

Introduction to Computational Science & Engineering (CSE)

16.0002 / 18.0002

Lecture 8:
Optimization II

Laurent Demanet (Math/EAPS)
Youssef Marzouk (AeroAstro)

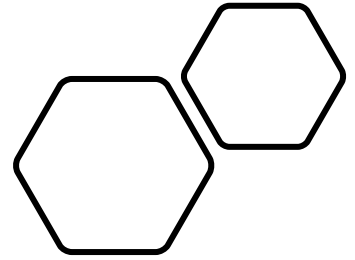
22 November 2021

Follow along: `tests_lec8.py`

**New and improved !
Now in person !!**

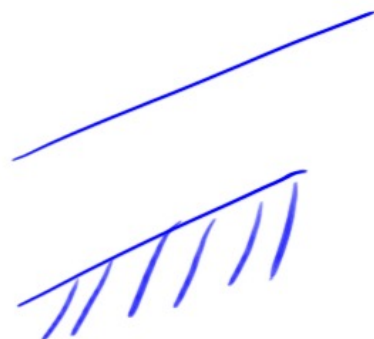


Today: Optimization
Gradient descent
+ plots in python



Unconstrained: $\min_{\underline{x}} J(\underline{x}) \quad \leftarrow$

Constrained: $\min_{\underline{x} \in S} J(\underline{x})$ $\underline{x} \in \mathbb{R}^2$
(some set) $\underline{x} = (x_0, x_1)$



$$a_0 x_0 + a_1 x_1 = b$$

linear equality

$$a_0 x_0 + a_1 x_1 \leq b$$

— inequality



$$x_0^2 + x_1^2 \leq 1$$

convex constraint



$$x_0^2 + x_1^2 \geq 1$$

nonconvex constraint.

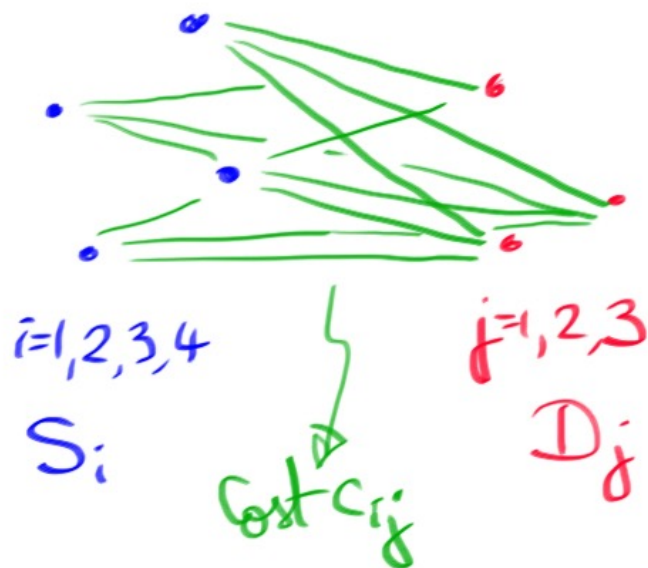
1)

