

M2 M090
2025-2026

MACHINE LEARNING FOR OPTIMIZATION

November 20, 2025

Today: Introduction, Automatic differentiation / Backpropagation

Logistics:

- Course webpage

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

- My email

clement.royer@lamsade.dauphine.fr

Schedule (subject to uncertainties ...)

Nov 20 1.45-5pm: Lecture

Nov 27 8.30-11.45pm: Lab

Dec 8 1.45-5pm: Lab

Dec 9 8.30-11.45pm: Lab

Dec 16 8.30-11.45pm: Lab



Assessment:

- Homework / Project (in groups of n students with $n \in \{1, 2\}$)

- Goal: Push Lab sessions a bit further

- Intended deadline: Winter 2026

(for sure after the holidays)

Key concepts in ML for optimization

Main question What can we do with ML tools and optimization solvers?

1) Try to learn something as you are solving an optimization problem and collecting data about the problem

Ex) In blackbox optimization, you collect a history of function values and you try to learn a (regression) model of the function to be optimized from these values

Giovarelli
Sohab
Vicente 2024:

Neural networks do not necessarily give the best models in blackbox optimization

2) Learn on past problem instances (and hope that it generalizes to new instances of the problem)

Ex) Learning hidden constraint classifier in blackbox optimization

Tools (from ML) that people are using to perform 1) and 2)

↳ No LLTs (for now)

↳ But neural networks!

- Classical ones (RNNs, CNNs, ...)

- Less common ones (Graph neural networks, neural network with implicit layers)

→ These networks are used to solve supervised learning tasks, but also reinforcement learning and imitation learning

Basic mathematical formulation of a neural net

$$z_0 \mapsto z_1 \mapsto \dots \mapsto z_L \quad \text{"L layers"}$$

$$\forall l=1..L, \quad z_l = \phi(A_l z_{l-1} + b_l) \quad \text{"l^{th} layer"}$$

$$\text{ReLU}(t) = \max(t, 0)$$

$$\text{ReLU}(z) = [\max(t_i, 0)]_i$$

Activation
function
(tanh, ReLU, ...)
applied componentwise

parameters to
be learned

Neural network (NN): Input $x = z_0$, get $\text{NN}(x) = z_L$
as output

Examples (of layers $z \mapsto y$)

- Fully connected layers (Basic one)

$$y_j = \phi\left(\sum_i A_{ij} z_i + b_j\right)$$

$$A = [A_{ij}]$$

$$b = \begin{bmatrix} b_1 \\ \vdots \end{bmatrix}$$

$$z = [z_i], y = [y_i]$$

- Convolutional layers (at the heart of CNNs)

$$z \in \mathbb{R}^{d_o \times d_o \times c_o}$$

Three-way
Tensor



$$y_{a,b,c} = \sigma\left(\sum_{i,j,h} A_{ijh,c} x_{a-i,b-j,h} + b_h\right)$$

$$\in \mathbb{R}^{q_o \times q_o \times c_i}$$

$$A \in \mathbb{R}^{q_o \times q_o \times c_o \times c_i}$$

$$b \in \mathbb{R}^{c_i}$$

Working
with
tensors:
Very important
in ML

- Recurrent layers (RNNs) / ResNet

ResNet: $y = z + B(\text{ReLU}(A z))$

$\uparrow$ residual connection

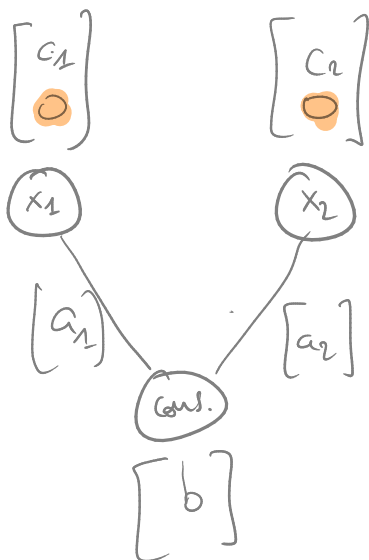
RNN: $z_l = f(z_{l-1}) \quad \forall l$ (same function f at every layer)

