



Weierstraß-Institut für
Angewandte Analysis und Stochastik



High dimensional logistic regression

Vladimir Spokoiny ,
HSE University and IITP RAS Moscow,
WIAS and HU Berlin

ICOMP 2025

1 Introduction

2 Tools

- Perturbed optimization
- Random matrix theory
- Deviation bounds for quadratic forms and fourth-order tensors
- Ridge penalty and a smooth operator

3 High-dimensional logistic regression

- Deterministic design
- Random design

Observed (Y_i, Ψ_i) with binary outputs (labels) Y_i and the corresponding feature vectors $\Psi_i \in \mathbb{R}^p$, $i = 1, \dots, n$.

Aim: predict (estimate) $f(\psi) = \mathbb{E}(Y \mid \psi)$.

Logistic regression (a linear model for the canonical parameter):

$$f(\psi) = \phi'(\langle \psi, \mathbf{v} \rangle), \quad \phi(t) = \log(1 + e^t), \quad \phi'(t) = \frac{e^t}{1 + e^t}.$$

MLE:

$$\tilde{\mathbf{v}} = \operatorname{argmin}_{\mathbf{v} \in \mathbb{R}^p} L(\mathbf{v}) = \operatorname{argmin}_{\mathbf{v} \in \mathbb{R}^p} \sum_{i=1}^n \{ \phi(\langle \Psi_i, \mathbf{v} \rangle) - Y_i \langle \Psi_i, \mathbf{v} \rangle \}.$$

In the limit of large samples in which p is fixed and $n \rightarrow \infty$, the MLE $\tilde{\mathbf{v}}$ obeys

$$\sqrt{n}(\tilde{\mathbf{v}} - \mathbf{v}) \xrightarrow{w} \mathcal{N}(0, F_{\mathbf{v}}),$$

where $F_{\mathbf{v}}$ is the Fisher information matrix at $\mathbf{v}$

$$F_{\mathbf{v}} \stackrel{\text{def}}{=} \mathbb{E}\{\boldsymbol{\Psi}\boldsymbol{\Psi}^{\top} \phi''(\langle \boldsymbol{\Psi}, \mathbf{v} \rangle)\}, \quad \phi(t) = \frac{e^t}{(1 + e^t)^2}.$$

see e.g.

- Lehmann and Romano. Testing statistical hypotheses. 2006.
- McCullagh and Nelder. Generalized linear models. 1989.
- Van der Vaart. Asymptotic statistics. 2000.

- **[Sur and Candès, 2019]** A modern maximum-likelihood theory for high-dimensional logistic regression,
- **[Candès and Sur, 2020]** The phase transition for the existence of the maximum likelihood estimate in high-dimensional logistic regression
- **[Montanari et al., 2025]** The generalization error of max-margin linear classifiers: Benign overfitting and high dimensional asymptotics in the overparametrized regime
- **[Kuchelmeister and van de Geer, 2024]** Finite Sample Rates for Logistic Regression with Small Noise or Few Samples
- **[Chardon et al., 2024]** Finite-sample performance of the maximum likelihood estimator in logistic regression
- **[Bach, 2024]** High-Dimensional Analysis of Double Descent for Linear Regression with Random Projections

Main issues:

- high dimension $p \gg n$; need of regularization;
- random and ill-posed design $\Psi\Psi^\top$;
- finite-sample and dimension-free accuracy guarantees.

Tools:

- Perturbed optimization;
- random matrix theory;
- deviation bounds for fourth order random tensors;
- regularization for nonlinear inverse problems.

1 Introduction

2 Tools

- Perturbed optimization
- Random matrix theory
- Deviation bounds for quadratic forms and fourth-order tensors
- Ridge penalty and a smooth operator

3 High-dimensional logistic regression

- Deterministic design
- Random design

Let $L(\mathbf{v})$ be a random function (loss, empirical risk, negative log-likelihood). Consider

$$\tilde{\mathbf{v}} = \underset{\mathbf{v}}{\operatorname{argmin}} L(\mathbf{v}); \quad \mathbf{v}^* = \underset{\mathbf{v}}{\operatorname{argmin}} \mathbb{E} L(\mathbf{v});$$

Includes MLE, LSE, LAD, Minimum Contrast, and many others procedures.

