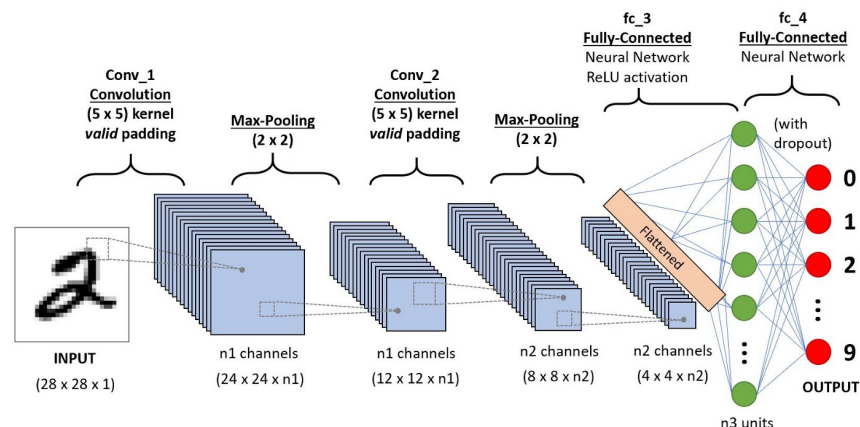


1 Definitions

Convolution Neural Network: A type of neural network that is particularly well-suited for image and video processing tasks. CNNs consist of the following building blocks:

1. **Convolutional Layer:** Involves sliding a small window of fixed size across the input data and computing the dot product between the kernel and the overlapping portion of the input at each position. The result of each dot product is then summed to produce a single output value, which is stored in a new output.
 - (a) **Kernel/Filter:** A small multidimensional array of weights that is used to scan over the input data during the convolution operation
 - (b) **Stride:** The number of pixels or steps by which a filter is moved across the layer input data
 - (c) **Channel:** An additional dimension within a layer of the network that represents a specific feature or attribute of the data. For example, an RGB image has three channels that correspond to the red, green, and blue color channels. The output of a convolutional layer usually has many different channels to allow the network to learn to detect different patterns in each channel. For example, the output of the first convolutional layer may learn to detect edges at different angles in each channel.
 - (d) **Padding:** Extra space added to the outsides of the layer input (filled with zeros or a copy of the data on the border) to make the output dimensions have the desired size.
2. **Pooling Layer:** Used to reduce the spatial dimensions of the data. Two popular types of pooling are max pooling and average pooling.
 - (a) **Max pooling:** Similar to a convolution, but in this case, the maximum value at each kernel location is selected. Analogous to picking the brightest pixel in a grid of small regions of the image and tossing the rest.
 - (b) **Average pooling:** The average value of all the pixels at each kernel location is selected. Analogous to smoothing out an image while reducing its size.
3. **Fully-connected layer:** Typically a CNN will end with one or more fully-connected layers. These layers effectively perform classification or regression based on the features coming from the convolutional layers.



2 CNNs and not the kind on TV

2.1 Conceptual Questions

1. What are some benefits of CNNs over fully connected (also called dense or linear) layers?

CNNs are good for image-related machine learning tasks because they learn the kernels that do feature engineering. Additionally, CNNs, through the use of convolutions and pooling, take advantage of local spatial coherence. This refers to the tendency of neighboring pixels in an image to have similar values or characteristics. CNNs are also parameter efficient.

2. How are the weights in a CNN updated during training?

Backpropagation, the same as fully connected layers.

3. What are the parameters of a convolutional layer?

The values in the kernels (a.k.a. weights or filters) and the bias term per output channel. The kernel sizes are num_input_channels by kernel_width by kernel_height for each output channel.

4. What are the hyperparameters of a convolutional layer?

Stride size, padding size, filter size.

5. Can convolutions only be applied to 2-D shapes like images?

Nope. Convolutions are very common on 1-D signals like sound or a stock market time series. You can also combine 2-D space with time and do convolutions for video. Or even higher dimensional data like 4-D cardiac MRI data.

2.2 Shapes and Parameters

Torch the Triceratops is working on his 07-280 programming homework and is looking into the MS-COCO dataset. His goal is to classify images that belong to one of ten possible classes (i.e. [cat, dog, bird, turtle, ..., horse]). The images come in RGB format (one channel for each color) and are downsampled to dimension 128×128 .

