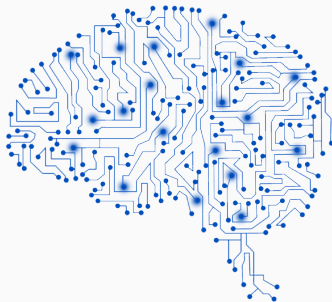


Course V – Introduction to Artificial Neural Networks for image classification II

Bruno Galerne

Thursday January 18, 2024



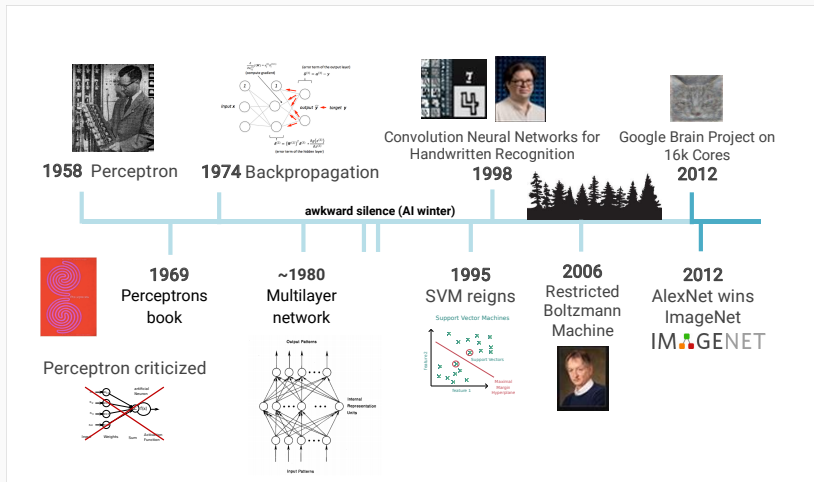
Most of the slides from **Charles Deledalle's** course "UCSD ECE285 Machine learning for image processing" (30 × 50 minutes course)



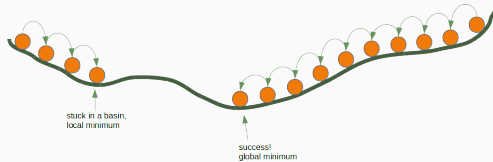
`www.charles-deledalle.fr/`

`https://www.charles-deledalle.fr/pages/teaching.php#learning`

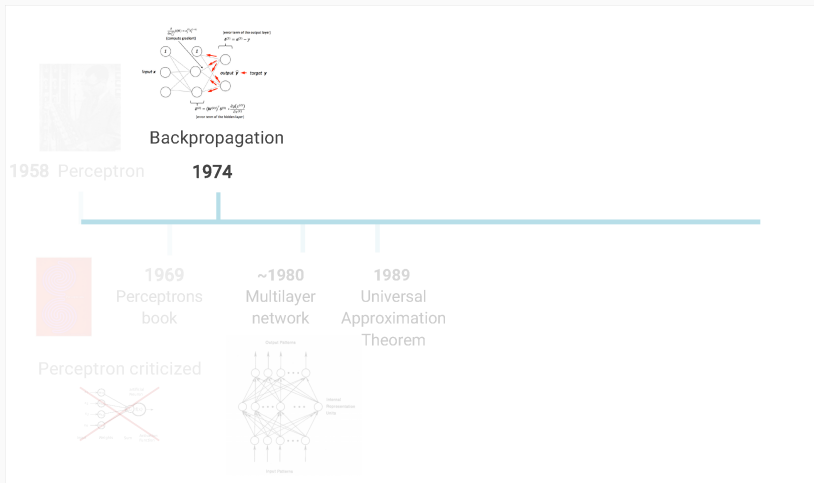
Timeline of (deep) learning



Backpropagation

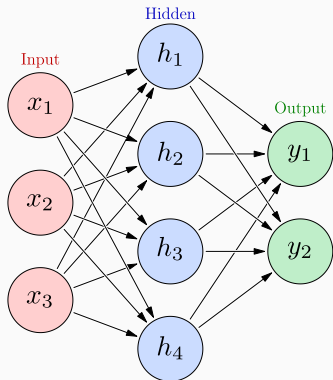


Learning with backpropagation



Feedforward Artificial Neural Network

Artificial neural network / Multilayer perceptron / NeuralNet



$$h_1 = g_1 (w_{11}^1 x_1 + w_{12}^1 x_2 + w_{13}^1 x_3 + b_1^1)$$

$$h_2 = g_1 (w_{21}^1 x_1 + w_{22}^1 x_2 + w_{23}^1 x_3 + b_2^1)$$

$$h_3 = g_1 (w_{31}^1 x_1 + w_{32}^1 x_2 + w_{33}^1 x_3 + b_3^1)$$

$$h_4 = g_1 (w_{41}^1 x_1 + w_{42}^1 x_2 + w_{43}^1 x_3 + b_4^1)$$

$$y_1 = g_2 (w_{11}^2 h_1 + w_{12}^2 h_2 + w_{13}^2 h_3 + w_{14}^2 h_4 + b_1^2)$$

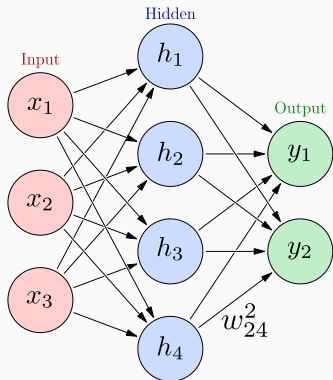
$$y_2 = g_2 (w_{21}^2 h_1 + w_{22}^2 h_2 + w_{23}^2 h_3 + w_{24}^2 h_4 + b_2^2)$$

w_{ij}^k synaptic weight between previous node j and next node i at layer k .

g_k are any activation function applied to each coefficient of its input vector.

Feedforward Artificial Neural Network

Artificial neural network / Multilayer perceptron / NeuralNet



$$h_1 = g_1 (w_{11}^1 x_1 + w_{12}^1 x_2 + w_{13}^1 x_3 + b_1^1)$$

$$h_2 = g_1 (w_{21}^1 x_1 + w_{22}^1 x_2 + w_{23}^1 x_3 + b_2^1)$$

$$h_3 = g_1 (w_{31}^1 x_1 + w_{32}^1 x_2 + w_{33}^1 x_3 + b_3^1)$$

$$h_4 = g_1 (w_{41}^1 x_1 + w_{42}^1 x_2 + w_{43}^1 x_3 + b_4^1)$$

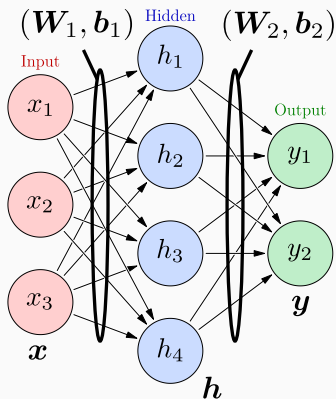
$$y_1 = g_2 (w_{11}^2 h_1 + w_{12}^2 h_2 + w_{13}^2 h_3 + w_{14}^2 h_4 + b_1^2)$$

$$y_2 = g_2 (w_{21}^2 h_1 + w_{22}^2 h_2 + w_{23}^2 h_3 + w_{24}^2 h_4 + b_2^2)$$

w_{ij}^k synaptic weight between previous node j and next node i at layer k .

g_k are any activation function applied to each coefficient of its input vector.

Artificial neural network / Multilayer perceptron / NeuralNet



$$h_1 = g_1 (w_{11}^1 x_1 + w_{12}^1 x_2 + w_{13}^1 x_3 + b_1^1)$$

$$h_2 = g_1 (w_{21}^1 x_1 + w_{22}^1 x_2 + w_{23}^1 x_3 + b_2^1)$$

$$h_3 = g_1 (w_{31}^1 x_1 + w_{32}^1 x_2 + w_{33}^1 x_3 + b_3^1)$$

$$h_4 = g_1 (w_{41}^1 x_1 + w_{42}^1 x_2 + w_{43}^1 x_3 + b_4^1)$$

$$\mathbf{h} = g_1 (\mathbf{W}_1 \mathbf{x} + \mathbf{b}_1)$$

$$y_1 = g_2 (w_{11}^2 h_1 + w_{12}^2 h_2 + w_{13}^2 h_3 + w_{14}^2 h_4 + b_1^2)$$

$$y_2 = g_2 (w_{21}^2 h_1 + w_{22}^2 h_2 + w_{23}^2 h_3 + w_{24}^2 h_4 + b_2^2)$$

$$\mathbf{y} = g_2 (\mathbf{W}_2 \mathbf{h} + \mathbf{b}_2)$$

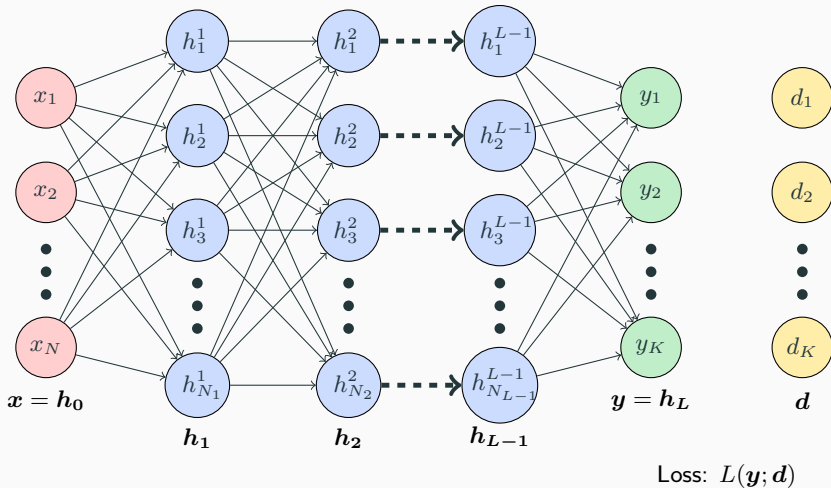
w_{ij}^k synaptic weight between previous node j and next node i at layer k .

g_k are any activation function applied to each coefficient of its input vector.

The matrices $\mathbf{W}_k$ and biases $\mathbf{b}_k$ are learned from labeled training data.