Aim: describe $\tilde{\mathbf{v}} - \mathbf{v}^*$ (estimation loss) and $L(\tilde{\mathbf{v}}) - L(\mathbf{v}^*)$ (excess).

Perturbed optimization: $L(\mathbf{v})$ is a perturbation of $f(\mathbf{v}) = \mathbb{E} L(\mathbf{v})$ by a stochastic component

$$\zeta(\mathbf{v}) = L(\mathbf{v}) - \mathbb{E} L(\mathbf{v}).$$

Consider a penalized MLE for a penalty function $\text{pen}_G(\mathbf{v})$

$$\tilde{\mathbf{v}}_G = \underset{\mathbf{v}}{\operatorname{argmin}} L_G(\mathbf{v}) = \underset{\mathbf{v}}{\operatorname{argmin}} \{L(\mathbf{v}) + \text{pen}_G(\mathbf{v})\}.$$

Typical examples: $\text{pen}_G(\mathbf{v}) = \|G\mathbf{v}\|^2/2$ and $\text{pen}_\lambda(\mathbf{v}) = \lambda\|\mathbf{v}\|_1$.

Consider three optimization problems

$$\tilde{\mathbf{v}}_G = \underset{\mathbf{v}}{\operatorname{argmin}} L_G(\mathbf{v}),$$

$$\mathbf{v}_G^* = \underset{\mathbf{v}}{\operatorname{argmin}} \mathbb{E} L_G(\mathbf{v}),$$

$$\mathbf{v}^* = \underset{\mathbf{v}}{\operatorname{argmin}} \mathbb{E} L(\mathbf{v}).$$

The penalized MLE $\tilde{\mathbf{v}}_G$ estimates rather $\mathbf{v}_G^*$ than $\mathbf{v}^*$.

Let $f(\mathbf{v})$ be a **smooth convex** function,

$$\mathbf{v}^* = \operatorname{argmin}_{\mathbf{v}} f(\mathbf{v}), \quad \mathbb{F} = \nabla^2 f(\mathbf{v}^*).$$

Let another function $g(\mathbf{v})$ satisfy for some vector $\mathbf{A}$

$$g(\mathbf{v}) - g(\mathbf{v}^*) = \langle \mathbf{v} - \mathbf{v}^*, \mathbf{A} \rangle + f(\mathbf{v}) - f(\mathbf{v}^*).$$

Define

$$\mathbf{v}^\circ \stackrel{\text{def}}{=} \operatorname{argmin}_{\mathbf{v}} g(\mathbf{v}), \quad g(\mathbf{v}^\circ) = \min_{\mathbf{v}} g(\mathbf{v}).$$

Aim: evaluate the quantities $\mathbf{v}^\circ - \mathbf{v}^*$ and $g(\mathbf{v}^\circ) - g(\mathbf{v}^*)$.

Consider $\mathbf{v}^* \stackrel{\text{def}}{=} \operatorname{argmin} f(\mathbf{v})$, $\mathbf{v}^\circ \stackrel{\text{def}}{=} \operatorname{argmin}\{f(\mathbf{v}) + \langle \mathbf{v}, \mathbf{A} \rangle\}$.

$(\mathcal{T}_3^*)$ $f(\mathbf{v})$ is strongly convex, $D^2 \leq \mathbb{F} \stackrel{\text{def}}{=}} \nabla^2 f(\mathbf{v}^*)$, and

$$\sup_{\mathbf{u}: \|D\mathbf{u}\| \leq r} \sup_{\mathbf{z} \in \mathbb{R}^p} \frac{|\langle \nabla^3 f(\mathbf{v}^* + \mathbf{u}), \mathbf{z}^{\otimes k} \rangle|}{\|D\mathbf{z}\|^3} \leq \tau_3.$$

Proposition ([Spokoiny, 2025b, Spokoiny, 2025a])

Assume $(\mathcal{T}_3^*)$ at $\mathbf{v}^*$ with D^2 , r , and τ_3 s.t.

$$r \geq \frac{3}{2} \|D \mathbb{F}^{-1} \mathbf{A}\|, \quad \tau_3 \|D \mathbb{F}^{-1} \mathbf{A}\| < \frac{4}{9}.$$

