

Automatic differentiation

Jill-Jênn Vie

Oct 3, 2025

Outline

Data $\mathbf{X} \in \mathbb{R}^{n \times d}$ targets $\mathbf{y} \in \mathcal{Y}^n$

Model $f : \mathbf{x} \mapsto \hat{\mathbf{y}} = f(\mathbf{x})$

Activation and loss $\mathcal{L}$

Why sigmoid? Why ReLU?

Optimizers

Why stochastic gradient descent?

Automatic differentiation

Why backpropagation? What does `loss.backward()` do?

Unsupervised, semi-supervised, self-supervised

What if we have very few labels $\mathbf{y}$ or only $\mathbf{X}$?



@untitled01ipynb

Optimization

Find the best parameters θ to reach a goal: making $f_{\theta}(x)$ close to y

Usually, minimize a differentiable loss function

Regression: continuous target

Ex. $y \in \mathbb{R}$ $\mathcal{L} = \sum_i (\hat{y}_i - y_i)^2$ squared error

Classification: discrete target

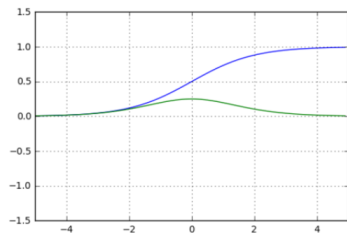
Ex. $y \in \{0, 1\}$

- ▶ Logistic regression $\hat{y} = \text{sigmoid}(\mathbf{w}^T \mathbf{x}) \in \{0, 1\}$
- ▶ Logistic loss: $\mathcal{L}_i(p, y) = y \log p + (1 - y) \log(1 - p)$

Ex. $y \in \{1, \dots, C\}$

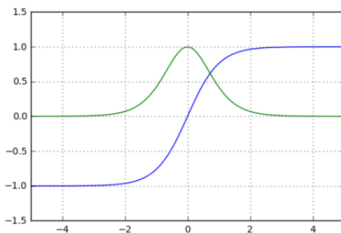
- ▶ Multinomial regression $\hat{y} = \text{softmax}(\mathbf{W}\mathbf{x} + \mathbf{b}) \in \{0, 1\}^C$
- ▶ Cross-entropy loss: $\mathcal{L}_i(\mathbf{p}, y) = \sum_{c=1}^C y_c \log p_c = \log p_y$

Activation / link functions; why sigmoid, softmax?



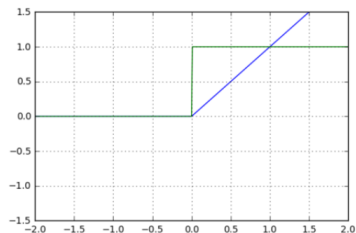
$$\text{sigm}(x) = \frac{1}{1 + e^{-x}}$$

$$\text{sigm}'(x) = \text{sigm}(x)(1 - \text{sigm}(x))$$



$$\tanh(x) = \frac{e^{2x} - 1}{e^{2x} + 1}$$

$$\tanh'(x) = 1 - \tanh(x)^2$$



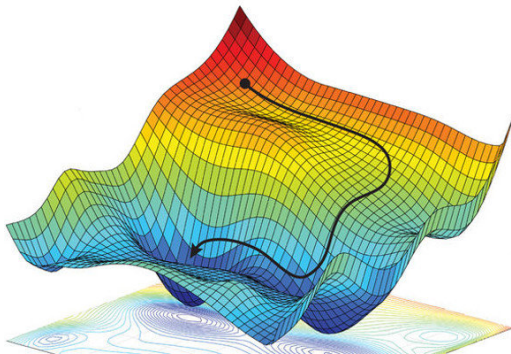
$$\text{relu}(x) = \max(0, x)$$

$$\text{relu}'(x) = 1_{x>0}$$

(Ollion & Grisel)

	Binary	Multiclass
f	softplus : $x \mapsto \log(1 + \exp(x))$	logsumexp ⁺ : $\mathbf{x} \mapsto \log(1 + \sum_c \exp(x_c))$
f'	sigmoid : $x \mapsto 1/(1 + \exp(-x))$	softmax : $\mathbf{x} \mapsto \exp(\mathbf{x}) / \sum_c \exp(x_c)$
f''	$x \mapsto \text{sigmoid}(x)(1 - \text{sigmoid}(x))$	$\mathbf{x} \mapsto \text{diag}(s(\mathbf{x})) - s(\mathbf{x})s(\mathbf{x})^T$

