

GEOMETRIC OPTIMIZATION

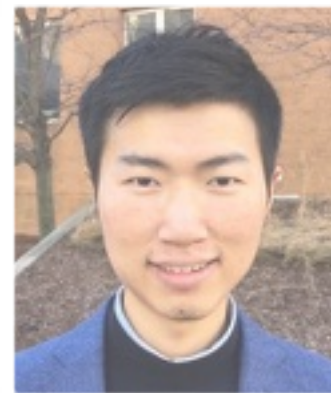
SUVRIT SRA

**Laboratory for Information and Decision Systems
Massachusetts Institute of Technology**

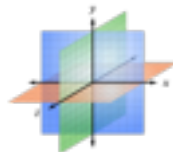
Includes work with:
Reshad Hosseini
Pouya H. Zadeh
Hongyi Zhang



π



► Vector spaces



► Manifolds



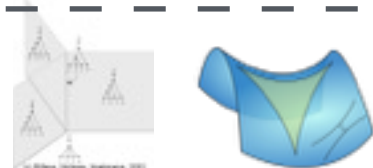
(hypersphere, orthogonal matrices, complicated surfaces)

► Convex sets



(probability simplex, semidefinite cone, polyhedra)

► Metric spaces



(tree space, Wasserstein spaces, $CAT(0)$, space-of-spaces)

Geometric Optimization

Machine Learning

Graphics

Robotics

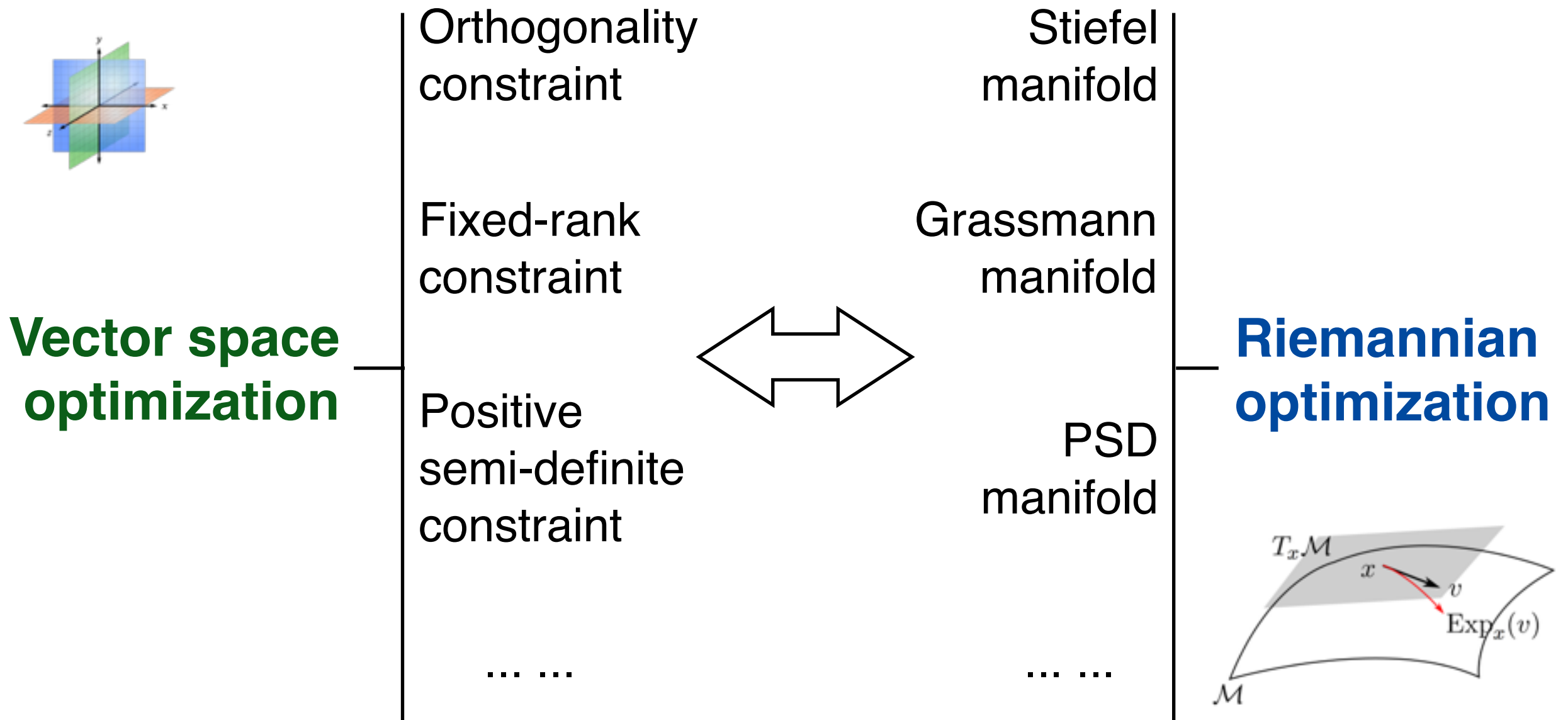
Vision

BCI

NLP

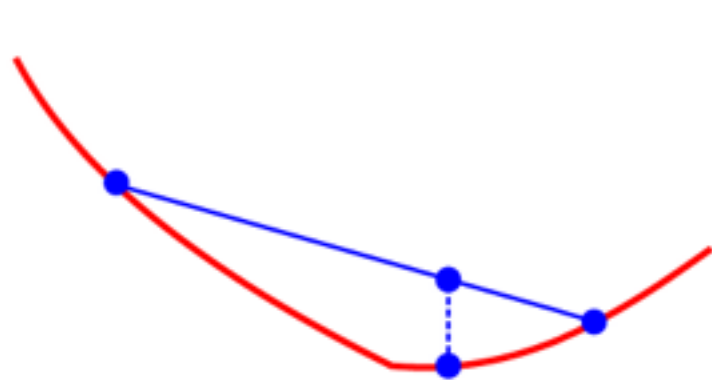
Statistics

Example: Riemannian optimization

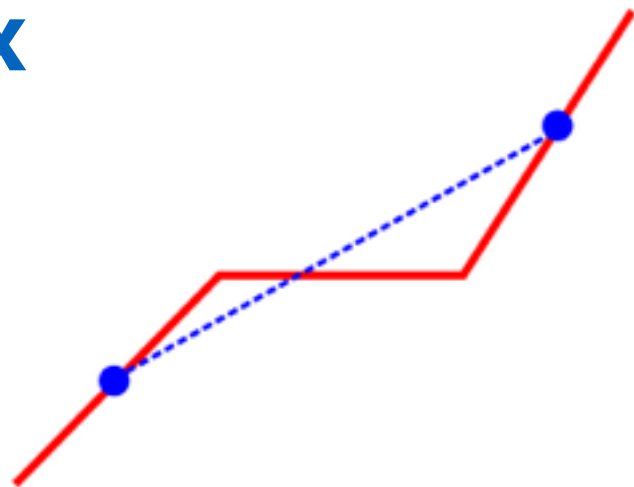


[Udriste, 1994; Absil et al., 2009]