Linear programming

ex. transportation

Supply

Demand



$$\min_x \sum_{ij} c_{ij} x_{ij}$$

$$\text{s.t. } \sum_j x_{ij} \leq S_i$$

$$\sum_i x_{ij} \geq D_j$$

$$(\sum D_j \leq \sum S_i)$$



2) Quadratic programming

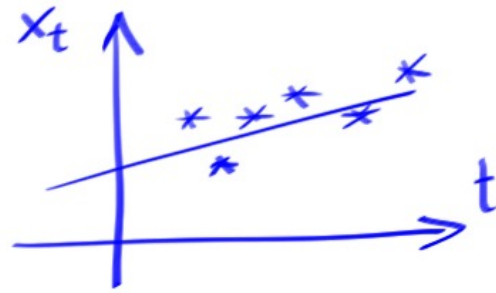
ex. Inverse problems

$$\min \sum_t (x_t - d_t)^2$$

$$\text{s.t. } Ax = b$$

← fits data d_i

← obeys model



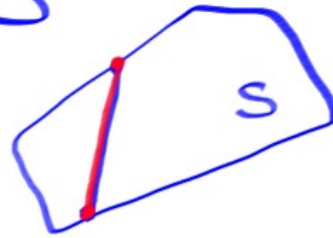
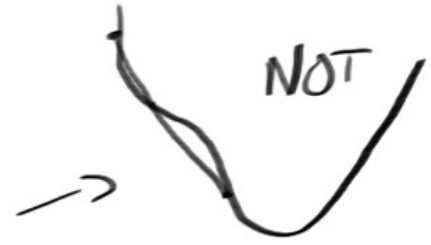
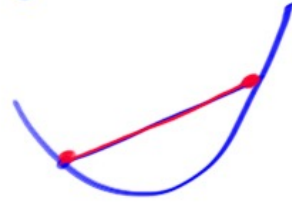
(here $\frac{d^2 x_t}{dt^2} = 0$)

3) Convex programming

$$\min J(x) \quad \leftarrow \text{convex } J$$

$$\text{s.t. } x \in S$$

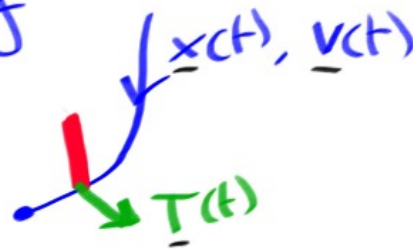
↑ Convex S



Every minimizer
is global!

