

# High Dimensional Statistics

Yanbo Tang

Imperial College London

March. 2025

# Content of this week

- LASSO sparse set recovery (PDW)
- Matrix concentration  $\|A\|_{\text{op}}$
- Spectral clustering on Random graphs
- Matrix Bernstein  $\|\sum A_i\|_{\text{op}}$
- Covariance estimation  $\frac{\sum_{i=1}^n x_i x_i^T}{n}$

## Theorem

Consider an  $S$ -sparse linear regression model for which the design matrix satisfies Assumptions 2 and 3. Then for any choice of regularization parameter such that

$$\lambda_n \geq \frac{2}{1 - \alpha} \left\| X_S^T \Pi_{S^\perp}(\mathbf{X}) \frac{\epsilon}{n} \right\|_\infty, \quad (1)$$

the Lasso program has the following properties:

- (a) **Uniqueness:** There is a unique optimal solution  $\hat{\theta}$ .
- (b) **No false inclusion:** This solution has its support set  $\hat{S}$  contained within the true support set  $S$ .
- (c)  **$\ell_\infty$ -bounds:** The error  $\hat{\theta} - \theta^*$  satisfies

$$\|\hat{\theta}_S - \theta_S^*\|_\infty \leq \left\| \left( \frac{X_S^T X_S}{n} \right)^{-1} X_S^T \frac{\epsilon}{n} \right\|_\infty + \left\| \left( \frac{X_S^T X_S}{n} \right)^{-1} \right\|_\infty \lambda_n, \quad (2)$$

where  $\|A\|_\infty = \max_{i=1, \dots, S} \sum_j |A_{i,j}|$  is the matrix  $\ell_\infty$ -norm.

- (d) **No false exclusion:** The Lasso includes all indices  $i \in S$  such that  $|\theta_i^*| > B(\lambda_n; \mathbf{X})$ , and hence is variable selection consistent if  $\min_{i \in S} |\theta_i^*| > B(\lambda_n; \mathbf{X})$ .

## Corollary

For a  $S$ -sparse linear model based on a noise vector  $\epsilon$  with zero-mean i.i.d.  $\sigma$ -sub-Gaussian entries, and a deterministic design matrix  $X$  that satisfies Assumptions 2 and 3, as well as the  $C$ -column normalization condition  $\max_{j=1,\dots,d} \|X_j\|_2/\sqrt{n} \leq C$ . Suppose that we solve the Lasso program with regularization parameter

$$\lambda_n = \frac{2C\sigma}{1-\alpha} \left\{ \sqrt{\frac{2\log(d-k)}{n}} + \delta \right\}$$

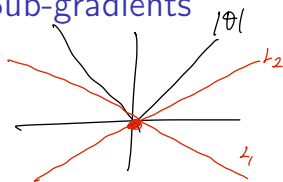
for some  $\delta > 0$ . Then the optimal solution  $\hat{\theta}$  is unique with its support contained within  $S$ , and satisfies the  $\ell_\infty$ -error bound

$$\|\hat{\theta}_S - \theta_S^*\|_\infty \leq \frac{\sigma}{\sqrt{c_{\min}}} \left( \sqrt{\frac{2\log s}{n}} + \delta \right) + \left\| \left( \frac{X_S^T X_S}{n} \right)^{-1} \right\|_\infty \lambda_n, \quad (3)$$

all with probability at least  $1 - 4e^{-n\delta^2/2}$ .

# Proof

# Sub-gradients



Convex function  $f: \mathbb{R}^d \rightarrow \mathbb{R}$   
 we say  $z \in \mathbb{R}^d$  is a sub-gradient  
 of  $f$  at  $(\theta)$

$$f(\theta + \Delta) \geq f(\theta) + \langle z, \Delta \rangle \quad \forall \Delta \in \mathbb{R}^d$$