Then $\|D(\mathbf{v}^\circ - \mathbf{v}^*)\| \leq (3/2) \|D \mathbb{F}^{-1} \mathbf{A}\|$ and moreover,

$$\|D^{-1} \mathbb{F}(\mathbf{v}^\circ - \mathbf{v}^* - \mathbb{F}^{-1} \mathbf{A})\| \leq \frac{3}{4} \tau_3 \|D \mathbb{F}^{-1} \mathbf{A}\|^2.$$

Model:

$$H = \mathbb{E} \Psi \Psi^\top, \quad H_n = \frac{1}{n} \sum_{i=1}^n \Psi_i \Psi_i^\top.$$

Whenever Ψ is subgaussian in all directions, and

$$n \geq K(d + \log(1/\delta)),$$

with probability at least $1 - \delta$, it holds

$$\frac{1}{2} \|\mathbf{u}\|_H^2 \leq \|\mathbf{u}\|_{H_n}^2 \leq 2 \|\mathbf{u}\|_H^2, \quad \forall \mathbf{u} \in \mathbb{R}^p.$$

[Koltchinskii and Lounici, 2017], [Vershynin, 2018], [Tropp, 2015].

Let $\Psi, \Psi_1, \Psi_2, \dots$ be i.i.d. in $\mathbb{R}^p$ and satisfy

$$n \log \mathbb{E} \exp \langle \Psi, \mathbf{u} \rangle \leq \frac{C_\Psi}{2} \|\mathbb{D} \mathbf{u}\|^2, \quad \mathbf{u} \in \mathbb{R}^p,$$

$$\mathrm{tr}(\mathbb{B}_4) \ll n^{1/2},$$

where $\mathbb{B}_4 \stackrel{\text{def}}{=} n \mathbb{D}^{-1} \mathbb{E} \{ |\phi^{(4)}(\langle \Psi, \mathbf{v}_\lambda^* \rangle)|^{1/2} \Psi \Psi^\top \} \mathbb{D}^{-1}$.

Then on a set Ω_n with $\mathbb{P}(\Omega_n) \geq 1 - 1/n$

$$\sup_{\|\mathbb{D} \mathbf{z}\|=1} \left| \sum_{i=1}^n \phi^{(4)}(\langle \Psi_i, \mathbf{v}_\lambda^* \rangle) \langle \Psi_i, \mathbf{z} \rangle^4 \right| \leq \frac{C_4 + \varepsilon_4}{n},$$

where

$$\varepsilon_4 \lesssim \sqrt{\frac{\mathrm{tr}(\mathbb{B}_4)}{n}} + \frac{\mathrm{tr}^2(\mathbb{B}_4)}{n} + \sqrt{\frac{\log n}{n}} + \frac{\log^2 n}{n}.$$

[Zhivotovskiy, 2024] and [Al-Ghattas et al., 2025].

Penalized MLE:

$$\tilde{\mathbf{v}}_G = \operatorname{argmin}_{\mathbf{v}} L_G(\mathbf{v}) = \operatorname{argmin}_{\mathbf{v}} \{ L(\mathbf{v}) + \text{pen}_G(\mathbf{v}) \}.$$

An important example of penalty choice is a ridge penalty $G^2 = g^2 \mathbb{I}_p$. It is **basis and coordinate-free** and enforces the “**benign overfitting**” phenomenon in high-dimensional regression; see [Bartlett et al., 2020], [Cheng and Montanari, 2022], [Noskov et al., 2025] and references therein.

A phase transition effect: if the operator $\mathbb{F} = \nabla^2 L(\mathbf{v}^*)$ is **smoother** than the signal then one can achieve **optimal accuracy**.

Does not apply if the operator $\mathbb{F}$ **is not smoother** than the signal.

1 Introduction

2 Tools

- Perturbed optimization
- Random matrix theory
- Deviation bounds for quadratic forms and fourth-order tensors
- Ridge penalty and a smooth operator

3 High-dimensional logistic regression

- Deterministic design
- Random design

Observed $(Y_i, \boldsymbol{\Psi}_i)$ with binary labels Y_i and features $\boldsymbol{\Psi}_i \in \mathbb{R}^p$, $i = 1, \dots, n$.

