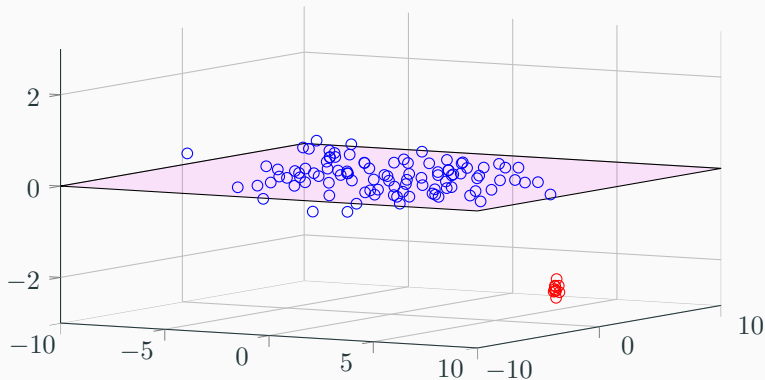


Flavors of sparse PCA

Modelization, optimization, and T-Rex 

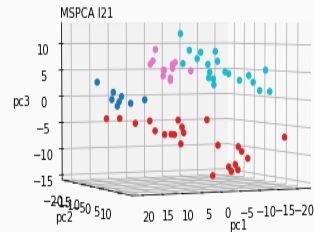
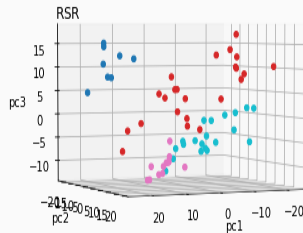
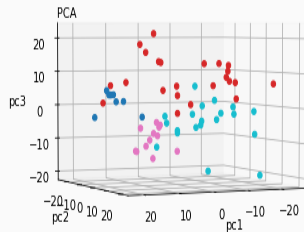
Subspace learning grounds

$$\mathbf{x}_i \simeq \mathbf{U}\mathbf{U}^\top \mathbf{x}_i, \text{ with } \mathbf{U} \in \text{St}(p, k) \triangleq \{\mathbf{U} \in \mathbb{R}^{p \times k} \mid \mathbf{U}^\top \mathbf{U} = \mathbf{I}\}$$



Can we find the best U?

Khan data: $n = 2308$ genes expression of tumors of $n = 63$ patients, with 4 cancer classes



Khan et al., "Classification and diagnostic prediction of cancers using gene expression profiling and artificial neural networks," Nature medicine, 2001

Can we define the best $\mathbf{U}$?

$$\underset{\mathbf{U} \in \text{St}(p, k)}{\text{minimize}} \quad f(\mathbf{U})$$

- **Design** the objective function f
- **Solve** the constrained minimization problem
- **Analyze** the expected performance, convergence speed, etc.
- **Apply** the result to some data

In this talk

✓ Design

- Generalize fitting costs from PCA to robust variants
- Motivating sparse PCA and sparsity promoting penalties

✓ Solve

- Quick panorama of optimization methods for $\text{St}(p, k)$

✗ Analyze

✗ Apply

✓ Try something else

- False discovery rate driven variable selection in SPCA

How to design a new “PCA”

Principal component analysis (PCA)

“Vanilla” PCA of **rank** k

- Singular value decomposition (**SVD**) of the data matrix $\mathbf{x} = [\mathbf{x}_1, \dots, \mathbf{x}_n] \in \mathbb{R}^{p \times n}$

$$\mathbf{X} \stackrel{\text{SVD}}{=} [\mathbf{U} | \mathbf{U}_\perp] \mathbf{D} \mathbf{V}^\top$$

- **Loading vectors** $\mathbf{U} \in \text{St}(p, k)$
- **Principal components** $\mathbf{x}_i^k = \mathbf{U}^\top \mathbf{x}_i \in \mathbb{R}^k$
- **Projected data** $\tilde{\mathbf{x}}_i = \mathbf{U} \mathbf{x}_i^k = \mathbf{U} \mathbf{U}^\top \mathbf{x}_i$

Beyond “vanilla” PCA

Generalizing PCA?