then  $z \in \partial f(\theta)$

{ we say a pair  $(\hat{\theta}, \hat{z})$  is primal-dual optimal

if 1)  $\hat{\theta}$  is a minimizer

2)  $\hat{z} \in \partial f(\hat{\theta})$

and

$$\frac{1}{n} X^T (X\hat{\theta} - y) + \lambda \hat{z} = 0$$

$$z_i \in [-1, 1]$$

# Primal Dual Witness

## Definition

**Primal–dual witness (PDW) construction:**

- 1 Set  $\hat{\theta}_{S^c} = 0$ .
- 2 Determine  $(\hat{\theta}_S, \hat{z}_S) \in \mathbb{R}^s \times \mathbb{R}^s$  by solving the *oracle subproblem*

$$\hat{\theta}_S \in \arg \min_{\theta_S \in \mathbb{R}^s} \left\{ \underbrace{\frac{1}{2n} \|y - X_S \theta_S\|_2^2 + \lambda_n \|\theta_S\|_1}_{=: f(\theta_S)} \right\}, \quad (4)$$

and then choosing  $\hat{z}_S \in \partial \|\hat{\theta}_S\|_1$  such that  $\nabla f(\theta_S)|_{\theta_S = \hat{\theta}_S} + \lambda_n \hat{z}_S = 0$ .

- 3 Solve for  $\hat{z}_{S^c} \in \mathbb{R}^{d-s}$  via the zero-subgradient equation, and check whether or not the *strict dual feasibility* condition  $\|\hat{z}_{S^c}\|_\infty < 1$  holds.

# Witness for LASSO $\lambda_{\min}\left(\frac{X_S^T X_S}{n}\right) > C_{\min} > 0$

## Lemma

If the lower eigenvalue condition holds, then success of the PDW construction implies that the vector  $(\hat{\theta}_S, 0) \in \mathbb{R}^d$  is the unique optimal solution of the Lasso.

Sketch if  $\tilde{\theta}$   $\|\tilde{\theta}\|_1 \leq \langle \hat{z}, \tilde{\theta} \rangle \leq \|\hat{z}\|_\infty \|\tilde{\theta}\|_1 = \|\tilde{\theta}\|_1$

$\hat{z}_i = -1, 1$   
if  $i \in S$

$\langle \hat{z}, \tilde{\theta} \rangle = \|\tilde{\theta}\|_1$   
 $\langle \hat{z}_S, \tilde{\theta}_S \rangle = \|\tilde{\theta}_S\|_1$   
 $\Rightarrow \tilde{\theta}_{S^c} = 0$



# Proof

# Proof of main theorem

$$\hat{\beta}_{SC} = \frac{1}{\lambda_n n} X_{SC}^T X_S (\hat{\theta}_S - \theta_S^0) + X_{SC}^T \left( \frac{\varepsilon}{\lambda_n n} \right) \quad (1)$$

$$\hat{\theta}_S - \theta_S^0 = (X_S^T X_S)^{-1} X_S^T \varepsilon - \lambda_n n (X_S^T X_S)^{-1} \hat{\beta}_S \quad (2)$$

$$\hat{\beta}_{SC} = \underbrace{X_S^T X_S (X_S^T X_S)^{-1}}_u \hat{\beta}_S + \underbrace{X_{SC}^T (I - X_S (X_S^T X_S)^{-1} X_S^T)}_{V_{SC}} \frac{\varepsilon}{\lambda_n n}$$

mutual incoherence gives  $\|V_{SC}\|_\infty = \alpha$

$$\|\hat{\beta}_S\|_\infty \leq 1$$

$$\|\hat{\beta}_{SC}\|_\infty \leq \|u\|_\infty + \|V_{SC}\|_\infty$$

$h^2$  penalty

$$\lambda_n \|\theta\|_2 \quad \|\theta\|_2 < \delta$$

$$\lambda_n \|f\|_{\mathcal{H}} \quad \|f\|_{\mathcal{H}} < \delta$$

# Matrix Concentration

## Some review

The Courant Fisher min-max theorem which states that:

$$\lambda_i(A) = \max_{\dim(E)=i} \min_{x \in S(E)} x^\top A x$$

where the maximum is taken over all  $i$ -dimensional subspaces  $E$  of  $\mathbb{R}^n$ .

