

# Optimal Compressive Covariance Sketching via Rank-One Sampling

Wenbin Wang and Ziping Zhao

School of Information Science and Technology, ShanghaiTech University, Shanghai, China

## Background and Motivation

- **High-dimensional Streaming Data:** Each time a snapshot of a data vector  $\mathbf{x} \in \mathbb{R}^d$  is generated, with  $d$  large.
- **Challenges in Modern Data Acquisition:**
  - **Data generation at unprecedented rate:** data samples are 1) not observable due to privacy or security constraints; 2) distributed at multiple locations; 3) online generated on the fly and can only be accessed once.
  - **Limited processing power at sensor platforms:** 1) **time-sensitive:** impossible to obtain a complete snapshot of the system; 2) **storage-limited:** cannot store the whole data set; 3) **power-hungry:** minimize the number of observations.
- **Covariance Sketching:** **Key Observation:** The covariance structure can be recovered without measuring the whole data stream.
- **Contribution:** We propose a **nonconvex problem** with an **efficient algorithm**, which can be proved to attain the **oracle statistical rate**.

## Sampling Model

- Consider  $n$  independent observations  $\{\mathbf{x}_t\}_{t=1}^n$ , each drawn from a zero-mean random vector  $\mathbf{x}$  with covariance matrix  $\Sigma^*$ . Given  $m$  sensing vectors  $\{\mathbf{a}_i\}_{i=1}^m$ , the quadratic measurement  $y_i$ ,  $i = 1, \dots, m$ , is given by

$$\mathbf{y} = \mathbf{A}_{\otimes} \text{vec}(\mathbf{S}) + \boldsymbol{\eta} = \mathcal{A}(\mathbf{S}) + \boldsymbol{\eta},$$

where  $\mathbf{y} [y_1, \dots, y_m]^\top$ ,  $\boldsymbol{\eta} [\eta_1, \dots, \eta_m]^\top$  are additive measurement noises,  $\mathbf{A}_{\otimes} = [(\mathbf{a}_1 \otimes \mathbf{a}_1) \cdots (\mathbf{a}_m \otimes \mathbf{a}_m)]^\top$  and  $\text{vec}(\mathbf{S})$  denotes the vectorization of  $\mathbf{S}$  obtained by stacking its columns, and  $\mathcal{A} : \mathbb{R}^{d \times d} \mapsto \mathbb{R}^m$  is a linear operator.

- **Assumptions:**
  - The sensing vectors  $\{\mathbf{a}_i\}_{i=1}^m$  are i.i.d. sub-Gaussian random variables with zero mean and identity covariance.
  - The measurement noises  $\{\eta_i\}_{i=1}^m$  are i.i.d. sub-exponential random variables with mean 0 and variance proxy  $\sigma^2$ .

## Problem Formulation

- We propose to estimate the sparse covariance matrices from quadratic measurements using the non-convex penalty

$$\min_{\Sigma \succ 0} \left\{ \frac{1}{2m} \|\mathbf{y} - \mathcal{A}(\Sigma)\|_F^2 - \tau \log \det \Sigma + \sum_{i,j} p_\lambda(|\Sigma_{ij}|) \right\}.$$

- Assumptions on the non-convex penalty function  $p_\lambda(\cdot)$ :
  - $p_\lambda(t)$  is non-decreasing on  $[0, +\infty)$  with  $p_\lambda(0) = 0$  and is differentiable on  $(0, +\infty)$ ;
  - $0 \leq p'_\lambda(t_1) \leq p'_\lambda(t_2) \leq \lambda$  for all  $t_1 \geq t_2 \geq 0$  and  $\lim_{t \rightarrow 0} p'_\lambda(t) = \lambda$ ;
  - There exists an  $\alpha > 0$  such that  $p'_\lambda(t) = 0$  for  $t \geq \alpha\lambda$ .

Prototypical examples of the non-convex penalty function  $p_\lambda(\cdot)$ :

$$\text{SACD: } p'_\lambda(t) = \begin{cases} \lambda, & \text{for } 0 < t \leq \lambda, \\ \frac{b\lambda - t}{b-1}, & \text{for } \lambda \leq t \leq b\lambda, \\ 0, & \text{for } t \geq b\lambda, \end{cases}$$

where  $b > 2$  is an additional tuning parameter.

$$\text{MCP: } p_\lambda(t) = \text{sign}(t) \lambda \cdot \int_0^{|t|} \left(1 - \frac{z}{\lambda b}\right)_+ dz$$

for some  $b > 0$ .

