - \widehat{\text{df}}) = O_{\mathbb{P}}(1)\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|/n$ .

PROOF OF THEOREM 3.3. Let  $\Omega_n$  be as in (7.32). Since  $I_{\Omega_n}\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^{-2} \leq 16/(C_*^2(\gamma, \mu)nR_*)$  by (7.34), using (7.21) completes the proof of (3.17). For the second part, random variables bounded in  $L_2$  are stochastically bounded so that (3.3) provides  $|\langle \mathbf{a}_0, \hat{\boldsymbol{\beta}}^{(\text{de-bias})} - \boldsymbol{\beta} \rangle| = O_{\mathbb{P}}(1)\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|/(n - \widehat{\text{df}})$ , and  $I_{\Omega_0}(1 - \widehat{\text{df}}/n)^{-1} \leq C_*(\gamma, \mu)^{-1}$  for  $\Omega_0$  in Lemma 7.4 provides  $(1 - \widehat{\text{df}}/n) = O_{\mathbb{P}}(1)$ .  $\square$

THEOREM 3.6. Let Assumption 3.1 be fulfilled. Let  $\hat{\boldsymbol{\beta}}^{(\text{de-bias})}$  be as in (3.15). Then, for any  $\mathbf{a}_0$  with  $\|\Sigma^{-1/2}\mathbf{a}_0\| = 1$  such that  $\langle \mathbf{a}_0, \mathbf{h} \rangle^2/R_* \xrightarrow{\mathbb{P}} 0$ ,

$$\sup_{t \in \mathbb{R}} \left[ \left| \mathbb{P}\left(\frac{\xi_0}{V_0^{1/2}} \leq t\right) - \Phi(t) \right| + \left| \mathbb{P}\left(\frac{\langle \mathbf{a}_0, \hat{\boldsymbol{\beta}}^{(\text{de-bias})} \rangle - \theta}{V_0^{1/2}/(n - \widehat{\text{df}})} \leq t\right) - \Phi(t) \right| \right] \rightarrow 0,$$

where  $V_0$  denotes any of the four quantities:  $\text{Var}_0[\xi_0]$ ,  $\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2$ ,  $\hat{V}(\theta)$  or  $\check{V}(\mathbf{a}_0)$ .

PROOF OF THEOREM 3.6. Let  $\Omega_n$  be as in (7.32). Let  $\delta_1^2$  be the quantity in (3.24), omitting the dependence in  $\mathbf{a}_0$  as it is clear from context. Since  $\delta_1^2 \leq 1$  by definition,  $\mathbb{E}[\delta_1^2] \leq \mathbb{E}[\Omega_n \delta_1^2] + \mathbb{P}(\Omega_n^c)$ . In  $\Omega_n$ , (7.35) provides a lower bound on the denominator of  $\delta_1^2$  so that  $\mathbb{E}[I_{\Omega_n} \delta_1^2] \leq \mathbb{E}[\|\nabla f(\mathbf{z}_0)\|_F^2] 16 / (n R_* C_*^2(\gamma, \mu))$ . By (7.23) and the bound (7.33) on  $\mathbb{P}(\Omega_n^c)$ , we obtain

$$(7.37) \quad \mathbb{E}[\delta_1^2] \leq \mathbb{E}[I_{\Omega_n} \delta_1^2] + \mathbb{P}(\Omega_n^c) \leq C_{23}(\gamma, \mu)(\mathbb{E}[\langle \mathbf{a}_0, \mathbf{h} \rangle^2 / R_*] + n^{-1}) + C_0(\gamma, \mu)n^{-1/2}$$

Furthermore,  $\langle \mathbf{a}_0, \mathbf{h} \rangle^2 / R_*$  is bounded in  $L_2$  thanks to  $\mathbb{E}[\langle \mathbf{a}_0, \mathbf{h} \rangle^4 / R_*^2] \leq \mathbb{E}[F_+^2 F^4] \leq C_{24}(\gamma, \mu)$  by (7.10) and (7.9). Since a sequence of random variables uniformly bounded in  $L_2$  is uniformly integrable, the assumption  $\langle \mathbf{a}_0, \mathbf{h} \rangle^2 / R_* \xrightarrow{\mathbb{P}} 0$  implies  $\mathbb{E}[\langle \mathbf{a}_0, \mathbf{h} \rangle^2 / R_*] \rightarrow 0$ , and thus  $\mathbb{E}[\delta_1^2] \rightarrow 0$ . This completes the proof that  $\mathbb{E}[\delta_1^2] \rightarrow 0$  and that  $\xi_0 / \text{Var}_0[\xi_0]^{1/2} \xrightarrow{d} N(0, 1)$  by Theorem 2.2. Next, by (3.16),  $(n - \widehat{\text{df}})\langle \mathbf{a}_0, \widehat{\boldsymbol{\beta}}^{(\text{de-bias})} - \boldsymbol{\beta} \rangle / \text{Var}_0[\xi_0]^{1/2} \xrightarrow{d} N(0, 1)$  also holds. It remains to prove  $V_0 / \text{Var}_0[\xi_0] \xrightarrow{\mathbb{P}} 1$  for all four possible choices for  $V_0$ . By (2.7),  $\mathbb{E}[\delta_1^2] \rightarrow 0$  implies  $\|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2 / \text{Var}_0[\xi_0] \xrightarrow{\mathbb{P}} 1$ , while

$$(7.38) \quad 0 \leq \frac{\widehat{V}(\theta)}{\|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2} - 1 = \frac{\|\widehat{\mathbf{H}} - \mathbf{I}_n\|_F^2 \langle \mathbf{a}_0, \mathbf{h} \rangle^2 / (n R_*)}{\|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2 / (n R_*)}.$$

Proposition 7.3 provides  $\|\widehat{\mathbf{H}} - \mathbf{I}_n\|_F^2 \leq n$  so that the numerator converges to 0 in probability thanks to assumption  $\langle \mathbf{a}_0, \mathbf{h} \rangle^2 / R_* \xrightarrow{\mathbb{P}} 0$ . The denominator is bounded from below by  $C_*^2(\gamma, \mu)/16$  in  $\Omega_n$  by (7.34) and  $\mathbb{P}(\Omega_n) \rightarrow 1$ . This proves  $\widehat{V}(\theta) / \text{Var}_0[\xi_0] \xrightarrow{\mathbb{P}} 1$  and  $\check{V}(\mathbf{a}_0) / \text{Var}_0[\xi_0] \xrightarrow{\mathbb{P}} 1$  follows by Lemma 3.5. Slutsky's theorem completes the proof as  $V_0 / \text{Var}_0[\xi_0] \xrightarrow{\mathbb{P}} 1$  for all four possible choices for  $V_0$ . As  $\Phi(t)$  is continuous, convergence in Kolmogorov distance is equivalent to convergence in distribution.  $\square$

THEOREM 3.7. *There exists an absolute constant  $C^* > 0$  such that the following holds. Let Assumption 3.1 be fulfilled,  $\widehat{\boldsymbol{\beta}}^{(\text{de-bias})}$  be as in (3.15). Then for any increasing sequence  $a_p \rightarrow +\infty$  (e.g.,  $a_p = \log \log p$ ), the subset*

$$(7.25) \quad \overline{\mathcal{S}} = \{\mathbf{v} \in S^{p-1} : \mathbb{E}[\langle \boldsymbol{\Sigma}^{1/2} \mathbf{v}, \mathbf{h} \rangle^2 / \|\boldsymbol{\Sigma}^{1/2} \mathbf{h}\|^2] \leq C^* / a_p\}$$

*of the unit sphere  $S^{p-1}$  in  $\mathbb{R}^p$  has relative volume  $|\overline{\mathcal{S}}| / |S^{p-1}| \geq 1 - 2e^{-p/a_p}$  and*

$$(7.26) \quad \sup_{\mathbf{a}_0 \in \boldsymbol{\Sigma}^{1/2} \overline{\mathcal{S}}} \sup_{t \in \mathbb{R}} \left[ \left| \mathbb{P}\left(\frac{\xi_0}{V_0^{1/2}} \leq t\right) - \Phi(t) \right| + \left| \mathbb{P}\left(\frac{\langle \mathbf{a}_0, \widehat{\boldsymbol{\beta}}^{(\text{de-bias})} - \boldsymbol{\beta} \rangle}{V_0^{1/2} / (n - \widehat{\text{df}})} \leq t\right) - \Phi(t) \right| \right] \rightarrow 0$$

*where  $V_0$  denotes any of the four quantities:  $\text{Var}_0[\xi_0]$ ,  $\|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2$ ,  $\widehat{V}(\theta)$  or  $\check{V}(\mathbf{a}_0)$ . Furthermore, with  $\mathbf{e}_j \in \mathbb{R}^p$  the  $j$ th canonical basis vector and  $\phi_{\text{cond}}(\boldsymbol{\Sigma}) = \|\boldsymbol{\Sigma}\|_{\text{op}} \|\boldsymbol{\Sigma}^{-1}\|_{\text{op}}$ , the asymptotic normality in (3.26) uniformly holds over at least  $(p - \phi_{\text{cond}}(\boldsymbol{\Sigma})a_p / C^*)$  canonical directions in the sense that  $J_p = \{j \in [p] : \mathbf{e}_j / \|\boldsymbol{\Sigma}^{-1/2} \mathbf{e}_j\| \in \boldsymbol{\Sigma}^{1/2} \overline{\mathcal{S}}\}$  has cardinality  $|J_p| \geq p - \phi_{\text{cond}}(\boldsymbol{\Sigma})a_p / C^*$ .*

PROOF OF THEOREM 3.7. We construct a subset  $\overline{\mathcal{S}}$  of the sphere such that  $\langle \mathbf{a}_0, \mathbf{h} \rangle^2 / R_* \xrightarrow{\mathbb{P}} 0$  uniformly over all  $\mathbf{a}_0 \in \boldsymbol{\Sigma}^{1/2} \overline{\mathcal{S}}$ . Let  $\mathbf{v}$  be uniformly distributed on the unit Euclidean sphere  $S^{p-1}$ , independently of  $(\mathbf{X}, \mathbf{y})$ , and denote by  $\nu$  its probability measure. The vector  $\sqrt{p}\mathbf{v}$  is sub-Gaussian in  $\mathbb{R}^p$  [47], Theorem 3.4.6, in the sense that for any nonzero vector  $\mathbf{u} \in \mathbb{R}^p$ ,  $\int \exp\{(\sqrt{p}\mathbf{v}^\top \mathbf{u})^2 / (C^* \|\mathbf{u}\|^2)\} d\nu(\mathbf{v}) \leq 2$  for some absolute constant  $C^* > 0$ . By Jensen's inequality and Fubini's theorem,

$$\int \exp\left\{\mathbb{E}\left[\frac{(\mathbf{v}^\top \boldsymbol{\Sigma}^{1/2} \mathbf{h})^2}{C^* \|\boldsymbol{\Sigma}^{1/2} \mathbf{h}\|^2}\right]\right\} d\nu(\mathbf{v}) \leq \mathbb{E}\left[\int \exp\left\{\frac{(\sqrt{p}\mathbf{v}^\top \mathbf{h})^2}{C^* \|\boldsymbol{\Sigma}^{1/2} \mathbf{h}\|^2}\right\} d\nu(\mathbf{v})\right] \leq 2.$$

