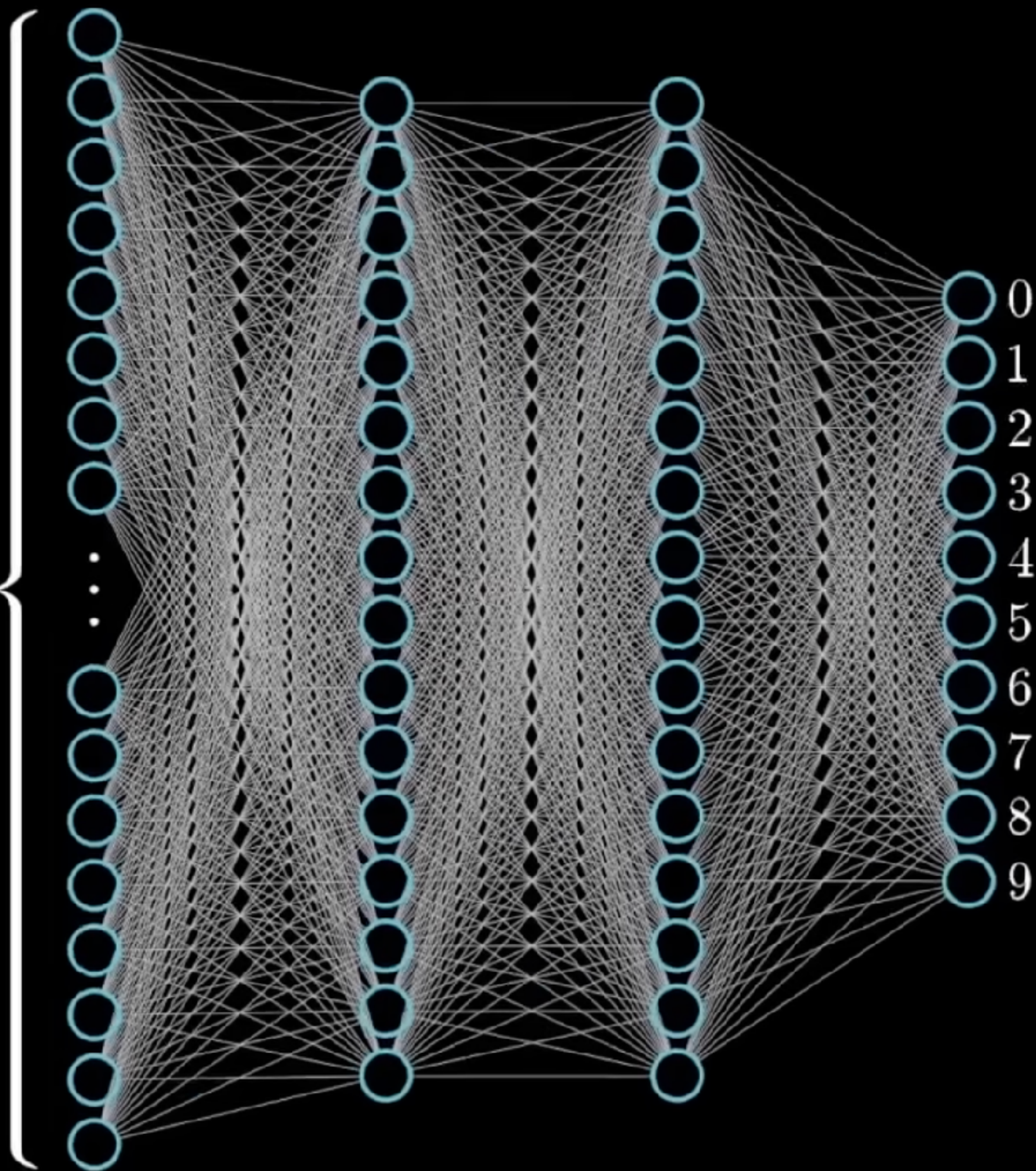


Reseau Neuronale

Reference:

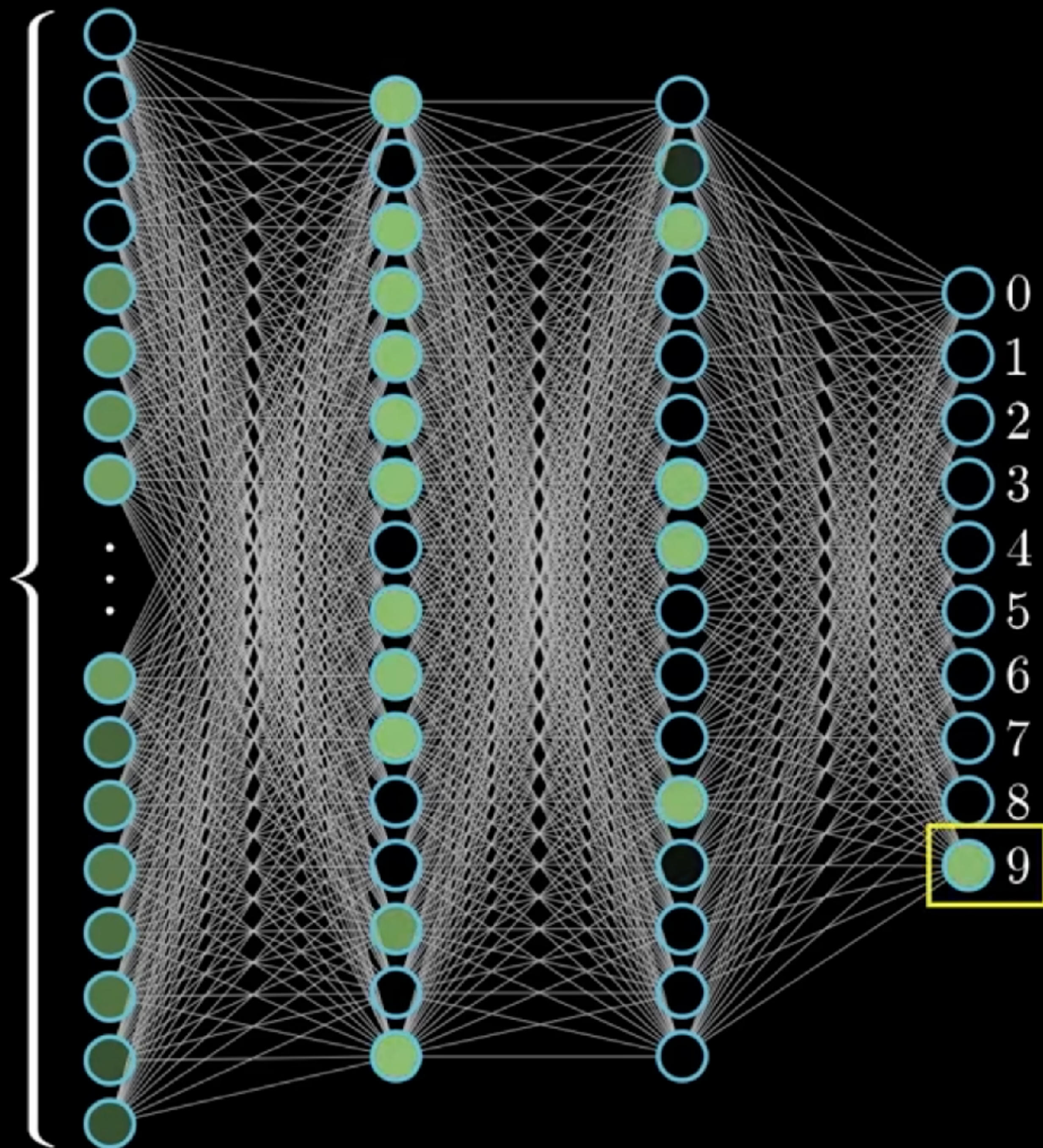
Youtube: 3Blue1Brown “But what is a neural network? | Deep learning chapter 1”