## Algorithm

### Algorithm 1: Majorization-Minimization Based Multistage Convex Relaxation

**Initialize**  $\tilde{\Sigma}^{(0)} = \mathbf{I}$

**for**  $k = 1, 2, \dots, K$  **do**

  update  $A_{ij}^{(k-1)} = p'_\lambda(|\tilde{\Sigma}_{ij}^{(k-1)}|)$ ;

  obtain  $\tilde{\Sigma}^{(k)}$  by solving

$$\min_{\Sigma \succ 0} \left\{ \frac{1}{2m} \|\mathbf{y} - \mathcal{A}(\Sigma)\|_F^2 - \tau \log \det \Sigma + \sum_{i,j} p_\lambda(|\Sigma_{ij}|) \right\},$$

$k = k + 1$ ;

**end for.**

### Algorithm 2: Proximal Newton Algorithm With Back-tracking Line Search

**Input**  $\Sigma^{k-1}$ ,  $\Lambda^k$ ,  $\varepsilon$

**Initialize**  $t = 0$ ,  $\Sigma_t = \Sigma^{k-1}$ ,  $\mu = 0.8$ ,  $\alpha = 0.3$

**Repeat**

$\Sigma_{t+\frac{1}{2}} \in \arg \min_{\Sigma \succ 0} \tilde{f}_t(\Sigma) + \|\Lambda \odot \Sigma\|_1$ ;

$\Delta_t = \Sigma_{t+\frac{1}{2}} - \Sigma_t$ ;

$\delta_t = \langle \nabla f(\Sigma_t), \Delta_t \rangle - \|\Lambda^k \odot \Sigma_t\|_1 + \|\Lambda^k \odot (\Sigma_t + \Delta_t)\|_1$ ;

$\beta = 1$ ,  $q = 0$ ;

**Repeat**

$\beta = \mu^q$ ,  $q = q + 1$ ;

**If**  $\Sigma_t + \beta \Delta_t \preceq 0$  **then**

      continue;

**end**

**until**  $\bar{F}(\Sigma_t + \beta \Delta_t) \leq \bar{F}(\Sigma_t) + \alpha \beta \delta_t$

$\Sigma_{t+1} = \Sigma_t + \beta \Delta_t$ ;

$t = t + 1$ ;

**until**  $\max_{i,j} \left| \langle \nabla f(\Sigma_{t+1}) + \Lambda^k \odot \Xi^k \rangle_{ij} \right| \leq \varepsilon$

**Output:**  $\Sigma^K = \Sigma_{t+1}$

## Main Results

### Essential Assumptions:

- The true covariance matrix  $\Sigma^*$  satisfies

$$0 < \frac{1}{\kappa} \leq \lambda_{\min}(\Sigma^*) \leq \lambda_{\max}(\Sigma^*) \leq \kappa < \infty,$$

for some constant  $\kappa \geq 1$ .

- There exist universal constants  $\alpha$  and  $\mu$  such that

$$\|\Sigma_S^*\|_{\min} = \min_{(i,j) \in \mathcal{S}} |\Sigma_{ij}^*| \geq (\alpha + \mu) \lambda,$$

where  $\mu \in (0, \alpha)$  satisfies  $p'_\lambda(\mu\lambda) \geq \frac{\lambda}{2}$ .

**Theorem 1:** Let  $\mathcal{S} = \{(i, j) \mid \Sigma_{ij}^* \neq 0\}$  and  $|\mathcal{S}| = s$ .

Define  $f(\Sigma) = \frac{1}{2m} \|\mathbf{y} - \mathcal{A}(\Sigma)\|_F^2 - \tau \log \det \Sigma$ .

Under some standard assumptions, the  $\varepsilon$ -optimal solution  $\tilde{\Sigma}^{(k)}$  ( $1 \leq k \leq K$ ) satisfies the following contraction property:

$$\left\| \tilde{\Sigma}^{(k)} - \Sigma^* \right\|_F \leq \frac{1}{\rho} \left( \underbrace{\|(\nabla f(\Sigma^*))_{\mathcal{S}}\|_F}_{\text{oracle rate}} + \underbrace{\varepsilon \sqrt{s}}_{\text{optimization error}} \right) + \underbrace{\delta \left\| \tilde{\Sigma}^{(k-1)} - \Sigma^* \right\|_F}_{\text{contraction}},$$

where  $\delta \in (0, 1)$  is the contraction parameter.