Gradient descent



$$\theta_{t+1} = \theta_t - \gamma \nabla_{\theta} \mathcal{L}$$

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

SGD in PyTorch

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

Variants

Batch gradient descent

Compute the gradient of loss on all examples and update parameters

Gradient descent

For each example $\mathbf{x}_i, y$ update parameters

Stochastic gradient descent

Sample examples $\mathbf{x}_i, y$ and update parameters

Minibatch gradient descent

Sample a batch of examples, estimate full gradient from batch and update parameters

$$\boldsymbol{\theta} = \boldsymbol{\theta} - \gamma \frac{n}{B} \sum_{i \in B} \mathcal{L}(\mathbf{x}_i, y)$$

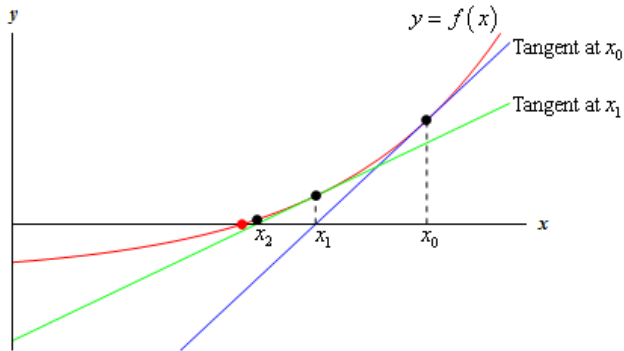
Why SGD?

To minimize a function

Find the zeroes of its derivative

How to find the zeroes of a function?

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


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

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

Quadratic convergence

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

In higher dimension

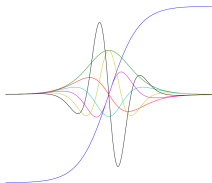
Let $g : \mathbb{R}^n \rightarrow \mathbb{R}$ twice differentiable, with n big

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

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

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

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



Existing methods and their limitations

Numerical differentiation

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

Round-off errors

Symbolic differentiation

Have to keep symbolic expressions at each step of the process

Automatic differentiation

Chain rule

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

Generalized chain rule

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

Reminder: Jacobians

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

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

Example: linear layer

W is a $m \times n$ matrix

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

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

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

Generalized multivariate chain rule

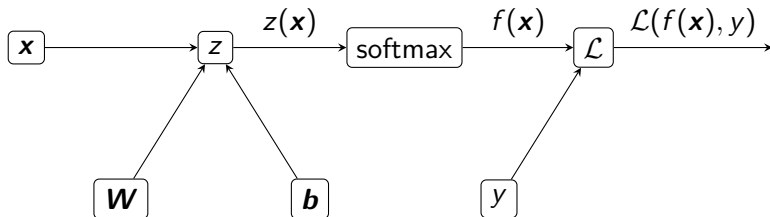
Consider differentiable $f : \mathbb{R}^m \rightarrow \mathbb{R}^k$, $g : \mathbb{R}^n \rightarrow \mathbb{R}^m$ and $\mathbf{a} \in \mathbb{R}^n$.

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

So the Jacobians verify:

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

Computation graph



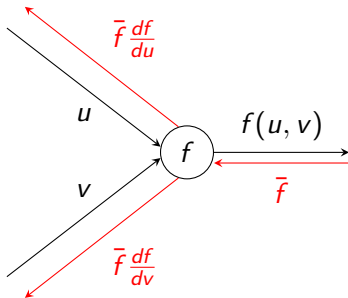
$$\begin{aligned}\frac{d\mathcal{L}}{db_1} &= \frac{d\mathcal{L}}{df} \frac{df}{db_1} = \frac{d\mathcal{L}}{df} \left(\frac{df}{dz} \frac{dz}{db_1} \right) \quad (\text{forward}) \\ &= \frac{d\mathcal{L}}{dz} \frac{dz}{db_1} = \left(\frac{d\mathcal{L}}{df} \frac{df}{dz} \right) \frac{dz}{db_1} \quad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