- Graph neural networks (GNNs)

Full description: $\left\{ \begin{array}{l} \text{Graph } G=(V, E) \text{ undirected graph} \\ \text{Embedding} \\ \xi(G, \cdot) : V \rightarrow \mathbb{R}^d \\ v \mapsto \xi(G, v) \end{array} \right.$

Map a vertex to a vector of features
(also possible for edges)

Ex) A vertex could represent a variable or a constraint in linear programming



$$\left\{ \begin{array}{l} \text{minimize} \quad [c_1]x_1 + [c_2]x_2 \\ x \in \mathbb{R}^2 \\ \text{s.t.} \quad [a_1]x_1 + [a_2]x_2 = [b] \end{array} \right. \text{Constraint}$$

$x \geq 0$

→ The GNN structure consists in computing functions of the graph and its embedding.

Message Passing Operator

Start with an embedding ξ_0

Compute the recursion $\forall v \in V$ (and possibly $\forall edge_{e \in E}$)

$$\forall t=0..T-1, \quad \xi^{t+1}(G, v) = \text{Comb} \left[\xi^t(G, v), \sum_{u \in N(v)} \xi^t(G, u) \right]$$

(ex)
→ Linear mapping
→ Neural network layer
(of convolutional/fully connected type)

↑
Neighbours of v in the graph

NB: As we will (probably) see, implementing GNNs is challenging!

Typical abstraction: The parameters of the NN are gathered in a (big) vector $w \in \mathbb{R}^d$

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

Training a NN: Given data $\{(x_i, y_i)\}_{i=1..m}$, find the best value of w , that is the value for which $NN(x_i, w) \approx y_i$
(want to output something close to y_i when given x_i as input)
 $\Rightarrow$ This is an optimization problem!

One typical formulation:

$$(*) \quad \underset{w \in \mathbb{R}^d}{\text{minimize}} \quad \frac{1}{m} \sum_{i=1}^m \text{loss}(NN(x_i, w), y_i)$$

↑
Output of the NN with input x_i and parameters w

↑
Loss function

(x) In the blackbox notebook

- linear regression

$$\underset{w \in \mathbb{R}^d}{\text{minimize}} \quad \frac{1}{n} \|Xw - y\|_2^2$$

- logistic regression

$$\underset{w \in \mathbb{R}^d}{\text{minimize}} \quad \frac{1}{n} \sum_{i=1}^n \log(1 + \exp(-y_i x_i^T w))$$

(e.g. $\text{loss}(a, b) = \frac{1}{2} \|a - b\|_2^2$

$$= \frac{1}{2} (a - b)^T (a - b)$$

$$= \frac{1}{2} \sum_{i=1}^m (a_i - b_i)^2$$

$x_i \mapsto x_i^T w$: "One-layer linear NN" (aka a linear model)

→ State of the art for problems of the form (x)

Apply some form of stochastic gradient technique
(SGD, Adam, Adagrad, ...)

that involve computing the derivative of NN with respect to the parameters w

→ Called backpropagation in general, but sometimes the derivative calculation (the most important part) is called backpropagation

→ Backpropagation is one of the reasons deep learning is successful!

Backpropagation / Automatic differentiation

General idea: Given a code that computes $f(w)$ where $f: \mathbb{R}^d \rightarrow \mathbb{R}$, design a code that evaluates the derivative of f at w , denoted by $\frac{\partial f}{\partial w}$.

Theorem (informally): If you apply backward-mode automatic differentiation (aka backpropagation), the cost of computing $\frac{\partial f}{\partial w}$ is at most 6 times the cost of computing $f(w)$.