- Find a **model** or **framework** that yields PCA as a solution

$$\underset{\mathbf{U} \in \text{St}(p, k)}{\text{minimize}} \quad f(\mathbf{U})$$

- Leverage corresponding **tools** and **extensions**

$$\underset{\mathbf{U} \in \text{St}(p, k)}{\text{minimize}} \quad \mathfrak{f}(\mathbf{U})$$

A major motivation: being robust to **heavy-tailed** distributions and **outliers**

Probabilistic and Bayesian PCA

- **Signal** plus **noise** model

$$\mathbf{x} \stackrel{d}{=} \mathbf{U}\mathbf{s} + \mathbf{n}$$

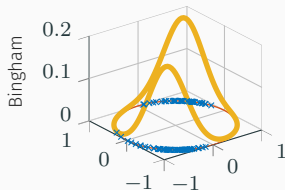
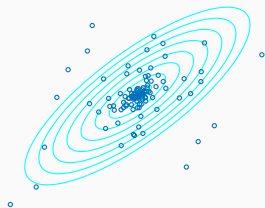
- **Probabilistic PCA** “MLE of $\mathbf{U}$ defines f ”

$$\mathbf{s} \sim \mathcal{N}(\mathbf{0}, \Sigma_s) \quad \text{and} \quad \mathbf{n} \sim \mathcal{N}(\mathbf{0}, \sigma^2 \mathbf{I})$$

$$\implies \mathbf{x} \sim \mathcal{N}(\mathbf{0}, \mathbf{U}\Sigma_s\mathbf{U}^\top + \sigma^2 \mathbf{I})$$

- **Extensions** by **changing assumptions**

- **Elliptical model** for $\mathbf{x}$, $\mathbf{s}$, or $\mathbf{n}$
- **Prior distribution** on $\mathbf{U}$ directional statistics



Geometric PCA

- **Euclidean distance** from $\mathbf{x}$ to $\text{span}(\mathbf{U})$

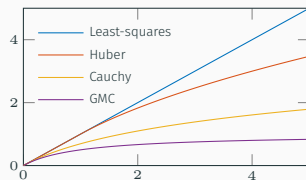
$$\text{dist}(\mathbf{U}, \mathbf{x}) = \sqrt{\mathbf{x}^\top \mathbf{x} - \mathbf{x}^\top \mathbf{U} \mathbf{U}^\top \mathbf{x}}$$

- **Geometric PCA**

$$\underset{\mathbf{U} \in \text{St}(p, k)}{\text{minimize}} \quad \sum_{i=1}^n \text{dist}^2(\mathbf{U}, \mathbf{x}_i)$$

- **Extensions** using **alternate distances**

- Robust costs: $f(\mathbf{U}) = \sum_{i=1}^n \rho(\text{dist}^2(\mathbf{U}, \mathbf{x}_i))$



Algebraic PCA

- **Low-rank approximation**

$$\begin{array}{ll} \underset{\mathbf{Y}}{\text{minimize}} & ||\mathbf{X} - \mathbf{Y}||_F^2 \\ \text{subject to} & \text{rank}(\mathbf{Y}) = k \end{array}$$

- **Extensions** using **alternate decompositions/structures**

- Low-rank plus sparse recovery (*Robust PCA*)
- Matrix completion (missing entries)
- Non-negative matrix factorization
- Additional structure in the principal components → SPCA?

Sparse PCA

Why Sparse PCA?

Recalling

- **Loading vectors** $\mathbf{U} = [\mathbf{u}_1, \dots, \mathbf{u}_k] \in \text{St}(p, k)$
- **Principal components**

$$[\tilde{\mathbf{x}}_i]_j = \mathbf{u}_j^\top \mathbf{z}_i \in \mathbb{R}$$

- The j^{th} -PC is a **linear combination** of initial variables **weighted by the loading vector**

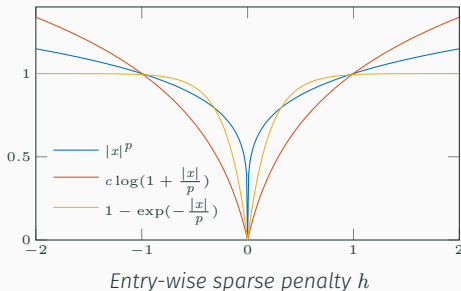
Loading vectors should be sparse

- Improved PC's **interpretation**
- Double duty: **dimension reduction & variable selection**

Sparse PCA from the optimization point of view

$$\underset{\mathbf{U} \in \text{St}(p,k)}{\text{minimize}} \quad f(\{\mathbf{x}_i\}_{i=1}^n, \mathbf{U}) + \lambda h(\mathbf{U})$$

- f is a **fitting cost** cf. many examples
- h is a **sparsity-promoting penalty** mimics ℓ_1 - or $\ell_{2,1}$ -norm



Optimization on $\text{St}(p, k)$

Existing solutions

- **Solution #1:** Riemannian optimization on $\text{St}(p, k)$

BOU23 Boumal, “An introduction to optimization on smooth manifolds,” Cambridge University Press, 2023

ABSO9 Absil, Mahony, Sepulchre, “Optimization algorithms on matrix manifolds,” Princeton Univ. Press, 2009

EDE98 Edelman, Arias, Smith, “The geometry of algorithms with orthogonality constraints,” SIMAX, 1998