Feedforward Artificial Neural Network

Recall the feedforward structure



Input Layer

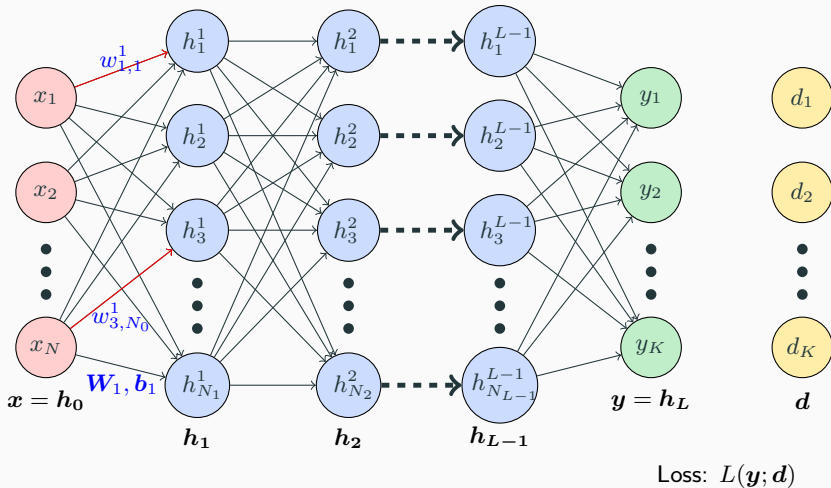
Hidden Layers

Output Layer

Label

Feedforward Artificial Neural Network

Recall the feedforward structure



Input Layer

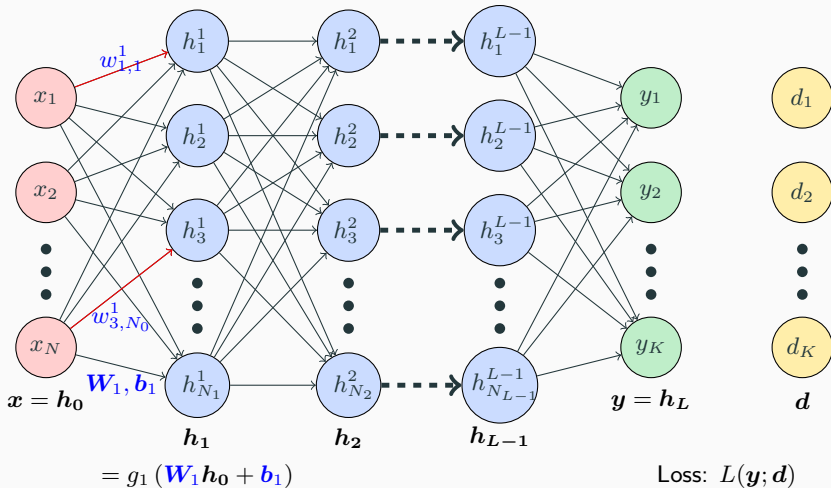
Hidden Layers

Output Layer

Label

Feedforward Artificial Neural Network

Recall the feedforward structure



Input Layer

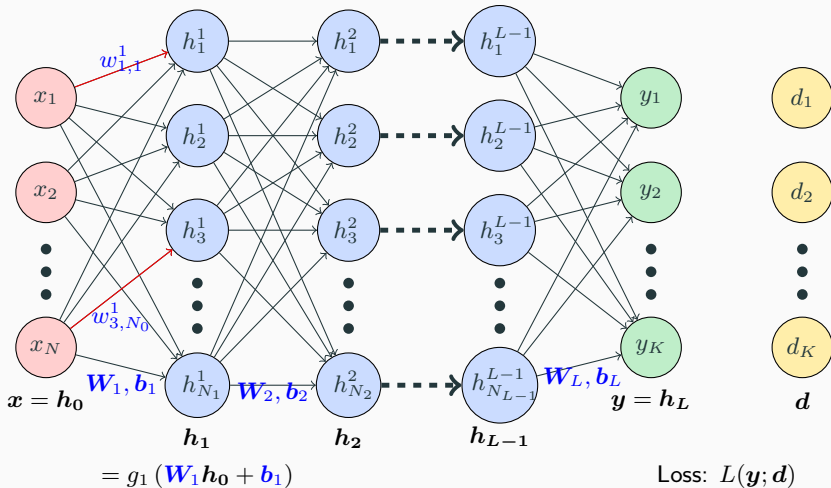
Hidden Layers

Output Layer

Label

Feedforward Artificial Neural Network

Recall the feedforward structure



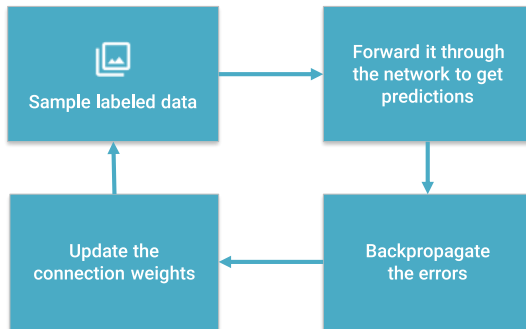
Input Layer

Hidden Layers

Output Layer

Label

Training process



Learns by generating an error signal that measures the difference between the predictions of the network and the desired values and then **using this error signal to change the weights** (or parameters) so that predictions get more accurate.

- The parameters of the neural network are

$$\mathbf{W} = (\mathbf{W}_1, \mathbf{b}_1, \mathbf{W}_2, \mathbf{b}_2, \dots, \mathbf{W}_L, \mathbf{b}_L)$$

- Training the network = minimizing the training loss $E(\mathbf{W})$

Objective: $\min_{\mathbf{W}} E(\mathbf{W})$

$$\Rightarrow \nabla E(\mathbf{W}) = \left(\frac{\partial E(\mathbf{W})}{\partial \mathbf{W}_1} \quad \frac{\partial E(\mathbf{W})}{\partial \mathbf{b}_1} \quad \dots \quad \frac{\partial E(\mathbf{W})}{\partial \mathbf{W}_L} \quad \frac{\partial E(\mathbf{W})}{\partial \mathbf{b}_L} \right)^T = 0$$

- **Solution:** no closed-form solutions $\Rightarrow$ use (stochastic) gradient descent.
- $\frac{\partial E(\mathbf{W})}{\partial \mathbf{W}_k}$ not really rigorous, we will use the notation

$$\nabla_{\mathbf{W}_k} E(\mathbf{W}) \quad \text{and} \quad \nabla_{\mathbf{b}_k} E(\mathbf{W}).$$

Minimizing training loss

For multilayer neural networks $\mathbf{W} \mapsto E(\mathbf{W})$ is non-convex

⇒ No guarantee of convergence.

Even if convergence occurs, the solution depends on the initialization and the step size/learning rate γ .

Nevertheless, really good minima or saddle points are reached in practice by

$$\mathbf{W}^{t+1} \leftarrow \mathbf{W}^t - \gamma \nabla E(\mathbf{W}^t), \quad \gamma > 0$$

Gradient descent can be expressed coordinate by coordinate as:

$$w_{i,j}^{k,t+1} \leftarrow w_{i,j}^{k,t} - \gamma \frac{\partial E(\mathbf{W}^t)}{\partial w_{i,j}^k}$$

for all weights $w_{i,j}^k$ linking a node j to a node i in the next layer k .

⇒ The algorithm to compute $\frac{\partial E(\mathbf{W})}{\partial w_{i,j}^k}$ for ANNs is called **backpropagation**.

Loss functions: Classical loss functions are:

For regression: $\mathbf{d}^i \in \mathbb{R}^K$

- Square error

$$E(\mathbf{W}) = \sum_{(\mathbf{x}^i, \mathbf{d}^i) \in \mathcal{T}} \frac{1}{2} \|\mathbf{y}^i - \mathbf{d}^i\|_2^2 = \sum_{(\mathbf{x}^i, \mathbf{d}^i) \in \mathcal{T}} \frac{1}{2} \sum_k (y_k^i - d_k^i)^2$$

For multi-class classification: $d^i \in \{1, \dots, K\}$, coded by $\mathbf{d}^i \in \{0, 1\}^K$,

- Cross-entropy with softmax as the last layer

$$E(\mathbf{W}) = - \sum_{(\mathbf{x}^i, \mathbf{d}^i)} \sum_{k=1}^K d_k^i \log y_k^i \quad \text{with} \quad \mathbf{y}^i = f(\mathbf{x}^i; \mathbf{W}) = \text{softmax}(\mathbf{a}^i) \in (0, 1)^K.$$

- Cross-entropy with softmax included in loss (PyTorch convention):

$\mathbf{y}^i = \mathbf{a}^i$ is the output of the last linear layer:

$$E(\mathbf{W}) = - \sum_{(\mathbf{x}^i, \mathbf{d}^i)} \left[a_{d^i} - \log \left(\sum_{k=1}^K \exp(a_k) \right) \right] \quad \text{with } d^i \text{ the class of } \mathbf{x}^i.$$

- The loss functions are of the form

$$E(\mathbf{W}) = \sum_{(\mathbf{x}^i, \mathbf{d}^i)} L(\mathbf{y}^i; \mathbf{d}^i)$$

- By linearity,

$$\nabla E(\mathbf{W}) = \sum_{(\mathbf{x}^i, \mathbf{d}^i)} \nabla L(\mathbf{y}^i; \mathbf{d}^i)$$

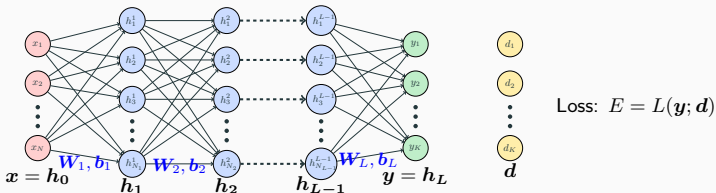
- There the neural net output $\mathbf{y}^i = f(\mathbf{x}^i; \mathbf{W})$ is a function of the input data $\mathbf{x}^i$ and the neural weights $\mathbf{W}$.
- We know the gradient of $L(\mathbf{y}^i; \mathbf{d}^i)$ with respect to the variable $\mathbf{y}$:
 - Regression/Square error:

$$L(\mathbf{y}; \mathbf{d}) = \frac{1}{2} \|\mathbf{y} - \mathbf{d}\|_2^2 \quad \Rightarrow \quad \nabla_{\mathbf{y}} L(\mathbf{y}; \mathbf{d}) = \mathbf{y} - \mathbf{d}$$

- Multi-class classification/cross-entropy:

$$L(\mathbf{y}; \mathbf{d}) = -y_{d^i} + \log \left(\sum_{k=1}^K \exp(y_k) \right) \quad \Rightarrow \quad \nabla_{\mathbf{y}} L(\mathbf{y}; \mathbf{d}) = \text{softmax}(\mathbf{y}) - \mathbf{d}^i.$$

ANN – Backpropagation



Forward pass

Initialization:

$$h_0 = x$$

for layer $k = 1$ **to** L **do**

Linear unit:

$$a_k = W_k h_{k-1} + b_k$$

Componentwise non-linear activation:

$$h_k = g_k(a_k)$$

end

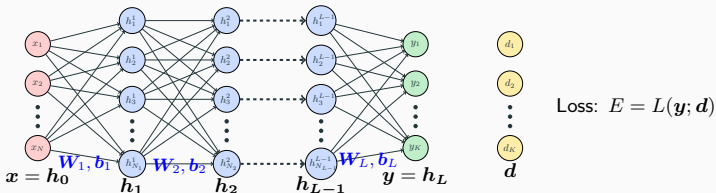
Output layer:

$$y = h_L$$

Compute loss:

$$E = L(y; d)$$

ANN – Backpropagation



Forward pass

Initialization:

$$\mathbf{h}_0 = \mathbf{x}$$

for layer $k = 1$ **to** L **do**

Linear unit:

$$\mathbf{a}_k = \mathbf{W}_k \mathbf{h}_{k-1} + \mathbf{b}_k$$

Componentwise non-linear activation:

$$\mathbf{h}_k = g_k(\mathbf{a}_k)$$

end

Output layer:

$$\mathbf{y} = \mathbf{h}_L$$

Compute loss:

$$E = L(\mathbf{y}; \mathbf{d})$$

Backward pass

Goal: Compute the gradient with respect to all parameters

$$\frac{\partial E}{\partial w_{i,j}^k} = ? \quad \frac{\partial E}{\partial b_i^k} = ?$$

