

MTI881

ASPECTS MATHÉMATIQUES DE L'APPRENTISSAGE PROFOND

SOLUTIONS DES EXERCICES

RÉDIGÉES PAR
GUILLAUME ROY-FORTIN

HIVER 2024

Ce document est mis à disposition selon les termes de la licence Creative Commons Attribution - Pas d'Utilisation Commerciale - Pas de Modification 4.0 International.



Table des matières

1	RAPPELS D'ALGÈBRE LINÉAIRE ET DE CALCUL DIFFÉRENTIEL	2
2	RÉSEAUX DE NEURONES PROFONDS	12
3	PROBLÈMES D'APPRENTISSAGE	17
4	UNIVERSALITÉ ET PUISSANCE EXPRESSIVE	22
5	RÔLE DE LA PROFONDEUR DANS LA PUISSANCE EXPRESSIVE	25
6	INTRODUCTION À L'APPRENTISSAGE PROFOND GÉOMÉTRIQUE	28
7	RÉDUCTION SPECTRALE DE LA DIMENSION	40
8	AUTOENCODEURS VARIATIONNELS	48

1 RAPPELS D'ALGÈBRE LINÉAIRE ET DE CALCUL DIFFÉRENTIEL

Exercice 1.1. Soit $\mathbf{u} = \begin{pmatrix} 2 \\ 1 \end{pmatrix}$, $\mathbf{v} = \begin{pmatrix} -1 \\ 1 \end{pmatrix}$ et $\mathbf{w} = \begin{pmatrix} 0 \\ 2 \end{pmatrix}$ des vecteurs de $\mathbb{R}^2$ ainsi que $\mathbf{a} = \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}$, $\mathbf{b} = \begin{pmatrix} 0 \\ 1 \\ 1 \end{pmatrix}$ des vecteurs de $\mathbb{R}^3$. On considère aussi $k \in \mathbb{R}$ un scalaire. Calculez le résultat des opérations suivantes et expliquez, le cas échéant, pourquoi celles-ci ne sont pas définies.

a. $3\mathbf{u} - 2\mathbf{v} + \mathbf{w}$

$$3\mathbf{u} - 2\mathbf{v} + \mathbf{w} = 3 \begin{pmatrix} 2 \\ 1 \end{pmatrix} - 2 \begin{pmatrix} -1 \\ 1 \end{pmatrix} + \begin{pmatrix} 0 \\ 2 \end{pmatrix} = \begin{pmatrix} 6 \\ 3 \end{pmatrix} - \begin{pmatrix} -2 \\ 2 \end{pmatrix} + \begin{pmatrix} 0 \\ 2 \end{pmatrix} = \begin{pmatrix} 8 \\ 3 \end{pmatrix}.$$

b. $(\mathbf{u}, \mathbf{u}) + (\mathbf{u}, \mathbf{v})$

$$(\mathbf{u}, \mathbf{u}) + (\mathbf{u}, \mathbf{v}) = 2 \cdot 2 + 1 \cdot 1 + 2(-1) + 1 \cdot 1 = 4.$$

c. $\|\mathbf{u}\|$

$$\|\mathbf{u}\| = \sqrt{(\mathbf{u}, \mathbf{u})} = \sqrt{2^2 + 1^2} = \sqrt{5}.$$

d. $3 + \mathbf{w}$

Cette opération n'est pas définie puisque $\mathbf{w}$ est un vecteur de $\mathbb{R}^2$ et 3 un scalaire.

e. $\mathbf{z} = \frac{1}{\|\mathbf{w}\|} \mathbf{w}$

On calcule d'abord

$$\|\mathbf{w}\| = \sqrt{(\mathbf{w}, \mathbf{w})} = \sqrt{0^2 + 2^2} = 2.$$

Ainsi,

$$\mathbf{z} = \frac{1}{\|\mathbf{w}\|} \mathbf{w} = \frac{1}{2} \begin{pmatrix} 0 \\ 2 \end{pmatrix} = \begin{pmatrix} 0 \\ 1 \end{pmatrix}.$$

f. $\|\mathbf{z}\|$

$$||\mathbf{z}|| = \sqrt{(\mathbf{z}, \mathbf{z})} = \sqrt{0^2 + 1^2} = 1.$$

On remarque donc que l'opération précédente a permis de normaliser le vecteur $\mathbf{z}$.

Exercice 1.2. Quelle est la valeur de l'angle θ entre les vecteurs $\mathbf{a}$ et $\mathbf{b}$ de l'exercice précédent?

On sait que l'angle aigu θ entre $\mathbf{a}$ et $\mathbf{b}$ satisfait

$$(\mathbf{a}, \mathbf{b}) = ||\mathbf{a}|| \cdot ||\mathbf{b}|| \cos \theta.$$

On calcule donc respectivement

$$\cdot ||\mathbf{a}|| = \sqrt{0^2 + 1^2 + 0^2} = 1$$

$$\cdot ||\mathbf{b}|| = \sqrt{0^2 + 1^2 + 1^2} = \sqrt{2}$$

$$\cdot (\mathbf{a}, \mathbf{b}) = 1$$

Ainsi, on a

$$\theta = \arccos\left(\frac{1}{\sqrt{2}}\right) = \frac{\pi}{4}.$$

Exercice 1.3. Soit $\mathbf{u} \in \mathbb{R}^m$ et on pose $\mathbf{n} = \frac{1}{||\mathbf{u}||} \mathbf{u}$. Montrer que $||\mathbf{n}|| = 1$.

En utilisant les propriétés du produit scalaire et de la norme, on obtient directement

$$(\mathbf{n}, \mathbf{n}) = \left(\frac{1}{||\mathbf{u}||} \mathbf{u}, \frac{1}{||\mathbf{u}||} \mathbf{u} \right) = \frac{1}{||\mathbf{u}||^2} (\mathbf{u}, \mathbf{u}) = \frac{||\mathbf{u}||^2}{||\mathbf{u}||^2} = 1.$$

Exercice 1.4. Soit $\mathfrak{B} = \{\mathbf{u}_1, \mathbf{u}_2\}$ une base de $\mathbb{R}^2$ avec $\mathbf{u}_1 = \begin{pmatrix} 1 \\ 1 \end{pmatrix}$ et $\mathbf{u}_2 = \begin{pmatrix} -1 \\ 2 \end{pmatrix}$. Exprimez chacun des vecteurs suivants comme combinaison linéaire des vecteurs de cette base.

$$\mathbf{a} = \begin{pmatrix} -3 \\ -3 \end{pmatrix}, \quad \mathbf{b} = \begin{pmatrix} 0 \\ 3 \end{pmatrix} \quad \text{et} \quad \mathbf{c} = \begin{pmatrix} 3 \\ 0 \end{pmatrix}.$$

On considère la matrice $A = (\mathbf{u}_1 \ \mathbf{u}_2)$ dont les colonnes sont les vecteurs de la base $\mathfrak{B}$. Les coordonnées d'un vecteur $\mathbf{b} \in \mathbb{R}^2$ dans cette base sont les coefficients de la combinaison linéaire de ces colonnes qui génère $\mathbf{b}$. Autrement dit, on peut les obtenir en solutionnant l'équation matricielle $A\mathbf{x} = \mathbf{b}$.

On calcule d'abord la matrice inverse de A , soit manuellement ou à l'aide de la TI. On obtient

$$A^{-1} = \begin{bmatrix} \frac{2}{3} & \frac{1}{3} \\ -\frac{1}{3} & \frac{1}{3} \end{bmatrix}$$

On obtient les coordonnées en multipliant les vecteurs fournis par l'inverse. Par exemple,

$$A^{-1}\mathbf{a} = \begin{pmatrix} -3 \\ 0 \end{pmatrix}$$

et on vérifie aisément que $\mathbf{a} = -3\mathbf{u}_1$. On écrit alors

$$(\mathbf{a})_{\mathfrak{B}} = \begin{pmatrix} -3 \\ 0 \end{pmatrix}.$$

De manière tout à fait similaire, on obtient enfin

$$(\mathbf{b})_{\mathfrak{B}} = \begin{pmatrix} 1 \\ 1 \end{pmatrix}, \quad (\mathbf{c})_{\mathfrak{B}} = \begin{pmatrix} 2 \\ 1 \end{pmatrix}.$$

Exercice 1.5. Soit $T : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ la transformation donnée par $T(x, y) = (3x + y, 2x - 5y)$. Montrer que T est une transformation linéaire.

Soit $\mathbf{u}, \mathbf{v} \in \mathbb{R}^2$ ainsi que des scalaires $\alpha, \beta \in \mathbb{R}$, avec $\mathbf{u} = \begin{pmatrix} u_1 \\ u_2 \end{pmatrix}$ et $\mathbf{v} = \begin{pmatrix} v_1 \\ v_2 \end{pmatrix}$. Alors,

$$\alpha \mathbf{u} + \beta \mathbf{v} = \begin{pmatrix} \alpha u_1 + \beta v_1 \\ \alpha u_2 + \beta v_2 \end{pmatrix}.$$

Ainsi,

$$\begin{aligned} T(\alpha \mathbf{u} + \beta \mathbf{v}) &= (3x + y, 2x - 5y) \\ &= (3(\alpha u_1 + \beta v_1) + \alpha u_2 + \beta v_2, 2(\alpha u_1 + \beta v_1) - 5(\alpha u_2 + \beta v_2)) \\ &= \alpha T(\mathbf{u}) + \beta T(\mathbf{v}). \end{aligned}$$

Remarquons que nous aurions également pu fournir la matrice canonique de T dont les colonnes sont l'image par la transformation T des vecteurs $\{\mathbf{e}_1, \mathbf{e}_2\}$ de la base canonique de $\mathbb{R}^2$ et vérifier que la transformation correspond bel et bien à la multiplication d'un vecteur par celle-ci.

Exercice 1.6. Soit les matrices

$$A = \begin{pmatrix} 1 & -1 & 0 \\ 1 & 0 & 2 \end{pmatrix}, \quad B = \begin{pmatrix} 1 & 2 \\ 0 & -1 \\ 1 & 3 \end{pmatrix} \quad \text{et} \quad C = \begin{pmatrix} a & b \\ c & d \end{pmatrix}$$

ainsi que les vecteurs $\mathbf{u} = \begin{pmatrix} 1 \\ 2 \end{pmatrix} \in \mathbb{R}^2$ et $\mathbf{v} = \begin{pmatrix} 1 \\ 0 \\ -2 \end{pmatrix} \in \mathbb{R}^3$.

Calculez le résultat des opérations suivantes et expliquez, le cas échéant, pourquoi celles-ci ne sont pas définies.

a. $A\mathbf{u}$

La multiplication d'un vecteur par une matrice est définie si la dimension du vecteur coïncide avec le nombre de colonnes de la matrice, ce qui n'est pas le cas ici.

b. $A\mathbf{v}$

$$A\mathbf{v} = \begin{pmatrix} 1 & -1 & 0 \\ 1 & 0 & 2 \end{pmatrix} \begin{pmatrix} 1 \\ 0 \\ -2 \end{pmatrix} = \begin{pmatrix} 1 \\ -3 \end{pmatrix}.$$

c. $B\mathbf{u}$

$$B\mathbf{u} = \begin{pmatrix} 1 & 2 \\ 0 & -1 \\ 1 & 3 \end{pmatrix} \begin{pmatrix} 1 \\ 2 \end{pmatrix} = \begin{pmatrix} 3 \\ -2 \\ 7 \end{pmatrix}.$$

d. $A^T \mathbf{u}$

$$A^T \mathbf{u} = \begin{pmatrix} 1 & 1 \\ -1 & 0 \\ 0 & 2 \end{pmatrix} \begin{pmatrix} 1 \\ 2 \end{pmatrix} = \begin{pmatrix} 3 \\ -1 \\ 4 \end{pmatrix}.$$

e. $A^T A$.

$$A^T A = \begin{pmatrix} 1 & 1 \\ -1 & 0 \\ 0 & 2 \end{pmatrix} \begin{pmatrix} 1 & -1 & 0 \\ 1 & 0 & 2 \end{pmatrix} = \begin{pmatrix} 2 & -1 & 2 \\ -1 & 1 & 0 \\ 2 & 0 & 4 \end{pmatrix}.$$

f. $C^T C$.

$$C^T C = \begin{pmatrix} a & c \\ b & d \end{pmatrix} \begin{pmatrix} a & b \\ c & d \end{pmatrix} = \begin{pmatrix} a^2 + c^2 & ab + cd \\ ab + cd & b^2 + d^2 \end{pmatrix}.$$

Exercice 1.7. Montrer que la matrice $A^{-1} = \begin{pmatrix} 1 & -1 \\ 0 & 1 \end{pmatrix}$ est l'inverse de la matrice $A = \begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix}$.

