

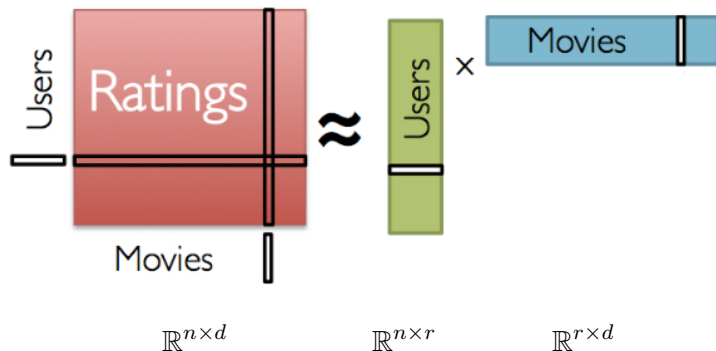
Estimating shared subspace with misalignment: the power and limitation of multiple data matrices

Yuepeng Yang

Mar 21, 2025

based on joint work with Cong Ma

Many data can be approximated as a low-rank matrix

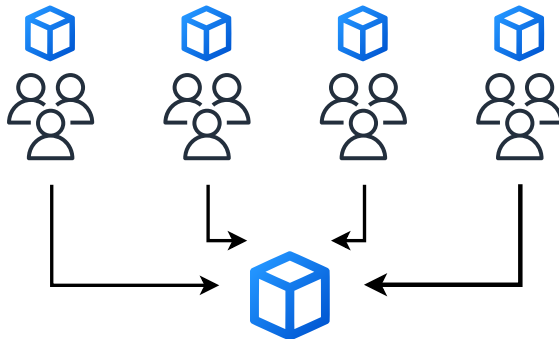


- Let $\mathbf{A} \in \mathbb{R}^{n \times d}$ be a rank- r matrix, then it can be written as

$$\mathbf{A} = \mathbf{U}\mathbf{V}^\top$$

for some $\mathbf{U} \in \mathbb{R}^{n \times r}$ such that $\mathbf{U}^\top \mathbf{U} = \mathbf{I}_r$ and $\mathbf{V} \in \mathbb{R}^{d \times r}$

- $\mathbf{U}$ is unique up to rotation and represent $\text{col}(\mathbf{A})$
- Given a noisy observation of $\mathbf{A}$, our goal is to estimate the column space $\mathbf{U}$



Learning a shared structure using multiple data sources

- Consider K data matrices $\mathbf{A}_1, \mathbf{A}_2, \dots, \mathbf{A}_K$ with shared column space

$$\mathbf{A}_k = \mathbf{U}\mathbf{V}_k^\top + \mathbf{E}_k$$

- $\mathbf{U} \in \mathbb{R}^{n \times r}$ is the shared column space
- $\mathbf{V}_k \in \mathbb{R}^{d_k \times r}$ are the loading matrices
- $\mathbf{E}_k \in \mathbb{R}^{n \times d_k}$ is entry-wise i.i.d. Gaussian noise

Can we estimate $\mathbf{U}$ from $\mathbf{A}_1, \dots, \mathbf{A}_K$ statistically efficiently?

Stack-SVD

- 1 Stack all data matrices into a single rank- r matrix

$$\mathbf{A} = [\mathbf{A}_1, \dots, \mathbf{A}_K] = \mathbf{U}[\mathbf{V}_1^\top \dots \mathbf{V}_K^\top] + [\mathbf{E}_1 \dots \mathbf{E}_K]$$

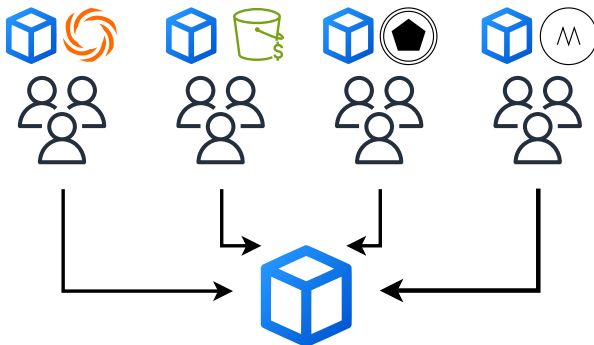
- 2 Compute $\hat{\mathbf{U}}$, the top- r left singular vectors of $\mathbf{A}$

