

OPTIMIZATION FOR DATA SCIENCE

September 29, 2025

Today (Session 1/8)

- Logistics
- Basics of optimization

Key resources :

Course webpage

<https://www.lamsade.dauphine.fr/~croyer/teachODS.html>

LOGISTICS

2022-2024: Common M2 MAGE/M2 MAGE Apprentissage

2024-2025: Split

2025-2026: Merge again

Grading

60% Exam (Early Feb. 2026) $\Rightarrow \approx$ like the exercises
40% Homework (groups of ≤ 2 students, extended lab session)

Schedule

8^h30 - 11^h45

8^h30 - 10^h: Lecture

10^h15 - 11^h45: Exercises

or 8^h30 - 11^h45: lab session

29/09] Basics of optimization

03/11] Gradient-based algorithms
10/11

Lab sessions

! 08/12] Stochastic gradient techniques
15/12

19/01] Regularization + Recap
21/01
26/01

BASICS OF OPTIMIZATION

↳ Optimization : "Taking the best decision out of a set of alternatives"

Ex) Vaccine dispatch in Wisconsin
Power generation (e.g. EDF)

) "classical optimization models"

Classify pictures of animals
Extract topics from documents

) modern optimization problems
data-driven

→ Mathematical formulation of an optimization problem

Goal (minimize/maximize)

↓
minimize
 $w \in \mathbb{R}^d$

$f(w)$

↑
objective function

s.t.

↑
"subject to"

$w \in F$

↓
constraint/feasible set
 $F \subseteq \mathbb{R}^d$

f : measure of how good our decision is
⇒ want the lowest possible value for f

$(f: \mathbb{R}^d \rightarrow \mathbb{R})$ → float / real number

w : decision variables (vector with d coordinates)

$$w = \begin{bmatrix} w_1 \\ \vdots \\ w_d \end{bmatrix}$$

F : conditions on the decision variables

⇒ For most of the course, $F = \mathbb{R}^d$

↳ 3 key aspects in optimization

- Theory: Analyze a problem, show existence of solutions, ...
- Computational: Build a solver to solve a certain class of problems (Gurobi, CVXPY, ...)

- Algorithms : $\rightarrow$ Build numerical procedures to solve a problem
 $\rightarrow$ Show theoretical guarantees
 $\rightarrow$ Validate practical performance

Examples of data science optimization problems

① Linear regression / linear least squares

Data $\{ (x_i, y_i) \}_{i=1..m}$ m data points
 $x_i \in \mathbb{R}^d$: vector of features (input)
 $y_i \in \mathbb{R}$: label (output)

Goal: Build the best linear model of your data

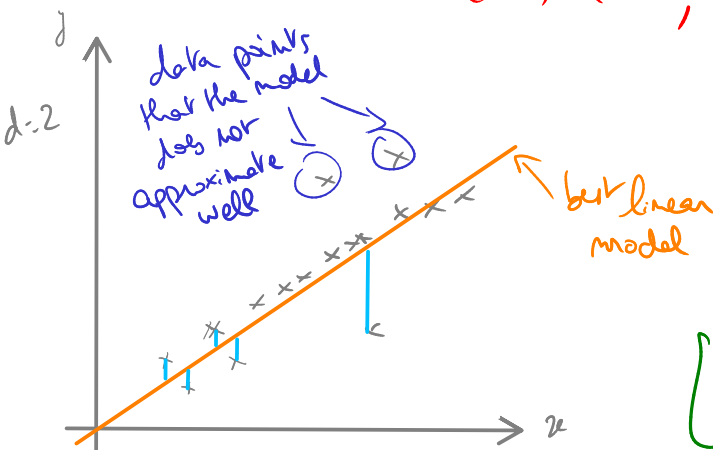
$\Leftrightarrow$ Seek $w \in \mathbb{R}^d$ such that

$$x_i^T w \approx y_i$$

"as close as possible"
(exact model may not exist)

$$\forall (a, b) \in (\mathbb{R}^d)^2, a^T b = \sum_{j=1}^d a_j b_j$$

Linear model of x_i :
combines the coordinates in a linear way, parameterized by w



$$\begin{bmatrix} x_1^T \\ \vdots \\ x_m^T \end{bmatrix} = X \in \mathbb{R}^{m \times d} \text{ (matrix with } m \text{ rows and } d \text{ columns)}$$