**Theorem 2:** Under some standard assumptions, let  $\mathbf{x}$  be a sub-Gaussian random vector with mean zero and covariance  $\Sigma^*$  and  $\{\mathbf{x}_i\}_{i=1}^n$  be a collection of i.i.d. samples drawn from  $\mathbf{x}$ , if  $\lambda \asymp \sqrt{\frac{\log d}{mn}}$ ,  $\tau \lesssim \sqrt{\frac{1}{mn}} \left\| (\Sigma^*)^{-1} \right\|_{\max}^{-1}$ ,  $\varepsilon \lesssim \sqrt{\frac{1}{mn}}$ , and  $K \gtrsim \log(\lambda \sqrt{mn}) \gtrsim \log \log d$ , then the  $\varepsilon$ -optimal solution  $\tilde{\Sigma}^{(K)}$  satisfies

$$\left\| \tilde{\Sigma}^{(K)} - \Sigma^* \right\|_F \lesssim \sqrt{\frac{s}{mn}}$$

with high probability.

## Numerical Simulations

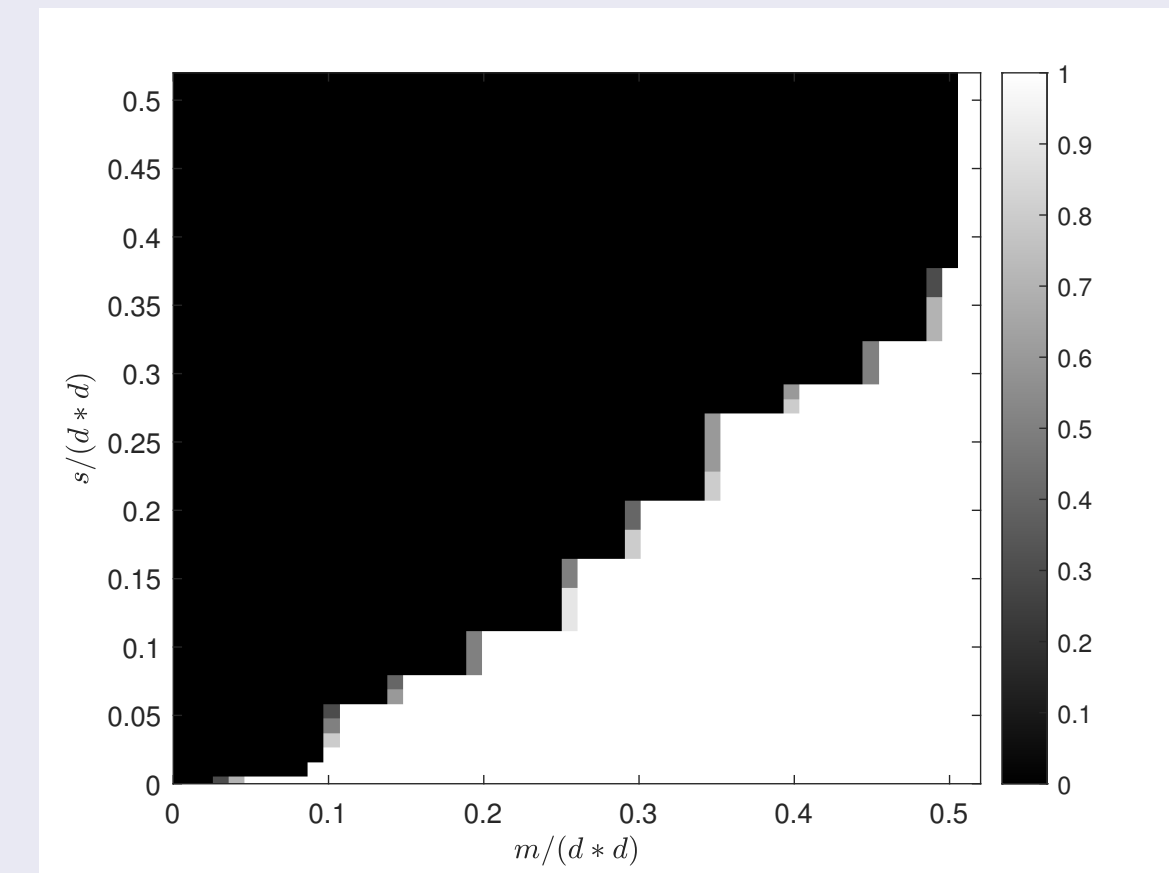
We use the MCP penalty, defined as

$$p_\lambda(t) = \text{sign}(t) \lambda \cdot \int_0^{|t|} \left(1 - \frac{z}{\lambda b}\right)_+ dz,$$