Hence, by Markov's inequality, for any positive  $x$ ,

$$\nu(\{\mathbf{v} \in S^{p-1} : \mathbb{E}[(\mathbf{v}^\top \boldsymbol{\Sigma}^{1/2} \mathbf{h})^2 / \|\boldsymbol{\Sigma}^{1/2} \mathbf{h}\|^2] > C^* x / p\}) \leq 2e^{-x}.$$

Setting  $x = p/a_p$ , we obtain that the subset  $\bar{S} \subset S^{p-1}$  defined by (3.25) has relative volume at least  $|\bar{S}|/|S^{p-1}| \geq 1 - 2e^{-p/a_p}$ , and for all  $\mathbf{a}_0 \in \boldsymbol{\Sigma}^{1/2} \bar{S}$ ,

$$(7.39) \quad \mathbb{E}[\langle \mathbf{a}_0, \mathbf{h} \rangle^2 / \|\boldsymbol{\Sigma}^{1/2} \mathbf{h}\|^2] \leq C^*/a_p.$$

Furthermore, the set  $\bar{S} \cap \{\boldsymbol{\Sigma}^{-1/2} \mathbf{e}_j / \|\boldsymbol{\Sigma}^{-1/2} \mathbf{e}_j\|, j \in [p]\}$  has cardinality at least  $p - \phi_{\text{cond}}(\boldsymbol{\Sigma})a_p/C^*$  due to

$$\sum_{j=1}^p \frac{1}{\|\boldsymbol{\Sigma}^{-1/2} \mathbf{e}_j\|^2} \mathbb{E}\left[\frac{\langle \mathbf{e}_j, \mathbf{h} \rangle^2}{\|\boldsymbol{\Sigma}^{1/2} \mathbf{h}\|^2}\right] \leq \|\boldsymbol{\Sigma}\|_{\text{op}} \mathbb{E}\left[\frac{\|\mathbf{h}\|^2}{\|\boldsymbol{\Sigma}^{1/2} \mathbf{h}\|^2}\right] \leq \phi_{\text{cond}}(\boldsymbol{\Sigma}).$$

To show that  $\sup_{\mathbf{a}_0 \in \boldsymbol{\Sigma}^{1/2} \bar{S}} \mathbb{E}[\delta_1^2(\mathbf{a}_0)] \rightarrow 0$ , thanks to (7.37) it is enough to prove that  $\mathbb{E}[\langle \mathbf{a}_0, \mathbf{h} \rangle^2 / R_*] \rightarrow 0$  uniformly over  $\mathbf{a}_0 \in \boldsymbol{\Sigma}^{1/2} \bar{S}$ . By the Cauchy–Schwarz inequality,

$$\begin{aligned} \mathbb{E}[\langle \mathbf{a}_0, \mathbf{h} \rangle^2 / R_*] &= \mathbb{E}[\langle \mathbf{a}_0, \mathbf{h} \rangle \|\boldsymbol{\Sigma}^{1/2} \mathbf{h}\| / R_* \langle \mathbf{a}_0, \mathbf{h} \rangle / \|\boldsymbol{\Sigma}^{1/2} \mathbf{h}\|] \\ &\leq \mathbb{E}[\|\boldsymbol{\Sigma}^{1/2} \mathbf{h}\|^4 / R_*^2]^{1/2} (C^*/a_p)^{1/2} \end{aligned}$$

for any  $\mathbf{a}_0 \in \boldsymbol{\Sigma}^{1/2} \bar{S}$  thanks to (7.39), while  $\mathbb{E}[\|\boldsymbol{\Sigma}^{1/2} \mathbf{h}\|^4 / R_*^2] \leq \mathbb{E}[F_+^2 F^4] \leq C_{25}(\gamma, \mu)$  by (7.10) and (7.9). This implies  $\sup_{\mathbf{a}_0 \in \boldsymbol{\Sigma}^{1/2} \bar{S}} \mathbb{E}[\langle \mathbf{a}_0, \mathbf{h} \rangle^2 / R_*] \rightarrow 0$  and  $\sup_{\mathbf{a}_0 \in \boldsymbol{\Sigma}^{1/2} \bar{S}} \mathbb{E}[\delta_1^2(\mathbf{a}_0)] \rightarrow 0$  hold.

By Theorem 2.2, this shows that  $\xi_0 / \text{Var}_0[\xi_0]^{1/2} \rightarrow^d N(0, 1)$  uniformly over  $\mathbf{a}_0 \in \boldsymbol{\Sigma}^{1/2} \bar{S}$ . Since the bounds (3.14), (7.38) are all uniform over all  $\mathbf{a}_0$  with  $\|\boldsymbol{\Sigma}^{-1/2} \mathbf{a}_0\| = 1$ , Slutsky's theorem implies  $V_0 / \text{Var}_0[\xi_0] \rightarrow^{\mathbb{P}} 1$ ,  $\xi_0 / V_0^{1/2} \rightarrow^d N(0, 1)$  and  $\{(n - \widehat{\text{df}}) \langle \mathbf{a}_0, \mathbf{h} \rangle + \mathbf{z}_0^\top (\mathbf{y} - \mathbf{X} \widehat{\boldsymbol{\beta}})\} / V_0^{1/2} \rightarrow^d N(0, 1)$  uniformly over  $\mathbf{a}_0 \in \boldsymbol{\Sigma}^{1/2} \bar{S}$  for all four possible choices for  $V_0$ , and convergence in Kolmogorov distance follows from convergence in distribution.  $\square$

**THEOREM 3.8.** *Under Assumption 3.1, there exists  $\Omega_n$  with  $\mathbb{P}(\Omega_n^c) \leq C_0(\gamma, \mu)n^{-1/2}$  and*

$$(3.27) \quad \mathbb{E}[I_{\Omega_n}(\langle \mathbf{a}_0, \widehat{\boldsymbol{\beta}} - \boldsymbol{\beta} \rangle + (n - \widehat{\text{df}})^{-1} \mathbf{z}_0^\top (\mathbf{y} - \mathbf{X} \widehat{\boldsymbol{\beta}}))^2] \leq R_* C_{26}(\gamma, \mu)/n$$

*If additionally  $g$  is a seminorm, then  $|\mathbf{z}_0^\top (\mathbf{y} - \mathbf{X} \widehat{\boldsymbol{\beta}})|/n = |\mathbf{a}_0^\top \boldsymbol{\Sigma}^{-1} \mathbf{X}^\top (\mathbf{y} - \mathbf{X} \widehat{\boldsymbol{\beta}})|/n \leq g(\boldsymbol{\Sigma}^{-1} \mathbf{a}_0)$  always holds by properties of the subdifferential of a norm. Consequently, if  $g(\boldsymbol{\Sigma}^{-1} \mathbf{a}_0)^2 / R_* \rightarrow 0$  then  $\langle \mathbf{a}_0, \mathbf{h} \rangle^2 / R_* \rightarrow^{\mathbb{P}} 0$  and the conclusions of Theorem 3.6 hold.*

**PROOF OF THEOREM 3.8.** The first statement of the theorem follows from Lemma 7.5. Finally, if  $g$  is a norm then the KKT conditions of  $\widehat{\boldsymbol{\beta}}$ ,  $|\mathbf{z}_0^\top (\mathbf{y} - \mathbf{X} \widehat{\boldsymbol{\beta}})| = n |(\boldsymbol{\Sigma}^{-1} \mathbf{a}_0)^\top \partial g(\widehat{\boldsymbol{\beta}})| \leq ng(\boldsymbol{\Sigma}^{-1} \mathbf{a}_0)$  since for a norm  $g(\mathbf{u}) = \sup_{\mathbf{v} \in \partial g(\mathbf{u})} \langle \mathbf{u}, \mathbf{v} \rangle$ .  $\square$

## APPENDIX A: INTEGRABILITY OF $\phi_{\min}^{-1}(\mathbf{X} \boldsymbol{\Sigma}^{-1/2} / \sqrt{n})$ WHEN $p/n \rightarrow \gamma \in (0, 1)$

In our regression model with Gaussian covariates, the matrix  $\mathbf{X} \boldsymbol{\Sigma}^{-1/2}$  has i.i.d.  $N(0, 1)$  entries, and the inverse of its smallest singular value enjoys the following integrability property as  $n, p \rightarrow +\infty$  with  $p/n \rightarrow \gamma \in (0, 1)$ .