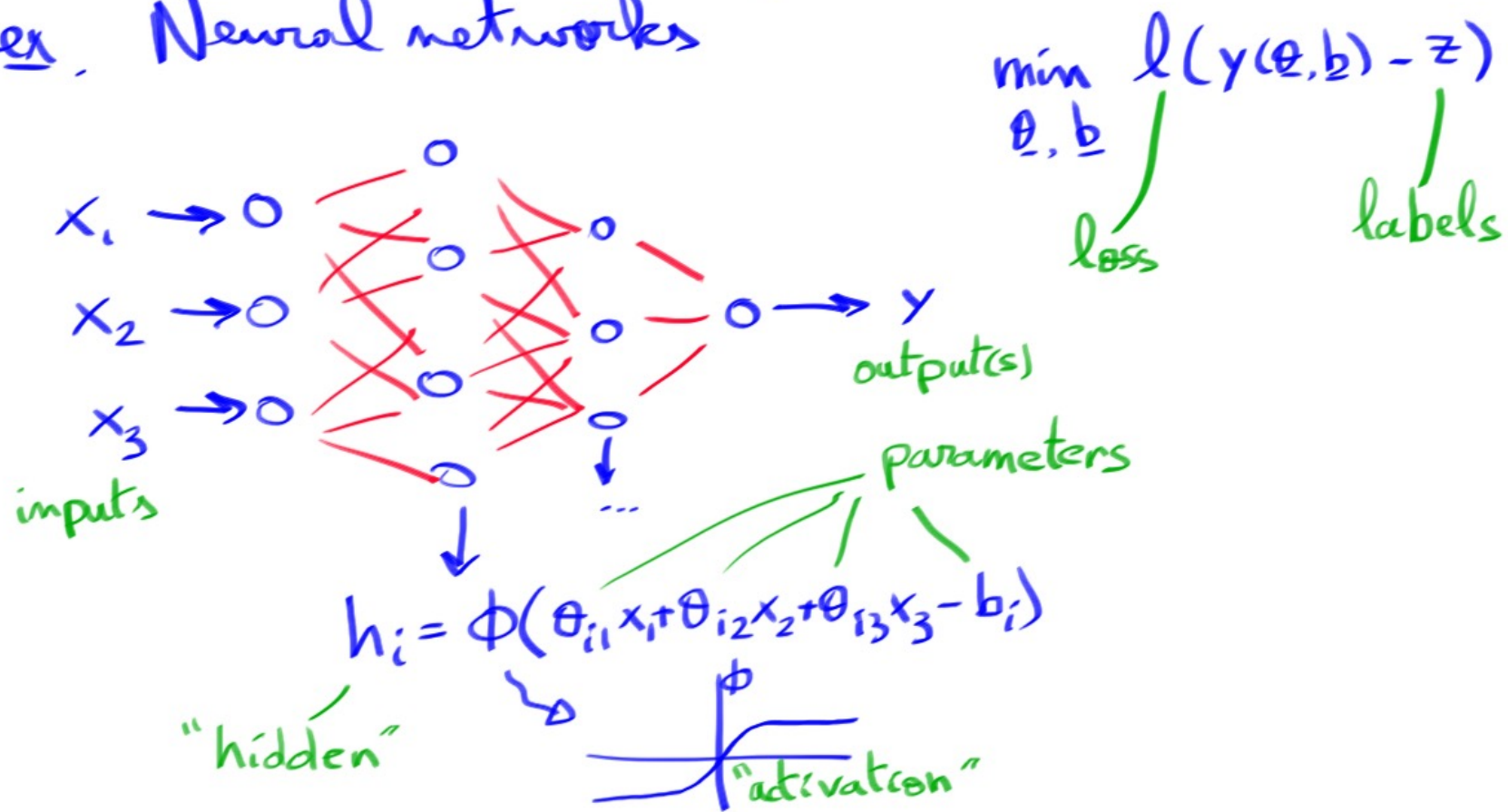
ex. Rocket landing

Convex
in $\underline{x}, \underline{v}, \underline{T}$



4) Nonlinear programming ("none")

ex. Neural networks



Gradient descent

Extremal points: $\nabla J = 0$

(global min
local min
saddle points)

$J(x_0, x_1)$

$$\Delta J = \frac{\partial J}{\partial x_0} \Delta x_0 + \frac{\partial J}{\partial x_1} \Delta x_1$$