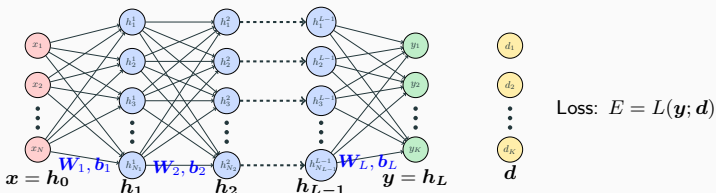
for all

$$k \in \{1, \dots, L\},$$

$$i \in \{1, \dots, N_k\},$$

$$j \in \{1, \dots, N_{k-1}\}.$$

ANN – Backpropagation

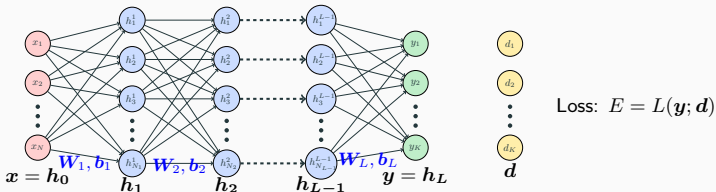


Going backward

- We know how to compute the loss function and its gradient:

$$\nabla_{h_L} E = \nabla L(y; d)$$

ANN – Backpropagation



Gradient with respect to last linear unit output $\mathbf{a}_L$

$$\mathbf{h}_L = g_L(\mathbf{a}_L)$$

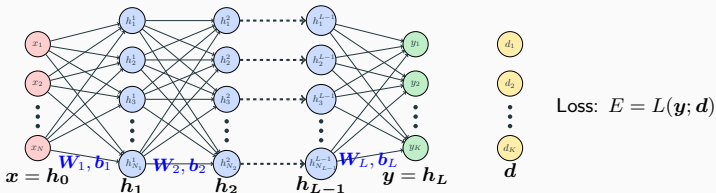
That is for all $i \in \{1, \dots, N_L\}$, $h_i^L = g_L(a_i^L)$. By the chain rule,

$$\frac{\partial E}{\partial a_i^L} = \frac{\partial E}{\partial h_i^L} \frac{\partial h_i^L}{\partial a_i^L} = [\nabla_{\mathbf{h}_L} E]_i g'_L(a_i^L)$$

Vector formula: $\nabla_{\mathbf{a}_L} E = \nabla_{\mathbf{h}_L} E \odot g'_L(\mathbf{a}_L)$

where $\odot$ is the componentwise product between vectors, ie Hadamard product.

ANN – Backpropagation



Gradient with respect to bias of last linear unit b_L

$$a_L = W_L h_{L-1} + b_L$$

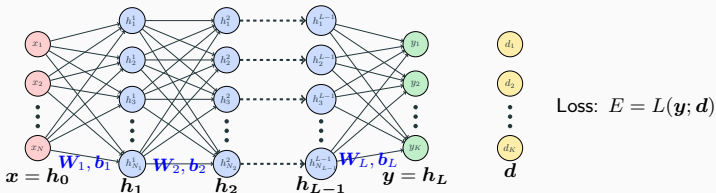
That is for all $i \in \{1, \dots, N_L\}$, $a_i^L = \sum_{j=1}^{N_{L-1}} w_{i,j}^L h_j^{L-1} + b_i^L$.

By the chain rule, for all $i \in \{1, \dots, N_L\}$,

$$\frac{\partial E}{\partial b_i^L} = \frac{\partial E}{\partial a_i^L} \underbrace{\frac{\partial a_i^L}{\partial b_i^L}}_{=1} = \frac{\partial E}{\partial a_i^L} = [\nabla_{a_L} E]_i$$

Vector formula: $\nabla_{b_L} E = \nabla_{a_L} E$

ANN – Backpropagation



Gradient with respect to weights of last linear unit W_L

$$a_L = W_L h_{L-1} + b_L$$

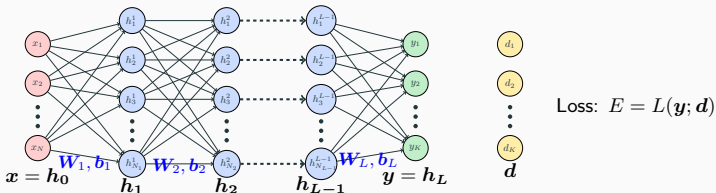
That is for all $i \in \{1, \dots, N_L\}$, $a_i^L = \sum_{j=1}^{N_{L-1}} w_{i,j}^L h_j^{L-1} + b_i^L$.

By the chain rule, for all $i \in \{1, \dots, N_L\}$ and $j \in \{1, \dots, N_{L-1}\}$

$$\frac{\partial E}{\partial w_{i,j}^L} = \frac{\partial E}{\partial a_i^L} \underbrace{\frac{\partial a_i^L}{\partial w_{i,j}^L}}_{=h_j^{L-1}} = \frac{\partial E}{\partial a_i^L} h_j^{L-1} = [\nabla_{a_L} E]_i [h_{L-1}]_j$$

Matrix formula: $\nabla_{W_L} E = \nabla_{a_L} E h_{L-1}^T$

ANN – Backpropagation

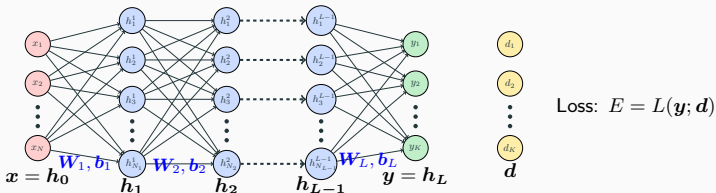


Gradients for last layer parameters

Given the gradient with respect to the output layer $\nabla_{h_L} E$, so far we can compute:

- $\nabla_{a_L} E = \nabla_{h_L} E \odot g'_L(a_L)$
- $\nabla_{b_L} E = \nabla_{a_L} E$
- $\nabla_{W_L} E = \nabla_{a_L} E h_{L-1}^T$

ANN – Backpropagation



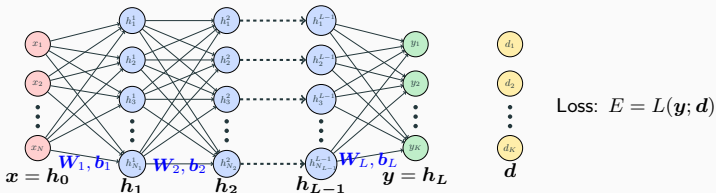
Gradients for last layer parameters

Given the gradient with respect to the output layer $\nabla_{h_L} E$, so far we can compute:

- $\nabla_{a_L} E = \nabla_{h_L} E \odot g'_L(a_L)$
- $\nabla_{b_L} E = \nabla_{a_L} E$
- $\nabla_{w_L} E = \nabla_{a_L} E h_{L-1}^T$

How can we compute the gradients for the parameters of layer $L - 1$?

ANN – Backpropagation



Gradients for last layer parameters

Given the gradient with respect to the output layer $\nabla_{h_L} E$, so far we can compute:

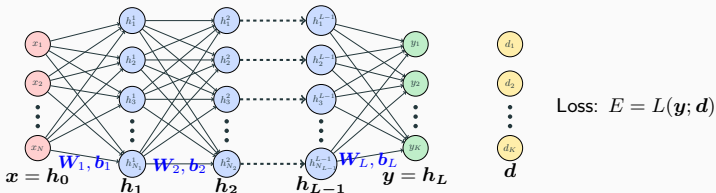
- $\nabla_{a_L} E = \nabla_{h_L} E \odot g'_L(a_L)$
- $\nabla_{b_L} E = \nabla_{a_L} E$
- $\nabla_{W_L} E = \nabla_{a_L} E h_{L-1}^T$

How can we compute the gradients for the parameters of layer $L - 1$?

We need the expression of the gradient with respect to the last but one hidden layer h_{L-1} ... and then the same formulas apply!

$$\nabla_{h_{L-1}} E = ?$$

ANN – Backpropagation

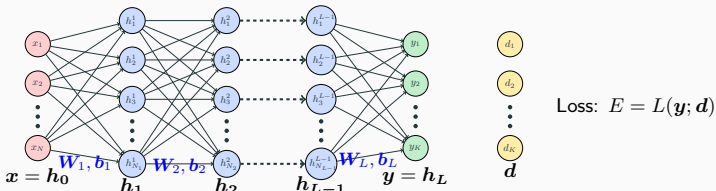


Gradient with respect to the last but one hidden layer h_{L-1}

Here, even to compute the scalar partial derivative $\frac{\partial E}{\partial h_j^{L-1}}$, we need to use differential calculus for multivariate functions since h_j^{L-1} appears in each component of $\mathbf{a}_L$:

$$\text{For all } i \in \{1, \dots, N_L\}, a_i^L = \sum_{j=1}^{N_{L-1}} w_{i,j}^L h_j^{L-1} + b_i^L.$$

ANN – Backpropagation



Gradient with respect to the last but one hidden layer h_{L-1}

Let us recall the derivative rule for composition with affine maps:

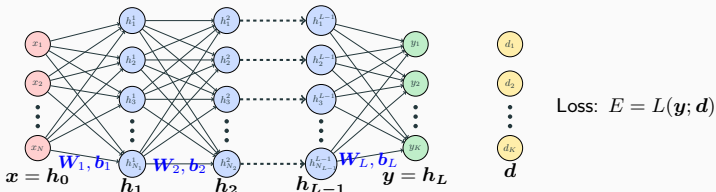
$$\text{For } \varphi(x) = f(Ax + b) \text{ one has } \nabla \varphi(x) = A^T \nabla f(Ax + b).$$

Using the decomposition

$$\begin{aligned} \mathbb{R}^{N_{L-1}} &\rightarrow \mathbb{R}^{N_L} \rightarrow \mathbb{R} \\ h_{L-1} &\mapsto a_L = W_L h_{L-1} + b_L \mapsto E \end{aligned}$$

$$\text{Vector formula: } \nabla_{h_{L-1}} E = W_L^T \nabla_{a_L} E$$

ANN – Backpropagation



Forward pass

Initialization:

$$\mathbf{h}_0 = \mathbf{x}$$

for layer $k = 1$ **to** L **do**

Linear unit:

$$\mathbf{a}_k = \mathbf{W}_k \mathbf{h}_{k-1} + \mathbf{b}_k$$

Componentwise non-linear activation:

$$\mathbf{h}_k = g_k(\mathbf{a}_k)$$

end

Output layer:

$$\mathbf{y} = \mathbf{h}_L$$

Compute loss:

$$E = L(\mathbf{y}; \mathbf{d})$$

Backward pass

Initialization: Gradient of output layer:

$$\nabla_{\mathbf{h}_L} E = \nabla L(\mathbf{y}; \mathbf{d})$$

for layer $k = L$ **to** 1 **do**

Componentwise gain of error:

$$\delta_k = \nabla_{\mathbf{a}_k} E = \nabla_{\mathbf{h}_k} E \odot g'_k(\mathbf{a}_k)$$

Gradient of layer bias:

$$\nabla_{\mathbf{b}_k} E = \delta_k$$

Gradient of weights:

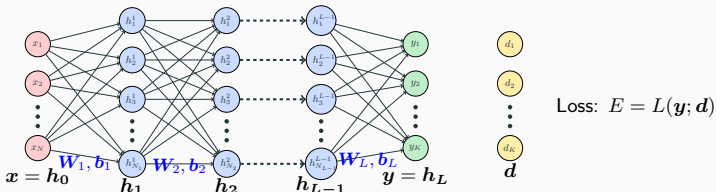
$$\nabla_{\mathbf{W}_k} E = \delta_k \mathbf{h}_{k-1}^T$$

Gradient of previous hidden layer:

$$\nabla_{\mathbf{h}_{k-1}} E = \mathbf{W}_k^T \delta_k$$

end

ANN – Backpropagation



Forward pass

Initialization:

$$\mathbf{h}_0 = \mathbf{x}$$

for layer $k = 1$ **to** L **do**

Linear unit:

$$\mathbf{a}_k = \mathbf{W}_k \mathbf{h}_{k-1} + \mathbf{b}_k \text{ (stored)}$$

Componentwise non-linear activation:

$$\mathbf{h}_k = g_k(\mathbf{a}_k) \text{ (stored)}$$

end

Output layer:

$$\mathbf{y} = \mathbf{h}_L$$

Compute loss:

$$E = L(\mathbf{y}; \mathbf{d})$$

Backward pass

Initialization: Gradient of output layer:

$$\nabla_{\mathbf{h}_L} E = \nabla L(\mathbf{y}; \mathbf{d})$$

for layer $k = L$ **to** 1 **do**

Componentwise gain of error:

$$\delta_k = \nabla_{\mathbf{a}_k} E = \nabla_{\mathbf{h}_k} E \odot g'_k(\mathbf{a}_k)$$

Gradient of layer bias:

$$\nabla_{\mathbf{b}_k} E = \delta_k$$

Gradient of weights:

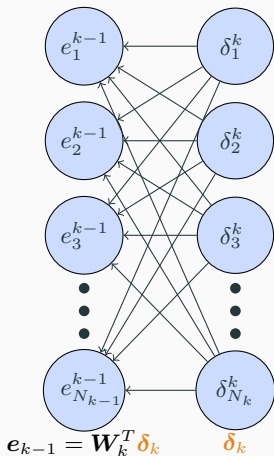
$$\nabla_{\mathbf{W}_k} E = \delta_k \mathbf{h}_{k-1}^T$$

Gradient of previous hidden layer:

$$\nabla_{\mathbf{h}_{k-1}} E = \mathbf{W}_k^T \delta_k$$

end

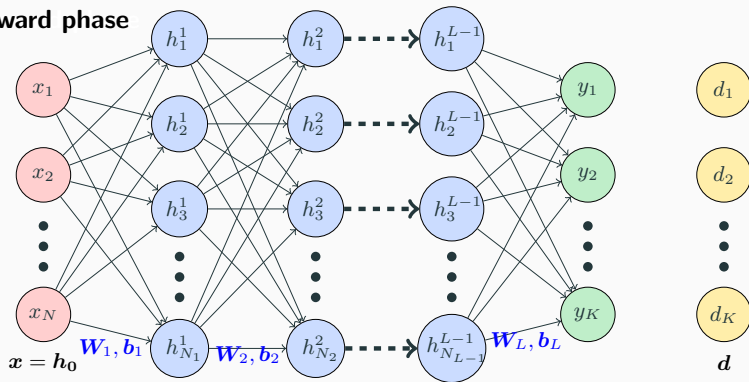
Error backpropagation



- Gradient of previous hidden layer:
$$e_{k-1} = \nabla_{h_{k-1}} E = \mathbf{W}_k^T \delta_k$$
- Multiplying by $\mathbf{W}_k^T$ corresponds to passing to the linear layer in reverse order.
- The error is backpropagated layer by layer to compute the gradient with respect to each layer parameters.