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

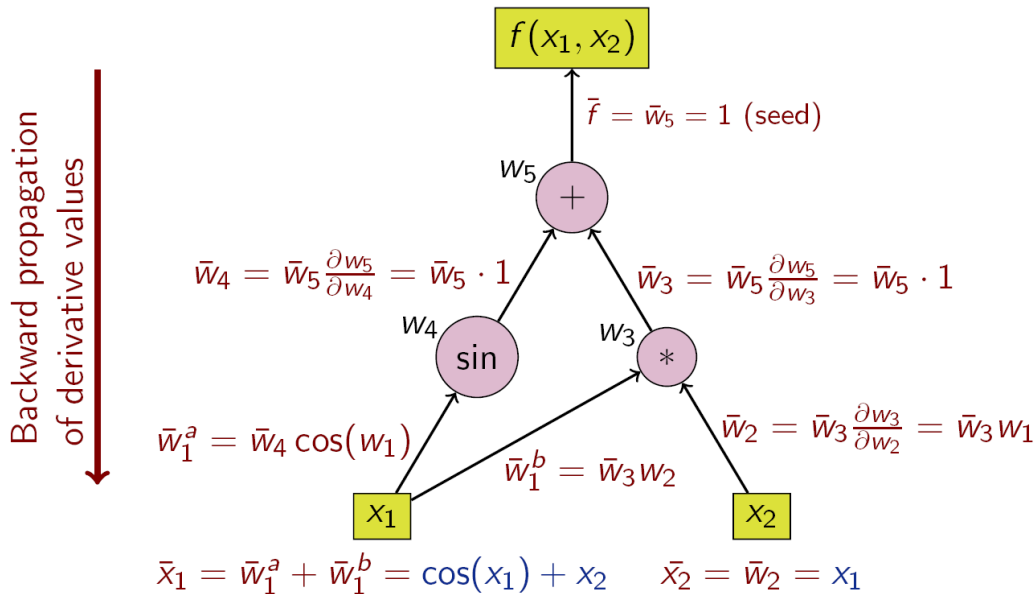
Given that $\mathbf{x} \in \mathbb{R}^d$, $z(\mathbf{x})$, $f(\mathbf{x}) \in \mathbb{R}^{d_2}$, $\mathcal{L}(f(\mathbf{x}), y) \in \mathbb{R}$,
which order is better? This is why **backpropagation**.

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

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

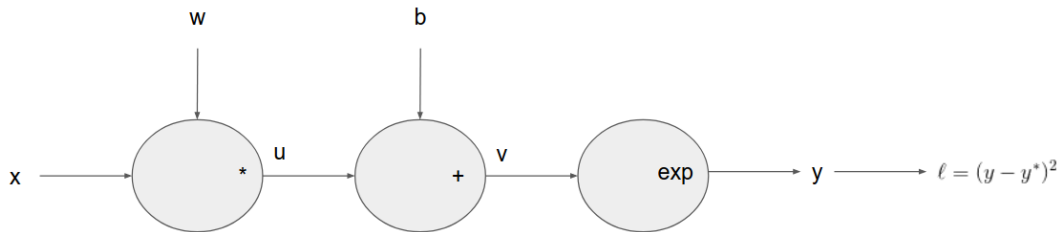


A complete example (Wikipedia)



Notebook: compute gradients

FORWARD PASS



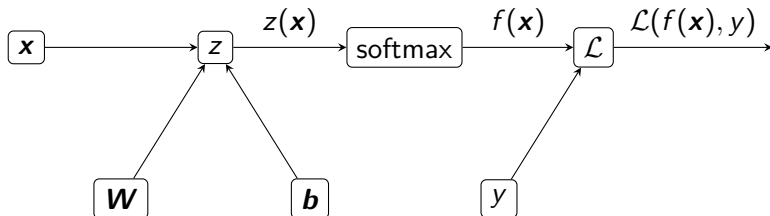
$$u = wx$$

$$v = u + b$$

$$y = e^v$$

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

Computation graph: classifier



Here we define $\mathcal{L}$ as cross-entropy:

$$\mathcal{L}(f(\mathbf{x}), y) = - \sum_{c=1}^K \mathbf{1}_{y=c} \log f(\mathbf{x})_c = - \log f(\mathbf{x})_y$$

Compute $\frac{d\mathcal{L}}{dz_c}$

Unsupervised, semi-supervised, self-supervised

What if we have very few labels $\mathbf{y}$ or only $\mathbf{X}$?

Next week: introduction to reinforcement learning

Dynamic programming's *Principle of Optimality* (Bellman, 1952)

value function $v_{\pi}(s)$
action-value function $q_{\pi}(s, a)$ } for policy $\pi : S \rightarrow A \rightarrow R \rightarrow S' \rightarrow A'$

Q-learning (1992), SARSA (1996), policy gradient (2000)