$$= \underline{\nabla J} \cdot \underline{\Delta x}$$

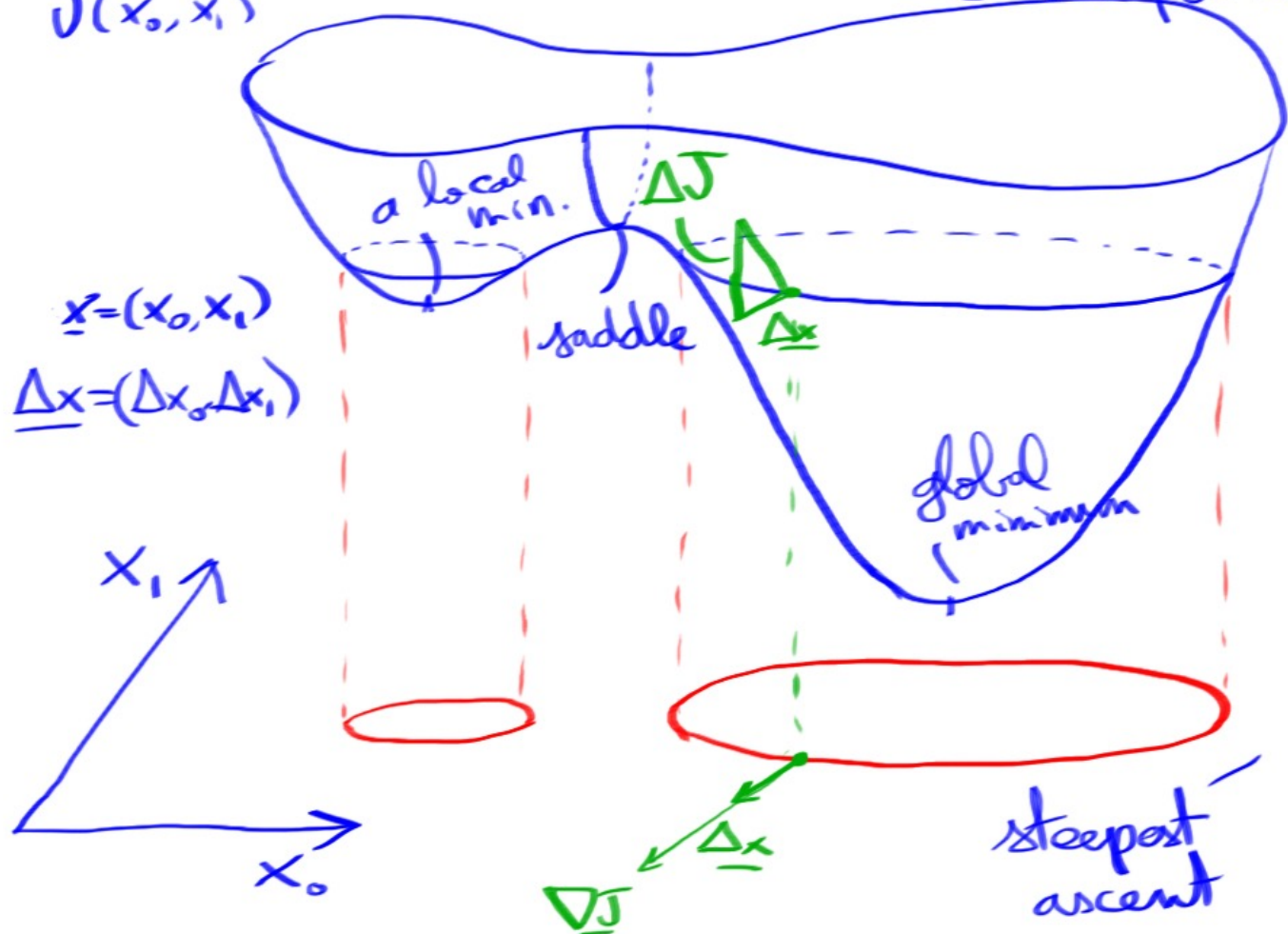
$$\left(\begin{bmatrix} \partial J / \partial x_0 \\ \partial J / \partial x_1 \end{bmatrix} \right) \cdot \begin{pmatrix} \theta \\ \Delta x \end{pmatrix}$$