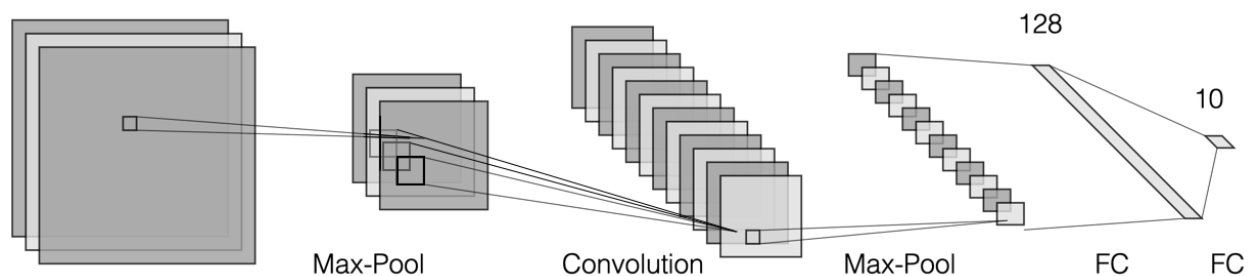
The following illustrates one such image from the MS-COCO dataset.



Torch constructs a convolutional neural network that has the following structure: the input is first max-pooled (not common) with a 2×2 filter with stride 2. The results are then sent to a convolutional layer that uses a 17×17 filter (uncommonly large) with stride 1 and 12 output channels. Those values are then passed through a max-pooling layer with a 3×3 filter with stride 3. The result is then flattened and passed through a fully-connected layer (with ReLU activation, not shown) with 128 hidden units, followed by a fully connected layer (softmax activation, not shown) with 10 hidden units. The final 10 hidden units represent the categorical probability for each of the ten classes. With enough labeled data, Torch plans on using some optimizer like SGD to train this model using backpropagation.

Note: Assume we have bias terms in all neural network layers unless explicitly stated otherwise.

3@128x128



1. What is the shape of the output tensor after max pooling?

With a stride of two the size is cut in half in each direction. Unlike convolution, pooling is applied independently per channel.

$3 \times 64 \times 64$

2. What is the shape of the output tensor after the convolutional layer?

Thinking in just one dimension, we can fit a window of width 17 within the larger width of 64 and then shift it by one 47 more times before it hits the right edge.

$12 \times 48 \times 48$

3. What is the shape of the output tensor after the second max pooling layer?

$$12 \times 16 \times 16$$

4. What is the general equation for the shape of the output tensor if a filter of size $F \times F$ with stride S and padding P is applied to a tensor with width W and height H ?

$$\text{output width} = \frac{W+2P-F}{S} + 1$$

$$\text{output height} = \frac{H+2P-F}{S} + 1$$

5. How many parameters are in this network for the first max pooling layer?

Zero. Max-pooling layers do not have any parameters, since a fixed operation is performed on the input.

6. How many parameters are in this network for the convolutional layer?

$$12 \times (3 \times 17 \times 17) + 12 = 10416$$

7. How many parameters are in this network for both fully connected layers?

$$(3072 \times 128 + 128) + (128 \times 10 + 10) = 394634$$

8. From these parameter calculations, what can you say about convolutional layers and fully connected layers in terms of parameter efficiency (the ratio between the number of parameters from some layer type and the total number of parameters)? Why do you think this is the case?

Total number of parameters = $10416 + 394634 = 405050$

Percent from convolutional components = $\frac{10416}{405050} = 2.57\%$

Percent from fully connected components = $\frac{394634}{405050} = 97.43\%$

Convolutional layers are much more parameter efficient, mainly because we are reusing the convolutional filter repeatedly for each convolutional layer (we only need to train one kernel per channel per layer). In comparison, the fully connected layer requires all nodes between two layers to be fully connected.

3 Visualizing CNNs

Explore the interactive visualization of a neural network architecture trained to recognize handwritten digits using the following link: https://adamharley.com/nn_vis/cnn/2d.html.

Take some time to experiment with the tool, observing the behavior of each layer and its effect on the feature maps. Then, respond to the questions below based on your observations.

This visualization is based on the work by A. W. Harley, "An Interactive Node-Link Visualization of Convolutional Neural Networks," in *ISVC*, pages 867-877, 2015.

1. When we hover over the pixels in the feature maps, we see a bunch of colored lines. What do these lines represent? Do you notice anything interesting as we move around the image?

Each line represents the weight that is learned by the model. For one of the feature maps, as we move across the image, the weights remain fixed. This is known as parameter sharing and helps reduce the computational complexity of a CNN.

2. How many filters does this network use in the first convolutional layer? Can we determine the size of each filter and the stride?

There are six 5x5 filters of stride 1 in the first convolutional layer.

3. Are there pooling layers in this network? If so, how many and what are their hyperparameters? How do these layers affect the size of the image? What is the purpose of the pooling layer?

There are 2 downsampling layers, one after each of the convolutional layers. Each one does 2x2 max pooling with stride 2. With these values, the dimension of the feature map is halved. By reducing the size of feature maps, pooling decreases the number of parameters and computational complexity, making the network more efficient. It can also help control overfitting.

4. How many parameters does one neuron in the output layer have? (Hint: the last fully connected layer has 100 nodes)

101

For those of you who are curious about this tool and want to read more, here is a paper on the topic: https://adamharley.com/nn_vis/harley_vis_isvc15.pdf.