- **Solution #2:** Majorization-Minimization ticks

KIE02 Kiers, “Setting up alternating least squares and iterative majorization algorithms for solving various matrix optimization problems,” Computational statistics & data analysis, 2002

BRE21 Breloy et al. “MM on the Stiefel Manifold With Application to Robust Sparse PCA”, IEEE TSP, 2021

- **Solution #3:** ADMM splitting tricks

BRE23 Brehier et al. “Robust and globally sparse PCA via MM and variable splitting,” ICASSP, 2023

UEM19 Uematsu, Fan, Chen, Lv, Lin, “SOFAR: Large-Scale Association Network Learning,” IEEE Trans. on IT, 2019

#1 Riemannian optimization

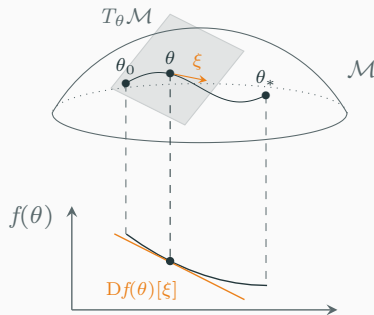
Tools from differential geometry

- Manifold $\mathcal{M}$, tangent spaces $T_\theta\mathcal{M}$
- Metric $\langle \cdot, \cdot \rangle_\theta$
- Retraction $R_\theta : T_\theta\mathcal{M} \rightarrow \mathcal{M}$

Riemannian gradient descent

$$\langle \text{grad}_{\mathcal{M}} f(\theta), \xi \rangle_\theta = Df(\theta)[\xi]$$

$$\theta_{i+1} = R_{\theta_i}(-t_i \text{grad}_{\mathcal{M}} f(\theta_i))$$



Pros: elegant, generic and flexible framework, links with information geometry

Cons: can be tedious and computationally expensive, no convincing proximal method (yet?)

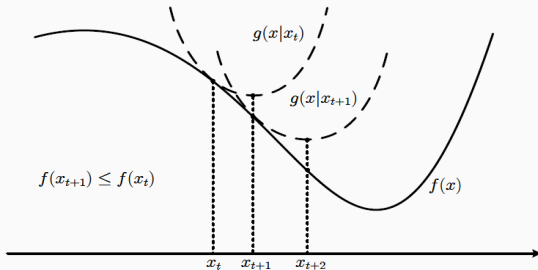
#2 Majorization-Minimization

Majorize f at $\mathbf{x}_t$ by a surrogate g

$$f(\mathbf{x}) \leq g(\mathbf{x}|\mathbf{x}_t)$$

Minimize g to iterate

$$\mathbf{x}_{t+1} = \operatorname{argmin}_{\mathbf{x} \in \mathcal{X}} g(\mathbf{x}|\mathbf{x}_t)$$



Pros: super fast (in time), covers a large class of functions

Cons: yet not universal, has to be tailored (or proxy the objective)

#2 Majorization-Minimization framework for $\text{St}(p, k)$

Proposal in BRE21

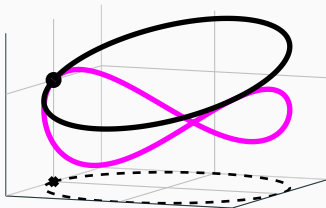
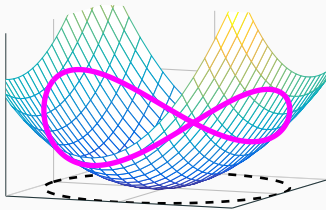
- A systematic trick to deal with orthonormality
- A catalog of surrogates for usual costs
- Applications to probabilistic and sparse PCA

Majorize f on $\text{St}(p, k)$ by a linear surrogate

$$g(\mathbf{U}|\mathbf{U}^t) = -2\Re\{\text{Tr}\{\mathbf{R}_{\mathbf{U}^t}^H \mathbf{U}\}\} + \text{const.}$$

Minimize g on $\text{St}(p, k) \Leftrightarrow$ Projection

$$\mathbf{U}^{t+1} = \mathcal{P}_{\text{St}}\{\mathbf{R}_{\mathbf{U}^t}^H\}$$



#3 ($\mathcal{M}$)-ADMM and distributed algorithms

Variable splitting

$$\begin{aligned} & \underset{\mathbf{U}, \mathbf{V}}{\text{minimize}} && f_{\mathbf{u}}(\mathbf{U}) + f_{\mathbf{v}}(\mathbf{V}) \\ & \text{subject to} && \mathbf{U} \in \text{St}(p, k) \\ & && \mathbf{U} = \mathbf{V} \end{aligned}$$

