

Automatic Differentiation

On peut calculer des dérivées partielles de différentes manières :

1. De façon symbolique, en fixant une des variables et en dérivant les autres soit à la main, soit par ordinateur.
2. De façon numérique, avec la formule $f'(x) \approx (f(x+h) - f(x))/h$.
3. De façon algorithmique, soit forward, soit reverse, c'est ce que nous verons ici.

Pour illustrer, nous utiliserons l'exemple de classification de points de deux formes de lunes.

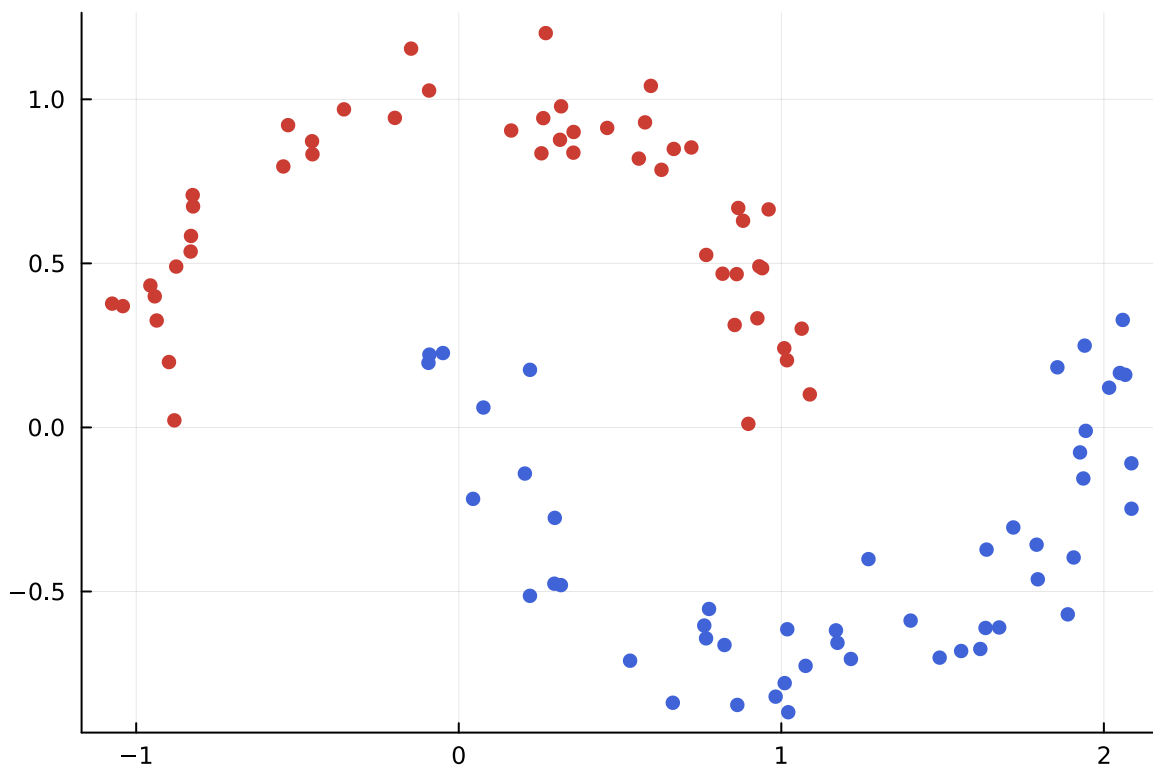
```
► (      x1      x2      , CategoricalArrays.CategoricalVector{Int64, UInt32, Int64, Categori
```

1	-0.198003	0.94344
2	0.220852	0.175566
3	1.90637	-0.396334
4	-0.875963	0.490281
5	-0.898384	0.19927
6	0.256341	0.835451
7	2.0663	0.160385
8	1.67547	-0.609588
9	1.85594	0.183258
10	0.866492	0.668861

... more

```
1 X_table, y_cat = MLJBase.make_moons(100, noise=0.1)
```

plot_w (generic function with 2 methods)



```
1 plot_w()
```

Nous travaillerons avec une matrice `X` contenant dans chaque ligne, les coordonnées d'un point.

```
X = 100x2 Matrix{Float64}:
-0.198003  0.94344
 0.220852  0.175566
 1.90637   -0.396334
-0.875963  0.490281
-0.898384  0.19927
 0.256341  0.835451
 2.0663    0.160385
 ⋮
 0.0444141 -0.217693
-0.824789  0.708108
 0.863596  -0.845744
-0.0939089 0.196828
 0.761188  -0.603782
-0.0919193 1.02678
```

```
1 X = Tables.matrix(X_table)
```

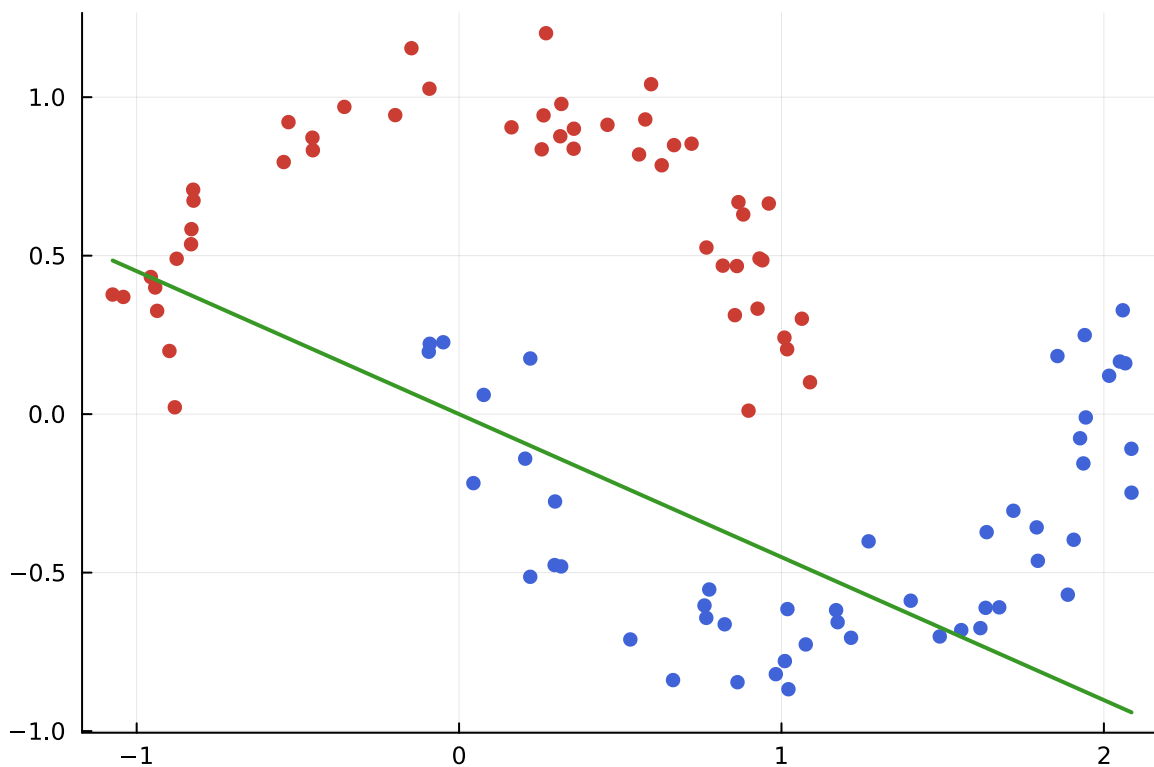
Le vecteur `y` contiendra 1 pour les points de la lune bleue et -1 pour les points de la lune rouge.

```
y =
► [-1.0, 1.0, 1.0, -1.0, -1.0, -1.0, 1.0, 1.0, 1.0, -1.0, 1.0, -1.0, 1.0, 1.0, 1.0, -1.0, 1.0,
1 y = 2(float.(y_cat.refs) .- 1.5)
```

Nous illustrons le calcul de dérivée automatique par l'entraînement du modèle linéaire $y \approx Xw$.
Commençons avec des poids aléatoires. Le modèle n'est pour le moment pas très précis comme il est aléatoire.

```
w = ▶ [0.366916, 0.813359]
```

```
1 w = rand(2)
```



```
1 plot_w(w)
```

En effet, les prédictions ne correspondent pas à y .

```
y_est =
```

```
▶ [0.694705, 0.223833, 0.377116, 0.077369, -0.167554, 0.773577, 0.888609, 0.118942, 0.830021]
```

```
1 y_est = X * w
```

On peut regrouper les erreurs des estimations de tous les points en les comparant avec y .

```
errors =
```

```
▶ [1.6947, -0.776167, -0.622884, 1.07737, 0.832446, 1.77358, -0.111391, -0.881058, -0.169971]
```

```
1 errors = y_est - y
```

Essayons de trouver des poids w qui minimisent la somme des carrés des erreurs (aka MSE) :

```
mean_squared_error = 1.7540528575277128
```

```
1 mean_squared_error = sum(errors.^2) / length(errors)
```

```
mse (generic function with 1 method)
```

```
1 mse(w, X, y) = sum((X * w - y).^2 / length(y))
```

```
1.754052857527713
```

```
1 mse(w, X, y)
```

Forward Differentiation ⇄

Commençons par définir la forward differentiation. Cette différentiation algorithmique se base sur l'observation que la chain rule permet de calculer la dérivée de n'importe quelle fonction dès lors qu'on connaît son gradient et la dérivée de chacun de ses paramètres. En d'autres mots, supposons qu'on doive calculer

$$\frac{\partial}{\partial x} f(g(x), h(x))$$

Supposons que la fonction f soit une fonction $f(a, b)$ simple (telle que $+$, $*$, $-$) dont on connaît la formule des dérivées partielles $\partial f / \partial a$ en fonction de a et $\partial f / \partial b$ en fonction de b : La chain rule nous donne

$$\frac{\partial}{\partial x} f(g(x), h(x)) = \frac{\partial f}{\partial a}(g(x), h(x)) \frac{\partial g}{\partial x} + \frac{\partial f}{\partial b}(g(x), h(x)) \frac{\partial h}{\partial x}$$

Pour calculer cette expression, ils nous faut les valeurs de $g(x)$ et $h(x)$ ainsi que les dérivées $\partial g / \partial x$ et $\partial h / \partial x$.