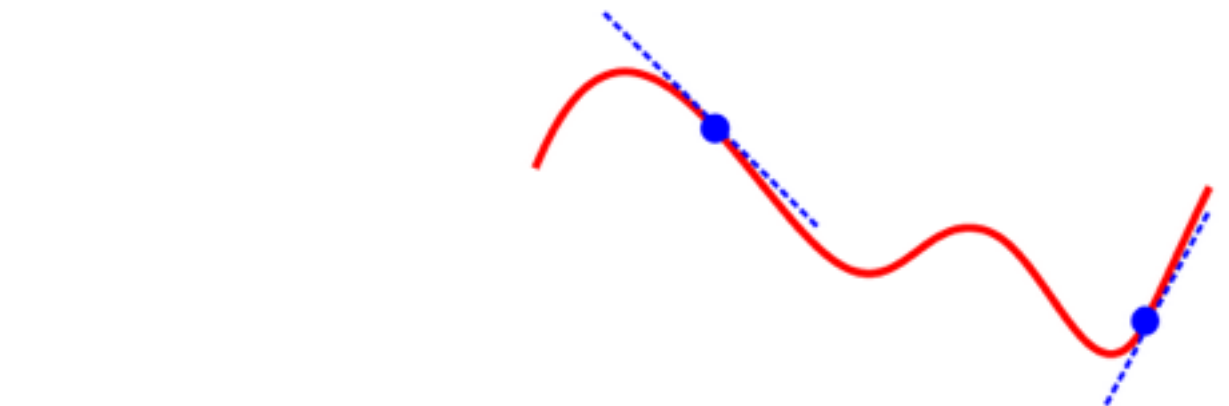
Classes of function in optimization



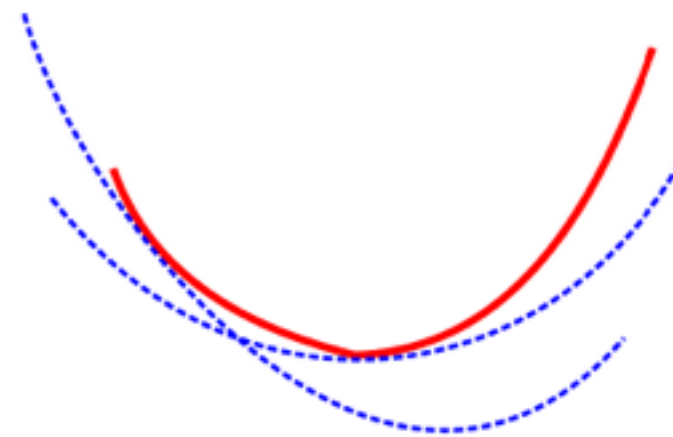
Convex



Lipschitz

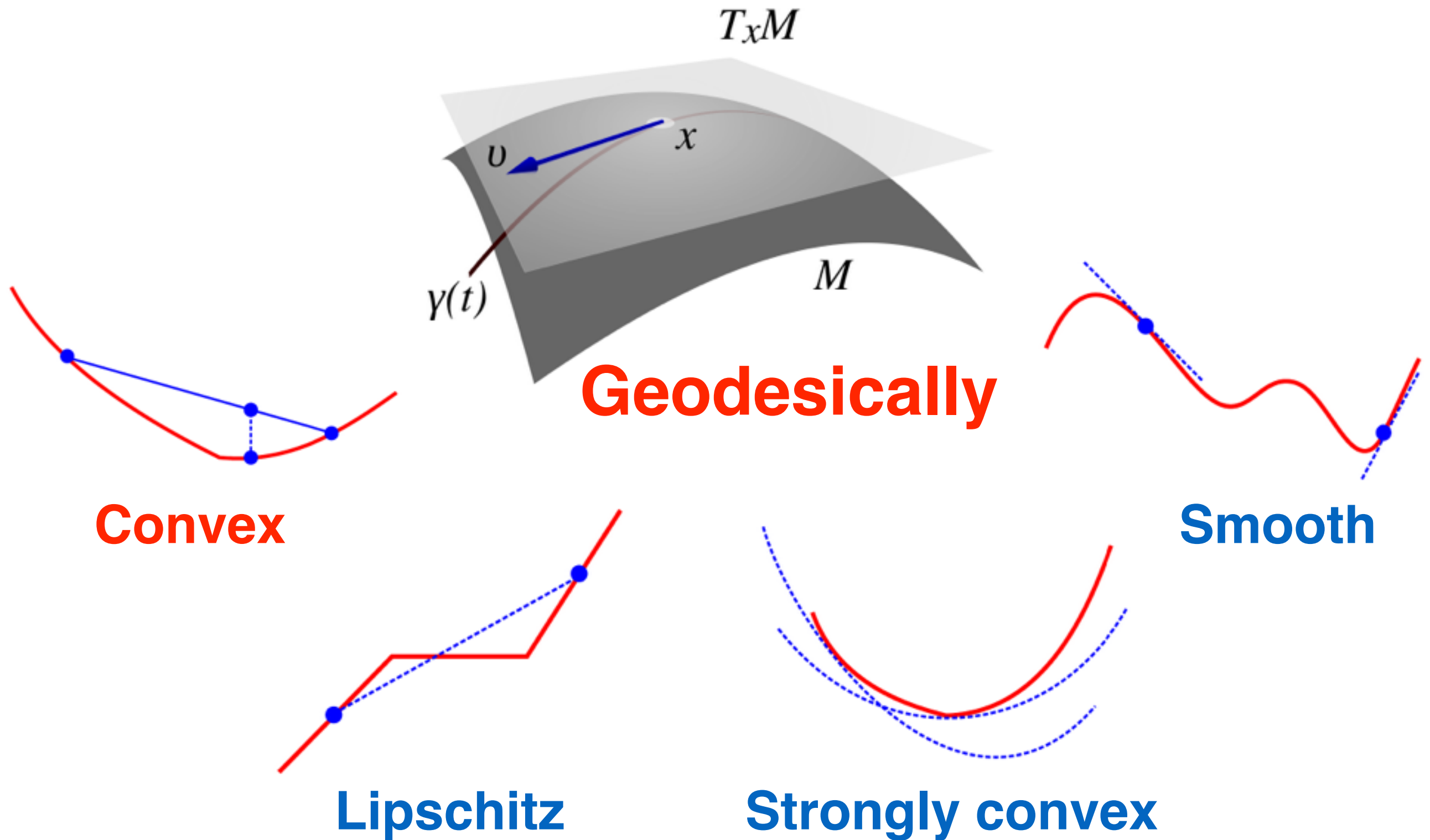


Smooth



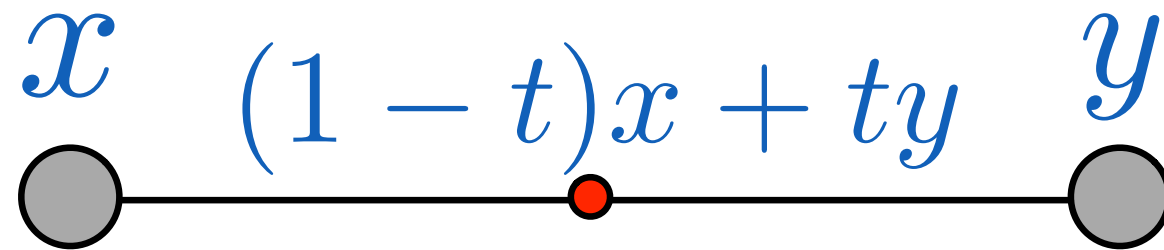
Strongly convex

Classes of function in optimization

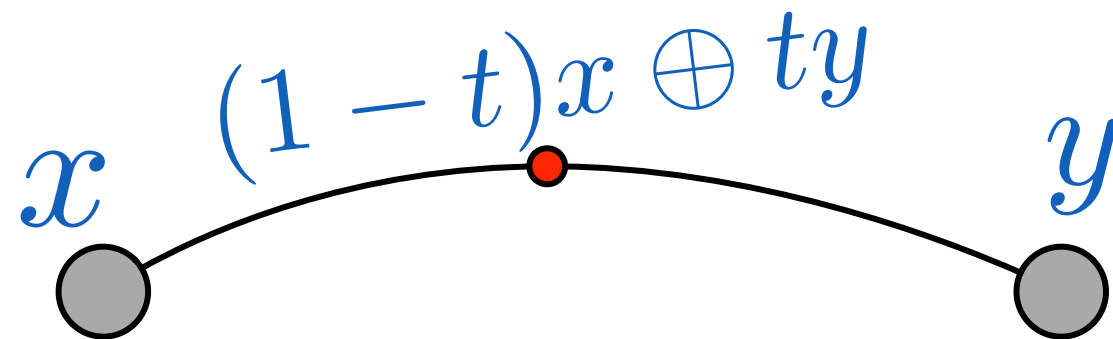


What is geodesic convexity?