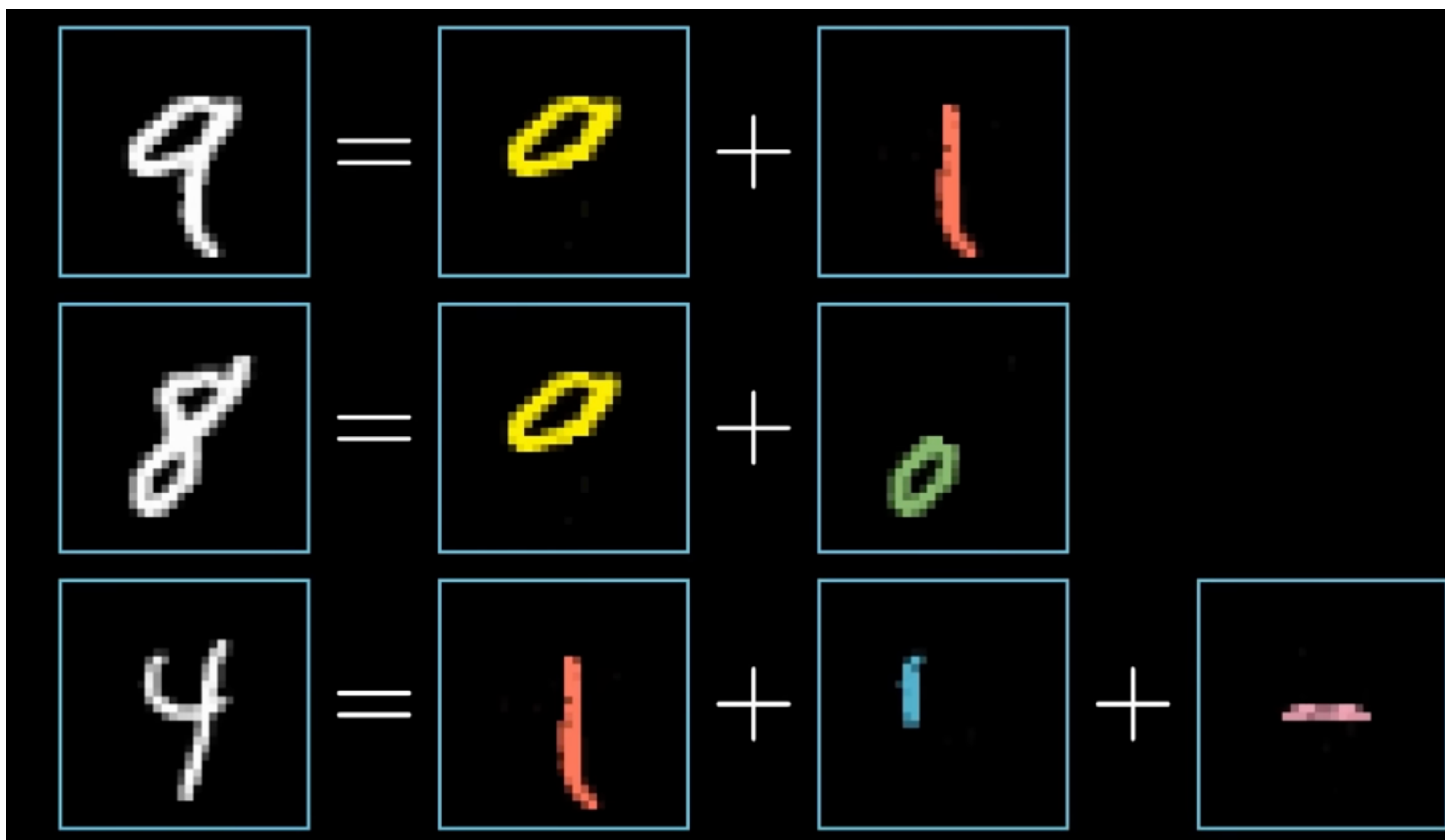
784





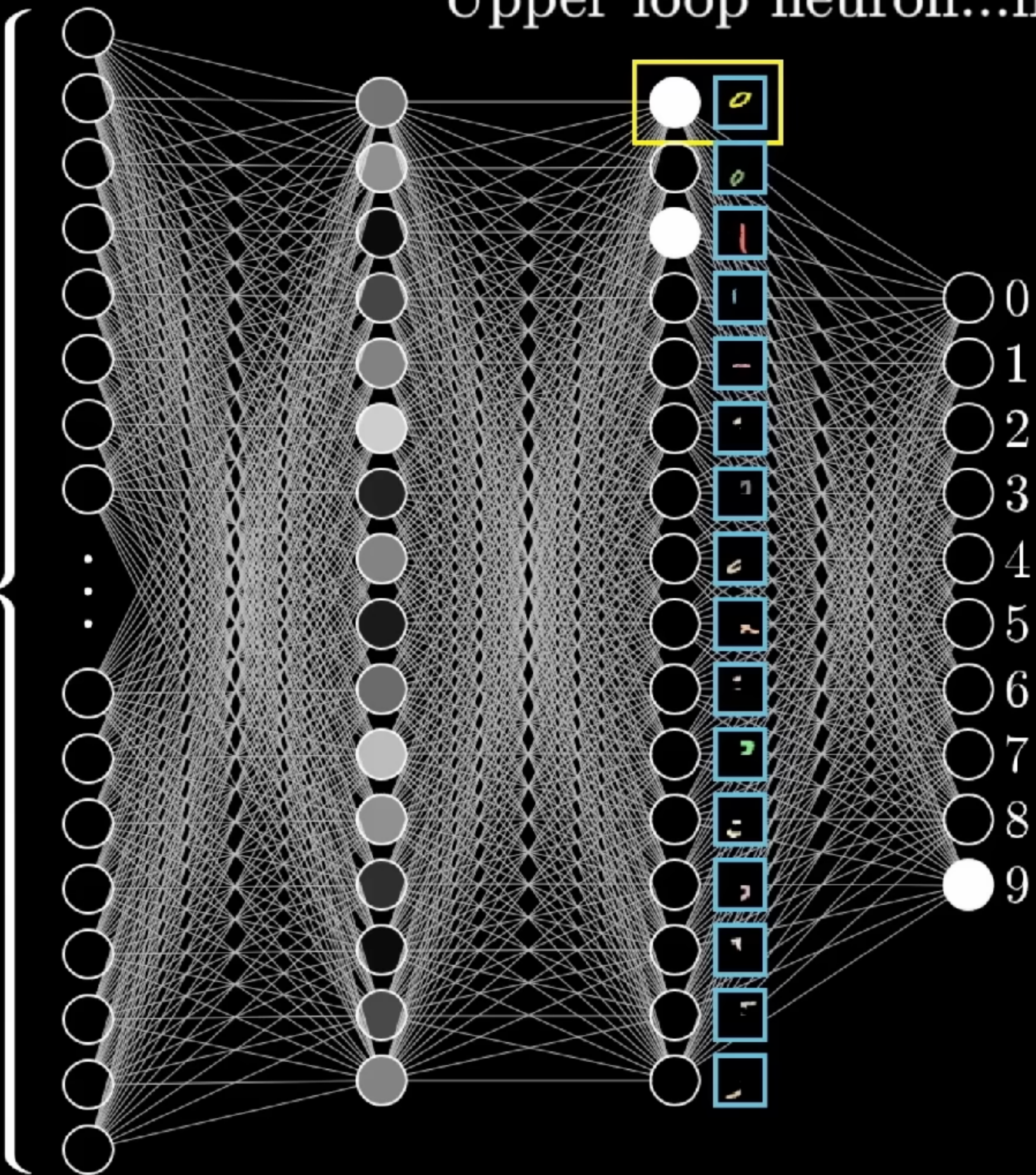
784

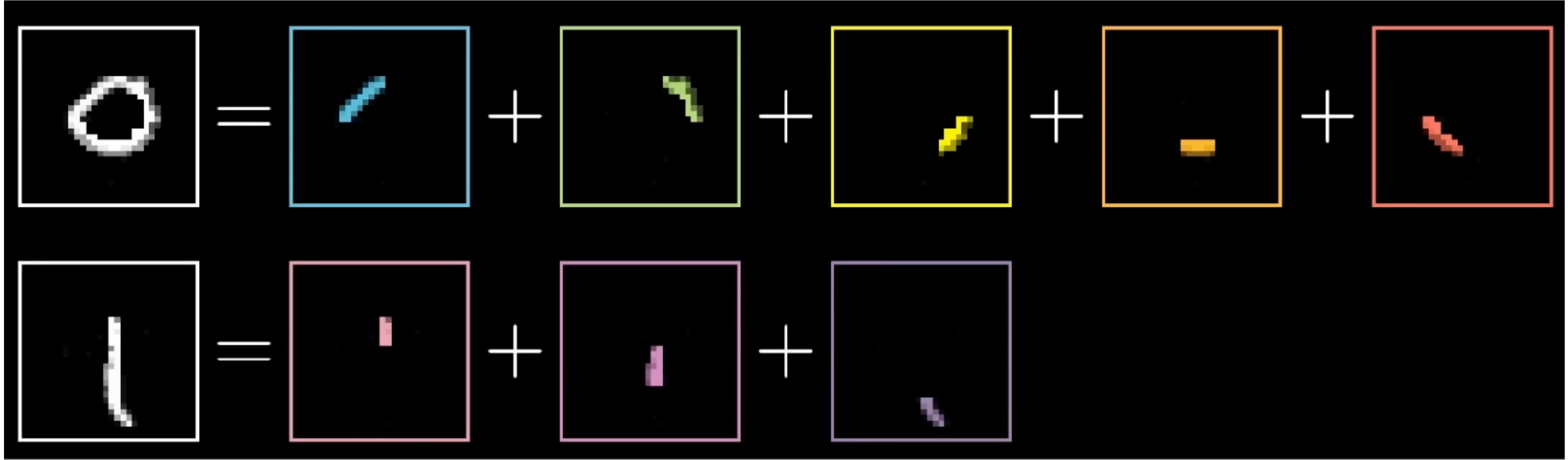