**PROPOSITION A.1.** *Let  $n > p$  and let  $\mathbf{G}$  be a matrix with  $n$  rows,  $p$  columns and i.i.d.  $N(0, 1)$  entries. Then  $\mathbf{G}^\top \mathbf{G}$  is a Wishart matrix and if  $n, p \rightarrow +\infty$  with  $p/n \rightarrow \gamma \in (0, 1)$  we have for any constant  $k$  not growing with  $n, p$ ,*

$$\lim_{p/n \rightarrow \gamma} \mathbb{E}[\phi_{\min}(\mathbf{G}^\top \mathbf{G}/n)^{-k}] = (1 - \sqrt{\gamma})^{-2k}$$

PROOF. Throughout the proof,  $p = p_n$  is an implicit function of  $n$ ; we omit the subscript for brevity. Since  $S_n = \phi_{\min}(\mathbf{G}^\top \mathbf{G}/n) \rightarrow (1 - \sqrt{\gamma})^2$  almost surely (cf. [36]), it is enough to show that the sequence of random variables  $(S_n^{-k})_{n \geq n_0}$  is uniformly integrable for some  $n_0 > 0$ , that is, that  $\sup_{n \geq n_0} \mathbb{E}[S_n^{-k} I_{\{S_n < \epsilon\}}] \rightarrow 0$  as  $\epsilon \rightarrow 0$ . For uniform integrability, we use the following argument from [22], Section 5. The matrix  $\mathbf{G}^\top \mathbf{G}$  is a Wishart matrix and the density of  $L = \phi_{\min}(\mathbf{G}^\top \mathbf{G})$  satisfies for  $\lambda \geq 0$ ,

$$f_L(\lambda) \leq \frac{\sqrt{\pi} 2^{-(n-p+1)/2} \Gamma(\frac{n+1}{2})}{\Gamma(\frac{p}{2}) \Gamma(\frac{n-p+1}{2}) \Gamma(\frac{n-p+2}{2})} \lambda^{(n-p-1)/2} e^{-\lambda/2} = \frac{\sqrt{\pi} \Gamma(\frac{n+1}{2})}{\Gamma(\frac{p}{2}) \Gamma(\frac{n-p+2}{2})} f_{\chi_{n-p+1}^2}(\lambda);$$

cf. [22], Section 5. The density of  $S_n = L/n = \phi_{\min}(\mathbf{G}^\top \mathbf{G}/n)$  that we are interested in is given by  $f_{S_n}(x) = n f_L(nx)$  for  $x \geq 0$ . Hence, if  $0 < \epsilon < (1 - \gamma)/2$ ,

$$\mathbb{E}[S_n^{-k} I_{\{S_n < \epsilon\}}] \leq \left[ \frac{\sqrt{\pi} \Gamma(\frac{n+1}{2}) (\frac{n}{2})^{(n-p+1)/2}}{\Gamma(\frac{p}{2}) \Gamma(\frac{n-p+1}{2}) \Gamma(\frac{n-p+2}{2})} \right] \int_0^\epsilon x^{(n-p-1)/2-k} e^{-nx/2} dx.$$

The mode of the integrand over  $[0, +\infty)$  is  $x_n^* = 1 - p/n - 1/n - 2k/n$ . Thanks to  $\epsilon < (1 - \gamma)/2$ , there exists some  $n_1 \geq 1$  such that for all  $n \geq n_1$ ,

$$(A.1) \quad n - p - 1 - 2k \geq n(1 - \gamma)/2,$$

$(1 - \gamma)/2$  is smaller than the mode  $x_n^*$  and the integral above is bounded by  $\epsilon^{(n-p-k+1)/2} \times e^{-n\epsilon/2}$ . Let  $\Lambda_n$  denote the bracket of the previous display. Then using Stirling's formula  $\Gamma(x+1) \asymp \sqrt{2\pi x} e^{-x} x^x$ , we have for some constants  $n_2, C_2(\gamma) > 0$  possibly depending on  $\gamma$ ,

$$\sup_{n \geq n_2} \frac{\log(\Lambda_n)}{(n - p + 1)/2} \leq C_2(\gamma)$$

because the main terms (coming from  $x^x$  in Stirling's formula) cancel each other. Then for any  $n \geq n_1 \vee n_2$ ,

$$(A.2) \quad \begin{aligned} \mathbb{E}[S_n^{-k} I_{\{S_n < \epsilon\}}] &\leq (\exp(C_2(\gamma))\epsilon)^{(n-p+1)/2} \epsilon^{-k} e^{-n\epsilon/2} \\ &\leq (\exp(C_2(\gamma))\epsilon)^{(n-p+1)/2-k} e^{kC_2(\gamma)-n\epsilon/2}. \end{aligned}$$

For  $n \geq n_1$ , (A.1) holds and if  $\epsilon < (\exp C_2(\gamma))^{-1}$  we have

$$\sup_{n \geq n_1 \vee n_2} \mathbb{E}[S_n^{-k} I_{\{S_n < \epsilon\}}] \leq (\exp(C_2(\gamma))\epsilon)^{n_1(1-\gamma)/4} e^{kC_2(\gamma)}$$

which converges to 0 as  $\epsilon \rightarrow 0$ . This shows uniform integrability of the sequence and proves the claim.  $\square$

## APPENDIX B: PROOF: $p > n$ WITHOUT STRONG CONVEXITY

LEMMA B.1. *Let  $\boldsymbol{\beta} \in \mathbb{R}^p$  and assume that  $p/n \leq \gamma$ . Then for any  $\kappa < 1$ ,*

$$\mathbb{P}\left(\inf_{t \in \mathbb{R}, \mathbf{u} \in \mathbb{R}^p: \|\mathbf{u}\|_0 \leq \kappa n, \mathbf{u} - t\boldsymbol{\beta} \neq \mathbf{0}} \left( \frac{\|\mathbf{X}(\mathbf{u} - t\boldsymbol{\beta})\|^2}{n \|\boldsymbol{\Sigma}^{1/2}(\mathbf{u} - t\boldsymbol{\beta})\|^2} \right) > \varphi(\gamma, \kappa)^2\right) \rightarrow 1,$$

for some constant  $\varphi(\gamma, \kappa) > 0$  depending only on  $\gamma, \kappa$ .

Lemma B.1 and its proof are straightforward extensions of [12], Proposition 2.10, which treats the case  $\boldsymbol{\Sigma} = \mathbf{I}_p, \boldsymbol{\beta} = \mathbf{0}$ .

PROOF. If  $V \subset \mathbb{R}^p$  is a subspace of dimension  $d = \lfloor \kappa n \rfloor + 1$  and  $\mathbf{G} = \mathbf{X}\mathbf{\Sigma}^{-1/2}$ , then by (A.2) with  $k = 0$ ,  $\epsilon \in (0, (1 - d/n)/2)$  and  $n$  large enough,

$$\mathbb{P}\left(\inf_{\mathbf{v} \in \mathbf{\Sigma}^{1/2}V: \|\mathbf{v}\|=1} \|\mathbf{G}\mathbf{v}\|^2/n < \epsilon\right) \leq \exp(C_2(\kappa') \log(\epsilon)(n - p + 1)/2 - n\epsilon/2)$$

for constant  $\kappa' = (\kappa + 1)/2$  thanks to  $1 > \kappa' \geq d/n$ . Applying this bound to the subspace  $V_B = \{\mathbf{u} - t\boldsymbol{\beta}, (\mathbf{u}, t) \in \mathbb{R}^{p+1} : \mathbf{u}_{B^c} = \mathbf{0}\}$  for  $B \subset [p]$  with  $|B| \leq \kappa n$  and using the union bound,

$$\begin{aligned} \mathbb{P}\left(\inf_{t \in \mathbb{R}, \mathbf{u} \in \mathbb{R}^p: \|\mathbf{u}\|_0 \leq \kappa n} \left(\frac{\|\mathbf{X}(\mathbf{u} - t\boldsymbol{\beta})\|^2}{n\|\mathbf{\Sigma}^{1/2}(\mathbf{u} - t\boldsymbol{\beta})\|^2}\right) < \epsilon\right) &\leq \binom{p}{\lfloor \kappa n \rfloor} e^{C_2(\kappa') \log(\epsilon)(n-d+1)/2 - n\epsilon/2} \\ &\leq e^{n \log(e\gamma) + C_2(\kappa') \log(\epsilon)(n-d+1)/2 - n\epsilon/2} \end{aligned}$$

using  $\binom{p}{q} \leq e^{q \log(ep/q)} \leq e^{n \log(ep/n)}$  with  $q = \lfloor \kappa n \rfloor \leq n$  and  $p/n \leq \gamma$ . Since  $d \leq \kappa n + 1$ , choosing  $\epsilon = 1 \wedge \exp(C_2(\kappa')^{-1}(1 - \kappa)^{-1}2 \log(e\gamma))$  the right-hand side of the previous display is bounded from above by  $e^{-n\epsilon/2}$ . This value of  $\epsilon$  provides  $\varphi(\gamma, \kappa)^2$ .  $\square$

**THEOREM 3.10.** *Let  $\gamma \geq 1$ ,  $\kappa \in (0, 1)$  be constants independent of  $\{n, p\}$ . Consider a sequence of regression problems with  $p/n \leq \gamma$  and invertible  $\mathbf{\Sigma}$ . Assume that the group Lasso estimator  $\hat{\boldsymbol{\beta}}$  in (3.33) satisfies*

$$(3.34) \quad \mathbb{P}(\|\hat{\boldsymbol{\beta}}\|_0 \leq \kappa n/2) \rightarrow 1.$$

*If  $\mathbf{a}_0$  is such that  $\|\mathbf{\Sigma}^{-1/2}\mathbf{a}_0\| = 1$  and  $\langle \mathbf{a}_0, \hat{\boldsymbol{\beta}} - \boldsymbol{\beta} \rangle^2/R_* \xrightarrow{\mathbb{P}} 0$  for the  $R_*$  in (3.23), then*

$$(3.35) \quad \sup_{t \in \mathbb{R}} |\mathbb{P}(\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^{-1}((n - \widehat{\text{df}})\langle \mathbf{a}_0, \hat{\boldsymbol{\beta}} - \boldsymbol{\beta} \rangle + \mathbf{z}_0^\top(\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}})) \leq t) - \Phi(t)| \rightarrow 0.$$

*Furthermore, for any  $a_p$  with  $a_p \rightarrow \infty$  and  $\bar{S}$  in (3.25), the relative volume bound given after (3.25) holds, and the asymptotic normality (3.35) holds uniformly over all  $\mathbf{a}_0 \in \mathbf{\Sigma}^{1/2}\bar{S}$  and uniformly over at least  $(p - \phi_{\text{cond}}(\mathbf{\Sigma})a_p/C^*)$  canonical directions in the sense that  $J_p = \{j \in [p] : \mathbf{e}_j/\|\mathbf{\Sigma}^{-1/2}\mathbf{e}_j\| \in \mathbf{\Sigma}^{1/2}\bar{S}\}$  has cardinality  $|J_p| \geq p - \phi_{\text{cond}}(\mathbf{\Sigma})a_p/C^*$ .*

