

Intelligence artificielle

Jill-Jênn Vie

8 mai 2025

Jill-Jênn Vie

Thiers, Lyon, Cachan, agrég', Orsay, Pix, Tokyo, Inria Lille, Inria Saclay, X

Jill-Jênn Vie

Thiers, Lyon, Cachan, agrég', Orsay, Pix, Tokyo, Inria Lille, Inria Saclay, X

ICPC coach Ecole polytechnique (gold SWERC 2024), ENS Paris-Saclay, Université de Lille, Institut Polytechnique de Paris, CentraleSupélec

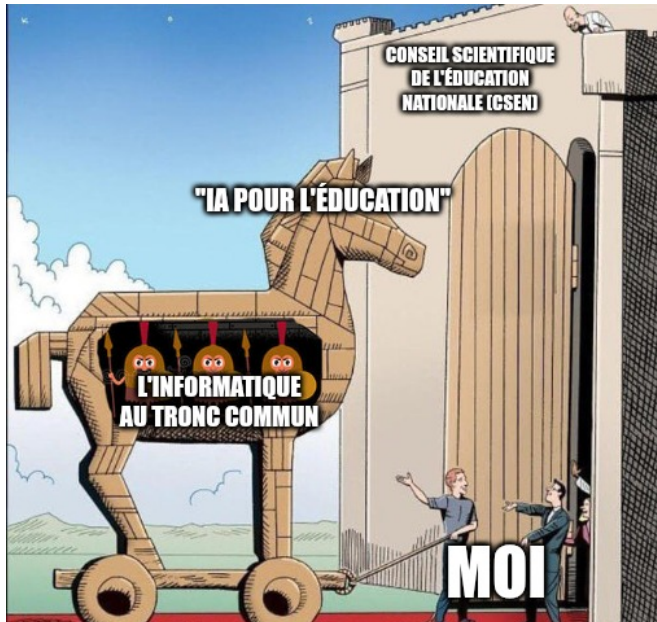


Équipe Soda à Inria Saclay : understanding society with data

IA en éducation, systèmes de recommandation (TP 1), tabular deep learning

scikit-learn : 3 millions de téléchargements par jour

J'ai rejoint le CSEN en 2013



Outline

Apprentissage

Cheval de troie 1

Cheval de troie 2

Démo : illustration de fairness (biais dans les décisions algorithmiques)

Bonus 1 : efficacité et impact environnemental des LLM

Bonus 2 : risques sociétaux des LLM, notamment pour l'éducation

IA

Une base de données : me renvoie les requêtes que je demande

Une IA : infère les points manquants par interpolation (« BayesDB »)

Apprend à partir d'exemples, ou collecte ses données par l'expérience (par renforcement)

Doit être capable de généraliser

Programme MPI / NSI

Apprentissage supervisé

- ▶ Algorithme des k plus proches voisins (NSI 1re)
- ▶ Matrice de confusion
- ▶ Arbres k -dimensionnels
- ▶ Arbre de décision (algorithme ID3)

Apprentissage non supervisé

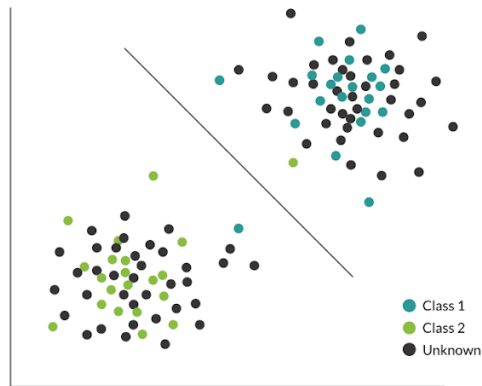
- ▶ Algorithme des k -moyennes
- ▶ Algorithme de classification hiérarchique ascendante

Jeux et heuristiques

- ▶ Jeux d'accessibilité à 2 joueurs sur un graphe
- ▶ Algorithme min-max avec une heuristique
- ▶ Élagage alpha-bêta
- ▶ Algorithme A* (ex. projet NSI)

Apprentissage supervisé

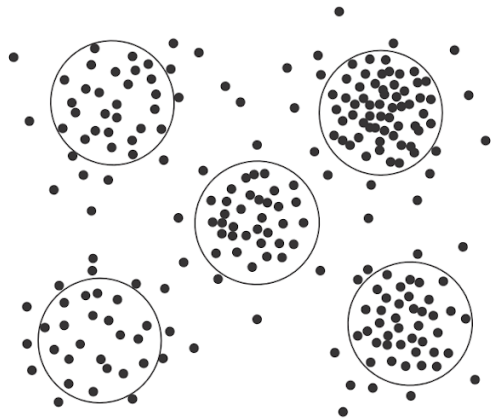
On observe des échantillons x, y
ex. classification, régression



```
model.fit(x_train, y_train)
model.predict(x_test)
```

Apprentissage non supervisé

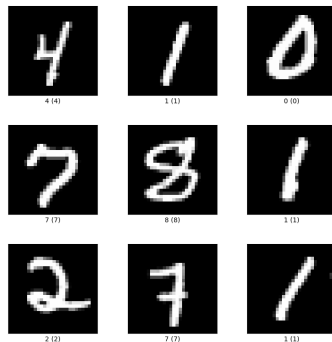
On observe seulement x
ex. texte brut, du son



```
x_tr = pca.fit_transform(x_train)
x_te = pca.transform(x_test)
```

Exemple : reconnaissance de chiffres manuscrits (TP 2)

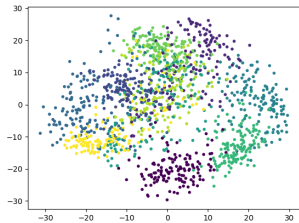
```
from sklearn.datasets import load_digits
```



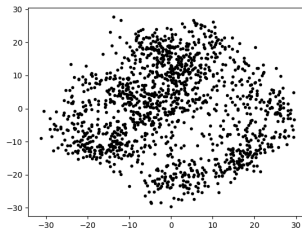
1797 images 8×8 en niveaux de gris, $X \in \mathbb{R}^{1797 \times 64}$, $y \in \{0, \dots, 9\}^{64}$

TP 2 : visualisation via projection en 2 dimensions

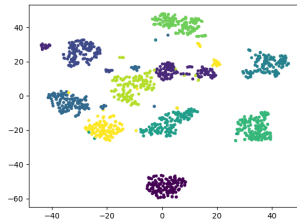
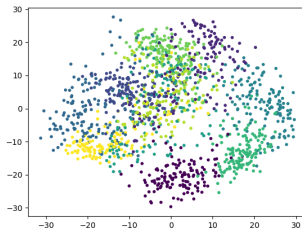
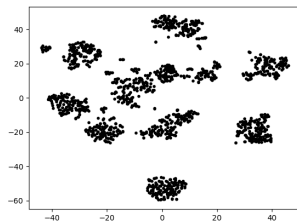
k plus proches voisins
(projection linéaire)



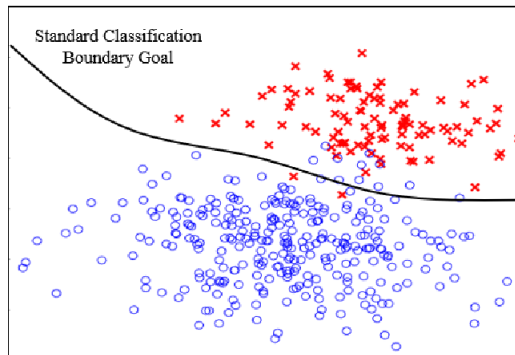
k -moyennes sur PCA
(projection linéaire)



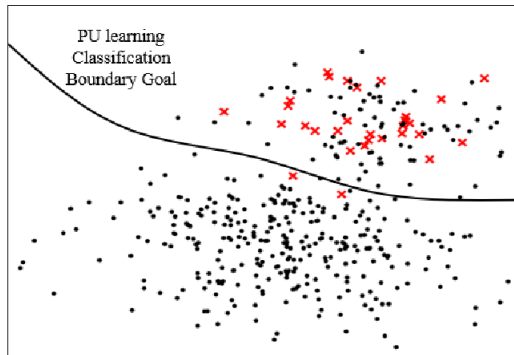
k -moyennes sur t-SNE
(non linéaire)



Apprentissage semi-supervisé : certains x n'ont pas de y



(a) Supervised classification learning problem with all labels known.