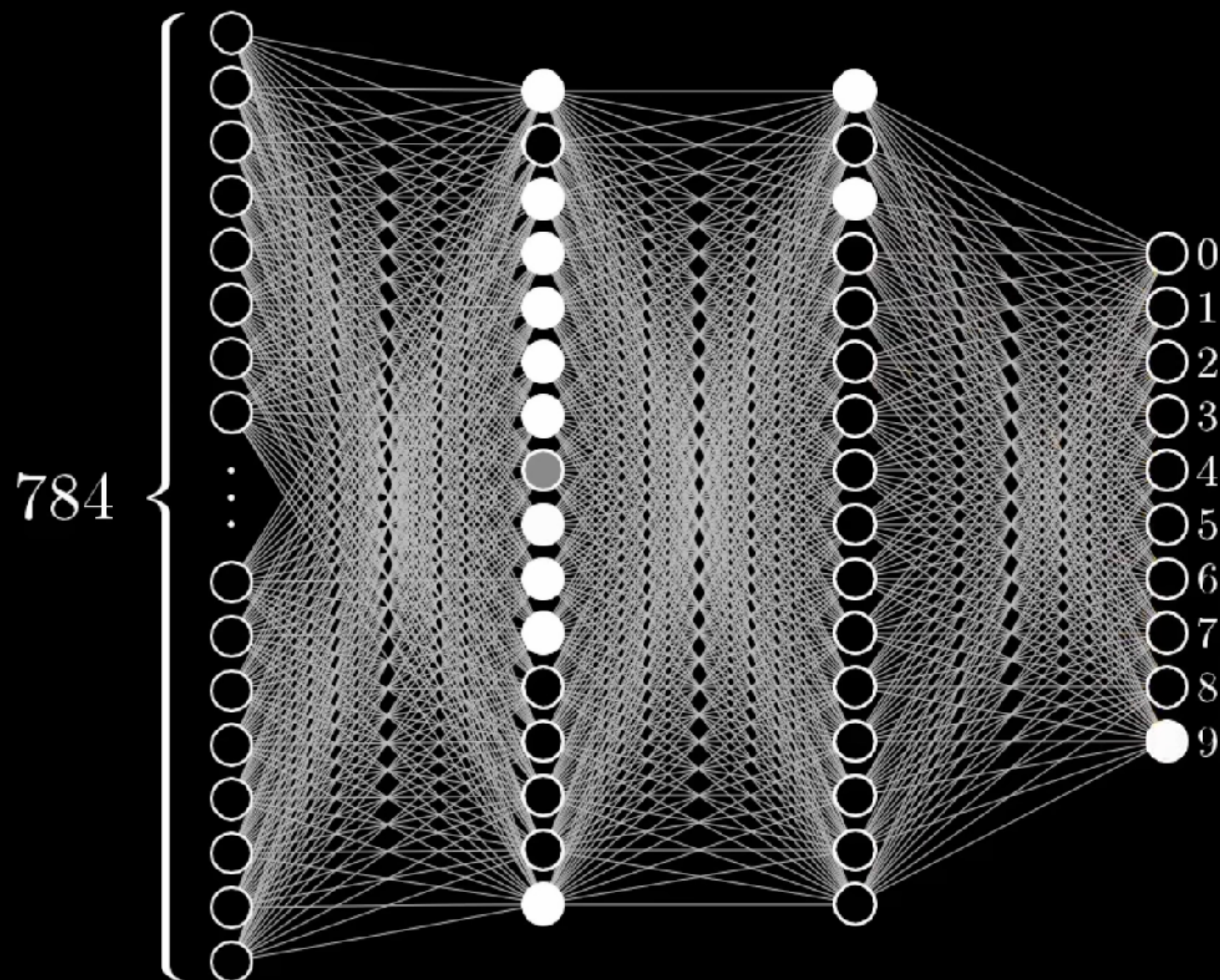




784



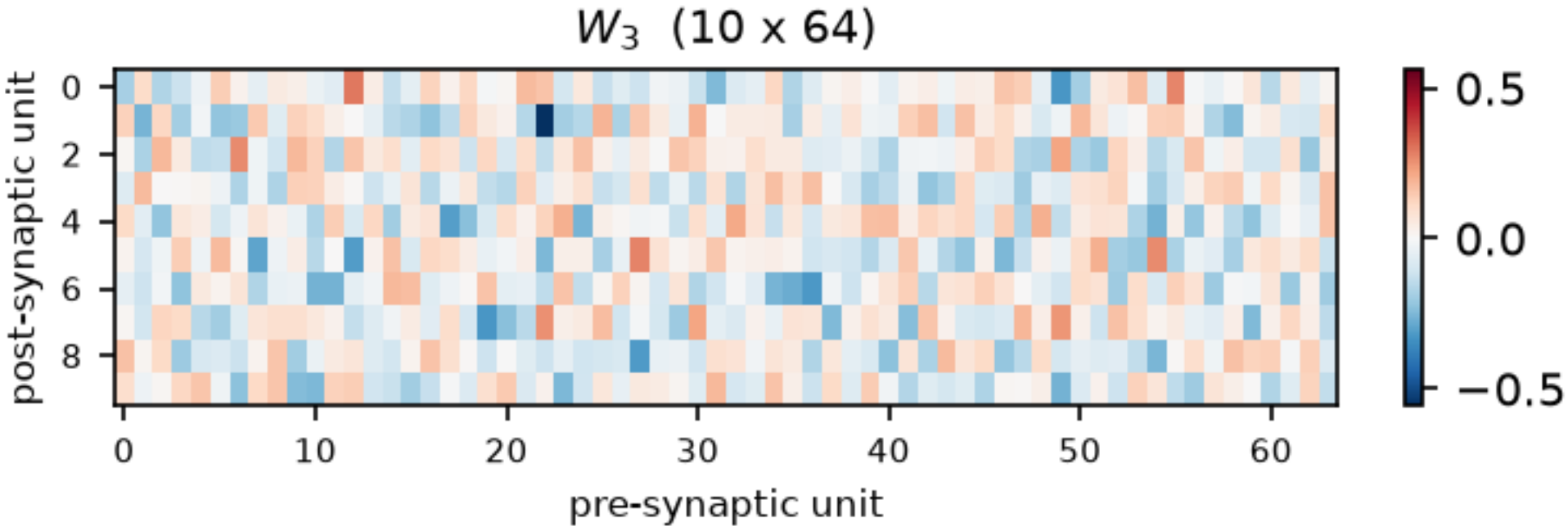