**PROOF OF THEOREM 3.10.** As in the rest of the paper,  $f(\mathbf{z}_0) = \mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}$  and we wish to apply Theorem 2.2 to  $\mathbf{z}_0$  conditionally on  $(\boldsymbol{\epsilon}, \mathbf{X}\mathbf{Q}_0)$ . Instead of applying Theorem 2.2 to  $f$ , and in order to avoid certain events of small probability where the sparse eigenvalues of  $\mathbf{X}$  are not well behaved, we will apply it to a different function. Consider  $F_+$  in (7.7) and the events,

$$\Omega_L = \{\|\hat{\boldsymbol{\beta}}\|_0 \leq \kappa n/2\},$$

$$\Omega_\chi = \{F_+ < 2, (F_+ - 1)_+ < 4\sqrt{\log(n)/n}\},$$

$$\Omega_E = \left\{ \min_{t \in \mathbb{R}, \mathbf{u} \in \mathbb{R}^p: \|\mathbf{u}\|_0 \leq \kappa n} \frac{\|\mathbf{X}(\mathbf{u} - t\boldsymbol{\beta})\|}{\|\mathbf{\Sigma}^{1/2}(\mathbf{u} - t\boldsymbol{\beta})\|} > \varphi\sqrt{n}, \|\mathbf{X}\mathbf{\Sigma}^{-1/2}\|_{\text{op}} < \sqrt{n}(2 + \sqrt{\gamma}) \right\},$$

where  $\varphi = \varphi(\gamma, \kappa)$  is the constant from Lemma B.1. Finally, let  $\Omega_{\text{KKT}}$  be the event (C.1) that the KKT conditions of  $\hat{\boldsymbol{\beta}}$  hold strictly, and set

$$\Omega \stackrel{\text{def}}{=} \Omega_L \cap \Omega_E \cap \Omega_{\text{KKT}} \cap \Omega_\chi.$$

We have  $\mathbb{P}(\Omega_L) \rightarrow 1$  by (3.34) and standard concentration bounds for  $\chi_n^2$  random variables [29], Lemma 1, give  $\mathbb{P}(\Omega_\chi) \rightarrow 1$ . Lemma B.1 and [20], Theorem II.13, provide  $\mathbb{P}(\Omega_E) \rightarrow 1$  and (C.1) gives  $\mathbb{P}(\Omega_{\text{KKT}}) = 1$ . These bounds imply  $\mathbb{P}(\Omega) \rightarrow 1$  by the union bound.

As the only randomness of the problem comes from  $(\boldsymbol{\varepsilon}, \mathbf{X})$ , we may choose the underlying probability space as  $\mathbb{R}^n \times \mathbb{R}^{n \times p}$ , so that  $\Omega$ ,  $\Omega_L$ ,  $\Omega_E$ ,  $\Omega_{\text{KKT}}$  are subsets of  $\mathbb{R}^n \times \mathbb{R}^{n \times p}$ . We next prove that  $\Omega$  is open as a subset of  $\mathbb{R}^n \times \mathbb{R}^{n \times p}$ . Indeed, because the KKT conditions are strict in  $\Omega$ ,  $\Omega$  is a disjoint union of sets of the form

$$(B.1) \quad \Omega_L \cap \Omega_E \cap \Omega_{\mathbf{X}} \cap \{\|\widehat{\boldsymbol{\beta}}_{G_k}\| > 0, k \in B\} \cap \{\|\mathbf{X}^\top(\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}})\| < n\lambda_k, k \in B^c\}$$

over all possible active groups  $B \subset \{1, \dots, K\}$ . The sets  $\Omega_E$ ,  $\Omega_{\mathbf{X}}$  are open as the inequalities are strict. In  $\Omega_E$ , the function  $(\boldsymbol{\varepsilon}, \mathbf{X}) \mapsto \widehat{\boldsymbol{\beta}}$  is locally Lipschitz by Lemma 7.1, hence continuous. By continuity, the preimage of the open set  $(0, +\infty)$  by the function  $\Omega_E \rightarrow \mathbb{R}$ ,  $(\boldsymbol{\varepsilon}, \mathbf{X}) \mapsto \|\widehat{\boldsymbol{\beta}}_{G_k}\|$  is open by continuity, and the preimage of the open set  $(-\infty, n\lambda_k)$  by the function  $\Omega_E \rightarrow \mathbb{R}$ ,  $(\boldsymbol{\varepsilon}, \mathbf{X}) \mapsto \|\mathbf{X}_{G_k}^\top(\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}})\|$  is also open, again by continuity. This shows that the set (B.1) is open for any fixed  $B \subset \{1, \dots, K\}$  so that  $\Omega$  is open as the union of sets of the form (B.1) over all  $B \subset \{1, \dots, K\}$  satisfying  $\sum_{k \in B} |G_k| \leq \kappa n/2$ . This proves that  $\Omega \subset \mathbb{R}^n \times \mathbb{R}^{n \times p}$  is open.

For  $F = 2 \max\{1, \|\boldsymbol{\Sigma}^{1/2}\mathbf{h}\|^2/(n\|\mathbf{X}\mathbf{h}\|^2)\}$  in Lemma 7.2, (7.12) is satisfied so that (7.13)–(7.14) hold. In  $\Omega$ , we thus have  $\|\boldsymbol{\Sigma}^{1/2}\mathbf{h}\|^2 \vee (\|\mathbf{X}\mathbf{h}\|^2/n) \leq F_+ F^2 R_* \leq 8\varphi^{-2} R_*$  and  $\|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|/\sqrt{n} \leq F_+^{1/2} \sigma + \sqrt{8}\varphi^{-1} R_*^{1/2} \leq 3\sqrt{2}\varphi^{-1} R_*$ . Furthermore,  $\|\mathbf{w}_0\|_2^2 \leq \varphi^{-1}/n$  in  $\Omega_E$  thanks to  $|\widehat{S}| \leq \kappa n/2$  and the explicit expression for  $\mathbf{w}_0$  in Proposition 4.2. In summary, we have in  $\Omega$ ,

$$(B.2) \quad \begin{aligned} \|\boldsymbol{\Sigma}^{1/2}\mathbf{h}\|^2 \vee (\|\mathbf{X}\mathbf{h}\|^2/n) &\leq 8\varphi^{-2} R_*, \\ \|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2/n &\leq 18\varphi^{-2} R_*, \\ \|\mathbf{w}_0\|^2 &\leq \varphi^{-2}/n \end{aligned}$$

which replace (7.10)–(7.11) in the present context. By the deterministic inequality (7.29), in  $\Omega$  we have  $\widehat{\mathbf{d}}\mathbf{f} \leq |\widehat{S}| \leq \kappa n/2$  since  $\widehat{\mathbf{H}}$  is rank at most  $|\widehat{S}|$  with operator norm at most one, so that

$$(B.3) \quad I_\Omega(1 - \kappa/2)^2/8 \leq \|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|^2/(nR_*) + \Delta_n^a + \Delta_n^b + \Delta_n^c.$$

Let  $(\boldsymbol{\varepsilon}, \mathbf{X})$ ,  $(\boldsymbol{\varepsilon}, \widetilde{\mathbf{X}})$  both in  $\Omega$ , let  $\widetilde{\boldsymbol{\varepsilon}} = \boldsymbol{\varepsilon}$  and let  $\mathbf{h}, \widetilde{\mathbf{h}}, \mathbf{f}, \widetilde{\mathbf{f}}$  be as in Lemma 7.1. Thanks to event  $\Omega_E$  and the fact that  $|\widehat{S}| \leq \kappa n/2$ , and similarly for  $\widetilde{\boldsymbol{\beta}}$ , we have  $\varphi^2 \|\boldsymbol{\Sigma}^{1/2}(\mathbf{h} - \widetilde{\mathbf{h}})\|^2 \leq \|\mathbf{X}(\mathbf{h} - \widetilde{\mathbf{h}})\|^2/(2n) + \|\widetilde{\mathbf{X}}(\mathbf{h} - \widetilde{\mathbf{h}})\|^2/(2n)$ . Thus, by (7.3),

$$n\varphi^2 \|\boldsymbol{\Sigma}^{1/2}(\mathbf{h} - \widetilde{\mathbf{h}})\|^2 \leq (\widetilde{\mathbf{h}} - \mathbf{h})^\top (\mathbf{X} - \widetilde{\mathbf{X}})^\top \boldsymbol{\varepsilon} + (\mathbf{h} - \widetilde{\mathbf{h}})^\top (\mathbf{X}^\top \mathbf{X} - \widetilde{\mathbf{X}}^\top \widetilde{\mathbf{X}})(\mathbf{h} + \widetilde{\mathbf{h}})/2.$$

Summing this inequality with the first line in (7.2), we find

$$(B.4) \quad \begin{aligned} n\varphi^2 \|\boldsymbol{\Sigma}^{1/2}(\mathbf{h} - \widetilde{\mathbf{h}})\|^2 + \|\mathbf{f} - \widetilde{\mathbf{f}}\|^2 \\ \leq (\widetilde{\mathbf{h}} - \mathbf{h})^\top (\mathbf{X} - \widetilde{\mathbf{X}})^\top \boldsymbol{\varepsilon} + (\mathbf{h} - \widetilde{\mathbf{h}})^\top (\mathbf{X}^\top \mathbf{X} - \widetilde{\mathbf{X}}^\top \widetilde{\mathbf{X}})(\mathbf{h} + \widetilde{\mathbf{h}})/2 \\ + (\widetilde{\mathbf{h}} - \mathbf{h})^\top (\mathbf{X} - \widetilde{\mathbf{X}})^\top \mathbf{f} + \mathbf{h}^\top (\mathbf{X} - \widetilde{\mathbf{X}})^\top (\mathbf{f} - \widetilde{\mathbf{f}}). \end{aligned}$$

Thanks to the bounds in (B.2), this implies  $\|\mathbf{f} - \widetilde{\mathbf{f}}\| \leq L\|(\mathbf{X} - \widetilde{\mathbf{X}})\boldsymbol{\Sigma}^{-1/2}\|_{\text{op}}$  if  $\{(\boldsymbol{\varepsilon}, \mathbf{X}), (\boldsymbol{\varepsilon}, \widetilde{\mathbf{X}})\} \subset \Omega$ , where  $L = C_{27}(\gamma, \kappa) R_*^{1/2}$ .

