

DERIVATIVE FREE OPTIMIZATION

CLASS 6

2025/2026

---

## Problem Statement

### Stochastic search algorithms - basics

#### Adaptive Evolution Strategies

- Mean Vector Adaptation

- Step-size control

- Covariance Matrix Adaptation

  - Rank-One Update

  - Cumulation—the Evolution Path

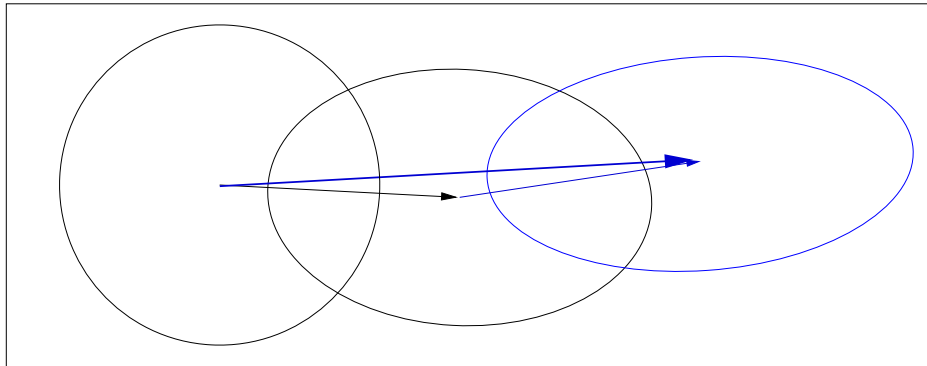
  - Rank- $\mu$  Update

# Cumulation

## The Evolution Path

### Evolution Path

Conceptually, the evolution path is the **search path** the strategy takes **over a number of iteration steps**. It can be expressed as a sum of consecutive *steps* of the mean ***m***.



An exponentially weighted sum of steps  $y_w$  is used

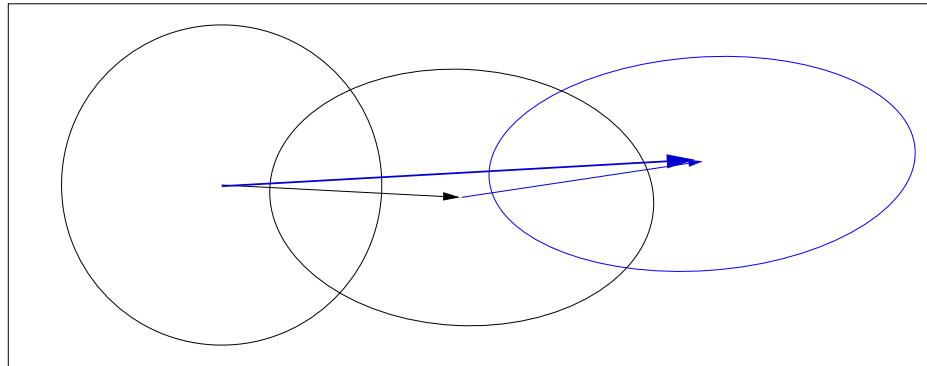
$$p_c \propto \sum_{i=0}^g \underbrace{(1 - c_c)^{g-i}}_{\text{exponentially fading weights}} y_w^{(i)}$$

# Cumulation

## The Evolution Path

### Evolution Path

Conceptually, the evolution path is the **search path** the strategy takes **over a number of iteration steps**. It can be expressed as a sum of consecutive *steps* of the mean ***m***.



An exponentially weighted sum of steps  $y_w$  is used

$$p_c \propto \sum_{i=0}^g \underbrace{(1 - c_c)^{g-i}}_{\text{exponentially fading weights}} y_w^{(i)}$$

The recursive construction of the evolution path (cumulation):

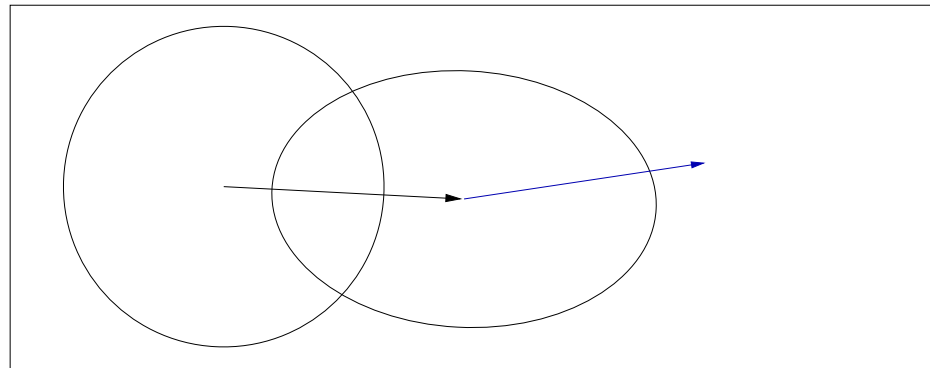
$$p_c \leftarrow \underbrace{(1 - c_c)}_{\text{decay factor}} p_c + \underbrace{\sqrt{1 - (1 - c_c)^2} \sqrt{\mu_w}}_{\text{normalization factor}} \underbrace{y_w}_{\text{input} = \frac{m - m_{\text{old}}}{\sigma}}$$

where  $\mu_w = \frac{1}{\sum w_i^2}$ ,  $c_c \ll 1$ . **History information** is accumulated in the evolution path.