Il s'agit simplement de vérifier que $A \cdot A^{-1} = A^{-1} \cdot A = I_2$. On a effectivement

$$\begin{pmatrix} 1 & 1 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} 1 & -1 \\ 0 & 1 \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} = I_2.$$

Exercice 1.8. Expliquez en vos mots la différence entre une transformation linéaire et une transformation affine.

Une transformation linéaire entre deux espaces vectoriels euclidiens en est une qui satisfait la propriété démontrée à l'exercice 1.5. De manière équivalente, elle peut être réalisée par la multiplication matricielle. Cela exclut donc toute translation non-nulle, puisqu'une transformation linéaire doit avoir $\mathbf{0}$ comme point-fixe, autrement dit $T(\mathbf{0}) = \mathbf{0}$.

Par ailleurs, une transformation affine W est une transformation linéaire suivie d'une translation. On peut donc l'écrire de manière vectorielle comme

$$W(\mathbf{x}) = A\mathbf{x} + \mathbf{b},$$

où A est la matrice canonique de la partie linéaire de la transformation et $\mathbf{b}$ le vecteur de translation.

Exercice 1.9. Soit la matrice de rotation

$$O_\theta = \begin{pmatrix} \cos\theta & \sin\theta \\ -\sin\theta & \cos\theta \end{pmatrix}.$$

Montrer que O_θ est une matrice orthogonale, pour tout $\theta \in [0, 2\pi)$.

On voudrait montrer que $O_\theta^T = O_\theta^{-1}$. On calcule donc le produit

$$\begin{aligned} O_\theta \cdot O_\theta^T &= \begin{pmatrix} \cos\theta & \sin\theta \\ -\sin\theta & \cos\theta \end{pmatrix} \begin{pmatrix} \cos\theta & -\sin\theta \\ \sin\theta & \cos\theta \end{pmatrix} \\ &= \begin{pmatrix} \cos^2\theta + \sin^2\theta & \sin\theta\cos\theta - \sin\theta\cos\theta \\ \sin\theta\cos\theta - \sin\theta\cos\theta & \cos^2\theta + \sin^2\theta \end{pmatrix} \\ &= \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} = I_2. \end{aligned}$$

Exercice 1.10. Soit $A = \begin{pmatrix} 1 & -2 \\ -2 & 1 \end{pmatrix}$. Vérifier que A est symétrique et qu'elle satisfait au théorème spectral, c'est-à-dire qu'elle admet une base orthonormale $\mathfrak{B} = \{\mathbf{u}_1, \mathbf{u}_2\}$ de vecteurs propres.

Indication : vous pouvez faire appel à un logiciel de calcul symbolique (TI ou Wolfram, notamment) pour calculer les vecteurs propres de A . Ensuite, il s'agit de vérifier qu'il s'agit bien de vecteurs propres et qu'ils sont orthogonaux.

Exercice 1.11. Soit les fonctions suivantes

$$f(x) = \sin(x), \quad g(x) = x^3 + 2x + 1 \quad \text{et} \quad h(x) = \ln(2x).$$

- a. Que vaut la composition $f \circ g$? Et $g \circ f$?

On a directement

$$(f \circ g)(x) = f(g(x)) = \sin(x^3 + 2x + 1)$$

et

$$(g \circ f)(x) = g(f(x)) = \sin^3(x) + 2\sin(x) + 1.$$

- b. Calculer la dérivée $(f \circ g)'(x)$.

En vertu de la règle de la chaîne, on a

$$(f \circ g)'(x) = f'(g(x))g'(x) = \cos(x^3 + 2x + 1)(3x^2 + 2).$$

- c. Calculer la dérivée $(g \circ f)'(x)$.

Similairement au numéro précédent, on a

$$(g \circ f)'(x) = g'(f(x))f'(x) = (3\sin^2(x) + 2)(\cos(x)) = 3\sin^2(x)\cos(x) + 2\cos(x).$$

- d. Calculer la dérivée $(h \circ g \circ f)'(x)$.

Il est plus simple ici de définir directement la fonction obtenue par la double composition et de la dériver ensuite. Nommons là $w(x)$. On a

$$w(x) = (h \circ g \circ f)(x) = h(g(f(x))) = \ln(2(\sin^3(x) + 2\sin(x) + 1)).$$

En utilisant les propriétés du logarithme, le facteur constant $\ln 2$ se détache et devient nul après la dérivation. On a ainsi

$$w'(x) = \frac{3\sin^2(x)\cos(x) + 2\cos(x)}{\sin^3(x) + 2\sin(x) + 1}.$$

- e. Calculer la dérivée $(h \circ f \circ g)'(x)$.

En procédant de manière similaire, on pose

$$w(x) = h(f(g(x))) = \ln(2 \sin(x^3 + 2x + 1)).$$

On calcule ensuite

$$w'(x) = \frac{(3x^2 + 2) \cos(x^3 + 2x + 1)}{\sin(x^3 + 2x + 1)} = (3x^2 + 2) \cot(x^3 + 2x + 1).$$

Exercice 1.12. Soit $x : \mathbb{R}^2 \rightarrow \mathbb{R}$ donnée par

$$x(u, v) = u^2 + 2v$$

et $y : \mathbb{R}^2 \rightarrow \mathbb{R}$ donnée par

$$y(u, v) = \sin(u + v).$$

Soit également $z : \mathbb{R}^2 \rightarrow \mathbb{R}$ une fonction de (x, y) donnée par

$$z(x, y) = (xy)^2 + x + 2y.$$

À l'aide de la règle de la chaîne, calculer les dérivées partielles suivantes :

$$\frac{\partial z}{\partial x}, \frac{\partial z}{\partial y}, \frac{\partial z}{\partial u}, \frac{\partial z}{\partial v}.$$

On calcule directement les dérivées partielles

$$\frac{\partial z}{\partial x} = 2xy^2 + 1, \quad \frac{\partial z}{\partial y} = 2x^2y + 2.$$

Aussi, en vertu de la règle de la chaîne en plusieurs variables, on a

$$\begin{aligned} \frac{\partial z}{\partial u} &= \frac{\partial z}{\partial x} \frac{\partial x}{\partial u} + \frac{\partial z}{\partial y} \frac{\partial y}{\partial u} \\ &= (2xy^2 + 1)2u + (2x^2y + 2) \cos(u + v) \\ &= 2u [2(u^2 + 2v) \sin^2(u + v) + 1] + [2(u^2 + 2v)^2 \sin(u + v) + 2] \cos(u + v). \end{aligned}$$

Pareillement,

$$\begin{aligned} \frac{\partial z}{\partial v} &= \frac{\partial z}{\partial x} \frac{\partial x}{\partial v} + \frac{\partial z}{\partial y} \frac{\partial y}{\partial v} \\ &= (2xy^2 + 1) \cdot 2 + (2x^2y + 2) \cos(u + v) \\ &= 2 [2(u^2 + 2v) \sin^2(u + v) + 1] + [2(u^2 + 2v)^2 \sin(u + v) + 2] \cos(u + v). \end{aligned}$$

Exercice 1.13. Soit $x : \mathbb{R}^2 \rightarrow \mathbb{R}$ donnée par

$$x(u, v) = u + v^2$$

et $y : \mathbb{R}^2 \rightarrow \mathbb{R}$ donnée par

$$y(u, v) = -\cos(u + v).$$

Soit également $z : \mathbb{R}^2 \rightarrow \mathbb{R}$ une fonction de (x, y) donnée par

$$z(x, y) = x^2 y + 2x + y.$$

Enfin, on pose $u(t) = t^2$ et $v(t) = 2t$. À l'aide de la règle de la chaîne, calculer la dérivée totale $z'(t)$.

Cet exercice est très similaire au précédent dans sa première partie. On calcule d'abord directement les dérivées partielles

$$\frac{\partial z}{\partial x} = 2xy + 2, \quad \frac{\partial z}{\partial y} = x^2 + 1.$$

Aussi, en vertu de la règle de la chaîne en plusieurs variables, on a

$$\begin{aligned} \frac{\partial z}{\partial u} &= \frac{\partial z}{\partial x} \frac{\partial x}{\partial u} + \frac{\partial z}{\partial y} \frac{\partial y}{\partial u} \\ &= (2xy + 2)1 + (x^2 + 1)\sin(u + v) \\ &= 2u[2(u^2 + 2v)\sin^2(u + v) + 1] + [2(u^2 + 2v)^2 \sin(u + v) + 2]\cos(u + v). \end{aligned}$$

Pareillement,

$$\begin{aligned} \frac{\partial z}{\partial v} &= \frac{\partial z}{\partial x} \frac{\partial x}{\partial v} + \frac{\partial z}{\partial y} \frac{\partial y}{\partial v} \\ &= (2xy^2 + 1) \cdot 2 + (2x^2 y + 2)\cos(u + v) \\ &= 2[2(u^2 + 2v)\sin^2(u + v) + 1] + [2(u^2 + 2v)^2 \sin(u + v) + 2]\cos(u + v). \end{aligned}$$

Enfin, il s'agit de remarquer que

$$\begin{aligned} z'(t) &= \frac{\partial z}{\partial t} \\ &= \frac{\partial z}{\partial x} \frac{\partial x}{\partial u} \frac{\partial u}{\partial t} + \frac{\partial z}{\partial x} \frac{\partial x}{\partial v} \frac{\partial v}{\partial t} \\ &\quad + \frac{\partial z}{\partial y} \frac{\partial y}{\partial u} \frac{\partial u}{\partial t} + \frac{\partial z}{\partial y} \frac{\partial y}{\partial v} \frac{\partial v}{\partial t}. \end{aligned}$$

On remarque finalement que

$$u'(t) = 2t, \quad v'(t) = 2.$$

Exercice 1.14. On considère la fonction

$$f(x, y) = x^2 + y^2$$

et le point initial $x_0 = (3, -2) \in \mathbb{R}^2$. Complétez les 2 premières itérations d'une descente de gradient visant à converger vers le minimum local (et global) de f en $(0, 0)$, avec le taux d'apprentissage fixé à $\eta = 1$.

La fonction donnée est $f(x, y) = x^2 + y^2$ dont le gradient est $\nabla f(x, y) = (2x, 2y)$. Notre point initial quant à lui est $(6, -4)$ et le pas est fixé à $\frac{1}{4}$.

Les itérations de descente de gradient sont donc les suivantes :

- 1. Point de départ $(6, -4)$. En ce point, on a $\nabla f(6, -4) = (12, -8)$. Notre second point est donc*

$$(6, -4) - \frac{1}{4}(12, -8) = (6, -4) - (3, -2) = (3, -2).$$

- 2. Point de départ $(3, -2)$. En ce point, on a $\nabla f(3, -2) = (6, -4)$. Notre second point est donc*

$$(3, -2) - \frac{1}{4}(6, -4) = (3, -2) - (1.5, -1) = (1.5, -1).$$

On remarque bien qu'on se rapproche graduellement du minimum local (et absolu) situé à l'origine.

2 RÉSEAUX DE NEURONES PROFONDS

Exercice 2.1. Soit un réseau de neurones complètement connexe $\mathcal{F}$ contenant 4 couches de largeurs respectives $d_1 = 100, d_2 = 50, d_3 = 10$ et $d_4 = 20$.

- a. Le réseau est représenté par la fonction $\mathcal{F}$ qui est définie sur quel domaine? Et à valeur dans quel espace?

Le réseau est représenté par une fonction

$$\mathcal{F} : \mathbb{R}^{100} \rightarrow \mathbb{R}^{20}.$$

- b. Quelle est la profondeur de ce réseau?

La profondeur du réseau est de $L = 4$.

- c. Combien de biais b_j^ℓ ce réseau compte-t-il? Et quel est le nombre total de poids w_{jk}^ℓ de $\mathcal{F}$?

Ce réseau contient

$$\sum_{i=1}^4 d_i = 100 + 50 + 10 + 20 = 180$$

neurones, donc tout autant de biais au total. Pour les poids, on peut procéder couche par couche.

