

High Dimensional Statistics

Yanbo Tang

Imperial College London

March. 2025

Content of this week

- Covariance estimation

$$\frac{X^T X}{n} = \hat{\Sigma}$$

$$1/\hat{\Sigma} = \hat{\Sigma}^{-1}$$

- PCA

- Matrix Regression

multiple

$$y = X\beta + \varepsilon \quad \begin{matrix} \text{vector} \\ \text{linear regression} \end{matrix}$$

$$Y = X \odot F + F_{\perp} \quad \text{matrix}$$

Some useful matrix inequalities

$\lambda \rightarrow$ singular values.

$$A = \sum \lambda_i u_i u_i^T$$

$$\|A\|_F^2 = \sum_{i=1}^r \lambda_i(A)^2 \quad (1)$$

Weyl's inequality $\max_i |\lambda_i(A) - \lambda_i(B)| \leq \|A - B\| \quad (2)$

Hoffman's $\sum_i (\lambda_i(A) - \lambda_i(B))^2 \leq \|A - B\|_F^2 \quad (3)$

$$\langle A, B \rangle = \text{tr}(A^T B)$$

Hölder's inequality $|\langle A, B \rangle| \leq \|A\|_p \|B\|_q \quad \frac{1}{q} + \frac{1}{p} = 1 \quad (4)$

where $\|A\|_p = \left(\sum_{i=1}^r \lambda_i(A)^p \right)^{1/p}$

$$|\langle A, B \rangle| = |\text{tr}(A^T B)| \leq \sum |\lambda_i(B)| \cdot \|A\|$$

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.

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 .

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

Covariance estimation

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

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

Recall that if X has zero mean, then Σ is the covariance matrix of X and Σ_m is the sample covariance matrix of X .

Covariance Estimation sub-Gaussian

Theorem

Let $Y \in \mathbb{R}^d$ be a random vector such that $\mathbb{E}[Y] = 0$, $\mathbb{E}[YY^\top] = I_d$ and $Y \sim \text{sub}G_d(1)$. Let $X_1, \dots, X_n$ be n independent copies of sub-Gaussian random vector $X = \Sigma^{1/2}Y$. Then $\mathbb{E}[X] = 0$, $\mathbb{E}[XX^\top] = \Sigma$ and $X \sim \text{sub}G_d(\|\Sigma\|_{\text{op}})$. Moreover,

$$\|\hat{\Sigma} - \Sigma\|_{\text{F}} \lesssim \|\Sigma\|_{\text{F}} \left(\sqrt{\frac{d + \log(1/\delta)}{n}} \vee \frac{d + \log(1/\delta)}{n} \right),$$

Mart

with probability $1 - \delta$.

General covariance estimation

Theorem

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

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

Then, for every positive integer m , we have

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

$$\approx \sqrt{\frac{d \log d}{n}}$$

Proof

$$E[\|X\|_2^2] = \text{tr}(\Sigma)$$

$$\|X\|_2^2 \leq k^2 \text{tr}(\Sigma)$$

$$E[\|\hat{\Sigma} - \Sigma\|] = \frac{1}{n} E[\|\sum_{i=1}^n (X_i X_i^T - \Sigma)\|] \leq \frac{1}{n} (\sigma \sqrt{\log(d)} + M \log(d))$$

$$\sigma^2 = n \|E(X X^T - \Sigma)^2\| = \left\| \sum_{i=1}^n E(X_i X_i^T - \Sigma)^2 \right\|$$

$$M : \|X X^T - \Sigma\| \leq M \quad \text{a.s.}$$

σ^2

$$E[(X X^T - \Sigma)^2] = E[(X X^T)^2] - \Sigma^2 \preceq E[(X X^T)^2]$$

$$(X X^T)^2 = \|X\|^2 X X^T \preceq k^2 \text{tr}(\Sigma) X X^T$$

$$\Rightarrow E[(X X^T)^2] \preceq k^2 \text{tr}(\Sigma) \Sigma$$

$$\Rightarrow \sigma^2 \leq k^2 n \text{tr}(\Sigma) \|\Sigma\|$$

Proof cont.