# Cumulation

## Utilizing the Evolution Path

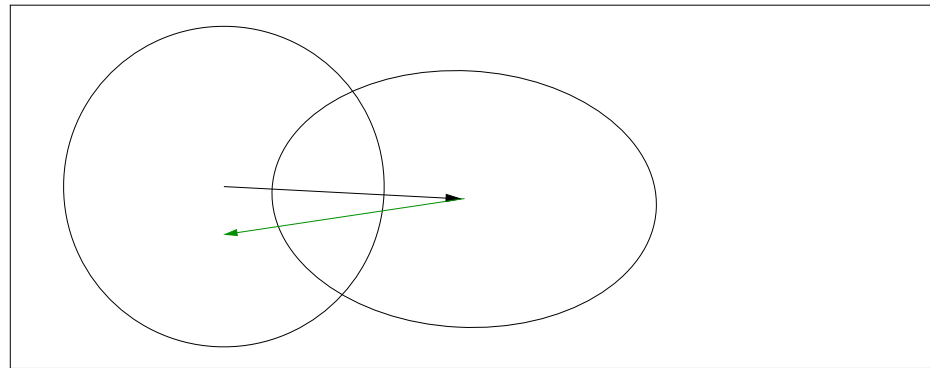
We used  $\mathbf{y}_w \mathbf{y}_w^T$  for updating  $\mathbf{C}$ . Because  $\mathbf{y}_w \mathbf{y}_w^T = -\mathbf{y}_w (-\mathbf{y}_w)^T$  the sign of  $\mathbf{y}_w$  is lost.



# Cumulation

## Utilizing the Evolution Path

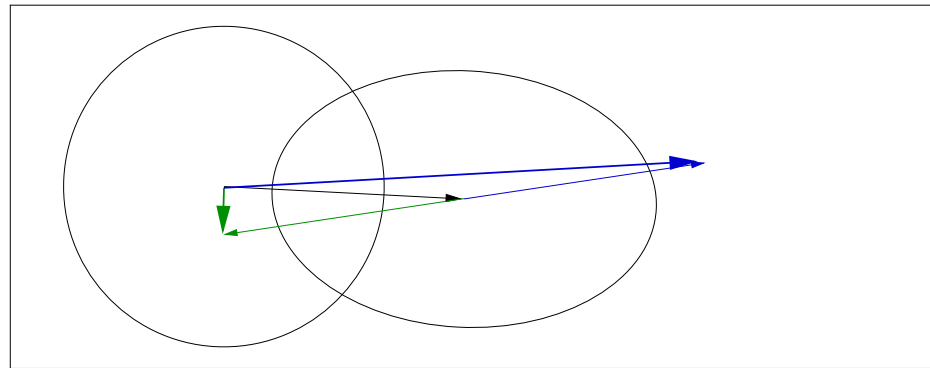
We used  $\mathbf{y}_w \mathbf{y}_w^T$  for updating **C**. Because  $\mathbf{y}_w \mathbf{y}_w^T = -\mathbf{y}_w (-\mathbf{y}_w)^T$  the sign of  $\mathbf{y}_w$  is lost.



# Cumulation

## Utilizing the Evolution Path

We used  $\mathbf{y}_w \mathbf{y}_w^T$  for updating  $\mathbf{C}$ . Because  $\mathbf{y}_w \mathbf{y}_w^T = -\mathbf{y}_w (-\mathbf{y}_w)^T$  the sign of  $\mathbf{y}_w$  is lost.



The sign information is (re-)introduced by using the *evolution path*.

$$\begin{aligned} \mathbf{p}_c &\leftarrow \underbrace{(1 - c_c)}_{\text{decay factor}} \mathbf{p}_c + \underbrace{\sqrt{1 - (1 - c_c)^2} \sqrt{\mu_w}}_{\text{normalization factor}} \mathbf{y}_w \\ \mathbf{C} &\leftarrow (1 - c_{\text{cov}}) \mathbf{C} + c_{\text{cov}} \underbrace{\mathbf{p}_c \mathbf{p}_c^T}_{\text{rank-one}} \end{aligned}$$

where  $\mu_w = \frac{1}{\sum w_i^2}$ ,  $c_c \ll 1$ .

Choice of normalizing constant for the path  $p_c$ .

$$p_c^{t+1} = (1 - c_c) p_c^t + \underbrace{\sqrt{1 - (1 - c_c)^2}}_{\alpha} \sqrt{\mu_w} y_w^{t+1}$$

The constant

is set such

that if  $f(x) = \text{rand}$  (i.e. each time we call  $f$  we get a random and independent number (say uniform))

$$\mu_w = \frac{1}{\sum w_i^2}$$

Proposition: Assume that  $p_c^t \sim \mathcal{N}(0, c_t)$ , if  $f$  is random then

$$p_c^{t+1} \sim \mathcal{N}(0, c_t)$$

because of the choice of the normalizing constant  $\alpha$ .



Using an **evolution path** for the **rank-one update** of the covariance matrix reduces the number of function evaluations to adapt to a straight ridge **from  $\mathcal{O}(n^2)$  to  $\mathcal{O}(n)$** .<sup>(3)</sup>

The overall model complexity is  $n^2$  but important parts of the model can be learned in time of order  $n$

---

<sup>3</sup> Hansen, Müller and Koumoutsakos 2003. Reducing the Time Complexity of the Derandomized Evolution Strategy with Covariance Matrix Adaptation (CMA-ES). *Evolutionary Computation*, 11(1), pp. 1-18

# Rank- $\mu$ Update

$$\begin{aligned} \mathbf{x}_i &= \mathbf{m} + \sigma \mathbf{y}_i, & \mathbf{y}_i &\sim \mathcal{N}_i(\mathbf{0}, \mathbf{C}), \\ \mathbf{m} &\leftarrow \mathbf{m} + \sigma \mathbf{y}_w & \mathbf{y}_w &= \sum_{i=1}^{\mu} w_i \mathbf{y}_{i:\lambda} \end{aligned}$$

The rank- $\mu$  update extends the update rule for large population sizes  $\lambda$  using  $\mu > 1$  vectors to update  $\mathbf{C}$  at each iteration step.