- De la première à la seconde, on a : $100 \cdot 50 = 5000$ poids.*
- De la seconde à la troisième, on a : $50 \cdot 10 = 500$ poids.*
- De la troisième à la dernière, on a : $10 \cdot 20 = 200$ poids.*

Ainsi, le réseau possède un total de 5700 poids.

- d. Combien de paramètres doit-on initialiser avant de faire une première propagation avant?

Avant de faire la première propagation avant, il faut initialiser tous les poids et les biais, soit donc un total de 5880 paramètres.

Exercice 2.2. Soit le réseau de neurones $\mathcal{F}$ représenté ci-dessous.



Pour l'initialiser, nous fixons tout d'abord les poids suivants :

$$w_{11}^2 = 2, w_{21}^2 = 0.5, w_{11}^3 = -1, w_{12}^3 = 1,$$

ainsi que les biais suivants

$$b_1^1 = 2, b_1^2 = 0, b_2^2 = 1, b_1^3 = -1.$$

Pour les besoins de l'exercice, on utilise la fonction identité comme fonction d'activation $\sigma(z) = z$.

- Placer les différents paramètres à l'endroit approprié sur le graphe du réseau.
- Quelle fonction $\mathcal{F}$ calcule-t-il? Fournir le domaine, l'ensemble d'arrivée ainsi que la définition explicite de $\mathcal{F}$.

On a d'abord

$$x \mapsto x + 2.$$

Puisque la fonction d'activation est l'identité, on peut calculer respectivement

$$\begin{aligned} \cdot a_1^1 &= \sigma(x + 2) = x + 2 \\ \cdot a_1^2 &= 2(x + 2) + 0 = 2x + 4 \\ \cdot a_2^2 &= \frac{1}{2}(x + 2) + 1 = \frac{1}{2}x + 2. \end{aligned}$$

Finalement,

$$\begin{aligned} a_1^3 &= w_{11}^3 a_1^2 + w_{21}^3 a_2^2 + b_1^3 \\ &= (-1)(2x + 4) + 1 \left(\frac{1}{2}x + 2 \right) - 1 \\ &= -2x - 4 + \frac{1}{2}x + 2 - 1 \\ &= -\frac{3}{2}x - 3. \end{aligned}$$

Donc, on a $\mathcal{F} : \mathbb{R} \rightarrow \mathbb{R}$ avec

$$\mathcal{F}(x) = \frac{-3x}{2} - 3.$$

On suppose maintenant que l'on voudrait faire apprendre à ce réseau à multiplier par deux. Plus formellement, la fonction idéale que cherche à approximer $\mathcal{F}$ est donnée par $y : \mathbb{R} \rightarrow \mathbb{R}$ avec $y(x) = 2x$. On suppose aussi que nos données d'entraînement sont les nombres entiers naturels, i.e. $X_n = n$, $n \in \mathbb{N}$.

- c. Donner la définition de la fonction de coût basée sur l'erreur quadratique moyenne (MSE), c'est-à-dire celle vue en classe.

On suppose la donnée d'entrée $x = n \in \mathbb{N}$ est fixée. Alors,

$$\begin{aligned} C(x) &= \frac{1}{2}(y(x) - a_1^3)^2 \\ &= \frac{1}{2}\left(2x - \left(\frac{-3x}{2} - 3\right)\right)^2 \\ &= \frac{1}{2}\left(2x + \frac{3}{2}x + 3\right)^2 \\ &= \frac{1}{2}\left(\frac{7}{2}x + 3\right)^2. \end{aligned}$$

- d. Calculer les sorties $a^L(X_n)$ du réseau pour $n = 1, 2$ de même que les sorties attendues $y(X_n)$ pour ces mêmes données d'entraînement.

· Pour $n = 1$, on a

$$a^3(1) = \mathcal{F}(1) = \frac{-3}{2}(1) - 3 = \frac{-9}{2}; \quad y(1) = 2.$$

· Pour $n = 2$, on a

$$a^3(2) = \mathcal{F}(2) = \frac{-3}{2}(2) - 3 = -6; \quad y(2) = 4.$$

- e. Calculer l'erreur du réseau sur ce mini-groupe de données.

On peut procéder directement à partir de la définition ponctuelle de la fonction d'erreur calculée précédemment. On obtient donc

$$\begin{aligned}
 C &= \frac{1}{2n} \sum_x C(x) = \frac{1}{2 \cdot 2} (C(1) + C(2)) \\
 &= \frac{1}{4} \left[(y(1) - a^2(1))^2 + (y(2) - a^2(2))^2 \right] \\
 &= \frac{1}{4} \left[\left(2 - \frac{9}{2} \right)^2 + (4 - -6)^2 \right] \\
 &= \frac{1}{4} \left(\left(\frac{13}{2} \right)^2 + 10^2 \right) \\
 &\approx 8.89.
 \end{aligned}$$

Exercice 2.3. Donner la définition de l'erreur associée au j -ème neurone de la couche ℓ . Expliquer en vos mots ce que celle-ci mesure à partir de la définition.

L'erreur du j -ème neurone de la couche ℓ est

$$\delta_j^\ell := \frac{\partial C}{\partial z_j^\ell},$$

où z est la préactivation de ce neurone. Elle mesure la possibilité d'améliorer le réseau en modifiant les paramètres en interaction avec ce neurone, i.e. les poids $w_{\cdot j}^\ell$ et le biais b_j^ℓ .

- *Si δ_j^ℓ est « petit », alors ces paramètres sont bien ajustés.*
- *Si δ_j^ℓ est « grand », alors au moins un de ces paramètres est mal ajusté et peut-être modifié de manière à faire significativement diminuer l'erreur C du réseau, via l'opération de descente de gradient.*

Exercice 2.4. Énoncer les quatre équations de la rétropropagation. De quelle manière pouvons-nous les utiliser pour calculer le gradient d'une fonction de coût?

Les 4 équations de la rétropropagation sont respectivement

· (RP1)

$$\delta_j^L = \frac{\partial C}{\partial a_j^\ell} \cdot \sigma'(z_j^\ell)$$

· (RP2)

$$\delta_j^\ell = \sum_k w_{kj}^{\ell+1} \cdot \delta_k^{\ell+1} \cdot \sigma'(z_j^\ell)$$

· (RP3)

$$\frac{\partial C}{\partial b_j^\ell} = \delta_j^\ell$$

· (RP4)

$$\frac{\partial C}{\partial w_{jk}^\ell} = \delta_j^\ell \cdot a_k^{\ell-1}$$

On utilise d'abord (RP1) pour calculer les erreurs δ_j^L associées aux neurones de la couche de sortie. Ensuite, on utilise successivement (RP2) pour calculer les erreurs δ_j^ℓ d'une couche à partir de celles $\delta_j^{\ell+1}$ de la couche qui la suit.

C'est en ce sens que nous pouvons parler de rétropropagation, puisque cette manière de procéder débute avec la toute dernière couche et propage vers l'arrière jusqu'à l'atteinte de la couche d'entrée.

Finalement, on utilise (RP3) et (RP4) pour calculer les dérivées partielles $\frac{\partial C}{\partial b_j^\ell}, \frac{\partial C}{\partial w_{jk}^\ell}$ qui composent le vecteur gradient ∇C .

3 PROBLÈMES D'APPRENTISSAGE

Exercice 3.1 Nous voulons poursuivre notre exploration de la fonction de coût d'entropie-croisée, donnée par

$$C = -\frac{1}{n} \sum_x y \ln(a) + (1 - y) \ln(1 - a),$$

où $a = a(x)$ est la sortie du réseau associée à l'entrée x et $y = y(x)$ est la sortie souhaitée pour x .

- a. Vérifier que la fonction de coût produit une sortie proche de 0 lorsque $y = 0.5$ et $a \approx 0.5$.

Si $y = \frac{1}{2}$ et $a \approx \frac{1}{2}$, alors on a $1 - a \approx \frac{1}{2}$ de telle sorte que

$$\ln(a) \approx \ln(1 - a) \approx \ln\left(\frac{1}{2}\right) = \ln(1) - \ln(2) = -\ln(2).$$

Ainsi, pour x fixé, on a

$$\begin{aligned} C_x &= -\left(\frac{1}{2}(-\ln(2)) + \frac{1}{2}(-\ln(2))\right) \\ &= \ln(2). \end{aligned}$$

Il s'agit de l'entropie croisée minimale dans ce cas.

- b. On considère un biais quelconque b_j^ℓ du réseau. Calculer la dérivée partielle

$$\frac{\partial C}{\partial b_j^\ell}.$$

On pose $a = \sigma(z)$ et $y = y(x)$. Ainsi

$$\begin{aligned}
 \frac{\partial C}{\partial b_j} &= \frac{\partial}{\partial b_j} \left(-\frac{1}{n} \sum_x y \ln(a) + (1-y) \ln(1-a) \right) \\
 &= -\frac{1}{n} \sum_x y \cdot \frac{1}{\sigma(z)} \cdot \sigma'(z) + \frac{1-y}{1-\sigma(z)} \cdot (-\sigma'(z)) \\
 &= -\frac{1}{n} \sum_x \frac{y\sigma'(z)(1-\sigma(z)) - (1-y)\sigma(z)\sigma'(z)}{\sigma(z)(1-\sigma(z))} \\
 &= -\frac{1}{n} \sum_x y - y\sigma(z) - \sigma(z) + y(z) \\
 &= \frac{1}{n} \sum_x (\sigma(z) - y).
 \end{aligned}$$

Dans le traitement plus haut, nous avons utilisé le fait que

$$\sigma'(z) = \sigma(z)(1 - \sigma(z)).$$

- c. Est-ce que l'expression obtenue permet d'éviter les problèmes d'apprentissage causés par l'usage conjoint de la fonction d'erreur quadratique (MSE) et de la sigmoïde de la section 3.1 ?

Oui, car la dérivée partielle est contrôlée par $(a^L - y)$ et n'est pas affectée par une éventuelle petite dérivée $\sigma'(z)$ de l'activation, comme c'était le cas en 3.1.

Exercice 3.2 On se propose d'explorer le concept d'initialisation aléatoire des paramètres en considérant le réseau de neurones suivant :



Soit également la distribution discrète W de support $\{0, 1\}$ et dont la fonction de masse $p(x)$ est uniforme, i.e

$$p(0) = p(1) = \frac{1}{2}.$$

On considère également la distribution discrète B de support $\{0, 1, 2\}$ et dont la fonction de masse $q(x)$ est également uniforme, c'est-à-dire donnée par

$$q(0) = q(1) = q(2) = \frac{1}{3}.$$

Rappel : la fonction de masse décrit les probabilités qu'une variable aléatoire suivant cette distribution prenne les différentes valeurs possibles de son support. Par exemple, dans les distributions que l'on vient de présenter, $p(0) = \frac{1}{2}$ veut dire qu'une variable aléatoire suivant la distribution W a une chance sur deux de prendre la valeur 0.

Finalement, on fixe la donnée d'entrée $x = 1$ et on n'utilise pas de fonction d'activation (autrement dit $\sigma(z) = z$). On suppose que seul le neurone du centre possède un biais b et on note par w_1 et w_2 les deux poids du réseau.

- a. On initialise le réseau en échantillonnant $w_1 \sim W$, $w_2 \sim W$ et $b \sim B$. Le réseau est maintenant une variable aléatoire Y . Donner le support ainsi que la fonction de masse de Y .

On a $w_i \sim W$, avec fonction de masse $p(x)$; $p(0) = p(1) = \frac{1}{2}$. Autrement dit,

$$P(w_i = 0) = P(w_i = 1) = \frac{1}{2}.$$

Aussi, $b \sim B$, avec fonction de masse $q(x)$; $q(0) = q(1) = q(2) = \frac{1}{3}$. Autrement dit,

$$P(b = 0) = P(b = 1) = P(b = 2) = \frac{1}{3}.$$

Puisque $\sigma = Id$, le réseau calcule

$$y(x) = w_2(w_1x + b) = w_2w_1x + w_2b.$$

Comme w_1, w_2, b sont des variables aléatoires, la sortie du réseau est aussi une variable aléatoire. Chacune des sorties possible survient avec une probabilité de $\frac{1}{2} \cdot \frac{1}{2} \cdot \frac{1}{3} = \frac{1}{12}$. En effectuant le dénombrement, on obtient que le support de la variable aléatoire de sortie Y est $\{0, 1, 2, 3\}$. Sa fonction de masse est donnée par

$$p(0) = \frac{7}{12}, p(1) = \frac{2}{12} = \frac{1}{6}, p(2) = \frac{2}{12} = \frac{1}{6}, \text{ et } p(3) = \frac{1}{12}.$$

- b. Calculer l'espérance de w_1 et de b .