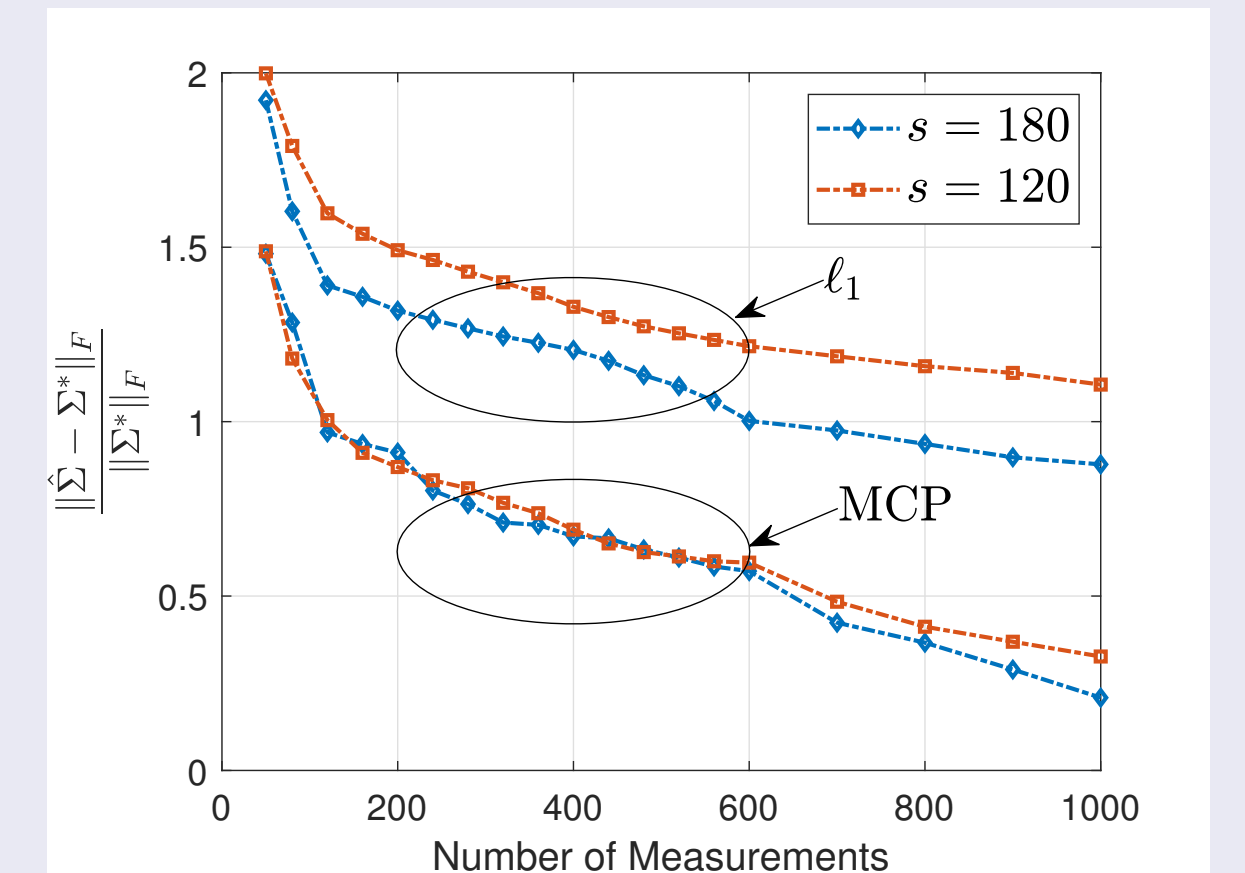
with  $b = 2$  across all trials. The regularization parameters  $\tau$  and  $\lambda$  are selected via five-fold cross-validation. The ground-truth covariance matrix  $\Sigma^*$  is generated using the built-in `sprandsym` function in MATLAB with  $s$  nonzero entries. We draw  $n = 50$  independent samples from the multivariate normal distribution  $\mathcal{N}(0, \Sigma^*)$ , and the noise variables  $\eta_i$  are sampled from a sub-exponential distribution with scale parameter  $\gamma$ , i.e.,  $\eta_i \sim \gamma \cdot \mathcal{N}(0, 1)$ . To evaluate recovery performance, we measure the success probability as visualized in the color-coded matrix in Fig. (a). To reduce the impact of limited sample size, we directly apply sketching to the true covariance matrix  $\Sigma^*$ . A recovery is considered successful if the relative Frobenius error satisfies

$$\frac{\|\Sigma - \Sigma^*\|_F}{\|\Sigma^*\|_F} \leq 10^{-3}.$$

Fig. (b) compares the proposed estimator with the  $\ell_1$ -norm-based method under a consistent noise level  $\gamma = 10^{-1}$ . As the number of measurements increases, the recovery error decreases, and our method consistently outperforms the  $\ell_1$ -based estimator.



(a) The rate of successful covariance reconstruction when  $d = 100$ .



(b) The FRE of the estimated covariance matrices for different sparsity levels with  $\gamma = 10^{-1}$