$$= \|\underline{\nabla J}\| \|\underline{\Delta x}\| \cos \theta$$

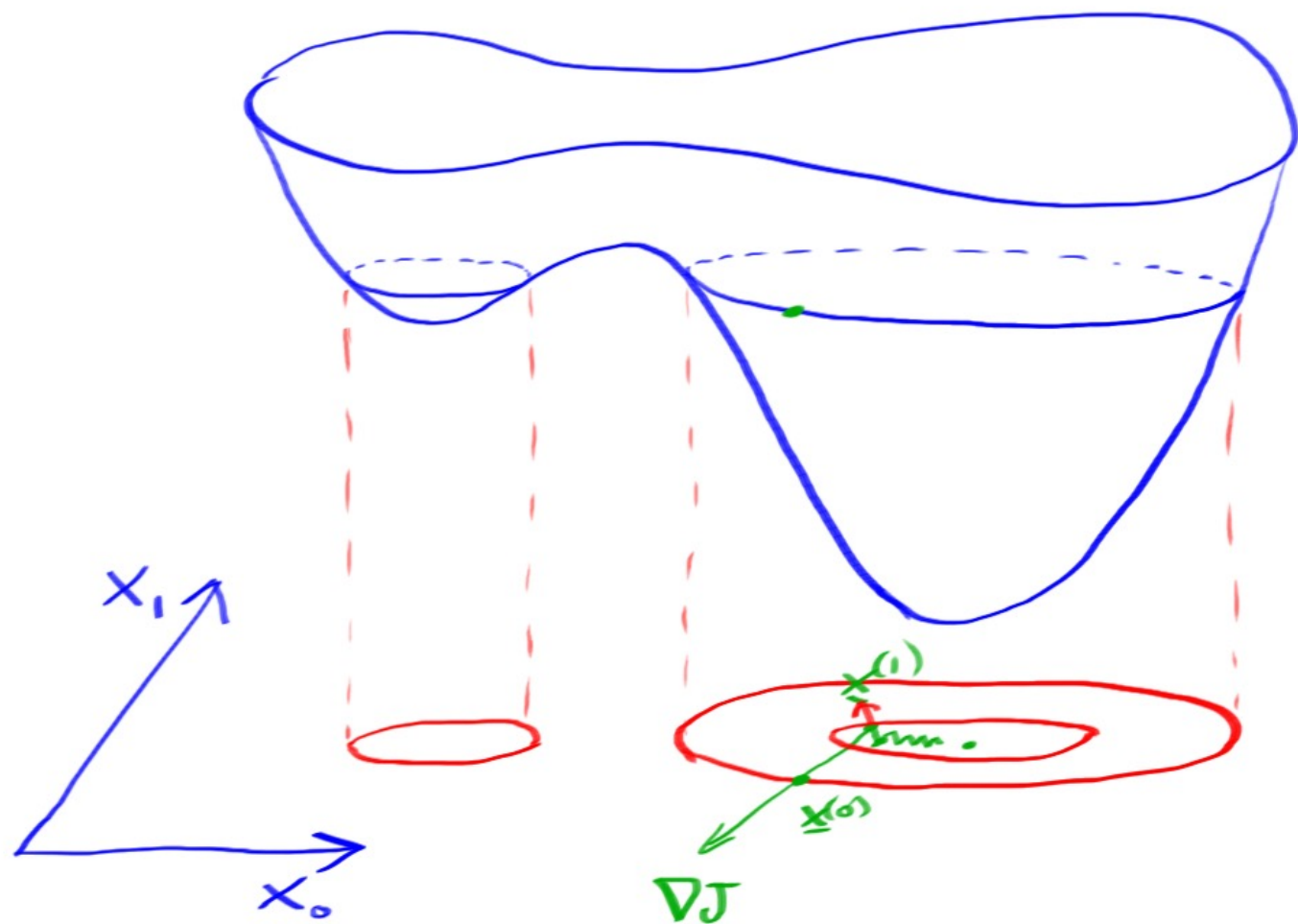
$$\underbrace{\quad}_{=1} \quad \underbrace{\quad}_{=1}$$

$$\underline{\Delta x} = \frac{\underline{\nabla J}}{\|\underline{\nabla J}\|}$$

$$\theta = 0$$



Gradient descent.



$$\underline{x}^{(0)} = (x_0, x_1)$$

$$\underline{x}^{(k+1)} = \underline{x}^{(k)} - \alpha \nabla J(\underline{x}^{(k)})$$

(step length
 $\alpha > 0$)

"line search"

Newton's method (optimization!)

Replace
with

$$\underline{x}^{(k+1)} = \underline{x}^{(k)} - \alpha \nabla J(\underline{x}^{(k)})$$

$$\underline{x}^{(k+1)} = \underline{x}^{(k)} - \underbrace{H^{-1}}_{\text{solution u to}} \nabla J(\underline{x}^{(k)})$$

solution u to

$$H u = \nabla J(\underline{x}^{(k)})$$

HESSIAN:

$$\begin{bmatrix} \frac{\partial^2 J}{\partial x_0^2} & \frac{\partial^2 J}{\partial x_0 \partial x_1} \\ \frac{\partial^2 J}{\partial x_0 \partial x_1} & \frac{\partial^2 J}{\partial x_1^2} \end{bmatrix}$$

$$\begin{bmatrix} \frac{\partial J}{\partial x_0} \\ \frac{\partial J}{\partial x_1} \end{bmatrix}$$