

Lecture 23

Convolutional neural networks (CNN)

NN for images

Continuous convolution:

$$[f \times g](\vec{z}) = \int d\vec{u} f(\vec{u}) g(\vec{z} - \vec{u})$$

Discrete convolution:

1D example

if $f(\vec{u})$ is defined at $\{-L, \dots, 0, \dots, L\}$:

$f(-L) = w_{-L}, \dots, f(L) = w_L$ filter/kernel

if $g(\vec{u})$ is defined at $\{-N, \dots, 0, \dots, N\}$:

$g(-N) = x_{-N}, \dots, g(N) = x_N$ input

Then $[w \times x](i) = w_{-L} x_{i+L} + \dots + w_0 x_i + \dots + w_L x_{i-L}$

(1) $i+L$ & $i-L$ must be within the $[-N, N]$ range

Similarly, one can define a cross-correlation:

$$[w \otimes x](i) = w_{-L} x_{i-L} + \dots + w_{-1} x_{i-1} + \dots + w_L x_{i+L}$$

(2)

(1) & (2) are the same if $w_{-k} = w_k$.

One can restrict Eq. (2) to non-negative indices:

$$[w \otimes x](i) = \sum_{u=0}^{L-1} w_u x_{i+u} \quad \text{filter of length } L$$

Ex.

$$\begin{array}{|c|c|c|c|c|c|c|} \hline 0 & 1 & 2 & 3 & 4 & 5 & 6 \\ \hline \end{array} \quad \text{input} \quad \times \quad \begin{array}{|c|c|} \hline 1 & 2 \\ \hline \end{array} \quad \text{kernel} = \begin{array}{|c|c|c|c|c|c|} \hline 2 & 5 & 8 & 11 & 14 & 17 \\ \hline \end{array} \quad \text{output}$$

Similarly, in 2D

$$[w \otimes x](i, j) = \sum_{u=0}^{H-1} \sum_{v=0}^{W-1} w_{u,v} x_{i+u, j+v}$$

2D $H \times W$ filter

Ex.

0	1	2
3	4	5
6	7	8

input

0	1
2	3

kernel
(feature)

=

19	25
37	43

output
(feature map)

In matrix notation,

$$\begin{pmatrix} w_1 & w_2 \\ w_3 & w_4 \end{pmatrix} \times \begin{pmatrix} x_1 & x_2 & x_3 \\ x_4 & x_5 & x_6 \\ x_7 & x_8 & x_9 \end{pmatrix} = \begin{pmatrix} w_1 x_1 + w_2 x_2 + w_3 x_4 + w_4 x_5 \dots \\ w_1 x_4 + w_2 x_5 + w_3 x_7 + w_4 x_8 \dots \end{pmatrix} \quad (11)$$

$$\textcircled{=} \begin{pmatrix} w_1 & w_2 & 0 & w_3 & w_4 & 0 & 0 & 0 & 0 \\ 0 & w_1 & w_2 & 0 & w_3 & w_4 & 0 & 0 & 0 \\ 0 & 0 & 0 & w_1 & w_2 & 0 & w_3 & w_4 & 0 \\ 0 & 0 & 0 & 0 & w_1 & w_2 & 0 & w_3 & w_4 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_9 \end{pmatrix}$$

↑
sparse weight matrices,
shared weights

↑
2x2
flattened

In general, $f_h \times f_w$ filter applied to an image of size $x_h \times x_w$ produces an output of size $(x_h - f_h + 1) \times (x_w - f_w + 1)$: use zero padding of the image (borders of zeros) so that the output retains the $x_h \times x_w$ size.

With symmetric padding of size p_h, p_w we have:

$$x_h \times x_w \xrightarrow{\text{padding}} (x_h + 2p_h) \times (x_w + 2p_w) \xrightarrow{\text{convolution}}$$

$$\rightarrow (x_h + 2p_h - f_h + 1) \times (x_w + 2p_w - f_w + 1) \rightarrow$$

$$\rightarrow x_h \times x_w$$

$$\begin{cases} 2p_h = f_h - 1 \\ 2p_w = f_w - 1 \end{cases}$$

Strided convolution: apply filters with step sizes > 1 or, equivalently, ~~step sizes~~ keep only $1/(s_h s_w)$ convolutions

$\uparrow \uparrow$
 step sizes

Output size:

$$\left\lfloor \frac{x_h + 2p_h - f_h + s_h}{s_h} \right\rfloor \times \left\lfloor \frac{x_w + 2p_w - f_w + s_w}{s_w} \right\rfloor$$

smaller than input

Multiple channels:

For ex., RGB channels in color images

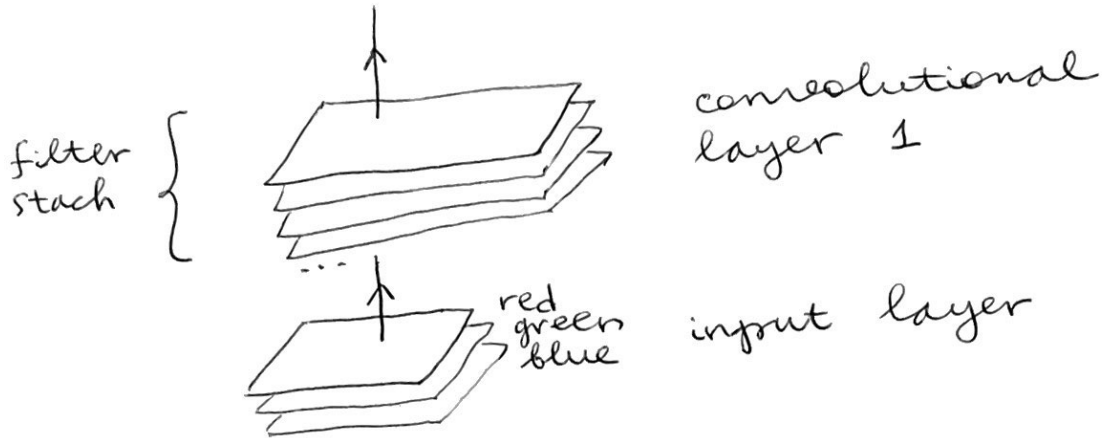
$$\text{Then } z_{i,j} = \underset{\substack{\uparrow \\ \text{bias}}}{b} + \sum_{c=0}^{C-1} \sum_{u=0}^{H-1} \sum_{v=0}^{W-1} w_{u,v,c} x_{si+u, sj+v, c}$$

(*) $s = s_h = s_w \leftarrow \text{stride}$

Eq. (*) provides a single filter ; with multiple filters , we have :

$$z_{i,j,d} = b_d + \sum_c \sum_u \sum_v w_{u,v,c,d} x_{s_i+u, s_j+v, c}$$

(**) filter index : $d=0, \dots, D-1$



pointwise convolution:

no "averaging" across locations

$$z_{i,j,d} = b_d + \sum_{c=0}^{G-1} x_{i,j,c} w_{0,0,c,d}$$

Same spatial dimensions : $H \times W$,
but the number of channels
goes from G to D .

pooling layers: convolutions preserve
spatial information (at reduced resolution)

Idea: "pool" incoming values by taking
a max (max pooling) or an average
(ave pooling).

Ex.

0	1	2
3	4	5
6	7	8

$\Rightarrow$
2x2
max
pooling

4	5
7	8

[pooling is usually done ~~on~~ separately]
on each feature channel.

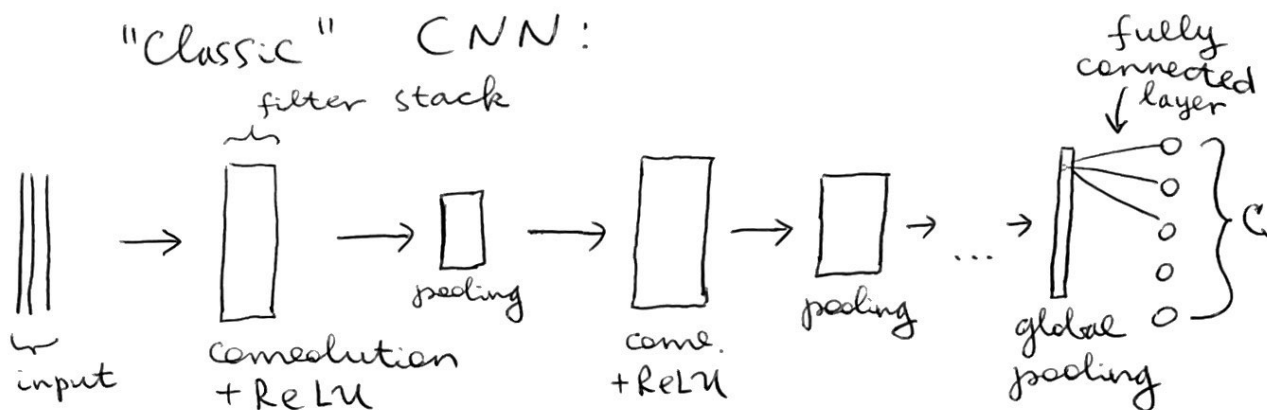
at the end of the CNN, one can do
global average pooling: (filter size = $H \times W$)

$$H \times W \times D \rightarrow 1 \times 1 \times D = D$$

$\uparrow$
#channels

This D-dim vector can be fed into a
softmax output layer with C nodes.

"Classic" CNN:



Deep CNNs may suffer from vanishing
or exploding grads.