Dual

```
1 begin
2     struct Dual{T}
3         value::T
4         derivative::T
5     end
6     Dual(x, y) = Dual{typeof(x)}(x, convert(typeof(x), y))
7 end
```

L'implémentation générique du produit de matrice va appeler `zero`:

```
1 Base.zero(::Dual{T}) where {T} = Dual(zero(T), zero(T))
```

Par linéarité de la dérivée:

```
1 begin
2   Base.:*(α::T, x::Dual{T}) where {T} = Dual(α * x.value, α * x.derivative)
3   Base.:*(x::Dual{T}, α::T) where {T} = Dual(x.value * α, x.derivative * α)
4   Base.:+(x::Dual{T}, y::Dual{T}) where {T} = Dual(x.value + y.value,
5     x.derivative + y.derivative)
6   Base.:-(x::Dual{T}, y::T) where {T} = Dual(x.value - y, x.derivative)
7   Base.:/(x::Dual, α::Number) = Dual(x.value / α, x.derivative / α)
8 end
```

Par la product rule $(fg)' = f'g + fg'$:

```
1 Base.:*(x::Dual{T}, y::Dual{T}) where {T} = Dual(x.value * y.value, x.value *
2   y.derivative + x.derivative * y.value)
```

Pour l'exponentiation, on peut juste se rabattre sur le produit qu'on a déjà défini :

```
1 Base.:^(x::Dual, n::Integer) = Base.power_by_squaring(x, n)
```

```
► OneHotVector(::UInt32): [true, false]
```

```
1 onehot(1, 1:2)
```

```
► [1.0, 0.0]
```

```
1 float.(onehot(1, 1:2))
```

```
► [0.0, 1.0]
```

```
1 float.(onehot(2, 1:2))
```

```
► [Dual(0.366916, 1.0), Dual(0.813359, 0.0)]
```

```
1 Dual.(w, onehot(1, 1:2))
```

```
► Dual(1.7540528575277128, -0.4139785483292133)
```

```
1 mse(Dual.(w, onehot(1, 1:2)), X, y)
```

forward_diff (generic function with 1 method)

```
1 function forward_diff(loss, w, X, y, i)
2   loss(Dual.(w, onehot(i, eachindex(w))), X, y).derivative
3 end
```

```
-0.4139785483292133
```

```
1 forward_diff(mse, w, X, y, 1)
```

forward_diff (generic function with 2 methods)

```
1 function forward_diff(loss, w, X, y)
2     [forward_diff(loss, w, X, y, i) for i in eachindex(w)]
3 end
```

Gradient descent ⇔

Cauchy-Schwarz inequality ⇔

$$\begin{aligned}\left(\sum_i x_i y_i\right)^2 &= \left(\sum_i x_i^2\right) \left(\sum_i y_i^2\right) \cos(\theta)^2 \\ \sum_i x_i y_i &= \sqrt{\sum_i x_i^2} \sqrt{\sum_i y_i^2} \cos(\theta) \\ \langle x, y \rangle &= \|x\|_2 \|y\|_2 \cos(\theta) \\ -\|x\|_2 \|y\|_2 &\leq \langle x, y \rangle \leq \|x\|_2 \|y\|_2\end{aligned}$$

Minimum atteint lorsque $x = -y$ et maximum atteint lorsque $x = y$. Dans les deux cas, x est **parallèle** à y , mais ce sont des **sens** différents.

Rappel dérivée directionnelle ⇔

Dérivée dans la direction d :

$$\langle d, \nabla f \rangle = d^\top \nabla f = d_1 \cdot \partial f / \partial x_1 + \dots + d_n \cdot \partial f / \partial x_n$$

Etant donné un gradient ∇f , la direction d telle que $\|d\|_2 = 1$ qui a une dérivée minimale est $d = -\nabla f$.

Gradient:

$\nabla = \blacktriangleright [-0.413979, 1.50562]$

```
1 ∇ = forward_diff(mse, w, X, y)
```

Gradient line search: pendant combien de temps doit-on suivre le gradient ?

`step_sizes = -1.0:0.2222222222222222:1.0`

```
1 step_sizes = range(-1, stop=1, length=num_η)
```

```
losses =  
► [0.585649, 0.625811, 0.791392, 1.08239, 1.49881, 2.04065, 2.70791, 3.50058, 4.41868, 5.462
```

```
1 losses = [mse(w +  $\eta$  *  $\nabla$ , X, y) for  $\eta$  in step_sizes]
```

```
best_idx = 1
```

```
1 best_idx = argmin(losses)
```

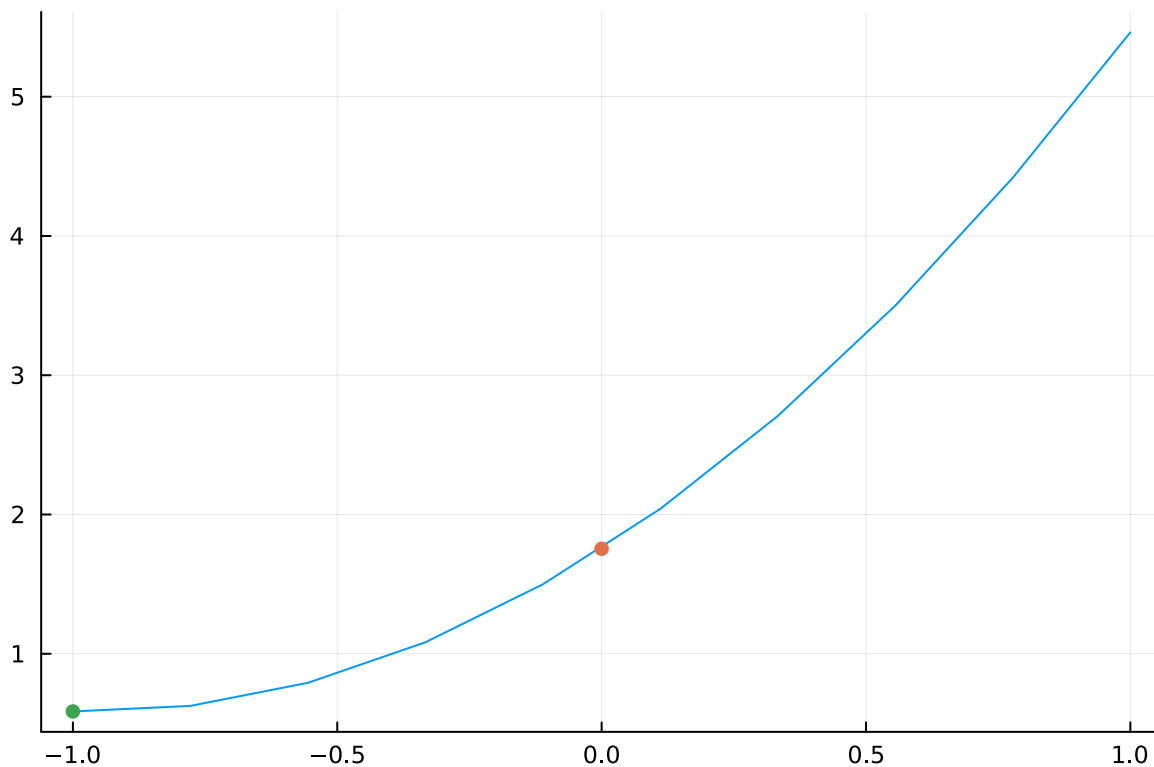
```
 $\eta_{\text{best}}$  = -1.0
```

```
1  $\eta_{\text{best}}$  = step_sizes[best_idx]
```

```
best_loss = 0.5856491142434972
```

```
1 best_loss = losses[best_idx]
```

 10



```
1 begin  
2   plot(step_sizes, losses, label = "")  
3   scatter!([0.0], [mse(w, X, y)], markerstrokewidth = 0, label = "")  
4   scatter!([ $\eta_{\text{best}}$ ], [best_loss], markerstrokewidth = 0, label = "")  
5 end
```

```
w_improved = ► [0.780895, -0.692262]
```

```
1 w_improved = w +  $\eta_{\text{best}}$  *  $\nabla$ 
```

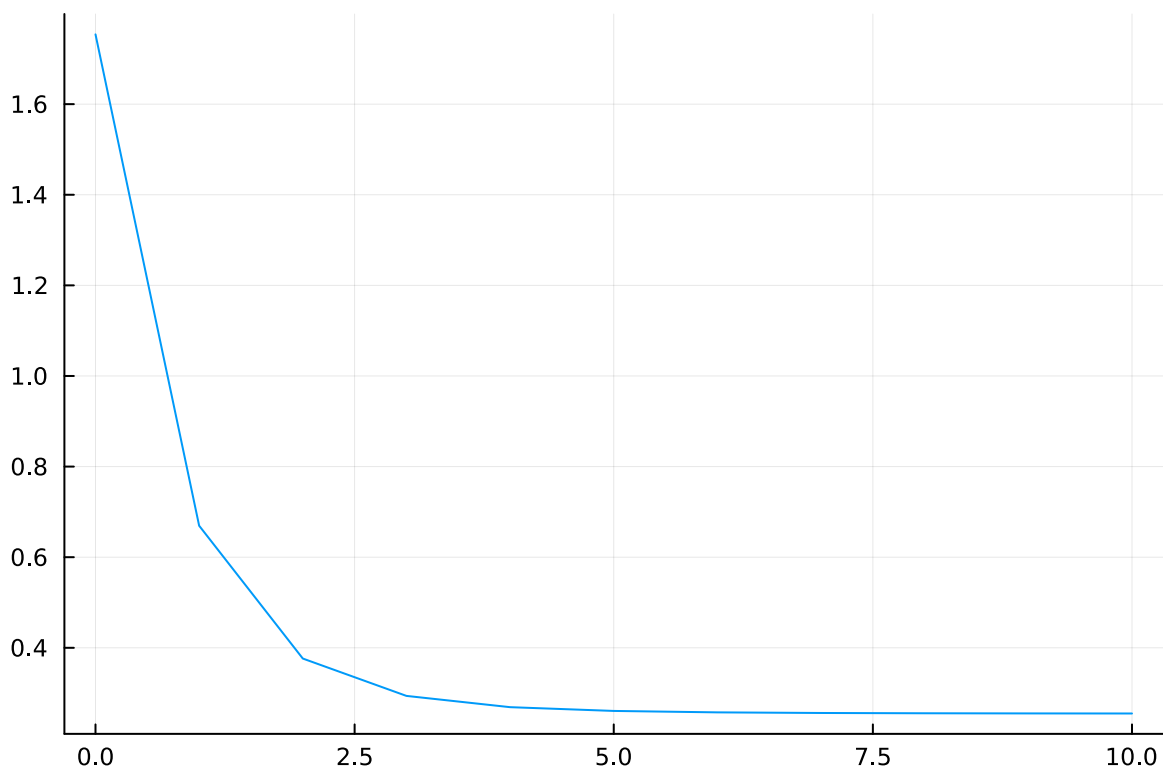
train! (generic function with 1 method)

```
1 function train!(diff, loss, w0, X, y, η, num_iters)
2     w = copy(w0)
3     training_losses = [loss(w, X, y)]
4     for _ in 1:num_iters
5         ∇ = diff(loss, w, X, y)
6         w .= w .- η .* ∇
7         push!(training_losses, loss(w, X, y))
8     end
9     return w, training_losses
10 end
```