Cyclic updates over $\{\mathbf{U}, \mathbf{V}, \mathbf{\Gamma}\}$ on **augmented Lagrangian** problem

$$\underset{\{\mathbf{U}_i\} \in \text{St}(p, k), \mathbf{V}}{\text{minimize}} \quad f_{\mathbf{u}}(\mathbf{U}) + f_{\mathbf{v}}(\mathbf{V}) + \gamma \|\mathbf{U} - \mathbf{V}\|_F^2 + \langle \mathbf{\Gamma}, \mathbf{U} - \mathbf{V} \rangle$$

Pros: generalization to distributed setting, splits problems into simple ones

Cons: relaxation trick, still requires to handle $\mathbf{U}$ with previous methods

Some experiments

Optimization benchmark

Robust subspace recovery

$$\underset{\mathbf{U} \in \text{St}(p, k)}{\text{minimize}} \quad \sum_{i=1}^n \rho_H(\text{dist}^2(\mathbf{U}, \mathbf{x}_i))$$

with **Huber loss**

$$\rho_H(t) = \begin{cases} t/\sqrt{T} & \text{if } t \leq T \\ 2\sqrt{t} - \sqrt{T} & \text{if } t > T \end{cases}$$

→ **compare iterative algorithms**

Different algorithms

(and computational bottlenecks)

LER17 Quadratic MM, data matrix version

rank- k SVD($p \times n$)

$$\mathbf{U}^{t+1} = \mathcal{P}_k\{\tilde{\mathbf{Z}}_t\}, \text{ with } [\tilde{\mathbf{Z}}_t]_{:,i} = \sqrt{\rho'(\text{dist}^2(\mathbf{U}^t, \mathbf{z}_i))} \mathbf{z}_i$$

MAR05 Fixed point heuristic, covariance matrix version

rank- k SVD($p \times p$)

$$\mathbf{U}^{t+1} = \mathcal{P}_k\{\mathbf{M}(\mathbf{U}^t)\}, \text{ with } \mathbf{M}(\mathbf{U}^t) = \tilde{\mathbf{Z}}_t \tilde{\mathbf{Z}}_t^H$$

MAN02 Steepest descent on Stiefel

$\times$ thin-SVDs($p \times k$)

$$\mathbf{U}^{t+1} = \mathcal{P}_{\text{St}}\{\mathbf{U}^t + \gamma \nabla_f(\mathbf{U}^t)\}, \text{ with the right } \gamma$$

MAN02 Newton method on Stiefel

$(p \times k)^2$ system

$$\mathbf{U}^{t+1} = \mathcal{P}_{\text{St}}\{\mathbf{U}^t + \mathbf{Y}\}, \text{ with } \mathbf{Y} = \text{cpoint}(\mathbf{U}^t, \nabla_f(\mathbf{U}^t), \mathbf{H}_f(\mathbf{U}^t))$$

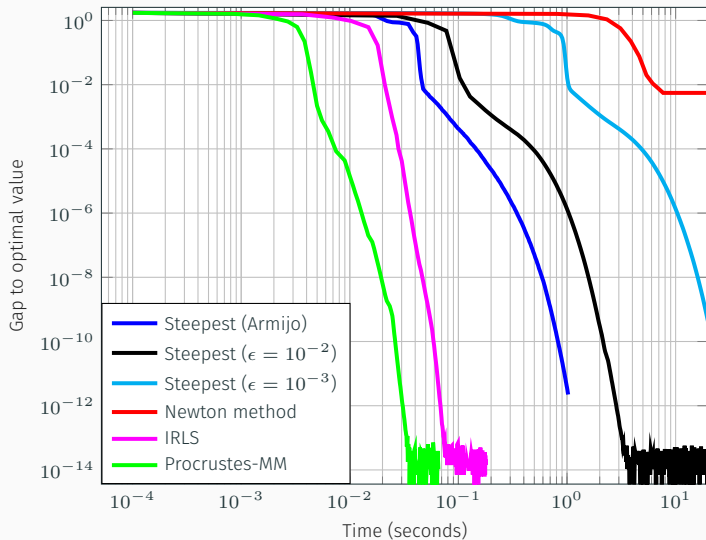
DIN06 Procrustes-MM

thin-SVD($p \times k$)