Logistic regression: with $\phi(t) = \log(1 + e^t)$,

$$\mathbb{E}(Y \mid \boldsymbol{\psi}) = f(\boldsymbol{\psi}) = \phi'(\langle \boldsymbol{\psi}, \mathbf{v} \rangle).$$

MLE:

$$\tilde{\mathbf{v}} = \operatorname{argmin}_{\mathbf{v} \in \mathbb{R}^p} L(\mathbf{v}) = \operatorname{argmin}_{\mathbf{v} \in \mathbb{R}^p} \sum_{i=1}^n \{ \phi(\langle \boldsymbol{\Psi}_i, \mathbf{v} \rangle) - Y_i \langle \boldsymbol{\Psi}_i, \mathbf{v} \rangle \}.$$

Penalized MLE: for the quadratic penalty $\|G\mathbf{v}\|^2/2$:

$$\tilde{\mathbf{v}}_G = \operatorname{argmin}_{\mathbf{v} \in \mathbb{R}^p} L_G(\mathbf{v}) = \operatorname{argmin}_{\mathbf{v} \in \mathbb{R}^p} L(\mathbf{v}) + \frac{1}{2} \|G\mathbf{v}\|^2.$$

A penalized MLE $\tilde{\mathbf{v}}_G$:

$$\tilde{\mathbf{v}}_G = \operatorname{argmin}_{\mathbf{v} \in \mathbb{R}^p} L_G(\mathbf{v}).$$

The truth $\mathbf{v}^*$:

$$\mathbf{v}^* = \operatorname{argmin}_{\mathbf{v} \in \mathbb{R}^p} \mathbb{E} L(\mathbf{v}),$$

The penalized truth:

$$\mathbf{v}_G^* = \operatorname{argmin}_{\mathbf{v} \in \mathbb{R}^p} \mathbb{E} L_G(\mathbf{v}).$$

The study is based on the decomposition

$$\tilde{\mathbf{v}}_G - \mathbf{v}^* = \underbrace{\tilde{\mathbf{v}}_G - \mathbf{v}_G^*}_{\text{stochastic term}} + \underbrace{\mathbf{v}_G^* - \mathbf{v}^*}_{\text{bias}}.$$

The Fisher matrix $\mathbb{F}(\boldsymbol{v})$ at $\boldsymbol{v}$ is given by

$$\mathbb{F}(\boldsymbol{v}) = \nabla^2 \mathbb{E} L(\boldsymbol{v}) = \sum_{i=1}^n \phi''(\langle \boldsymbol{\Psi}_i, \boldsymbol{v} \rangle) \boldsymbol{\Psi}_i \boldsymbol{\Psi}_i^\top.$$

We also write

$$\mathbb{F}_G(\boldsymbol{v}) = \mathbb{F}(\boldsymbol{v}) + G^2, \quad \mathbb{F}_G = \mathbb{F}_G(\boldsymbol{v}_G^*) = \mathbb{F}(\boldsymbol{v}_G^*) + G^2.$$

Alternatively one can define $\mathbb{F}_G = \mathbb{F}_G(\boldsymbol{v}^*)$.

Our conditions rely on a metric tensor D which defines a local vicinity of $\boldsymbol{v}_G^*$. We assume $\mathbb{F} \leq D^2 \leq \mathbb{F}_G$ and $D^2 = \mathbb{F}_G$ is a default choice.

For checking the conditions, we need some regularity of the design

$\Psi_1, \dots, \Psi_n$.

(Ψ) With a metric tensor D satisfying

$\mathbb{E}(\mathbf{v}_G^*) \leq D^2 \leq \mathbb{E}_G(\mathbf{v}_G^*)$, for some $\delta, \delta_0 > 0$

$$\sup_{\mathbf{z} \in \mathbb{R}^p} \frac{1}{\|D\mathbf{z}\|^4} \left| \sum_{i=1}^n \langle \Psi_i, \mathbf{z} \rangle^4 \phi^{(4)}(\langle \Psi_i, \mathbf{v}_G^* \rangle) \right| \leq \delta^2,$$

$$\max_{i \leq n} \|D^{-1} \Psi_i\| \leq \delta_0.$$

Even for a Gaussian regression, dimension-free guarantees are quite involved; e.g. [Bartlett et al., 2020], [Cheng and Montanari, 2022], [Noskov et al., 2025] and references therein.