## Some review

The Courant Fisher min-max theorem which states that:

$$\lambda_i(A) = \max_{\dim(E)=i} \min_{x \in S(E)} x^\top A x$$

where the maximum is taken over all  $i$ -dimensional subspaces  $E$  of  $\mathbb{R}^n$ . Using this we can characterize the operator norm or the maximum singular value as follows:

$$\|A\|_{op} := \max_{x \in \mathbb{R}^n \setminus \{0\}} \frac{\|Ax\|_2}{\|x\|_2} = \max_{x \in S^{n-1}} \sqrt{x^\top A x}.$$

Equivalently, the operator norm of  $A$  can be computed by maximizing the quadratic form  $\langle Ax, y \rangle$  over all unit vectors  $x, y$ :

$$\|A\| = \max_{x \in S^{n-1}, y \in S^{m-1}} \langle Ax, y \rangle.$$

# Controlling the operator norm

## Lemma

Let  $A$  be an  $m \times n$  matrix and  $\varepsilon \in [0, 1)$ . Then, for any  $\varepsilon$ -net  $\mathcal{N}$  of the sphere  $S^{n-1}$ , we have

$$\sup_{x \in \mathcal{N}} \|Ax\|_2 \leq \|A\| \leq \frac{1}{1 - \varepsilon} \sup_{x \in \mathcal{N}} \|Ax\|_2.$$

Proof: fix  $x$  for which  $\|A\| = \|Ax\|_2$

choose  $x_0 \in \mathcal{N}$  where  $\|x - x_0\|_2 \leq \varepsilon$

$$\|Ax - Ax_0\|_2 \leq \|A\| \|x - x_0\|_2 \leq \varepsilon \|A\|$$

$\leq \varepsilon$

$$\|Ax_0\|_2 \geq \|Ax\|_2 - \|Ax - Ax_0\|_2 \geq \|A\| - \varepsilon \|A\| \geq (1 - \varepsilon) \|A\|$$

$$\|A\| \leq \frac{\|Ax_0\|_2}{1 - \varepsilon} \leq \sup_{x \in \mathcal{N}} \frac{\|Ax\|_2}{1 - \varepsilon}$$

# Proof

# The actual lemma

## Lemma

Let  $A$  be an  $m \times n$  matrix and  $\varepsilon \in [0, 1)$ . Then, for any  $\varepsilon$ -net  $\mathcal{N}$  of the sphere  $S^{n-1}$  and any  $\varepsilon$ -net  $\mathcal{M}$  of the sphere  $S^{m-1}$ , we have

$$\sup_{x \in \mathcal{N}} \sup_{y \in \mathcal{M}} y^\top A x \leq \|A\|_{\infty, 2} \leq \frac{1}{1 - 2\varepsilon} \sup_{x \in \mathcal{N}} \sup_{y \in \mathcal{M}} y^\top A x.$$



# Random sub-Gaussian matrices

## Theorem

Let  $A$  be an  $m \times n$  random matrix whose entries  $A_{ij}$  are independent, mean zero,  $\sigma$ -sub-gaussian random variables. Then, for any  $t > 0$  we have

$$\|A\| \leq C\sigma (\sqrt{m} + \sqrt{n} + t)$$

with probability at least  $1 - 2 \exp(-t^2)$  for some constant  $C > 0$ .

$$\|D\| \leq \|D\|_F = O(\sqrt{mn})$$

## Step 1: Approximation.

Choose  $\varepsilon = 1/4$  then have  $\varepsilon$ -net  $N(x)$

$\varepsilon$ -net  $M(y)$

$$|N| \leq q^n$$

$$|M| \leq q^m$$

$$\left(1 + \frac{2}{\varepsilon}\right)^d$$

$$\|A\| \leq 2 \max_{x \in N, y \in M} \langle Ax, y \rangle$$

## Step 2: Concentration.

fix  $x \in N$  and  $y \in M$

$$\left( \sum_{i=1}^n \sum_{j=1}^m A_{ij} x_i y_j \right)$$

$\langle Ax, y \rangle = \sum_{i=1}^n \sum_{j=1}^m A_{ij} x_i y_j$  is a sum of independent sub-Gaussians.