Let $\mathbf{E}_k \sim \mathcal{N}(0, \sigma^2)$ and for simplicity take $n = d_1 = \dots = d_K$
Stack-SVD is minimax optimal:

$$\|\hat{\mathbf{U}}\hat{\mathbf{U}}^\top - \mathbf{U}\mathbf{U}^\top\| \lesssim \sigma \sqrt{\frac{n}{K}}$$

In the setting where the column space is completely shared,
we can use multiple data matrices to better estimate $\mathbf{U}$

Shared and unique component



Learning a shared structure using multiple data sources

Assume each data matrix is the sum of a component with shared column space and a component with unique column space

$$A_k = \underbrace{UV_k^\top}_{\text{rank}-r, \text{ shared component}} + \underbrace{U_k W_k^\top}_{\text{rank}-r_k, \text{ unique component}} + \underbrace{E_k}_{\text{Gaussian noise}}$$

- $U \in \mathbb{R}^{n \times r}$ is the shared column space
- $U_k \in \mathbb{R}^{n \times r_k}$ are the unique column spaces, $U_k \perp U$
- $V_k \in \mathbb{R}^{d_k \times r}$ and $W_k \in \mathbb{R}^{d_k \times r_k}$ are the loading matrices
- $E_k \in \mathbb{R}^{n \times d_k}$ is entry-wise i.i.d. Gaussian noise

Can we estimate $\mathbf{U}$ given $\mathbf{A}_k$ in presence of the unique components?

- Can we simply use Stack-SVD?
- Is the estimation of $\mathbf{U}$ fundamentally different when unique components exist?

Can we estimate $\mathbf{U}$ given $\mathbf{A}_k$ in presence of the unique components?

- Can we simply use Stack-SVD?
No. It fails even in the noiseless case.
- Is the estimation of $\mathbf{U}$ fundamentally different when unique components exist?
Yes.

AJIVE is a simple, non-iterative two-step spectral algorithm:

- 1 Estimate the shared + unique column space each $\mathbf{A}_k$:

Let $\tilde{\mathbf{U}}_k$ the top- $(r + r_k)$ left singular vectors of $\mathbf{A}_k$

- 2 Estimate the shared column space:

Let $\hat{\mathbf{U}}$ be the top- r eigenvectors of $\sum_{k=1}^K \tilde{\mathbf{U}}_k \tilde{\mathbf{U}}_k^\top$

- Angle-based: $\hat{\mathbf{U}}$ minimizes its angles from $\{\text{col}(\mathbf{A}_k)\}_{k=1}^K$
- AJIVE exactly recovers $\mathbf{U}$ in noiseless setting

Performance guarantee of AJIVE:

- Assumptions on the misalignment of the unique components
- The power and limitation of using multiple data matrices
- The role of the unique component

What does the “shared subspaces” mean?

JIVE is not yet well-defined

$$\mathbf{A}_k = \underbrace{\mathbf{U}\mathbf{V}_k^\top}_{\substack{\text{rank-}r \\ \text{shared component}}} + \underbrace{\mathbf{U}_k\mathbf{W}_k^\top}_{\substack{\text{rank-}r_k \\ \text{unique component}}}$$

- **Faithfulness:** $\text{col}(\mathbf{U}) \subset \cap_{k=1}^K \text{col}(\mathbf{A}_k)$

equiv. $\text{rank}(\mathbf{A}_k) = r + r_k$

- **Exhaustiveness:** $\cap_{k=1}^K \text{col}(\mathbf{A}_k) \subset \text{col}(\mathbf{U})$

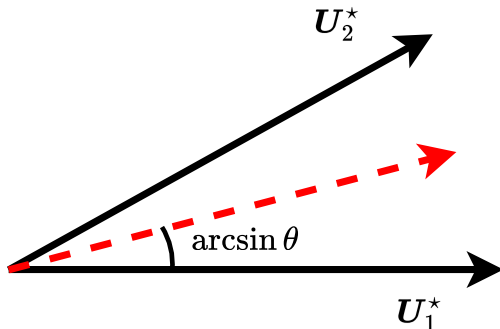
equiv.1: $\left\| \underbrace{\frac{1}{K} \sum_{k=1}^K \mathbf{U}_k \mathbf{U}_k^\top}_{\text{Avg. Proj. Mat}} \right\| < 1$ equiv.2: $\cap_{k=1}^K \text{col}(\mathbf{U}_k) = \emptyset$