Convexity



Geodesic convexity



$$f((1-t)x \oplus ty) \leq (1-t)f(x) + tf(y)$$

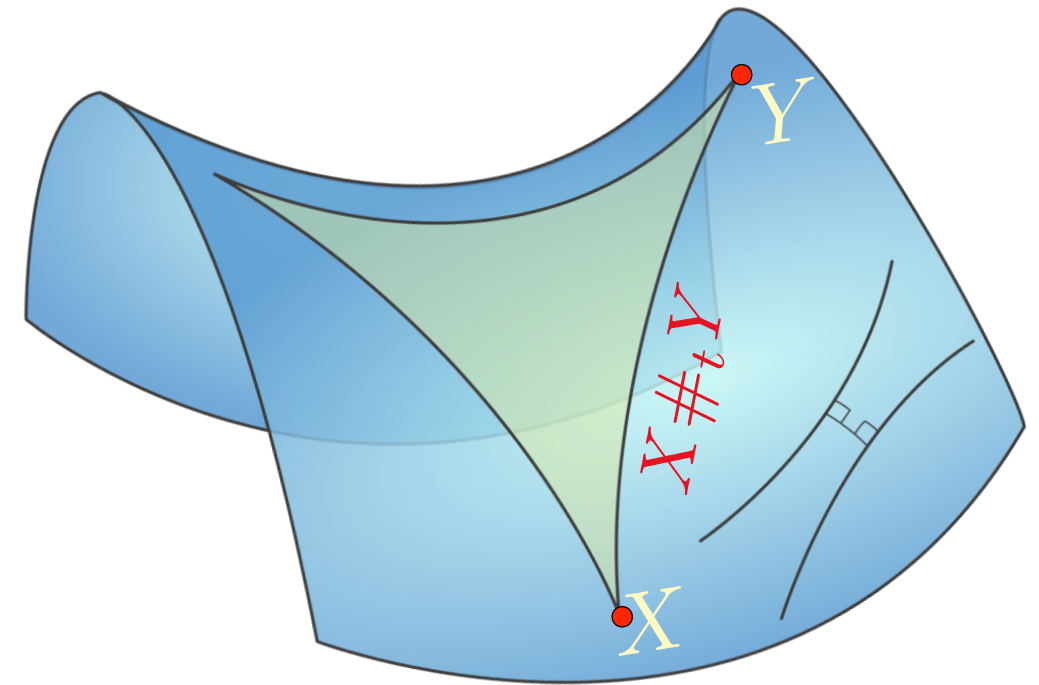
on a Riemannian manifold $f(y) \geq f(x) + \langle g_x, \text{Exp}_x^{-1}(y) \rangle_x$

Metric spaces & curvature: [Menger; Alexandrov; Busemann; Bridson, Häflinger; Gromov; Perelman]

Example: Positive definite matrices

Geodesic

$$\begin{aligned} X \#_t Y &:= X^{\frac{1}{2}} (X^{-\frac{1}{2}} Y X^{-\frac{1}{2}})^t X^{\frac{1}{2}} \\ &= (1-t)X \oplus tY \end{aligned}$$



Examples

$$f(X) = \begin{cases} \log \det(X), & \log \operatorname{tr}(X), \\ \operatorname{tr}(X^\alpha), & \|X^\alpha\|. \end{cases} \quad (\alpha > 0)$$

Exercise: verify for above examples

$$f(X \#_t Y) \leq (1-t)f(X) + tf(Y)$$

PSD manifold and g-convexity

Recognizing, constructing, and optimizing g-convex functions



[Sra, Hosseini (2013,2015)]

- *[Wiesel 2012]*
- *[Rápcsák 1984]*
- *[Udriste 1994]*

Corollaries

$$X \mapsto \log \det(B + \sum_i A_i^* X A_i)$$

$$X \mapsto \log \text{per}(B + \sum_i A_i^* X A_i)$$

$$\delta_R^2(X, Y), \quad \delta_S^2(X, Y)$$

(jointly g-convex)

Many more theorems and corollaries

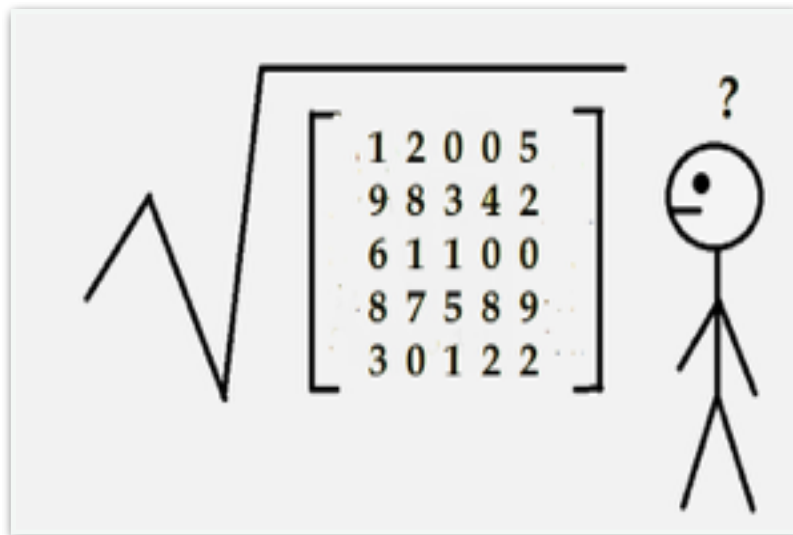
One-D version known as: **Geometric Programming**
www.stanford.edu/~boyd/papers/gp_tutorial.html

[Boyd, Kim, Vandenberghe, Hassibi (2007). 61 pp.]

Examples

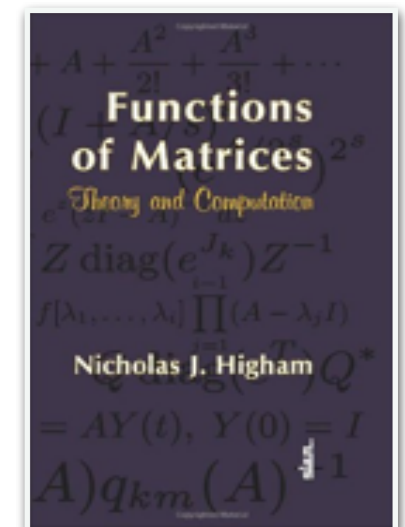
$$X \succcurlyeq 0$$

Matrix square root



Broadly applicable

Key to 'logm'



Matrix square root



Nonconvex optimization through the Euclidean lens

[Jain, Jin, Kakade, Netrapalli; Jul 2015]

$$\min_{X \in \mathbb{R}^{n \times n}} \|M - X^2\|_F^2$$

Gradient descent

$$X_{t+1} \leftarrow X_t - \eta(X_t^2 - M)X_t - \eta X_t(X_t^2 - M)$$

Simple algorithm; linear convergence; **nontrivial** analysis

Matrix square root

Geodesic

$$X \#_t Y := X^{\frac{1}{2}} \left(X^{-\frac{1}{2}} Y X^{-\frac{1}{2}} \right)^t X^{\frac{1}{2}}$$

Midpoint

$$A^{\frac{1}{2}} = A \#_{\frac{1}{2}} I$$

Matrix square root



Nonconvex optimization through **non-Euclidean** lens

[Sra; Jul 2015]

$$\min_{X \succ 0} \quad \delta_S^2(X, A) + \delta_S^2(X, I)$$

Fixed-point iteration

$$X_{k+1} \leftarrow [(X_k + A)^{-1} + (X_k + I)^{-1}]^{-1}$$

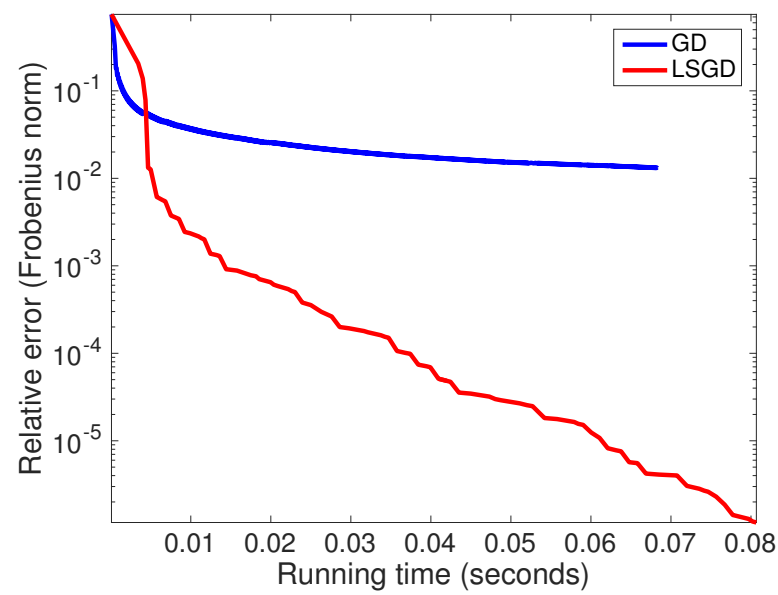
Simple method; linear convergence; 1/2 page analysis!