# Rank- $\mu$ Update

$$\begin{aligned} \mathbf{x}_i &= \mathbf{m} + \sigma \mathbf{y}_i, & \mathbf{y}_i &\sim \mathcal{N}_i(\mathbf{0}, \mathbf{C}), \\ \mathbf{m} &\leftarrow \mathbf{m} + \sigma \mathbf{y}_w & \mathbf{y}_w &= \sum_{i=1}^{\mu} w_i \mathbf{y}_{i:\lambda} \end{aligned}$$

The rank- $\mu$  update extends the update rule for large population sizes  $\lambda$  using  $\mu > 1$  vectors to update  $\mathbf{C}$  at each iteration step.

The matrix

$$\mathbf{C}_{\mu} = \sum_{i=1}^{\mu} w_i \mathbf{y}_{i:\lambda} \mathbf{y}_{i:\lambda}^T$$

computes a weighted mean of the outer products of the best  $\mu$  steps and has rank  $\min(\mu, n)$  with probability one.

# Rank- $\mu$ Update

$$\begin{aligned} \mathbf{x}_i &= \mathbf{m} + \sigma \mathbf{y}_i, & \mathbf{y}_i &\sim \mathcal{N}_i(\mathbf{0}, \mathbf{C}), \\ \mathbf{m} &\leftarrow \mathbf{m} + \sigma \mathbf{y}_w & \mathbf{y}_w &= \sum_{i=1}^{\mu} w_i \mathbf{y}_{i:\lambda} \end{aligned}$$

The rank- $\mu$  update extends the update rule for large population sizes  $\lambda$  using  $\mu > 1$  vectors to update  $\mathbf{C}$  at each iteration step.

The matrix

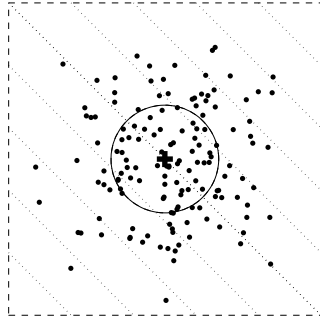
$$\mathbf{C}_{\mu} = \sum_{i=1}^{\mu} w_i \mathbf{y}_{i:\lambda} \mathbf{y}_{i:\lambda}^T$$

computes a weighted mean of the outer products of the best  $\mu$  steps and has rank  $\min(\mu, n)$  with probability one.

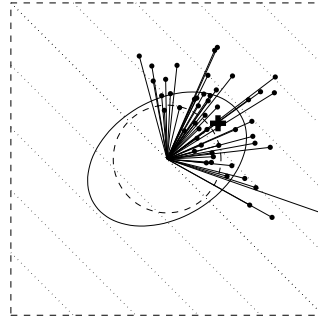
The rank- $\mu$  update then reads

$$\mathbf{C} \leftarrow (1 - c_{\text{cov}}) \mathbf{C} + c_{\text{cov}} \mathbf{C}_{\mu}$$

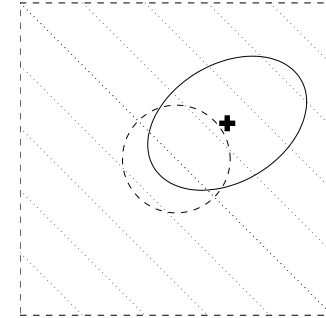
where  $c_{\text{cov}} \approx \mu_w / n^2$  and  $c_{\text{cov}} \leq 1$ .



$$\mathbf{x}_i = \mathbf{m} + \sigma \mathbf{y}_i, \quad \mathbf{y}_i \sim \mathcal{N}(\mathbf{0}, \mathbf{C})$$



$$\begin{aligned} \mathbf{C}_\mu &= \frac{1}{\mu} \sum \mathbf{y}_{i:\lambda} \mathbf{y}_{i:\lambda}^T \\ \mathbf{C} &\leftarrow \frac{1}{(1 - 1/\mu) \times \mathbf{C} + 1/\mu \times \mathbf{C}_\mu} \end{aligned}$$



$$\mathbf{m}_{\text{new}} \leftarrow \mathbf{m} + \frac{1}{\mu} \sum \mathbf{y}_{i:\lambda}$$

sampling of  
 $\lambda = 150$  solutions  
 where  $\mathbf{C} = \mathbf{I}$  and  
 $\sigma = 1$

calculating  $\mathbf{C}$  where  
 $\mu = 50$ ,  $w_1 = \dots =$   
 $w_\mu = \frac{1}{\mu}$ , and  
 $\mathbf{C}_{\text{cov}} = \mathbf{I}$

new distribution

## The rank- $\mu$ update

- ▶ increases the possible learning rate in large populations  
roughly from  $2/n^2$  to  $\mu_w/n^2$
- ▶ can reduce the number of necessary iterations roughly from  $\mathcal{O}(n^2)$  to  $\mathcal{O}(n)$  <sup>(4)</sup>  
given  $\mu_w \propto \lambda \propto n$

Therefore the rank- $\mu$  update is the primary mechanism whenever a large population size is used

say  $\lambda \geq 3n + 10$

---

<sup>4</sup> Hansen, Müller, and Koumoutsakos 2003. Reducing the Time Complexity of the Derandomized Evolution Strategy with Covariance Matrix Adaptation (CMA-ES). *Evolutionary Computation*, 11(1), pp. 1-18

$\lambda$  larger than default is typically good if:

→  $f$  is multimodal

→ we can easily have access to different CPUs → we can parallelize the  $\lambda$  evaluation

## The rank- $\mu$ update

- ▶ increases the possible learning rate in large populations  
roughly from  $2/n^2$  to  $\mu_w/n^2$
- ▶ can reduce the number of necessary iterations roughly from  $\mathcal{O}(n^2)$  to  $\mathcal{O}(n)$  <sup>(4)</sup>  
given  $\mu_w \propto \lambda \propto n$

Therefore the rank- $\mu$  update is the primary mechanism whenever a large population size is used

say  $\lambda \geq 3n + 10$

## The rank-one update

- ▶ uses the evolution path and reduces the number of necessary function evaluations to learn straight ridges from  $\mathcal{O}(n^2)$  to  $\mathcal{O}(n)$ .