The misalignment parameter θ

Misalignment: $\|\frac{1}{K} \sum_{k=1}^K \mathbf{U}_k \mathbf{U}_k^\top\| = 1 - \theta$

θ measures

- how hard it is to represent $\{\text{col}(\mathbf{U}_k)\}_{k=1}^K$ with one direction
- how far away the unique column spaces are from being shared



Naive approach (Wedin / Davis-Kahan) yields a suboptimal rate:

$$\|\hat{\mathbf{U}}_k \hat{\mathbf{U}}_k^\top - \mathbf{U} \mathbf{U}^\top\| \lesssim \frac{\sigma \sqrt{n}}{\theta}$$

We want to give a statistical analysis of AJIVE that

- shows the power of using multiple data matrices
- understand the role of the misalignment parameter θ

For simplicity suppose $n = d_1 = \dots = d_K$, $r = r_1 = \dots = r_K$,

Theorem 1: Upper bound (simplified for $K\theta > 1$)

Assuming $\sigma_{r+r_k}(\mathbf{A}_k) \geq 1$, $\sigma\sqrt{n} \leq \sqrt{\theta}$,

$$\|\hat{\mathbf{U}}\hat{\mathbf{U}}^\top - \mathbf{U}\mathbf{U}^\top\| \lesssim \sigma\sqrt{\frac{n}{K} + \frac{r}{K\theta}} + \frac{\sigma^2 n}{\theta}$$

- $\sigma\sqrt{\frac{n}{K} + \frac{r}{K\theta}}$: first-order error that scales with $\frac{1}{\sqrt{K}}$
- $\frac{\sigma^2 n}{\theta}$: non-diminishing second-order error

AJIVE is first-order minimax optimal

Theorem 2: Minimax lower bound

Assuming $\sigma_{r+r_k}(\mathbf{A}_k) \geq 1$,

$$\inf_{\tilde{\mathbf{U}}} \sup_{\{\mathbf{A}_k\}} \left\| \tilde{\mathbf{U}} \tilde{\mathbf{U}}^\top - \mathbf{U} \mathbf{U}^\top \right\| \gtrsim \sigma \sqrt{\frac{n}{K} + \frac{r}{K\theta}}$$

- $\sigma \sqrt{\frac{n}{K}}$: expected error without unique component
- $\sigma \sqrt{\frac{r}{K\theta}}$: additional mild error due to misalignment

AJIVE is first-order minimax optimal

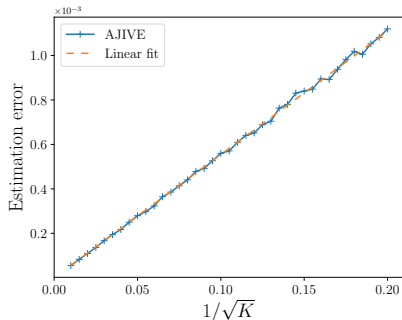
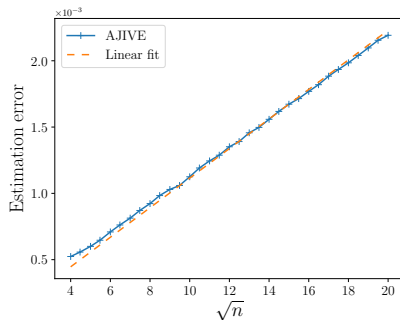
Theorem 3: Minimax lower bound

Assuming $\sigma_{r+r_k}(\mathbf{A}_k) \geq 1$,

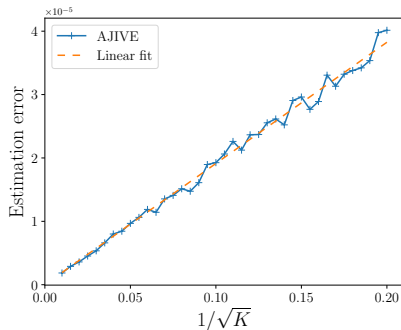
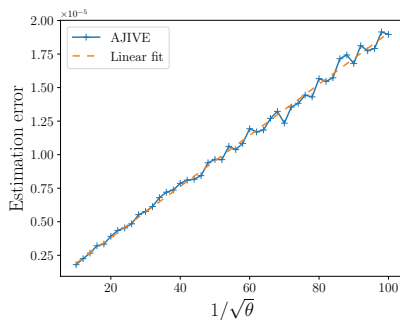
$$\inf_{\tilde{\mathbf{U}}} \sup_{\{\mathbf{A}_k\}} \left\| \tilde{\mathbf{U}} \tilde{\mathbf{U}}^\top - \mathbf{U} \mathbf{U}^\top \right\| \gtrsim \sigma \sqrt{\frac{n}{K} + \frac{r}{K\theta}}$$