As in these papers, consider a ridge (Tikhonov) penalization $\lambda \|\mathbf{v}\|^2/2$. Let observations (labels) Y_i follow logistic regression with parameter $\mathbf{v}^*$, that is, the pairs $(Y_i, \boldsymbol{\Psi}_i)$ are i.i.d. with

$$\mathbb{P}(Y_i = 1 \mid \boldsymbol{\Psi}_i) = \sigma(\langle \boldsymbol{\Psi}_i, \mathbf{v}^* \rangle), \quad \sigma(t) = \phi'(t) = e^t / (1 + e^t).$$

We assume $\boldsymbol{\Psi}, \boldsymbol{\Psi}_1, \boldsymbol{\Psi}_2, \dots \in \mathbb{R}^p$ to be i.i.d. zero mean random vectors drawn from some distribution, e.g. $\boldsymbol{\Psi} \sim \mathcal{N}(0, \Sigma)$.

The central object of our study is the empirical Fisher information

$$\mathbb{F}_\lambda(\mathbf{v}) = \sum_{i=1}^n \phi''(\langle \boldsymbol{\Psi}_i, \mathbf{v} \rangle) \boldsymbol{\Psi}_i \boldsymbol{\Psi}_i^\top + \lambda \mathbb{I}_p,$$

which is now random.

If the dimension p becomes larger than the sample size n , The properties of this random matrix become crucial; see e.g.

[Bartlett et al., 2020], [Cheng and Montanari, 2022],
[Sur and Candès, 2019], [Kuchelmeister and van de Geer, 2024],
[Bach, 2024].

In our results, we fix a metric tensor $\mathbb{D}$. A canonical choice is

$$\mathbb{D}^2 = \mathbb{D}_\lambda^2 \stackrel{\text{def}}{=} n \mathbb{E} \{ \Psi \Psi^\top \phi''(\langle \Psi, \mathbf{v}_\lambda^* \rangle) \} + \lambda \mathbb{I}_p.$$

We impose a sub-Gaussian condition on Ψ : for some fixed constant C_Ψ

$$n \log \exp \langle \Psi, \mathbf{u} \rangle \leq \frac{C_\Psi}{2} \|\mathbb{D}_\lambda \mathbf{u}\|^2, \quad \mathbf{u} \in \mathbb{R}^p. \quad (1)$$

The estimator $\tilde{\mathbf{v}}_\lambda$ solves the empirical problem

$$\tilde{\mathbf{v}}_\lambda = \operatorname{argmin}_{\mathbf{v} \in \mathbb{R}^p} L_\lambda(\mathbf{v}) = \operatorname{argmin}_{\mathbf{v} \in \mathbb{R}^p} \left[- \sum_{i=1}^n \left\{ Y_i \langle \boldsymbol{\Psi}_i, \mathbf{v} \rangle - \phi(\langle \boldsymbol{\Psi}_i, \mathbf{v} \rangle) \right\} + \frac{\lambda}{2} \|\mathbf{v}\|^2 \right]$$

The whole analysis is based on the decomposition

$$\tilde{\mathbf{v}}_\lambda - \mathbf{v}^* = (\tilde{\mathbf{v}}_\lambda - \mathbf{v}_\lambda^*) + (\mathbf{v}_\lambda^* - \mathbf{v}^*) \quad (3)$$

for the penalized truth

$$\mathbf{v}_\lambda^* = \operatorname{argmin}_{\mathbf{v} \in \mathbb{R}^p} \left[-\mathbb{E} \sum_{i=1}^n \left\{ \mathbb{E}(Y_i \mid \boldsymbol{\Psi}_i) \langle \boldsymbol{\Psi}_i, \mathbf{v} \rangle - \phi(\langle \boldsymbol{\Psi}_i, \mathbf{v} \rangle) \right\} + \frac{\lambda}{2} \|\mathbf{v}\|^2 \right].$$

The first term $\tilde{\mathbf{v}}_\lambda - \mathbf{v}_\lambda^*$ in (3) describes the stochastic error caused by using the noisy data Y_i , $i \leq n$. The bias term $\mathbf{v}_\lambda^* - \mathbf{v}^*$ is due to the ridge penalty $\lambda \|\mathbf{v}\|^2/2$ in (2).