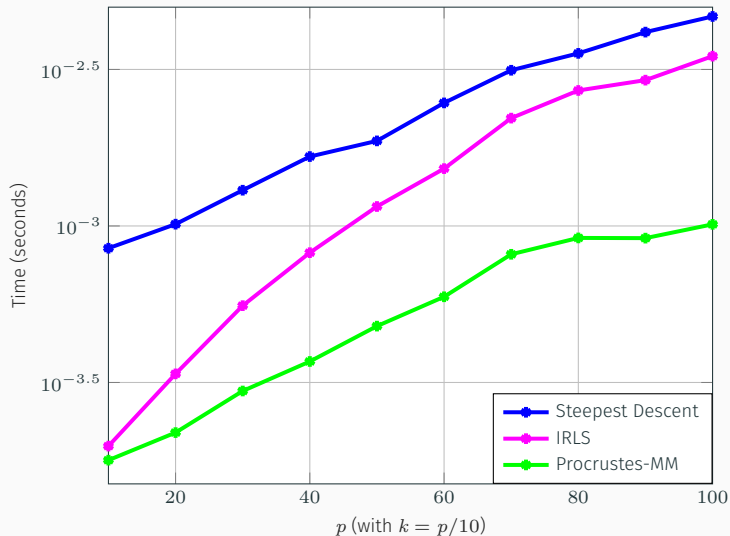
$$\mathbf{U}^{t+1} = \mathcal{P}_{\text{St}}\{\mathbf{M}(\mathbf{U}^t) \mathbf{U}^t\}$$

Objective value (-optimal value) versus CPU time

$(p = 30, k = 5, n = 100)$



Average CPU time of an iteration versus size and rank



Robustness and sparsity

$$\underset{\mathbf{U} \in \text{St}(p, k)}{\text{minimize}} \quad \sum_{i=1}^n \rho(\text{dist}^2(\mathbf{U}, \mathbf{x}_i)) + \lambda h(\mathbf{U})$$

- **Fitting**

- **Least-squares** $\rho(x) = x$
- **Huber Loss** $\rho_H(x)$

- **Penalty**

- **“PCA”** $\lambda = 0$
- **SPCA** $\lambda \neq 0 + h$ is smooth- ℓ_1

Subspace recovery

Sampling

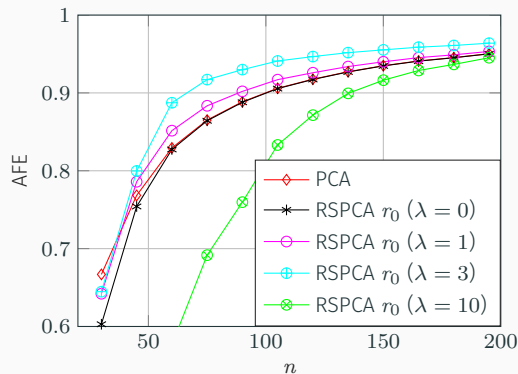
$$\{\mathbf{x}_i\}_{i=1}^n = \{\{\mathbf{x}_i^{\text{in}}\}_{i=1}^{n_{\text{in}}}, \{\mathbf{x}_i^{\text{out}}\}_{i=n_{\text{in}}+1}^n\}$$

$$\mathbf{x}^{\text{in}} \sim \mathcal{CN}(\mathbf{0}, \text{SNR} \times \mathbf{U}\mathbf{U}^H + \mathbf{I})$$

$$\mathbf{x}^{\text{out}} \sim \mathcal{CN}(\mathbf{0}, \text{ONR} \times \mathbf{U}_{\perp}\mathbf{U}_{\perp}^H + \mathbf{I})$$

Metric

$$\text{AFE} = \mathbb{E} \left[\text{Tr}\{\hat{\mathbf{U}}^H \mathbf{U}\mathbf{U}^H \hat{\mathbf{U}}\} / k \right]$$



$p = 100, k = 10, \text{SNR} = 10, \text{dense loadings}$

Subspace recovery

Sampling

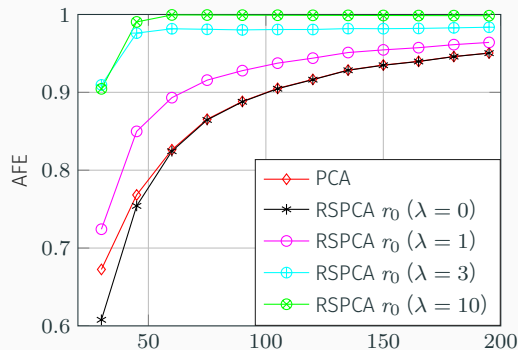
$$\{\mathbf{x}_i\}_{i=1}^n = \{\{\mathbf{x}_i^{\text{in}}\}_{i=1}^{n_{\text{in}}}, \{\mathbf{x}_i^{\text{out}}\}_{i=n_{\text{in}}+1}^n\}$$

$$\mathbf{x}^{\text{in}} \sim \mathcal{CN}(\mathbf{0}, \text{SNR} \times \mathbf{U}\mathbf{U}^H + \mathbf{I})$$

$$\mathbf{x}^{\text{out}} \sim \mathcal{CN}(\mathbf{0}, \text{ONR} \times \mathbf{U}_{\perp}\mathbf{U}_{\perp}^H + \mathbf{I})$$