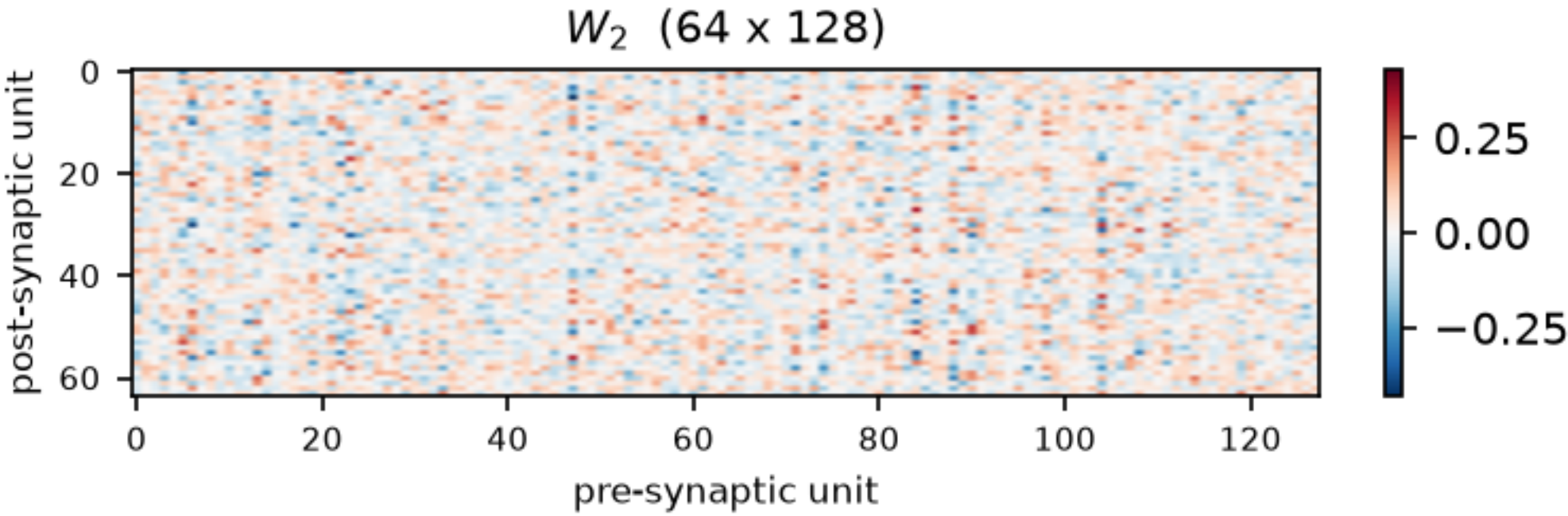
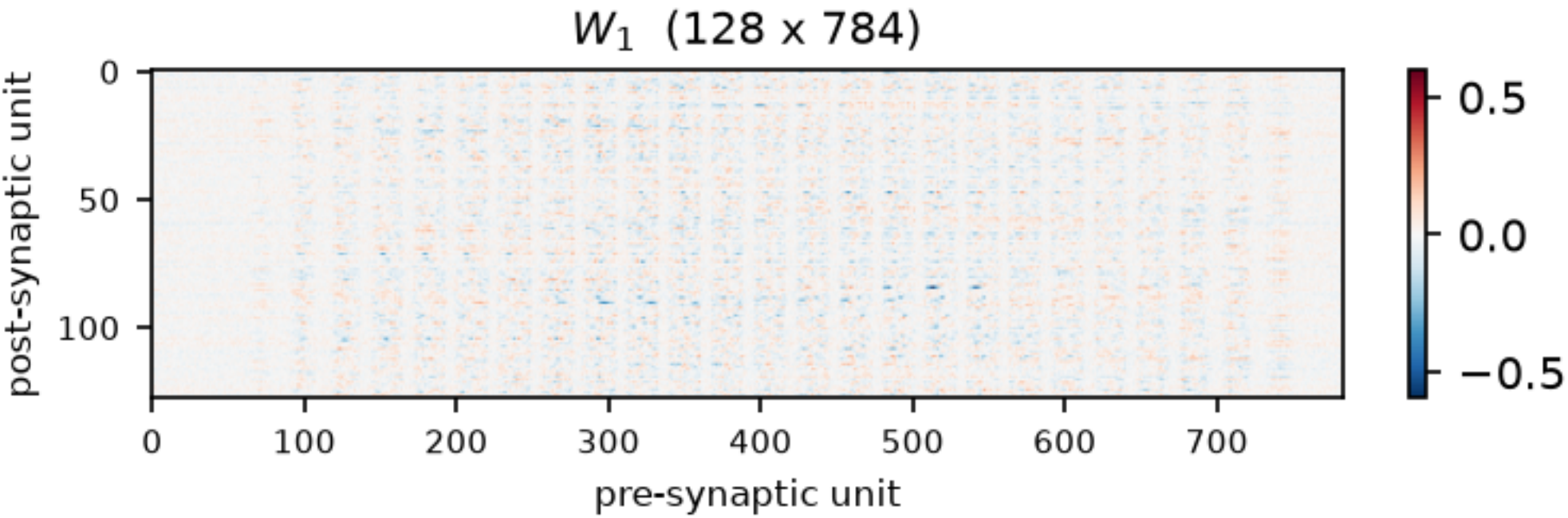
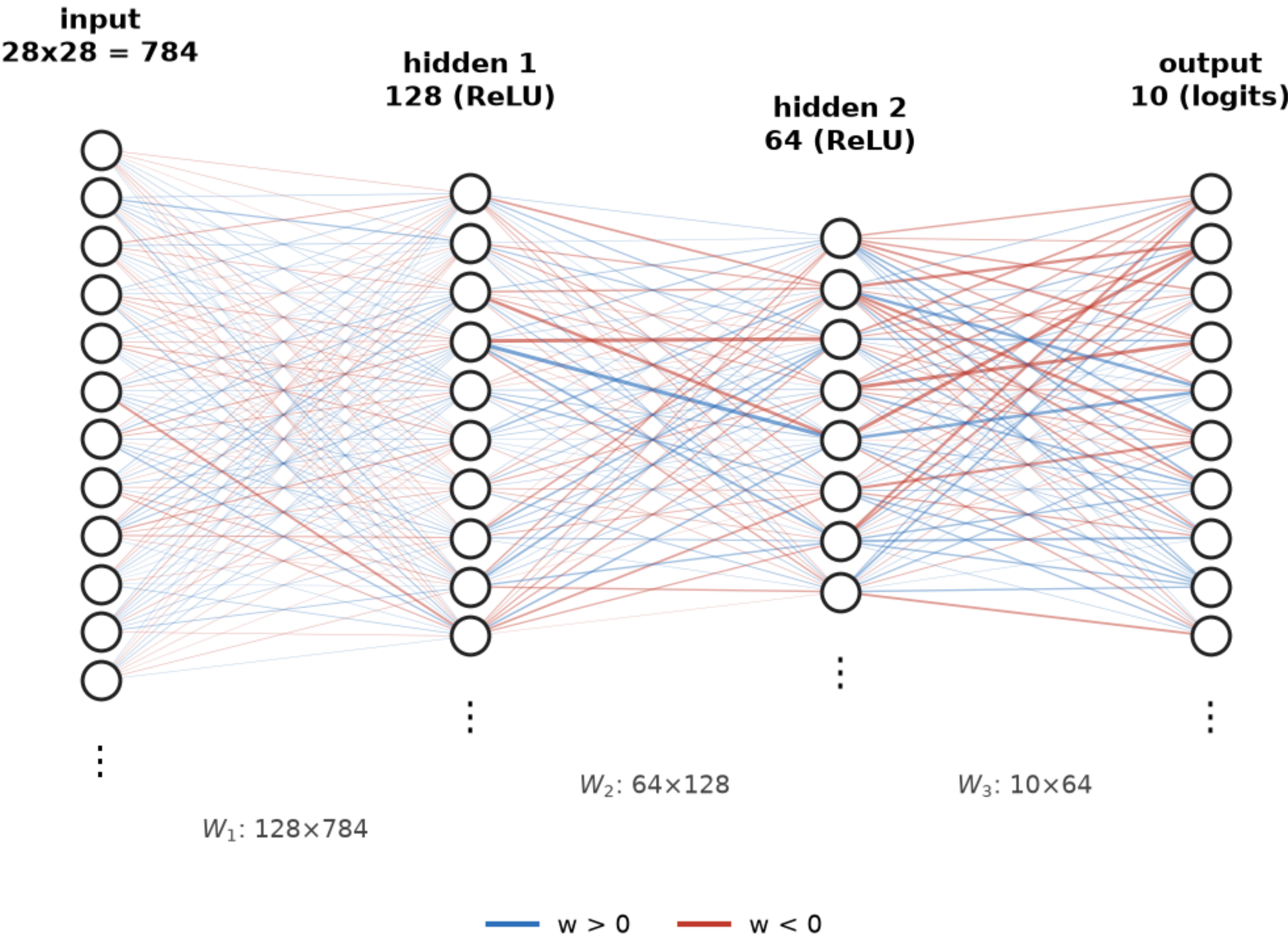