Error backpropagation

Forward phase



Input Layer

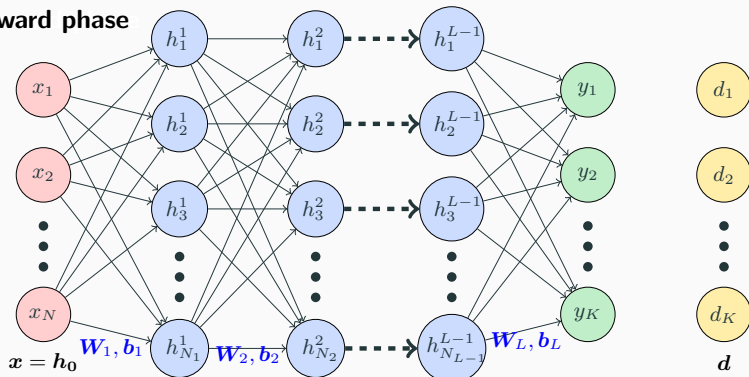
Hidden Layers

Output Layer

Label

Error backpropagation

Forward phase



Input Layer

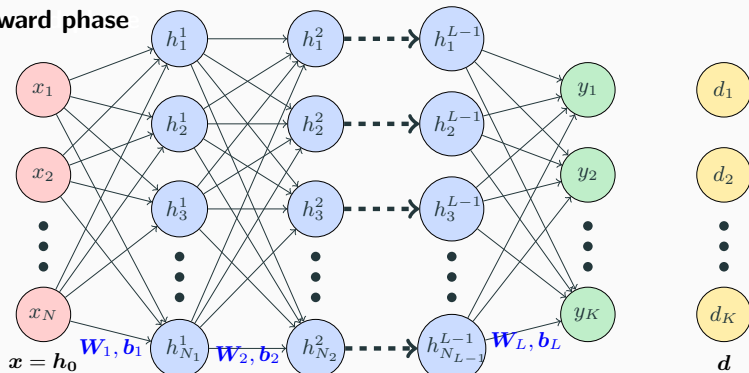
Hidden Layers

Output Layer

Label

Error backpropagation

Forward phase



Input Layer

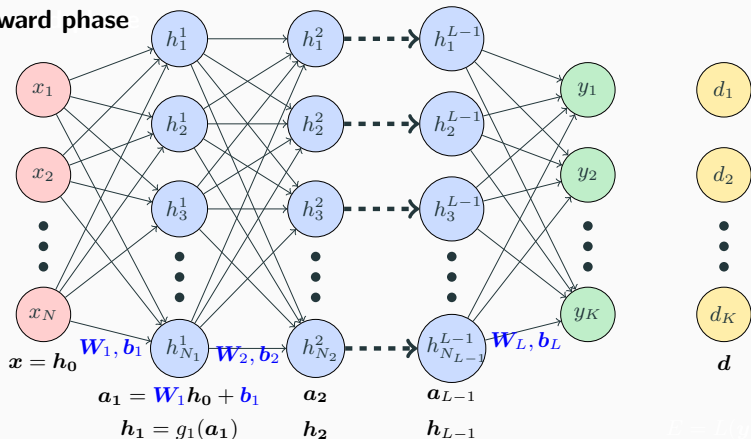
Hidden Layers

Output Layer

Label

Error backpropagation

Forward phase



Input Layer

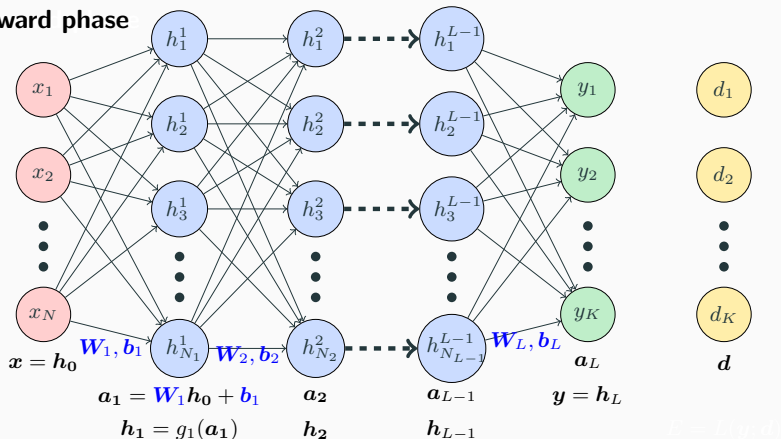
Hidden Layers

Output Layer

Label

Error backpropagation

Forward phase



Input Layer

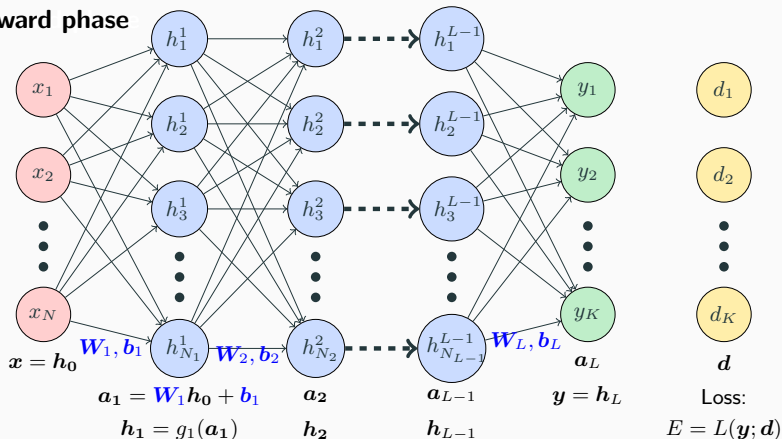
Hidden Layers

Output Layer

Label

Error backpropagation

Forward phase



Input Layer

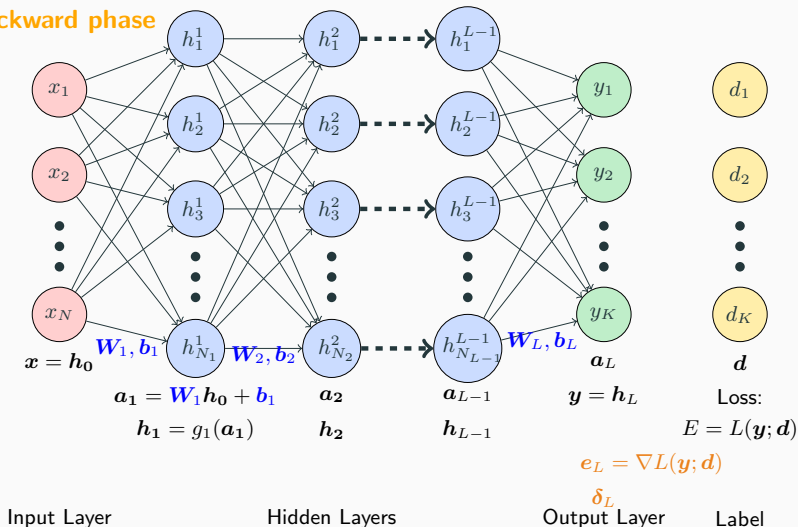
Hidden Layers

Output Layer

Label

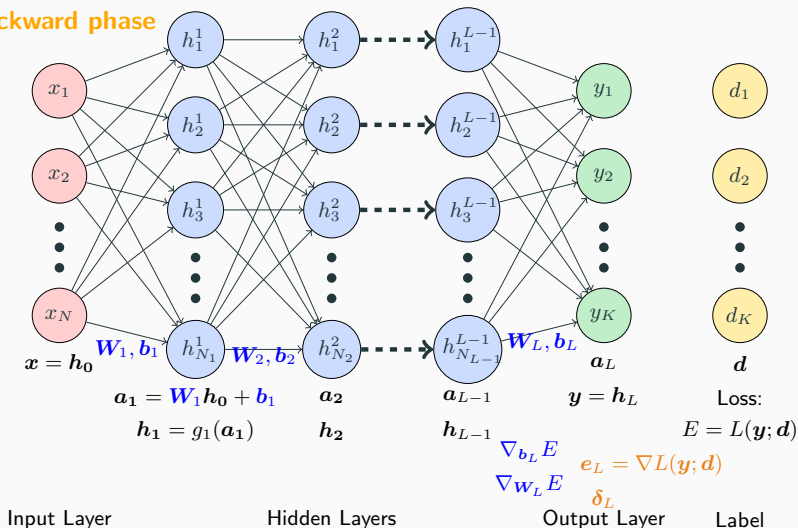
Error backpropagation

Backward phase



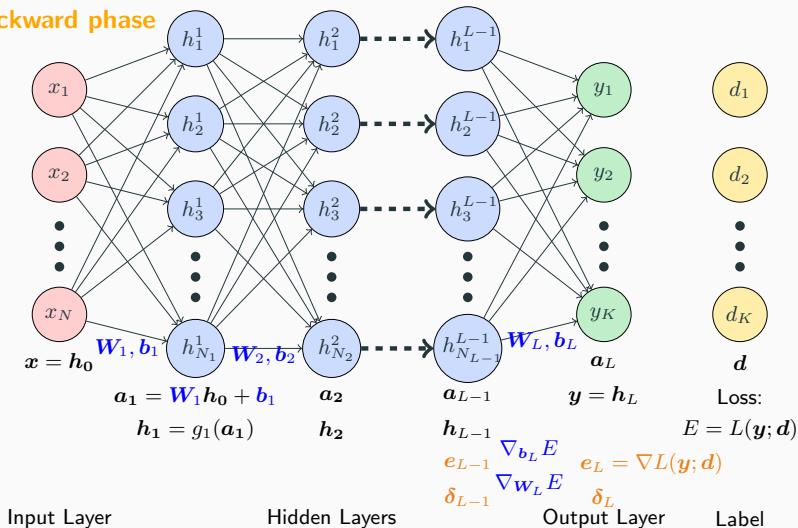
Error backpropagation

Backward phase



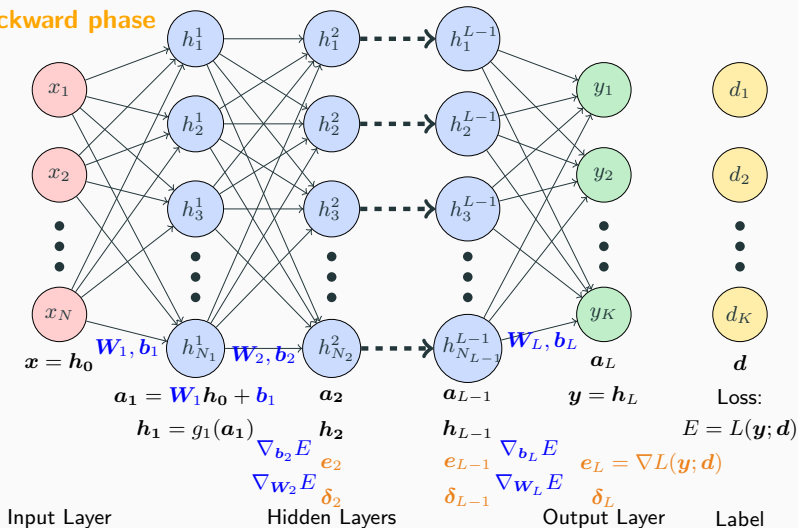
Error backpropagation

Backward phase



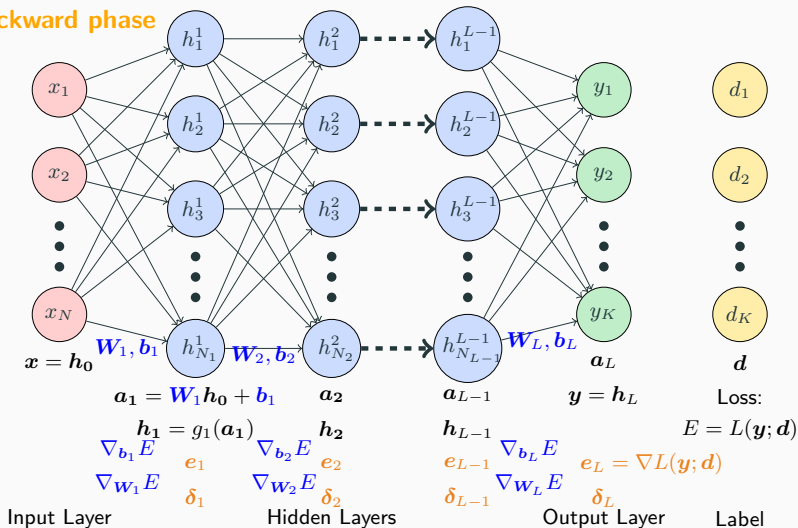
Error backpropagation

Backward phase



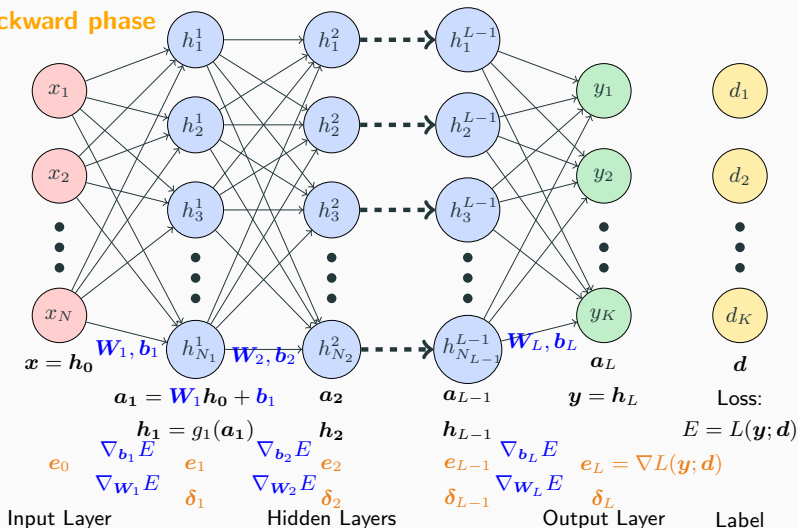
Error backpropagation

Backward phase



Error backpropagation

Backward phase



Error backpropagation in practice

Training loss:

$$E(\mathbf{W}) = \sum_{(\mathbf{x}^i, \mathbf{d}^i) \in \mathcal{T}} L(\mathbf{y}^i; \mathbf{d}^i)$$

- The backpropagation procedure computes $\nabla_{\mathbf{W}} L(\mathbf{y}^i; \mathbf{d}^i) = \nabla_{\mathbf{W}} L(f(\mathbf{x}^i; \mathbf{W}); \mathbf{d}^i)$.
- This has to be done for each data point $\mathbf{x}^i \in \mathcal{T}$.
- By linearity, the final gradient $\nabla E(\mathbf{W})$ is the sum of all individual gradients $\nabla_{\mathbf{W}} L(\mathbf{y}^i; \mathbf{d}^i)$.
- These gradients are summed sequentially (no need to store each individual gradients).
- In general we do not compute the exact gradient...

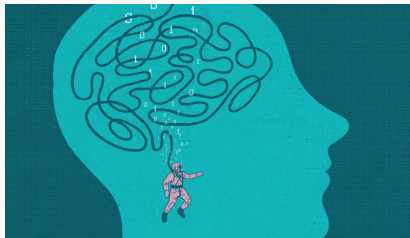
Error backpropagation in practice

Batch loss:

$$E(\mathbf{W}) \approx \sum_{(\mathbf{x}^i, \mathbf{d}^i) \in \mathcal{S}} L(\mathbf{y}^i; \mathbf{d}^i), \quad \text{with } \mathcal{S} \subset \mathcal{T}$$

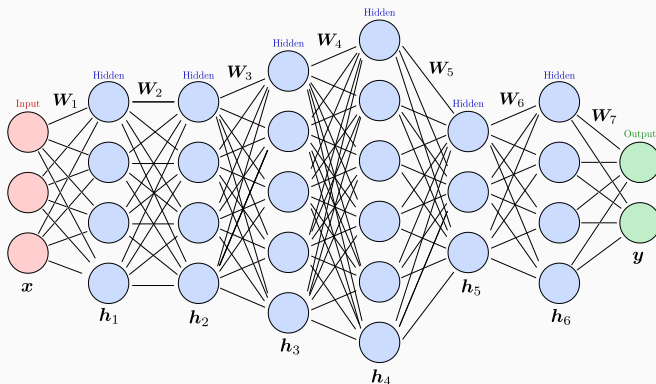
- The backpropagation has to be done for each visited data point $\mathbf{x}^i \in \mathcal{S}$ of the batch.
- The gradient for each point $\mathbf{x}^i$ is added to the running gradient = current gradient estimation.
- Once the noisy estimated gradient is used as a gradient step, one needs to set the gradients to zero: See PyTorch `zero_grad()` procedure.

Deep learning



(Source: Dan Page)

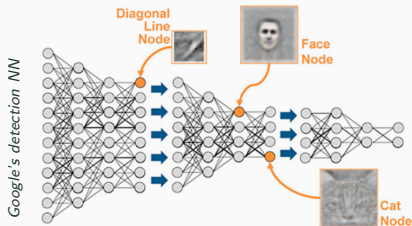
ANNs recap



- Interconnections of neurons (summations and activations),
- Feedforward architecture: input $\rightarrow$ hidden layer(s) $\rightarrow$ output,
- Parameterized by a collection of weights W_k (and biases),
- Backprop: parameters learned from a collection of labeled data.

What is deep learning?

- Representation learning using artificial neural networks
 - Learning good features automatically from raw data.
 - Exceptionally effective at learning patterns.
- Learning representations of data with multiple levels of abstraction
 - hierarchy of layers that mimic the neural networks of our brain,
 - cascade of non-linear transforms.



- If you provide the system with tons of information, it begins to understand it and responds in useful ways.

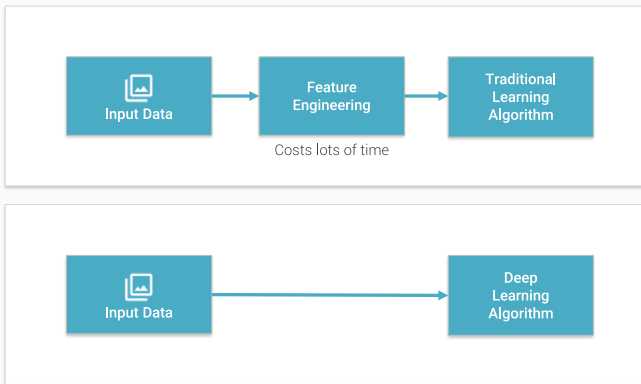
How to teach a machine?



Good representations are often very complex to define.