$$\| \sigma^2 \| \leq C \sum_{i=1}^n \sum_{j=1}^m \sigma^2 x_i^2 y_j^2 \leq C \sigma^2$$

$$P(\langle Ax, y \rangle \geq u) \leq \exp\left(-\frac{u^2}{C \sigma^2}\right)$$

### Step 3: Union bound.

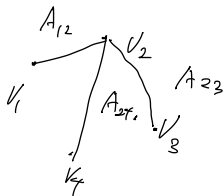
$$\text{if } \max_{x \in N, y \in M} \langle Ax, y \rangle \geq u$$

$$\begin{aligned} P(\max_{x \in N, y \in M} \langle Ax, y \rangle \geq u) &\leq \sum_{x \in N, y \in M} P(\langle Ax, y \rangle \geq u) \\ &\leq qnm \cdot 2 \exp\left(-\frac{u^2}{C\sigma^2}\right) \end{aligned}$$

$$u = C\sigma(\sqrt{n} + \sqrt{m} + t)$$

$$\begin{aligned} \text{then solve to get} \\ \leq 2 \exp(-t^2) \end{aligned}$$

# Application stochastic block model



Erdős-Rényi Random graph  $G(n, p)$   
 $A_{ij} \sim \text{Bern}(p)$

---

$G(n, p, q)$   $E[A] = \begin{bmatrix} p & p & q & q \\ p & p & q & q \\ q & q & p & p \\ q & q & p & p \end{bmatrix}$   $p > q$   
 $n/2$  per group  
 2 groups

$A = E[A] + R$   $\|E[A]\| \approx n, \|R\| \leq C\sqrt{n}$

$E[A]$

$\lambda_1 = \frac{(p+q)}{2} \cdot n$

$\lambda_2 = \frac{(p-q)}{2} \cdot n$

$\lambda_3 = 0$

$\vdots$

$u_1 = \begin{bmatrix} 1 \\ 1 \\ 1 \\ 1 \end{bmatrix}, u_2 = \begin{bmatrix} 1 \\ 1 \\ -1 \\ -1 \end{bmatrix}$

$u_2$

# Illustration

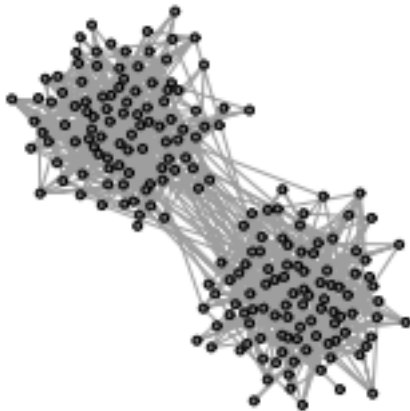


Figure 1: Taken from High Dimensional Probability by Roman Vershynin, with  $n = 200$ ,  $p = 1/20$  and  $q = 1/200$ .

## Theorem

Let  $S$  and  $T$  be symmetric matrices with the same dimensions. Fix  $i$  and assume that the  $i$ -th largest eigenvalue of  $S$  is well separated from the rest of the spectrum:

$$\min_{j:j \neq i} |\lambda_i(S) - \lambda_j(S)| = \delta > 0.$$

Then the angle between the eigenvectors of  $S$  and  $T$  corresponding to the  $i$ -th largest eigenvalues (as a number between 0 and  $\pi/2$ ) satisfies

$$\sin \angle(v_i(S), v_i(T)) \leq \frac{2\|S - T\|}{\delta}.$$

The conclusion of the Davis-Kahan theorem implies that the unit eigenvectors  $v_i(S)$  and  $v_i(T)$  are close to each other up to a sign, namely

$$\exists \theta \in \{-1, 1\} : \|v_i(S) - \theta v_i(T)\|_2 \leq \frac{2^{3/2}\|S - T\|}{\delta}.$$

Cont.