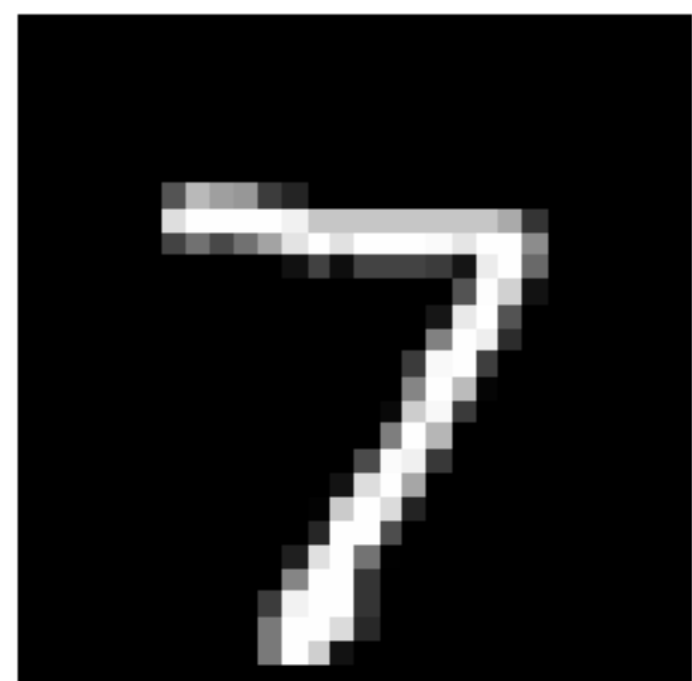


MLP for MNIST: 784 → 128 → 64 → 10

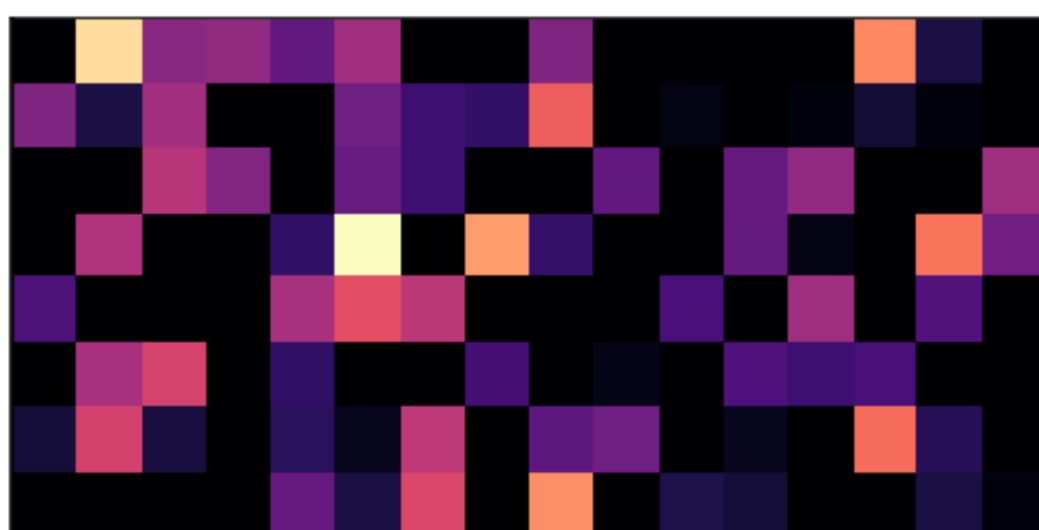
Network graph (units subsampled)