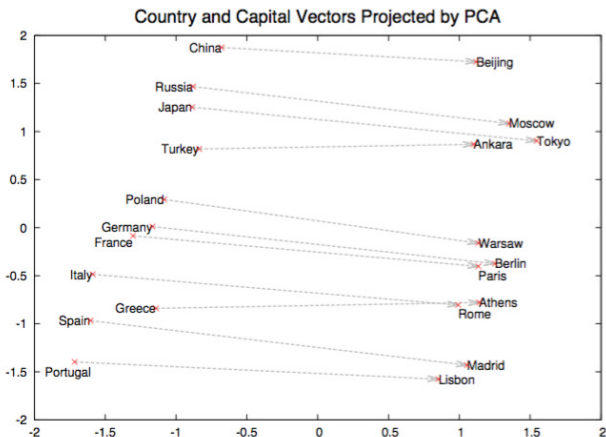
(b) Positive Unlabeled learning problem with only a percentage of known positive labels.

« PU » learning pour « positive unlabeled »

Apprentissage auto-supervisé

Prédire un mot à partir des mots autour (word2vec, Mikolov et al., 2013)

→ TP 3



Prédire le token suivant (GPT) → aussi TP 3

Apprentissage en contexte

Traditional ML: `fit(train)`, `predict(test)`

Transfer learning: `pretrain`, `finetune(train)`, `predict(test)` (e.g. word embeddings)

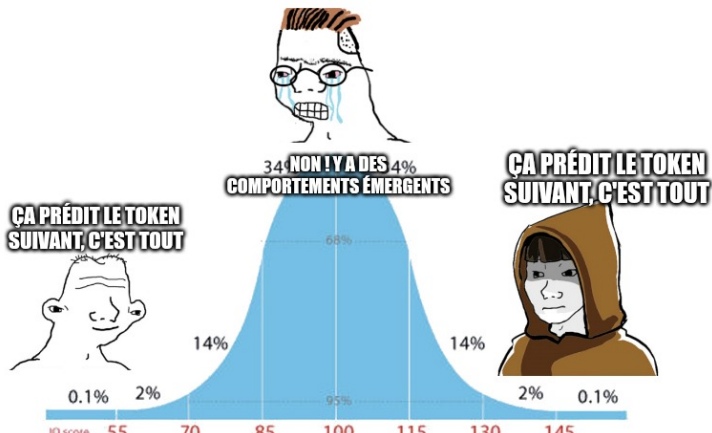
In-context learning: `pretrain (LLM)`, `predict(test, train)`

Apprentissage en contexte

Traditional ML: fit(train), predict(test)

Transfer learning: pretrain, finetune(train), predict(test) (e.g. word embeddings)

In-context learning: pretrain (LLM), predict(test, train)



LLMs are few shot learners

The three settings we explore for in-context learning

Zero-shot

The model predicts the answer given only a natural language description of the task. No gradient updates are performed.

```
1 Translate English to French: ← task description
2 cheese => ..... ← prompt
```

One-shot

In addition to the task description, the model sees a single example of the task. No gradient updates are performed.

```
1 Translate English to French: ← task description
2 sea otter => loutre de mer ← example
3 cheese => ..... ← prompt
```

Few-shot

In addition to the task description, the model sees a few examples of the task. No gradient updates are performed.

```
1 Translate English to French: ← task description
2 sea otter => loutre de mer ← examples
3 peppermint => menthe poivrée ←
4 plush girafe => girafe peluche ←
5 cheese => ..... ← prompt
```

Accurate predictions on small data with a tabular foundation model

<https://doi.org/10.1038/s41586-024-08328-6>

Received: 17 May 2024

Accepted: 31 October 2024

Published online: 8 January 2025

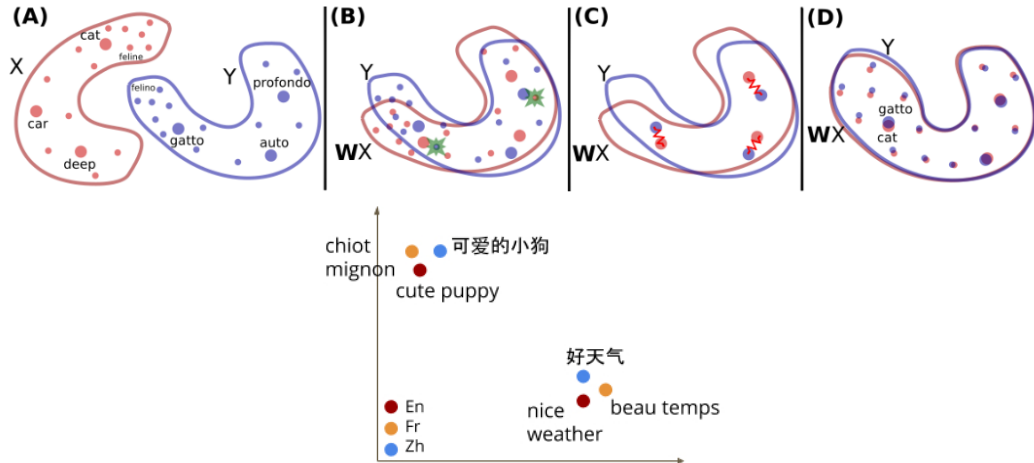
Open access

 Check for updates

Noah Hollmann^{1,2,3,7}, Samuel Müller^{1,7}, Lennart Purucker¹, Arjun Krishnakumar¹, Max Körfer¹, Shi Bin Hoo¹, Robin Tibor Schirrmeyer^{4,5} & Frank Hutter^{1,3,6}

Tabular data, spreadsheets organized in rows and columns, are ubiquitous across scientific fields, from biomedicine to particle physics to economics and climate science^{1,2}. The fundamental prediction task of filling in missing values of a label column based on the rest of the columns is essential for various applications as diverse as biomedical risk models, drug discovery and materials science. Although deep learning has revolutionized learning from raw data and led to numerous high-profile success stories^{3–5}, gradient-boosted decision trees^{6–9} have dominated tabular data for the past 20 years. Here we present the Tabular Prior-data Fitted Network (TabPFN), a tabular foundation model that outperforms all previous methods on datasets with up to 10,000 samples by a wide margin, using substantially less training time. **In 2.8 s, TabPFN outperforms an ensemble of the strongest baselines tuned for 4 h in a classification setting.** As a generative transformer-based foundation model, this model also allows fine-tuning, data generation, density estimation and learning reusable embeddings. TabPFN is a learning algorithm that is itself learned across millions of synthetic datasets, demonstrating the power of this approach for algorithm development. By improving modelling abilities across diverse fields, TabPFN has the potential to accelerate scientific discovery and enhance important decision-making in various domains.

Multilingual word embeddings (unsupervised)



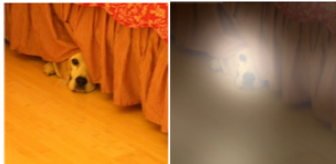
Guillaume Lample et al. (2018). « Word translation without parallel data ». In : *International Conference on Learning Representations*. URL : <https://arxiv.org/abs/1710.04087>

Légende automatique et attention (Xu and Bengio, 2015)

« C'est une boîte noire »



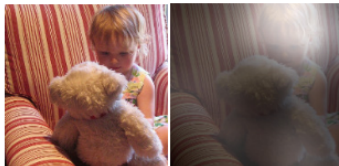
A woman is throwing a frisbee in a park.



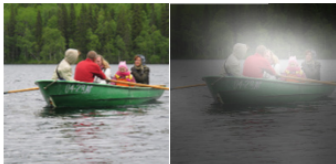
A dog is standing on a hardwood floor.



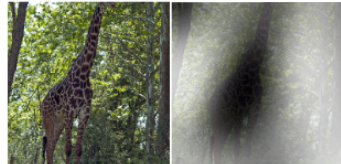
A stop sign is on a road with a mountain in the background.



A little girl sitting on a bed with a teddy bear.

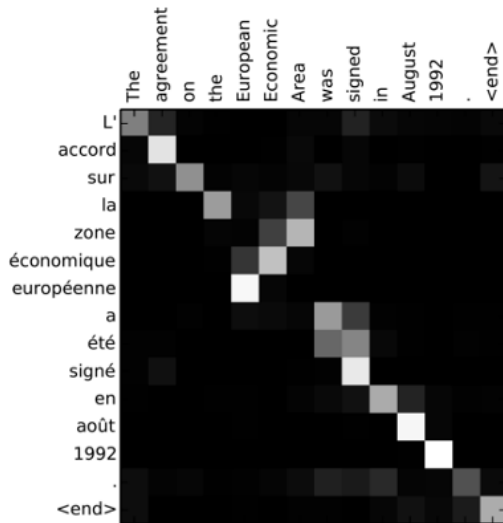


A group of people sitting on a boat in the water.



A giraffe standing in a forest with trees in the background.