$$\|X - \bar{X}\| \leq \|X\|_2 + \|\bar{X}\|$$

A probability bound

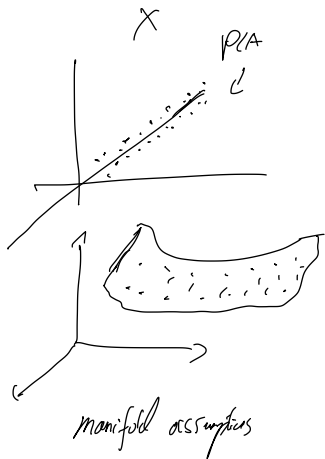
Lemma

Under the same assumption as the previous theorem, we have

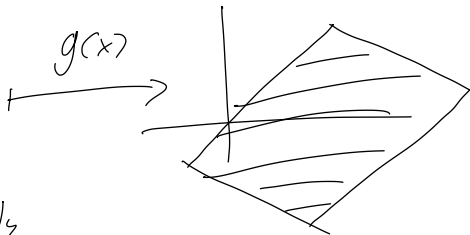
$$\|\hat{\Sigma} - \Sigma\| \lesssim \left(\sqrt{\frac{K^2 d (\log d + \log(2/\delta))}{n}} + \frac{K^2 d (\log d + \log(2/\delta))}{n} \right) \|\Sigma\|$$

with probability at least $1 - \delta$.

PCA



GANs
VAE



Eckart-Young-Mirsky

Lemma

Let A be a rank- r matrix with singular value decomposition

$$A = \sum_{i=1}^r \lambda_i u_i v_i^\top,$$

where $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_r > 0$ are the ordered singular values of A . For any $k < r$, let $A_k = \sum_{i=1}^k \lambda_i u_i v_i^\top$. Then for any matrix B such that $\text{rank}(B) \leq k$, it holds

$$\|A - A_k\|_F \leq \|A - B\|_F.$$

Moreover,

$$\textcircled{1} \quad \|A - A_k\|_F^2 = \sum_{j=k+1}^r \lambda_j^2.$$

Proof ① is by definition since $A - A_k = \sum_{j=k+1}^n \lambda_j u_j u_j^T$

$$\Rightarrow \|A - A_k\|_F^2 = \sum_{j=k+1}^n \lambda_j^2$$

$\forall B$ of rank $\leq k$
 $\sigma_1 \geq \sigma_2 \geq \dots$
 are singular values of B

$$\begin{aligned} \|A - B\|_F^2 &= \sum_{j=1}^k (\lambda_j - \sigma_j)^2 \\ &= \sum_{j=1}^k (\lambda_j - \sigma_j)^2 + \sum_{j=k+1}^n \lambda_j^2 \\ &\geq \sum_{j=k+1}^n \lambda_j^2 \end{aligned}$$

Spiked covariance

For a fixed direction $v \in S^{d-1}$, and consider the sequence $Y_1, \dots, Y_n \sim N(0, I_d)$, then the vectors $v^\top Y_i v$ all live in the one dimensional space spanned by v . Typically we would have some noise in our observations so that they will be closer to:

$$X_i = v^\top Y_i v + Z_i,$$

for a noise vector $Z_i \sim N(0, \sigma^2 I_d)$. Note that the covariance matrix of X_i will be:

$$\Sigma = E[XX^\top] = vv^\top + \sigma^2 I_d.$$

Spiked covariance

Definition

A covariance matrix $\Sigma \in \mathbb{R}^{d \times d}$ is a spiked covariance matrix if:

$$\Sigma = \theta \mathbf{v} \mathbf{v}^\top + I_d,$$

for $\theta > 0$ and $\mathbf{v} \in S^{d-1}$. The vector $\mathbf{v}$ is called the spike.