input (28x28 = 784)
true label: 7

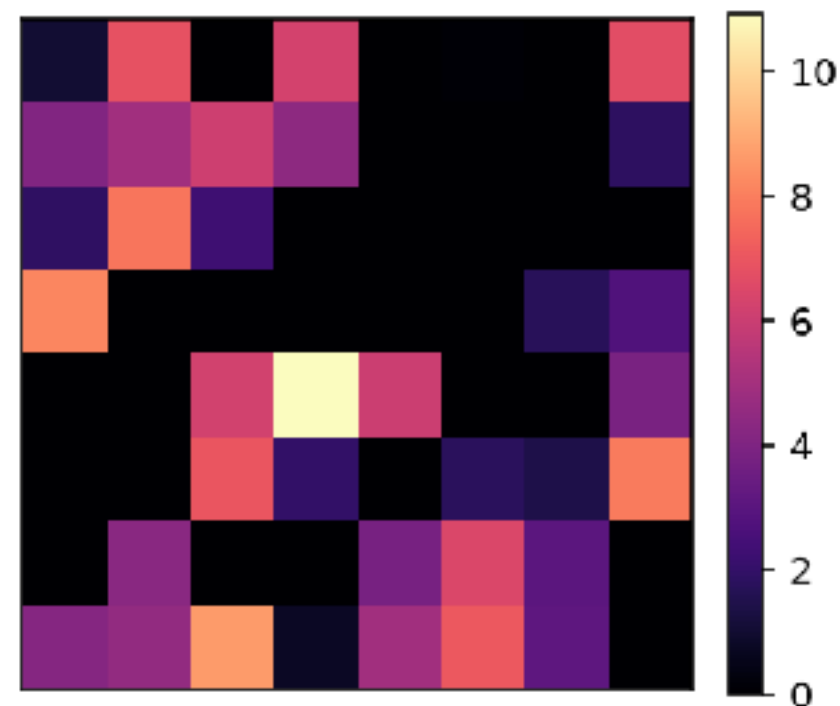


$h_1 = \text{ReLU}(W_1x + b_1)$ -- 128 units
active: 55%

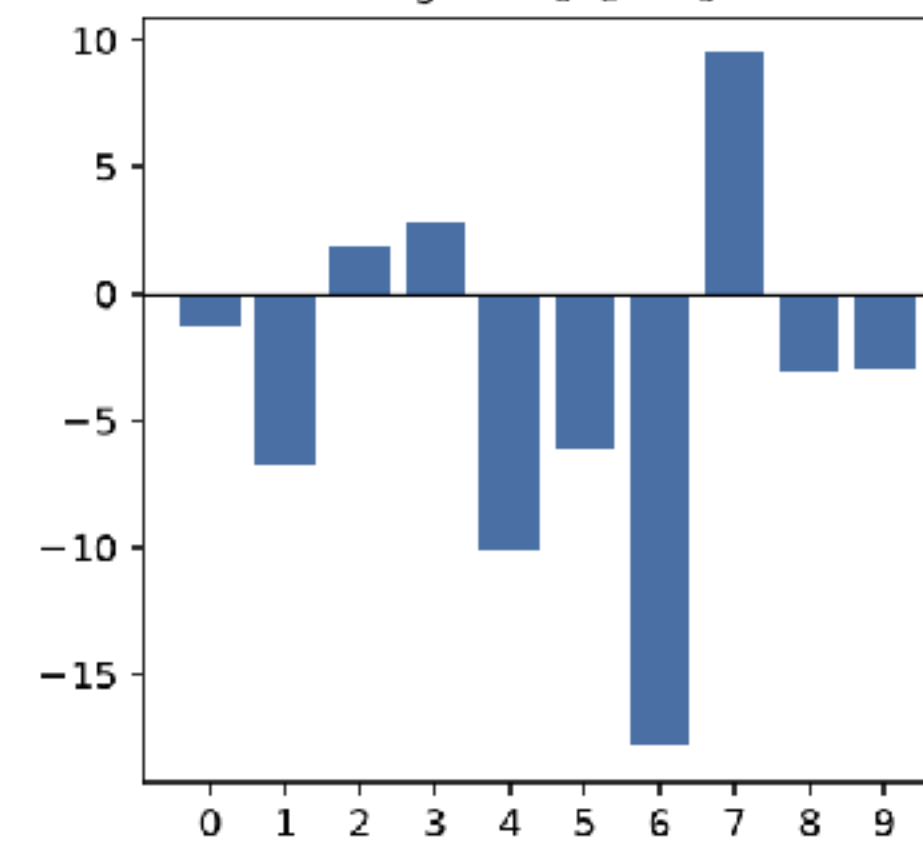


One digit through the network

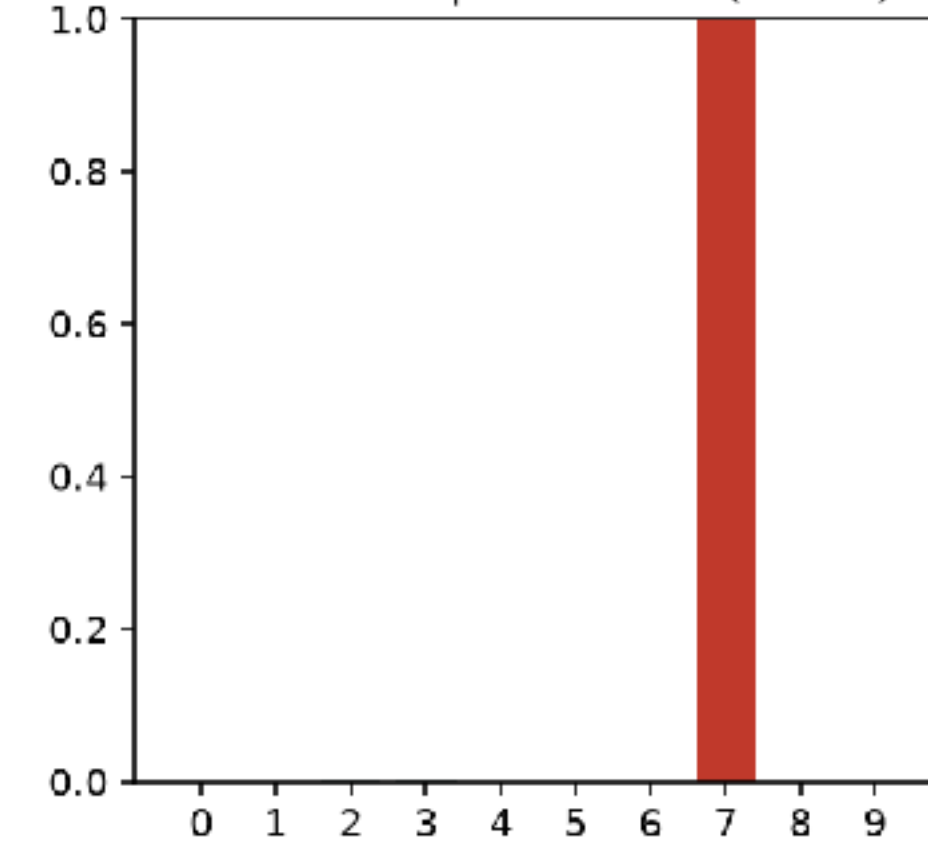
$h_2 = \text{ReLU}(W_2h_1 + b_2)$ -- 64 units
active: 56%



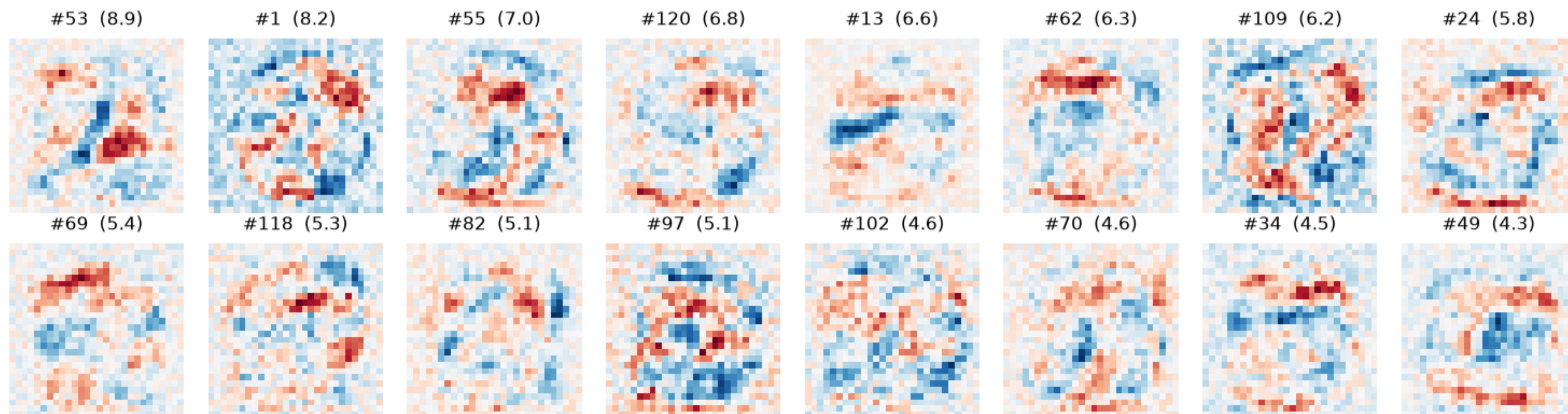
logits $W_3h_2 + b_3$



softmax -> prediction: 7 (99.8%)

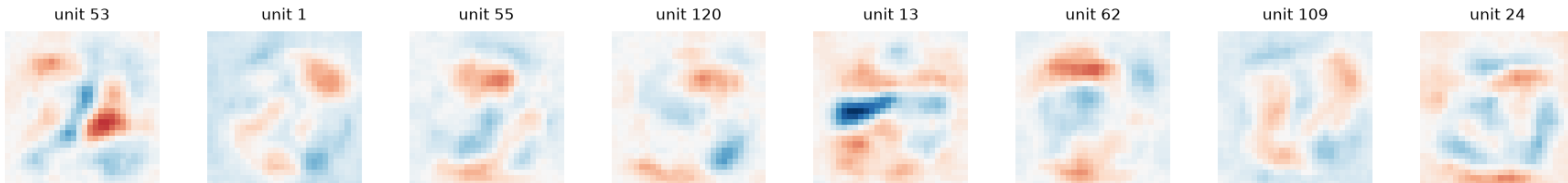