Attention (Bahdanau, Cho and Bengio, 2014, Test of Time 2024)



k plus proches voisins (Russell et Norvig, 2020, p. 687)

Algorithme k plus proches voisins parmi n en d dimensions

pour tout i de 1 à n **faire**

 calculer la distance de $\mathbf{x}$ à $\mathbf{x}_i$

identifier les k plus proches voisins de $\mathbf{x}$

calculer la moyenne de leurs valeurs, ou la classe majoritaire

► Complexité de trouver les voisins de $\mathbf{x}$ dans X ? $O(nd + kn)$

Pour une version recommandation dans des bases de données, cf. (Abiteboul et al., 2011, p. 365)

k plus proches voisins (Russell et Norvig, 2020, p. 687)

Algorithme k plus proches voisins parmi n en d dimensions

pour tout i de 1 à n **faire**

 calculer la distance de $\mathbf{x}$ à $\mathbf{x}_i$

 identifier les k plus proches voisins de $\mathbf{x}$

 calculer la moyenne de leurs valeurs, ou la classe majoritaire

- ▶ Complexité de trouver les voisins de $\mathbf{x}$ dans X ? $O(nd + kn)$
- ▶ Avec *median of medians* : $O(nd)$

Pour une version recommandation dans des bases de données, cf. (Abiteboul et al., 2011, p. 365)

k plus proches voisins (Russell et Norvig, 2020, p. 687)

Algorithme k plus proches voisins parmi n en d dimensions

pour tout i de 1 à n **faire**

 calculer la distance de $\mathbf{x}$ à $\mathbf{x}_i$

 identifier les k plus proches voisins de $\mathbf{x}$

 calculer la moyenne de leurs valeurs, ou la classe majoritaire

- ▶ Complexité de trouver les voisins de $\mathbf{x}$ dans X ? $O(nd + kn)$
- ▶ Avec *median of medians* : $O(nd)$
- ▶ Peut être accéléré avec un K - d tree en $O((n + K) \log n)$

Pour une version recommandation dans des bases de données, cf. (Abiteboul et al., 2011, p. 365)

k plus proches voisins (Russell et Norvig, 2020, p. 687)

Algorithme k plus proches voisins parmi n en d dimensions

pour tout i de 1 à n **faire**

 calculer la distance de $\mathbf{x}$ à $\mathbf{x}_i$

identifier les k plus proches voisins de $\mathbf{x}$

calculer la moyenne de leurs valeurs, ou la classe majoritaire

- ▶ Complexité de trouver les voisins de $\mathbf{x}$ dans X ? $O(nd + kn)$
- ▶ Avec *median of medians* : $O(nd)$
- ▶ Peut être accéléré avec un K - d tree en $O((n + K) \log n)$
- ▶ Product quantization : $O(n + k \log k \log \log n)$ (algorithme d'approximation)

Pour une version recommandation dans des bases de données, cf. (Abiteboul et al., 2011, p. 365)

Product Quantization (2010) Vector Databases (2022)

Une structure de données pour calculer des k plus proches voisins $\operatorname{argmax}_{i \in S} q^T x_i$
(MIPS : *maximum inner product search*)

- ▶ K -dim tree
- ▶ en grande dimension, k plus proches voisins approchés (ex. FAISS) : *product quantization* (Jégou, Douze and Schmid, 2010), des objets proches doivent avoir des hash proches¹ (*locality-sensitive hashing*)

Variante : k plus proches voisins pondérés

Similarité $w_i = \mathbf{x}_{test}^T \mathbf{x}_{train_i} / N_z$ (renormalisés pour sommer à 1)

Je prédis pour $\mathbf{x}_{test}$: $\hat{y} = \sum_i w_i y_{train_i}$

Tiens mais l'attention c'est quoi ? $\operatorname{softmax}(QK^T)V$

Une combinaison linéaire de valeurs V à poids correspondant à des produits scalaires (similarités) entre la requête Q et les clés K

Comprendre les KNN, c'est une bonne étape pour comprendre l'attention

¹À ne pas utiliser pour hacher des mots de passe SVP

Récapitulatif et bonus

Apprentissage supervisé : caractéristiques X et étiquettes y

Apprentissage non supervisé : apprendre une représentation des données X , ex. réduire la dimension

Apprentissage par renforcement : apprendre en interagissant avec un environnement

Apprentissage semi-supervisé : certains X n'ont pas de y associé

Apprentissage auto-supervisé : prédire des parties de X à partir de lui-même

Apprentissage en contexte : prédire y_{test} à partir de $X_{test}, X_{train}, y_{train}$, pas d'entraînement, seulement de l'inférence

Méthodes paramétriques (réseaux de neurones)

Méthodes non paramétriques (k plus proches voisins, arbres de décision)

Formalisation apprentissage supervisé

Échantillons : $(\mathbf{x}_i, y_i) \in \mathbf{X} \times \mathcal{Y}$

Les caractéristiques $\mathbf{x}_i$ et les étiquettes y_i .

Habituellement $\mathbf{X} = \mathbb{R}^d$

Si $\mathcal{Y}$ est discret : classification, sinon régression.

Modèle $f_\theta : \mathbf{X} \rightarrow \mathcal{Y}$ ayant des paramètres θ

Une fonction d'erreur $\mathcal{L}$ à minimiser qui dépend des paramètres θ .

But : trouver $\theta = \operatorname{argmin}_\theta \mathcal{L}(\theta, \mathbf{x}_i, y_i)$

Régression linéaire (ATTENTION, HORS PROGRAMME !!!)

Modèle $f_{\theta}(\mathbf{x}_i) = \mathbf{x}_i^T \boldsymbol{\theta} \simeq y_i$

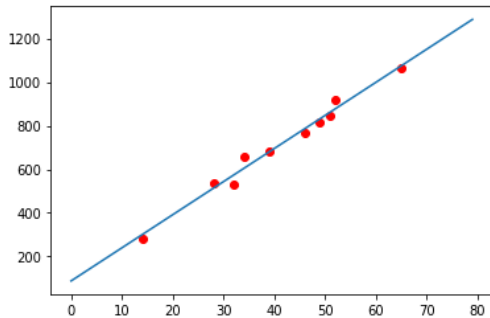
Erreur : moindres carrés

$$\mathcal{L} = \sum_i ||\mathbf{x}_i^T \boldsymbol{\theta} - y_i||^2$$

$$\frac{\partial \mathcal{L}}{\partial \boldsymbol{\theta}} = \sum_i 2\mathbf{x}_i(\mathbf{x}_i^T \boldsymbol{\theta} - y_i) = 0$$

$$\left(\sum_i \mathbf{x}_i \mathbf{x}_i^T \right) \boldsymbol{\theta} = \sum_i \mathbf{x}_i y_i$$

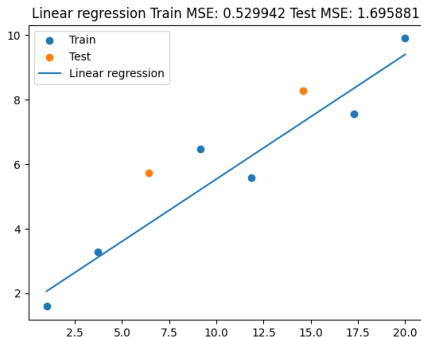
$$\boldsymbol{\theta} = (X^T X)^{-1} X^T \mathbf{y}$$



Risque de surapprentissage

Fitting

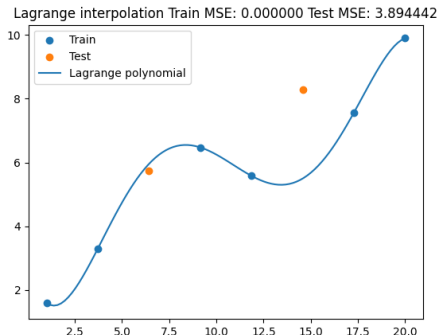
Polynomial degree 1 $Y = wX + b$



Linear regression
`scipy.stats.linreg`

Overfitting

Polynomial degree 6 $Y = \sum_{k=0}^6 w_k X^k$



Lagrange interpolation
`scipy.interpolation.lagrange`

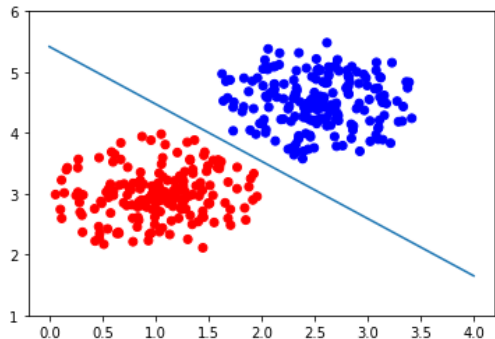
Régression logistique (classification)

$$f_{\theta}(\mathbf{x}_i) = \sigma(\mathbf{x}_i^T \boldsymbol{\theta}) = \Pr(Y_i = 1)$$

où $\sigma : x \mapsto 1/(1 + e^{-x})$

Pratique car $\sigma' = \sigma(1 - \sigma)$

Pas de formule close pour $\boldsymbol{\theta}$ donc
apprentissage de $\boldsymbol{\theta}$ par méthode de
Newton



Algorithme du perceptron (Rosenblatt, 1957) est une couche linéaire (remplacer σ par la fonction en escalier Heaviside)

Variante : reconnaissance de caractères

Remplacer σ par un softmax : régression logistique multinomiale

Réseau de neurones

$$\mathbf{x}^{(0)} = \mathbf{x}$$

$$\mathbf{x}^{(\ell+1)} = \sigma(\mathbf{W}^{(\ell)} \mathbf{x}^{(\ell)} + \mathbf{b}^{(\ell)}) \quad \ell = 0, \dots, L-2$$

$$y = \mathbf{x}^{(L)} = \mathbf{W}^{(L-1)} \mathbf{x}^{(L-1)} + \mathbf{b}^{(L-1)}$$

La ℓ -ième couche a d_ℓ neurones. Couche $\ell = 0$ d'entrée, $\ell = L$ de sortie.