$$V_1(\hat{\Sigma}) = \hat{\mathbf{v}}$$

Davis Kahan reminder

Theorem

(Davis-Kahan) 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

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

which implies

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

Guarantees

Corollary

Let $Y \in \mathbb{R}^d$ be a random vector such that $\mathbb{E}[Y] = 0$, $\mathbb{E}[YY^T] = I_d$ and $Y \sim \text{subG}_d(1)$. Let $X_1, \dots, X_n$ be n independent copies of sub-Gaussian random vector $X = \Sigma^{1/2}Y$ so that $\mathbb{E}[X] = 0$, $\mathbb{E}[XX^T] = \Sigma$ and $X \sim \text{subG}_d(\|\Sigma\|_{\text{op}})$. Assume further that $\Sigma = \theta vv^T + I_d$ satisfies the spiked covariance model. Then, the largest eigenvector $\hat{v}$ of the empirical covariance matrix $\hat{\Sigma}$ satisfies,

$$\min_{e \in \{\pm 1\}} \|e\hat{v} - v\|_2 \lesssim \frac{1 + \theta}{\min(\theta, 1)} \left(\sqrt{\frac{d + \log(1/\delta)}{n}} \vee \frac{d + \log(1/\delta)}{n} \right)$$

with probability $1 - \delta$.

Proof By Davis-Kahan $\|V_1(\tilde{\Sigma}) - V_1(\Sigma)\|_2 \leq 2 \frac{\sigma_2}{\delta} \|\Sigma - \tilde{\Sigma}\|$

$1+\theta, \theta, 0$ $\delta = \min(1, \theta)$

$\|\Sigma - \tilde{\Sigma}\|$

Sparse PCA

If we assume that v in the spiked covariance model is k -sparse: $|v|_0 = k$. Therefore, a natural candidate to estimate v is given by $\hat{v}$ defined by

$$\hat{v}^\top \hat{\Sigma} \hat{v} = \max_{u \in S^{d-1}, |u|_0 = k} u^\top \hat{\Sigma} u.$$

It is the case that $\lambda_{\max}^k(\hat{\Sigma}) = \hat{v}^\top \hat{\Sigma} \hat{v}$ is the largest of all leading eigenvalues among all $k \times k$ sub-matrices of $\hat{\Sigma}$ so that the maximum is indeed attained.

$$\frac{\log\left(\frac{ed}{k}\right)}{n}$$



Sparse PCA

Theorem

Let $Y \in \mathbb{R}^d$ be a random vector such that $\mathbb{E}[Y] = 0$, $\mathbb{E}[YY^\top] = I_d$ and $Y \sim \text{subG}_d(1)$. Let $X_1, \dots, X_n$ be n independent copies of sub-Gaussian random vector $X = \Sigma^{1/2}Y$ so that $\mathbb{E}[X] = 0$, $\mathbb{E}[XX^\top] = \Sigma$ and $X \sim \text{subG}_d(\|\Sigma\|_{\text{op}})$. Assume further that $\Sigma = \theta vv^\top + I_d$ satisfies the spiked covariance model for v such that $|v|_0 = k \leq d/2$. Then, the k -sparse largest eigenvector $\hat{v}$ of the empirical covariance matrix satisfies,

$$\begin{aligned} & \min_{\varepsilon \in \{\pm 1\}} \|\varepsilon \hat{v} - v\|_2 \\ & \lesssim \frac{1 + \theta}{\min(\theta, 1)} \left(\sqrt{\frac{k \log(ed/k) + \log(1/\delta)}{n}} \vee \frac{k \log(ed/k) + \log(1/\delta)}{n} \right) \end{aligned}$$

with probability $1 - \delta$.