---

<sup>4</sup> Hansen, Müller, and Koumoutsakos 2003. Reducing the Time Complexity of the Derandomized Evolution Strategy with Covariance Matrix Adaptation (CMA-ES). *Evolutionary Computation*, 11(1), pp. 1-18

## The rank- $\mu$ update

- ▶ increases the possible learning rate in large populations  
roughly from  $2/n^2$  to  $\mu_w/n^2$
- ▶ can reduce the number of necessary iterations roughly from  $\mathcal{O}(n^2)$  to  $\mathcal{O}(n)$  <sup>(4)</sup>  
given  $\mu_w \propto \lambda \propto n$

Therefore the rank- $\mu$  update is the primary mechanism whenever a large population size is used

say  $\lambda \geq 3n + 10$

## The rank-one update

- ▶ uses the evolution path and reduces the number of necessary function evaluations to learn straight ridges from  $\mathcal{O}(n^2)$  to  $\mathcal{O}(n)$ .

Rank-one update and rank- $\mu$  update can be combined

---

<sup>4</sup> Hansen, Müller, and Koumoutsakos 2003. Reducing the Time Complexity of the Derandomized Evolution Strategy with Covariance Matrix Adaptation (CMA-ES). *Evolutionary Computation*, 11(1), pp. 1-18

# Summary of Equations

## The Covariance Matrix Adaptation Evolution Strategy

Input:  $\mathbf{m} \in \mathbb{R}^n$ ,  $\sigma \in \mathbb{R}_+$ ,  $\lambda$

Initialize:  $\mathbf{C} = \mathbf{I}$ , and  $\mathbf{p}_\mathbf{C} = \mathbf{0}$ ,  $\mathbf{p}_\sigma = \mathbf{0}$ ,

Set:  $c_\mathbf{C} \approx 4/n$ ,  $c_\sigma \approx 4/n$ ,  $c_1 \approx 2/n^2$ ,  $c_\mu \approx \mu_w/n^2$ ,  $c_1 + c_\mu \leq 1$ ,  
 $d_\sigma \approx 1 + \sqrt{\frac{\mu_w}{n}}$ , and  $w_{i=1\dots\lambda}$  such that  $\mu_w = \frac{1}{\sum_{i=1}^\mu w_i^2} \approx 0.3 \lambda$

While not terminate

$\mathbf{x}_i = \mathbf{m} + \sigma \mathbf{y}_i$ ,  $\mathbf{y}_i \sim \mathcal{N}_i(\mathbf{0}, \mathbf{C})$ , for  $i = 1, \dots, \lambda$  sampling

$\mathbf{m} \leftarrow \sum_{i=1}^\mu w_i \mathbf{x}_{i:\lambda} = \mathbf{m} + \sigma \mathbf{y}_w$  where  $\mathbf{y}_w = \sum_{i=1}^\mu w_i \mathbf{y}_{i:\lambda}$  update mean

$\mathbf{p}_\mathbf{C} \leftarrow (1 - c_\mathbf{C}) \mathbf{p}_\mathbf{C} + \underbrace{\mathbb{1}_{\{\|\mathbf{p}_\sigma\| < 1.5\sqrt{n}\}}}_{\text{purple bracket}} \sqrt{1 - (1 - c_\mathbf{C})^2} \sqrt{\mu_w} \mathbf{y}_w$  cumulation for  $\mathbf{C}$

$\mathbf{p}_\sigma \leftarrow (1 - c_\sigma) \mathbf{p}_\sigma + \sqrt{1 - (1 - c_\sigma)^2} \sqrt{\mu_w} \mathbf{C}^{-\frac{1}{2}} \mathbf{y}_w$  cumulation for  $\sigma$

$\mathbf{C} \leftarrow (1 - c_1 - c_\mu) \mathbf{C} + c_1 \mathbf{p}_\mathbf{C} \mathbf{p}_\mathbf{C}^T + c_\mu \sum_{i=1}^\mu w_i \mathbf{y}_{i:\lambda} \mathbf{y}_{i:\lambda}^T$  update  $\mathbf{C}$

$\sigma \leftarrow \sigma \times \exp\left(\frac{c_\sigma}{d_\sigma} \left(\frac{\|\mathbf{p}_\sigma\|}{\mathbb{E}\|\mathcal{N}(\mathbf{0}, \mathbf{I})\|} - 1\right)\right)$  update of  $\sigma$