$x_i \in \mathbb{R}^d$
 $x_i^T \in \mathbb{R}^{1 \times d}$

Optimization formulation

minimize
 $w \in \mathbb{R}^d$

$$\frac{1}{2m} \|Xw - y\|^2 = \frac{1}{2m} \sum_{i=1}^m (x_i^T w - y_i)^2$$

$y \in \mathbb{R}^m, y = \begin{bmatrix} y_1 \\ \vdots \\ y_m \end{bmatrix}$

$\forall v \in \mathbb{R}^m, \|v\|^2 = \sum_{i=1}^m v_i^2$ (ℓ_2 norm)

→ Minimize the MSE (mean-squared error), i.e.
the average of the squares of the model errors

→ simple, interpretable model

→ Smooth, convex optimization problem $\Rightarrow$ "Easy" to solve

→ Statistical perspective: Maximum Likelihood estimation
(assuming Gaussian noise)

(2) Neural network classification

Data: $\{(x_i, y_i)\}_{i=1..m}$

$x_i \in \mathbb{R}^{d_x}$

$y_i \in \{0, 1\}^{d_y}$

d_y : Number of classes

$$[y_i]_j = \begin{cases} 1 & \text{if data point } i \text{ is in class } j \\ 0 & \text{otherwise} \end{cases}$$

Model: Feed-forward neural architecture

$$NN: x^0 \in \mathbb{R}^{d_0} \mapsto x^1 \in \mathbb{R}^{d_1} \mapsto \dots \mapsto x^L \in \mathbb{R}^{d_L}$$

$d_0 = d_x$ $d_L = d_y$

L : number of layers

$$\forall l = 1, \dots, L,$$

$$x^l = \sigma(\underline{w^l x^{l-1} + b^l})$$

↑
Linear transformation

Nonlinear activation function

$$\text{ex) } \sigma(v) = \begin{bmatrix} \tanh(v_1) \\ \tanh(v_p) \end{bmatrix}$$

$$w^l \in \mathbb{R}^{d_{l-1} \times d_l}$$

$$b^l \in \mathbb{R}^{d_l}$$

w (model parameter): concatenation of (w^l, b^l)

$$w \in \mathbb{R}^d$$

$$d = d_1 d_0 + d_1 + d_2 d_1 + d_2 + \dots + d_L d_{L-1} + d_L$$

Goal: Given my model $NN(\cdot; w)$, find w such that
 $\forall i = 1 \dots n, \quad NN(x_i; w) \approx y_i$

(Key $\neq$ with linear regression:

• Model is nonlinear

• $y_i \in \{0, 1\}^{d_y} \Rightarrow$ closeness to y_i measured differently than in a regression context

Optimization problem

minimize
 $w \in \mathbb{R}^d$

$$\frac{1}{n} \sum_{i=1}^n \left\{ \log \left(\sum_{j=1}^{d_y} \exp([NN(x_i; w)]_j) \right) - \sum_{j=1}^{d_y} [y_i]_j [NN(x_i; w)]_j \right\}$$

Finite-sum
structure
(average over
data points)

Negative Log-likelihood
function

(accounts for the fact that
 $NN(x_i; w)$ is not necessarily a
0/1 vector)