Later, we evaluate each difference using general results on perturbed optimization. Given λ , define

$$\bar{F}_\lambda(\mathbf{v}) \stackrel{\text{def}}{=} n\mathbb{E}\{\phi''(\langle \Psi, \mathbf{v} \rangle) \Psi \Psi^\top\} + \lambda \mathbb{I}_p.$$

Also, set $\bar{F}_\lambda \stackrel{\text{def}}{=} \bar{F}_\lambda(\mathbf{v}_\lambda^*)$ and define

$$\mathbf{S} \stackrel{\text{def}}{=} - \sum_{i=1}^n (Y_i - \mathbb{E}Y_i) \Psi_i.$$

We show that

$$\tilde{\mathbf{v}}_\lambda - \mathbf{v}_\lambda^* \approx -\bar{F}_\lambda^{-1} \mathbf{S}, \quad \mathbf{v}_\lambda^* - \mathbf{v}^* \approx -\lambda \bar{F}_\lambda^{-1} \mathbf{v}^*,$$

and provide some finite sample accuracy guarantees.

With $\mathbb{D}_\lambda^2 \stackrel{\text{def}}{=} \bar{\mathbb{F}}_\lambda$, define

$$\mathbb{B}_2 \stackrel{\text{def}}{=} n \mathbb{D}_\lambda^{-1} \mathbb{E} \{ \phi''(\langle \Psi, \mathbf{v}_\lambda^* \rangle) \Psi \Psi^\top \} \mathbb{D}_\lambda^{-1}.$$

Proposition

Let the design distribution of the Ψ_i 's satisfy (1). Assume $\text{tr } \mathbb{B}_2 \ll n$. Then on a random set Ω_n with $\mathbb{P}(\Omega_n) \geq 1 - 1/n$, it holds

$$\bar{\mathbb{F}}_\lambda(1 - \varepsilon_\lambda) \leq \mathbb{F}_\lambda \leq \bar{\mathbb{F}}_\lambda(1 + \varepsilon_\lambda),$$

where

$$\varepsilon_\lambda \lesssim \sqrt{\frac{\mathbf{c}_\Psi \text{tr } \mathbb{B}_2}{n}} + \sqrt{\frac{\log(n+p)}{n}} \ll 1.$$

Theorem

With $\xi_\lambda \stackrel{\text{def}}{=} \|\mathbb{D}_\lambda^{-1} \mathbf{S}\|$, $b_\lambda \stackrel{\text{def}}{=} \|\mathbb{D}_\lambda(\mathbf{v}_\lambda^* - \mathbf{v}^*)\|$, and $\tau_3 = c_3 n^{-1/2}$ for a constant c_3 depending on C_Ψ only, it holds on Ω_n

$$\|\xi_\lambda\| \leq \sqrt{\text{tr } \mathbb{B}_2} + 2\sqrt{\log n}.$$

Furthermore, $b_\lambda \leq \frac{3}{2} \|\lambda \bar{\mathbb{F}}_\lambda^{-1/2} \mathbf{v}^*\|$, and

$$\begin{aligned} & \left\| \mathbb{D}_\lambda(\tilde{\mathbf{v}}_\lambda - \mathbf{v}^*) + \mathbb{D}_\lambda^{-1} \mathbf{S} + \lambda \bar{\mathbb{F}}_\lambda^{-1} \mathbf{v}^* \right\| \\ & \leq \frac{\varepsilon_\lambda}{1 - \varepsilon_\lambda} (\xi_\lambda + b_\lambda) + \frac{\tau_3}{(1 - \varepsilon_\lambda)^3} \left(\xi_\lambda^2 + \frac{3}{4} \|\lambda \bar{\mathbb{F}}_\lambda^{-1/2} \mathbf{v}^*\|^2 \right). \end{aligned}$$

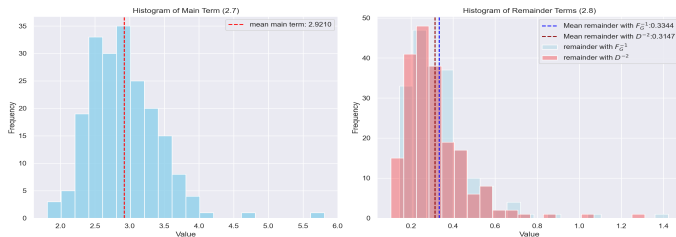


Abbildung: Distribution of the leading term and the remainder for $n = 100$.

Left: the norm of the leading term $\|\mathbb{F}_G^{-1} \nabla \zeta\|_\infty$ and $\|D^{-2} \nabla \zeta\|_\infty$.