σ est la fonction d'activation (non linéaire), ex. $\sigma = \text{ReLU} = \max(\mathbf{0}, \mathbf{x})$. Le code :

```
1 from torch import nn
2 mlp = nn.Sequential(
3     nn.Linear(d0, d1),
4     nn.ReLU(),
5     nn.Linear(d1, d2),
6     ...
7     nn.Linear(dL-1, dL)
8 )
```

Mais du coup ? Je peux le représenter comme $d_0 \rightarrow d_1 \rightarrow \dots \rightarrow d_L$

Réseau de neurones

$$\mathbf{x}^{(0)} = \mathbf{x} \in \mathbb{R}^{d_0}$$

$$\mathbf{x}^{(\ell+1)} = \sigma(\mathbf{W}^{(\ell)}\mathbf{x}^{(\ell)} + \mathbf{b}^{(\ell)}) \in \mathbb{R}^{d_\ell} \quad \ell = 0, \dots, L-2$$

$$y = \mathbf{x}^{(L)} = \mathbf{W}^{(L-1)}\mathbf{x}^{(L-1)} + \mathbf{b}^{(L-1)} \in \mathbb{R}^{d_L}$$

La ℓ -ième couche a d_ℓ neurones. Couche $\ell = 0$ d'entrée, $\ell = L$ de sortie.

σ est la fonction d'activation (non linéaire), ex. $\sigma = \text{ReLU} = \max(\mathbf{0}, \mathbf{x})$. Le code :

```
1 from torch import nn
2 mlp = nn.Sequential(
3     nn.Linear(d0, d1),
4     nn.ReLU(),
5     nn.Linear(d1, d2),
6     ...
7     nn.Linear(dL-1, dL)
8 )
```

Mais du coup ? Je peux le représenter comme $d_0 \rightarrow d_1 \rightarrow \dots \rightarrow d_L$

Cette notation simplifie cet exposé

Régression linéaire ou logistique : $n \rightarrow 1$

Perceptron pour la reconnaissance des chiffres :

Cette notation simplifie cet exposé

Régression linéaire ou logistique : $n \rightarrow 1$

Perceptron pour la reconnaissance des chiffres : $n \rightarrow 10$

L'apprentissage de représentations

Autoencodeur $n \rightarrow \dots \rightarrow k \rightarrow \dots \rightarrow n$

Analyse en composantes principales PCA :

Cette notation simplifie cet exposé

Régression linéaire ou logistique : $n \rightarrow 1$

Perceptron pour la reconnaissance des chiffres : $n \rightarrow 10$

L'apprentissage de représentations

Autoencodeur $n \rightarrow \dots \rightarrow k \rightarrow \dots \rightarrow n$

Analyse en composantes principales PCA : $n \rightarrow k \rightarrow n$ (sans les σ non linéaires)

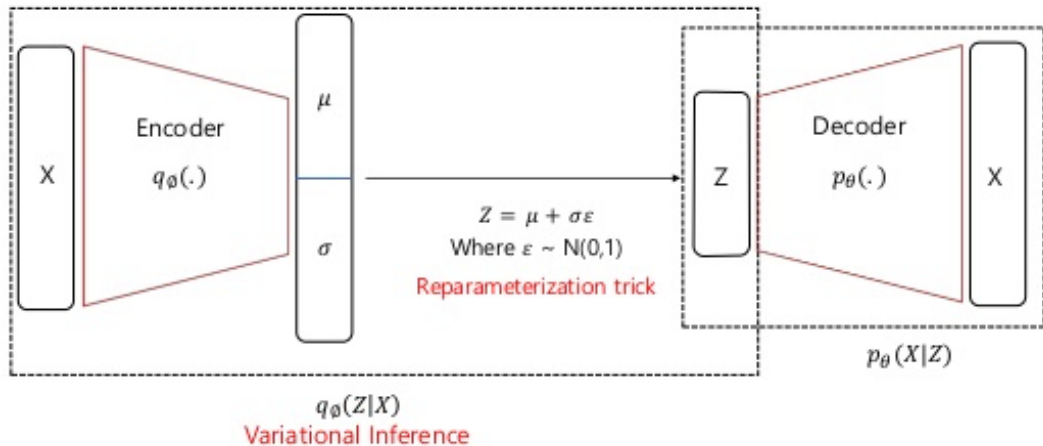
Ainsi, une PCA est un autoencodeur linéaire à une couche

(k -SVD : meilleure approximation de rang k , on tronque aux k plus grandes valeurs singulières, cf. TP 1 et TP 3)

Extrait du MOOC de Hinton (2012)

The image is a composite of three parts. On the left is a 10x10 grid of handwritten digits from the MNIST dataset, with the digit '5' highlighted in the bottom-left cell. Above this grid is a small 2x5 grid of numbers 0 through 9. In the center is a portrait of Geoffrey Hinton. On the right is a diagram illustrating a Restricted Boltzmann Machine (RBM) architecture. It shows a stack of layers: a top layer of random noise, a hidden layer, and a bottom layer of random noise. Bidirectional arrows connect the top layer to the hidden layer, and the hidden layer to the bottom layer. A single downward arrow points from the bottom layer to a square image of the digit '5'.

Variational autoencoders (Kingma 2014)



Contrôler la génération par une classe, ou un prompt

Generative Adversarial Networks (Goodfellow, 2014, Test of Time 2024)

Un étudiant de master à l'université de Kyoto (Yanghua Jin) a fait ceci :



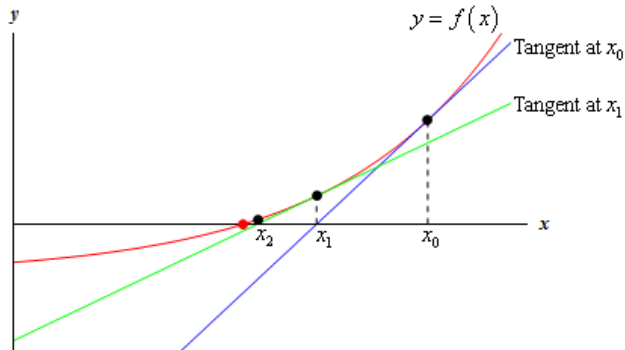
Why SGD?

To minimize a function

Find the zeroes of its derivative

How to find the zeroes of a function?

Newton's method: find x such that $f(x) = 0$


$$f : \mathbb{R} \rightarrow \mathbb{R} \text{ differentiable, } \nexists x, f'(x) = 0$$

$$x_{t+1} = x_t - \frac{f(x_t)}{f'(x_t)}$$

Quadratic convergence

$$\exists C > 0, |x_{t+1} - \ell| \leq C|x_t - \ell|^2$$

In higher dimension

Let $g : \mathbb{R}^n \rightarrow \mathbb{R}$ twice differentiable, with n big

In higher dimension

Let $g : \mathbb{R}^n \rightarrow \mathbb{R}$ twice differentiable, with n big

What is the size of $g'(\mathbf{x})$? Usually noted $\frac{\partial g}{\partial \mathbf{x}}$ or $\nabla_{\mathbf{x}} g$.