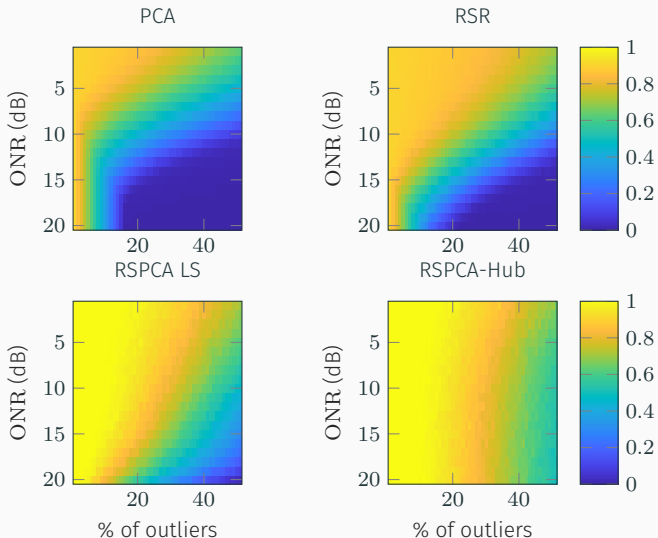
Metric

$$\text{AFE} = \mathbb{E} \left[\text{Tr}\{\hat{\mathbf{U}}^H \mathbf{U}\mathbf{U}^H \hat{\mathbf{U}}\} / k \right]$$



$n, p = 100, k = 10, \text{SNR} = 10, \text{sparse loadings}$

Robustness to outliers



FDR controlled SPCA

Sparse PCA: EV vs SP trade-off

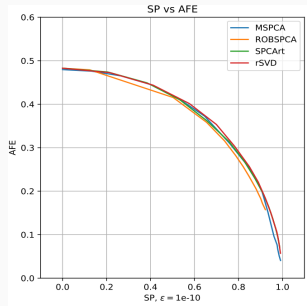
Limitation

- SPCA is only driven by **explained variance** (EV)
- Not always relevant for variable selection
- Might capture high EV **outliers**

New point of view in this work

- Links with **SPCA**
- Focuses on **variable selection** step
- Reformulates the optimization with promising criterion

FDR = “% of allowed false variable selection”



Back to the roots of Sparse PCA

Inspiration:

Zou, Hastie, Tibshirani, “Sparse principal component analysis,” J. of computational and graphical statistics, 2006

2-step **SPCA** given **pre-computed plug-in PCs** $\tilde{\mathbf{X}}$

$$\underset{\mathbf{U} \in \text{St}(p, k)}{\text{minimize}} \quad \|\tilde{\mathbf{X}} - \mathbf{U}^\top \mathbf{X}\|_F^2 + \lambda \|\mathbf{U}\|_1$$

Relaxation of orthogonality constraint $\rightarrow$ series of k **elastic-net** problems

$$\underset{\beta_j}{\text{minimize}} \quad \|[\tilde{\mathbf{X}}]_{j,:}^\top - \mathbf{X}^\top \beta_j\|_F^2 + \lambda \|\beta_j\|^2 + \lambda_1 \|\beta_j\|_1$$

Sparse loading vectors $\mathbf{U}_{\text{SPCA}} = [\beta_1, \dots, \beta_k]$

T-Rex PCA: motivations

SPCA boils down to solve k problems

$$\underset{\beta_j}{\text{minimize}} \quad ||[\tilde{\mathbf{X}}]_{j,:}^\top - \mathbf{X}^\top \beta_j||_F^2 + \lambda ||\beta_j||^2 + \lambda_1 ||\beta_j||_1$$

and **tune** λ_1 to achieve a **good SP-vs-EV**

Can we change the paradigm?

- move away from this trade-off
- have guarantees on false-discovery-rate
- be more computationally efficient

We are going to ask all of this to the **T-Rex** 🦖

False Discovery Rate (FDR) and True Positive Rate (TPR)

- # samples: n
- # candidate variables: $p > n$
- Active variables: $\mathcal{A} \subseteq \{1, \dots, p\}$ ← the “true” support of our loadings
- Selected variables: $\hat{\mathcal{A}} \subseteq \{1, \dots, p\}$

$$\text{FDR} := \mathbb{E}[\text{FDP}] := \mathbb{E}\left[\frac{|\hat{\mathcal{A}} \setminus \mathcal{A}|}{1 \vee |\hat{\mathcal{A}}|}\right] \quad \text{and} \quad \text{TPR} := \mathbb{E}[\text{TPP}] := \mathbb{E}\left[\frac{|\mathcal{A} \cap \hat{\mathcal{A}}|}{1 \vee |\mathcal{A}|}\right]$$