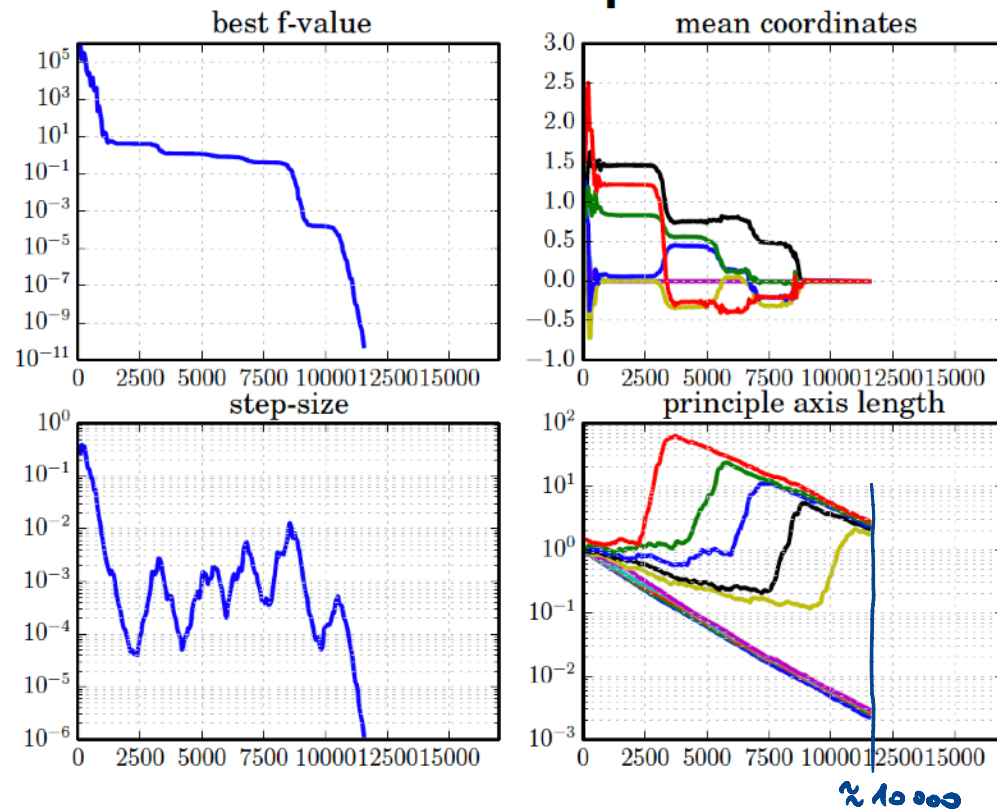
Not covered on this slide: termination, restarts, useful output, boundaries and encoding

## Rank-one and Rank-mu updates

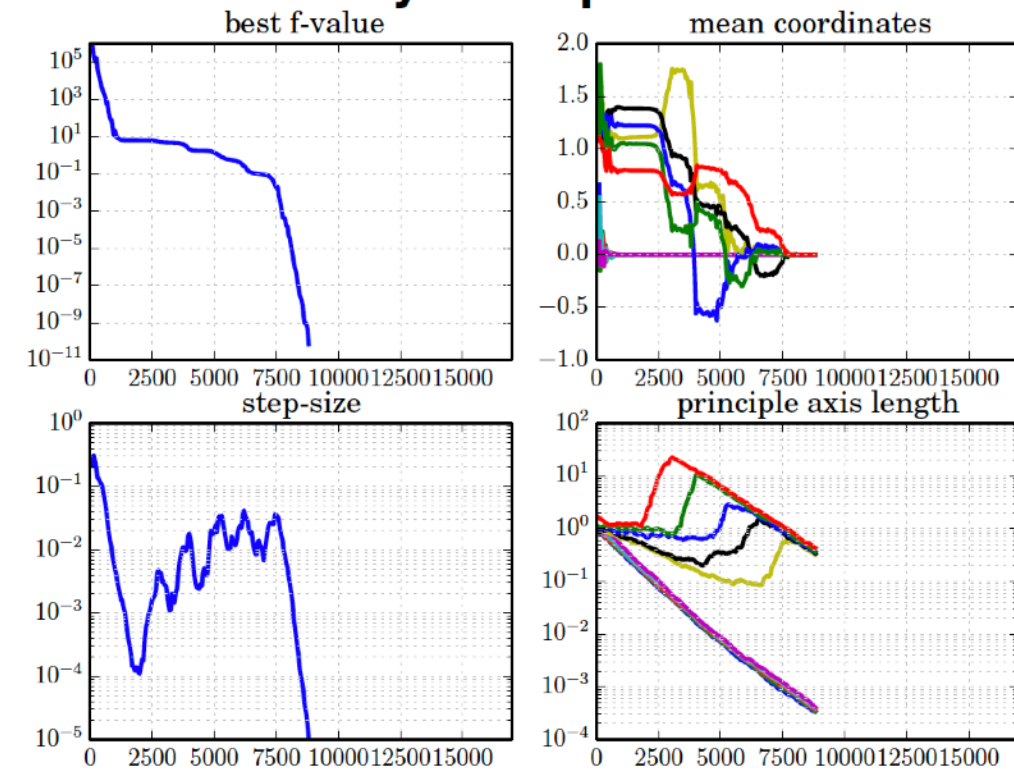
$$f_{\text{TwoAxes}}(x) = \sum_{i=1}^5 x_i^2 + 10^6 \sum_{i=6}^{10} x_i^2$$