Model

$NN(x_i; w) = x^L$ when given

$x^0 = x_i$ as input

w for parameters

→ High-dimensional, highly nonconvex problem

→ Yet it can be solved using stochastic gradient algorithms

1) Solutions of optimization problems

From now on, we consider

$$(P) \underset{w \in \mathbb{R}^d}{\text{minimize}} f(w)$$

$$f: \mathbb{R}^d \rightarrow \mathbb{R}$$

Def. $w^* \in \mathbb{R}^d$ is a solution/a global minimum of (P)

$$\text{if } \forall w \in \mathbb{R}^d, f(w^*) \leq f(w)$$

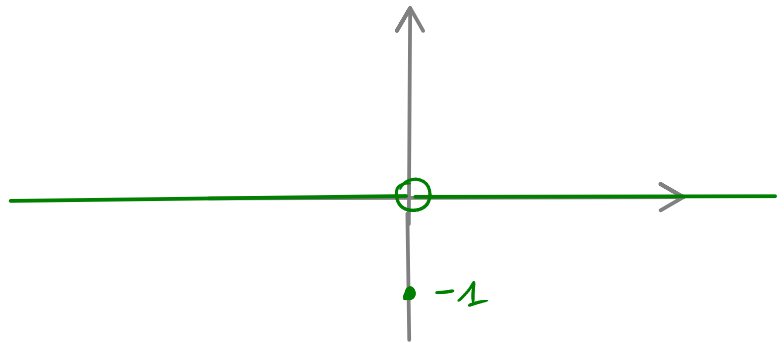
→ Finding global minima is (NP-) hard!

→ In practice,
we:

$$\text{Ex) } d=1 \quad f(w) = \begin{cases} 0 & \text{if } w \neq 0 \\ -1 & \text{if } w = 0 \end{cases}$$

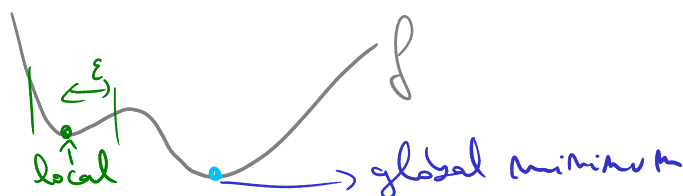
- use classes of
functions f for
which minimization is
easier

- use weaker notions
of minima



Def. $\bar{w} \in \mathbb{R}^d$ is a local minimum of (P) if

$$\exists \varepsilon > 0 \text{ such that } \forall w \in \mathbb{R}^d, \|\bar{w} - w\| \leq \varepsilon, f(\bar{w}) \leq f(w)$$



- Any global minimum is a local minimum, but the converse is not true
- Local minima are also difficult to find in general
- but it is easier to check that a point is a local minimum, especially for certain classes of functions

C^1 functions

"continuously differentiable"

$f: \mathbb{R}^d \rightarrow \mathbb{R}$ is called C^1 if $\forall w \in \mathbb{R}^d, \exists \nabla f(w) \in \mathbb{R}^d$ that represents the local variations in f around w

Gradient of f at w

Intuition: If $v \in \mathbb{R}^d$ is close to w , then

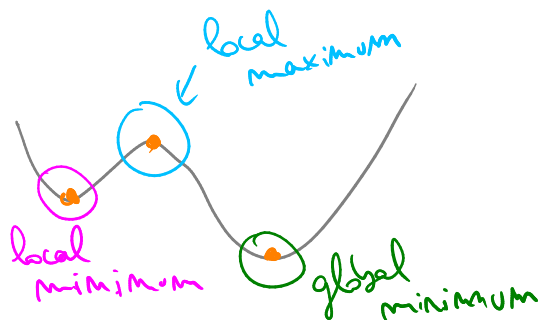
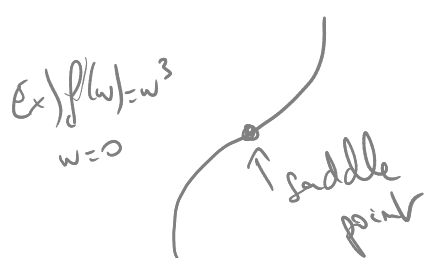
$$f(v) \approx f(w) + \nabla f(w)^T (v - w)$$