À partir de la définition de l'espérance mathématique, on calcule

$$\mathbb{E}(w_1) = 0 \cdot P(w_1 = 0) + 1 \cdot P(w_1 = 1) = \frac{1}{2}.$$

De manière similaire,

$$\mathbb{E}(b_1) = 0 \cdot P(b = 0) + 1 \cdot P(b = 1) + 2 \cdot P(b = 2) = 1.$$

- c. Calculer la variance de w_1 et de b .

À partir de la définition de la variance, on a

$$\begin{aligned} \text{Var}(w_1) &= \mathbb{E}(w_1^2) - (\mathbb{E}(w_1))^2 \\ &= (0^2 \cdot p(0) + 1^2 \cdot p(1)) - \left(\frac{1}{2}\right)^2 = 1 - \frac{1}{4} \\ &= \frac{3}{4}. \end{aligned}$$

Similairement,

$$\begin{aligned} \text{Var}(b) &= \mathbb{E}(b^2) - (\mathbb{E}(b))^2 \\ &= (0^2 \cdot p(0) + 1^2 \cdot p(1) + 2^2 \cdot p(2)) - 1^2 = \frac{5}{3} - 1 \\ &= \frac{2}{3}. \end{aligned}$$

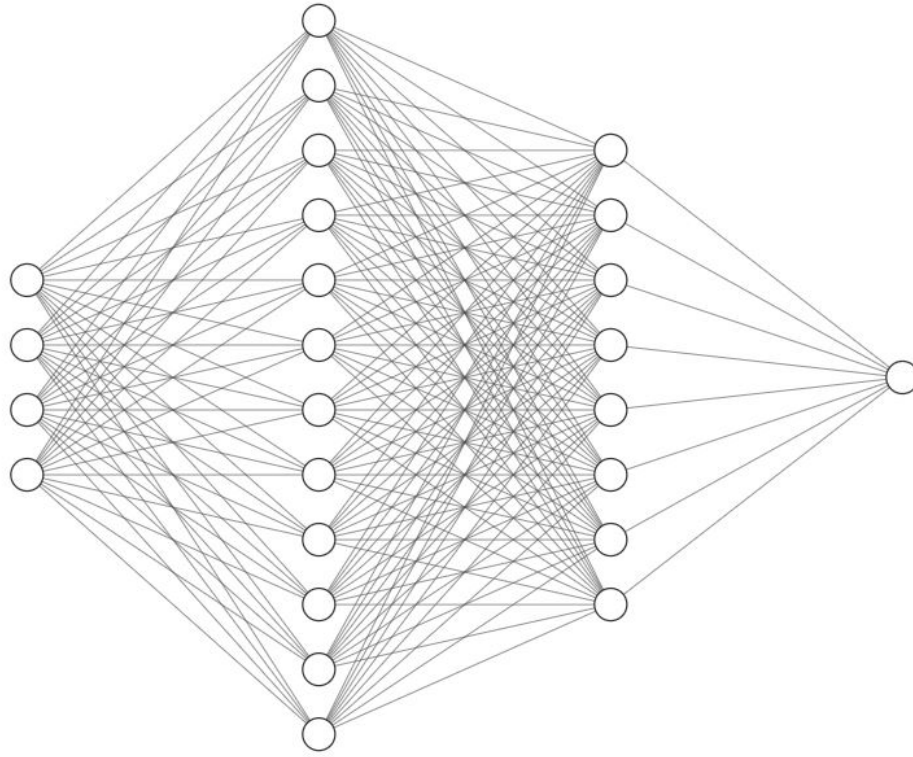
- d. Est-ce que l'initialisation de ce réseau aléatoire satisfait le troisième critère du théorème de Hanin-Rolnick?

Pour satisfaire le 3ème critère du théorème de Hanin-Rolnick, il faut que la variance d'un paramètre associé à la couche ℓ soit égale à

$$\frac{2}{fan-in} = \frac{2}{d_{\ell-1}}.$$

Ici, $d \equiv 1$ pour les 3 couches et il faudrait donc des variances de $\frac{2}{1} = 2$, ce qui n'est pas le cas après les calculs du numéro précédents.

Exercice 3.3 Soit le réseau profond suivant, qui prend en entrée un vecteur de $\mathbb{R}^4$ et produit un nombre réel en sortie.



On désire initialiser aléatoirement les paramètres de ce réseau de manière à respecter les critères du théorème de Hanin-Rolnick. Pour chaque poids w_{jk}^ℓ et biais b_j^ℓ du réseau, fournir une distribution probabiliste qui soit telle que ces conditions d'application du théorème sont remplies.

Ce réseau a $L = 4$ couches avec $d_1 = 4, d_2 = 12, d_3 = 8, d_4 = 1$.

- *Pour les paramètres associés à la couche $\ell = 2$, i.e. b_j^2 et w_{ij}^2 , il faut que*

$$Var = \frac{2}{\text{"fan-in"}} = \frac{2}{d_1} = \frac{2}{4} = \frac{1}{2}.$$

- *Pour les paramètres associés à la couche $\ell = 3$, i.e. b_j^3 et w_{ij}^3 , il faut que*

$$Var = \frac{2}{\text{"fan-in"}} = \frac{2}{d_2} = \frac{2}{12} = \frac{1}{6}.$$

- *Enfin, pour les paramètres associés à la couche $\ell = 4$, i.e. b_j^4 et w_{ij}^4 , il faut que*

$$Var = \frac{2}{\text{"fan-in"}} = \frac{2}{d_3} = \frac{2}{8} = \frac{1}{4}.$$

Pour tous ces paramètres, on peut choisir une distribution normale de variance approprié, soit $\frac{1}{2}$, $\frac{1}{6}$ ou $\frac{1}{4}$ selon le cas.

4 UNIVERSALITÉ ET PUISSANCE EXPRESSIVE

Exercice 4.1 Répondre par *vrai* ou *faux* à chacun des énoncés suivants portant sur le théorème de Cybenko. Lorsque vous répondez *faux*, justifiez.

On rappelle d'abord ici le type de réseaux utilisés dans le théorème de Cybenko :

$$S = \left\{ G(x) = \sum_{j=1}^D \alpha_j \sigma(w^j \cdot x + b_j) : w^j \in \mathbb{R}^n; \alpha_j, b_j \in \mathbb{R} \right\}.$$

- a. Pour toute fonction continue $f^*(x)$, il existe un réseau $G(x) \in S$ tel que $f^*(x) = G(x)$.

Faux. On peut approximer arbitrairement près, mais rien ne garantit qu'un réseau $G(x)$ existe tel que $G(x) \equiv f^(x)$.*

- b. L'énoncé du théorème nous dit que plus $\epsilon > 0$ est petit, plus l'approximation d'une fonction idéale f^* par un réseau $G(x) \in S$ est précise.

Vrai.

- c. Il est possible d'appliquer tel quel le théorème de Cybenko en choisissant ReLU comme fonction d'activation, i.e. en fixant $\sigma(z) = \max(0, z)$.

Faux, car $\sigma = \text{ReLU}$ n'est pas sigmoïdale :

$$\lim_{z \rightarrow \infty} \text{ReLU}(z) = \infty.$$

- d. On peut approximer n'importe quelle fonction continue à un degré arbitrairement précis avec un réseau $G(x) \in S$ qui contient un nombre fixe de neurones.

Faux. Pour améliorer l'approximation, il faut typiquement augmenter le nombre de neurones de la couche du milieu. Autrement dit

$$N = N(\epsilon) \rightarrow \infty \text{ si } \epsilon \rightarrow 0.$$

- e. Les réseaux $G(x) \in S$ du théorème de Cybenko sont profonds mais peu larges.

Faux. Ils sont peu profonds $L = 3$ mais larges.

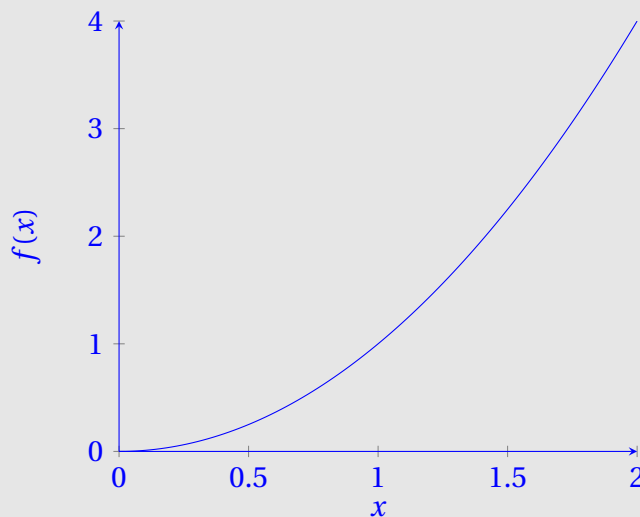
Exercice 4.2 On rappelle qu'une *fonction simple* est une combinaison linéaire de fonctions indicatrices, c'est-à-dire qu'elles sont de la forme

$$s(x) = \sum_{i=1}^N \alpha_i \chi_{(a_i, b_i)},$$

et où $\chi_{(a_i, b_i)}$ est la fonction indicatrice de l'intervalle (a_i, b_i) , c'est-à-dire la fonction qui prend la valeur 1 si $x \in (a_i, b_i)$ et 0 sinon.

Soit maintenant la fonction $f : [0, 2] \rightarrow \mathbb{R}$ donnée par $f(x) = x^2$.

a. Approximez f par une fonction simple qui contient un morceau.



On veut donc approximer cette fonction avec une expression de la forme

$$s(x) = \sum_{i=1}^N \alpha_i \chi_{(a_i, b_i)},$$

où $\chi_{(a_i, b_i)}$ est la fonction indicatrice de l'intervalle (a_i, b_i) . On commence donc avec un seul morceau et on cherche

$$s_1(x) = \alpha_1.$$

Quelle amplitude α_1 choisir? Il semble raisonnable de choisir la valeur moyenne de f sur $[0, 2]$:

$$\alpha_1 = \frac{1}{2-0} \int_0^2 x^2 dx = \frac{4}{3}.$$

Donc, $s_1(x) = \frac{4}{3}$.

b. Approximez f par une fonction simple qui contient deux morceaux.

Même chose, mais avec $s_2(x) = \alpha_1 \chi_{(0,1)} + \alpha_2 \chi_{(1,2)}$. Les coefficients α_1 et α_2 correspondant à la valeur moyenne respective de f sur les intervalles $(0,1)$ et $(1,2)$. On calcule

$$\alpha_1 = \frac{1}{1-0} \int_0^1 x^2 dx = \frac{1}{3}, \quad \alpha_2 = \frac{1}{2-1} \int_1^2 x^2 dx = \frac{7}{3}.$$

On obtient donc la fonction simple d'ordre 2 suivante

$$s_2(x) = \frac{1}{3} \chi_{(0,1)} + \frac{7}{3} \chi_{(1,2)}.$$

c. Approximez f par une fonction simple qui contient quatre morceaux.