# Rank-one and Rank-mu update - default pop size

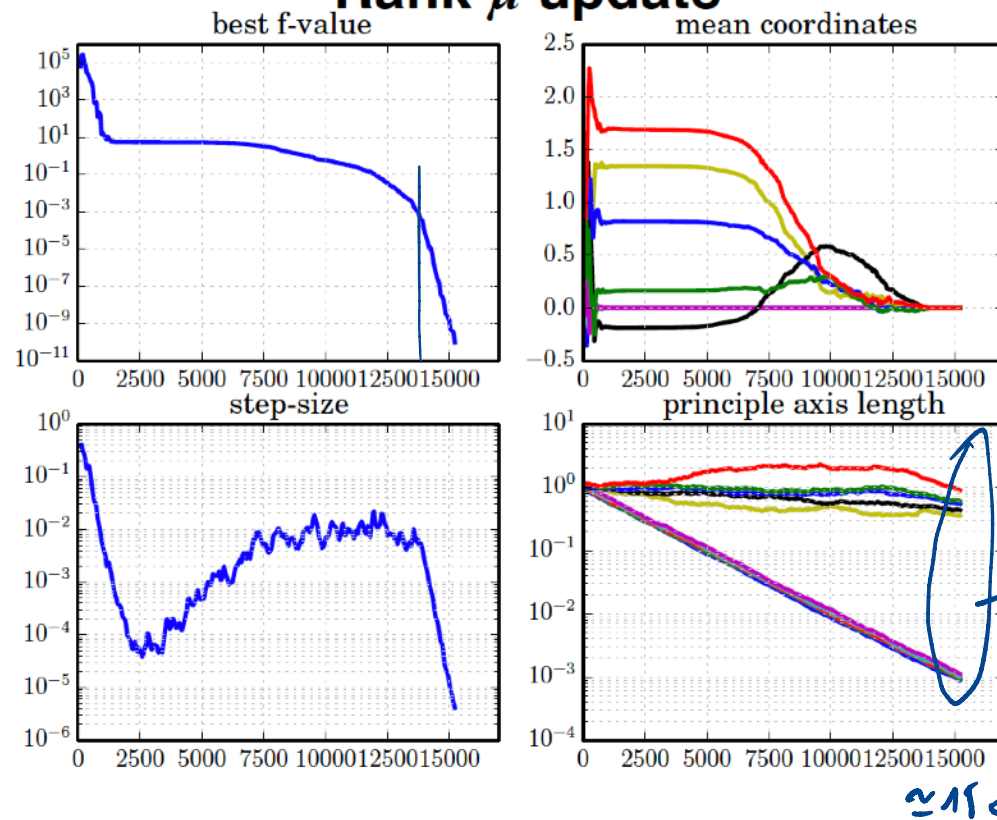
## Rank-one update



## Hybrid update



## Rank- $\mu$ update



$$f_{\text{TwoAxes}}(x) = \sum_{i=1}^5 