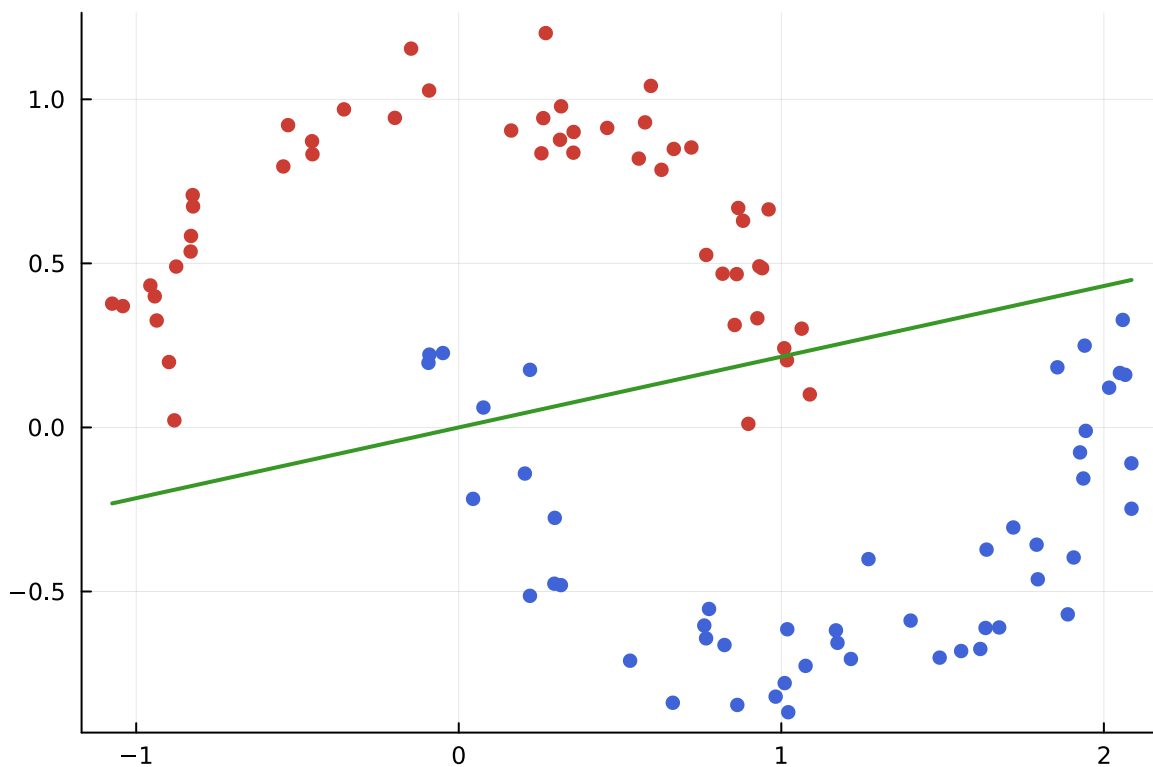
► ([0.258667, -1.20017], [1.75405, 0.669498, 0.376316, 0.293834, 0.26897, 0.260656, 0.257498])

```
1 w_trained, training_losses = train!(forward_diff, mse, w, X, y, 0.7, num_iters)
```

num_iters = 10



```
1 plot(0:num_iters, training_losses, label = "")
```



```
1 plot_w(w_trained)
```

Kernel trick ⇔

lift (generic function with 1 method)

```
1 lift(x) = [1.0, x[1], x[2], x[1]^2, x[1] * x[2], x[2]^2, x[1]^3, x[1]^2 * x[2],
            x[1] * x[2]^2, x[2]^3]
```

X_lift =

100x10 Matrix{Float64}:

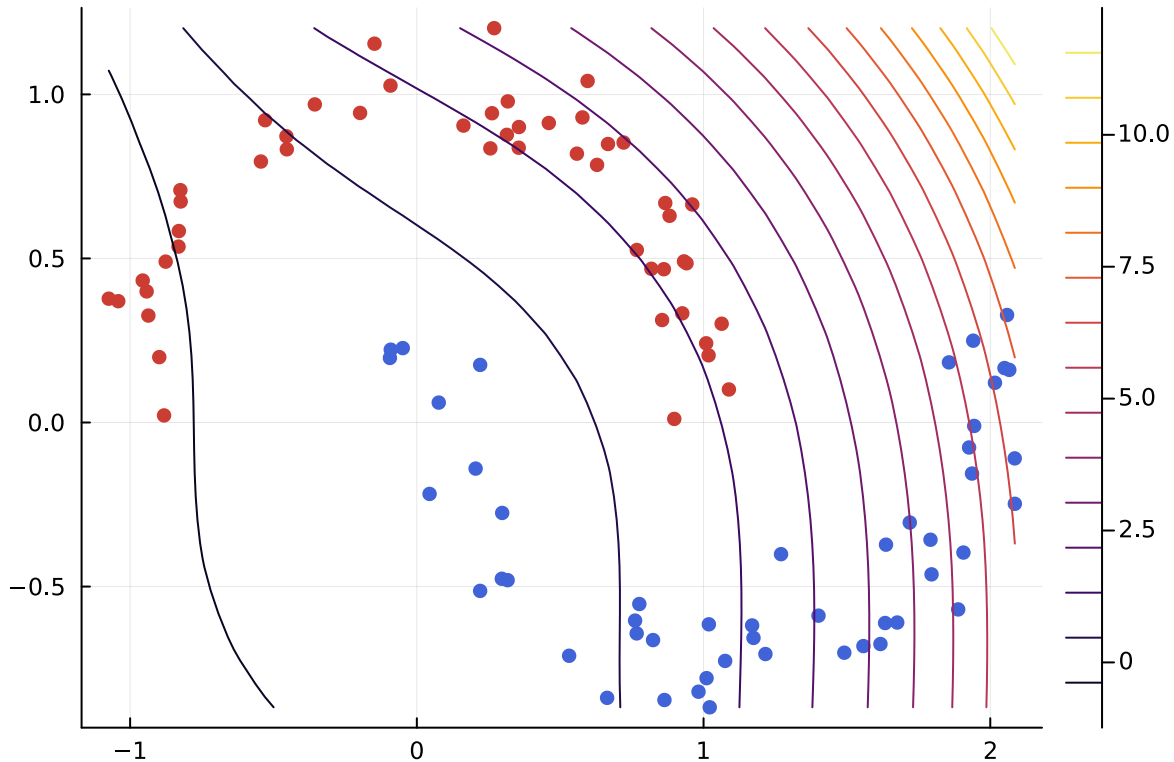
1.0	-0.198003	0.94344	0.039205	...	0.0369876	-0.176238	0.839735
1.0	0.220852	0.175566	0.0487758	...	0.00856339	0.00680746	0.00541159
1.0	1.90637	-0.396334	3.63424	...	-1.44037	0.299454	-0.0622565
1.0	-0.875963	0.490281	0.767311	...	0.376198	-0.21056	0.117851
1.0	-0.898384	0.19927	0.807095	...	0.16083	-0.0356736	0.00791275
1.0	0.256341	0.835451	0.0657108	...	0.0548982	0.178921	0.583126
1.0	2.0663	0.160385	4.26958	...	0.684776	0.0531519	0.00412562
⋮				⋮			
1.0	0.0444141	-0.217693	0.00197261	...	-0.000429423	0.00210479	-0.0103165
1.0	-0.824789	0.708108	0.680277	...	0.481709	-0.413563	0.355057
1.0	0.863596	-0.845744	0.745798	...	-0.630754	0.617715	-0.604946
1.0	-0.0939089	0.196828	0.00881888	...	0.0017358	-0.00363814	0.00762532
1.0	0.761188	-0.603782	0.579407	...	-0.349835	0.277493	-0.220111
1.0	-0.0919193	1.02678	0.00844916	...	0.00867541	-0.096908	1.0825

```
1 X_lift = reduce(vcat, transpose.(lift.(eachrow(X))))
```

```
w_lift =
```

```
► [0.0369498, 0.34274, 0.228468, 0.199867, 0.395805, 0.512978, 0.578332, 0.161477, 0.608745
```

```
1 w_lift = rand(size(X_lift, 2))
```