$$f = \max\left(\frac{p-q}{2}, q\right) \cdot n =: \mu \cdot n$$

$$\exists \theta \in \{-1, 1\} : \|V_2(E(A)) - \theta V_2(A)\|_2 \leq \frac{C\sqrt{n}}{\mu \cdot n} = \frac{C}{\mu\sqrt{n}}$$
$$(1 - \exp(-n))$$

$$\|V_2(E(A)) - \theta V_2(A)\|_2 \leq \frac{C}{\mu}$$

$$\Rightarrow \sum_{j=1}^n |V_2(E(A))_j - \theta V_2(A)_j|^2 \leq \frac{C}{\mu^2}$$

$$\text{total number of mistakes} \leq \frac{C}{\mu^2}$$



## Bernstein inequality for matrices

$$\chi_i \sim \mathcal{N}(0, \Sigma)$$

Recall that we were able to use the fact that

$$E[\exp(t(X + Y))] = E[\exp(tX) \exp(tY)],$$

and use Chernoff's method if  $X$  and  $Y$  are real valued random variables. To generalize this approach we need to define what a *matrix exponential* means:

# Bernstein inequality for matrices

Recall that we were able to use the fact that

$$E[\exp(t(X + Y))] = E[\exp(tX) \exp(tY)],$$

and use Chernoff's method if  $X$  and  $Y$  are real valued random variables. To generalize this approach we need to define what a *matrix exponential* means:

## Definition

For a function  $f : \mathbb{R} \rightarrow \mathbb{R}$  and an  $n \times n$  symmetric matrix

$$X = \sum_{i=1}^n \lambda_i u_i u_i^\top,$$

define

$$f(X) := \sum_{i=1}^n f(\lambda_i) u_i u_i^\top.$$

## Matrix power series

For a convergent power series expansion of  $f$  about  $x_0$ :

$$f(x) = \sum_{k=1}^{\infty} a_k (x - x_0)^k.$$

It is the case that series of matrix terms converges, and

$$f(X) = \sum_{k=1}^{\infty} a_k (X - x_0 I)^k.$$

As an example, for each  $n \times n$  symmetric matrix  $X$  we have

$$e^X = I + X + \frac{X^2}{2!} + \frac{X^3}{3!} + \dots$$

~~$\exp(A+B) = \exp(A) \cdot \exp(B)$~~

$$A \cdot B \neq B \cdot A$$

# Generalization of exponential inequality

## Theorem

(Golden-Thompson inequality). For any  $n \times n$  symmetric matrices  $A$  and  $B$ , we have

$$\operatorname{tr}(e^{A+B}) \leq \operatorname{tr}(e^A e^B).$$

Unfortunately, Golden-Thompson inequality does not hold for three or more matrices: in general, the inequality  $\operatorname{tr}(e^{A+B+C}) \leq \operatorname{tr}(e^A e^B e^C)$  may fail.

## Theorem

(Lieb's inequality). Let  $H$  be an  $n \times n$  symmetric matrix. Define the function on matrices

$$f(X) := \operatorname{tr} \exp(H + \log X).$$

Then  $f$  is concave on the space on positive definite  $n \times n$  symmetric matrices.

## Lemma

(Lieb's inequality for random matrices). Let  $H$  be a fixed  $n \times n$  symmetric matrix and  $Z$  be a random  $n \times n$  symmetric matrix. Then

$$\mathbb{E} \operatorname{tr} \exp(H + Z) \leq \operatorname{tr} \exp(H + \mathbb{E} Z).$$

This follows by using Jensen's inequality.

# Matrix Bernstein

## Theorem

(Matrix Bernstein's inequality). Let  $X_1, \dots, X_N$  be independent, mean zero,  $n \times n$  symmetric random matrices, such that  $\|X_i\| \leq K$  almost surely for all  $i$ . Then, for every  $t \geq 0$ , we have

$$\mathbb{P} \left\{ \left\| \sum_{i=1}^N X_i \right\| \geq t \right\} \leq 2n \exp \left( -\frac{t^2/2}{\sigma^2 + Kt/3} \right).$$

Here  $\sigma^2 = \left\| \sum_{i=1}^N \mathbb{E} X_i^2 \right\|$  is the norm of the matrix variance of the sum.