For a given  $(\boldsymbol{\varepsilon}, \mathbf{X}\mathbf{Q}_0)$ , we define  $U_0 = \{\mathbf{z}_0 \in \mathbb{R}^n : (\boldsymbol{\varepsilon}, \mathbf{X}\mathbf{Q}_0 + \mathbf{z}_0\mathbf{a}_0^\top) \in \Omega\}$ . In  $U_0$ , the function  $f(\mathbf{z}_0) = \mathbf{X}(\widehat{\boldsymbol{\beta}} - \boldsymbol{\beta}) - \boldsymbol{\varepsilon}$  is  $L$ -Lipschitz. By Kirszbraun's theorem, there exists a function  $F : \mathbb{R}^n \rightarrow \mathbb{R}^n$  that is an extension of  $f$ , that is,  $F(\mathbf{z}_0) = f(\mathbf{z}_0)$  for  $\mathbf{z}_0 \in U_0$ , and such that  $F$  is  $L$ -Lipschitz in the whole  $\mathbb{R}^n$ . Note that both function  $F$  and  $f$  implicitly depend on  $(\boldsymbol{\varepsilon}, \mathbf{X}\mathbf{Q}_0)$ . Since  $\Omega$  is open,  $U_0$  is also open, and thus conditionally on  $(\mathbf{X}\mathbf{Q}_0, \boldsymbol{\varepsilon})$ ,

$$(B.5) \quad \nabla f(\mathbf{z}_0) = \nabla F(\mathbf{z}_0), \quad \text{for all } \mathbf{z}_0 \in U_0.$$

(Without the openness of  $\Omega$  established above, equality of the gradients would be unclear).

Since  $F : \mathbb{R}^n \rightarrow \mathbb{R}^n$  is such that  $F(\mathbf{z}_0) = f(\mathbf{z}_0)$  in  $\Omega$ , by (B.3),

$$(B.6) \quad \begin{aligned} (1 - \kappa/2)^2 I_\Omega &\leq I_\Omega[\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2/(nR_*) + \Delta_n^a + \Delta_n^b + \Delta_n^c] \\ &= I_\Omega[\|F(\mathbf{z}_0)\|^2/(nR_*) + \Delta_n^a + \Delta_n^b + \Delta_n^c]. \end{aligned}$$

Taking conditional expectations and multiplying both sides by  $\delta_1^2 \stackrel{\text{def}}{=} \mathbb{E}_0[\|\nabla F(\mathbf{z}_0)\|_F^2]/\{\mathbb{E}_0[\|\nabla F(\mathbf{z}_0)\|_F^2] + \mathbb{E}_0[\|F(\mathbf{z}_0)\|^2]\}$ , we find

$$\delta_1^2(1 - \kappa/2)^2 \mathbb{E}_0[I_\Omega] \leq \mathbb{E}_0[\|\nabla F(\mathbf{z}_0)\|_F^2/(nR_*)] + \mathbb{E}_0[I_\Omega(\Delta_n^a + \Delta_n^b + \Delta_n^c)],$$

due to  $\delta_1^2 \mathbb{E}_0[I_\Omega\|F(\mathbf{z}_0)\|^2] \leq \mathbb{E}_0[\|\nabla F(\mathbf{z}_0)\|_F^2]$  for the first term and  $\delta_1^2 \leq 1$  for the second. Using  $\delta_1^2 \leq 1$  and  $1 = I_\Omega + I_{\Omega^c}$ ,

$$\begin{aligned} \mathbb{E}[\delta_1^2](1 - \kappa/2)^2 &\leq \mathbb{E}[I_\Omega\|\nabla F(\mathbf{z}_0)\|_F^2/(nR_*)] + \mathbb{E}[I_\Omega(\Delta_n^a + \Delta_n^b + \Delta_n^c)] \\ &\quad + \mathbb{P}[\Omega^c](1 + L^2/R_*), \end{aligned}$$

where we used that  $\|\nabla F(\mathbf{z}_0)\|_F^2 \leq n\|\nabla F(\mathbf{z}_0)\|_{\text{op}}^2 \leq nL^2$  in  $\Omega^c$  since  $F$  is  $L$ -Lipschitz. We now prove that the three terms on the right-hand side converge to 0. For the third term,  $L^2/R_* \leq C_{28}(\gamma, \kappa)$  and  $\mathbb{P}(\Omega^c) \rightarrow 0$  as  $\Omega$  has probability approaching one. For the first term, since  $F$  is  $L$ -Lipschitz,  $\|\nabla F(\mathbf{z}_0)\|_F^2 \leq nL^2$  almost surely so that the sequence of random variables  $I_\Omega\|\nabla F(\mathbf{z}_0)\|_F^2/(R_*n)$  is uniformly integrable. Thanks to uniform integrability, if we can prove  $\|\nabla F(\mathbf{z}_0)\|_F^2/(R_*n) \xrightarrow{\mathbb{P}} 0$ , then  $\mathbb{E}[I_\Omega\|\nabla F(\mathbf{z}_0)\|_F^2/(R_*n)] \rightarrow 0$  holds. We use that  $I_\Omega \nabla f(\mathbf{z}_0) = I_\Omega \nabla F(\mathbf{z}_0)$  by (B.5), and that in  $\Omega$  the gradients of  $f$  with respect to  $\mathbf{z}_0$  are given in Proposition 4.2, so that by (B.2),

$$\begin{aligned} I_\Omega\|\nabla F(\mathbf{z}_0)\|_F/(R_*n)^{1/2} &= I_\Omega\|\nabla f(\mathbf{z}_0)\|_F/(R_*n)^{1/2} \\ &\leq I_\Omega[\|\mathbf{w}_0\|\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\| + \|\mathbf{I}_n - \hat{\mathbf{H}}\|_F|\langle \mathbf{a}_0, \mathbf{h} \rangle|]/(R_*n)^{1/2} \\ &\leq C_{29}(\gamma, \kappa)(n^{-1/2} + |\langle \mathbf{a}_0, \mathbf{h} \rangle|/R_*^{1/2}), \end{aligned}$$

which converges to 0 in probability thanks to assumption  $\langle \mathbf{a}_0, \hat{\boldsymbol{\beta}} - \boldsymbol{\beta} \rangle^2/R_* \xrightarrow{\mathbb{P}} 0$ . Thanks to uniform integrability, this proves  $\mathbb{E}[I_\Omega\|\nabla F(\mathbf{z}_0)\|_F^2/(\sigma^2 n)] \rightarrow 0$ . It remains to show  $\mathbb{E}[I_\Omega(\Delta_n^a + \Delta_n^b + \Delta_n^c)] \rightarrow 0$ . By definition of  $\Delta_n^b$  in (7.27), thanks to  $\Omega_\chi$  and (B.2) we have  $I_\Omega\Delta_n^b \leq 18\varphi^{-2}(F_+ - 1) \leq C_{30}(\gamma, \kappa)\sqrt{\log(n)/n}$ . For  $\Delta_n^a$  in (7.26), let  $\boldsymbol{\Pi} : \mathbb{R}^n \rightarrow \mathbb{R}^n$  be the convex projection onto the Euclidean ball of radius  $\sqrt{18\varphi^{-2}R_*}$ , then  $\boldsymbol{\Pi}(\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}) = \mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}$  in  $\Omega$  by (B.2) so that

$$(B.7) \quad \begin{aligned} \mathbb{E}[I_\Omega\Delta_n^a] &= \mathbb{E}[I_\Omega\{(1 - \widehat{\text{df}}/n) - \boldsymbol{\varepsilon}^\top \boldsymbol{\Pi}(\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}})/(n\sigma^2)\}^2]\sigma^2/R_* \\ &\leq \mathbb{E}[\|\boldsymbol{\Pi}(\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}})\|^2]/(n^2 R_*) + \sigma^2/(nR_*) \\ &\leq 18\varphi^{-2}/n + 1/n \end{aligned}$$

by applying Proposition 2.1 to the function  $\boldsymbol{\varepsilon} \mapsto \boldsymbol{\Pi}(\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}})$ , which is 1-Lipschitz as the composition of two 1-Lipschitz functions (cf. Proposition 7.3(i)). For  $\Delta_n^c$  in (7.28), let  $\mathbf{g}, \mathbf{a}_*$  be as in Lemma 7.6 and set  $\mathbf{u}_* = \boldsymbol{\Sigma}^{-1}\mathbf{a}_*$ ,  $\mathbf{Q}_* = \mathbf{I}_p - \mathbf{u}_*\mathbf{a}_*^\top$  and note that  $(\mathbf{g}, \mathbf{u}_*, \mathbf{Q}_*) = (\mathbf{z}_0, \mathbf{u}_0, \mathbf{Q}_0)$  for  $\mathbf{a}_0 = \mathbf{a}_*$ . Let also  $\mathbf{w}_*$  be the  $\mathbf{w}_0$  from Proposition 4.2 for  $\mathbf{a}_0 = \mathbf{a}_*$ . As above for  $\mathbf{a}_0$ , for a fixed  $(\boldsymbol{\varepsilon}, \mathbf{X}\mathbf{Q}_*)$  the function  $\mathbf{g} \mapsto \mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}$  is  $L$ -Lipschitz in  $U_* = \{\mathbf{g} \in \mathbb{R}^n : (\boldsymbol{\varepsilon}, \mathbf{X}\mathbf{Q}_* + \mathbf{g}\mathbf{a}_*^\top) \in \Omega\}$  by (B.4) for the value of  $L$  given after (B.4). Furthermore,  $\|\mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}\|^2 \leq 18n\varphi^{-2}R_*$  in  $\Omega$ . By Kirszbraun's theorem, there exists an extension  $F_* : \mathbb{R}^n \rightarrow \mathbb{R}^n$  implicitly depending on  $(\boldsymbol{\varepsilon}, \mathbf{X}\mathbf{Q}_*)$  such that  $F_*(\mathbf{g}) = \mathbf{y} - \mathbf{X}\hat{\boldsymbol{\beta}}$  in  $\Omega$

and  $\|F_*(\mathbf{g})\|^2 \leq 18n\varphi^{-2}R_*$  by composing the extension given by Kirszbraun's theorem by the convex projection onto the Euclidean ball of radius  $(18n\varphi^{-2}R_*)^{1/2}$ . By Proposition 2.1, with respect to  $\mathbf{g}$  conditionally on  $(\boldsymbol{\varepsilon}, \mathbf{X}\mathbf{Q}_*)$ ,

$$\begin{aligned} \mathbb{E}[I_\Omega((n - \widehat{\text{df}})\langle \mathbf{a}_*, \mathbf{h} \rangle + \langle \mathbf{w}_* + \mathbf{g}, \mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}} \rangle)^2] &= \mathbb{E}[I_\Omega(\text{div } F_*(\mathbf{g}) + \langle \mathbf{g}, F_*(\mathbf{g}) \rangle)^2] \\ &\leq 18n\varphi^{-2}R_* + nL^2. \end{aligned}$$

For the value of  $L$  given after (B.4) and using the bound (B.2) to control  $\langle \mathbf{w}_*, \mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}} \rangle$  in  $\Omega$ , this gives  $\mathbb{E}[I_\Omega \Delta_n^c] \leq C_{31}(\gamma, \kappa)n^{-1}$ .