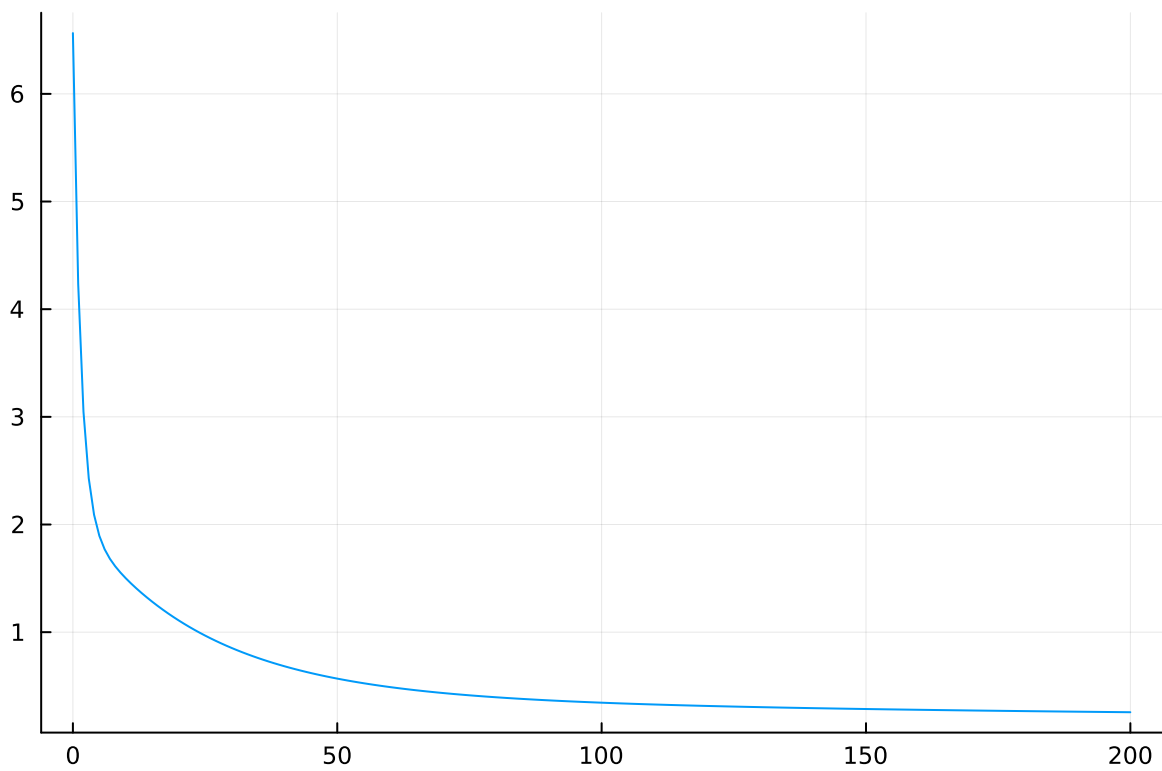
```
1 plot_w(w_lift)
```

$\eta_{lift} =$ 0.01

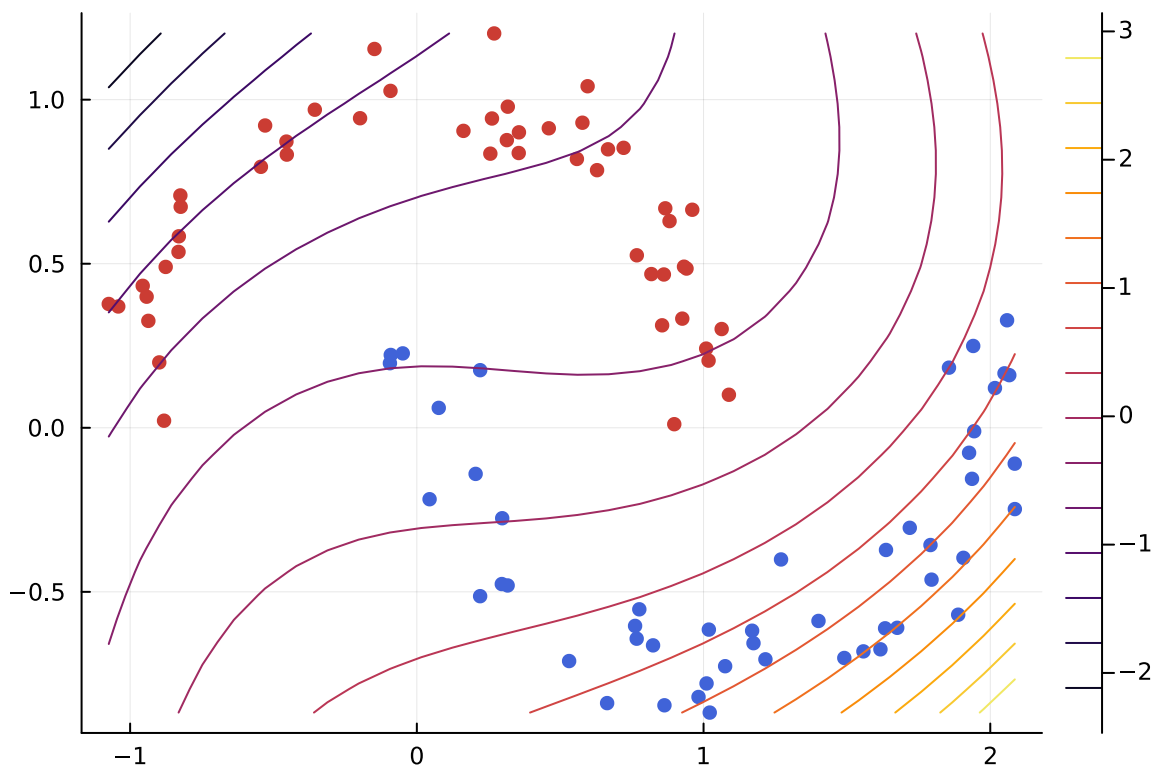
num_iters_lift = 200

```
► ([-0.237495, 0.0163016, -0.689397, -0.182725, -0.0727689, 0.0976839, 0.216581, -0.147087,
```

```
1 w_lift_trained, training_losses_lift = train!(forward_diff, mse, w_lift, X_lift, y,  
   $\eta_{lift}$ , num_iters_lift)
```



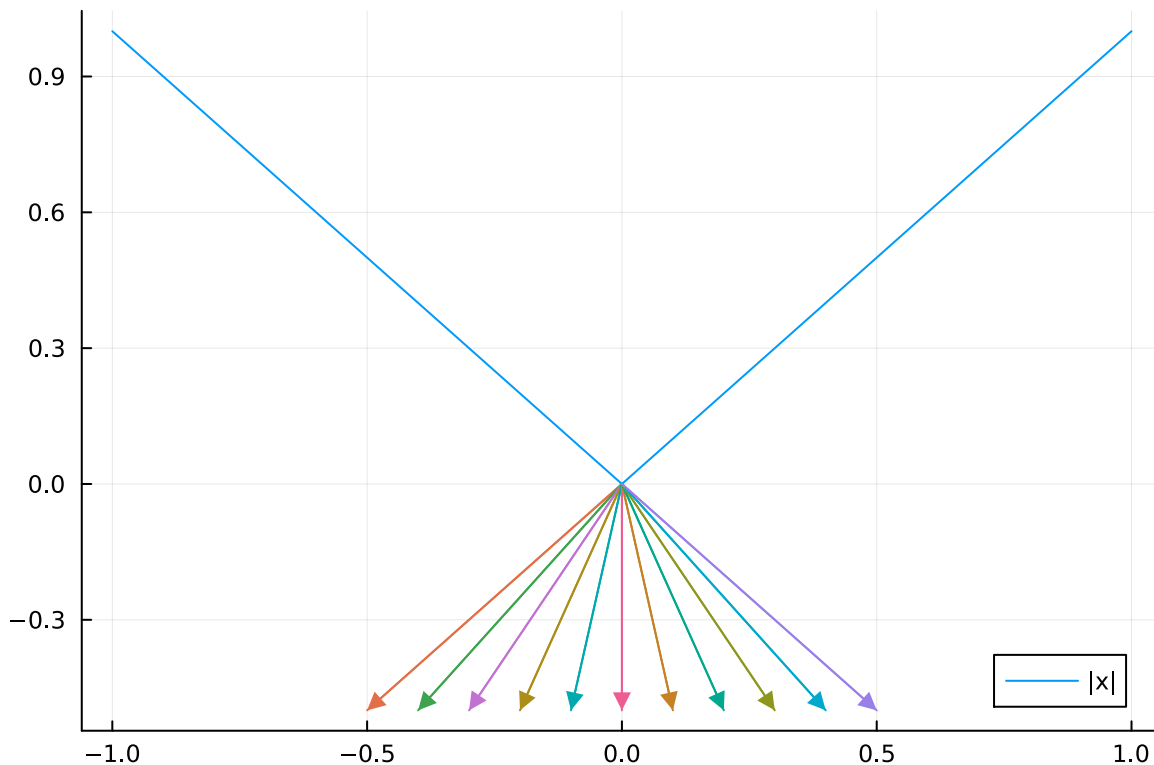
```
1 plot(0:num_iters_lift, training_losses_lift, label = "")
```



```
1 plot_w(w_lift_trained)
```

L1 norm $\Leftrightarrow$

La fonction $|x|$ n'est pas différentiable lorsque $x = 0$. Si on s'approche par la gauche (c'est à dire $x < 0$, la fonction est $-x$) donc la dérivée vaut -1 . Si on s'approche par la droite (c'est à dire $x > 0$, la fonction est x) donc la dérivée vaut 1 . Il n'y a pas de gradient valide ! Par contre, n'importe quel nombre entre -1 et 1 est un **subgradient** valide ! Alors que le gradient est la normale à la tangente **unique**, le subgradient est un élément du **cone tangent**.



```
1 let
2   x = range(-1, stop = 1, length = 11)
3   p = plot(x, abs, axis = :equal, label = "|x|")
4   for λ in range(0, stop = 1, length = 11)
5     plot!([0, λ/2 - (1 - λ)/2], [0, -1/2], arrow = Plots.arrow(:closed), label
6           = "")
7   end
8 end
```

⚠ Skipped xaxis arg equal

⚠ Skipped yaxis arg equal

⚠ Skipped zaxis arg equal

L1_loss (generic function with 1 method)

```
1 L1_loss(w, X, y) = sum(abs.(X * w - y)) / length(y)
```

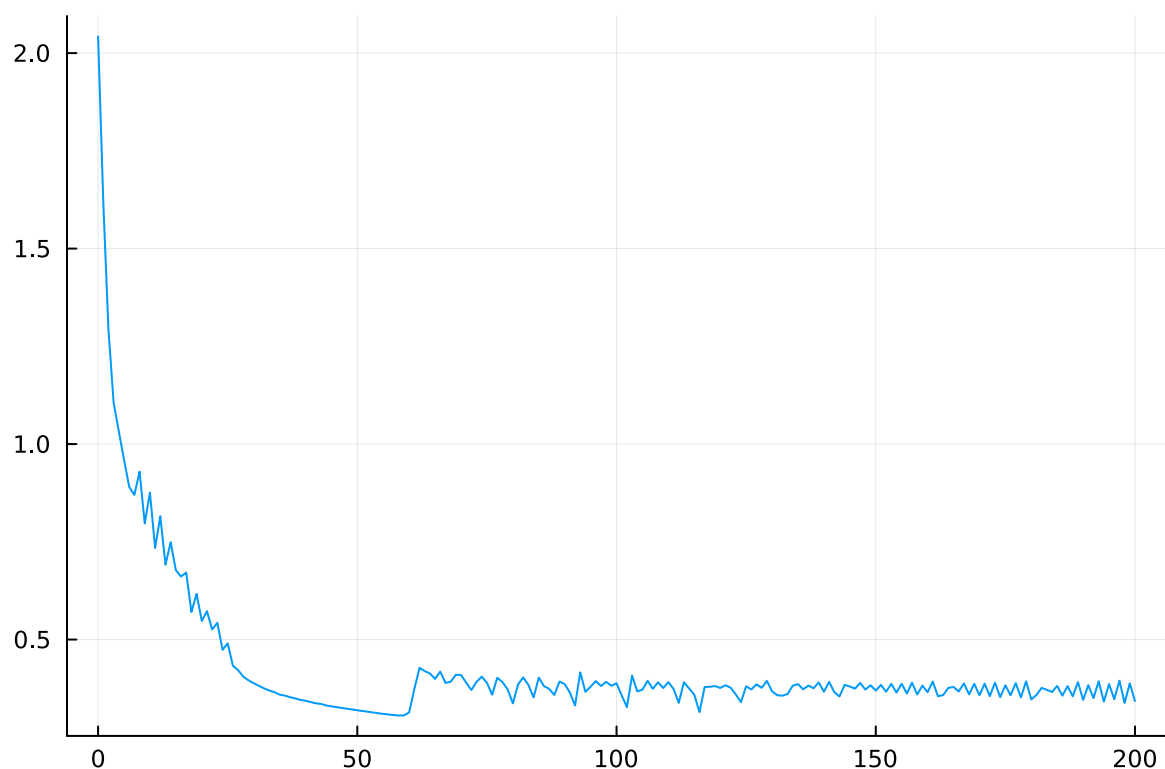
```
1 function Base.abs(d::Dual)
2     if d.value < 0
3         return -1.0 * d
4     else
5         return d
6     end
7 end
```