Linear function of v

Theorem (for $f(C^1)$): $\left[\bar{w} \in \mathbb{R}^d \text{ local minimum of } f \right] \Rightarrow \left[\nabla f(\bar{w}) = 0_{\mathbb{R}^d} \right]$

"First-order optimality conditions"

→ Only an implication, converse not true in general. A point zero gradient can be a minimum, a maximum or a saddle point



• points with $\nabla f() = 0$

→ If we have access to ∇f , we can check whether $\nabla f(w) = 0$

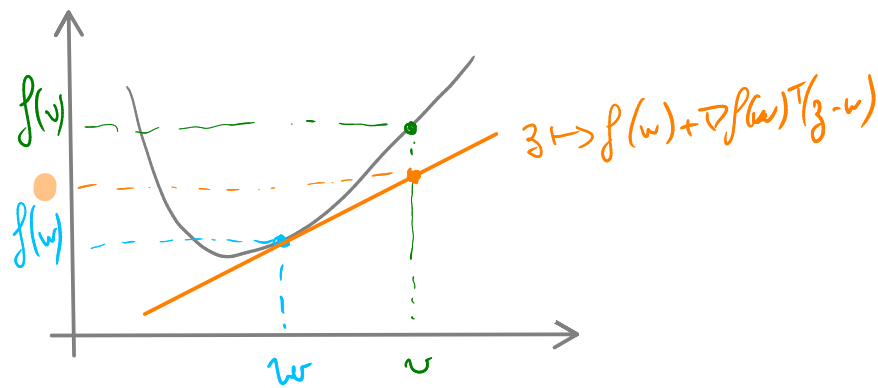
or not $\Rightarrow$ Basic idea behind the gradient descent algorithm!

C^1 , convex functions

Def: $f: C^1$ is convex if

$$\forall (v, w) \in (\mathbb{R}^d)^2, \quad \underline{f(v) \geq f(w) + \nabla f(w)^T (v - w)}$$

Linear function of v
(Approximation/Tangent of f at w)



Ex) - Any linear function

$$w \mapsto a^T w + b$$

is convex $a \in \mathbb{R}^d, b \in \mathbb{R}$

$$\bullet w \mapsto \frac{1}{2} \|w\|^2 \text{ convex}$$

$$\bullet w \mapsto \frac{1}{2m} \|Xw - y\|^2$$

$\forall X \in \mathbb{R}^{n \times d}, \forall y \in \mathbb{R}^n$, is convex

Key properties of C^1 , convex functions

$f: C^1$
convex

• Every local minimum of $\underset{w \in \mathbb{R}^d}{\text{minimize}} f(w)$ is a global minimum.

• $[w^* \text{ global minimum of } \underset{w \in \mathbb{R}^d}{\text{minimize}} f(w)] \Leftrightarrow [\nabla f(w^*) = 0_{nd}]$

C^1 , strongly convex functions

Def: $f \in C^1$ is μ -strongly convex (for $\mu > 0$) if

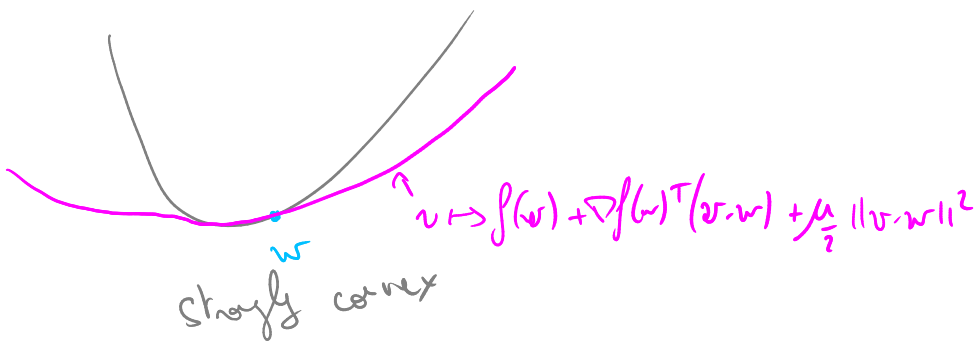
$$\forall (v, w) \in (\mathbb{R}^d)^2, \quad f(v) \geq f(w) + \nabla f(w)^T (v - w) + \frac{\mu}{2} \|v - w\|^2$$