Idea: add normalization layers to the
CNN architecture.

Batch normalization (BN):

ensures the distribution of activations has $\text{mean} = 0$ & $\text{var} = 1$ when averaged over a minibatch.

← vector of activations in a layer

$$\vec{\mu}_B = \frac{1}{|B|} \sum_{\vec{z} \in B} \vec{z}, \quad \sigma_B^2 = \frac{1}{|B|} \sum_{\vec{z} \in B} (\vec{z} - \vec{\mu}_B)^2$$

B is the minibatch of size $|B|$ containing example n (datapoint)

$$\vec{z}_n' = \frac{\vec{z}_n - \vec{\mu}_B}{\sqrt{\sigma_B^2 + \epsilon}} \quad \text{standardized}$$

$\epsilon > 0 = \text{const}$

$$\vec{z}_n'' = \vec{\gamma} \times \vec{z}_n' + \vec{\beta} \quad \text{shifted \& scaled}$$

↑
element-wise

[output of the BN layer]

$\vec{\gamma}$ & $\vec{\beta}$ are fitted prms.

In the context of CNN, $\vec{\mu}_B$ is averaged across both spatial locations and examples $\Rightarrow \dim\{\vec{\mu}_B\} = \# \text{ channels}$.

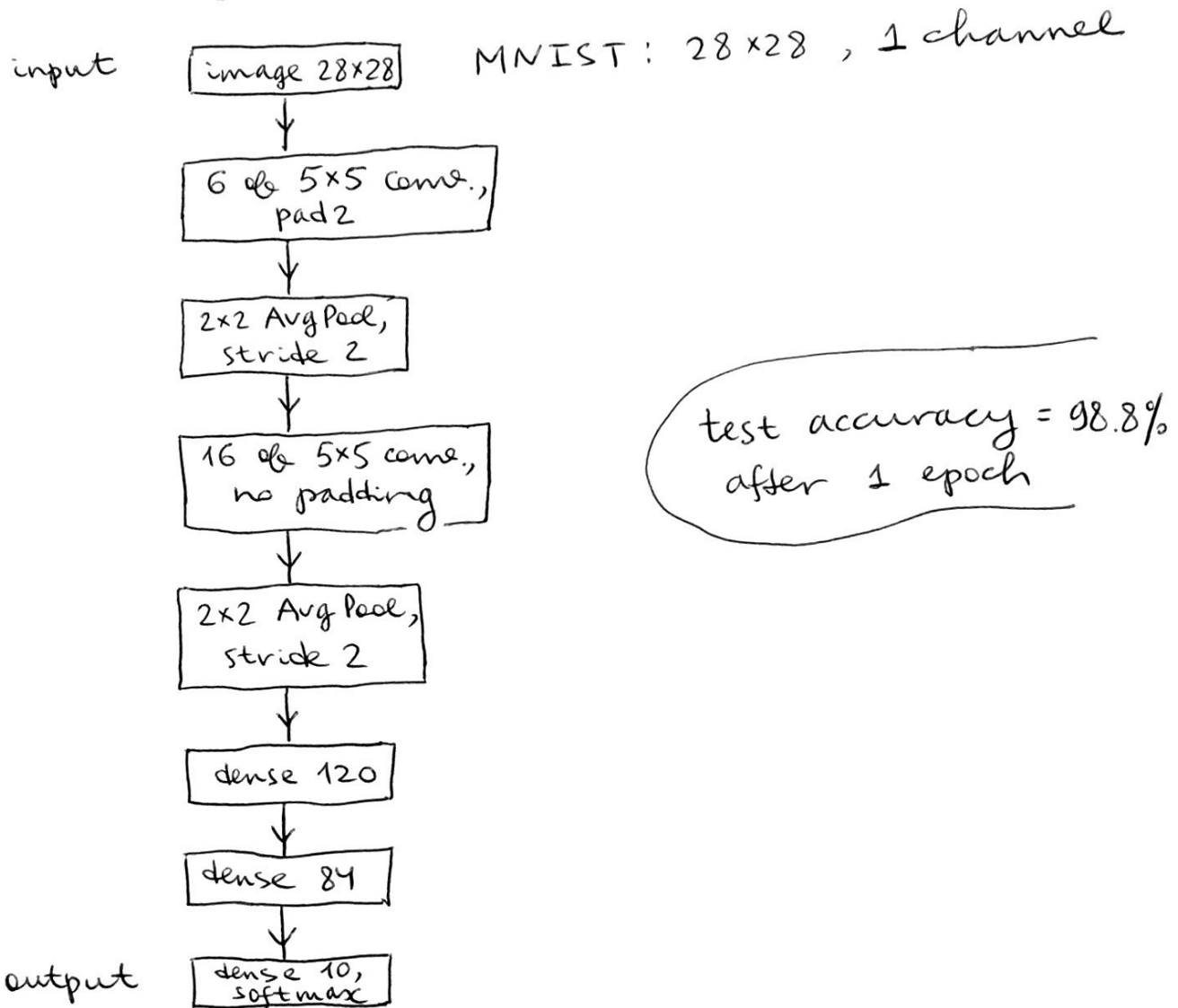
At test time, first compute $\vec{\mu}_\ell$ & σ_ℓ^2
on the entire training set
and then 'freeze' these prms for testing.