Global optimality thanks to geodesic convexity

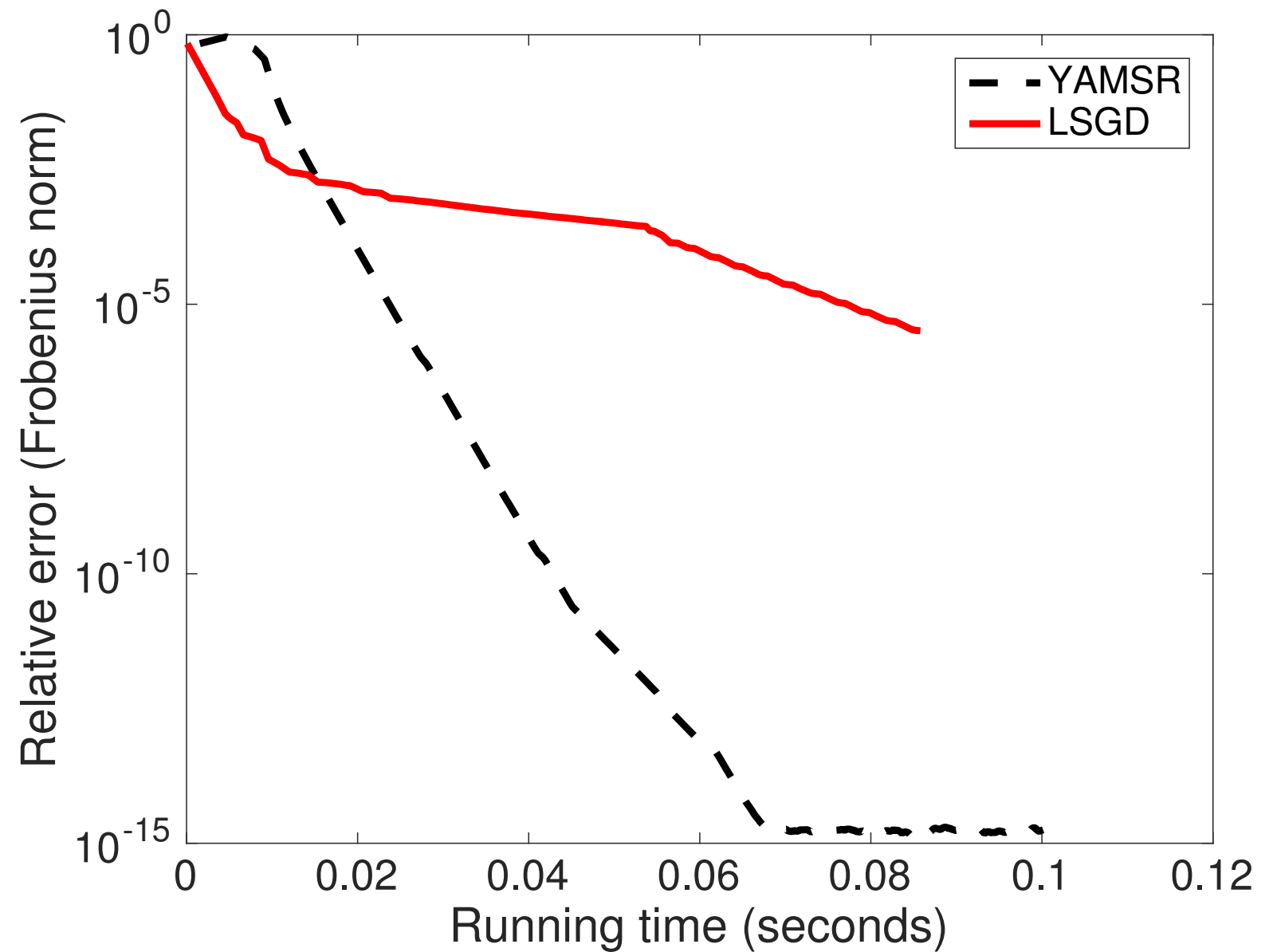
$$\delta_S^2(X, Y) := \frac{1}{2} \log \det \left(\frac{X+Y}{2} \right) - \frac{1}{2} \log \det(XY)$$

(Curiously, this is the Jensen-Shannon divergence between multivariate Gaussians!)

Matrix square root

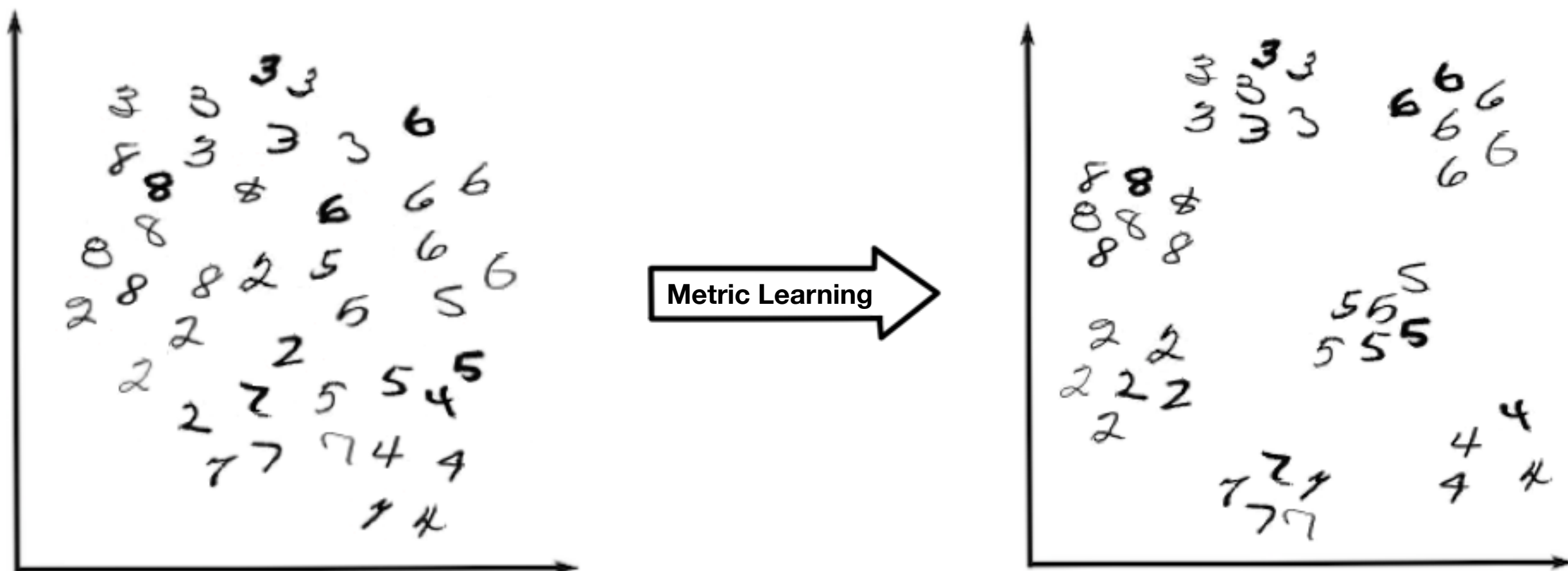


50×50 matrix $I + \beta U U^T$
 $\kappa \approx 64$



Metric learning

What does a metric learning method do?