→ [μ -strongly convex] $\Rightarrow$ convex

→ Examples: $w \mapsto \frac{\mu}{2} \|w\|^2$ μ -strongly convex

$w \mapsto a^T w + b$ not strongly convex

Quadratic function of v
Quadratic term



Key property of C^1 strongly convex functions.

If $f \in C^1$, μ -strongly convex, minimize $f(w)$ has a
 $w \in \mathbb{R}^d$
unique global minimum, and it is the only point $w^* \in \mathbb{R}^d$
such that $\nabla f(w^*) = 0_{\mathbb{R}^d}$

$$\nabla f(w^*) = \underbrace{0}_{\begin{bmatrix} 0 \\ \vdots \\ 0 \end{bmatrix} \in \mathbb{R}^d} \quad (\Leftrightarrow) \quad \forall j=1 \dots d, \quad [\nabla f(w^*)]_j = 0 \in \mathbb{R}$$

$$\Leftrightarrow \|\nabla f(w^*)\| = 0$$

$$\|v\| = \sqrt{\sum_{j=1}^d v_j^2}$$

C^2 functions

Def: f is C^2 if it is C^1 and $\forall w \in \mathbb{R}^d, \exists \nabla^2 f(w) \in \mathbb{R}^{d \times d}$ that reflects the (second-order) changes of f around w

Hessian matrix of f at w Square matrix of size $d \times d$

Intuition: $f(v) \approx f(w) + \nabla f(w)^T (v-w) + \underbrace{\frac{1}{2} \underbrace{(v-w)^T}_{\in \mathbb{R}^{1 \times d}} \underbrace{\nabla^2 f(w)}_{\in \mathbb{R}^{d \times d}} \underbrace{(v-w)}_{\in \mathbb{R}^d}}_{\in \mathbb{R}}$

Theorem: $f \in C^2, (P) \underset{w \in \mathbb{R}^d}{\text{minimize}} f(w)$

2nd order necessary conditions

i) $[\bar{w} \text{ local minimum of } (P)] \Rightarrow [\nabla f(\bar{w}) = 0_{\mathbb{R}^d} \text{ and } \nabla^2 f(\bar{w}) \succeq 0]$

ii) $[\nabla f(\bar{w}) = 0_{\mathbb{R}^d} \text{ and } \nabla^2 f(\bar{w}) \succ 0] \Rightarrow [\bar{w} \text{ local minimum of } (P)]$

2nd order sufficient conditions

$\forall A \in \mathbb{R}^{d \times d}, A \succeq 0 =$ "A is positive semi-definite"

Löwner order

$=$ A symmetric ($A_{ij} = A_{ji} \quad \forall (i,j)$)
and $v^T A v \geq 0 \quad \forall v \in \mathbb{R}^d$

⚠ $A \geq 0$ does not imply that $A_{ij} \geq 0 \forall (i,j)$ $A = \begin{bmatrix} 2 & -1 \\ -1 & 2 \end{bmatrix}$

$A > 0$ = "A positive definite"

= $A \geq 0$ and $v^T A v > 0 \forall v \in \mathbb{R}^d, v \neq 0_{\mathbb{R}^d}$