Multiple regression

$$\begin{matrix} \text{BMI} & \text{HR} & \text{Cholesterol} \\ \begin{pmatrix} Y_{11} & Y_{12} & Y_{13} \\ Y_{21} & Y_{22} & Y_{23} \\ \vdots & \vdots & \vdots \end{pmatrix} & = & X \textcircled{4} + E \rightarrow \text{min} \\ & & \text{Sub} h: \end{matrix}$$

$$Y_1 = X_1 \beta + \varepsilon_1$$

$$Y_2 = X_2 \beta + \varepsilon_2$$

⋮

Introduction

A simple question to ask is can we extend the classical regression set up to matrices? Specifically, is it useful to consider a model such that:

$$\mathbb{Y} = \mathbb{X}\Theta^* + E,$$

where $Y \in \mathbb{R}^{n \times T}$, $\mathbb{X} \in \mathbb{R}^{n \times d}$ and $\Theta \in \mathbb{R}^{d \times T}$ is the unknown matrix of coefficients for some noise matrix $E \sim \text{subG}_{n \times T}(\sigma)$.

Definition

We call a $n \times m$ matrix A $\text{subG}_{n \times m}(\sigma)$ if for every $x \in S^{m-1}$ and $y \in S^{n-1}$ if:

$$y^\top Ax \sim \text{subG}(\sigma^2).$$



Direct observation model

1

linear regression $\hat{\beta} \sim \mathcal{N}(\mu, (X^T X)^{-1} \sigma^2)$

We are going to make our life simple (at first) and assume that we have an orthogonal design (ORT condition) for our design matrix, i.e., $X^T X = nI_d$. Under the ORT assumption,

$$\frac{1}{n} X^T Y = \Theta^* + \frac{1}{n} X^T E.$$

Which can be written as an equation in $\mathbb{R}^{d \times T}$ called the *sub-Gaussian matrix model (sGMM)*:

$$y = \Theta^* + F,$$

where $y = \frac{1}{n} X^T Y$ and $F = \frac{1}{n} X^T E \sim \text{subG}_{d \times T}(\sigma^2/n)$.

$$A \in \mathbb{R}^{T \times d}$$

If Θ^* is sparse then it is possible to estimate it from a single observation. Consider the SVD of Θ^* :

$$\Theta^* = \sum_j \lambda_j u_j v_j^\top.$$

and define $\|\Theta^*\|_0 := |\lambda|_0$. Therefore, if we knew u_j and v_j , we could simply estimate the λ_j s thresholding. It turns out that estimating the eigenvectors by themselves is sufficient. Consider the SVD of the observed matrix y :

$$y = \sum_j \hat{\lambda}_j \hat{u}_j \hat{v}_j^\top.$$

Singular value thresholding

Definition The **singular value thresholding** estimator with threshold $\tau_n \geq 0$ is defined by

$$\hat{\Theta}^{\text{SVT}} = \sum_j \hat{\lambda}_j \mathbb{I}(|\hat{\lambda}_j| > \tau_n) \hat{u}_j \hat{v}_j^\top.$$

Singular value thresholding

Definition The **singular value thresholding** estimator with threshold $2\tau \geq 0$ is defined by

$$\hat{\Theta}^{\text{SVT}} = \sum_j \hat{\lambda}_j \mathbb{I}(|\hat{\lambda}_j| > \tau_n) \hat{u}_j \hat{v}_j^\top.$$

Lemma

Let A be a $d \times T$ random matrix such that $A \sim \text{subG}_{d \times T}(\sigma^2)$. Then

$$\|A\|_{\text{op}} \leq 4\sigma \sqrt{\log(12)(d \vee T)} + 2\sigma \sqrt{2 \log(1/\delta)}$$

with probability $1 - \delta$.

Theorem

Consider the multivariate linear regression model under the assumption **ORT** or, equivalently, the sub-Gaussian matrix model. Then, the singular value thresholding estimator $\hat{\Theta}^{SVT}$ with threshold

$$\tau_n = 8\sigma\sqrt{\frac{\log(12)(d \vee T)}{n}} + 4\sigma\sqrt{\frac{2\log(1/\delta)}{n}},$$

satisfies

$$\begin{aligned} \frac{1}{n} \|\bar{X} \hat{\Theta}^{SVT} - \bar{X} \Theta^*\|_F^2 &= \|\hat{\Theta}^{SVT} - \Theta^*\|_F^2 \leq 36 \text{rank}(\Theta^*) \tau_n^2 \\ &\lesssim \frac{\sigma^2 \text{rank}(\Theta^*)}{n} (d \vee T + \log(1/\delta)). \end{aligned}$$

with probability $1 - \delta$.