This proves  $(1 - \kappa/2)^2 \mathbb{E}[\delta_1^2] \rightarrow 0$ . Consequently,  $\Xi_0 = \mathbf{z}_0^\top F(\mathbf{z}_0) - \text{div } F(\mathbf{z}_0)$  satisfies  $|\sup_t |\mathbb{P}(\Xi_0/\|F(\mathbf{z}_0)\| \leq t) - \Phi(t)| \rightarrow 0$  by (2.8). Since  $\xi_0 = \mathbf{z}_0^\top f(\mathbf{z}_0) - \text{div } f(\mathbf{z}_0)$  is equal to  $\Xi_0$  on the event  $\Omega$  because  $F$  is an extension of  $f$ , we have  $|\mathbb{P}(\xi_0/\|f(\mathbf{z}_0)\| \leq t) - \mathbb{P}(\Xi_0/\|F(\mathbf{z}_0)\| \leq t)| \leq 2\mathbb{P}(\Omega^c) \rightarrow 0$ , so that  $\sup_t |\mathbb{P}(\xi_0/\|f(\mathbf{z}_0)\| \leq t) - \Phi(t)| \rightarrow 0$  as well. The conclusion (3.35) is obtained by controlling the term  $\mathbf{w}_0^\top (\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}})/\|\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}\|$  by  $\|\mathbf{w}_0\|$ , which is bounded as in (B.2) in  $\Omega$ .

It remains to show that (3.35) holds uniformly over all  $\mathbf{a}_0 \in \Sigma^{1/2}\bar{S}$  and to derive the properties of  $\bar{S}$ . The proof of the relative volume bound on  $\bar{S}$  and the lower bound on the cardinality of  $\{j \in [p] : \mathbf{e}_j/\|\Sigma^{-1/2}\mathbf{e}_j\| \in \Sigma^{1/2}\bar{S}\}$  is the same as in the proof of Theorem 3.7 given around (7.39), and for  $\mathbf{a}_0 \in \Sigma^{1/2}\bar{S}$  inequality (7.39) holds. For such  $\mathbf{a}_0$ ,  $\mathbb{E}[I_\Omega \langle \mathbf{a}_0, \mathbf{h} \rangle^2 / R_*] \leq 8\varphi^{-2} \mathbb{E}[\langle \mathbf{a}_0, \mathbf{h} \rangle^2 / \|\Sigma^{1/2}\mathbf{h}\|^2] \leq 8\varphi^{-2}C^*/a_p \rightarrow 0$  by (B.2) for the first inequality and (7.39) for the second.  $\square$

## APPENDIX C: STRICT KKT CONDITIONS WITH PROBABILITY ONE FOR THE GROUP LASSO

LEMMA C.1. Consider a design matrix  $\mathbf{X} \in \mathbb{R}^{n \times p}$  and a response vector  $\mathbf{y} \in \mathbb{R}^n$  for which the joint distribution of  $(\mathbf{X}, \mathbf{y})$  admits a density with respect to the Lebesgue measure. Consider a partition of  $\{1, \dots, p\}$  into groups  $(G_1, \dots, G_K)$  and any minimizer

$$\widehat{\boldsymbol{\beta}} \in \arg \min_{\mathbf{b} \in \mathbb{R}^p} \frac{1}{2n} \|\mathbf{X}\mathbf{b} - \mathbf{y}\|^2 + \|\mathbf{b}\|_{\text{GL}}, \quad \|\mathbf{b}\|_{\text{GL}} \stackrel{\text{def}}{=} \sum_{k=1, \dots, K} \lambda_k \|\mathbf{b}_{G_k}\|_2$$

for some deterministic  $\lambda_1, \dots, \lambda_K > 0$ . There exists an open set  $U \subset \mathbb{R}^{n \times (1+p)}$  such that  $\mathbb{P}((\mathbf{y}, \mathbf{X}) \in U) = 1$  and the KKT conditions are strict in  $\{(\mathbf{y}, \mathbf{X}) \in U\}$  in the sense that

$$(C.1) \quad \{(\mathbf{y}, \mathbf{X}) \in U\} \subset \{\forall k = 1, \dots, K, \widehat{\boldsymbol{\beta}}_{G_k} = 0 \Rightarrow \|\mathbf{X}_{G_k}^\top (\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}})\|_2 < n\lambda_k\}.$$

Finally,  $\widehat{B} = \{k \in [K] : \|\widehat{\boldsymbol{\beta}}_{G_k}\| > 0\}$  is constant in a small neighborhood of any point in  $U$ .

PROOF OF LEMMA C.1. Consider a fixed  $B \subset \{1, \dots, K\}$  and its complementary set  $B^c$ , and consider the group Lasso estimator  $\widehat{\boldsymbol{\beta}}(B)$  with the additional constraint  $\mathbf{b}_{G_k} = 0$  for every  $k \in B^c$ . Now consider a group  $k \in B^c$ . Since the joint distribution of  $(\mathbf{X}, \mathbf{y})$  has a density with respect to the Lebesgue measure, the conditional distribution of  $\mathbf{X}_{G_k}$  given  $(\mathbf{y}, (\mathbf{X}\mathbf{e}_j)_{j \notin G_k})$  also admits a density with respect to the Lebesgue measure. Conditionally, on  $(\mathbf{y}, (\mathbf{X}\mathbf{e}_j)_{j \notin G_k})$ , two cases may appear:

- (i) If  $\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}(B) = 0$ , the KKT condition for group  $G_k$  hold strictly since  $\lambda_k \neq 0$ .
- (ii) If  $\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}(B) \neq 0$ , the distribution of  $\mathbf{X}_{G_k}$  given  $(\mathbf{y}, (\mathbf{X}\mathbf{e}_j)_{j \notin G_k})$  and the distribution of  $\mathbf{X}_{G_k}^\top (\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}(B))$  given  $(\mathbf{y}, (\mathbf{X}\mathbf{e}_j)_{j \notin G_k})$  both have a density with respect to the Lebesgue measure. The sphere of radius  $n\lambda_k$  has measure 0 for any continuous distribution, hence

$$\mathbb{P}(\|\mathbf{X}_{G_k}^\top (\mathbf{y} - \mathbf{X}\widehat{\boldsymbol{\beta}}(B))\|_2 \neq n\lambda_k | \mathbf{y}, (\mathbf{X}\mathbf{e}_j)_{j \notin G_k}) = 1.$$

Finally, the unconditional probability  $\mathbb{P}(\|X_{G_k}^\top(\mathbf{y} - X\hat{\boldsymbol{\beta}}(B))\|_2 \neq n\lambda_k)$  is also one. Let  $U = \bigcap_{B \subset \{1, \dots, K\}} \bigcap_{k \notin B} \{(\mathbf{y}, X) : \|X_{G_k}^\top(\mathbf{y} - X\hat{\boldsymbol{\beta}}(B))\|_2 \neq n\lambda_k\}$ . Then  $\mathbb{P}((\mathbf{y}, X) \in U) = 1$  as a finite intersection of events of probability one and (C.1) holds. The set  $U$  is open as a finite intersection of open sets, since  $\{(\mathbf{y}, X) : \|X_{G_k}^\top(\mathbf{y} - X\hat{\boldsymbol{\beta}}(B))\|_2 \neq n\lambda_k\}$  is open by continuity of  $(\mathbf{y}, X) \mapsto X^\top(\mathbf{y} - X\hat{\boldsymbol{\beta}}(B))$  by the claim following (7.2).

Next, to show that  $\hat{B}$  is constant in a neighborhood of every point in  $U$ , set  $U_\delta = \bigcap_{B \subset \{1, \dots, K\}} \bigcap_{k \notin B} \{(\mathbf{y}, X) : \|\|X_{G_k}^\top(\mathbf{y} - X\hat{\boldsymbol{\beta}}(B))\|_2/(n\lambda_k) - 1\| > \delta\}$  for all  $\delta > 0$ . We have  $U = \bigcup_{\delta > 0} U_\delta$  and the set  $U_\delta$  is open by continuity of  $(\mathbf{y}, X) \mapsto \|\|X_{G_k}^\top(\mathbf{y} - X\hat{\boldsymbol{\beta}}(B))\|_2/(n\lambda_k) - 1\|$ , which follows from the continuity of  $(\mathbf{y}, X) \mapsto X^\top(\mathbf{y} - X\hat{\boldsymbol{\beta}}(B))$  by the claim following (7.2). For any  $(\bar{\mathbf{y}}, \bar{X}) \in U$ , there exists some  $\delta > 0$  with  $(\bar{\mathbf{y}}, \bar{X}) \in U_\delta$ . Let  $\bar{B} = \{k \in [K] : \|\bar{\boldsymbol{\beta}}_{G_k}\| > 0\}$ . By continuity of  $(\mathbf{y}, X) \mapsto X^\top(\mathbf{y} - X\hat{\boldsymbol{\beta}})$  thanks to the claim following (7.2), there exists a neighborhood  $\mathcal{N}$  of  $(\bar{\mathbf{y}}, \bar{X})$  with  $\mathcal{N} \subset U_\delta$  such that for all  $(\mathbf{y}, X) \in \mathcal{N}$ ,  $\|X_{G_k}^\top(\mathbf{y} - X\hat{\boldsymbol{\beta}})\|/(n\lambda_k) < 1 - \delta/2$  for  $k \notin \bar{B}$  and  $\|X_{G_k}^\top(\mathbf{y} - X\hat{\boldsymbol{\beta}})\|/(n\lambda_k) > 1 - \delta/2$  for  $k \in \bar{B}$ . Since  $\mathcal{N} \subset U_\delta$ ,  $\|X_{G_k}^\top(\mathbf{y} - X\hat{\boldsymbol{\beta}})\|/(n\lambda_k) > 1 - \delta/2$  implies  $\|X_{G_k}^\top(\mathbf{y} - X\hat{\boldsymbol{\beta}})\|/(n\lambda_k) = 1$ , so that  $\hat{B} = \bar{B}$  in  $\mathcal{N}$ .  $\square$