η_{L1} = 0.1

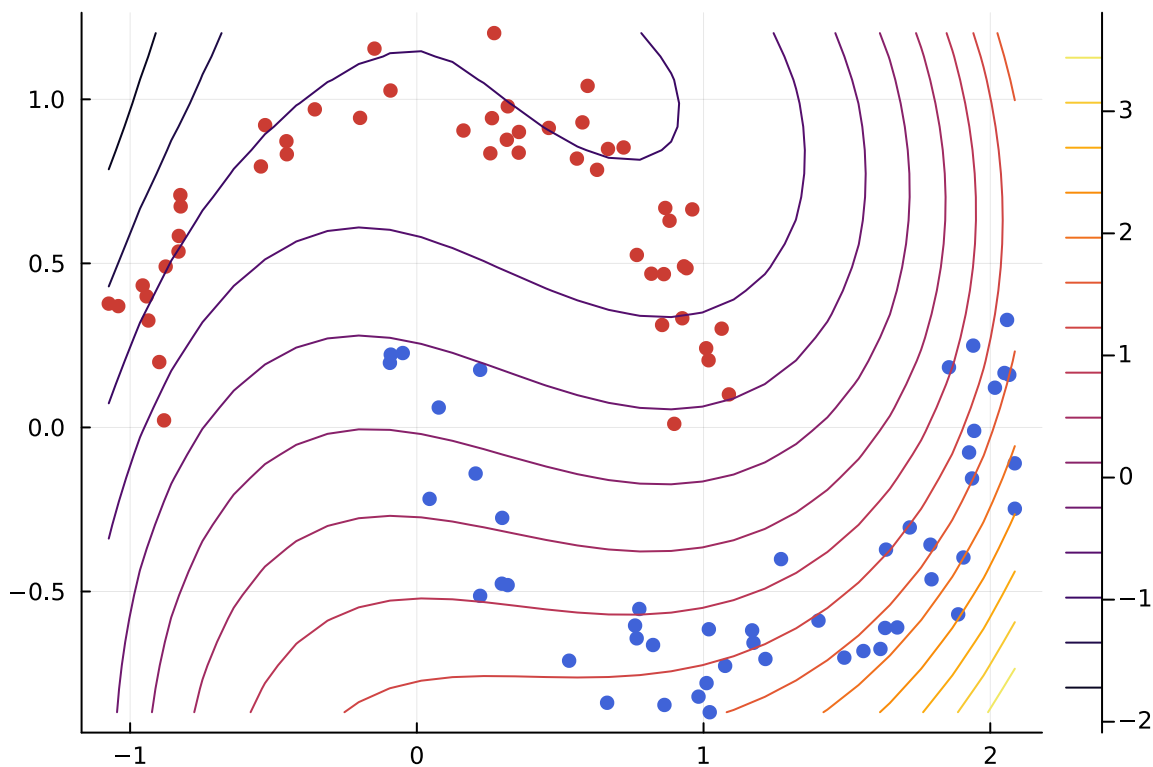
num_iters_L1 = 200

► ([0.0969498, -0.213861, -1.39789, -0.533729, -0.275933, 0.208445, 0.501265, 0.117361, 0.361861, 0.0969498])

```
1 w_trained_L1, training_losses_L1 = train!(forward_diff, L1_loss, w_lift, X_lift, y,
   $\eta_{L1}$ , num_iters_L1)
```



```
1 plot(0:num_iters_L1, training_losses_L1, label = "")
```



```
1 plot_w(w_trained_L1)
```

Reverse diff ⇄

Le désavantage de la forward differentiation, c'est qu'il faut recommencer tout le calcul pour calculer la dérivée par rapport à chaque variable. La *reverse differentiation*, aussi appelée *backpropagation*, résout ce problème en calculant la dérivée par rapport à toutes les variables en une fois !

Chain rule ⇄

Exemple univarié ⇄

Commençons par un exemple univarié pour introduire le fait qu'il existe un choix dans l'ordre de la multiplication des dérivées. La liberté introduite par ce choix donne lieu à la différence entre la différentiation *forward* et *reverse*.

Supposons qu'on veuille dériver la fonction $\tan(\cos(\sin(x)))$ pour $x = \pi/3$. La Chain Rule nous donne :

$$\begin{aligned}
(\tan(\cos(\sin(x))))' &= (\tan(x))'|_{x=\cos(\sin(x))} (\cos(\sin(x)))' \\
&= (\tan(x))'|_{x=\cos(\sin(x))} (\cos(x))'|_{x=\sin(x)} (\sin(x))' \\
&= \frac{1}{\cos^2(\cos(\sin(x)))} (-\sin(\sin(x))) \cos(x)
\end{aligned}$$

La dérivée pour $x = \pi/3$ est donc :

$$\begin{aligned}
(\tan(\cos(\sin(x))))'|_{x=\pi/3} &= (\tan(x))'|_{x=\cos(\sin(\pi/3))} (\cos(x))'|_{x=\sin(\pi/3)} (\sin(x))'|_{x=\pi/3} \\
&= \frac{1}{\cos^2(\cos(\sin(\pi/3)))} (-\sin(\sin(\pi/3))) \cos(\pi/3)
\end{aligned}$$

Pour calculer ce produit de 3 nombres, on a 2 choix.

La première possibilité (qui correspond à forward diff) est de commencer par calculer le produit

$$\begin{aligned}
(\cos(\sin(x)))'|_{x=\pi/3} &= (\cos(x))'|_{x=\sin(\pi/3)} (\sin(x))'|_{x=\pi/3} \\
&= (-\sin(\sin(\pi/3))) \cos(\pi/3)
\end{aligned}$$

puis de le multiplier avec $(\tan(x))'|_{x=\cos(\sin(\pi/3))} = \frac{1}{\cos^2(\cos(\sin(\pi/3)))}$.

La deuxième possibilité (qui correspond à reverse diff) est de commencer par calculer le produit

$$\begin{aligned}
(\tan(\cos(x)))'|_{x=\sin(\pi/3)} &= (\tan(x))'|_{x=\cos(\sin(\pi/3))} (\cos(x))'|_{x=\sin(\pi/3)} \\
&= \frac{1}{\cos^2(\cos(\sin(\pi/3)))} (-\sin(\sin(\pi/3)))
\end{aligned}$$

puis de le multiplier avec $\cos(\pi/3)$.

Vous remarquerez que dans l'équation ci-dessus, comme mis en évidence en rouge, les valeurs auxquelles les dérivées doivent être évaluées dépendent de $\sin(\pi/3)$. L'approche utilisée par reverse diff de multiplier de gauche à droite ne peut donc pas être effectuée sans prendre en compte la valeur qui doit être évaluée de droite à gauche.

Pour appliquer reverse diff, il faut donc commencer par une *forward pass* de droite à gauche qui calcule $\sin(\pi/3)$ puis $\cos(\sin(\pi/3))$ puis $\tan(\cos(\sin(\pi/3)))$. On peut ensuite faire la *backward pass* qui multiplie les dérivées de gauche à droite. Afin d'être disponibles pour la backward pass, les valeurs calculées lors de la forward pass doivent être **stockées** ce qui implique un **coût mémoire**. En revanche, comme forward diff calcule la dérivée dans le même sens que l'évaluation, les dérivées et évaluations peuvent être calculées en même temps afin de ne pas avoir besoin de stocker les évaluations. C'est effectivement ce qu'on a implémenté avec Dual précédemment.

Au vu de ce coût mémoire supplémentaire de reverse diff par rapport à forward diff, ce dernier paraît préférable en pratique. On va voir maintenant que dans le cas multivarié, dans certains cas, ce désavantage est contrebalancé par une meilleure complexité temporelle qui rend reverse diff indispensable !