Rows of W_1 reshaped to 28x28 — the 16 hidden-1 units most activated by this digit
(blue = excitatory pixel, red = inhibitory; value in brackets = activation)

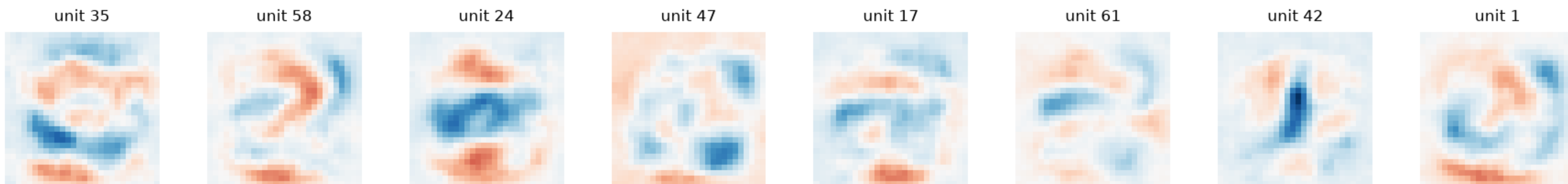


hidden 1 (128 units)

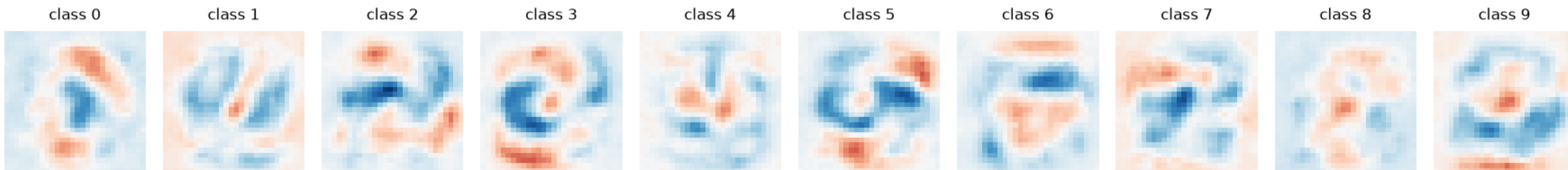
**Activation maximisation: input optimised to drive each unit
gradient ascent with L2 + total-variation penalties; the result need not look like a digit, and does not have to**



hidden 2 (64 units)