(Source: Caner Hazırbaş)

Deep learning: No more feature engineering



Learn the latent factors/features behind the data generation

(Source: Lucas Masuch)

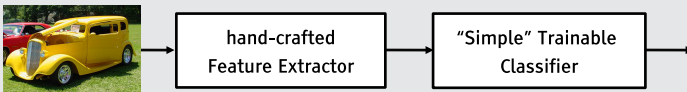
Inspired by the Brain

- The first hierarchy of neurons that receives information in the visual cortex are sensitive to specific edges while brain regions further down the visual pipeline are sensitive to more complex structures such as faces.
- Our brain has lots of neurons connected together and the **strength of the connections** between neurons represents **long term knowledge**.
- **One learning algorithm hypothesis**: all significant mental algorithms are learned except for the learning and reward machinery itself.

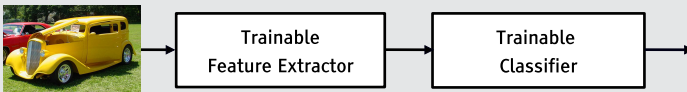
(Source: Lucas Masuch)

Deep learning: No more feature engineering

- The traditional model of pattern recognition (since the late 50's)
 - Fixed/engineered features (or fixed kernel) + trainable classifier



- End-to-end learning / Feature learning / Deep learning
 - Trainable features (or kernel) + trainable classifier



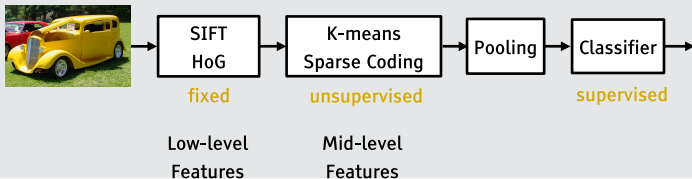
(Source: Yann LeCun & Marc'Aurelio Ranzato)

Deep learning: No more feature engineering

- Non-deep learning architecture for pattern recognition
 - Speech recognition: early 90's – 2011

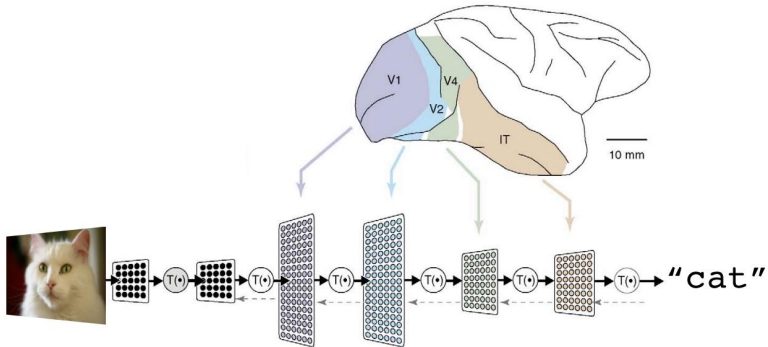


- Object Recognition: 2006 – 2012



(Source: Yann LeCun & Marc Aurelio Ranzato)

Deep learning – Basic architecture



A deep neural network consists of a **hierarchy of layers**, whereby each layer **transforms the input data** into more abstract representations (e.g. edge -> nose -> face). The output layer combines those features to make predictions.

Trainable feature hierarchy

- Hierarchy of representations with increasing levels of abstraction.
- Each stage is a kind of trainable feature transform.

Image recognition

- Pixel → edge → texon → motif → part → object

Text

- Character → word → word group → clause → sentence → story

Speech

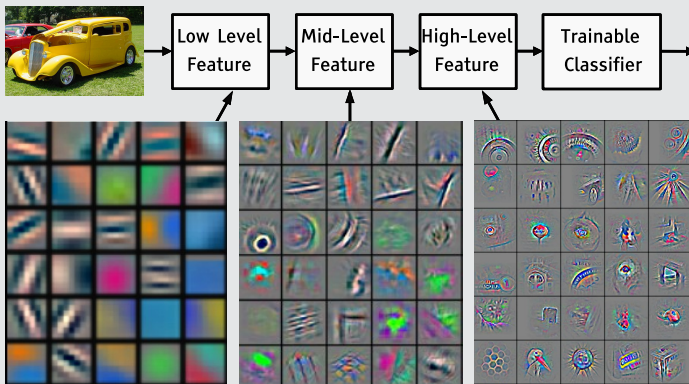
- Sample → spectral band → sound → ... → phone → phoneme → word

Deep Learning addresses the problem of learning hierarchical representations.

(Source: Yann LeCun & Marc'Aurelio Ranzato)

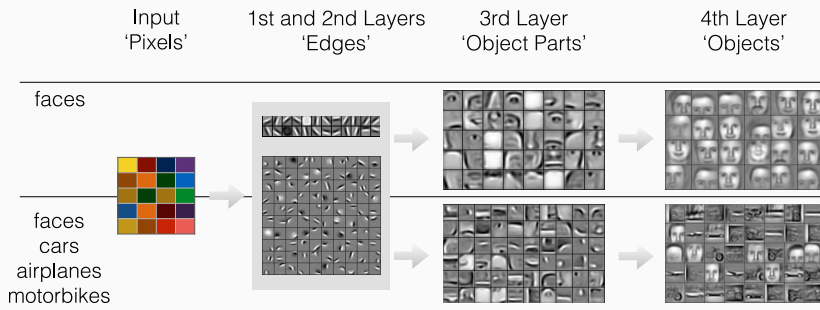
Deep learning – Feature hierarchy

- It's **deep** if it has **more than one stage** of non-linear feature transformation



Feature visualization of convolutional net
trained on ImageNet from (Zeiler & Fergus, 2013)

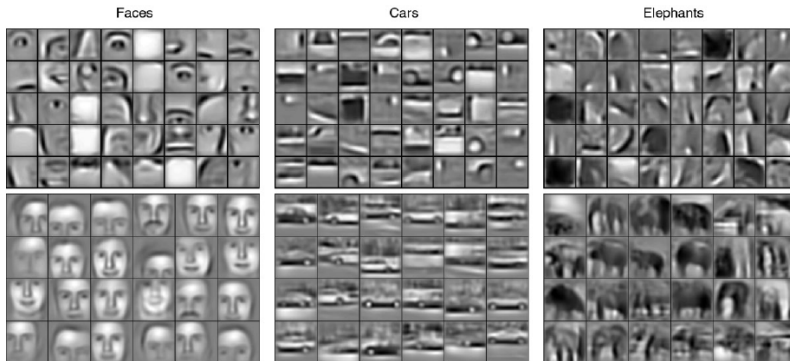
Deep learning – Feature hierarchy



Each layer progressively extracts higher level features of the input until the final layer essentially makes a decision about what the input shows. The more layers the network has, the higher level features it will learn.

(Source: Andrew Ng & Lucas Masuch & Caner Hazırbaş)

Deep learning – Feature hierarchy



- Mid- to high-level features are specific to the given task.
- Low-level representations contain less specific features.
(can be shared for many different applications)

Deep learning – Training

Today's trend: make it deeper and deeper

- 2012: 8 layers (AlexNet – Krizhevsky *et al.*, 2012)
- 2014: 19 layers (VGG Net – Simonyan & Zisserman, 2014)
- 2014: 22 layers (GoogLeNet – Szegedy *et al.*, 2014)
- 2015: 152 layers (ResNet – He *et al.*, 2015)
- 2016: 201 layers (DenseNet – Huang *et al.*, 2017)

But remember, with back-propagation:

- We got stuck at local optima or saddle points
- The learning time does not scale well
 - it is very slow for deep networks and can be unstable.

How did networks get so deep? First, why does backprop fail?

Deep learning – Gradient vanishing problems

Back-propagation and gradient vanishing problems

$$\text{Update } w_{i,j} \leftarrow w_{i,j} - \gamma \frac{\partial E(\mathbf{W})}{\partial w_{i,j}} \quad \text{with} \quad \frac{\partial E(\mathbf{W})}{\partial w_{i,j}} = \sum_{(\mathbf{x}, \mathbf{d}) \in \mathcal{T}} \delta_i h_j$$

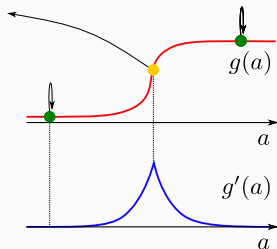
$$\text{where } \delta_i = g'(a_i) \times \begin{cases} y_i - d_i & \text{if } i \text{ is an output node} \\ \sum_l w_{l,i} \delta_l & \text{otherwise} \end{cases}$$

- With deep networks, the **gradient vanishes quickly**.
- Unfortunately, this arises even though we are **far from a solution**.
- The updates become insignificant, which leads to **slow training rates**.
- This strongly depends on the **shape of $g'(a)$** .
- The gradient may also explode leading to instabilities:
→ **gradient exploding problem**.

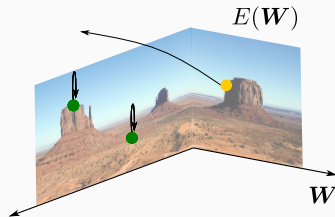
Deep learning – Gradient vanishing problem

As the network gets deeper, the landscape of E becomes:

- very hilly → lots of stationary points,
- with large plateaus → gradient vanishing problem,
- and delimited by cliffs → gradient exploding problem.



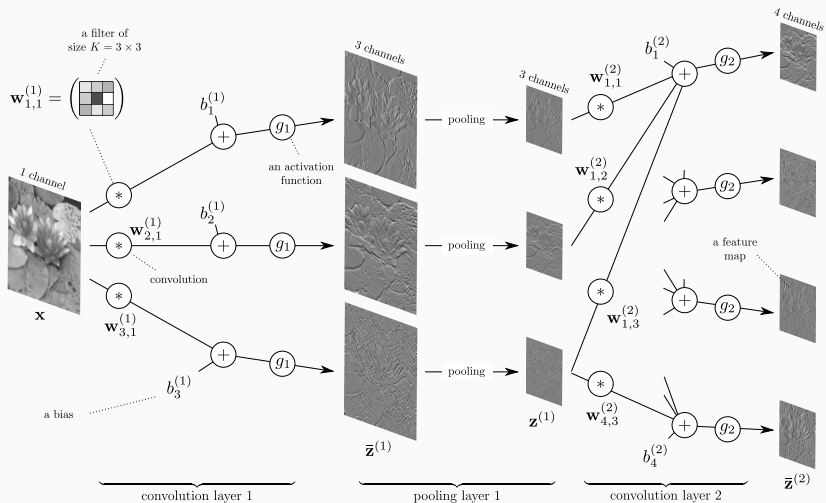
Activation function



Cost landscape

So, what has changed? (see later for recipes...)

CNN for image processing



Convolution: One channel

For an image $x = x(i, j)$ having **one channel**, the convolution with a kernel κ of size $[-s, s] \times [-s, s]$

- the **convolution** is:

$$x * \kappa(i, j) = \sum_{(k, \ell) \in [-s, s] \times [-s, s]} \kappa(k, \ell) x(i - k, j - \ell)$$

$x * \kappa$ = local average of x with the weight of κ in **opposite position**

- the **cross-correlation** is:

$$x \otimes \kappa(i, j) = \sum_{(k, \ell) \in [-s, s] \times [-s, s]} \kappa(k, \ell) x(i + k, j + \ell).$$

$x \otimes \kappa$ = local average of x with the weight of κ in **same position**

- Need to deal with boundary issues for pixels at the border: zero-padding (0 if outside border), valid positions only (do not compute at border, smaller output images), mirror symmetry at border...