**PROPOSITION 4.2.** *The following holds for almost every  $(\bar{\mathbf{y}}, \bar{X}) \in \mathbb{R}^{n \times (1+p)}$ . The set  $\bar{B} = \{k \in [K] : \|\bar{\boldsymbol{\beta}}_{G_k}\| > 0\}$  of active groups is the same for all minimizers  $\bar{\boldsymbol{\beta}}$  of (3.33) at  $(\bar{\mathbf{y}}, \bar{X})$  and  $\hat{B} = \bar{B}$  for all  $(\mathbf{y}, X)$  in a sufficiently small neighborhood of  $(\bar{\mathbf{y}}, \bar{X})$ . If additionally  $\bar{X}_S^\top \bar{X}_S$  is invertible where  $\bar{S} = \bigcup_{k \in \bar{B}} G_k$  then the map  $(\mathbf{y}, X) \mapsto \hat{\boldsymbol{\beta}}$  is Lipschitz in a sufficiently small neighborhood of  $(\bar{\mathbf{y}}, \bar{X})$ . In this neighborhood, we have*

$$\begin{aligned} [\nabla \hat{\boldsymbol{\beta}}(\mathbf{z}_0)]_{\hat{S}^c} &= 0, \quad [\nabla \hat{\boldsymbol{\beta}}(\mathbf{z}_0)]_{\hat{S}}^\top = (X_S^\top X_S + \mathbf{M})^{-1}[(\mathbf{a}_0)_{\hat{S}}^\top (\mathbf{y} - X\hat{\boldsymbol{\beta}})^\top - \langle \mathbf{a}_0, \mathbf{h} \rangle X_S^\top], \\ \hat{\mathbf{H}} &= X_S (X_S^\top X_S + \mathbf{M})^{-1} X_S^\top \text{ and (3.9) holds with } \mathbf{w}_0 = X_S (X_S^\top X_S + \mathbf{M})^{-1} (\mathbf{a}_0)_{\hat{S}}. \end{aligned}$$

**PROOF OF PROPOSITION 4.2.** By Lemma C.1,  $\hat{B}$  and  $\hat{S}$  are constant in a sufficiently small neighborhood of almost every  $(\bar{\mathbf{y}}, \bar{X})$ . The additional assumption that  $\bar{X}_S^\top \bar{X}_S$  is invertible provides that  $X_S^\top X_S$  is invertible by continuity of the smallest eigenvalue in a small enough compact neighborhood of  $(\bar{\mathbf{y}}, \bar{X})$ , and in this neighborhood  $(\mathbf{y}, X) \mapsto \hat{\boldsymbol{\beta}}$  is Lipschitz by the sentence following (7.4), and thus almost everywhere differentiable by Rademacher's theorem. The formulae for  $\nabla \hat{\boldsymbol{\beta}}(\mathbf{z}_0)$ ,  $\hat{\mathbf{H}}$  and  $\mathbf{w}_0$  involving the matrix  $\mathbf{M}$  in (4.5) are then obtained by differentiating the KKT conditions restricted to  $\hat{S}$  in this neighborhood, that is,  $X_{G_k}^\top(\mathbf{y} - X\hat{\boldsymbol{\beta}}) = n\lambda_k \hat{\boldsymbol{\beta}}_{G_k} / \|\hat{\boldsymbol{\beta}}_{G_k}\|$  for all  $k \in \hat{B}$ .  $\square$

#### APPENDIX D: PROOF OF THEOREM 2.3

**PROOF OF THEOREM 2.3.** With  $\mathbb{E}[\|\bar{\boldsymbol{\mu}} + \bar{\mathbf{A}}^\top \mathbf{z}\|^2] = \|\bar{\boldsymbol{\mu}}\|^2 + \|\bar{\mathbf{A}}\|_F^2$  in mind, consider

$$\begin{aligned} \widehat{\text{Var}}[\xi] &= \|f(\mathbf{z}) - (\bar{\boldsymbol{\mu}} + \bar{\mathbf{A}}^\top \mathbf{z})\|^2 + \text{trace}[\{\nabla f(\mathbf{z}) - \bar{\mathbf{A}}\}^2] \\ &\quad + 2(f(\mathbf{z}) - \bar{\boldsymbol{\mu}} - \bar{\mathbf{A}}^\top \mathbf{z})^\top (\bar{\boldsymbol{\mu}} + \bar{\mathbf{A}}^\top \mathbf{z}) + 2\text{trace}[\{\nabla f(\mathbf{z}) - \bar{\mathbf{A}}\} \bar{\mathbf{A}}] \\ &\quad + (\|\bar{\boldsymbol{\mu}} + \bar{\mathbf{A}}^\top \mathbf{z}\|^2 - \|\bar{\boldsymbol{\mu}}\|^2 - \|\bar{\mathbf{A}}\|_F^2) + \|\bar{\boldsymbol{\mu}}\|^2 + \|\bar{\mathbf{A}}\|_F^2 + \text{trace}[\bar{\mathbf{A}}^2]. \end{aligned}$$

By the triangle and Cauchy–Schwarz inequalities,

$$\begin{aligned} &\mathbb{E}[\|f(\mathbf{z})\|^2 + \text{trace}[\{\nabla f(\mathbf{z})\}^2] - \text{Var}[\xi]] \\ &\leq \mathbb{E}[\|f(\mathbf{z}) - (\bar{\boldsymbol{\mu}} + \bar{\mathbf{A}}^\top \mathbf{z})\|^2] + \mathbb{E}[\|\nabla f(\mathbf{z}) - \bar{\mathbf{A}}\|_F^2] \end{aligned}$$

$$\begin{aligned}
& + 2\{\mathbb{E}[\|f(\mathbf{z}) - (\bar{\boldsymbol{\mu}} + \bar{\mathbf{A}}^\top \mathbf{z})\|^2] + \mathbb{E}[\|\nabla f(\mathbf{z}) - \bar{\mathbf{A}}\|_F^2]\}^{1/2} \{\|\bar{\boldsymbol{\mu}}\|^2 + 2\|\bar{\mathbf{A}}\|_F^2\}^{1/2} \\
& + \mathbb{E}[\|\bar{\boldsymbol{\mu}} + \bar{\mathbf{A}}^\top \mathbf{z}\|^2 - \|\bar{\boldsymbol{\mu}}\|^2 - \|\bar{\mathbf{A}}\|_F^2] + \|\bar{\boldsymbol{\mu}}\|^2 + \|\bar{\mathbf{A}}\|_F^2 + \text{trace}(\bar{\mathbf{A}}^2) - \text{Var}[\xi].
\end{aligned}$$

We have  $\mathbb{E}[\|f(\mathbf{z}) - (\bar{\boldsymbol{\mu}} + \bar{\mathbf{A}}^\top \mathbf{z})\|^2] \leq \mathbb{E}[\|\nabla f(\mathbf{z}) - \bar{\mathbf{A}}\|_F^2] \leq \bar{\epsilon}_{1,2}^2 \text{Var}[\xi]/2$  by the Gaussian Poincaré inequality,  $\mathbb{E}[\|\bar{\boldsymbol{\mu}} + \bar{\mathbf{A}}^\top \mathbf{z}\|^2 - \|\bar{\boldsymbol{\mu}}\|^2 - \|\bar{\mathbf{A}}\|_F^2] = \|\bar{\mathbf{A}}\bar{\boldsymbol{\mu}}\|^2 + 2\text{trace}\{(\bar{\mathbf{A}}\bar{\mathbf{A}}^\top)^2\} \leq \|\bar{\mathbf{A}}\|_{\text{op}}^2 C_0^2 \text{Var}[\xi]$ , and  $0 \leq 1 - \{\|\bar{\boldsymbol{\mu}}\|^2 + \|\bar{\mathbf{A}}\|_F^2 + \text{trace}(\bar{\mathbf{A}}^2)\} / \text{Var}[\xi] \leq \bar{\epsilon}_{1,2}^2$  as in (2.11). Thus, (2.14) holds and the conclusions follow.  $\square$

**Funding.** P.C. Bellec’s Research was partially supported by the NSF Grants DMS-1811976 and DMS-1945428.

C.-H. Zhang’s research was partially supported by the NSF Grants DMS-1721495, IIS-1741390, CCF-1934924, DMS-2052949 and DMS-2210850.

## REFERENCES

- [1] BAYATI, M. and MONTANARI, A. (2012). The LASSO risk for Gaussian matrices. *IEEE Trans. Inf. Theory* **58** 1997–2017. [MR2951312](#) <https://doi.org/10.1109/TIT.2011.2174612>
- [2] BELLEC, P. and TSYBAKOV, A. (2017). Bounds on the prediction error of penalized least squares estimators with convex penalty. In *Modern Problems of Stochastic Analysis and Statistics. Springer Proc. Math. Stat.* **208** 315–333. Springer, Cham. [MR3747672](#) [https://doi.org/10.1007/978-3-319-65313-6\\_13](https://doi.org/10.1007/978-3-319-65313-6_13)
- [3] BELLEC, P. C. (2018). The noise barrier and the large signal bias of the lasso and other convex estimators. Available at: [arXiv:1804.01230](#), <https://arxiv.org/pdf/1804.01230.pdf>.
- [4] BELLEC, P. C. (2020). Out-of-sample error estimate for robust m-estimators with convex penalty. arXiv preprint. Available at [arXiv:2008.11840](#).
- [5] BELLEC, P. C., LECUÉ, G. and TSYBAKOV, A. B. (2018). Slope meets Lasso: Improved oracle bounds and optimality. *Ann. Statist.* **46** 3603–3642. [MR3852663](#) <https://doi.org/10.1214/17-AOS1670>
- [6] BELLEC, P. C. and SHEN, Y. (2022). Derivatives and residual distribution of regularized m-estimators with application to adaptive tuning. In *Conference on Learning Theory* 1912–1947. PMLR.
- [7] BELLEC, P. C. and ZHANG, C.-H. (2021). Second-order Stein: SURE for SURE and other applications in high-dimensional inference. *Ann. Statist.* **49** 1864–1903. [MR4319234](#) <https://doi.org/10.1214/20-aos2005>
- [8] BELLEC, P. C. and ZHANG, C.-H. (2022). De-biasing the lasso with degrees-of-freedom adjustment. *Bernoulli* **28** 713–743. [MR4389062](#) <https://doi.org/10.3150/21-BEJ1348>
- [9] BELLONI, A., CHERNOZHUKOV, V. and HANSEN, C. (2014). High-dimensional methods and inference on structural and treatment effects. *J. Econ. Perspect.* **28** 29–50.
- [10] BICKEL, P. J., KLAASSEN, C. A. J., RITOV, Y. and WELLNER, J. A. (1993). *Efficient and Adaptive Estimation for Semiparametric Models. Johns Hopkins Series in the Mathematical Sciences.* Johns Hopkins Univ. Press, Baltimore, MD. [MR1245941](#)
- [11] BICKEL, P. J., RITOV, Y. and TSYBAKOV, A. B. (2009). Simultaneous analysis of lasso and Dantzig selector. *Ann. Statist.* **37** 1705–1732. [MR2533469](#) <https://doi.org/10.1214/08-AOS620>
- [12] BLANCHARD, J. D., CARTIS, C. and TANNER, J. (2011). Compressed sensing: How sharp is the restricted isometry property? *SIAM Rev.* **53** 105–125. [MR2785881](#) <https://doi.org/10.1137/090748160>
- [13] BLUMENSATH, T. and DAVIES, M. E. (2009). Iterative hard thresholding for compressed sensing. *Appl. Comput. Harmon. Anal.* **27** 265–274. [MR2559726](#) <https://doi.org/10.1016/j.acha.2009.04.002>
- [14] BOGACHEV, V. I. (1998). *Gaussian Measures. Mathematical Surveys and Monographs* **62**. Amer. Math. Soc., Providence, RI. [MR1642391](#) <https://doi.org/10.1090/surv/062>
- [15] BU, Z., KLUSOWSKI, J., RUSH, C. and SU, W. (2019). Algorithmic analysis and statistical estimation of slope via approximate message passing. In *Advances in Neural Information Processing Systems* 9361–9371.
- [16] CAI, T. T. and GUO, Z. (2017). Confidence intervals for high-dimensional linear regression: Minimax rates and adaptivity. *Ann. Statist.* **45** 615–646. [MR3650395](#) <https://doi.org/10.1214/16-AOS1461>
- [17] CELENTANO, M. and MONTANARI, A. (2022). Fundamental barriers to high-dimensional regression with convex penalties. *Ann. Statist.* **50** 170–196. [MR4382013](#) <https://doi.org/10.1214/21-aos2100>
- [18] CELENTANO, M., MONTANARI, A. and WEI, Y. (2020). The lasso with general gaussian designs with applications to hypothesis testing. arXiv preprint. Available at [arXiv:2007.13716](#).