C^2 , convex functions

$f \in C^2$

i) f convex if $\nabla^2 f(w) \geq 0 \forall w \in \mathbb{R}^d$

ii) f μ -strongly convex (with $\mu > 0$)

if $\nabla^2 f(w) \geq \mu I \forall w \in \mathbb{R}^d$

$\Leftrightarrow \nabla^2 f(w) - \mu I \geq 0$

$I = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \in \mathbb{R}^{d \times d}$
(identity matrix "eye")

TAKETAWAYS

For convex functions: $\nabla^2 f(w) \geq 0$ true $\forall w$

• 1st order condition $\Rightarrow$ 2nd order necessary conditions

• local $\Rightarrow$ global $\Leftrightarrow \nabla f() = 0$

For μ -strongly convex:

- $\exists$ unique global minimum
- Special class of convex functions

C^1 / C^2 : classes of functions for which we have conditions linked to global minima

Exercise (number 1 in Exercise sheet 1)

Linear least squares

$$X = \begin{bmatrix} x_1^T \\ \vdots \\ x_m^T \end{bmatrix} \quad y = \begin{bmatrix} y_1 \\ \vdots \\ y_m \end{bmatrix}$$

$$\{(x_i, y_i)\}_{i=1..m} \quad x_i \in \mathbb{R}^d, y_i \in \mathbb{R}$$

$$(P) \quad \underset{w \in \mathbb{R}^d}{\text{minimize}} \quad f(w) = \frac{1}{2m} \|Xw - y\|^2 = \frac{1}{2m} \sum_{i=1}^m (x_i^T w - y_i)^2$$

a) Suppose that there exists w^* such that $\underbrace{X}_{\substack{m \times d \\ m \times 1}} \underbrace{w^*}_{\substack{d \times 1 \\ m \times 1}} = \underbrace{y}_{\substack{m \times 1 \\ m \times 1}}$.
What is the link between w^* and problem (P)?

$$\left\{ \begin{array}{l} \cdot f(w^*) = \frac{1}{2m} \|Xw^* - y\|^2 = \frac{1}{2m} \|y - y\|^2 = \frac{1}{2m} \|0\|^2 = 0 \\ \cdot \forall w \in \mathbb{R}^d, \quad f(w) = \frac{1}{2m} \sum_{i=1}^m (x_i^T w - y_i)^2 \geq 0 \quad (\text{average of squares} \geq 0) \end{array} \right.$$

$\rightarrow \forall w \in \mathbb{R}^d, f(w^*) \leq f(w)$
Hence w^* is a global minimum of (P)

$$b) \quad \nabla f(w) = \frac{1}{m} \underbrace{X^T}_{\substack{d \times m \\ d \times 1}} \underbrace{(Xw - y)}_{\substack{m \times 1 \\ m \times 1}} \quad \forall w \in \mathbb{R}^d$$

If $\bar{w}$ is a local minimum of (P), what condition involving $\nabla f(\bar{w})$ holds? Does this condition hold for w^* ?

$$[\bar{w} \text{ local minimum of (P)}] \Rightarrow [\nabla f(\bar{w}) = 0_{\mathbb{R}^d}]$$

If w^* is the point from a),

$$\nabla f(w^*) = \frac{1}{n} X^T (Xw^* - y) = \frac{1}{n} X^T (y - y) = \frac{1}{n} X^T 0_{n \times d} = 0_{1 \times d}$$

c) $\nabla^2 f(w) = \frac{X^T X}{n}$ for every $w \in \mathbb{R}^d$ (constant because f is a quadratic function)

i) $\frac{1}{n} X^T X \succeq 0$, what does this imply on f and (P) ?

$$\nabla^2 f(w) \succeq 0 \quad \forall w \Rightarrow f \text{ is convex}$$

Then, any local minimum of (P) is global

ii) If we had $\frac{1}{n} X^T X \succeq \mu I$ with $\mu > 0$, what would that imply on f and (P) ?

$$\nabla^2 f(w) \succeq \mu I \quad \forall w \Rightarrow f \text{ is } \mu\text{-strongly convex}$$

Then, (P) has a unique global minimum