↑ layer index

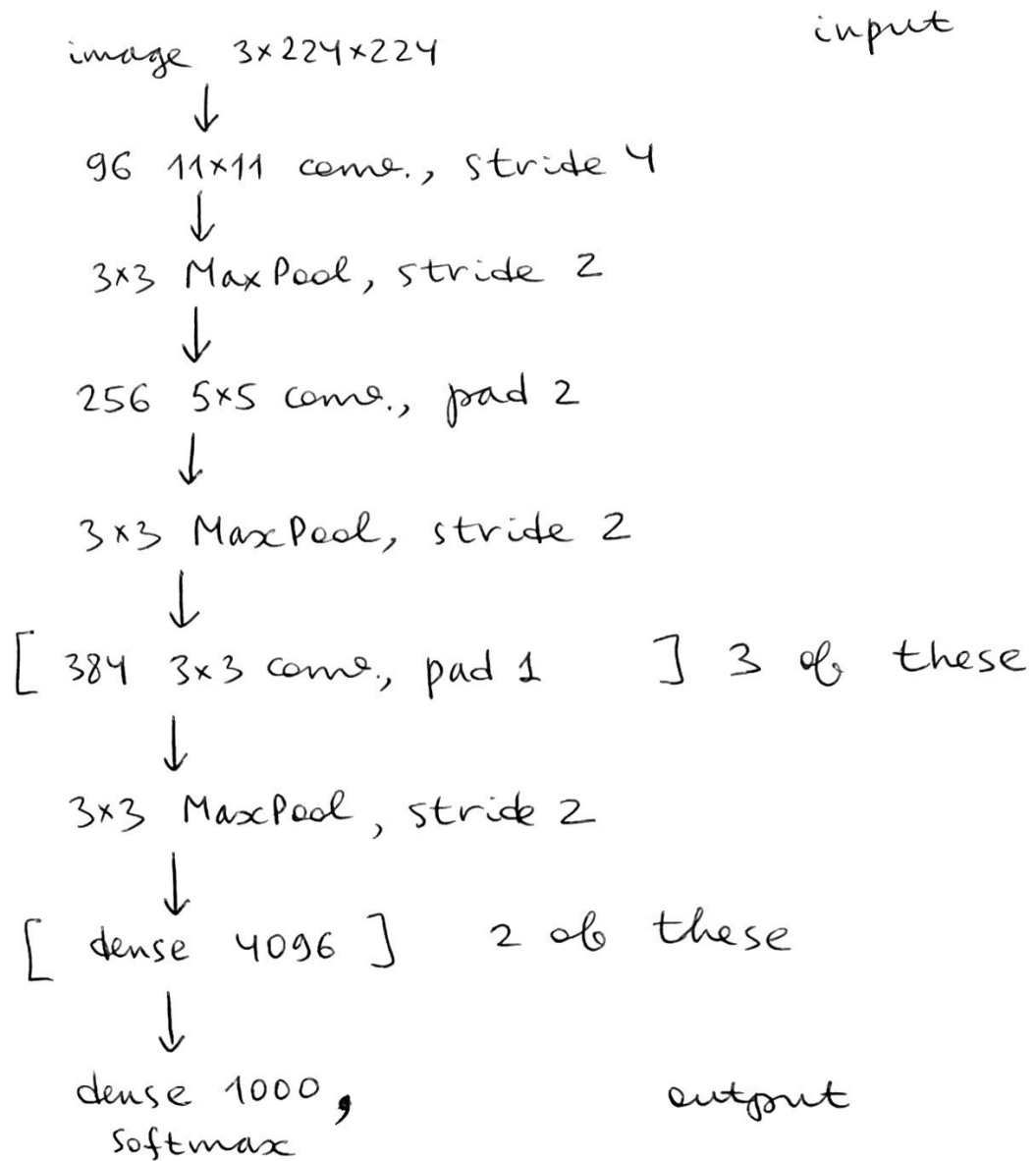
Adaptive grad clipping: instead of BN layers, use grad clipping with $c = \text{fitting prm.}$

CNN architectures

LeNet5



AlexNet



~60 M prms (mostly due to dense layers)

ImageNet : ~1M labeled examples

Top 5 error rate was reduced
from ~26% (SOTA) to ~15%.