x_i^2 + 10^6 \sum_{i=6}^{10} x_i^2$$

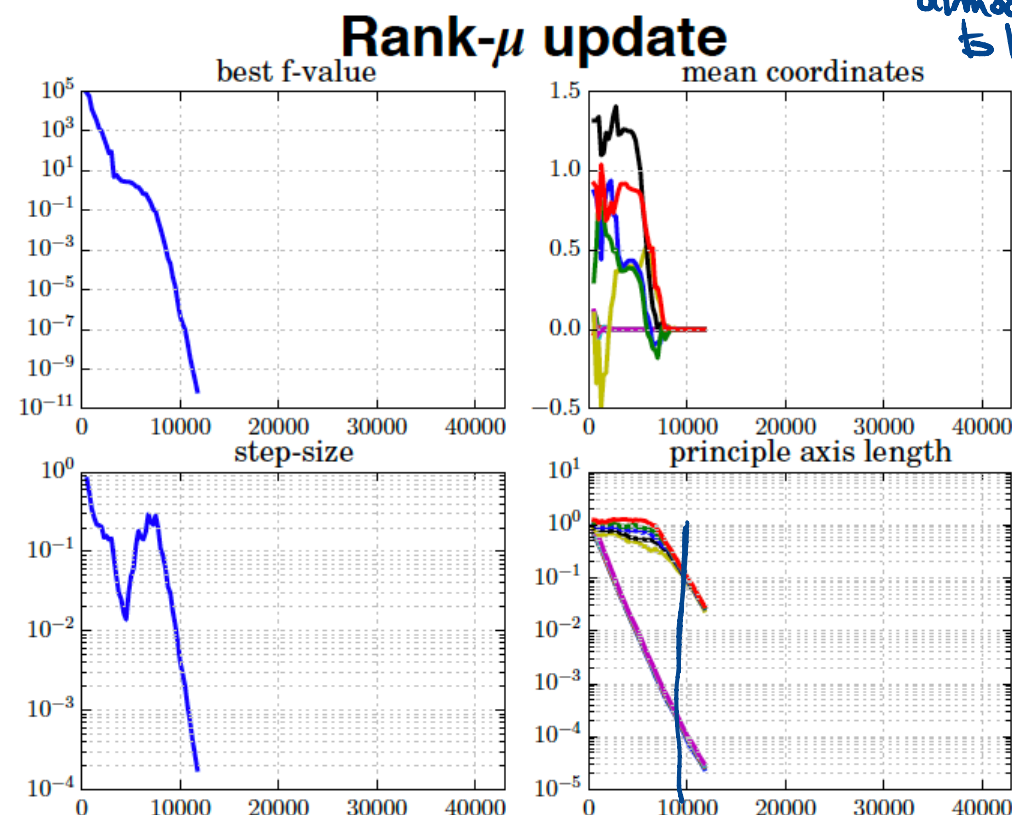
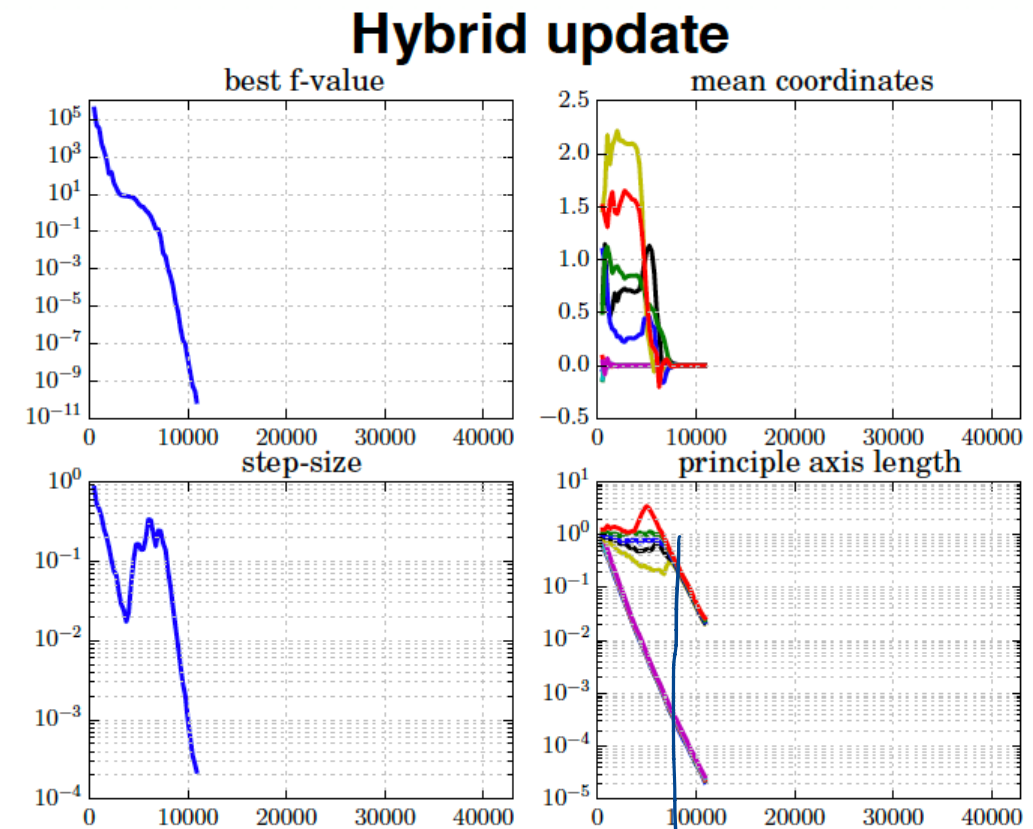
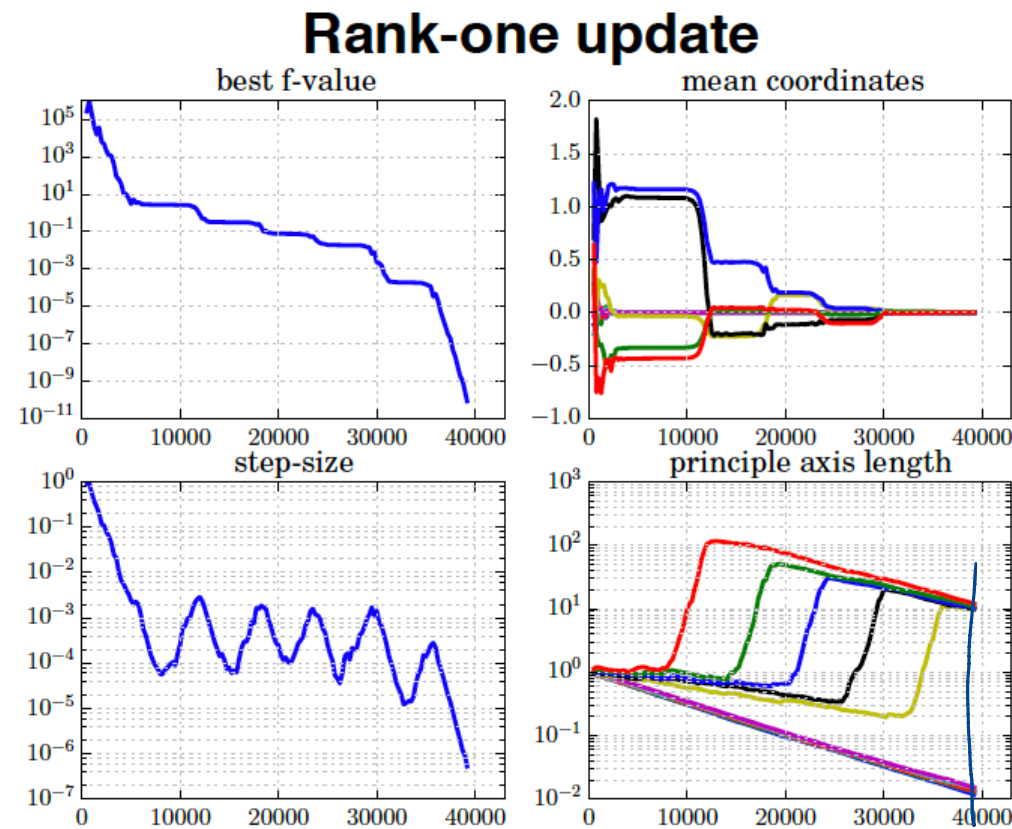
$\lambda = 10$  (default for  $N = 10$ )