Même chose, mais cette fois-ci avec $n = 4$ morceaux. On partitionne l'intervalle ainsi $(0, \frac{1}{2})$, $(\frac{1}{2}, 1)$, $(1, \frac{3}{2})$ et enfin $(\frac{3}{2}, 2)$. On aura

$$s_4(x) = \alpha_1 \chi_{(0, \frac{1}{2})} + \alpha_2 \chi_{(\frac{1}{2}, 1)} + \alpha_3 \chi_{(1, \frac{3}{2})} + \alpha_4 \chi_{(\frac{3}{2}, 2)}.$$

On calcule les amplitudes de manière similaire aux exemples précédents. On a

$$\alpha_1 = \frac{1}{\frac{1}{2} - 0} \int_0^{\frac{1}{2}} x^2 dx = \frac{1}{12}; \quad \alpha_2 = \frac{1}{1 - \frac{1}{2}} \int_{\frac{1}{2}}^1 x^2 dx = \frac{7}{12}$$

ainsi que

$$\alpha_3 = \frac{1}{\frac{3}{2} - 1} \int_1^{\frac{3}{2}} x^2 dx = \frac{19}{12}; \quad \alpha_4 = \frac{1}{2 - \frac{3}{2}} \int_{\frac{3}{2}}^2 x^2 dx = \frac{37}{12}.$$

On obtient finalement la fonction simple d'ordre 4 suivante

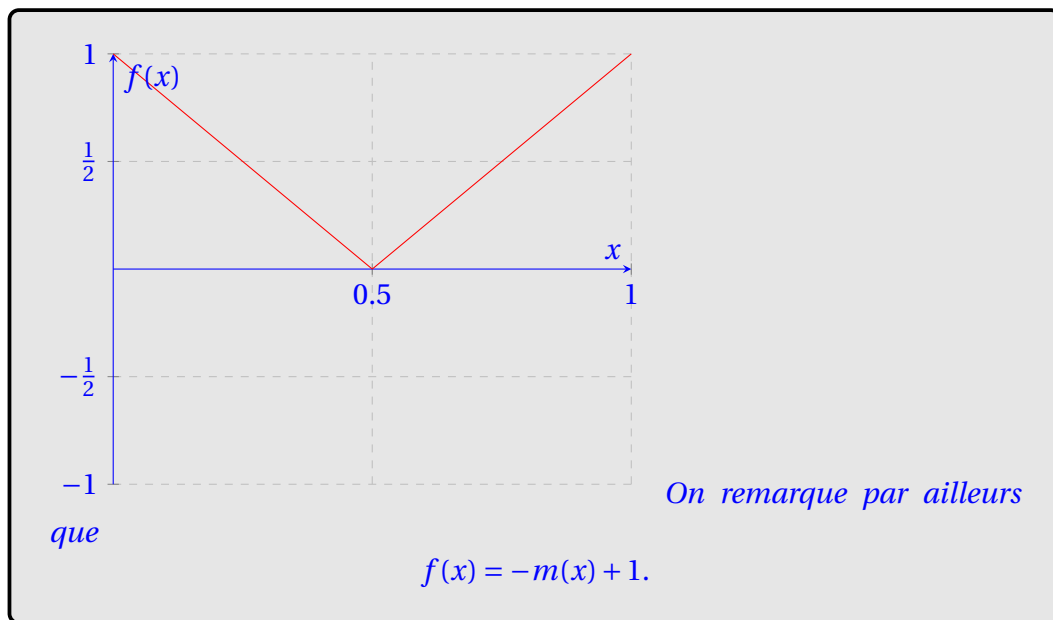
$$s_4(x) = \frac{1}{12} \chi_{(0, \frac{1}{2})} + \frac{7}{12} \chi_{(\frac{1}{2}, 1)} + \frac{19}{12} \chi_{(1, \frac{3}{2})} + \frac{37}{12} \chi_{(\frac{3}{2}, 2)}.$$

5 RÔLE DE LA PROFONDEUR DANS LA PUISSANCE EXPRESSIVE

Exercice 5.1 Soit la fonction

$$f(x) = \begin{cases} 1 - 2x & \text{si } 0 \leq x \leq \frac{1}{2} \\ 2x - 1 & \text{si } \frac{1}{2} \leq x \leq 1 \end{cases}$$

a. Tracez le graphe de f .



b. Créez un réseau de neurones qui calcule exactement f , i.e. dont la sortie est $a^L = f(x)$.

On fixe d'abord $\sigma = \text{ReLU}$. Aussi, nous avons vu que

$$m(x) = \sigma(2\sigma(x) - 4\sigma(x - 1/2)).$$

Ainsi, $f(x) = -m(x) + 1$ peut être calculé en rajoutant une couche de 1 neurone au réseau qui calculait $m(x)$ avec un poids de $w = -1$ et un biais associé au nouveau neurone de 1.. Ainsi

$$a^L = \sigma(-m(x) + 1) = f(x).$$

c. Approximez f par une fonction simple qui contient quatre morceaux.

Il s'agit de considérer

$$s_4(x) = \alpha_1 \chi(0, \frac{1}{4}) + \alpha_2 \chi(\frac{1}{2}, \frac{1}{4}) + \alpha_3 \chi(\frac{1}{2}, \frac{3}{4}) + \alpha_4 \chi(\frac{3}{4}, 1),$$

avec

$$\alpha_1 = \frac{3}{4}, \alpha_2 = \frac{1}{4}, \alpha_3 = \frac{1}{4}, \alpha_4 = \frac{3}{4}.$$

- d. Tracez le graphe de la composition $f^{(2)}(x) = (f \circ f)(x)$.

Il s'agit d'une fonction en dents de scie définie sur $(0,1)$ et qui prend la valeur 1 en $0, 1/2, 1$ et 0 en $1/4$ et $3/4$.

- e. Créez un réseau de neurones qui calcule exactement $f^{(2)}(x)$.

Il s'agit simplement de coller horizontalement deux fois le RN obtenu en (b).

- f. Les fonctions $f^{(n)}$ sont linéaires par morceaux. Combien de morceaux possède $f^{(n)}$?

- $f^{(1)}$ possède 2 morceaux (on se restreint à l'intervalle $[0, 1]$.)
- $f^{(2)}$ possède $4 = 2^2$ morceaux.

Exercice 5.2 On se place dans le contexte du test de classification à n -points dont nous avons discuté en classe en prélude au théorème de Telgarsky. On rappelle que la discrétisation $\tilde{f}$ d'une fonction $f : [0, 1] \rightarrow \mathbb{R}$ est donnée par

$$\tilde{f}(x) = \begin{cases} 1 & \text{si } f(x) \geq \frac{1}{2} \\ 0 & \text{si } f(x) < \frac{1}{2} \end{cases}$$

L'erreur $E(f)$ associé à f est alors définie comme la proportion des n -points qui ont été mal classés par cette discrétisation, c'est-à-dire

$$E(f) = \frac{1}{2^n} \sum \chi \{ \tilde{f}(x_i) \neq y_i \}.$$

On considère la fonction $f(x) = x^2$ sur l'intervalle $[0, 1]$. Calculez $E(f)$ pour le test de classification avec $n = 1, 2, 3$.

La discrétisation de $f(x)$ est

$$\tilde{f}(x) = \begin{cases} 1 & \text{si } x^2 \geq \frac{1}{2} \\ 0 & \text{si } x^2 < \frac{1}{2} \end{cases}$$

Le seuil est donc ici d'environ 0.7. L'erreur $\mathcal{E}(f)$ de $f(x) = x^2$ au test de classification est la proportion des points qui ne sont pas traversés par la discrétisation $\tilde{f}$. On a donc

$$\mathcal{E}(f) = \begin{cases} \frac{1}{2} & \text{si } n = 1 \\ \frac{1}{4} & \text{si } n = 2 \\ \frac{4}{8} = \frac{1}{2} & \text{si } n = 4. \end{cases}$$

Exercice 5.3 Expliquez en vos mots la signification et la portée théorique du théorème de Telgarsky.

6 INTRODUCTION À L'APPRENTISSAGE PROFOND GÉOMÉTRIQUE

Exercice 6.1. On suppose qu'un réseau de convolution reçoit en entrée une image de format de 32×32 pixels.

- a. Combien de neurones possède la couche d'entrée de ce réseau?

On a

$$d_1 = 32 \times 32 = 1024.$$

- b. Si vous créez un filtre h de taille 5×5 , quelles dimensions aura la sortie obtenue par la convolution de l'image initiale avec h ?

La sortie aura alors une dimension de $28 \times 28 = 784$.

Note : on suppose ici que la couche convolutionnelle ne fait pas de padding et a un stride (pas) unitaire.

- c. Quelle propriété structurelle des images les distingue des graphes et nous permet d'implémenter naturellement une couche de convolution sur celles-ci?

Organisation en grille, c'est-à-dire structure euclidienne sous-jacente.

Exercice 6.2. On rappelle que le *laplacien*, noté Δ , est un opérateur différentiel qui agit sur des fonctions $f(x_1, x_2, \dots, x_n)$ de n variables de la manière suivante

$$\Delta f(x) = - \sum_{i=1}^n \frac{\partial^2 f}{\partial x_i^2}(x).$$

Par exemple, pour une fonction de deux variables $f(x, y)$, le laplacien devient

$$\Delta f = \frac{\partial^2 f}{\partial x^2} + \frac{\partial^2 f}{\partial y^2}.$$

Note : le signe négatif devant la somme est une convention qui nous permet de s'assurer d'avoir des valeurs propres non-négatives.

- a. Que devient le laplacien lorsque $n = 1$, c'est-à-dire lorsqu'on le fait agir sur des fonctions d'une seule variable?

Si $n = 1$, on a

$$\Delta = -\frac{d^2}{dx^2},$$

soit l'opérateur de dérivée seconde pour les fonctions d'une variable.

- b. Vérifier que $\sin(\lambda x)$ et $\cos(\lambda x)$ sont des fonctions propres de Δ et identifier la valeur propre associée.

On calcule directement

$$\Delta(\sin(\lambda x)) = -\frac{d^2}{dx^2} \sin(\lambda x) = -\frac{d}{dx} \lambda \cos(\lambda x) = \lambda^2 \sin(\lambda x),$$

et donc la fonction $\sin(\lambda x)$ est une fonction propre de valeur propre λ^2 . De manière tout à fait similaire, on a

$$\Delta(\cos(\lambda x)) = -\frac{d^2}{dx^2} \cos(\lambda x) = \frac{d}{dx} \lambda \sin(\lambda x) = \lambda^2 \cos(\lambda x).$$

- c. Soit $f(x, y) = x^2 + y^2$. Calculer Δf .

Pour une fonction de deux variables, le laplacien devient

$$\Delta f(x, y) = -\left(\frac{\partial^2 f}{\partial x^2} + \frac{\partial^2 f}{\partial y^2}\right).$$

Puisque,

$$\frac{\partial^2 f}{\partial x^2} = \frac{\partial^2 f}{\partial y^2} = 2,$$

on a donc que $\Delta f(x, y) = 4$.

- d. On considère la famille de fonctions

$$\varphi_{m,n}(x, y) = \sin(mx) \sin(ny),$$

indexées par $(m, n) \in \mathbb{Z}^2$, c'est-à-dire que (m, n) est n'importe quel couple d'entiers relatifs. Montrer que $\varphi_{m,n}$ est une fonction propre du laplacien et donner la valeur propre correspondante.

On calcule

$$\frac{\partial \varphi}{\partial x} = m \cos(mx) \sin(ny); \quad \frac{\partial^2 \varphi}{\partial x^2} = -m^2 \sin(mx) \sin(ny).$$

Les dérivées par rapport à la variable y se calculent de manière similaire

$$\frac{\partial \varphi}{\partial y} = n \sin(mx) \cos(ny); \quad \frac{\partial^2 \varphi}{\partial y^2} = -n^2 \sin(mx) \sin(ny).$$

Ainsi,

$$\begin{aligned} \Delta \varphi_{m,n} &= - \left(\frac{\partial^2 \varphi}{\partial x^2} + \frac{\partial^2 \varphi}{\partial y^2} \right) = (m^2 + n^2) \sin(mx) \sin(ny) \\ &= \lambda_{m,n} \varphi_{m,n}. \end{aligned}$$

Donc, les éléments de la famille $\varphi_{m,n}$ sont des fonctions propres du Δ avec valeur propre $\lambda_{m,n} = (m^2 + n^2)$.

Exercice 6.3. Un *champ vectoriel* est une fonction dont la sortie est un vecteur, par opposition aux fonctions usuelles dont les sorties sont des nombres (ou scalaires). Par exemple, l'application $F : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ donnée par

$$F(x, y) = (2x, 2y)$$

est un champ vectoriel.

- a. Calculer le gradient ∇f de la fonction $f(x, y) = x^2 + y^2$. Que remarquez-vous?

On a directement

$$\nabla f = (2x, 2y)$$

qui est donc un champ vectoriel sur $\mathbb{R}^2$.

On souligne donc que le gradient transforme des fonctions en champs vectoriels. L'opérateur de *divergence*, noté div agit quant à lui dans la direction opposée et transforme des champs vectoriels en fonctions. Plus précisément, étant donné un champ vectoriel F , la divergence de ce champ est donnée par

$$\text{div } F = \sum_{i=1}^n \frac{\partial}{\partial x_i} = \frac{\partial}{\partial x_1} + \frac{\partial}{\partial x_2} + \dots + \frac{\partial}{\partial x_n}.$$

- b. Calculer la divergence du champ vectoriel F tel que défini au début de cette question. Que remarquez-vous?