Notation: $\frac{\partial f}{\partial w} \equiv \frac{\partial f}{\partial w}(w)$ fixed value rather than a function

$$f: \mathbb{R} \rightarrow \mathbb{R} \quad w \mapsto f(w) \quad \frac{\partial f}{\partial w} \in \mathbb{R} \text{ (or } \mathbb{R}^{1 \times 1}, \text{ or } \mathbb{R}^{1 \times 1 \times 1}, \dots)$$

$$f: \mathbb{R}^d \rightarrow \mathbb{R} \quad w \mapsto f(w) \quad \frac{\partial f}{\partial w} \in \mathbb{R}^{1 \times d} \quad \text{"row vector"} \\ [v_1 \dots v_d]$$

$$\frac{\partial f}{\partial w} = \nabla f(w)^T \quad \nabla f(w): \text{gradient} \\ \in \mathbb{R}^d = \mathbb{R}^{d \times 1} \quad \begin{bmatrix} v_1 \\ \vdots \\ v_d \end{bmatrix}$$

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

$$w \mapsto f(w) \in \mathbb{R}^q$$

$$\frac{\partial f}{\partial w} \in \mathbb{R}^{q \times d}$$

"Jacobian matrix"

$$f: \mathbb{R}^{d_1 \times \dots \times d_n} \rightarrow \mathbb{R}^{q_1 \times \dots \times q_s}$$

$$w \mapsto f(w)$$

$$\frac{\partial f}{\partial w} \in \mathbb{R}^{(q_1 \times \dots \times q_s) \times (d_1 \times \dots \times d_n)}$$

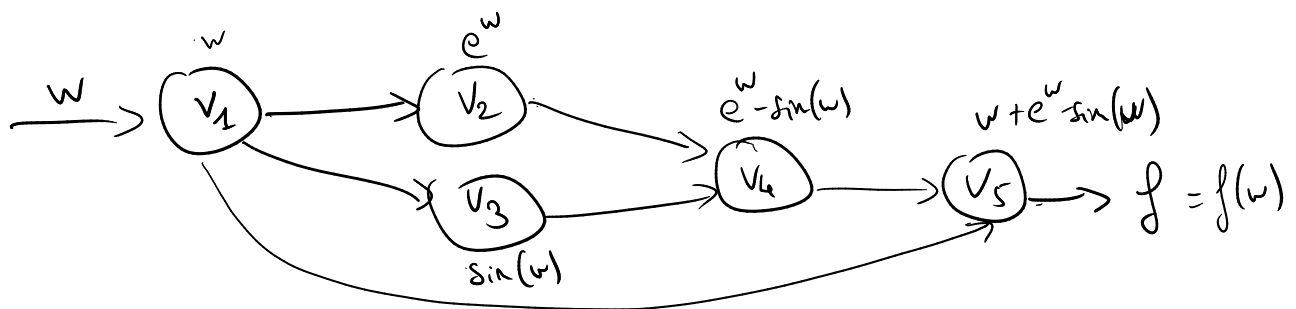
(equivalently $\mathbb{R}^{q_1 \times \dots \times q_s \times d_1 \times \dots \times d_n}$)

Example of applying automatic differentiation

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

$$w \mapsto w + e^w - \sin(w)$$

① Build a computational graph for computing $f(w)$



$\Rightarrow$ this graph is always a DAG (Directed Acyclic Graph)

② Compute $f(w)$ by doing a "forward pass" on the graph

Actual calculation of $f = f(w)$

$$v_1 = w$$

$$v_2 = e^{v_1}$$

$$v_3 = \sin(v_1)$$

$$v_4 = v_2 - v_3$$

$$v_5 = v_1 + v_4$$

$$f = v_5$$

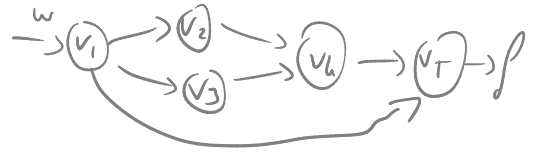
⊕

Store the dependencies between the v_i s that are linked by an arc

$$\frac{\partial v_i}{\partial v_j} \text{ for arc } (j,i)$$

→ After the forward pass, every node i contains $v_i(w)$ (a value)

③ Go through the graph backwards ("backward pass") to compute $\frac{\partial f}{\partial v_i} \forall i$:



stored at node v_5 → $\frac{\partial f}{\partial v_5} = 1$

$\frac{\partial f}{\partial v_4} = \frac{\partial f}{\partial v_5} \frac{\partial v_5}{\partial v_4} = 1 \times 1 = 1$

Information available after the forward pass

→ "Chain rule"