Euclidean metric learning

Pairwise constraints

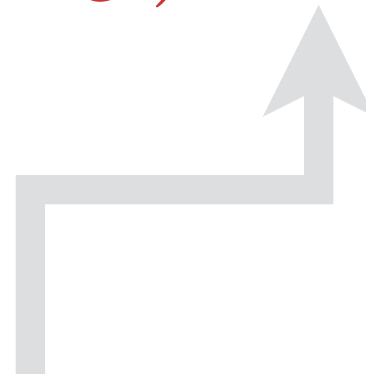
$\mathcal{S} := \{(\mathbf{x}_i, \mathbf{x}_j) \mid \mathbf{x}_i \text{ and } \mathbf{x}_j \text{ are in the same class}\}$

$\mathcal{D} := \{(\mathbf{x}_i, \mathbf{x}_j) \mid \mathbf{x}_i \text{ and } \mathbf{x}_j \text{ are in different classes}\}$

Goal

given pairwise constraints learn Mahalanobis distance

$$d_{\mathbf{A}}(\mathbf{x}, \mathbf{y}) := (\mathbf{x} - \mathbf{y})^T \mathbf{A} (\mathbf{x} - \mathbf{y})$$



Positive definite matrix $\mathbf{A}$

Metric learning methods

MMC

[Xing, Jordan, Russell, Ng 2002]

LMNN

[Weinberger, Saul 2005]

ITML

[Davis, Kulis, Jain, Sra, Dhillon 2007]

tons of other methods!

$$\min_{\mathbf{A} \succeq 0} \sum_{(\mathbf{x}_i, \mathbf{x}_j) \in \mathcal{S}} d_{\mathbf{A}}(\mathbf{x}_i, \mathbf{x}_j)$$

$$\text{such that } \sum_{(\mathbf{x}_i, \mathbf{x}_j) \in \mathcal{D}} \sqrt{d_{\mathbf{A}}(\mathbf{x}_i, \mathbf{x}_j)} \geq 1$$

$$\min_{\mathbf{A} \succeq 0} \sum_{(\mathbf{x}_i, \mathbf{x}_j) \in \mathcal{S}} \left[(1 - \mu) d_{\mathbf{A}}(\mathbf{x}_i, \mathbf{x}_j) + \mu \sum_l (1 - y_{il}) \xi_{ijl} \right]$$

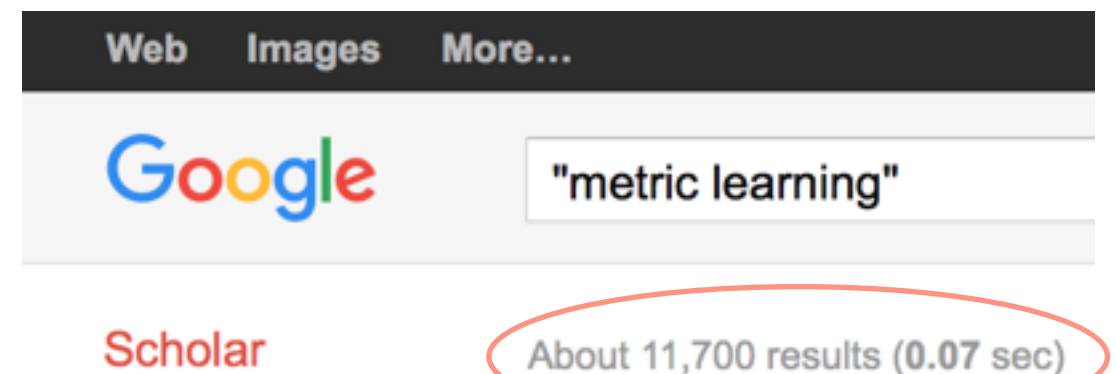
$$d_{\mathbf{A}}(\mathbf{x}_i, \mathbf{x}_l) - d_{\mathbf{A}}(\mathbf{x}_i, \mathbf{x}_j) \geq 1 - \xi_{ijl}$$

$$\xi_{ijl} \geq 0$$

$$\min_{\mathbf{A} \succeq 0} D_{\text{ld}}(\mathbf{A}, \mathbf{A}_0)$$

$$\text{such that } \begin{aligned} d_{\mathbf{A}}(\mathbf{x}, \mathbf{y}) &\leq u, & (\mathbf{x}, \mathbf{y}) \in \mathcal{S}, \\ d_{\mathbf{A}}(\mathbf{x}, \mathbf{y}) &\geq l, & (\mathbf{x}, \mathbf{y}) \in \mathcal{D} \end{aligned}$$

$$D_{\text{ld}}(\mathbf{A}, \mathbf{A}_0) := \text{tr}(\mathbf{A} \mathbf{A}_0^{-1}) - \log \det(\mathbf{A} \mathbf{A}_0^{-1}) - d$$



A simple new way for metric learning

Euclidean idea

$$\min_{\mathbf{A} \succeq 0} \sum_{(\mathbf{x}_i, \mathbf{x}_j) \in \mathcal{S}} d_{\mathbf{A}}(\mathbf{x}_i, \mathbf{x}_j) - \lambda \sum_{(\mathbf{x}_i, \mathbf{x}_j) \in \mathcal{D}} d_{\mathbf{A}}(\mathbf{x}_i, \mathbf{x}_j)$$

New idea

$$\min_{\mathbf{A} \succeq 0} \sum_{(\mathbf{x}_i, \mathbf{x}_j) \in \mathcal{S}} d_{\mathbf{A}}(\mathbf{x}_i, \mathbf{x}_j) + \sum_{(\mathbf{x}_i, \mathbf{x}_j) \in \mathcal{D}} d_{\mathbf{A}^{-1}}(\mathbf{x}_i, \mathbf{x}_j)$$

Equivalently solve

$$\min_{\mathbf{A} \succ 0} h(\mathbf{A}) := \text{tr}(\mathbf{A}\mathbf{S}) + \text{tr}(\mathbf{A}^{-1}\mathbf{D})$$

$$\mathbf{S} := \sum_{(\mathbf{x}_i, \mathbf{x}_j) \in \mathcal{S}} (\mathbf{x}_i - \mathbf{x}_j)(\mathbf{x}_i - \mathbf{x}_j)^T,$$

$$\mathbf{D} := \sum_{(\mathbf{x}_i, \mathbf{x}_j) \in \mathcal{D}} (\mathbf{x}_i - \mathbf{x}_j)(\mathbf{x}_i - \mathbf{x}_j)^T$$

[Habibzadeh, Hosseini, Sra, ICML 2016]

A simple new way for metric learning

$$X \#_t Y := X^{\frac{1}{2}} (X^{-\frac{1}{2}} Y X^{-\frac{1}{2}})^t X^{\frac{1}{2}}$$

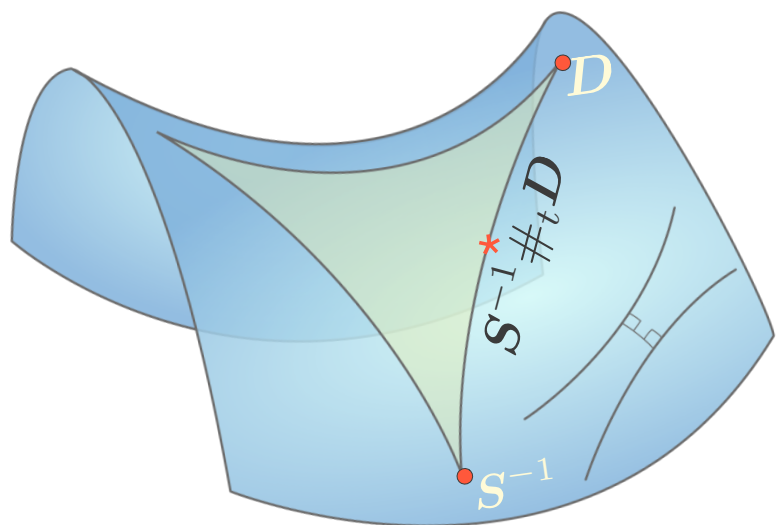
Closed form solution!