Proof

$$\textcircled{4}^* \quad \lambda_1 \geq \lambda_2 \geq \dots$$

$$Y \quad \hat{\lambda}_1 = \hat{\lambda}_2 = \hat{\lambda}_3 = \dots$$

define $S = \{j : |\hat{\lambda}_j| > \tau_n\}$ $\|F\| \leq \tau_n/2$ w.p 1- δ

$$|\hat{\lambda}_j - \lambda_j| \leq \|F\| \leq \tau_n/2 \quad \text{by wgl's}$$

$$\textcircled{2} \leq \sum_{j \in S^c} \lambda_j^2$$

$$S \subset \{j : |\lambda_j| > \tau_n/2\} \quad S^c \subset \{j : |\lambda_j| \leq 3\tau_n/2\}$$

$$|\hat{\lambda}_j| = |\hat{\lambda}_j - \lambda_j + \lambda_j| \geq |\lambda_j| - |\lambda_j - \hat{\lambda}_j| \geq \tau_n - \tau_n/2 \geq \tau_n/2$$

$$\textcircled{4} = \sum_{j \in S} \lambda_j u_j v_j^T$$

$$\|\hat{\Theta} - \Theta^*\|_F^2 \leq 2 \|\hat{\Theta} - \bar{\Theta}\|_F^2 + 2 \|\bar{\Theta} - \Theta^*\|_F^2 \quad \text{Young's ineq.}$$

$$\textcircled{1} \quad \|\hat{\Theta} - \bar{\Theta}\|_F^2 \leq \text{rank}(\hat{\Theta} - \bar{\Theta}) \|\hat{\Theta} - \bar{\Theta}\|_F^2 \leq 2|S| \|\hat{\Theta} - \bar{\Theta}\|_F^2$$

$$\|\hat{\Theta} - \bar{\Theta}\| \leq \|\hat{\Theta} \cdot Y\| + \|Y - \Theta^*\| + \|\Theta^* - \bar{\Theta}\| \leq \max_{j \in S^c} |\hat{\lambda}_j| + \tau_n/2 + \max_{j \in S^c} |\lambda_j|$$

Without ORT

$\leq 3\mathcal{C}$

We can then consider penalizing the rank of the matrix instead of direct truncation in cases when we don't have the ORT assumption. Let $\hat{\Theta}^{\text{RK}}$ be any solution to the following minimization problem:

$$\min_{\Theta \in \mathbb{R}^{d \times T}} \left\{ \frac{1}{n} \|Y - X\Theta\|_F^2 + \tau_n^2 \text{rank}(\Theta) \right\}.$$

Guarantees

Theorem

Consider the multivariate linear regression model (5.1). Then, the estimator by rank penalization $\hat{\Theta}^{RK}$ with regularization parameter τ_n^2 , where τ_n is defined in as in the previous theorem, satisfies

$$\frac{1}{n} \|X\hat{\Theta}^{RK} - X\Theta^*\|_F^2 \leq 2 \operatorname{rank}(\Theta^*) \tau^2 \lesssim \frac{\sigma^2 \operatorname{rank}(\Theta^*)}{n} \left(d\sqrt{T} + \log(1/\delta) \right),$$

with probability $1 - \delta$.

Proof

Why can we solve this problem efficiently

Concluding remarks

- Concentration inequalities
e.g. Chernoff bound

$\rightarrow$ Isoperimetric inequalities

Talagrand
Talagrand. Massart.



$d \rightarrow \infty$
Concentration on spheres.
on gaussian processes.

- Complexity $\sup \|X\|_{\infty}$
 $\sup \langle g, f \rangle$
 $f \in F$.

Vershynin

- Random matrix theory Terence Tao.

- minimax rates

minimizing the worst case
inf sup

Rigollet
chapter 4/