$v_5 = v_4 + v_4$

$\frac{\partial f}{\partial v_3} = \frac{\partial f}{\partial v_4} \frac{\partial v_4}{\partial v_3} = 1 \times (-1) = -1$

$v_4 = v_2 - v_3$

$\frac{\partial f}{\partial v_2} = \frac{\partial f}{\partial v_4} \times \frac{\partial v_4}{\partial v_2} = 1 \times 1 = 1$

$v_2 = e^{v_1}$
 $v_3 = \sin v_1$

$$\frac{\partial f}{\partial v_1} = \sum_{j: (1,j) \text{ arc}} \frac{\partial f}{\partial v_j} \times \frac{\partial v_j}{\partial v_1} = \frac{\partial f}{\partial v_2} \times \frac{\partial v_2}{\partial v_1} + \frac{\partial f}{\partial v_3} \times \frac{\partial v_3}{\partial v_1} + \frac{\partial f}{\partial v_5} \times \frac{\partial v_5}{\partial v_1}$$

$$= 1 \times e^{v_1} + -1 \times \cos(v_1) + 1 \times 1$$

$$= e^{v_1} - \cos(v_1) + 1$$

$$\frac{\partial f}{\partial w} = \frac{\partial f}{\partial v_1} \times \frac{\partial v_1}{\partial w} = e^{v_1} - \cos(v_1) + 1 = e^w - \cos(w) + 1$$

→ The approach extends to functions with multiple variables (and multiple outputs), as well as functions defined in a sequential

way (like NN)

$$x \rightarrow z_0 \mapsto \dots \mapsto z_L = NN(x)$$

→ AD computes $\frac{\partial z_L}{\partial z_{L-1}} \neq 1$

→ Easy to compute $\frac{\partial NN}{\partial x}$ by AD, and you also get the derivatives with respect to all parameters by doing so

Example: Two-layer neural network

$$f(x, y, W_1, W_2) = \frac{1}{2} \|W_2 \sigma(W_1 x) - y\|_2^2$$

Neural network part
 $x \mapsto W_2 \sigma(W_1 x)$
 h : hidden dimension

$$x \in \mathbb{R}^{d_x}, y \in \mathbb{R}^{d_y}, W_1 \in \mathbb{R}^{h \times d_x}, W_2 \in \mathbb{R}^{d_y \times h}$$

$$\sigma: \mathbb{R}^h \rightarrow \mathbb{R}^h$$

$$v \mapsto [\tanh(v_i)]_{i=1..h}$$

Goal: Compute $\frac{\partial f}{\partial x}$ and $\frac{\partial f}{\partial W_1}$

Forward pass (looks like the actual "forward" function in $\left\{ \begin{array}{l} \text{PyTorch} \\ \text{Jax} \end{array} \right\}$)

$$\mathbb{R}^h \rightarrow z_1 = W_1 x$$

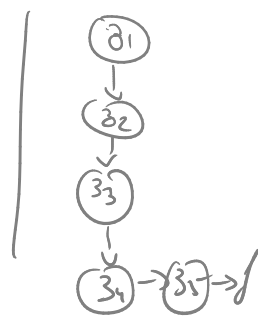
$$\mathbb{R}^h \rightarrow z_2 = \sigma(z_1)$$

$$\mathbb{R}^{d_y} \rightarrow z_3 = W_2 z_2$$

$$\mathbb{R}^{d_y} \rightarrow z_4 = z_3 - y$$

$$\mathbb{R} \rightarrow z_5 = \frac{1}{2} \|z_4\|_2^2$$

$$f = z_5$$



Backward pass

$$\frac{\partial f}{\partial z_5} = 1$$

$$\frac{\partial f}{\partial z_4} = \frac{\partial f}{\partial z_5} \frac{\partial z_5}{\partial z_4} = 1 \times z_4^T$$

$$z_4 \in \mathbb{R}^{d_y}$$

$$z_4^T \in \mathbb{R}^{1 \times d_y}$$

$$\frac{\partial f}{\partial z_3} = \frac{\partial f}{\partial z_4} \times \frac{\partial z_4}{\partial z_3} = z_4^T \in \mathbb{R}^{1 \times d_y}$$