$$\nabla h(\mathbf{A}) = 0 \quad \Leftrightarrow \quad \mathbf{S} - \mathbf{A}^{-1} \mathbf{D} \mathbf{A}^{-1} = 0$$

$$\mathbf{A} = \mathbf{S}^{-1} \#_{\frac{1}{2}} \mathbf{D}$$

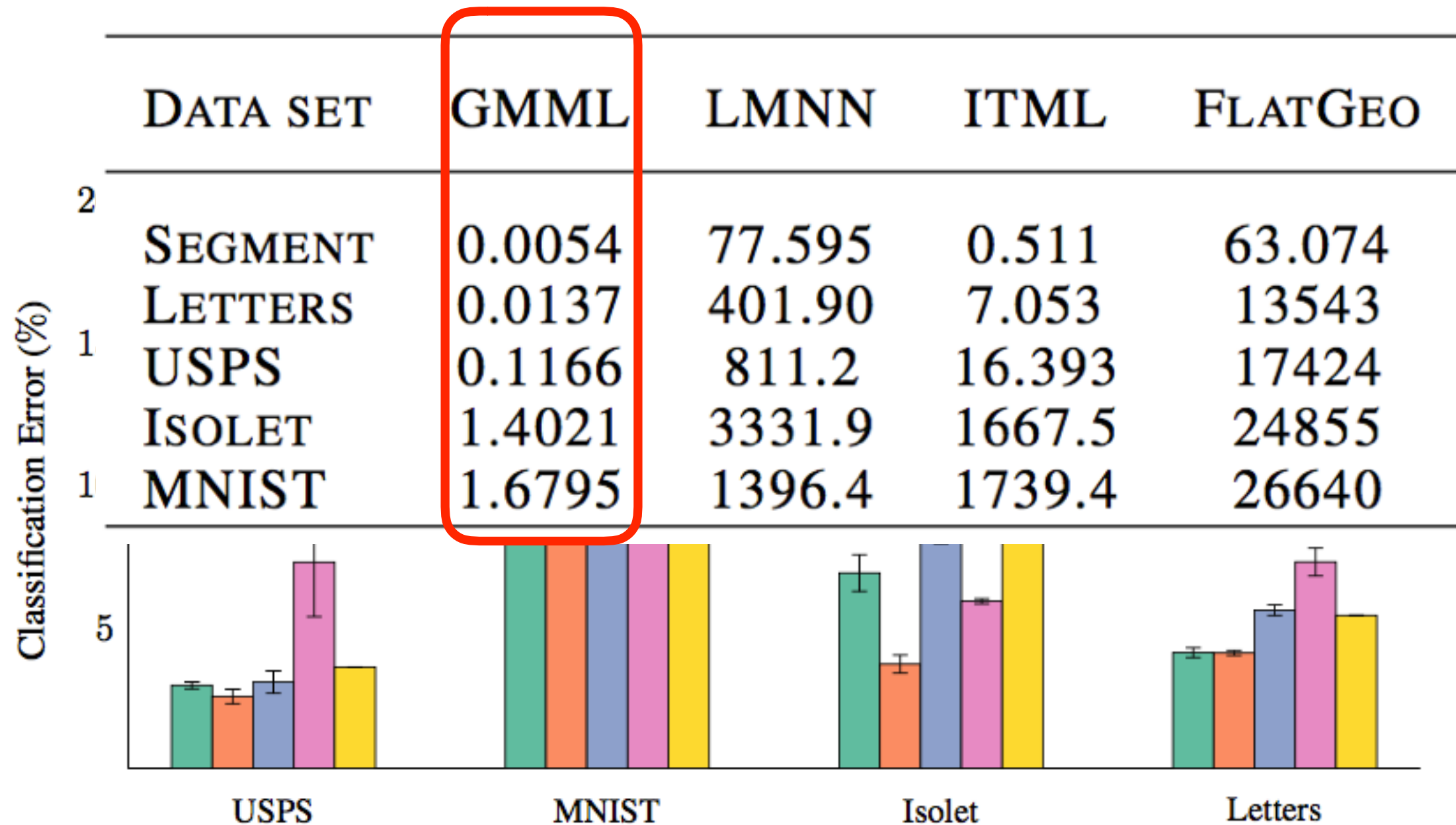
More generally

$$\min_{\mathbf{A} \succ 0} (1 - t) \delta_R^2(\mathbf{S}^{-1}, \mathbf{A}) + t \delta_R^2(\mathbf{D}, \mathbf{A})$$



$$\mathbf{S}^{-1} \#_t \mathbf{D}$$

Experiments



[Habibzadeh, Hosseini, Sra ICML 2016]

Brascamp-Lieb Constant

$$\int_{\mathbb{R}^n} \prod_{i=1}^m f_i(B_i x)^{p_i} dx \leq D^{-1/2} \prod_{i=1}^m \left(\int_{\mathbb{R}^{n_i}} f_i(y) dy \right)^{p_i}$$

$$D := \inf \left\{ \frac{\det(\sum_i p_i B_i^* X_i B_i)}{\prod_i (\det X_i)^{p_i}} \mid X_i \succ 0, n_i \times n_i, \right\}$$

$$p_i > 0, f_i \geq 0 \quad \sum_{i=1}^m p_i n_i = n$$

powerful inequality; includes Hölder, Loomis-Whitney, Young's, many others!

Brascamp-Lieb constant

$$\min_{X_1, \dots, X_m \succ 0} \log \det \left(\sum_i p_i B_i^* X_i B_i \right) - \sum_i p_i \log \det X_i$$

- Arises in geometric complexity theory:
[Garg, Gurvits, Oliveira, Wigderson; Jul 2016]

Exercise

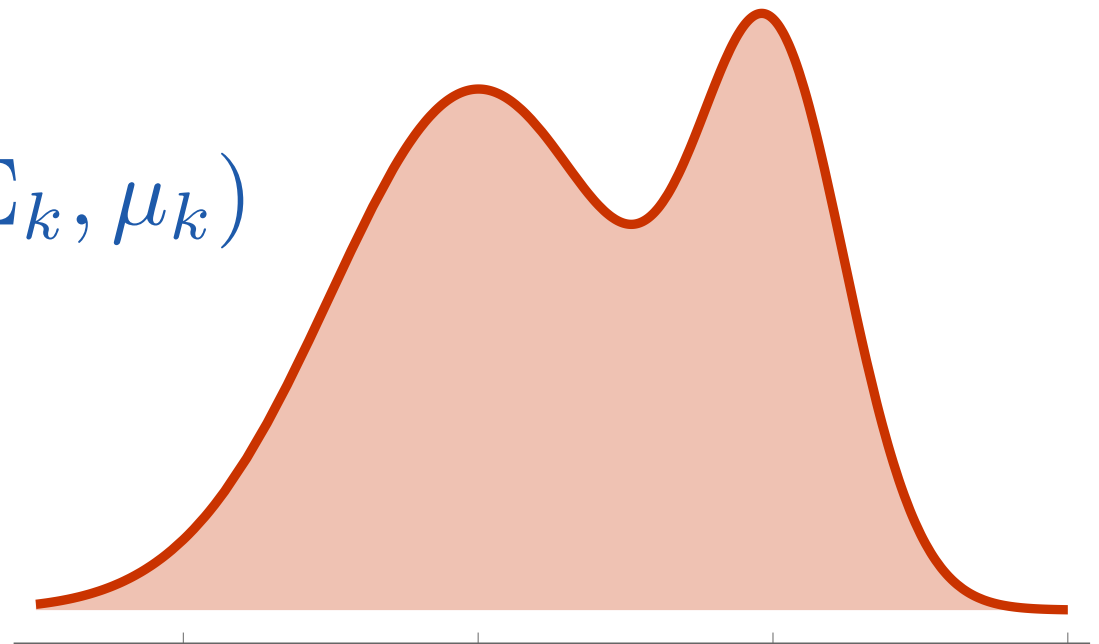
Prove this is a g-convex opt problem

- G-convexity yields transparent algorithms & complexity analysis for global optimum.