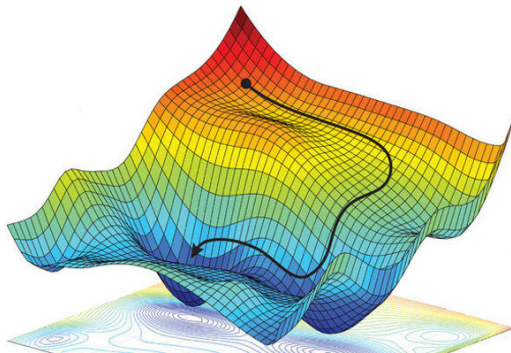
In higher dimension

Let $g : \mathbb{R}^n \rightarrow \mathbb{R}$ twice differentiable, with n big

What is the size of $g'(\mathbf{x})$? Usually noted $\frac{\partial g}{\partial \mathbf{x}}$ or $\nabla_{\mathbf{x}} g$.

$$\mathbf{x}_{t+1} = \mathbf{x}_t - \underbrace{g''(\mathbf{x}_t)^{-1}}_{\in \mathbb{R}^{n \times n}, O(n^3)} \underbrace{g'(\mathbf{x}_t)}_{\in \mathbb{R}^n}$$

Gradient descent



$$\mathbf{x}_{t+1} = \mathbf{x}_t - \gamma \nabla_{\mathbf{x}} \mathcal{L}(\mathbf{x}_t)$$

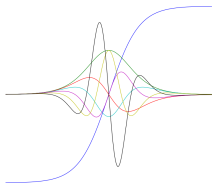
```
1 for each epoch:  
2   for  $\mathbf{x}$ ,  $y$  in dataset:  
3     compute gradients  $\frac{\partial \mathcal{L}}{\partial \theta}(f_{\theta}(\mathbf{x}), y)$  # also noted  $\nabla_{\theta} \mathcal{L}$   
4      $\theta \leftarrow \theta - \gamma \nabla_{\theta} \mathcal{L}$  #  $\gamma$  is the learning rate
```

SGD in PyTorch

```
1 import torch
2
3 # define model, n_epochs, trainloader
4 criterion = torch.nn.CrossEntropyLoss()
5 optimizer = torch.optim.SGD(model.parameters(), lr=0.001)
6
7 for _ in range(n_epochs):
8     for batch_inputs, batch_labels in trainloader:
9         outputs = model(batch_inputs)
10        loss = criterion(outputs, batch_labels)
11
12        optimizer.zero_grad()
13        loss.backward()
14        optimizer.step()
```

How to compute gradients automatically? What's in `loss.backward()`?

```
1 from autograd import elementwise_grad as egrad
2 import matplotlib.pyplot as plt
3 x = np.linspace(-7, 7, 200)
4 plt.plot(x, tanh(x),
5          x, egrad(tanh)(x),                # first
6          x, egrad(egrad(tanh))(x),         # second
7          x, egrad(egrad(egrad(tanh)))(x),  # third
8          x, egrad(egrad(egrad(egrad(tanh))))(x), # fourth
9          x, egrad(egrad(egrad(egrad(egrad(tanh)))))(x), # fifth
10         x, egrad(egrad(egrad(egrad(egrad(egrad(tanh)))))(x)) # sixth
11 plt.show()
```



Existing methods and their limitations

Numerical differentiation

$$\frac{f(x+h) - f(x)}{h}$$

Round-off errors

Existing methods and their limitations

Numerical differentiation

$$\frac{f(x+h) - f(x)}{h}$$

Round-off errors

Symbolic differentiation

Have to keep symbolic expressions at each step of the process

Existing methods and their limitations

Numerical differentiation

$$\frac{f(x+h) - f(x)}{h}$$

Round-off errors

Symbolic differentiation

Have to keep symbolic expressions at each step of the process

Automatic differentiation

Chain rule

$$(f \circ g)' = g' \cdot (f' \circ g)$$

Generalized chain rule

$$\frac{df_1}{dx} = \frac{df_1}{df_2} \frac{df_2}{df_3} \cdots \frac{df_n}{dx}$$

Reminder: Jacobians

Consider differentiable $f : \mathbb{R}^n \rightarrow \mathbb{R}^m$, its Jacobian $J_f \in \mathbb{R}^{m \times n}$ contains its first-order partial derivatives $(J_f)_{ij} = \frac{\partial f_i}{\partial x_j}$:

$$J_f = \begin{bmatrix} \frac{\partial f}{\partial x_1} & \cdots & \frac{\partial f}{\partial x_n} \end{bmatrix} = \begin{bmatrix} \nabla^T f_1 \\ \vdots \\ \nabla^T f_m \end{bmatrix} = \begin{bmatrix} \frac{\partial f_1}{\partial x_1} & \cdots & \frac{\partial f_1}{\partial x_n} \\ \vdots & \ddots & \vdots \\ \frac{\partial f_m}{\partial x_1} & \cdots & \frac{\partial f_m}{\partial x_n} \end{bmatrix}$$

Reminder: Jacobians

Consider differentiable $f : \mathbb{R}^n \rightarrow \mathbb{R}^m$, its Jacobian $J_f \in \mathbb{R}^{m \times n}$ contains its first-order partial derivatives $(J_f)_{ij} = \frac{\partial f_i}{\partial x_j}$:

$$J_f = \begin{bmatrix} \frac{\partial f}{\partial x_1} & \cdots & \frac{\partial f}{\partial x_n} \end{bmatrix} = \begin{bmatrix} \nabla^T f_1 \\ \vdots \\ \nabla^T f_m \end{bmatrix} = \begin{bmatrix} \frac{\partial f_1}{\partial x_1} & \cdots & \frac{\partial f_1}{\partial x_n} \\ \vdots & \ddots & \vdots \\ \frac{\partial f_m}{\partial x_1} & \cdots & \frac{\partial f_m}{\partial x_n} \end{bmatrix}$$

Example: linear layer

W is a $m \times n$ matrix

If $f(\mathbf{x}) = W\mathbf{x}$

$f_i = W_i \mathbf{x} = \sum_j W_{ij} x_j$ where W_i is i th row of W

$(J_f)_{ij} = \frac{\partial f_i}{\partial x_j} = W_{ij}$ so $J_f = W$

Generalized multivariate chain rule

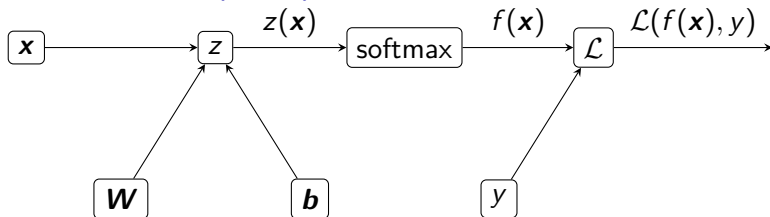
Consider differentiable $f : \mathbb{R}^m \rightarrow \mathbb{R}^k$, $g : \mathbb{R}^n \rightarrow \mathbb{R}^m$ and $\mathbf{a} \in \mathbb{R}^n$.

$$D_{\mathbf{a}}(f \circ g) = D_{g(\mathbf{a})}f \circ D_{\mathbf{a}}g$$

So the Jacobians verify:

$$J_{f \circ g} = (J_f \circ g)J_g$$

Graphe orienté acyclique (DAG) de calcul dérivable

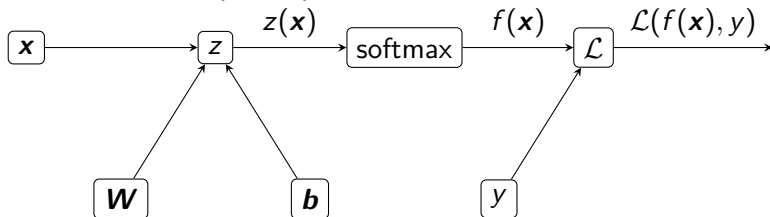


$$\begin{aligned}\frac{d\mathcal{L}}{db_1} &= \frac{d\mathcal{L}}{df} \frac{df}{db_1} = \frac{d\mathcal{L}}{df} \left(\frac{df}{dz} \frac{dz}{db_1} \right) \quad (\text{forward}) \\ &= \frac{d\mathcal{L}}{dz} \frac{dz}{db_1} = \left(\frac{d\mathcal{L}}{df} \frac{df}{dz} \right) \frac{dz}{db_1} \quad (\text{backward})\end{aligned}$$

Properly written: $J_{\mathcal{L} \circ \text{softmax} \circ z} = (J_{\mathcal{L}} \circ f)(J_{\text{softmax}} \circ z)J_z$

Given that $\mathbf{x} \in \mathbb{R}^d$, $z(\mathbf{x}), f(\mathbf{x}) \in \mathbb{R}^{d_2}$, $\mathcal{L}(f(\mathbf{x}), y) \in \mathbb{R}$,
which order is better?