- Images have generally several channels (eg RGB).
- By computing several convolutions we can stack the result as a single multi-channel output.

In machine learning
“convolution”
means
“cross-correlation + bias”

- Recall that “linear” layers are affine map $x \mapsto Wx + b$, so convolution layers are a specific case where $x \mapsto Wx$ is a cross-correlation.

Convolution layer with c_{in} input channels and c_{out} output channels:

- Input image $\mathbf{x}$ with c_{in} **channels**: values $\mathbf{x}(i, j) \in \mathbb{R}^{c_{\text{in}}}$
- Output image $\mathbf{y}$ with c_{out} **channels**.
- Kernel: κ such that for all $(k, \ell) \in [-s, s] \times [-s, s]$

$$\kappa(k, \ell) \in \mathbb{R}^{c_{\text{out}} \times c_{\text{in}}}, \quad \text{is a } c_{\text{out}} \times c_{\text{in}} \text{ matrix}$$

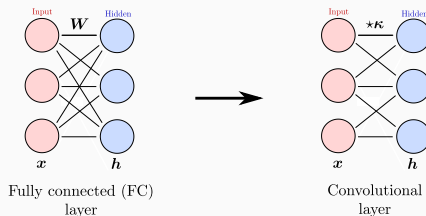
- Bias: $\mathbf{b} \in \mathbb{R}^{c_{\text{out}}}$

$$\begin{aligned} \mathbf{y}(i, j) &= \text{Conv}(\mathbf{x}; \kappa, \mathbf{b})(i, j) \\ &= \left[\sum_{(k, \ell) \in [-s, s] \times [-s, s]} \kappa(k, \ell) \mathbf{x}(i + k, j + \ell) \right] + \mathbf{b} \in \mathbb{R}^{c_{\text{out}}} \end{aligned}$$

- Number of parameters: $(2s + 1)^2 \times c_{\text{in}} \times c_{\text{out}}$ for κ and c_{out} for $\mathbf{b}$

What are CNNs?

- Essentially neural networks that use convolution in place of general matrix multiplications at least for the first layers.



- CNNs are designed to process the data in the form of multidimensional arrays/tensors (e.g., 2D images, 3D video/volumetric images).
- Composed of series of stages: **convolutional** layers and **pooling** layers.
- Units connected to local regions in the feature maps of the previous layer.
- Do not only mimic the brain connectivity but also the **visual cortex**.

CNNs are composed of three main ingredients:

- ① **Local receptive fields**
 - hidden units connected only to a small region of their input,
- ② **Shared weights**
 - same weights and biases for all units of a hidden layer,
- ③ **Pooling**
 - condensing hidden layers.

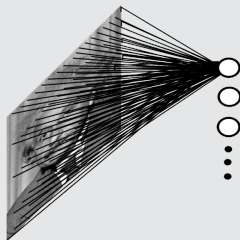
but also

- ④ **Redundancy:** more units in a hidden layer than inputs,
- ⑤ **Sparsity:** units should not all fire for the same stimulus.

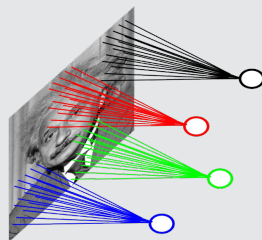
All take inspiration from the **visual cortex.**

Local receptive fields → Locally connected layer

- Each unit in a hidden layer can see only a small neighborhood of its input,
- Captures the concept of spatiality.



Fully connected



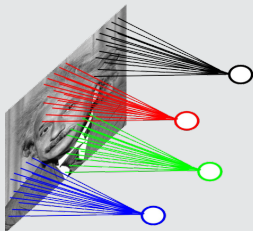
Locally connected

For a 200×200 image and 40,000 hidden units

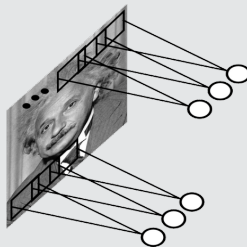
- Fully connected: 1.6 billion parameters,
- Locally connected (10×10 fields): 4 million parameters.

Self-similar receptive fields $\rightarrow$ Shared weights

- Detect features regardless of position (translation invariance),
- Use convolutions to learn simple input patterns.



Locally connected



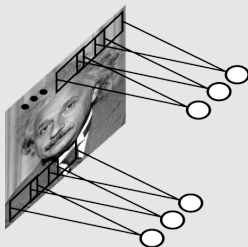
Shared weights

For a 200×200 image and 40,000 hidden units

- Locally connected (10×10 fields): 4 million parameters,
- & Shared weights: 100 parameters (independent of image size).

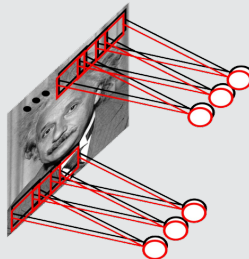
Specialized cells $\rightarrow$ Filter bank

- Use a filter bank to detect multiple patterns at each location,
- Multiple convolutions with different kernels,
- Result is a 3d array, where each slice is a feature map.



Shared weights

(1 input $\rightarrow$ 1 feature map)



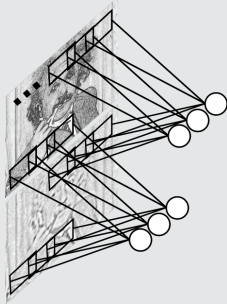
Filter bank

(1 input $\rightarrow$ 2 feature maps)

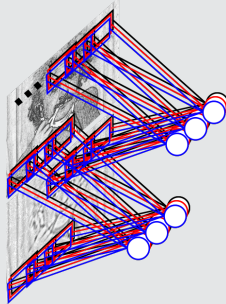
- 10×10 fields & 10 output features: 1,000 parameters.

Hierarchy → inputs of deep layers are themselves 3d arrays

- Learn to filter each channel such that their sum detects a relevant feature,
- Repeat as many times as the desired number of output features should be.



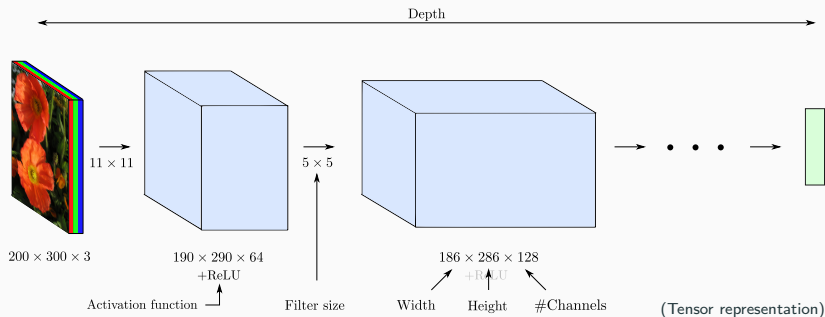
Multi-input filter
(2 inputs → 1 feature map)



Multi-input filter bank
(2 inputs → 3 feature maps)

- **Remark:** these are not 3d convolutions, but sums of 2d convolutions.
- 10×10 fields & 10 inputs & 10 outputs: 10,000 parameters.

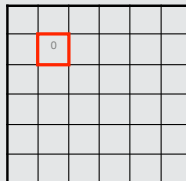
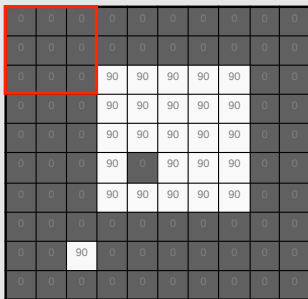
Overcomplete $\rightarrow$ increase the number of channels



- **Redundancy:** increase the number of channels between layers.
- **Padding:** $n \times n$ conv + *valid* $\rightarrow$ width and height decrease by $n - 1$.
- Can we control even more the number of simple cells?

Controlling the number of simple cells → Stride

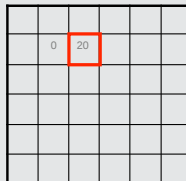
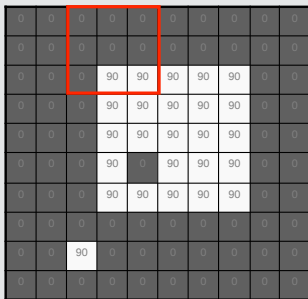
3×3 boxcar strided convolution with stride $s = 2$



- Slide the filter by s pixels step by step, not one by one,
- The interval s is called **stride** (usually $s = 2$),
- $n \times n$ conv + *valid* → width/height decrease to $\lceil \frac{w-n+1}{s} \rceil$ and $\lceil \frac{h-n+1}{s} \rceil$,
- Equivalent in subsampling/decimating the standard convolution,
- Trade-off between computation and degradation of performance.

Controlling the number of simple cells → Stride

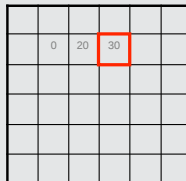
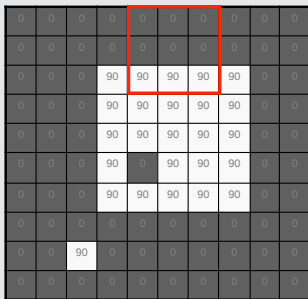
3×3 boxcar strided convolution with stride $s = 2$



- Slide the filter by s pixels step by step, not one by one,
- The interval s is called **stride** (usually $s = 2$),
- $n \times n$ conv + *valid* → width/height decrease to $\lceil \frac{w-n+1}{s} \rceil$ and $\lceil \frac{h-n+1}{s} \rceil$,
- Equivalent in subsampling/decimating the standard convolution,
- Trade-off between computation and degradation of performance.

Controlling the number of simple cells → Stride

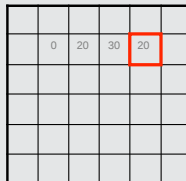
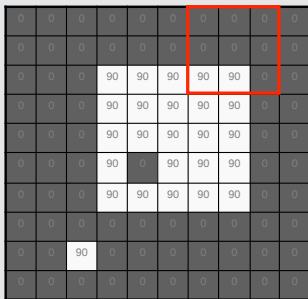
3×3 boxcar strided convolution with stride $s = 2$



- Slide the filter by s pixels step by step, not one by one,
- The interval s is called **stride** (usually $s = 2$),
- $n \times n$ conv + *valid* → width/height decrease to $\lceil \frac{w-n+1}{s} \rceil$ and $\lceil \frac{h-n+1}{s} \rceil$,
- Equivalent in subsampling/decimating the standard convolution,
- Trade-off between computation and degradation of performance.

Controlling the number of simple cells → Stride

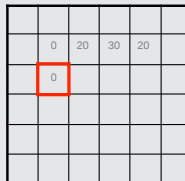
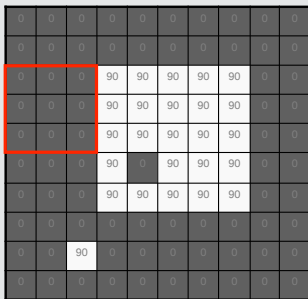
3×3 boxcar strided convolution with stride $s = 2$



- Slide the filter by s pixels step by step, not one by one,
- The interval s is called **stride** (usually $s = 2$),
- $n \times n$ conv + *valid* → width/height decrease to $\lceil \frac{w-n+1}{s} \rceil$ and $\lceil \frac{h-n+1}{s} \rceil$,
- Equivalent in subsampling/decimating the standard convolution,
- Trade-off between computation and degradation of performance.