Gaussian mixture models

$$p_{\text{mix}}(x) := \sum_{k=1}^K \pi_k p_{\mathcal{N}}(x; \Sigma_k, \mu_k)$$

$$\max \prod_i p_{\text{mix}}(x_i)$$



Expectation maximization (EM): default choice

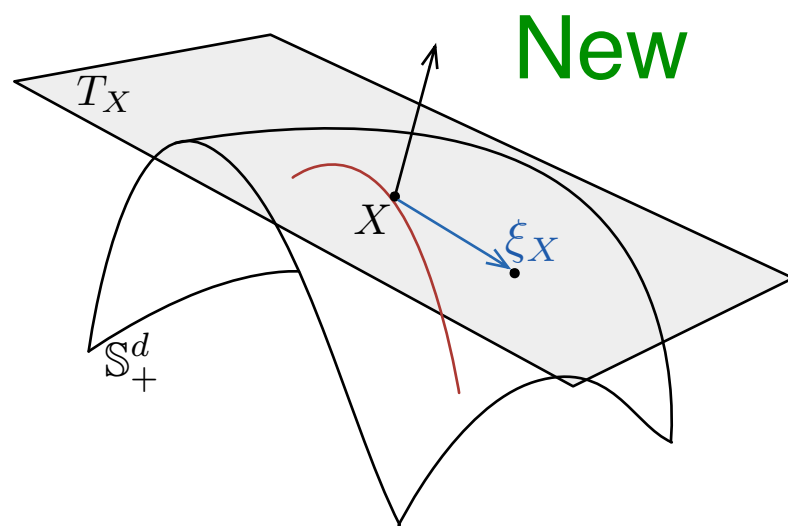
$$p_{\mathcal{N}}(x; \Sigma, \mu) \propto \frac{1}{\sqrt{\det(\Sigma)}} \exp \left(-\frac{1}{2} (x - \mu)^T \Sigma^{-1} (x - \mu) \right)$$

Gaussian mixture models

- **Nonconvex** – difficult, possibly several local optima
- **GMMs** – Recent surge of theoretical results
- **In Practice** – EM still default choice

Difficulty: Positive definiteness constraint on Σ_k

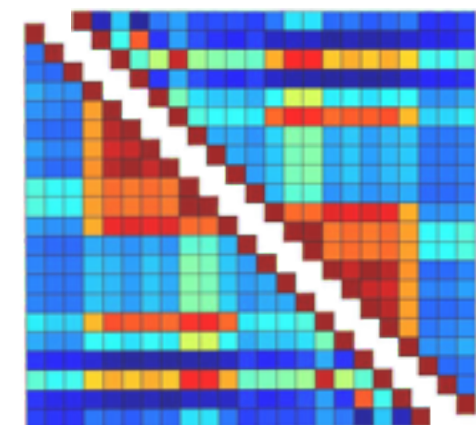
Geometric opt



Unconstrained, Cholesky

Folklore

LL^T



[Hosseini, Sra NIPS 2015]

Failure of “obvious” LL^T

sep.	EM	CG- LL^T
0.2	52s // 12.7	614s // 12.7
1	160s // 13.4	435s // 13.5
5	72s // 12.8	426s // 12.8

$$\|\mu_i - \mu_j\| \geq \text{sep} \max_{ij} \{\text{tr} \Sigma_i, \text{tr} \Sigma_j\}$$

$d=20$
simulation

Failure of Riemannian optimization

K	EM	Riem-CG
2	17s // 29.28	947s // 29.28
5	202s // 32.07	5262s // 32.07
10	2159s // 33.05	17712s // 33.03



manopt.org

Riemannian opt. toolbox

*d=35
images
dataset*

What's wrong?



log-likelihood for one component

$$-\frac{n}{2} \log \det \Sigma - \frac{1}{2} \sum_{i=1}^n (x_i - \mu)^T \Sigma^{-1} (x_i - \mu)$$

Euclidean convex problem
Not geodesically convex



Reformulate as g-convex

$$y_i = \begin{bmatrix} x_i \\ 1 \end{bmatrix} \quad S = \begin{bmatrix} \Sigma + \mu\mu^T & \mu \\ \mu^T & 1 \end{bmatrix}$$

$$\max_{S \succ 0} \hat{\mathcal{L}}(S) := \sum_{i=1}^n \log q_{\mathcal{N}}(y_i; S),$$

Thm. The modified log-likelihood is g-convex. Local max of modified mixture LL is local max of original.

Success of geometric optimization

K	EM	Riem-CG	L-RBFGS
2	17s // 29.28	18s // 29.28	14s // 29.28
5	202s // 32.07	140s // 32.07	117s // 32.07
10	2159s // 33.05	1048s // 33.06	658s // 33.06

Riem-CG (manopt) savings:

947 → **18**; 5262 → **140**; 17712 → **1048**

*d=35
images
dataset*



github.com/utvisionlab/mixest

First-order algorithms

[Zhang, Sra, COLT 2016]

first-order g-convex optimization

$$\min_{x \in \mathcal{X} \subset \mathcal{M}} f(x)$$

$\mathcal{X}$ g-convex set; f g-convex func; $\mathcal{M}$ Riemannian manifold



oracle access to exact or stochastic (sub)gradients

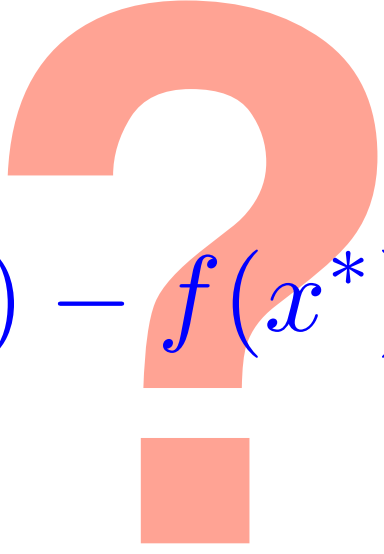
$$x \leftarrow \text{Exp}_x(-\eta \nabla f(x))$$

analog to: $x \leftarrow x - \eta \nabla f(x)$

In particular, we study the **global complexity** of **first-order g-convex optimization**

Global Complexity

Gradient Descent
Stochastic Gradient Descent
Coordinate Descent
Accelerated Gradient Descent
Fast Incremental Gradient
... ..

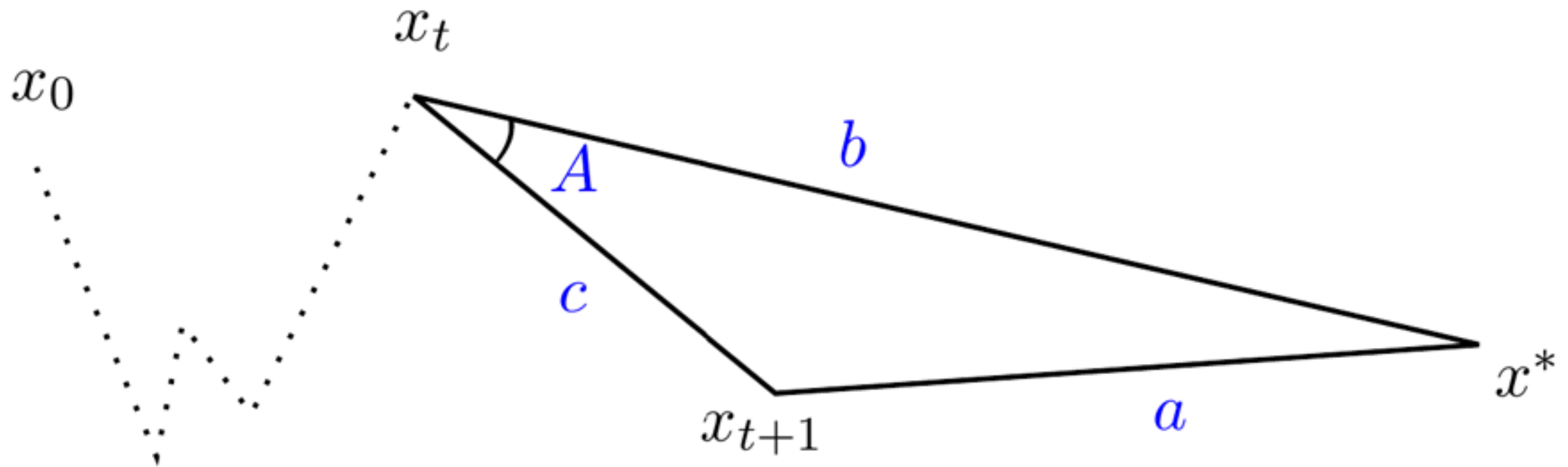

$$\mathbb{E}[f(x_a) - f(x^*)] \leq ?$$

Convex Optimization

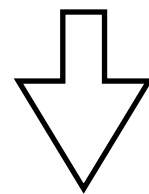
G-Convex Optimization

The Euclidean **law of cosines** is essential to bound $d^2(x_{t+1}, x^*)$ in analysis of usual convex opt. methods