Goal

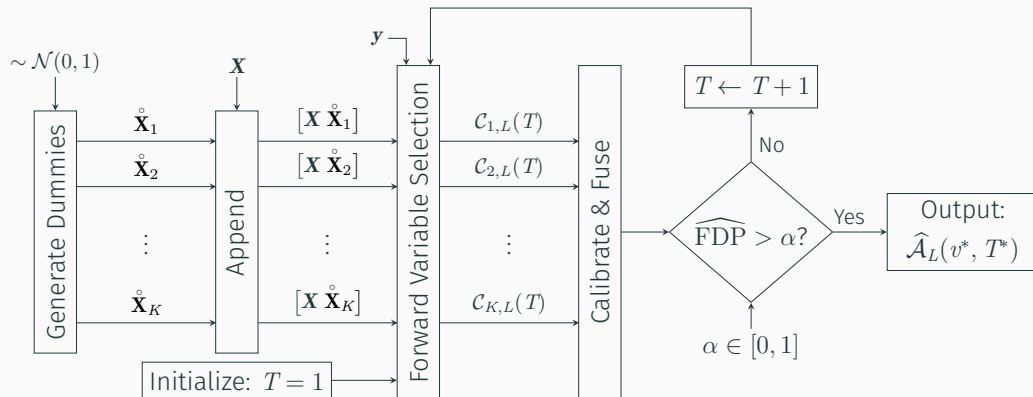
$$\max_{\hat{\mathcal{A}}} \text{TPR} \quad \text{s.t.} \quad \text{FDR} \leq \alpha,$$

where $\alpha \in [0, 1]$ is the user-defined target FDR level.

Terminating-Random Experiments (T-Rex): the ingredients

- **LARS algorithm**
 - Solves **LASSO** or **Elastic-Net** iteratively
 - Produces **a sequence of selected variables**
- **Terminated random experiment**
 - Append L dummies to $\mathbf{X} := [\mathbf{X}; \overset{\circ}{\mathbf{X}}]$ usually $\sim \mathcal{N}(0, 1)$
 - Apply **LARS** and stop when T **dummies are selected**
- **T-Rex selector**
 - Run **term. rand. exp. in parallel** and **vote for** $\hat{\mathcal{A}}$ threshold $v \in (0.5, 1]$
 - **FDP can be estimated** at this step
 - **Calibrate** T and v to **maximize TPP** and **achieve desired FDR target**

T-Rex Selector



- $L \in \mathbb{N}_+$ generated dummies
- $T \in \{1, \dots, L\}$ included dummies before stopping
- $v \in [0.5, 1)$ voting level
- $\Phi_{T,L}(j)$ relative occurrence of variable j
- $\hat{\mathcal{A}}_L(v, T) := \{j : \Phi_{T,L}(j) > v\}$

Back to Sparse PCA

Inspiration:

Zou, Hastie, Tibshirani, "Sparse principal component analysis," J. of computational and graphical statistics, 2006

2-step **SPCA** given **pre-computed plug-in PCs** $\tilde{\mathbf{X}}$

$$\underset{\mathbf{U} \in \text{St}(p, k)}{\text{minimize}} \quad \|\tilde{\mathbf{X}} - \mathbf{U}^\top \mathbf{X}\|_F^2 + \lambda \|\mathbf{U}\|_1$$

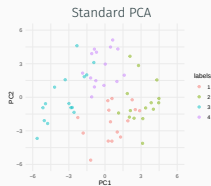
Relaxation of orthogonality constraint $\rightarrow$ series of k **elastic-net** problems

~~$$\underset{\beta_j}{\text{minimize}} \quad \|[\tilde{\mathbf{X}}]_{j,:}^\top - \mathbf{X}^\top \beta_j\|_F^2 + \lambda \|\beta_j\|^2 + \lambda_1 \|\beta_j\|_1 \rightarrow \text{🦖}$$~~

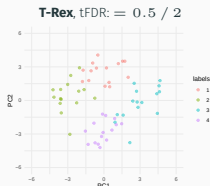
Sparse loading vectors $\mathbf{U}_{\text{SPCA}} = [\beta_1, \dots, \beta_k]$

Cleaning PCA with the T-Rex Selector

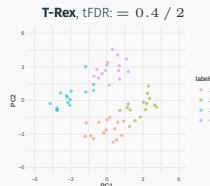
Setup (see setup in R pack. scPCA): 4 Classes, $n = 60$, $p = 150$, $p_1 = 10$.



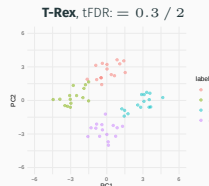
(a) FDP = 1, TPP = 1, #Var = 150.