$\approx 7500$  to learn the  $H^{-1}$

$\text{end} \approx 10^6$

$\approx 11,000$

# Rank-one and Rank-mu update - larger pop size



almost 40,000  
to learn  $H^{-1}$

less than 10,000 evols to  
learn  $H^{-1}$

$$f_{\text{TwoAxes}}(x) = \sum_{i=1}^5 x_i^2 + 10^6 \sum_{i=6}^{10} x_i^2$$

$$\lambda = 50$$

The hybrid update is as best as  
the best among rank one & rank- $\mu$ .

less than 10,000

# Experimentum Crucis (0)

What did we want to achieve?

- ▶ reduce any convex-quadratic function

$$f(\mathbf{x}) = \mathbf{x}^T \mathbf{H} \mathbf{x}$$

to the sphere model

$$f(\mathbf{x}) = \mathbf{x}^T \mathbf{x}$$

e.g.  $f(\mathbf{x}) = \sum_{i=1}^n 10^{6 \frac{i-1}{n-1}} x_i^2$

without use of derivatives

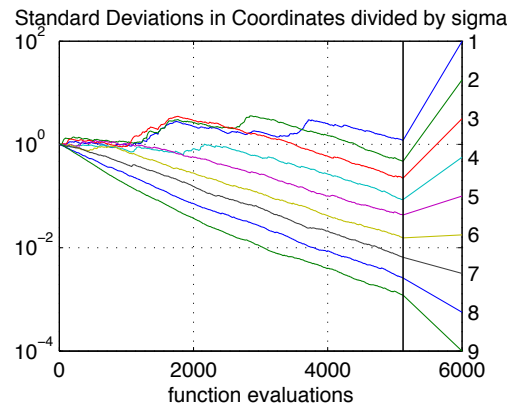
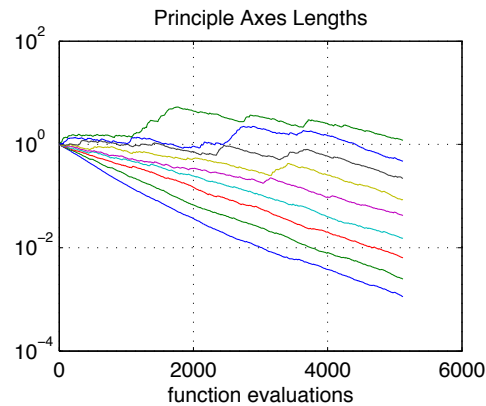
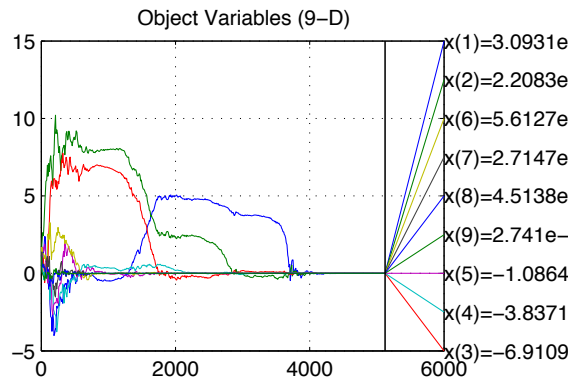
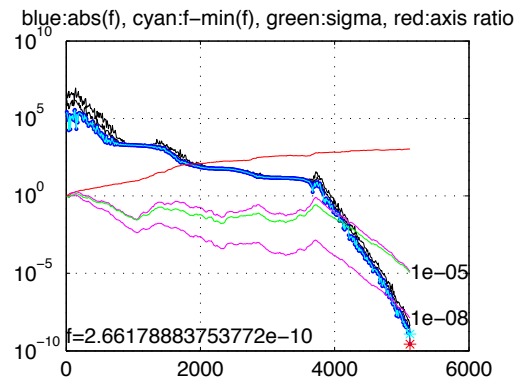
- ▶ lines of equal density align with lines of equal fitness

$$\mathbf{C} \propto \mathbf{H}^{-1}$$

in a stochastic sense

# Experimentum Crucis (1)

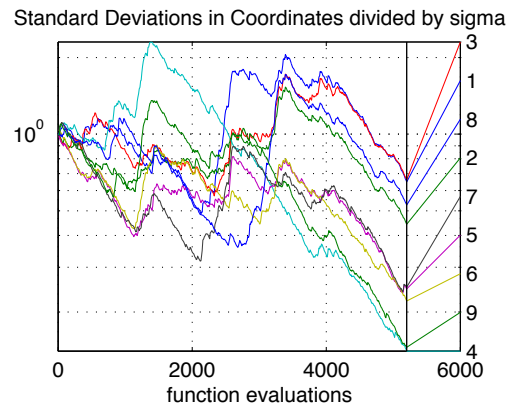
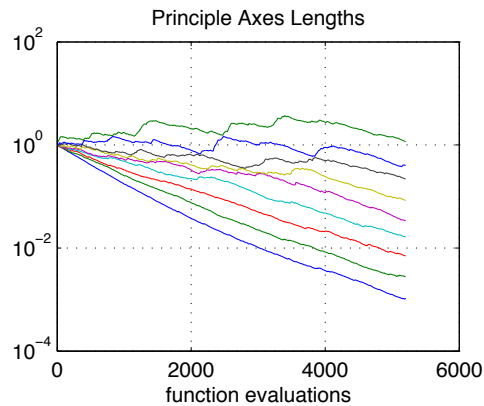
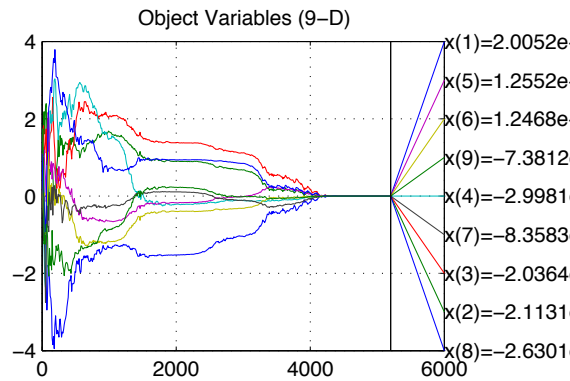
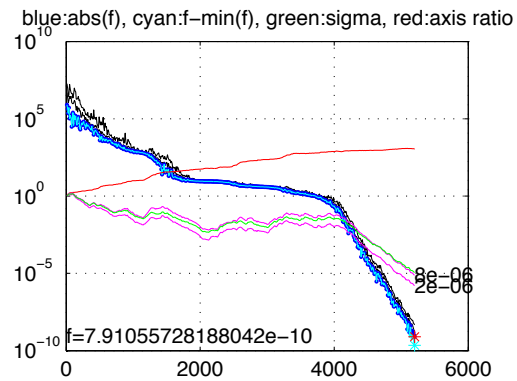
$f$  convex quadratic, separable



$$f(\mathbf{x}) = \sum_{i=1}^n 10^{\alpha \frac{i-1}{n-1}} x_i^2, \alpha = 6$$

# Experimentum Crucis (2)

$f$  convex quadratic, as before but non-separable (rotated)



$C \propto H^{-1}$  for all  
 $g, H$

$$f(\mathbf{x}) = g(\mathbf{x}^T \mathbf{H} \mathbf{x}), \quad g: \mathbb{R} \rightarrow \mathbb{R} \text{ strictly increasing}$$