Controlling the number of simple cells → Stride

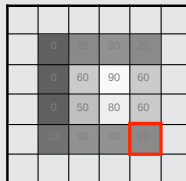
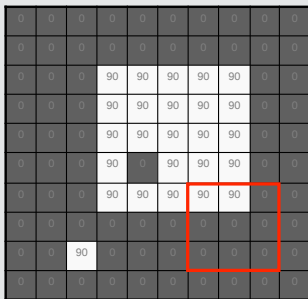
3×3 boxcar strided convolution with stride $s = 2$



- Slide the filter by s pixels step by step, not one by one,
- The interval s is called **stride** (usually $s = 2$),
- $n \times n$ conv + *valid* → width/height decrease to $\lceil \frac{w-n+1}{s} \rceil$ and $\lceil \frac{h-n+1}{s} \rceil$,
- Equivalent in subsampling/decimating the standard convolution,
- Trade-off between computation and degradation of performance.

Controlling the number of simple cells → Stride

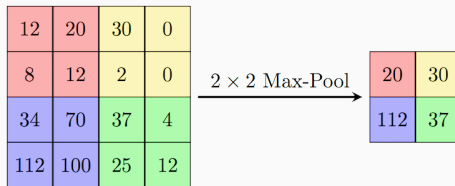
3×3 boxcar strided convolution with stride $s = 2$



- Slide the filter by s pixels step by step, not one by one,
- The interval s is called **stride** (usually $s = 2$),
- $n \times n$ conv + *valid* → width/height decrease to $\lceil \frac{w-n+1}{s} \rceil$ and $\lceil \frac{h-n+1}{s} \rceil$,
- Equivalent in subsampling/decimating the standard convolution,
- Trade-off between computation and degradation of performance.

Pooling layer

- Used after each convolution layer to mimic **complex cells**,
- Unlike striding, reduce the size by **aggregating** inputs:
 - Partition the image in a grid of $z \times z$ windows (usually $z = 2$),
 - **max-pooling**: take the max in the window

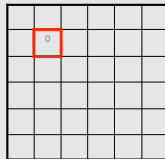
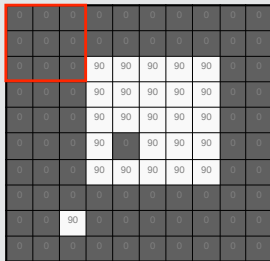


- **average-pooling**: take the average

Pooling layer

Variant: **Overlapping pooling** (Krizhevsky, Sutskever, Hinton, 2012)

- Perform pooling in a $z \times z$ sliding window,
- Use striding every s pixels,
- Typically: use max-pooling with $z = 3$ and $s = 2$:

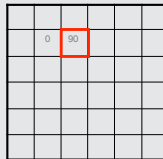
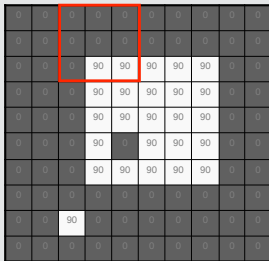


- When $z = s$: standard pooling (non-overlapping),

Pooling layer

Variant: **Overlapping pooling** (Krizhevsky, Sutskever, Hinton, 2012)

- Perform pooling in a $z \times z$ sliding window,
- Use striding every s pixels,
- Typically: use max-pooling with $z = 3$ and $s = 2$:

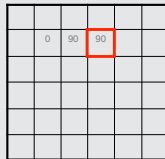
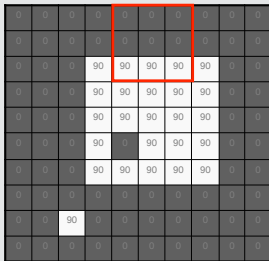


- When $z = s$: standard pooling (non-overlapping),

Pooling layer

Variant: **Overlapping pooling** (Krizhevsky, Sutskever, Hinton, 2012)

- Perform pooling in a $z \times z$ sliding window,
- Use striding every s pixels,
- Typically: use max-pooling with $z = 3$ and $s = 2$:

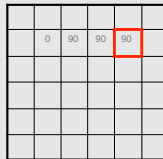
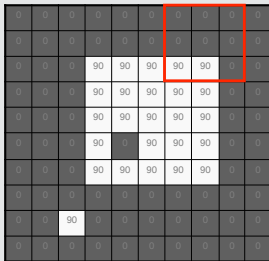


- When $z = s$: standard pooling (non-overlapping),

Pooling layer

Variant: **Overlapping pooling** (Krizhevsky, Sutskever, Hinton, 2012)

- Perform pooling in a $z \times z$ sliding window,
- Use striding every s pixels,
- Typically: use max-pooling with $z = 3$ and $s = 2$:

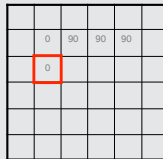
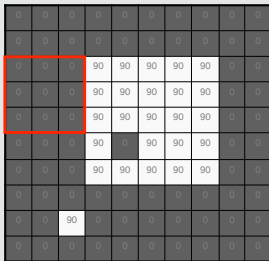


- When $z = s$: standard pooling (non-overlapping),

Pooling layer

Variant: **Overlapping pooling** (Krizhevsky, Sutskever, Hinton, 2012)

- Perform pooling in a $z \times z$ sliding window,
- Use striding every s pixels,
- Typically: use max-pooling with $z = 3$ and $s = 2$:

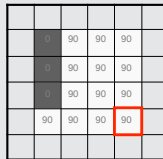
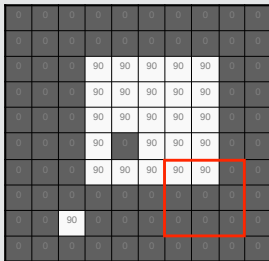


- When $z = s$: standard pooling (non-overlapping),

Pooling layer

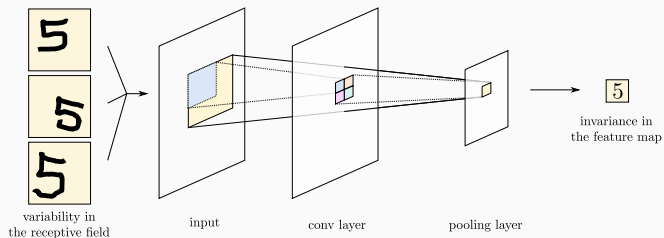
Variant: **Overlapping pooling** (Krizhevsky, Sutskever, Hinton, 2012)

- Perform pooling in a $z \times z$ sliding window,
- Use striding every s pixels,
- Typically: use max-pooling with $z = 3$ and $s = 2$:



- When $z = s$: standard pooling (non-overlapping),

Pooling layer



- Makes the output **unchanged** even if the input is a little bit changed,
- Allows some invariance/robustness with respect to the exact position,
- Simplifies/Condenses/Summarizes the output from hidden layers,
- **Increases the effective receptive fields** (with respect to the first layer.)

CNNs parameterization

Setting up a **convolution layer** requires choosing

- Filter size: $n \times n$
- #output channels: C
- Stride: s
- Padding: p

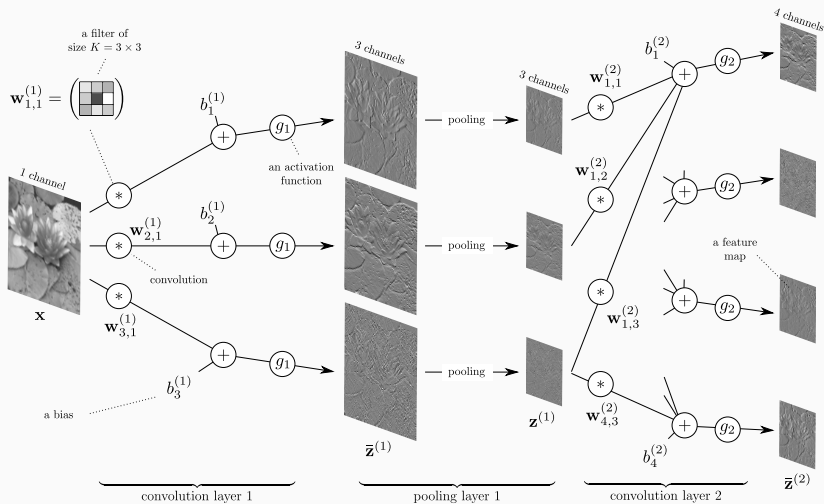
The filter weights κ and the bias b are learned by backprop.

Setting up a **pooling layer** requires choosing

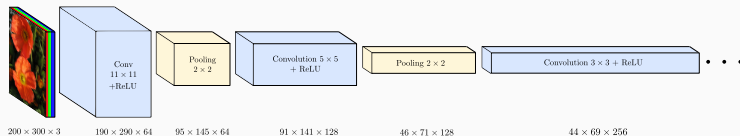
- Pooling size: $z \times z$
- Aggregation rule: max-pooling, average-pooling, ...
- Stride: s
- Padding: p

No free parameters to be learned here.

All concepts together



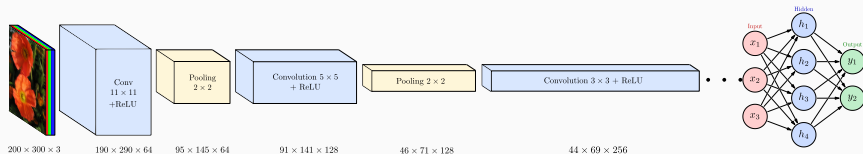
All concepts together with tensor representation



CNN: Alternate:

Conv + ReLU + pooling

All concepts together with tensor representation



CNN: Alternate:

Conv + ReLU + pooling

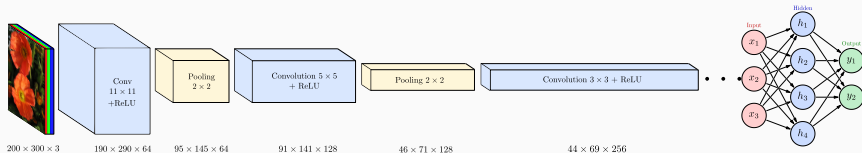
End of network:

Plug a standard neural network:

Fully connected hidden layers

(linear) + ReLU

All concepts together with tensor representation



Full network:

CNN: Alternate:
Conv + ReLU + pooling

End of network:

Plug a standard neural network:

Fully connected hidden layers
(linear) + ReLU

- **CNN:** Extract features specific to spatial data
- **Fully connected part:** Use CNN features for specific regression/classification task
- **Training:** Learn regression/classification and feature extraction **jointly**

Questions?

Slides from Charles Deledalle

Sources, images courtesy and acknowledgment

K. Chatfield

P. Gallinari,

C. Hazırbaş

A. Horodniceanu

Y. LeCun

V. Lepetit

L. Masuch

A. Ng

M. Ranzato



Go through the PyTorch tutorial:

“Deep Learning with PyTorch: A 60 Minute Blitz”

https://pytorch.org/tutorials/beginner/deep_learning_60min_blitz.html

Each part is a notebook with a “Run in Google Colab” button.

After that:

- ① Reproduce the results of the previous session using PyTorch: Define and train a neural network for multiclass logistic regression (linear+cross-entropy loss) on the toy 2D datasets.
- ② Add a first hidden layer with ReLU and train this new model.
- ③ Observe that the performances are better than with a random hidden layer.