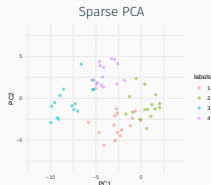
(b)



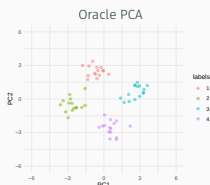
(c)



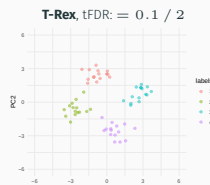
(d)



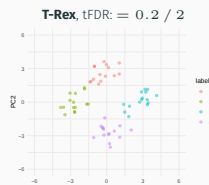
(h) FDP $\approx$ 0.9, TPP = 0.5,
#Var = 48.



(g) FDP = 0, TPP = 1, #Var = 10.

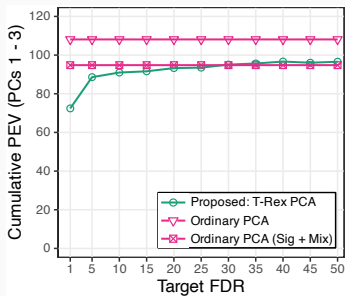
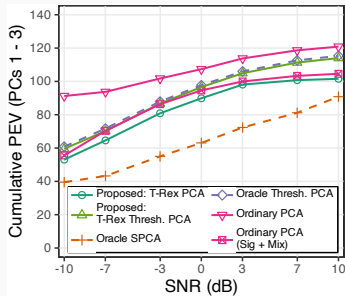
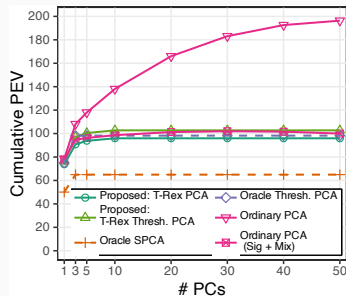


(f) FDP $\approx$ 0.09, TPP = 1,
#Var = 11.

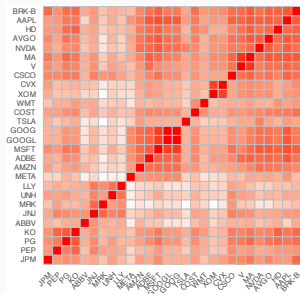


(e)

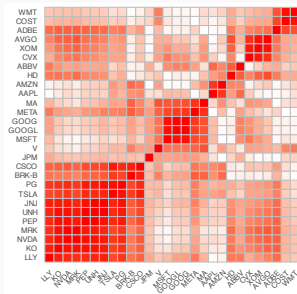
More experiments



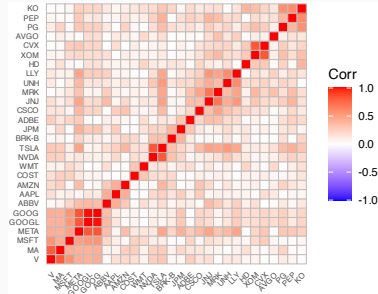
Experiments on the 28 most influential S&P 500 stocks



Raw covariance matrix



T-Rex SPCs removed



PCs removed

Covered so far

✓ Design

- Generalize fitting costs from PCA to robust variants
- Motivating sparse PCA and sparsity promoting penalties

✓ Solve

- Quick panorama of optimization methods for $\text{St}(p, k)$

✗ Analyze

✗ Apply

✓ Try something else

- False discovery rate driven variable selection in SPCA

Thanks!

Papers in this talk

Machkour, J., Breloy, A., Muma, M., Palomar, D. P., & Pascal, F. (2024, April). Sparse PCA with false discovery rate controlled variable selection. In ICASSP 2024-2024 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP) (pp. 9716-9720). IEEE.

Brehier, H., Breloy, A., El Korso, M. N., & Kumar, S. (2023, June). Robust and globally sparse PCA via majorization-minimization and variable splitting. In ICASSP 2023-2023 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP) (pp. 1-5). IEEE.

Collas, A., Bouchard, F., Breloy, A., Ginolhac, G., Ren, C., & Ovarlez, J. P. (2021). Probabilistic PCA from heteroscedastic signals: geometric framework and application to clustering. IEEE Transactions on Signal Processing, 69, 6546-6560.

Breloy, A., Kumar, S., Sun, Y., & Palomar, D. P. (2021). Majorization-minimization on the Stiefel manifold with application to robust sparse PCA. IEEE Transactions on Signal Processing, 69, 1507-1520.

T-Rex demo page https://cran.r-project.org/web/packages/TRexSelector/vignettes/TRexSelector_usage_and_simulations.html