Graphe orienté acyclique (DAG) de calcul dérivable



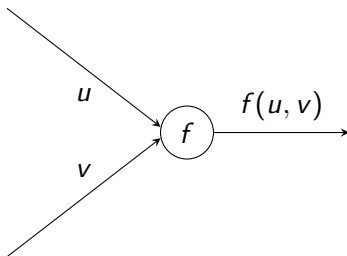
$$\begin{aligned}\frac{d\mathcal{L}}{db_1} &= \frac{d\mathcal{L}}{df} \frac{df}{db_1} = \frac{d\mathcal{L}}{df} \left(\frac{df}{dz} \frac{dz}{db_1} \right) \quad (\text{forward}) \\ &= \frac{d\mathcal{L}}{dz} \frac{dz}{db_1} = \left(\frac{d\mathcal{L}}{df} \frac{df}{dz} \right) \frac{dz}{db_1} \quad (\text{backward})\end{aligned}$$

Properly written: $J_{\mathcal{L} \circ \text{softmax} \circ z} = (J_{\mathcal{L}} \circ f)(J_{\text{softmax}} \circ z)J_z$

Given that $\mathbf{x} \in \mathbb{R}^d$, $z(\mathbf{x})$, $f(\mathbf{x}) \in \mathbb{R}^{d_2}$, $\mathcal{L}(f(\mathbf{x}), y) \in \mathbb{R}$,
which order is better? This is why **backpropagation**.

Reverse accumulation (in $\mathbb{R}$)

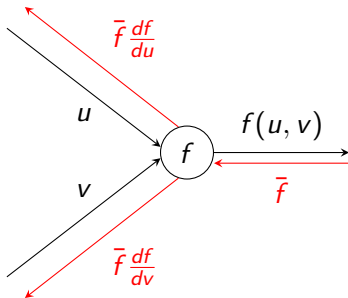
Let us note the adjoint $\bar{f} \triangleq \frac{d\mathcal{L}}{df}$.



Éviter de recalculer deux fois la même chose (mémoïsation)

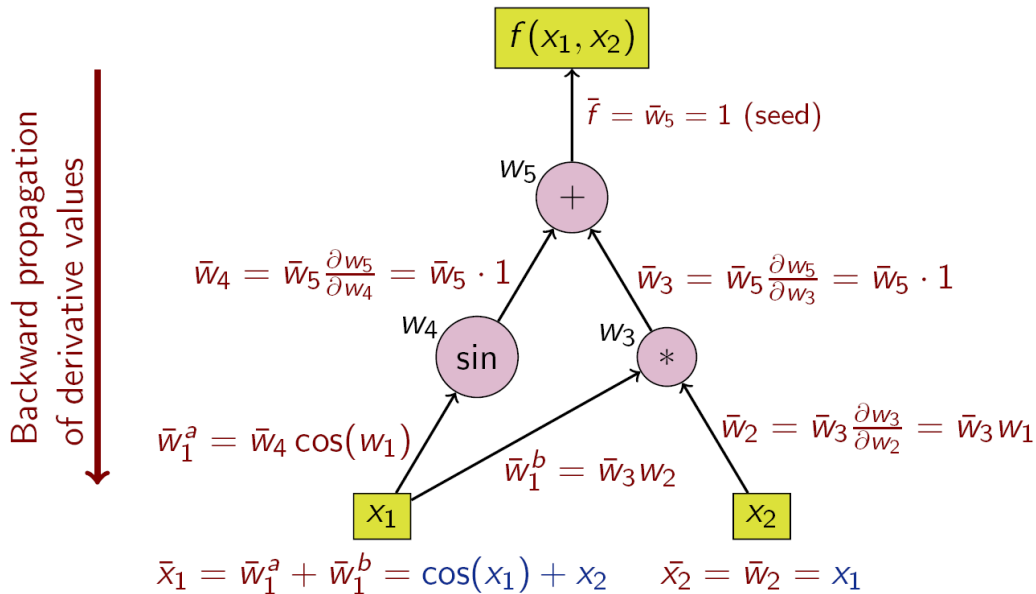
Reverse accumulation (in $\mathbb{R}$)

Let us note the adjoint $\bar{f} \triangleq \frac{d\mathcal{L}}{df}$.



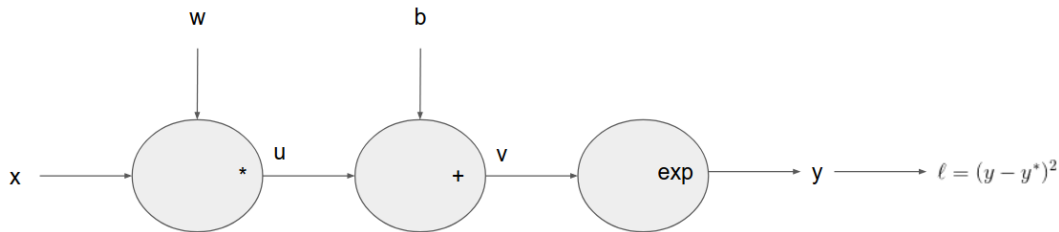
Éviter de recalculer deux fois la même chose (mémoïsation)

A complete example (Wikipedia)



Notebook: compute gradients

FORWARD PASS



$$u = wx$$

$$v = u + b$$

$$y = e^v$$

$$\ell = (y - y^*)^2$$

Question (piège ?)

Pourquoi A* c'est de l'IA ?

Question (piège ?)

Pourquoi A^* c'est de l'IA ?

Est-ce que l'algorithme de Dijkstra c'est de l'IA aussi ?

Parce que les deux sont Best-first search

Algorithme Recherche du meilleur en premier

Mettre la *source* dans la file

tant que la file n'est pas vide **faire**

Extraire le nœud u ayant la priorité $f(u)$ minimale

si u est la *cible* **alors renvoyer** dist, prec

pour tout voisin v du nœud u **faire**

 candidat $\leftarrow \text{dist}(u) + \text{poids arête } w_{uv}$

si c'est un meilleur candidat i.e. candidat $< \text{dist}(\text{voisin})$ **alors**

 dist(voisin v) $\leftarrow$ candidat

 prec(voisin v) $\leftarrow$ nœud u

 Ajouter v à la file avec priorité $f(v)$

Valeurs de priorité $f(u)$

- ▶ Dijkstra : parcourir par plus courte distance(source, nœud u)
- ▶ *Greedy best-first search* : $h(u)$ distance à vol d'oiseau à l'arrivée
- ▶ A* : par $\text{dist}(u) + h(u)$ croissant.
- ▶ Prim (arbre couvrant) : distance à l'arbre en cours

Principe d'optimalité de Bellman

An optimal policy has the property that whatever the initial state and initial decision are, the remaining decisions must constitute an optimal policy with regard to the state resulting from the first decision.

Équation de Bellman

Soit un état s , on choisit une action a qui nous rapporte une récompense $R(s, a)$ et nous met dans un état s' . V indique la récompense moyenne obtenue si l'on agit de façon optimale.

$$V(s) = \max_a R(s, a) + \gamma V(s')$$

Ça vous rappelle quelque chose ?

Principe d'optimalité de Bellman

An optimal policy has the property that whatever the initial state and initial decision are, the remaining decisions must constitute an optimal policy with regard to the state resulting from the first decision.

Équation de Bellman

Soit un état s , on choisit une action a qui nous rapporte une récompense $R(s, a)$ et nous met dans un état s' . V indique la récompense moyenne obtenue si l'on agit de façon optimale.

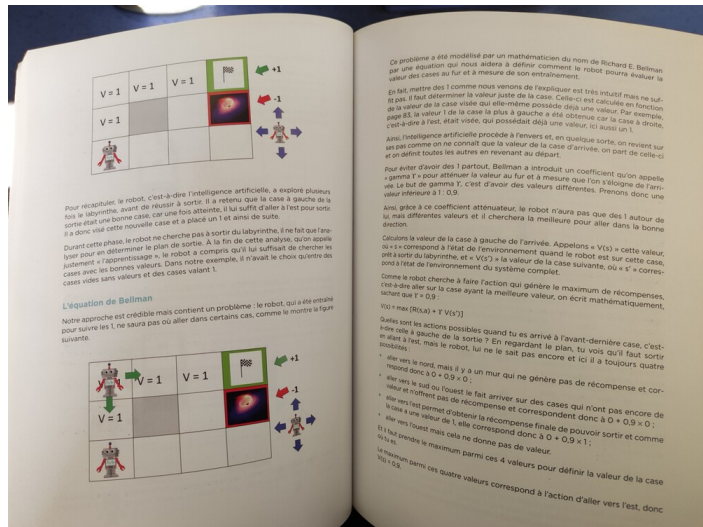
$$V(s) = \max_a R(s, a) + \gamma V(s')$$

Ça vous rappelle quelque chose ?