# Moment generating function of random matrices

## Lemma

$$\begin{bmatrix} 1 & 0 \\ 0 & 2 \end{bmatrix} \preceq \begin{bmatrix} 2 & 0 \\ 0 & 3 \end{bmatrix}$$

(Moment generating function). Let  $X$  be an  $n \times n$  symmetric mean zero random matrix such that  $\|X\| \leq K$  almost surely. Then

$$\mathbb{E} \exp(\lambda X) \preceq \exp(g(\lambda) \mathbb{E} X^2) \quad \text{where} \quad g(\lambda) = \frac{\lambda^2/2}{1 - |\lambda|K/3},$$

provided that  $|\lambda| < 3/K$ .

# Proof *Reduction to MGF*



# Matrix Bernstein - Step 1: Reduction to MGF

$$S := \sum_{i=1}^N X_i$$

$$\|S\| = \max_{i=1, \dots, n} |\lambda_i(S)| = \max(\lambda_{\max}(S), -\lambda_{\min}(S))$$

do  $\max \lambda$

$$\begin{aligned} \mathbb{P}(\lambda_{\max}(S) \geq t) &= \mathbb{P}(e^{\lambda \lambda_{\max}(S)} \geq e^{\lambda t}) \\ &\leq e^{-\lambda t} \mathbb{E}[e^{\lambda \lambda_{\max}(S)}] \end{aligned}$$

$$\mathbb{E} = \mathbb{E}[e^{\lambda S}] \leq \mathbb{E}[\text{tr exp}(\lambda S)]$$

## Step 2: Application of Lieb's inequality

$$EZA]$$

$$EZ EZA | B ] ]$$

$$E \leq E \operatorname{tr} \exp \left( \sum_{i=1}^{n-1} \lambda X_i + \lambda X_n \right)$$

Condition on  $(X_i)_{i=1}^{n-1}$  apply lemma with  $H := \sum_{i=1}^{n-1} \lambda X_i$

and take  $Z := \lambda X_n$

$$\leq E \operatorname{tr} \exp \left( \sum_{i=1}^{n-1} \lambda X_i + \log E Z e^{\lambda X_n} \right)$$

$$\leq \operatorname{tr} \exp \left[ \sum_{i=1}^n \log (E Z e^{\lambda X_i}) \right]$$

$$\leq \operatorname{tr} \exp (g(\lambda) Z) \quad Z := \sum_{i=1}^n E Z X_i^2]$$

$$\leq n \cdot \lambda_{\max} (\exp (g(\lambda) Z))$$

$$= n \cdot \exp (g(\lambda) \lambda_{\max} (Z))$$

$$= n \cdot \exp (g(\lambda) \sigma^2)$$

then play back into ~~the~~ and minimize get the result.

## Step 3: Using the MGF bound

# Expectation

## Lemma

*(Matrix Bernstein's inequality: expectation). Let  $X_1, \dots, X_N$  be independent, mean zero,  $n \times n$  symmetric random matrices, such that  $\|X_i\| \leq K$  almost surely for all  $i$ .*

$$\mathbb{E} \left\| \sum_{i=1}^N X_i \right\| \lesssim \left\| \sum_{i=1}^N \mathbb{E} X_i^2 \right\|^{1/2} \sqrt{1 + \log n} + K(1 + \log n).$$

## General covariance estimation

We can estimate the second moment matrix  $\Sigma = \mathbb{E}XX^T$  by its sample version

$$\Sigma_m = \frac{1}{m} \sum_{i=1}^m X_i X_i^T.$$

Recall that if  $X$  has zero mean, then  $\Sigma$  is the covariance matrix of  $X$  and  $\Sigma_m$  is the sample covariance matrix of  $X$ .

# General covariance estimation

## Theorem

*(General covariance estimation). Let  $X$  be a random vector in  $\mathbb{R}^n$ ,  $n \geq 2$ . Assume that for some  $K \geq 1$ ,*

$$\|X\|_2 \leq K(\mathbb{E}\|X\|_2)^{1/2} \quad \text{almost surely.} \quad (5.16)$$

*Then, for every positive integer  $m$ , we have*

$$\mathbb{E}\|\Sigma_m - \Sigma\| \leq C \left( \sqrt{\frac{K^2 n \log n}{m}} + \frac{K^2 n \log n}{m} \right) \|\Sigma\|.$$

# Proof

Proof cont.