- $\sigma \sqrt{\frac{n}{K}}$: expected error without unique component
- $\sigma \sqrt{\frac{r}{K\theta}}$: additional mild error due to misalignment

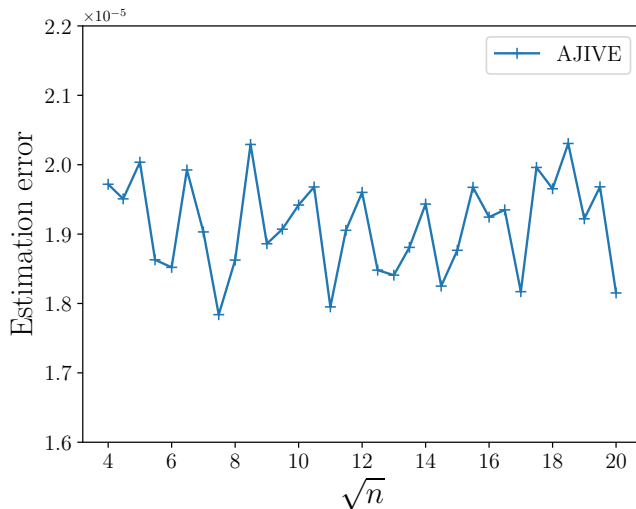
Large θ (0.5): $\|\hat{U}\hat{U}^\top - UU^\top\| \propto \sqrt{n/K}$



Small θ (0.0001 to 0.01): $\|\hat{\mathbf{U}}\hat{\mathbf{U}}^\top - \mathbf{U}\mathbf{U}^\top\| \propto \sqrt{1/K\theta}$



Empirical results: small θ (0.0001)



- Using multiple data matrices is powerful even in presence of the misaligned unique components ($\propto 1/\sqrt{K}$)
- AJIVE is first-order minimax optimal
- The extra error induced by unique components scales with intrinsic dimension r

Revisiting the upper bound of AJIVE:

Assuming $\sigma_{r+r_k}(\mathbf{A}_k) \geq 1$, $\sigma\sqrt{n} \leq \sqrt{\theta}$, and $K\theta \geq 1$

$$\|\hat{\mathbf{U}}\hat{\mathbf{U}}^\top - \mathbf{U}\mathbf{U}^\top\| \lesssim \sigma\sqrt{\frac{n}{K} + \frac{r}{K\theta}} + \underbrace{\frac{\sigma^2 n}{\theta}}_{\text{non-diminishing second-order error}}$$

Revisiting the upper bound of AJIVE:

Assuming $\sigma_{r+r_k}(\mathbf{A}_k) \geq 1$, $\sigma\sqrt{n} \leq \sqrt{\theta}$, and $K\theta \geq 1$

$$\|\hat{\mathbf{U}}\hat{\mathbf{U}}^\top - \mathbf{U}\mathbf{U}^\top\| \lesssim \sigma\sqrt{\frac{n}{K} + \frac{r}{K\theta}} + \underbrace{\frac{\sigma^2 n}{\theta}}_{\text{non-diminishing second-order error}}$$

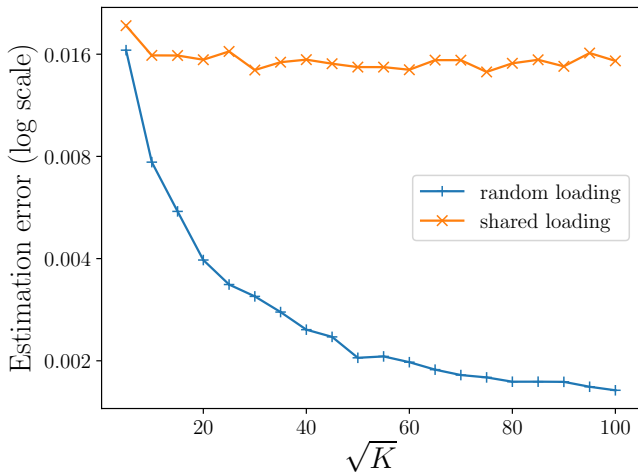
Is this term a technical artifact or a fundamental limitation?