Exemple multivarié ⇔

Prenons maintenant un exemple multivarié, supposons qu'on veuille calculer le gradient de la fonction $f(g(h(x_1, x_2)))$ qui compose 3 fonctions f , g et h . Le gradient est obtenu via la chain rule comme suit :

$$\begin{aligned}\frac{\partial}{\partial x_1} f(g(h(x_1, x_2))) &= \frac{\partial f}{\partial g} \frac{\partial g}{\partial h} \frac{\partial h}{\partial x_1} \\ \frac{\partial}{\partial x_2} f(g(h(x_1, x_2))) &= \frac{\partial f}{\partial g} \frac{\partial g}{\partial h} \frac{\partial h}{\partial x_2} \\ \nabla_{x_1, x_2} f(g(h(x_1, x_2))) &= \begin{bmatrix} \frac{\partial f}{\partial g} \frac{\partial g}{\partial h} \frac{\partial h}{\partial x_1} & \frac{\partial f}{\partial g} \frac{\partial g}{\partial h} \frac{\partial h}{\partial x_2} \end{bmatrix} \\ &= \begin{bmatrix} \frac{\partial f}{\partial g} \end{bmatrix} \begin{bmatrix} \frac{\partial g}{\partial h} \end{bmatrix} \begin{bmatrix} \frac{\partial h}{\partial x_1} & \frac{\partial h}{\partial x_2} \end{bmatrix}\end{aligned}$$

On voit que c'est le produit de 3 matrices. Forward diff va exécuter ce produit de droite à gauche :

$$\begin{aligned}\nabla_{x_1, x_2} f(g(h(x_1, x_2))) &= \begin{bmatrix} \frac{\partial f}{\partial g} \end{bmatrix} \begin{bmatrix} \frac{\partial g}{\partial h} \frac{\partial h}{\partial x_1} & \frac{\partial g}{\partial h} \frac{\partial h}{\partial x_2} \end{bmatrix} \\ &= \begin{bmatrix} \frac{\partial f}{\partial g} \end{bmatrix} \begin{bmatrix} \frac{\partial g}{\partial x_1} & \frac{\partial g}{\partial x_2} \end{bmatrix} \\ &= \begin{bmatrix} \frac{\partial f}{\partial g} \frac{\partial g}{\partial x_1} & \frac{\partial f}{\partial g} \frac{\partial g}{\partial x_2} \end{bmatrix} \\ &= \begin{bmatrix} \frac{\partial f}{\partial x_1} & \frac{\partial f}{\partial x_2} \end{bmatrix}\end{aligned}$$

L'idée de reverse diff c'est d'effectuer le produit de gauche à droite:

$$\begin{aligned}\nabla_{x_1, x_2} f(g(h(x_1, x_2))) &= \begin{bmatrix} \frac{\partial f}{\partial g} \frac{\partial g}{\partial h} \end{bmatrix} \begin{bmatrix} \frac{\partial h}{\partial x_1} & \frac{\partial h}{\partial x_2} \end{bmatrix} \\ &= \begin{bmatrix} \frac{\partial f}{\partial h} \end{bmatrix} \begin{bmatrix} \frac{\partial h}{\partial x_1} & \frac{\partial h}{\partial x_2} \end{bmatrix} \\ &= \begin{bmatrix} \frac{\partial f}{\partial x_1} & \frac{\partial f}{\partial x_2} \end{bmatrix}\end{aligned}$$

Pour calculer $\partial f / \partial x_1$ via forward diff, on part donc de $\partial x_1 / \partial x_1 = 1$ et $\partial x_2 / \partial x_1 = 0$ et on calcule ensuite $\partial h / \partial x_1$, $\partial g / \partial x_1$ puis $\partial f / \partial x_1$. Effectuer la reverse diff est un peu moins intuitif. L'idée est de partir de la dérivée du résultat par rapport à lui-même $\partial f / \partial f = 1$ et de calculer $\partial f / \partial g$ puis $\partial f / \partial h$ et ensuite $\partial f / \partial x_1$. L'avantage de reverse diff c'est qu'il n'y a que la dernière

étape qui est spécifique à x_1 . Tout jusqu'au calcul de $\partial f / \partial h$ peut être réutilisé pour calculer $\partial f / \partial x_2$, il n'y a plus qu'à multiplier ! Reverse diff est donc plus efficace pour calculer le gradient d'une fonction qui a une seule output par rapport à beaucoup de paramètres, comme détaillé dans la discussion à la fin de ce notebook.

Forward pass : Construction de l'expression graph

Pour implémenter reverse diff, il faut construire l'expression graph pour garder en mémoire les valeurs des différentes expressions intermédiaires afin de pouvoir calculer les dérivées locales $\partial f / \partial g$ et $\partial g / \partial h$. Le code suivant définit un noeud de l'expression graph. Le field derivative correspond à la valeur de $\partial f_{\text{final}} / \partial f_{\text{node}}$ où f_{final} est la dernière fonction de la composition et f_{node} est la fonction correspondant au node.

Node

```
1 begin
2     mutable struct Node
3         op::Union{Nothing,Symbol}
4         args::Vector{Node}
5         value::Float64
6         derivative::Float64
7     end
8     Node(op, args, value) = Node(op, args, value, NaN)
9     Node(value) = Node(nothing, Node[], value)
10 end
```

L'opérateur overloading suivant sera suffisant pour construire l'expression graph dans le cadre de ce notebook, vous l'étendrez pendant la séance d'exercice.

```
1 begin
2     Base.zero(x::Node) = Node(0.0)
3     Base.*(x::Node, y::Node) = Node(:*, [x, y], x.value * y.value)
4     Base.+(x::Node, y::Node) = Node(:+, [x, y], x.value + y.value)
5     Base.-(x::Node, y::Node) = Node(:-, [x, y], x.value - y.value)
6     Base./(x::Node, y::Number) = x * Node(inv(y))
7     Base.^(x::Node, n::Integer) = Base.power_by_squaring(x, n)
8     Base.sin(x::Node) = Node(:sin, [x], sin(x.value))
9     Base.cos(x::Node) = Node(:cos, [x], cos(x.value))
10 end
```

On crée les leafs correspondant aux variables x_1 et x_2 de valeurs 1 et 2 respectivement. Les valeurs correspondent aux valeurs de x_1 et x_2 auxquelles on veut dériver la fonction. On a choisi 1 et 2 pour pouvoir les reconnaître facilement dans le graphe.

```
x_nodes = ▶[Node(nothing, [], 1.0, NaN), Node(nothing, [], 2.0, NaN)]
```

```
1 x_nodes = Node.([1, 2])
```

```
expr = ▶Node(:cos, [Node(:sin, [... more], 0.909297, NaN)], 0.6143002821164822, NaN)
```

```
1 expr = cos(sin(prod(x_nodes)))
```

_nodes! (generic function with 1 method)

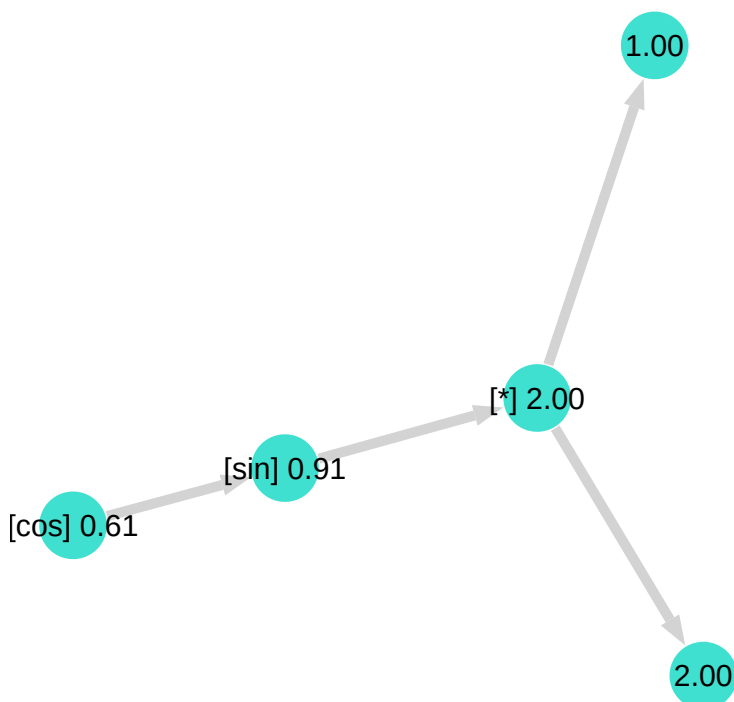
_edges (generic function with 1 method)

graph (generic function with 1 method)

Pour le visualiser, on le convertit en graphe en utilisant la structure de données de Graphs.jl pour pouvoir utiliser gplot.

```
▶({5, 4} directed simple Int64 graph, ["[cos] 0.61", "[sin] 0.91", "[*] 2.00", "1.00", "2.00"])
```

```
1 expr_graph, labels = graph(expr)
```



```
1 gplot(expr_graph, node_label = labels)
```

Combinaison des dérivées ⇔

Que faire si plusieurs expressions dépendent d'une même variable ? Considérons l'exemple

$f(x) = \sin(x) \cos(x)$ qui correspond à $f(g, h) = gh$, $g(x) = \sin(x)$ et $h(x) = \cos(x)$. La chain rule donne

$$f'(x) = \frac{\partial f}{\partial g} g'(x) + \frac{\partial f}{\partial h} h'(x)$$

Une fois la valeur $\partial f / \partial g$ calculée, on peut la multiplier par $g'(x)$ pour avoir la première partie de $f'(x)$. Idem pour h . Ces deux contributions seront calculée séparément lors de la backward pass sur le noeud g et h . On voit par la formule de la chain rule que ces deux contributions doivent être sommées. Lors de la backward pass, on initialise donc toutes les dérivées à 0. Pour chaque contribution, on ajoute la dérivée avec +=. On s'assure ensuite qu'on ne procède pas à la backward pass sur un noeud avant qu'il ait fini d'accumuler les contributions de toutes les expressions qui en dépendent via un *tri topologique*.

```
x_node = ▶ Node(nothing, [], 1.0, NaN)
```

```
1 x_node = Node(1)
```