- [19] CHATTERJEE, S. (2009). Fluctuations of eigenvalues and second order Poincaré inequalities. *Probab. Theory Related Fields* **143** 1–40. MR2449121 <https://doi.org/10.1007/s00440-007-0118-6>
- [20] DAVIDSON, K. R. and SZAREK, S. J. (2001). Local operator theory, random matrices and Banach spaces. In *Handbook of the Geometry of Banach Spaces, Vol. I* 317–366. North-Holland, Amsterdam. MR1863696 [https://doi.org/10.1016/S1874-5849\(01\)80010-3](https://doi.org/10.1016/S1874-5849(01)80010-3)
- [21] DONOHO, D. and MONTANARI, A. (2016). High dimensional robust M-estimation: Asymptotic variance via approximate message passing. *Probab. Theory Related Fields* **166** 935–969. MR3568043 <https://doi.org/10.1007/s00440-015-0675-z>
- [22] EDELMAN, A. (1988). Eigenvalues and condition numbers of random matrices. *SIAM J. Matrix Anal. Appl.* **9** 543–560. MR0964668 <https://doi.org/10.1137/0609045>
- [23] EL KAROUI, N., BEAN, D., BICKEL, P. J., LIM, C. and YU, B. (2013). On robust regression with high-dimensional predictors. *Proc. Natl. Acad. Sci. USA* **110** 14557–14562.
- [24] FAN, J. and LI, R. (2001). Variable selection via nonconcave penalized likelihood and its oracle properties. *J. Amer. Statist. Assoc.* **96** 1348–1360. MR1946581 <https://doi.org/10.1198/016214501753382273>
- [25] FENG, L. and ZHANG, C.-H. (2019). Sorted concave penalized regression. *Ann. Statist.* **47** 3069–3098. MR4025735 <https://doi.org/10.1214/18-AOS1759>
- [26] JAVANMARD, A. and MONTANARI, A. (2014). Confidence intervals and hypothesis testing for high-dimensional regression. *J. Mach. Learn. Res.* **15** 2869–2909. MR3277152
- [27] JAVANMARD, A. and MONTANARI, A. (2014). Hypothesis testing in high-dimensional regression under the Gaussian random design model: Asymptotic theory. *IEEE Trans. Inf. Theory* **60** 6522–6554. MR3265038 <https://doi.org/10.1109/TIT.2014.2343629>
- [28] JAVANMARD, A. and MONTANARI, A. (2018). Debiasing the Lasso: Optimal sample size for Gaussian designs. *Ann. Statist.* **46** 2593–2622. MR3851749 <https://doi.org/10.1214/17-AOS1630>
- [29] LAURENT, B. and MASSART, P. (2000). Adaptive estimation of a quadratic functional by model selection. *Ann. Statist.* **28** 1302–1338. MR1805785 <https://doi.org/10.1214/aos/1015957395>
- [30] LECUÉ, G. and MENDELSON, S. (2018). Regularization and the small-ball method I: Sparse recovery. *Ann. Statist.* **46** 611–641. MR3782379 <https://doi.org/10.1214/17-AOS1562>
- [31] LEI, L., BICKEL, P. J. and EL KAROUI, N. (2018). Asymptotics for high dimensional regression  $M$ -estimates: Fixed design results. *Probab. Theory Related Fields* **172** 983–1079. MR3877551 <https://doi.org/10.1007/s00440-017-0824-7>
- [32] LOUREIRO, B., GERBELOT, C., CUI, H., GOLDT, S., KRZAKALA, F., MEZARD, M. and ZDEBOROVÁ, L. (2021). Learning curves of generic features maps for realistic datasets with a teacher-student model. *Adv. Neural Inf. Process. Syst.* **34** 18137–18151.
- [33] MIOLANE, L. and MONTANARI, A. (2021). The distribution of the Lasso: Uniform control over sparse balls and adaptive parameter tuning. *Ann. Statist.* **49** 2313–2335. MR4319252 <https://doi.org/10.1214/20-aos2038>
- [34] MOHAMED, N. (2020). Scaled minimax optimality in high-dimensional linear regression: A non-convex algorithmic regularization approach. arXiv preprint. Available at [arXiv:2008.12236](https://arxiv.org/abs/2008.12236).
- [35] NICULESCU, C. P. and PERSSON, L.-E. (2006). *Convex Functions and Their Applications. CMS Books in Mathematics/Ouvrages de Mathématiques de la SMC* **23**. Springer, New York. MR2178902 <https://doi.org/10.1007/0-387-31077-0>
- [36] SILVERSTEIN, J. W. (1985). The smallest eigenvalue of a large-dimensional Wishart matrix. *Ann. Probab.* **13** 1364–1368. MR0806232
- [37] STEIN, C. M. (1981). Estimation of the mean of a multivariate normal distribution. *Ann. Statist.* **9** 1135–1151. MR0630098
- [38] SUN, T. and ZHANG, C.-H. (2012). Scaled sparse linear regression. *Biometrika* **99** 879–898. MR2999166 <https://doi.org/10.1093/biomet/ass043>
- [39] SUN, T. and ZHANG, C.-H. (2013). Sparse matrix inversion with scaled lasso. *J. Mach. Learn. Res.* **14** 3385–3418. MR3144466
- [40] SUR, P. and CANDÈS, E. J. (2019). A modern maximum-likelihood theory for high-dimensional logistic regression. *Proc. Natl. Acad. Sci. USA* **116** 14516–14525. MR3984492 <https://doi.org/10.1073/pnas.1810420116>
- [41] TALAGRAND, M. (2011). *Mean Field Models for Spin Glasses. Volume I: Basic Examples. Ergebnisse der Mathematik und Ihrer Grenzgebiete. 3. Folge. A Series of Modern Surveys in Mathematics [Results in Mathematics and Related Areas. 3rd Series. A Series of Modern Surveys in Mathematics]* **54**. Springer, Berlin. MR2731561 <https://doi.org/10.1007/978-3-642-15202-3>
- [42] THRAMPOULIDIS, C., ABBASI, E. and HASSIBI, B. (2015). Lasso with non-linear measurements is equivalent to one with linear measurements. In *Advances in Neural Information Processing Systems* 3420–3428.

- [43] THRAMPOULIDIS, C., ABBASI, E. and HASSIBI, B. (2018). Precise error analysis of regularized  $M$ -estimators in high dimensions. *IEEE Trans. Inf. Theory* **64** 5592–5628. [MR3832326](#) <https://doi.org/10.1109/TIT.2018.2840720>
- [44] TIBSHIRANI, R. J. (2013). The lasso problem and uniqueness. *Electron. J. Stat.* **7** 1456–1490. [MR3066375](#) <https://doi.org/10.1214/13-EJS815>
- [45] VAITER, S., DELEDALLE, C., PEYRÉ, G., FADILI, J. and DOSSAL, C. (2012). The degrees of freedom of the group lasso. arXiv preprint. Available at [arXiv:1205.1481](#).
- [46] VAN DE GEER, S., BÜHLMANN, P., RITOV, Y. and DEZEURE, R. (2014). On asymptotically optimal confidence regions and tests for high-dimensional models. *Ann. Statist.* **42** 1166–1202. [MR3224285](#) <https://doi.org/10.1214/14-AOS1221>
- [47] VERSHYNIN, R. (2018). *High-Dimensional Probability: An Introduction with Applications in Data Science*. Cambridge Series in Statistical and Probabilistic Mathematics **47**. Cambridge Univ. Press, Cambridge. [MR3837109](#) <https://doi.org/10.1017/9781108231596>
- [48] ZHANG, C.-H. (2010). Nearly unbiased variable selection under minimax concave penalty. *Ann. Statist.* **38** 894–942. [MR2604701](#) <https://doi.org/10.1214/09-AOS729>
- [49] ZHANG, C.-H. (2011). Statistical inference for high-dimensional data. *Math. Forsch. Oberwolfach* **48** 28–31.
- [50] ZHANG, C.-H. and HUANG, J. (2008). The sparsity and bias of the LASSO selection in high-dimensional linear regression. *Ann. Statist.* **36** 1567–1594. [MR2435448](#) <https://doi.org/10.1214/07-AOS520>
- [51] ZHANG, C.-H. and ZHANG, S. S. (2014). Confidence intervals for low dimensional parameters in high dimensional linear models. *J. R. Stat. Soc. Ser. B. Stat. Methodol.* **76** 217–242. [MR3153940](#) <https://doi.org/10.1111/rssb.12026>