On a

$$\begin{aligned}\operatorname{div} F &= \sum_{i=1}^2 \frac{\partial}{\partial x_i} = \frac{\partial F_1}{\partial x} + \frac{\partial F_2}{\partial y} \\ &= \frac{\partial}{\partial x}(2x) + \frac{\partial}{\partial y}(2y) = 2 + 2 = 4.\end{aligned}$$

Donc, $-\operatorname{div} F = -\operatorname{div} \nabla f = -4$, tout comme lorsque nous avons calculé le laplacien de $f(x, y)$. Ce n'est pas un hasard comme nous allons le démontrer au prochain exercice.

c. Montrer que $\Delta = -\operatorname{div} \nabla$, i.e. montrer que

$$-\operatorname{div} \nabla f = -\sum_{i=1}^n \frac{\partial^2 f}{\partial x_i^2}.$$

À partir des définitions respectives, on a

$$-\operatorname{div} \nabla f = -\sum_{i=1}^n \frac{\partial}{\partial x_i} (\nabla f)_i = -\sum_{i=1}^n \frac{\partial}{\partial x_i} \left(\frac{\partial f}{\partial x_i} \right) = -\sum_{i=1}^n \frac{\partial^2 f}{\partial x_i^2} = \Delta f.$$

Exercice 6.4. Expliquez en vos mots en quoi le théorème de Hilbert-Schmidt est analogue au théorème spectral pour les matrices symétriques.

Il permet d'établir que les fonctions propres du laplacien forment une base orthonormale de l'espace $L^2(X)$, où $X \subset \mathbb{R}^n$ est le domaine des fonction sur lequel agit le laplacien.

Exercice 6.5. Soit $X \subset \mathbb{R}$. On rappelle que l'espace $L^2(X)$ contient l'ensemble des fonctions de carré intégrable à valeurs réelles et définies sur X . Autrement dit, $u \in L^2(X)$ si et seulement si

$$\int_X |f(x)|^2 dx < \infty.$$

L'espace $L^2(X)$ est un espace dit hilbertien parce qu'il porte un *produit scalaire*, en tout point analogue à celui que vous utilisez usuellement avec les vecteurs de $\mathbb{R}^n$. Ce produit scalaire est défini avec l'aide de l'intégrale. En effet, pour tout $u, v \in L^2(X)$, on a

$$(u, v)_X := \int_X u(x) \cdot v(x) dx.$$

On considère maintenant les 4 fonctions suivantes

$$f(x) = e^{-x^2}, \quad g(x) = \frac{1}{x^2 + 1}, \quad u(x) = e^{-x}, \quad h(x) = x.$$

Note : pour calculer les intégrales nécessaires pour répondre aux prochaines sous-questions, vous pouvez et devriez utiliser la TI ou un logiciel de calcul symbolique!

- a. Parmi les 4 fonctions présentées ci-haut, dire lesquelles font partie de $L^2(\mathbb{R})$ et justifier sa réponse.

Pour appartenir à $L^2(\mathbb{R})$, une fonction doit satisfaire

$$\int_{\mathbb{R}} |f|^2 dx < \infty.$$

On vérifie aisément que $f, g \in L^2(\mathbb{R})$, mais que $u, h \notin L^2(\mathbb{R})$.

- b. Parmi les 4 fonctions présentées ci-haut, dire lesquelles font partie de $L^2[0, 1]$ et justifier sa réponse.

Toutes les fonctions, i.e. f, g, u et h , sont dans $L^2([0, 1])$ car pour chacune de ces fonctions, on a

$$\int_0^1 |f|^2 dx < \infty.$$

- c. Calculer le produit scalaire suivant

$$(h, u)_{[0,1]}.$$

Il s'agit de calculer l'intégrale

$$(h, u)_{[0,1]} = \int_0^1 h(x)u(x) dx = \int_0^1 xe^{-x} dx \approx 0.26.$$

- d. Calculer le produit scalaire suivant

$$(f, g)_{\mathbb{R}}.$$

Il faut cette fois-ci calculer l'intégrale

$$(f, g)_{\mathbb{R}} = \int_{\mathbb{R}} f(x)g(x) dx = \int_{\mathbb{R}} e^{-x^2} \frac{1}{1+x^2} dx \approx 1.34.$$

Exercice 6.6. On rappelle que la *transformée de Fourier* $\hat{f}$ d'une fonction $f : \mathbb{R} \rightarrow \mathbb{R}$ est la fonction définie par

$$\mathcal{F}\{f\} = \hat{f}(\xi) = \int_{\mathbb{R}} f(t)e^{2\pi i \xi t} dt.$$

De plus, la *transformée inverse de Fourier* est obtenu via

$$\mathcal{F}^{-1}\{\hat{f}\} = f(x) = \frac{1}{\sqrt{2\pi}} \int_{\mathbb{R}} \hat{f}(\xi) e^{-2\pi i \xi t} d\xi.$$

- a. Écrire l'expression de la transformée de Fourier comme un produit scalaire sur $\mathbb{R}$.

On a

$$\hat{f}(\xi) = \mathcal{F}(f)(\xi) = \int_{\mathbb{R}} f(t) e^{2\pi i \xi t} dt = (f(t), e^{2\pi i \xi t})_{t \in \mathbb{R}},$$

où la dernière notation indique que la variable libre sur laquelle est pris le produit scalaire est $t \in \mathbb{R}$.

- b. Écrire l'expression de la transformée de Fourier inverse comme un produit scalaire sur $\mathbb{R}$.

On a

$$f(x) = \mathcal{F}^{-1}(\hat{f})(x) = \int_{\mathbb{R}} \hat{f}(\xi) e^{-2\pi i \xi t} d\xi = (\hat{f}(\xi), e^{-2\pi i \xi t})_{\xi \in \mathbb{R}},$$

où la dernière notation indique que la variable libre sur laquelle est pris le produit scalaire est $\xi \in \mathbb{R}$.

- c. Vérifier que les exponentielles complexes qui apparaissent dans ces produits scalaires sont des fonctions propres du laplacien en 1 dimension, agissant sur la variable t pour $\mathcal{F}$ et sur ξ pour la transformée inverse $\mathcal{F}^{-1}$.

Il s'agit d'un calcul direct, dans les deux directions. On

$$\begin{aligned} \Delta_t(e^{2\pi i \xi t}) &= \frac{d^2}{dt^2}(e^{2\pi i \xi t}) \\ &= (2\pi i \xi)^2 e^{2\pi i \xi t} \\ &= (-4\pi^2 \xi^2) e^{2\pi i \xi t} = \lambda \varphi_s. \end{aligned}$$

Du côté de la transformée inverse, on a plutôt

$$\begin{aligned} \Delta_\xi(e^{-2\pi i \xi t}) &= \frac{d^2}{d\xi^2}(e^{-2\pi i \xi t}) \\ &= (-2\pi i t)^2 e^{-2\pi i \xi t} \\ &= (-4\pi^2 t^2) e^{-2\pi i \xi t} = \lambda \varphi_t. \end{aligned}$$

Exercice 6.7 Soit le graphe G donné par les sommets $V = \{1, 2, 3\}$ et la matrice d'adjacence

des poids

$$W = \begin{pmatrix} 0 & 1 & 2 \\ 1 & 0 & -1 \\ 2 & -1 & 0 \end{pmatrix}.$$

- Dessiner le le graphe G en indiquant les poids reliés à chacune des arêtes qu'il possède.
- Une fonction sur G est une fonction définie sur les sommets, soit une application $f : V \rightarrow \mathbb{R}$. L'ensemble $L^2(V)$ est l'ensemble de toutes ces fonctions, qui sont en fait des vecteurs de $\mathbb{R}^n$ et on écrit alors

$$L^2(V) \cong \mathbb{R}^n.$$

Quelle est la valeur de n ici?

On a l'isomorphisme suivant

$$L^2(V) \cong \mathbb{R}^n,$$

avec $n = |V| = 3$.

- On rappelle que le gradient est une application

$$\nabla : L^2(V) \rightarrow L^2(\mathcal{E})$$

donnée par

$$(\nabla u)_{i,j} = u(i) - u(j),$$

et où $L^2(\mathcal{E})$ est l'ensemble des fonctions définies sur les arêtes de G . Autrement dit, le gradient transforme des fonctions prenant comme un argument un sommet $i \in V$ en fonctions prenant en argument une arête $(i, j) \in \mathcal{E}$. On considère maintenant la fonction $u \in L^2(V)$, donnée par $u = (1, -1, 2)^T$. Calculer ∇u .

On a

$$\mathcal{E} = \{(1, 2); (2, 1); (1, 3); (3, 1); (2, 3); (3, 2)\}.$$

et le gradient est l'opérateur $\nabla : L^2(V) \rightarrow L^2(\mathcal{E})$. Ainsi, pour $u = \begin{pmatrix} 1 \\ -1 \\ 2 \end{pmatrix}$, on a

$$\cdot (\nabla u)_{(1,2)} = u(1) - u(2) = 1 - (-1) = 2.$$

$$\cdot (\nabla u)_{(1,3)} = u(1) - u(3) = 1 - 2 = -1.$$

$$\cdot (\nabla u)_{(2,3)} = u(2) - u(3) = -1 - 2 = -3.$$

On complète finalement en remarquant que, par antisymétrie, on a

$$(\nabla u)_{(i,j)} = -(\nabla u)_{(j,i)}.$$

- d. L'opérateur de divergence est quant à lui un opérateur $\text{div} : L^2(\mathcal{E}) \rightarrow L^2(V)$ qui transforme des fonctions définies sur les arêtes en fonctions définies sur les sommets. Cet opérateur agit sur $F \in L^2(\mathcal{E})$ par

$$\text{div} F(i) = \sum_{j:(i,j) \in \mathcal{E}} W_{ij} \cdot F_{ij}.$$

Calculer $-\text{div} \nabla u$.

Pour $F \in L^2(\mathcal{E})$, on a

$$\text{div} F(i) = \sum_{j:(i,j) \in \mathcal{E}} W_{ij} \cdot F_{ij}.$$

Ainsi,

$$-\text{div} \nabla u(1) = - \sum_{j:(1,j) \in \mathcal{E}} W_{1j} \cdot \nabla_{1j} u = -(W_{12} \cdot \nabla_{12} u + W_{13} \cdot \nabla_{13} u) = -(2-2) = 0.$$

Similairement,

$$-\text{div} \nabla u(2) = - \sum_{j:(2,j) \in \mathcal{E}} W_{2j} \cdot \nabla_{2j} u = -(W_{21} \cdot \nabla_{21} u + W_{23} \cdot \nabla_{23} u) = -(2+3) = -1.$$

Enfin,

$$-\text{div} \nabla u(3) = - \sum_{j:(3,j) \in \mathcal{E}} W_{3j} \cdot \nabla_{3j} u = -(W_{31} \cdot \nabla_{31} u + W_{32} \cdot \nabla_{32} u) = -2 + 3 = 1.$$

On peut donc conclure que

$$-\text{div} \nabla u = \begin{pmatrix} 0 \\ -1 \\ 1 \end{pmatrix}.$$

- e. Par analogie avec le laplacien usuel, le g -laplacien $\Delta : L^2(V) \rightarrow L^2(V)$ est défini par

$$\Delta u = -\text{div} \nabla u,$$

ce qui revient à

$$\Delta u(i) = - \sum_{j:(i,j) \in \mathcal{E}} W_{ij} (u(i) - u(j)).$$

Vérifier que vous obtenez la même réponse qu'en (d) en utilisant plutôt cette définition pour calculer directement Δu .

On a par définition

$$\Delta u(i) = - \sum_{j:(i,j) \in \mathcal{E}} W_{ij}(u(i) - u(j))$$

avec $u = \begin{pmatrix} 1 \\ -1 \\ 2 \end{pmatrix}$. Nous allons calculer sommet par sommet, individuellement.

D'abord,

$$\begin{aligned} \Delta u(1) &= -(W_{12}(u(1) - u(2)) + W_{13}(u(1) - u(3))) \\ &= -(1(1 - (-1)) + 2(1 - 2)) = -(2 - 2) = 0. \end{aligned}$$

On passe au sommet suivant :

$$\begin{aligned} \Delta u(2) &= -(W_{21}(u(2) - u(1)) + W_{23}(u(2) - u(3))) \\ &= -(1(-2) + (-1) \cdot (-3)) = -1. \end{aligned}$$

Enfin,

$$\begin{aligned} \Delta u(3) &= -(W_{31}(u(3) - u(1)) + W_{32}(u(3) - u(2))) \\ &= -(2 \cdot (2 - 1) - 1 \cdot (2 - (-1))) = 1. \end{aligned}$$

On obtient donc bel et bien le même vecteur de sortie qu'avec l'autre méthode de calcul, soit $\begin{pmatrix} 0 \\ -1 \\ 1 \end{pmatrix}$.