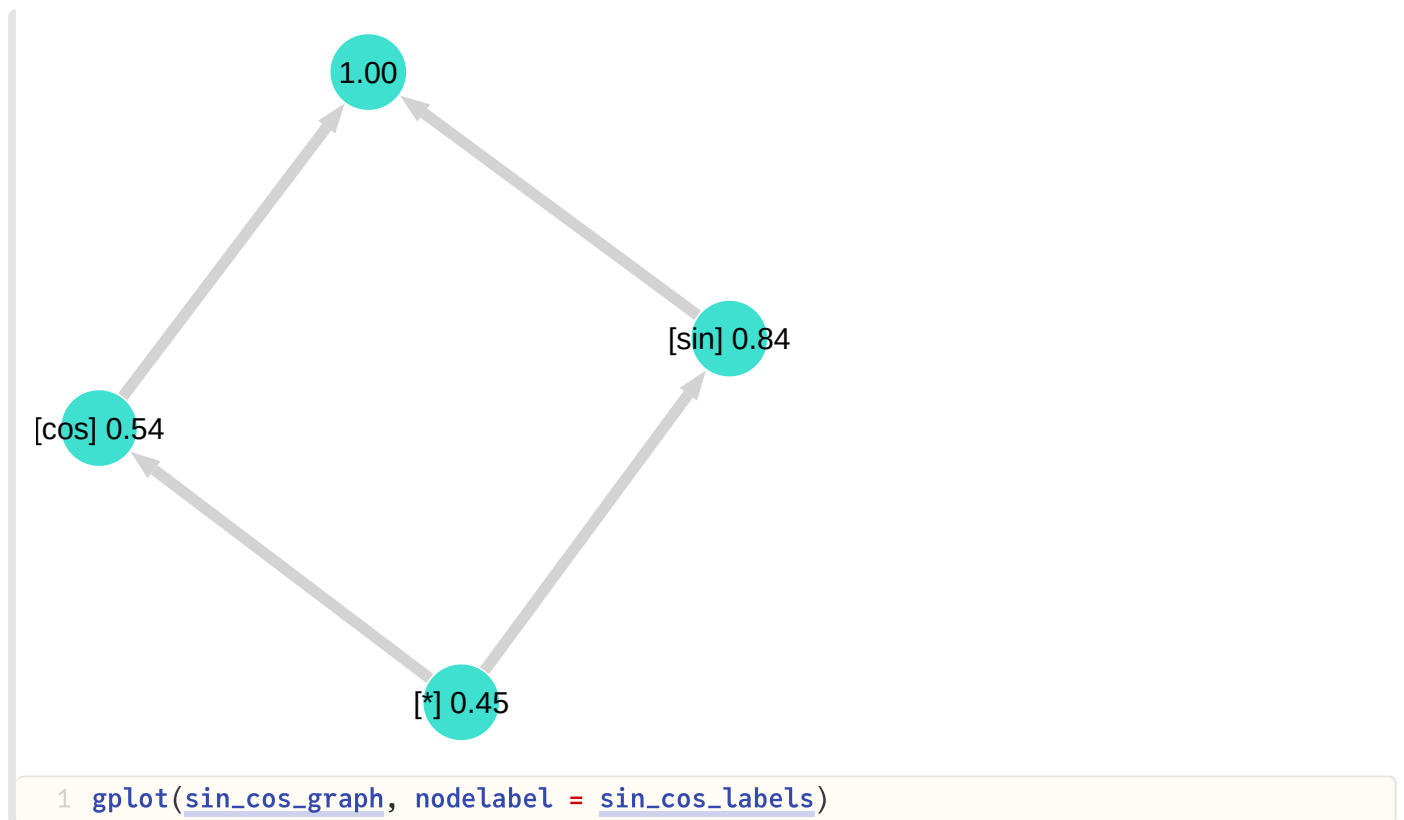
```
sin_cos =
```

```
▶ Node(:,*, [Node(:,sin, [ ... more], 0.841471, NaN), Node(:,cos, [ ... more], 0.540302, NaN)]), 0.454
```

```
1 sin_cos = sin(x_node) * cos(x_node)
```

```
▶ ({4, 4} directed simple Int64 graph, ["[*] 0.45", "[sin] 0.84", "1.00", "[cos] 0.54"])
```

```
1 sin_cos_graph, sin_cos_labels = graph(sin_cos)
```



Backward pass : Calcul des dérivées ⇔

La fonction suivante propage la dérivée $\partial f_{\text{final}} / \partial f_{\text{node}}$ à la dérivée des arguments de la fonction f_{node} . Comme les arguments peuvent être utilisés par d'autres fonction, on somme la dérivée avec += .

_backward! (generic function with 1 method)

```
1 function _backward!(f::Node)
2     if isnothing(f.op)
3         return
4     elseif f.op == :+
5         for arg in f.args
6             arg.derivative += f.derivative
7         end
8     elseif f.op == :- && length(f.args) == 2
9         f.args[1].derivative += f.derivative
10        f.args[2].derivative -= f.derivative
11    elseif f.op == :* && length(f.args) == 2
12        f.args[1].derivative += f.derivative * f.args[2].value
13        f.args[2].derivative += f.derivative * f.args[1].value
14    elseif f.op == :sin
15        f.args[1].derivative += f.derivative * cos(f.args[1].value)
16    elseif f.op == :cos
17        f.args[1].derivative -= f.derivative * sin(f.args[1].value)
18    else
19        error("Operator `$(f.op)` not supported yet")
20    end
21 end
```

La fonction `_backward!` ne doit être appelée que sur un noeud pour lequel `f.derivative` a déjà été calculé. Pour cela, `_backward!` doit avoir été appelé sur tous les noeuds qui représentent des fonctions qui dépendent directement ou indirectement du résultat du noeud. Pour trouver l'ordre dans lequel appeler `_backward!`, on utilise donc un tri topologique (nous reviendrons sur les tris topologiques dans la partie graphe).

backward! (generic function with 1 method)

```
1 function backward!(f::Node)
2     topo = typeof(f)[]
3     topo_sort!(Set{typeof(f)}(), topo, f)
4     reverse!(topo)
5     for node in topo
6         node.derivative = 0
7     end
8     f.derivative = 1
9     for node in topo
10        _backward!(node)
11    end
12    return f
13 end
```

topo_sort! (generic function with 1 method)

```
1 function topo_sort!(visited, topo, f::Node)
2     if !(f in visited)
3         push!(visited, f)
4         for arg in f.args
5             topo_sort!(visited, topo, arg)
6         end
7         push!(topo, f)
8     end
9 end
```

► Node(:cos, [Node(:sin, [... more], 0.909297, -0.789072)], 0.6143002821164822, 1.0)

```
1 backward!(expr)
```

On a maintenant l'information sur les dérivées de x_nodes :

► [Node(nothing, [], 1.0, 0.65674), Node(nothing, [], 2.0, 0.32837)]

```
1 x_nodes
```

Comparaison avec Forward Diff dans l'exemple moon ↔

Revenons sur l'exemple utilisé pour illustrer la forward diff et essayons de calculer la même dérivée mais à présent en utilisant reverse diff.

w_nodes = ► [Node(nothing, [], 0.366916, NaN), Node(nothing, [], 0.813359, NaN)]

```
1 w_nodes = Node.(w)
```

mse_expr = ► Node(:*, [Node(:*, [... more], 1.0, 1.0), Node(nothing, [], 1.0, 1.0)], 1.0, 1.0)

```
1 mse_expr = backward!(mse(Node.(ones(1)), Node.(2ones(1, 1)), Node.(3ones(1))))
```

reverse_diff (generic function with 1 method)

```
1 function reverse_diff(loss, w, X, y)
2     w_nodes = Node.(w)
3     expr = loss(w_nodes, Node.(X), Node.(y))
4     backward!(expr)
5     return [w.derivative for w in w_nodes]
6 end
```

► [-0.413979, 1.50562]

```
1 @time forward_diff(mse, w, X, y)
```



0.000019 seconds (25 allocations: 13.359 KiB)



```
► [-0.413979, 1.50562]
```

```
1 @time reverse_diff(mse, w, X, y)
```

```
0.105477 seconds (198.98 k allocations: 10.394 MiB, 99.85% compilation time)
```

On l'exécute une seconde fois, pour éliminer le temps de compilation :

```
► [-0.413979, 1.50562]
```

```
1 @time reverse_diff(mse, w, X, y)
```

```
0.000160 seconds (4.14 k allocations: 245.500 KiB)
```

On remarque que reverse diff est plus lent ! Il y a un certain coût mémoire lorsqu'on construit l'expression graph. Pour cette raison, si on veut calculer plusieurs dérivées consécutives pour différentes valeurs de x_1 et x_2 , on a intérêt à garder le graphe et à uniquement changer la valeur des variables plutôt qu'à reconstruire le graphe à chaque fois qu'on change les valeurs. Alternativement, on peut essayer de condenser le graphe en exprimant les opérations sur des larges matrices ou même tenseurs, c'est l'approche utilisée par pytorch ou tensorflow.

num_data = 32

num_features = 32

num_hidden = 32

mse2 (generic function with 1 method)

```
1 mse2(W1, W2, X, y) = sum((X * W1 * W2 - y).^2 / length(y))
```

bench (generic function with 1 method)

```
1 function bench(num_data, num_features, num_hidden)
2     X = rand(num_data, num_features)
3     W1 = rand(num_features, num_hidden)
4     W2 = rand(num_hidden)
5     y = rand(num_data)
6
7     @time for i in axes(W1, 1)
8         for j in axes(W1, 2)
9             mse2(
10                 Dual.(W1, onehot(i, axes(W1, 1)) * onehot(j, axes(W1, 2)))',
11                 W2,
12                 X,
13                 y,
14             )
15         end
16     end
17     expr = @time mse2(Node.(W1), Node.(W2), Node.(X), Node.(y))
18     @time backward!(expr)
19     return
20 end
```

```
1 bench(num_data, num_features, num_hidden)
```



```
0.008739 seconds (16.38 k allocations: 19.133 MiB)
0.000176 seconds (17.55 k allocations: 737.734 KiB)
0.000661 seconds (37 allocations: 314.578 KiB)
```



Comment choisir entre forward et reverse diff ? ⇔

Suppose that we need to differentiate a composition of functions: $(f_n \circ f_{n-1} \circ \dots \circ f_2 \circ f_1)(w)$. For each function, we can compute a jacobian given the value of its input. So, during a forward pass, we can compute all jacobians. We now just need to take the product of these jacobians:

$$J_n J_{n-1} \dots J_2 J_1$$

While the product of matrices is associative, its computational complexity depends on the order of the multiplications! Let $d_i \times d_{i-1}$ be the dimension of J_i .