Bellman-Ford

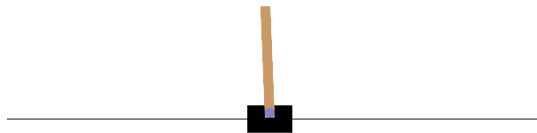
$$V(u) = \max_v -w_{u,v} + V(v) \quad V = -d$$

Trouvé dans une médiathèque



Exemples d'environnements d'apprentissage par renforcement

CartPole-v1

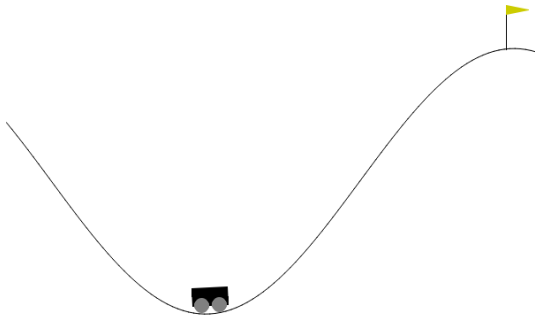


$$S = (x, \dot{x}, \theta, \dot{\theta}) \in \mathbb{R}^4$$

$$A \in \{\leftarrow, \rightarrow\}$$

$$R = 1$$

MountainCar-v2



$$S = (x, \dot{x}) \in \mathbb{R}^2$$

$$A \in \{\leftarrow, 0, \rightarrow\}$$

$$R = -1$$

Dans le même bouquin, quelques pages plus loin

Code une voiture autonome avec Python

Dans cet exercice, tu vas jouer en ligne et entraîner une voiture autonome. Comment vas-tu procéder ? Avant le début de la course sur le circuit, tu pourras activer un enregistrement qui récupérera les multiples images de la vidéo de ton jeu qui serviront à générer un modèle basé sur un réseau de neurones à convolution. Te rappelles-tu des réseaux spécialisés pour la reconnaissance d'images ?

Nous nous appuyerons sur les outils développés par NVIDIA, une société américaine spécialisée dans la fabrication de cartes graphiques qui développe depuis 2017 de nouvelles cartes pour l'intelligence artificielle.

Les voitures autonomes sont l'un des domaines de recherche et d'activité les plus dynamiques. Ce qui semblait relever de la



Gold mine problem (Bellman, 1952)

Two gold mines, Anaconda (amount x) and Bonanza (amount y).

- ▶ Anaconda: probability p_1 to collect r_1x $1 - p_1$ to break the machine forever
- ▶ Bonanza: probability q_1 to collect r_2y $1 - q_1$ to break the machine forever

Parameters of the problem: $p_1, q_1, r_1, r_2 \in [0, 1]$, $x, y \in \mathbb{R}^+$.

What is the optimal action, that maximizes the amount of extracted gold?

Gold mine problem (Bellman, 1952)

Two gold mines, Anaconda (amount x) and Bonanza (amount y).

- ▶ Anaconda: probability p_1 to collect r_1x $1 - p_1$ to break the machine forever
- ▶ Bonanza: probability q_1 to collect r_2y $1 - q_1$ to break the machine forever

Parameters of the problem: $p_1, q_1, r_1, r_2 \in [0, 1]$, $x, y \in \mathbb{R}^+$.

What is the optimal action, that maximizes the amount of extracted gold?

Let $f(x, y)$ be the maximum expected gold extracted.

Solution to the gold mine problem

$$f(x, y) = \max \left\{ \begin{array}{l} p_1(r_1x + f((1 - r_1)x, y)) \\ q_1(r_2y + f(x, (1 - r_2)y)) \end{array} \right\}.$$

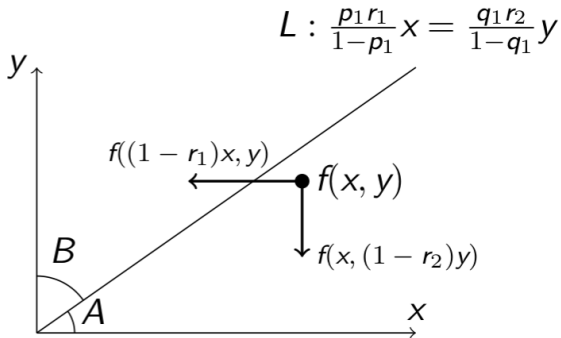
Solution

Si $\frac{p_1 r_1}{1 - p_1} x > \frac{q_1 r_2}{1 - q_1} y$,

Choisir A

Sinon

Choisir B.

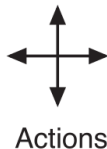
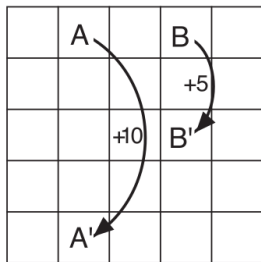


Policy evaluation on Gridworld (not a DAG)

$$S \in \{(i,j)\}_{1 \leq i,j \leq 5}$$

$$A \in \{\leftarrow, \rightarrow, \uparrow, \downarrow\}$$

$$R = \begin{cases} 10 & \text{if leaving from A} \\ 5 & \text{if leaving from B} \\ -1 & \text{if hitting a wall} \\ 0 & \text{otherwise.} \end{cases}$$



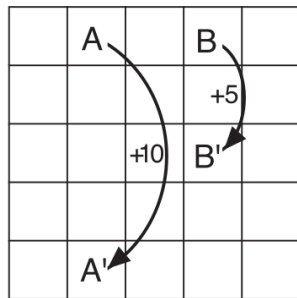
3.3	8.8	4.4	5.3	1.5
1.5	3.0	2.3	1.9	0.5
0.1	0.7	0.7	0.4	-0.4
-1.0	-0.4	-0.4	-0.6	-1.2
-1.9	-1.3	-1.2	-1.4	-2.0

Gridworld

$$\pi_{\text{uniform}} = (1/4, 1/4, 1/4, 1/4)$$

$$V_{\pi_{\text{uniform}}}$$

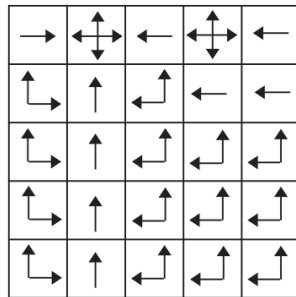
Policy improvement on Gridworld



Gridworld

22.0	24.4	22.0	19.4	17.5
19.8	22.0	19.8	17.8	16.0
17.8	19.8	17.8	16.0	14.4
16.0	17.8	16.0	14.4	13.0
14.4	16.0	14.4	13.0	11.7

v_*



π_*

La programmation dynamique de Richard Bellman (1920-1984)

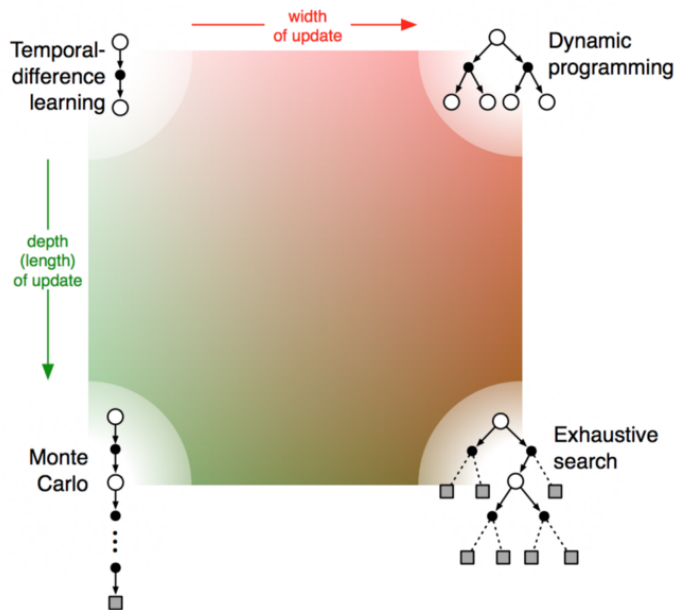


1952 : Richard Bellman publie *Programmation dynamique*

1952 : Autocode, le premier langage de programmation

Il avait la présence d'esprit de résoudre des problèmes (de maths appliquées) par prises de décision optimales successives (*multistage decision processes*), de façon constructive

La programmation dynamique a mené à l'apprentissage par renforcement



This basic equation can be used to study optimal trajectories, including launching, geodesics, paths of minimum length through discrete networks, r -th shortest paths through networks, soft landing on the moon and other planets, chemical process control, reactor control, and many other problems and problems.

Questions of computational solution will be discussed in the next chapter.

4.17 The Brachistochrone

As an example of a problem in the calculus of variations which can be explicitly resolved, either by classical techniques, or by means of the functional equation technique which we shall pursue here, let us consider the problem of determining a curve connecting the two points P and Q with the property that a particle sliding along this curve under the influence of gravity will arrive at Q in a minimum time.

Without loss of generality, let P be the origin, and let the axes be oriented as indicated below.