Two experiment settings

$$\mathbf{A}_k = \underbrace{\mathbf{U}\mathbf{V}_k^\top}_{\substack{\text{rank-}r \\ \text{shared component}}} + \underbrace{\mathbf{U}_k\mathbf{W}_k^\top}_{\substack{\text{rank-}r_k \\ \text{unique component}}} + \underbrace{\mathbf{E}_k}_{\text{Gaussian noise}}$$

- Random loading: Let $\mathbf{V}_k$ and $\mathbf{W}_k$ be i.i.d. random orthonormal matrices
- Shared loading: Let $\mathbf{V}$ and $\mathbf{W}$ be random orthonormal matrices. Let $\mathbf{V}_k = \mathbf{V}$ and $\mathbf{W}_k = \mathbf{W}$ for all k

Empirical results: non-diminishing error is real



Shared vs random loading on $\|\hat{U}\hat{U}^\top - UU^\top\|$ vs K .

AJIVE

- ① Let $\tilde{U}_k$ be the top- $(r + r_k)$ left singular vectors of A_k .
 - ② Let $\hat{U}$ be the top- r eigenvectors of $\sum_{k=1}^K \tilde{U}_k \tilde{U}_k^\top$.
- SVD is biased: $\mathbb{E}[\tilde{U}\tilde{U}^\top] \neq UU^\top + U_k U_k^\top$
 - Such bias can be aligned (e.g. with shared loading), inducing a non-diminishing error in the second step of AJIVE

AJIVE

- ① Let $\tilde{U}_k$ be the top- $(r + r_k)$ left singular vectors of A_k .
 - ② Let $\hat{U}$ be the top- r eigenvectors of $\sum_{k=1}^K \tilde{U}_k \tilde{U}_k^\top$.
- SVD is biased: $\mathbb{E}[\tilde{U}\tilde{U}^\top] \neq UU^\top + U_k U_k^\top$
 - Such bias can be aligned (e.g. with shared loading), inducing a non-diminishing error in the second step of AJIVE

The non-diminishing error does exist for AJIVE, but is it fundamental to the shared subspace estimation problem?

Non-convex loss (neg. log lik. under Gaussian noise)

$$\ell(\tilde{\mathbf{U}}, \{\tilde{\mathbf{U}}_k\}, \{\tilde{\mathbf{V}}_k\}, \{\tilde{\mathbf{W}}_k\}) = \sum_{k=1}^K \|\mathbf{A}_k - \mathbf{U}\mathbf{V}_k^\top - \mathbf{U}_k\mathbf{W}_k^\top\|_{\text{F}}^2$$

- Let $\hat{\mathbf{U}}$ be one-step alternating minimization of the non-convex loss starting from the ground truth $\mathbf{U}$
- Non-diminishing error still provably exists:

$$\sup_{\mathbf{U}, \{\mathbf{U}_k\}, \{\mathbf{V}_k\}, \{\mathbf{W}_k\}} \|\hat{\mathbf{U}}\hat{\mathbf{U}}^\top - \mathbf{U}\mathbf{U}^\top\| \gtrsim \sigma^4 n^2$$

In the Neyman and Scott's problem, the law of a sequence of ind. observations $(\mathbf{A}_k)$ is governed by two sets of parameters:

- structural parameters, which appears in the law of every observation $(\mathbf{U})$
- incidental parameters, which only appears in each individual observation $(\{\mathbf{U}_k\}, \{\mathbf{V}_k\}, \{\mathbf{W}_k\})$

MLE can be inconsistent for the structural parameters

- The JIVE framework and AJIVE algorithm
- Power: AJIVE is first-order optimal
- Limitation: Non-diminishing error appears to be fundamental

- The JIVE framework and AJIVE algorithm
- Power: AJIVE is first-order optimal
- Limitation: Non-diminishing error appears to be fundamental

Future directions:

- Fully settling the non-diminishing error
- Going beyond the JIVE framework

- Feng, Q., Jiang, M., Hannig, J., and Marron, J. (2018). Angle-based joint and individual variation explained. *Journal of multivariate analysis*, 166:241–265.
- Lock, E. F., Hoadley, K. A., Marron, J. S., and Nobel, A. B. (2013). Joint and individual variation explained (jive) for integrated analysis of multiple data types. *The annals of applied statistics*, 7(1):523.