Forward diff: from right to left ⇔

If the product is computed from right to left:

$$\begin{array}{ll}
J_{1,2} = J_2 J_1 & \Omega(d_2 d_1 d_0) \\
J_{1,3} = J_3 J_{1,2} & \Omega(d_3 d_2 d_0) \\
J_{1,4} = J_4 J_{1,3} & \Omega(d_4 d_3 d_0) \\
\vdots & \\
J_{1,n} = J_n J_{1,(n-1)} & \Omega(d_n d_{n-1} d_0)
\end{array}$$

we have a complexity of

$$\Omega\left(\sum_{i=2}^n d_i d_{i-1} d_0\right).$$

Reverse diff: from left to right $\Leftrightarrow$

Reverse differentiation corresponds to multiplying the adjoint from right to left or equivalently the original matrices from left to right. This means computing the product in the following order:

$$\begin{array}{ll}
J_{(n-1),n} = J_n J_{n-1} & \Omega(d_n d_{n-1} d_{n-2}) \\
J_{(n-2),n} = J_{(n-1),n} J_{n-2} & \Omega(d_n d_{n-2} d_{n-3}) \\
J_{(n-3),n} = J_{(n-2),n} J_{n-3} & \Omega(d_n d_{n-3} d_{n-4}) \\
\vdots & \\
J_{1,n} = J_{2,n} J_1 & \Omega(d_n d_1 d_0)
\end{array}$$

We have a complexity of

$$\Omega\left(\sum_{i=1}^{n-1} d_n d_i d_{i-1}\right).$$

Mixed : from inward to outward $\Leftrightarrow$

Suppose we multiply starting from some d_k where $1 < k < n$. We would then first compute the left side:

$$\begin{array}{ll}
J_{k+1,k+2} = J_{k+2} J_{k+1} & \Omega(d_{k+2} d_{k+1} d_k) \\
J_{k+1,k+3} = J_{k+3} J_{k+1,k+2} & \Omega(d_{k+3} d_{k+2} d_k) \\
\vdots & \\
J_{k+1,n} = J_n J_{k+1,n-1} & \Omega(d_n d_{n-1} d_k)
\end{array}$$

then the right side:

$$\begin{array}{ll}
J_{k-1,k} = J_k J_{k-1} & \Omega(d_k d_{k-1} d_{k-2}) \\
J_{k-2,k} = J_{k-1,k} J_{k-2} & \Omega(d_k d_{k-2} d_{k-3}) \\
\vdots & \\
J_{1,k} = J_{2,k} J_1 & \Omega(d_k d_1 d_0)
\end{array}$$

and then combine both sides:

$$J_{1,n} = J_{k+1,n} J_{1,k} \quad \Omega(d_n d_k d_0)$$

we have a complexity of

$$\Omega(d_n d_k d_0 + \sum_{i=1}^{k-1} d_k d_i d_{i-1} + \sum_{i=k+2}^n d_i d_{i-1} d_k).$$

Comparison ⇔

We see that we should find the minimum d_k and start from there. If the minimum is attained at $k = n$, this corresponds multiplying from left to right, this is reverse differentiation. If the minimum is attained at $k = 0$, we should multiply from right to left, this is forward differentiation. Otherwise, we should start from the middle, this would mean mixing both forward and reverse diff.

What about neural networks ? In that case, d_0 is equal to the number of entries in W_1 added with the number of entries in W_2 while d_n is 1 since the loss is scalar. We should therefore clearly multiply from left to right hence do reverse diff.

Pour la séance d'exercices: ⇔

tanh ⇔

Passer un réseau de neurone avec fonction d'activation tanh

```
W1 = 2x32 Matrix{Float64}:
 0.335894  0.652456  0.865473  0.475114  ...  0.699115  0.34574  0.339306  0.546143
 0.484145  0.394062  0.978982  0.626543  ...  0.233263  0.529536  0.211592  0.12284
```

```
1 W1 = rand(size(X, 2), num_hidden)
```

```
W2 =
▶ [0.799448, 0.0244941, 0.829904, 0.0229698, 0.826625, 0.0779315, 0.879647, 0.223648, 0.749
```

```
1 W2 = rand(num_hidden)
```

```
y_est_tanh =
```

```
► [5.18026, 3.55378, 11.1326, -4.3213, -6.605, 8.17954, 13.682, 8.96737, 13.2021, 11.1509, 1
```

```
1 y_est_tanh = tanh.(X * W1) * W2
```

ReLU ⇔

Passer un réseau de neurone avec fonction d'activation ReLU

```
relu (generic function with 1 method)
```

```
1 function relu(x)
2     if x < 0
3         return 0
4     else
5         return x
6     end
7 end
```

```
y_est_relu =
```

```
► [5.88146, 3.62417, 15.7205, 0.385251, 0.0200185, 9.35975, 21.8158, 11.8086, 19.9146, 14.05
```

```
1 y_est_relu = relu.(X * W1) * W2
```

One-Hot encoding et cross-entropy ⇔

```
Y = 100×2 BitMatrix:
```

```
1 0
0 1
0 1
1 0
1 0
1 0
0 1
⋮
0 1
1 0
0 1
0 1
0 1
1 0
```

```
1 Y = unique!(sort(y_cat.refs))' .== y_cat.refs
```

```
W = 2×2 Matrix{Float64}:
```

```
0.336865 0.150351
0.782704 0.407011
```

```
1 W = rand(size(X, 2), size(Y, 2))
```

```
Y_1 = 100×2 Matrix{Float64}:
 0.671734  0.354221
 0.211814  0.104663
 0.331977  0.125312
 0.0886631 0.0678478
-0.146665 -0.0539678
 0.740263  0.378579
 0.821598  0.375948
 ⋮
-0.155427 -0.0819257
 0.276396  0.1642
-0.371052 -0.214385
 0.122423  0.0659918
-0.216165 -0.131301
 0.772699  0.40409
```

```
1 Y_1 = X * W
```

```
Y_2 = 100×2 Matrix{Float64}:
 1.95763  1.42507
 1.23592  1.11034
 1.39372  1.1335
 1.09271  1.0702
 0.863583 0.947463
 2.09649  1.46021
 2.27413  1.45637
 ⋮
 0.856049 0.92134
 1.31837  1.17845
 0.690008 0.807038
 1.13023  1.06822
 0.805602 0.876954
 2.1656  1.49794
```

```
1 Y_2 = exp.(X * W)
```

```
sums = 100×1 Matrix{Float64}:
 3.38269889172229
 2.3462544691723632
 2.5272228684428626
 2.162914860959751
 1.8110458933638958
 3.556695814420749
 3.7305026007341096
 ⋮
 1.7773896516049998
 2.4968197375395933
 1.4970458074310753
 2.198450239265491
 1.682556199293044
 3.663542062194829
```

```
1 sums = sum(Y_2, dims=2)
```

```
Y_est = 100×2 Matrix{Float64}:
 0.578718  0.421282
 0.526762  0.473238
 0.551483  0.448517
 0.505204  0.494796
 0.476842  0.523158
 0.589448  0.410552
 0.609604  0.390396
 ⋮
 0.481633  0.518367
 0.52802   0.47198
 0.460913  0.539087
 0.514104  0.485896
 0.478797  0.521203
 0.591123  0.408877
```

```
1 Y_est = Y_2 ./ sums
```

```
cross = 100×2 Matrix{Float64}:
 0.578718  0.0
 0.0        0.473238
 0.0        0.448517
 0.505204  0.0
 0.476842  0.0
 0.589448  0.0
 0.0        0.390396
 ⋮
 0.0        0.518367
 0.52802   0.0
 0.0        0.539087
 0.0        0.485896
 0.0        0.521203
 0.591123  0.0
```

```
1 cross = Y_est .* Y
```

```
cross_entropies = 100×1 Matrix{Float64}:
 0.5469397931445459
 0.7481572667199263
 0.8018090042096699
 0.6827936846369236
 0.7405694483928652
 0.5285686597688829
 0.9405945229540451
 ⋮
 0.6570714821598023
 0.6386218229034386
 0.6178785894776497
 0.7217608882547931
 0.6516150074480459
 0.5257315261790176
```

```
1 cross_entropies = -log.(sum(Y_est .* Y, dims=2))
```

```
► [0.54694, 0.748157, 0.801809, 0.682794, 0.740569, 0.528569, 0.940595, 0.735758, 0.922027,
```

```
1 -log.(getindex.(Ref(Y_est), axes(Y_est, 1), y_cat.refs))
```

Acknowledgements and further readings

- Dual est inspiré de [ForwardDiff](#)
- Node est inspiré de [micrograd](#)
- Une bonne intro à l'automatic differentiation est disponible [ici](#)

Utils

```
1 import MLJBase, Colors, Tables
```

```
Precompiling MLJBase...
 905.6 ms ✓ Accessors → AccessorsStaticArraysExt
 911.0 ms ✓ BangBang → BangBangStaticArraysExt
1020.7 ms ✓ CategoricalArrays → CategoricalArraysRecipesBaseExt
1241.5 ms ✓ HypergeometricFunctions
1556.5 ms ✓ StatsFuns
 687.0 ms ✓ StatsFuns → StatsFunsInverseFunctionsExt
1474.9 ms ✓ StatsFuns → StatsFunsChainRulesCoreExt
6405.1 ms ✓ MLUtils
5018.8 ms ✓ Distributions
1321.4 ms ✓ Distributions → DistributionsChainRulesCoreExt
1402.7 ms ✓ Distributions → DistributionsTestExt
2655.7 ms ✓ ScientificTypes
1972.8 ms ✓ CategoricalDistributions
7013.0 ms ✓ StatisticalMeasuresBase
5811.9 ms ✓ MLJBase
15 dependencies successfully precompiled in 20 seconds. 143 already precompiled.
Precompiling CategoricalArraysJSONExt...
 470.3 ms ✓ CategoricalArrays → CategoricalArraysJSONExt
1 dependency successfully precompiled in 0 seconds. 16 already precompiled.
```

```
1 using Graphs, GraphPlot, Printf
```

```
Precompiling Graphs...
1009.1 ms ✓ ArnoldiMethod
3513.2 ms ✓ Graphs
2 dependencies successfully precompiled in 5 seconds. 28 already precompiled.
Precompiling GraphPlot...
2140.0 ms ✓ Compose
1232.7 ms ✓ GraphPlot
2 dependencies successfully precompiled in 3 seconds. 40 already precompiled.
```

1 using Plots, OneHotArrays, PlutoUI

Precompiling Plots...

```
868.4 ms ✓ Cairo_jll
1365.6 ms ✓ Qt6Base_jll
654.8 ms ✓ HarfBuzz_jll
911.1 ms ✓ Qt6ShaderTools_jll
844.7 ms ✓ libass_jll
919.7 ms ✓ Pango_jll
869.1 ms ✓ libdecor_jll
1317.0 ms ✓ FFMPEG_jll
859.3 ms ✓ GLFW_jll
777.5 ms ✓ FFMPEG
962.2 ms ✓ GR_jll
2879.7 ms ✓ Qt6Declarative_jll
611.2 ms ✓ Qt6Wayland_jll
4202.9 ms ✓ GR
11749.7 ms ✓ PlotUtils
2806.0 ms ✓ PlotThemes
4216.7 ms ✓ RecipesPipeline
63479.9 ms ✓ Plots
2878.3 ms ✓ Plots → UnitfulExt
19 dependencies successfully precompiled in 83 seconds. 165 already precompiled.
```

Precompiling OneHotArrays...

```
510.0 ms ✓ Adapt → AdaptStaticArraysExt
1074.5 ms ✓ LLVMExtra_jll
5756.9 ms ✓ LLVM
1501.6 ms ✓ UnsafeAtomicsLLVM
3455.4 ms ✓ KernelAbstractions
1119.3 ms ✓ KernelAbstractions → LinearAlgebraExt
```