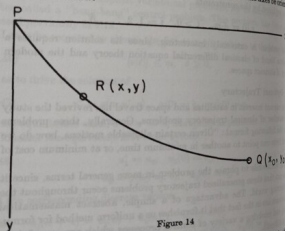


Figure 14

Omitting the gravitational constant, the problem is that of minimizing the functional

$$(4.56) \quad J(y) = \int_0^{x_0} \frac{(1+y'^2)^{1/2}}{y} dx,$$

over all curves satisfying the end conditions

$$(4.57) \quad y(0) = 0, \quad y(x_0) = y_0.$$

THE BRACHISTOCURVE

To treat the problem by dynamic programming techniques, we let $f(x, y)$ denote the minimum time required to go from the generic point $R(x, y)$ to (x_0, y_0) . The equation in (4.56) is in this case

$$(4.58) \quad f(x, y) = \min_{y'} \left[\left(\frac{1+y'^2}{y} \right)^{1/2} \Delta + f(x + \Delta, y + y' \Delta) + O(\Delta^2) \right].$$

Passing to the limit as $\Delta \rightarrow 0$, this yields the nonlinear partial differential equation

$$(4.59) \quad 0 = \min_{y'} \left[\left(\frac{1+y'^2}{y} \right)^{1/2} + f_x + y' f_y \right],$$

equivalent to the two equations

$$(4.60) \quad (a) \quad 0 = \left(\frac{1+y'^2}{y} \right)^{1/2} + f_x + y' f_y,$$

$$(b) \quad 0 = \frac{y'}{[y(1+y'^2)]^{1/2}} + f_y.$$

To eliminate the function f , and thus obtain an equation for y' , we differentiate (4.60a) with respect to x and take the partial derivative of (4.60a) with respect to y , obtaining

$$(4.61) \quad (a) \quad \frac{d}{dx} \left[\frac{y'}{[y(1+y'^2)]^{1/2}} \right] + f_{xx} + f_{xy} y' = 0,$$

$$(b) \quad -\frac{1}{2} (1+y'^2)^{-3/2} + f_{xy} + f_{yy} y' = 0.$$

Thus

$$(4.62) \quad \frac{d}{dx} \left[\frac{y'}{[y(1+y'^2)]^{1/2}} \right] + \frac{1}{2} (1+y'^2)^{-3/2} = 0.$$

A first integral is

$$(4.63) \quad \left(\frac{1+y'^2}{y} \right)^{1/2} - \frac{y'^2}{[(1+y'^2)y]^{1/2}} = k,$$

or

$$(4.64) \quad \frac{1}{[y(1+y'^2)]^{1/2}} = k.$$

This is equivalent to Snell's law for the propagation of light. It is interesting to interpret this well-known physical principle as an optimal policy in a multistage decision process.

The integration can now be readily completed to obtain the standard representation of a brachistochrone in parametric form:

$$(4.65) \quad \begin{aligned} y &= (1 - \cos t)/2k^2, \\ x &= c_1 + (t - \sin t)/2k^2. \end{aligned}$$

* We have deliberately kept y' as the slope, instead of v , in order to manipulate the Euler equation.

Les grands modèles de langage (LLM) : quels objectifs ?

1. Transformer : prédire le mot suivant

Transformers / are / a / new / machine / [learning]

Transformers / are / a / new / machine / learning / [architecture]

2. Données de démonstration d'experts :

Query: put the first letters in uppercase in "optimizing human learning"

Answer: Optimizing Human Learning

3. Données de comparaison d'experts :

Query: write a poem

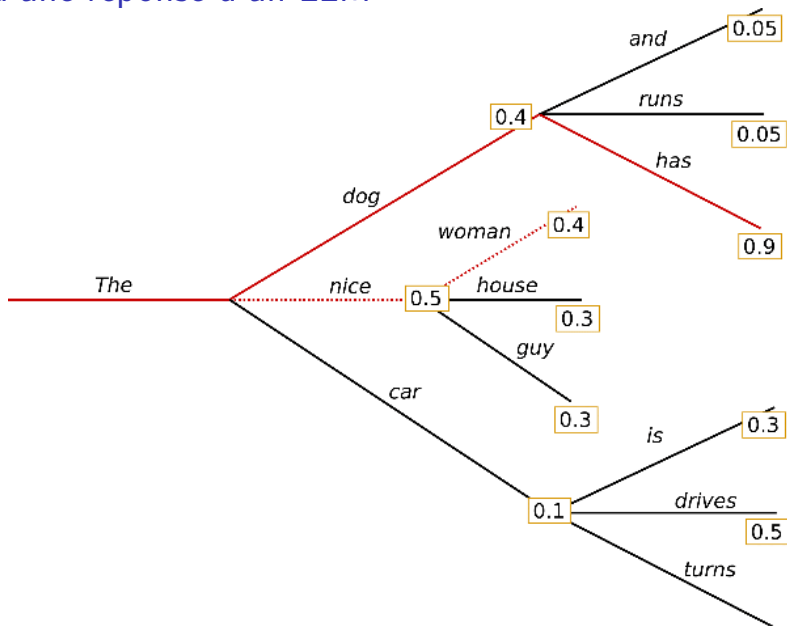
Answer 1: Roses are red

Answer 2: Once upon a time, a prince in a castle

Où la réponse 2 est votée favorablement par les experts

ChatGPT est donc entraîné pour plaire aux experts et généraliser à de nouvelles tâches

Recherche d'une réponse d'un LLM



La tâche est dans les données d'entrée

Machine de Turing : le programme est dans les données

Machine de Turing universelle : exécuter le i -ème programme

François Chollet dit ça beaucoup mieux :

My mental model for LLMs is that they work as a repository of vector programs. When prompted, they will fetch the program that your prompt maps to and « execute » it on the input at hand. LLMs are a way to store and operationalize millions of useful mini-programs via passive exposure to human-generated content.

If a LLM is like a database of millions of vector programs, then a prompt is like a search query in that database [...] this « program database » is continuous and interpolative — it's not a discrete set of programs.

Sources: <https://arcprize.org/blog/oai-o3-pub-breakthrough>,
<https://simonwillison.net/2023/Oct/25/francois-chollet/>

Take home message

Nombreuses applications de l'IA, beaucoup d'impacts sociétaux, essentiel pour faire rêver les jeunes (ou les inquiéter)

On peut regretter que l'informatique et l'IA ne s'entendent pas bien alors que les deux s'intéressent à des algorithmes qui optimisent un objectif

Je mise tout sur les *vector databases* et Bellman pour les réconcilier

Analyse (du second degré)

Le ou la prof de maths fait les fonctions à plusieurs variables, jacobienues, vraiment à la fin de l'année parce que flemme et de toute façon il n'a pas l'intention de le coder

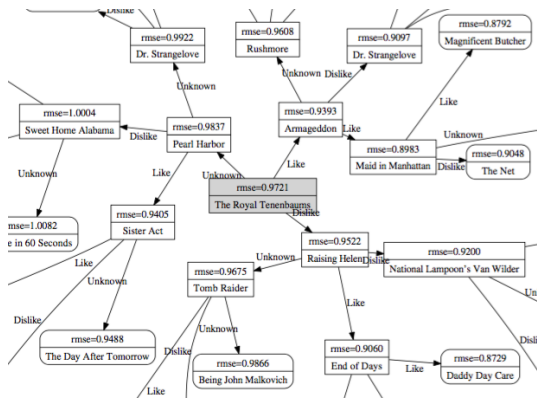
Le ou la prof d'info fait l'IA à la fin parce que « l'IA ça marche pas et de toute façon ça ne tombe jamais aux concours » (pourtant, X 2025) et 1) pf faut installer numpy sur les ordi 2) en SQL, C, OCaml, l'IA c'est infernal

TP 1 – Systèmes de recommandation : KNN, SVD, ALS, SGD

TP 2 – Reconnaissance de motifs OU Gridworld OU Quoridor

TP 3 – Traitement automatique des langues : chaîne de Markov, GPT-based LLM

On pourrait adapter le TP 1 en posant des questions informatives sur les goûts (à la ID3)



Merci pour votre attention

Si vous souhaitez parler de l'environnement de concours NonOS, n'hésitez pas, je ne mords pas

-  Abiteboul, Serge et al. (2011). *Web Data Management*. Cambridge University Press. URL : <http://webdam.inria.fr/Jorge/files/wdm.pdf>.
-  Lample, Guillaume et al. (2018). « Word translation without parallel data ». In : *International Conference on Learning Representations*. URL : <https://arxiv.org/abs/1710.04087>.
-  Mikolov, Tomas et al. (2013). « Distributed Representations of Words and Phrases and their Compositionality ». In : *Advances in Neural Information Processing Systems*. Sous la dir. de C.J. Burges et al. T. 26. Curran Associates, Inc. URL : <https://proceedings.neurips.cc/paper/2013/file/9aa42b31882ec039965f3c4923ce901b-Paper.pdf>.
-  Russell, Stuart et Peter Norvig (2020). *Artificial Intelligence: A Modern Approach*.