$$I \in \mathbb{R}^{d_y \times d_y}$$

$$\begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}$$

$$\overbrace{z_3}^{\in \mathbb{R}^{d_y}} = W_2 \overbrace{z_2}^{\in \mathbb{R}^{d_h}}$$

$$\frac{\partial z_3}{\partial z_2} = W_2$$

$$\frac{\partial f}{\partial z_2} = \frac{\partial f}{\partial z_3} \times \frac{\partial z_3}{\partial z_2} = z_4^T W_2$$

$$\underbrace{1 \times d_y}_{1 \times h} \quad \underbrace{d_y \times h}_{1 \times h}$$

$$\frac{\partial f}{\partial z_1} = \frac{\partial f}{\partial z_2} \times \frac{\partial z_2}{\partial z_1} = z_4^T W_2 \Sigma$$

$$\underbrace{h \times h}$$

$$\Sigma = \begin{bmatrix} \sigma_1 & 0 \\ 0 & \sigma_2 \end{bmatrix}$$

$$\sigma_i = \tanh'([z_1]_i)$$

$$= 1 + \tanh^2([z_1]_i)$$

$$a) \frac{\partial f}{\partial x} = \frac{\partial f}{\partial z_1} \times \frac{\partial z_1}{\partial x} = z_4^T W_2 \Sigma W_1 = (W_2 \sigma(W, x) - y)^T W_2 \Sigma W_1$$

$$\underbrace{1 \times d_x}$$

$$b) \frac{\partial f}{\partial W_1} = \frac{\partial f}{\partial z_1} \times \frac{\partial z_1}{\partial W_1} = z_4^T W_2 \Sigma \frac{\partial z_1}{\partial W_1} = \Sigma W_2^T z_4 x^T$$

$$\underbrace{h \times h}_{h \times d_x} \quad \underbrace{d_y}_{h \times 1} \quad \underbrace{1 \times d_x}_{1 \times d_x}$$

$$W_1 \in \mathbb{R}^{h \times d_x}$$

Implicit layers in neural networks

↳ Standard layer choices in NNs are explicit

$$z_l = \tanh(A z_l + b)$$

which enables backpropagation.

↳ Implicit layer approach: Define a layer in terms of joint conditions involving the input and the output

→ De-couple what the layer does from the actual calculation

Examples

- Algebraic equation as a layer:

z_l computed from z_{l-1} , so that $c(z_l, z_{l-1}) = 0$
(Deep Equilibrium Models)

- Differential equations

$z_l = \varphi(1)$ where φ solves an ordinary differential equation

$$\begin{cases} \frac{\partial \varphi}{\partial t}(t) = f(\varphi(t)) & \forall t \in [0, 1] \\ \varphi(0) = z_{l-1} \end{cases}$$

(Neural ODE, 2018-)

- Optimization problem

$$z_l \in \underset{\substack{\uparrow \\ \text{set of} \\ \text{solutions}}}{\text{argmin}} \underbrace{f(z; z_{l-1})}_{\substack{\downarrow \\ \text{objective} \\ \text{function}}} \text{ s.t. } z \in \underbrace{\mathcal{Z}(z_{l-1})}_{\substack{\uparrow \\ \text{feasible} \\ \text{set}}}$$

"Optimization Layers"

Used in 2 ways:

① For problems that are easy to solve

$$\forall z \in \mathbb{R}^q, \text{ReLU}(z) = \left[\max(z_i, 0) \right]_{i=1..q}$$

$$= \underset{y}{\text{argmin}} \frac{1}{2} \|y - z\|_2^2 \text{ s.t. } y \geq z, y \geq 0$$

② For embedding optimization solvers in a ML pipeline

z_{l-1} : parameters/data for the solver

z_l : solution/output information from the solver

$\Rightarrow$ Want to train a NN on a task which involves solving an optimization problem through a solver (which you do not have explicit description of)

Key: Computing $\frac{\partial z_l}{\partial z_{l-1}}$!

$\rightarrow$ Involves optimality conditions for the optimization problem
 $\rightarrow$ Has to be done "by hand" because solvers are not part of the classical AD rules