Right: the errors $\|\tilde{\mathbf{v}}_G - \mathbf{v}^* + \mathbb{F}_G^{-1} \nabla \zeta\|_\infty$ and $\|\tilde{\mathbf{v}}_G - \mathbf{v}^* + D^{-2} \nabla \zeta\|_\infty$ for $n = 100$.

The magnitude of the remainder is much smaller than the leading term, thus indicating a very good quality of the provided expansions.

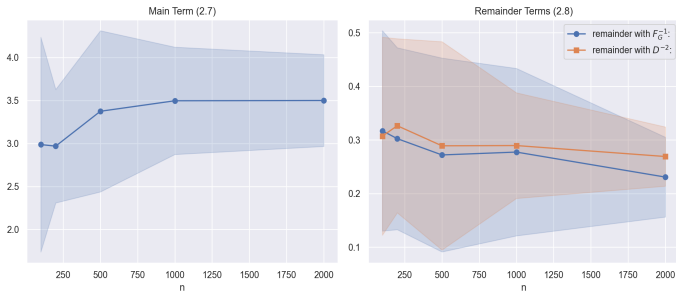


Abbildung: Comparison of the leading term and the remainder for different n .

- (Parametric) lower bounds.
- LASSO-type procedures. **Issue** $\|v\|_1$ is non-smooth.
Idea: replace by a smooth penalty ensuring $p \asymp \#(\text{active set})$.
- Other models including
 - dimension reduction, manifold learning
 - density estimation
 - error-in-operator models
 - matrix completion
 - inference for stochastic processes
 - statistical nonlinear inverse problems, PDEs
 - ...
- adaptation, parameter/penalty choice
- higher-order expansions, estimation of smooth functionals.

-  Al-Ghattas, O., Chen, J., and Sanz-Alonso, D. (2025).
Sharp concentration of simple random tensors.
arxiv.org/abs/2502.16916.
-  Bach, F. (2024).
High-dimensional analysis of double descent for linear regression with random projections.
SIAM Journal on Mathematics of Data Science, 6(1):26–50.
-  Bartlett, P. L., Long, P. M., Lugosi, G., and Tsigler, A. (2020).
Benign overfitting in linear regression.
Proceedings of the National Academy of Sciences, 117(48):30063–30070.
-  Candès, E. J. and Sur, P. (2020).
The phase transition for the existence of the maximum likelihood estimate in high-dimensional logistic regression.
The Annals of Statistics, 48(1):27–42.
-  Chardon, H., Lerasle, M., and Mourtada, J. (2024).
Finite-sample performance of the maximum likelihood estimator in logistic regression.
-  Cheng, C. and Montanari, A. (2022).
Dimension free ridge regression.
<https://arxiv.org/abs/2210.08571>.
-  Koltchinskii, V. and Lounici, K. (2017).
Concentration inequalities and moment bounds for sample covariance operators.
Bernoulli, 23(1):110 – 133.



Kuchelmeister, F. and van de Geer, S. (2024).

Finite sample rates for logistic regression with small noise or few samples.

Sankhya A.



Montanari, A., Ruan, F., Sohn, Y., and Yan, J. (2025).

The generalization error of max-margin linear classifiers: Benign overfitting and high dimensional asymptotics in the overparametrized regime.

The Annals of Statistics, 53(2):822 – 853.



Noskov, F., Puchkin, N., and Spokoiny, V. (2025).

Dimension-free bounds in high-dimensional linear regression via error-in-operator approach.

arxiv.org/abs/2502.15437.



Spokoiny, V. (2025a).

Marginal minimization and sup-norm expansions in perturbed optimization.

arxiv.org/2505.02562.



Spokoiny, V. (2025b).

Sharp bounds in perturbed smooth optimization.

arxiv.org/2504.11834.



Sur, P. and Candès, E. J. (2019).

A modern maximum-likelihood theory for high-dimensional logistic regression.

Proceedings of the National Academy of Sciences, 116(29):14516–14525.



Tropp, J. A. (2015).

An introduction to matrix concentration inequalities.

Foundations and Trends in Machine Learning, 8(1-2):1–230.



Vershynin, R. (2018).

High-Dimensional Probability: An Introduction with Applications in Data Science.

Number 47 in Cambridge Series in Statistical and Probabilistic Mathematics. Cambridge University Press.



Zhivotovskiy, N. (2024).

Dimension-free bounds for sums of independent matrices and simple tensors via the variational principle.

Electronic Journal of Probability, 29(none):1 – 28.