- f. Puisque le g-laplacien transforme linéairement des vecteurs de $L^2(V) \cong \mathbb{R}^n$ en vecteurs de $\mathbb{R}^n$, il peut-être représenté par une matrice. En effet, on a

$$-\Delta = D - W,$$

où W est la matrice adjacente des poids et

$$D = \text{diag}\{\deg(1), \deg(2), \dots, \deg(n)\}.$$

On rappelle $\deg(i) = \sum_j W_{ij}$ retourne la somme de tous les poids qui connectent avec le sommet i . Utiliser cette définition matricielle pour calculer une nouvelle fois Δu .

Dans notre contexte, la définition matricielle du laplacien devient $-\Delta = D - W$, avec la matrice diagonale des degrés obtenue via

$$D = \text{diag}\{\deg(1), \deg(2), \deg(3)\}.$$

Le calcul des degrés nous donne d'abord

- $\deg(1) = W_{12} + W_{13} = 3$,
- $\deg(2) = W_{23} + W_{21} = 0$,
- $\deg(3) = W_{31} + W_{32} = 1$.

Ainsi,

$$-\Delta = \begin{pmatrix} 3 & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 1 \end{pmatrix} - \begin{pmatrix} 0 & 1 & 2 \\ 1 & 0 & -1 \\ 2 & -1 & 0 \end{pmatrix} = \begin{pmatrix} 3 & -1 & 2 \\ -1 & 0 & 1 \\ -2 & 1 & 1 \end{pmatrix}.$$

Une simple multiplication matricielle permet enfin une nouvelle fois de confirmer que

$$-\Delta u = - \begin{pmatrix} 3 & -1 & 2 \\ -1 & 0 & 1 \\ -2 & 1 & 1 \end{pmatrix} \begin{pmatrix} 1 \\ -1 \\ 2 \end{pmatrix} = \begin{pmatrix} 0 \\ -1 \\ 1 \end{pmatrix}.$$

Il est donc clair qu'il existe plusieurs manières distinctes d'envisager le laplacien sur un graphe. Les premières définitions ont l'avantage d'être similaire à la version usuelle continue du laplacien, tandis que la dernière est définitivement plus pratique du point de vue calculatoire. Il s'agit en effet de réduire directement le travail de cet opérateur à la simple multiplication matricielle.

Exercice 6.8 On veut vérifier que le g-laplacien est une matrice symétrique. Soit D une matrice diagonale et W une matrice symétrique. Montrer que $D - W$ est symétrique.

Les éléments hors de la diagonale sont ceux de $-W$ qui est elle-même symétrique. Donc,

$$\Delta = -(D - W)$$

est symétrique. Sinon, en utilisant les propriétés de la transposée, on a

$$\Delta^T = -(D - W)^T = -(D^T - W^T) = -(D - W) = \Delta,$$

où on a respectivement utilisé le fait que tant W que D sont des matrices symétriques.

Exercice 6.9 Calculer la base orthogonale Φ de vecteurs propres du g-laplacien $\Delta = -(D - W)$ de l'exercice 6.7.

Note : pour calculer les vecteurs propres, utilisez un logiciel de calcul!

On veut donc calculer 3 vecteurs propres de

$$\Delta = \begin{pmatrix} -3 & 1 & -2 \\ 1 & 0 & -1 \\ 2 & -1 & -1 \end{pmatrix}$$

Avec un logiciel, on calcule directement

$$\cdot \varphi_1 = \left(-\frac{3+\sqrt{7}}{1+\sqrt{7}}, \frac{2}{1+\sqrt{7}}, 1 \right)^T$$

$$\cdot \varphi_2 = \left(-\frac{3-\sqrt{7}}{1+\sqrt{7}}, \frac{-2}{1+\sqrt{7}}, 1 \right)^T$$

$$\cdot \varphi_3 = (1, 1, 1)^T$$

Donc, une base orthogonale de vecteurs propres est donnée par

$$\Phi = \{\varphi_1, \varphi_2, \varphi_3\}.$$

Exercice 6.10 Nous avons développé une g-théorie de Fourier qui nous a permis de définir la transformée de Fourier et son inverse par

$$\mathcal{F}(f)(k) = \hat{f}(k) = \sum_{i=1}^n f(i) \cdot \phi_k(i)$$

et

$$\mathcal{F}^{-1}(f)(i) = f(i) = \sum_{k=1}^n \hat{f}(k) \cdot \phi_k(i),$$

respectivement. Dans les expressions ci-haut, ϕ_k est le k -ème vecteur propre de Δ . En notation matricielle, on a

$$\hat{f} = \Phi^T f, \quad f = \Phi \hat{f}.$$

- a. On travaille toujours avec le graphe G de l'exercice 6.7. Soit un signal d'entrée $f \in L^2(V)$ donné par $f = (1, 1, 3)^T$. Calculer $\hat{f}$.

On pose Φ la matrice dont les colonnes sont les vecteurs propres de l'exercice précédent et $f = \begin{pmatrix} 1 \\ 1 \\ 3 \end{pmatrix}$. Alors, la transformée de Fourier de f est

$$\hat{f} = \Phi^T f = \begin{pmatrix} \varphi_1 \cdot f \\ \varphi_2 \cdot f \\ \varphi_3 \cdot f \end{pmatrix}.$$

- b. Soit maintenant un filtre $h \in L^2(V)$, donné par $h = (1, 1, 0)$. Calculer $\hat{h}$.

Similairement,

$$\hat{h} = \Phi^T h = \begin{pmatrix} \varphi_1 \cdot h \\ \varphi_2 \cdot h \\ \varphi_3 \cdot h \end{pmatrix}.$$

- c. Calculer la convolution $f * h$.

On utilise finalement le théorème de la convolution pour obtenir

$$\begin{aligned} f * h &= \mathcal{F}^{-1} = (\hat{f} \bullet \hat{h}) \\ &= \Phi \begin{pmatrix} (\varphi_1 \cdot f) \cdot (\varphi_1 \cdot h) \\ (\varphi_2 \cdot f) \cdot (\varphi_2 \cdot h) \\ (\varphi_3 \cdot f) \cdot (\varphi_3 \cdot h) \end{pmatrix}. \end{aligned}$$

7 RÉDUCTION SPECTRALE DE LA DIMENSION

Exercice 7.1 On considère la matrice symétrique A donnée par

$$A = \begin{pmatrix} 1 & 1 & 2 \\ 1 & 2 & -1 \\ 2 & -1 & 1 \end{pmatrix}.$$

- a. Soit le vecteur $u = \begin{pmatrix} 1 \\ 2 \\ 1 \end{pmatrix}$. Calculer le quotient de Rayleigh $\mathcal{R}_A(u)$.

On rappelle que le quotient de Rayleigh est défini par

$$\mathcal{R}_A(u) = \frac{u^T A u}{u^T u}.$$

On calcule d'abord $Au = \begin{pmatrix} 5 \\ 5 \\ 1 \end{pmatrix}$. Donc

$$u^T A u = 1 \cdot 5 + 2 \cdot 5 + 1 \cdot 1 = 16.$$

De plus,

$$u^T u = \|u\|^2 = 1^2 + 2^2 + 1^2 = 6.$$

Donc, le quotient de Rayleigh demandé est $\mathcal{R}_A(u) = \frac{16}{6} = \frac{8}{3}$.

- b. À l'aide d'un logiciel de calcul, trouvez les modes propres (i.e. les vecteurs et valeurs propres) de la matrice A .

On obtient

$$\lambda_1 = 3, \lambda_2 \approx 2.56, \lambda_3 \approx -1.56$$

et des vecteurs propres correspondants sont

$$\varphi_1 = \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix}, \varphi_2 = \begin{pmatrix} -1 \\ -3.56 \\ 1 \end{pmatrix}, \varphi_3 = \begin{pmatrix} -1 \\ 0.56 \\ 1 \end{pmatrix}.$$

On remarque enfin qu'on peut vérifier que ces vecteurs propres sont orthogonaux entre eux.

- c. Pour chacun de ces vecteurs propres φ_i , $i = 1, 2, 3$, calculer le quotient de Rayleigh $\mathcal{R}_A(\varphi_i)$. Que remarquez-vous?

On calcule directement

$$A\varphi_1 = 3\varphi_1 = \begin{pmatrix} 3 \\ 0 \\ 3 \end{pmatrix}$$

et donc $\varphi_1^T A\varphi_1 = 6$ de telle sorte que

$$\mathcal{R}_A(\varphi_1) = \frac{\varphi_1^T A\varphi_1}{\varphi_1^T \varphi_1} = \frac{6}{2} = 3 = \lambda_1.$$

De manière générale, si on a $Au = \lambda u$, alors on obtient

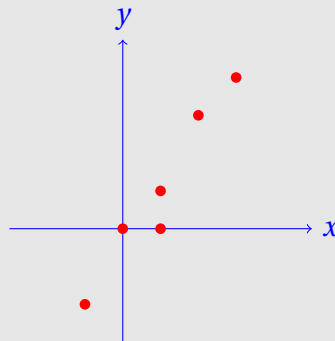
$$\mathcal{R}_A(u) = \frac{u^T Au}{u^T u} = \frac{u^T \lambda u}{u^T u} = \lambda \frac{u^T u}{u^T u} = \lambda.$$

Exercice 7.2 On considère le nuage de points $x^{(i)}$, $i = 1, 2, \dots, k$ suivant dans $\mathbb{R}^2$:

$$N = \{(1, 1); (2, 3); (0, 0); (-1, -2), (3, 4); (1, 0)\}.$$

- a. Représenter graphiquement le nuage de points dans le plan.

Notre nuage de points est le suivant :



- b. Procéder à l'analyse en composantes principales de N . Combien de composantes principales possède-t-il ?

On forme d'abord la matrice

$$X = \begin{pmatrix} x^{(1)} \\ x^{(2)} \\ \dots \\ x^{(6)} \end{pmatrix} = \begin{pmatrix} 1 & 1 \\ 2 & 3 \\ 0 & 0 \\ -1 & -2 \\ 3 & 4 \\ 1 & 0 \end{pmatrix}.$$

On peut ensuite considérer la matrice carrée de dimension 2

$$X^T X = \begin{pmatrix} 1 & 2 & 0 & -1 & 3 & 1 \\ 1 & 3 & 0 & -2 & 4 & 0 \end{pmatrix} \begin{pmatrix} 1 & 1 \\ 2 & 3 \\ 0 & 0 \\ -1 & -2 \\ 3 & 4 \\ 1 & 0 \end{pmatrix} = \begin{pmatrix} 16 & 21 \\ 21 & 30 \end{pmatrix}.$$

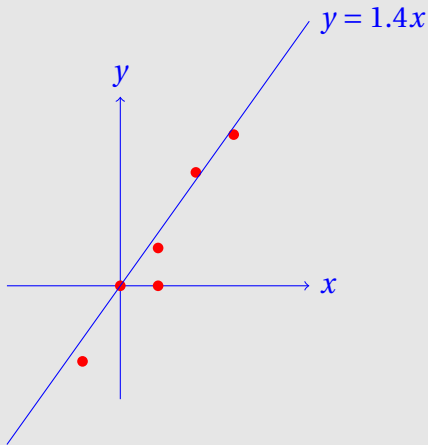
*Les composantes principales $w^{(1)}$ et $w^{(2)}$ sont les espaces propres de $X^T X$.
On a*

$$w^{(1)} = \begin{pmatrix} 0.7 \\ 1 \end{pmatrix}, \quad w^{(2)} = \begin{pmatrix} -1.4 \\ 1 \end{pmatrix}.$$

- c. Utiliser la première composante principale pour définir l'application de réduction de la dimension

$$\Pi_1 : \mathbb{R}^2 \rightarrow \mathbb{R}.$$

L'espace engendré par la première composante principale est la droite L d'équation $y = 1.4 \cdot x$.



L'application de réduction de dimension est la projection orthogonale sur cette droite. Plus précisément, on a

$$\Pi_1 : \mathbb{R}^2 \rightarrow L \cong \mathbb{R}$$

avec

$$x^{(i)} \mapsto \text{proj}_L x^{(i)}.$$

- d. Utiliser la projection Π_1 pour projeter le points $x^{(3)}$ sur un espace de dimension 1.