$$x_{t+1} = x_t - \eta_t g_t$$



$$a^2 = b^2 + c^2 - 2bc \cos(A)$$



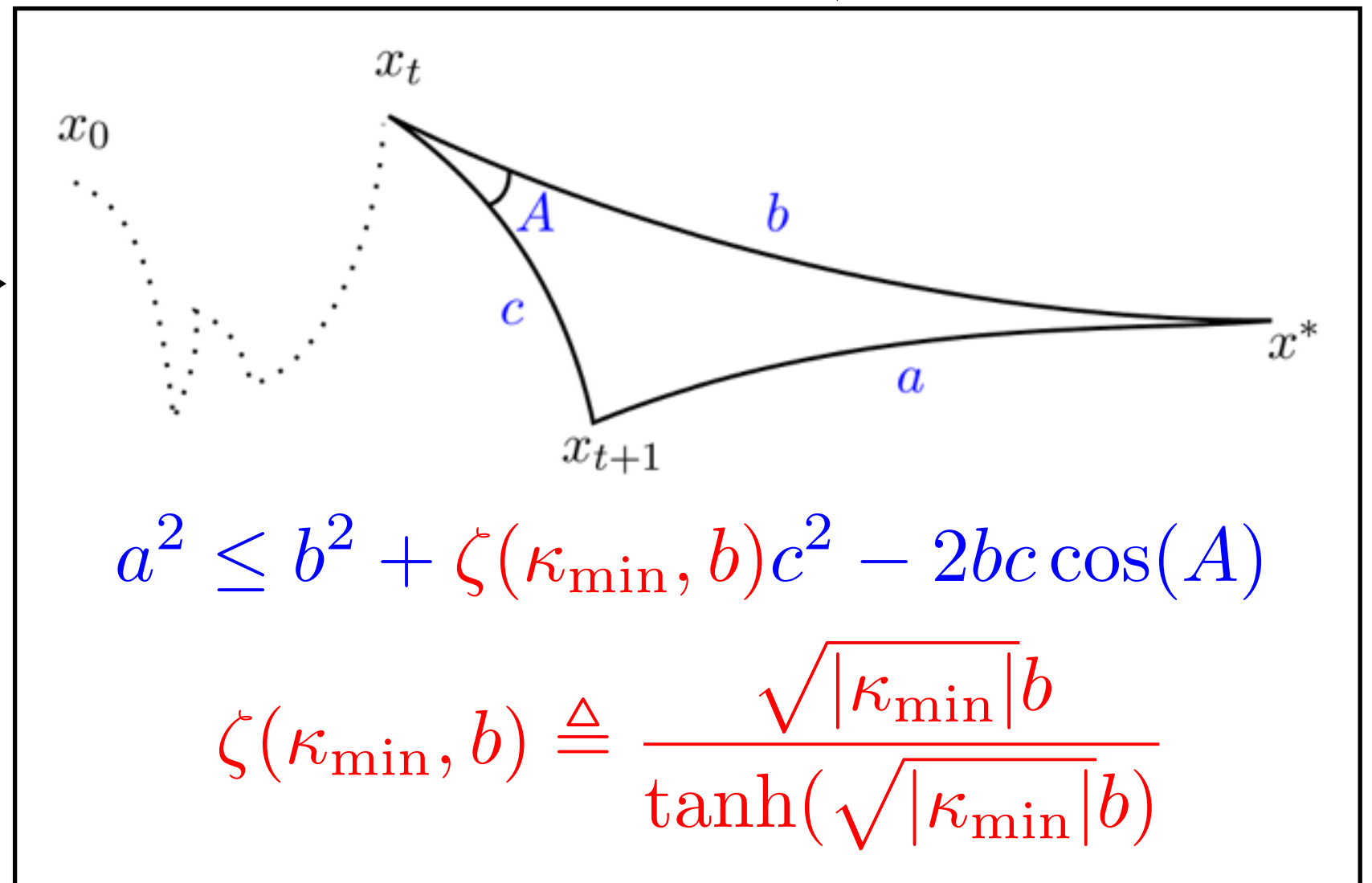
$$\|x_{t+1} - x^*\|^2 = \|x_t - x^*\|^2 + \eta_t^2 \|g_t\|^2 - 2\eta_t \langle g_t, x_t - x^* \rangle$$

We develop a corresponding **inequality** to bound $d^2(x_{t+1}, x^*)$ on manifolds

$$\cosh(-\kappa a) = \cosh(-\kappa b) \cosh(-\kappa c) + \sinh(-\kappa b) \sinh(-\kappa c) \cos(A)$$

Toponogov's theorem

Grönwall's inequality



Rediscovered and generalizes result of:
C.-E. Raoult, McCann, Schmuckenschläger,
Invent. Math. (2001).

Convergence rates depend on **lower bounds**
on the **sectional curvature**

(Sub)gradient

Lipschitz

**Strongly convex
/ smooth**

**Strongly convex
& smooth**

convex

g-convex

$$O\left(\sqrt{\frac{1}{t}}\right)$$

$$O\left(\sqrt{\frac{\zeta_{\max}}{t}}\right)$$

$$O\left(\frac{1}{t}\right)$$

$$O\left(\frac{\zeta_{\max}}{t}\right)$$

$$O\left(\left(1 - \frac{\mu}{L_g}\right)^t\right)$$

$$O\left(\left(1 - \min\left\{\frac{1}{\zeta_{\max}}, \frac{\mu}{L_g}\right\}\right)^t\right)$$

**Stochastic
(sub)gradient**

... ..

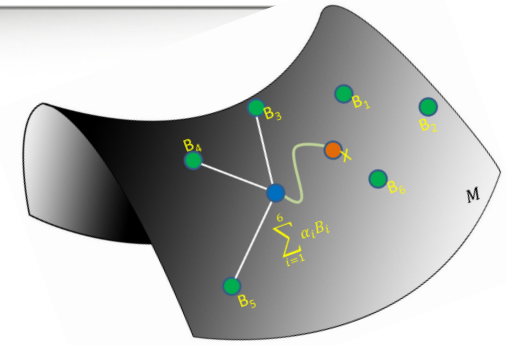
$$\zeta_{\max} \triangleq \frac{\sqrt{|\kappa_{\min}|} D}{\tanh\left(\sqrt{|\kappa_{\min}|} D\right)}$$

See paper for other interesting results *[Zhang, Sra, COLT 2016]*

Is that all?

Riemannian finite-sum problems

$$\min_{x \in \mathcal{M}} f(x) := \frac{1}{n} \sum_{i=1}^n f_i(x)$$



- $\mathcal{M}$ is a Riemannian manifold
- g-convex and g-nonconvex 'f' allowed
- First global complexity results for stochastic methods on Riemannian manifolds
- Riemannian SVRG

[Zhang, Reddi, Sra, NIPS 2016]

Stochastic eigenvector computation

$$\min_{x^T x = 1} -x^T \left(\sum_{i=1}^n z_i z_i^T \right) x$$

compute dominant eigenvector

Online PCA, stochastic power method, SVRG accelerated, etc.

[Garber, Hazan 2015; Jin, Kakade, Musco, Netrapalli, Sidford 2015; Shamir 2015, 2016]

Key insight:

A non-convex problem on sphere; “satisfies” manifold version of Polyak-Łojasiewicz inequality; RSVRG analysis delivers simple proof of linear convergence

[Zhang, Reddi, Sra, NIPS 2016]

Invitation for more!

- Faster stochastic optimization
- Beyond Riemannian manifolds
- Faster gradient based algorithms
- Handling more complex constraints
- Lower bound theory
- ...

Thanks!