Puisque $x^{(3)}$ appartient à la droite L , sa projection sur celle-ci est triviale. On a donc

$$\text{proj}_L x^{(3)} = (0, 0).$$

Exercice 7.3 On considère un nuage contenant m points de $\mathbb{R}^n$.

- a. Au plus, combien de composantes principales possède-t-il?

Au plus n composantes principales.

- b. On veut définir une projection Π_k pour réduire la dimension vers $\mathbb{R}^k$. Quelles sont les valeurs possibles pour k ?

Il faut bien sûr que k soit strictement plus petit que n .

Exercice 7.4 On considère le nuage contenant les 4 points suivants

$$(1, 2); (2, 1); (0, 0); (1, 0).$$

On se propose de développer l'algorithme de Belkin-Niyagi pour réduire la dimension de ces points en les projetant sur un espace de dimension 1.

- a. Calculer les distances entre les points et déterminer si chaque paire potentielle fait partie de l'ensemble $\mathcal{E}$ des arêtes en fixant le seuil à $\epsilon = 2$.

On calcule les distances

$$d(i, j) = \|x^{(i)} - x^{(j)}\|$$

entre $x^{(i)}$ et $x^{(j)}$ pour toute paire i, j avec $i \neq j$. On a

- $d(1, 2) = ((1-2)^2 + (2-1)^2)^{\frac{1}{2}} = \sqrt{2} \approx 1.41 \leq \epsilon = 2 \quad \checkmark \quad (1, 2) \in \mathcal{E}.$
- $d(1, 3) = (2^2 + 1^2)^{\frac{1}{2}} = \sqrt{5} > \epsilon = 2.$
- $d(1, 4) = ((1-1)^2 + (2-0)^2)^{\frac{1}{2}} = 2 \leq \epsilon = 2 \quad \checkmark \quad (1, 4) \in \mathcal{E}.$
- $d(2, 3) = \sqrt{5} > \epsilon.$
- $d(2, 4) = ((2-1)^2 + (1-0)^2)^{\frac{1}{2}} = \sqrt{2} \approx 1.41 \leq \epsilon = 2 \quad \checkmark \quad (2, 4) \in \mathcal{E}.$
- $d(3, 4) = 1 < \epsilon \quad \checkmark \quad (3, 4) \in \mathcal{E}.$

- b. En utilisant les noyaux de la chaleur, assigner un poids à chacun des éléments de $\mathcal{E}$.

On rappelle que la technique du noyau de la chaleur pour fournir un poids entre deux sommets connectés d'un graphe vise à établir une analogie avec la quantité de chaleur qui pourrait se transmettre entre les deux. On définit la quantité suivante

$$W_{ij} = e^{-\frac{d^2(i,j)}{t}},$$

où $t > 0$ est le temps écoulé de diffusion de la chaleur. On fixe ici

$$t = 1$$

pour simplifier les calculs. On a donc

- $W_{12} = e^{-2} \approx 0.14.$
- $W_{14} = e^{-4} \approx 0.02.$
- $W_{24} = e^{-2} \approx 0.14.$
- $W_{34} = e^{-1} \approx 0.37.$

- c. Calculer les degrés des sommets pour construire la matrice diagonale des degrés, la matrice d'adjacence des poids et enfin trouver les modes propres du laplacien normalisé Δ sur ce graphe.

On calcule d'abord les degrés :

- $\deg(1) = 0.14 + 0.02 = 0.16.$
- $\deg(2) = 0.14 + 0.14 = 0.28.$
- $\deg(3) = 0.37.$
- $\deg(4) = 0.02 + 0.14 + 0.37 = 0.53.$

La matrice D est simplement la matrice diagonale de dimension 4 telle que

$$D_{ii} = \deg(i), \quad i = 1, \dots, 4.$$

Donc $\Delta = D - W$ et on considère l'équation aux valeurs propres normalisée sur ce graphe

$$\Delta\varphi = \lambda D\varphi.$$

On note d'abord que cette équation peut-être réécrite ainsi

$$(D^{-1}\Delta)\varphi = \lambda\varphi.$$

On calcule d'abord

$$\begin{aligned} D^{-1}\Delta &= \begin{pmatrix} (0,16)^{-1} & 0 & 0 & 0 \\ 0 & (0,28)^{-1} & 0 & 0 \\ 0 & 0 & (0,37)^{-1} & 0 \\ 0 & 0 & 0 & (0,57)^{-1} \end{pmatrix} \\ &\cdot \begin{pmatrix} 0,16 & -0,14 & 0 & -0,02 \\ -0,14 & 0,28 & 0 & -0,14 \\ 0 & 0 & 0,37 & -0,37 \\ -0,02 & -0,14 & -0,37 & 0,53 \end{pmatrix} \\ &= \begin{pmatrix} 1 & -0,5 & 0 & -0,04 \\ -0,88 & 1 & 0 & -0,26 \\ 0 & 0 & 1 & -0,7 \\ -0,13 & -0,5 & -1 & 1 \end{pmatrix}. \end{aligned}$$

On obtient

$$0 = \lambda_0 < \lambda_1 \approx 0.46 < \lambda_2 \approx 1.6 < \lambda_3 = 2.$$

Les vecteurs propres correspondants sont

$$\varphi_0 = \begin{pmatrix} 0.2 \\ 0.3 \\ 0.5 \\ 0.7 \end{pmatrix}, \quad \varphi_1 = \begin{pmatrix} 0.5 \\ 0.5 \\ -0.6 \\ -0.4 \end{pmatrix}, \quad \varphi_2 = \begin{pmatrix} 0.5 \\ -0.6 \\ 0.5 \\ -0.4 \end{pmatrix}, \quad \varphi_3 = \begin{pmatrix} 0.1 \\ -0.3 \\ -0.5 \\ 0.7 \end{pmatrix}.$$

d. Définir la projection spectrale ainsi obtenue.

Suivant l'algorithme de Belkin-Niyogi, on définit la projection $\Pi_1 : \mathbb{R}^2 \rightarrow \mathbb{R}$ par

$$x^{(i)} \mapsto \varphi_1(i).$$

8 AUTOENCODEURS VARIATIONNELS

Exercice 8.1 Calculer l'entropie H des distributions probabilistes suivantes. On rappelle que l'entropie d'une distribution discrète est donnée par

$$H(p) = - \sum_i p(x_i) \log_2(p(x_i)).$$

Pour la version continue supportée sur $\mathbb{R}$, il faut remplacer la somme par une intégrale et on obtient

$$H(p) = - \int_{\mathbb{R}} p(x) \log_2(p(x)) dx.$$

- a. La distribution du lancer d'une pièce de monnaie équilibrée.

On définit les événements

$$x_1 = \text{"pile"}, x_2 = \text{"face"}.$$

On a évidemment

$$p(x_1) = p(x_2) = \frac{1}{2}.$$

Donc, l'entropie H de p est donnée par

$$H(p) = \frac{1}{2} \log_2(2) + \frac{1}{2} \log_2(2) = 1.$$

- b. La distribution d'un lancer de dé à 6 faces.

On définit l'événement

$$x(i) = \text{"obtenir } i\text{"},$$

et on a $p(x_i) = \frac{1}{6}$. L'entropie H de la distribution p est donnée par

$$H(p) = - \sum_{i=1}^6 p(x_i) \log_2 p(x_i) = 6 \left(\frac{1}{6} \log_2(6) \right) \approx 2.6.$$

- c. La distribution $p(x)$ décrite par la fonction de masse suivante

$$p(x_1) = \frac{1}{8}, p(x_2) = \frac{3}{8}, p(x_3) = \frac{1}{2}.$$

On calcule directement

$$\begin{aligned} H(p) &= - \sum_{i=1}^3 p(x_i) \log(p(x_i)) \\ &= \frac{1}{8} \log_2(8) = \frac{3}{8} \log_2\left(\frac{8}{3}\right) + \frac{1}{2} \log_2(2) \\ &= \frac{1}{8} \cdot 3 + \frac{3}{8} \cdot (1.41) + \frac{1}{2} \cdot 2 \\ &\approx 1,4. \end{aligned}$$

d. Une distribution discrète uniforme supportée sur n événements, i.e.

$$p(x_i) = \frac{1}{n}, \text{ pour } i = 1, 2, \dots, n.$$

On utilise l'uniformité de la distribution pour directement calculer

$$H(p) = - \sum_{i=1}^n p(x_i) \log_2(p(x_i)) = n \left(\frac{1}{n} \log_2(n) \right) = \log_2(n).$$

Dans un tel contexte de distribution uniforme, on constate donc que l'entropie est monotone croissante en n .

e. Une gaussienne $N(0, 1)$, c'est-à-dire une loi normale centrée réduite.

La densité $p(x)$ de la normale centrée réduite $N(0, 1)$ est

$$p(x) = \frac{1}{\sqrt{2\pi}} e^{-\frac{x^2}{2}}.$$

Ainsi, l'entropie de cette gaussienne est

$$- \int_{\mathbb{R}} p(x) \log(p(x)) dx \approx 1.42.$$

Exercice 8.2 Soit les distributions discrètes p, q, r supportées sur l'ensemble $\{1, 2, 3\}$ et respectivement données par les fonctions de masse suivantes

$$p(1) = p(2) = p(3) = \frac{1}{3},$$

$$q(1) = q(3) = \frac{1}{6}, q(2) = \frac{2}{3}$$

et

$$r(1) = r(2) = \frac{1}{12}, r(3) = \frac{5}{6}$$

On rappelle que l'entropie croisée d'une distribution q par rapport à une distribution p est donnée par

$$H(p, q) = - \sum_i p(x_i) \log_2(q(x_i)).$$

De plus, la divergence de Kullback-Leibler est donnée par

$$D_{KL}(p||q) = H(p, q) - H(p).$$

Calculez les quantités demandées.

a. $H(p, q)$.

On calcule directement

$$\begin{aligned} H(p, q) &= -E_{x \sim p}(\log_2 q) = - \sum_{i=1}^3 p(x_i) \log_2(q(x_i)) \\ &= \frac{1}{3} \left(\log_2 12 + \log_2 12 + \log_2 \frac{6}{5} \right) \\ &\approx 2.48. \end{aligned}$$

b. $H(p, r)$.

Cette fois-ci, on a

$$\begin{aligned} H(p, q) &= -E_{x \sim p}(\log_2 q) = - \sum_{i=1}^3 p(x_i) \log_2(q(x_i)) \\ &= \frac{1}{3} \left(\log_2 6 + \log_2 6 + \log_2 \frac{3}{2} \right) \\ &\approx 1.92. \end{aligned}$$

c. $H(q, p)$.

Cette fois-ci, on a

$$\begin{aligned} H(p, q) &= -E_{x \sim q}(\log_2 p) = - \sum_{i=1}^3 q(x_i) \log_2(p(x_i)) \\ &= \frac{1}{6} \log_2(3) + \frac{2}{3} \log_2(3) + \frac{1}{6} \log_2(3) \\ &\approx 1.58. \end{aligned}$$

d. $D_{KL}(p||p)$.

On calcule directement la KL-divergence

$$D_{KL}(p||p) = H(p, p) - H(p) = H(p) - H(p) = 0.$$

e. $D_{KL}(p||q)$.

Ici, on a

$$D_{KL}(p||q) = H(p, q) - H(p) = 1.92 - \log_2(3) \approx 0.36$$

f. $D_{KL}(q||p)$.

On a

$$D_{KL}(q||p) = H(q, p) - H(q) \approx 0.33$$

g. $D_{KL}(p||r)$.

Finalement,

$$D_{KL}(p||r) = H(p, r) - H(p) \approx 0.9.$$

Exercice 8.3 Donnez les 4 distributions théoriques sous-jacentes au développement d'un autoencodeur variationnel et indiquez de quelle manière elles sont reliées entre elles par le théorème de Bayes.

On a les distributions théoriques suivantes

- $p(x)$: évidence.
- $p(x|z)$: vraisemblance.
- $p(z|x)$: postérieure.
- $p(z)$: antérieure.

Dans ce contexte, on considère x comme une donnée et z comme une variable latente. C'est le théorème de Bayes qui nous permet de relier ces distributions théoriques, via la formule

$$p(z|x) = \frac{p(x|z)p(z)}{p(x)}.$$