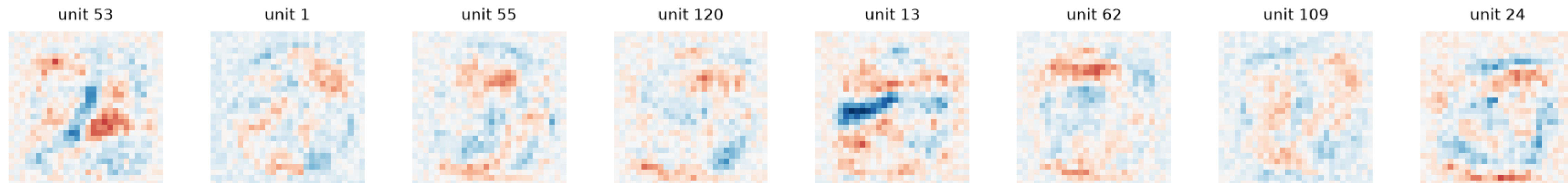
output (10 logits)



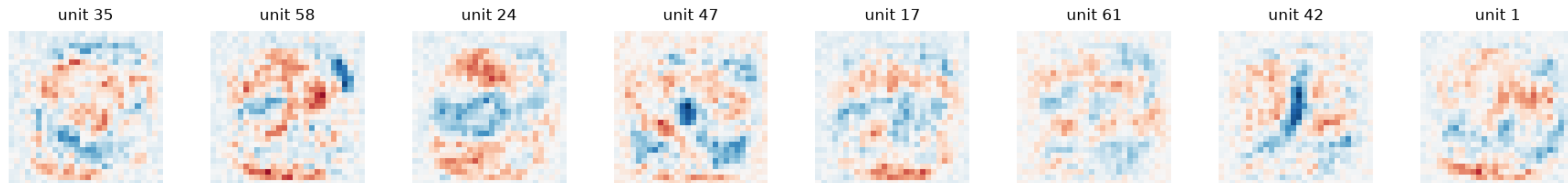
Linearised receptive fields: W_1 , W_2W_1 , $W_3W_2W_1$

blue = pixel pushes the unit up, red = pushes it down (ReLU gates set to identity -- exact only for a linear net)

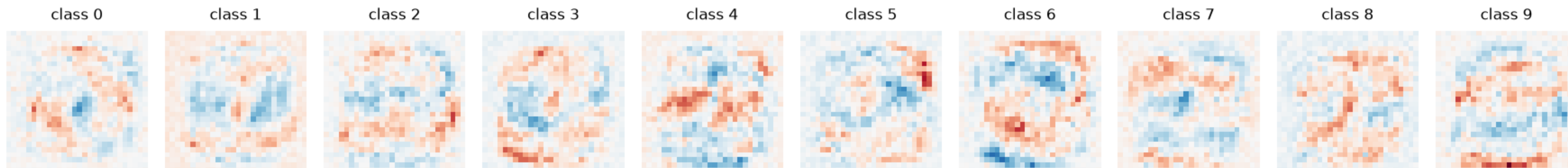
hidden 1 (128 units)



hidden 2 (64 units)

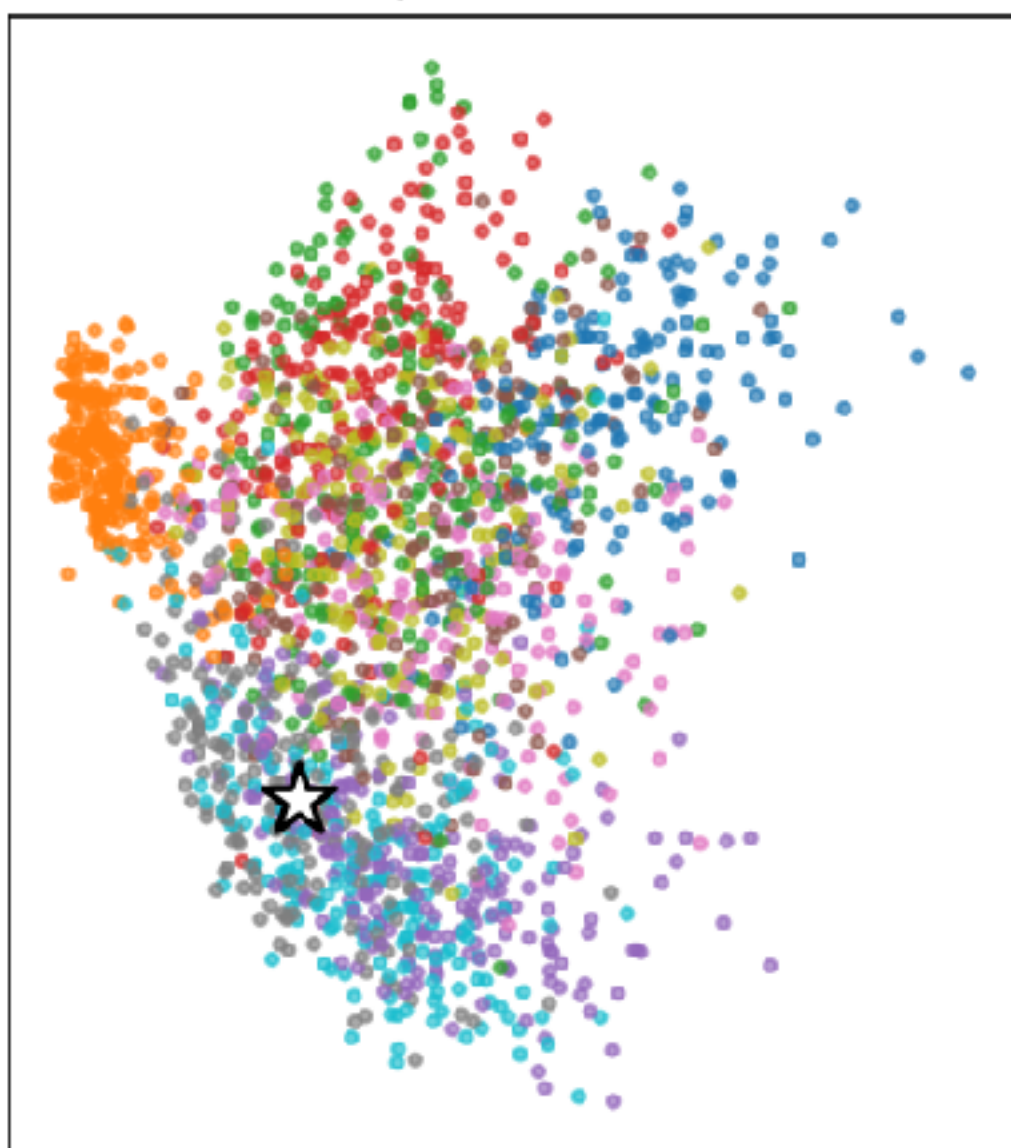


output (10 logits)

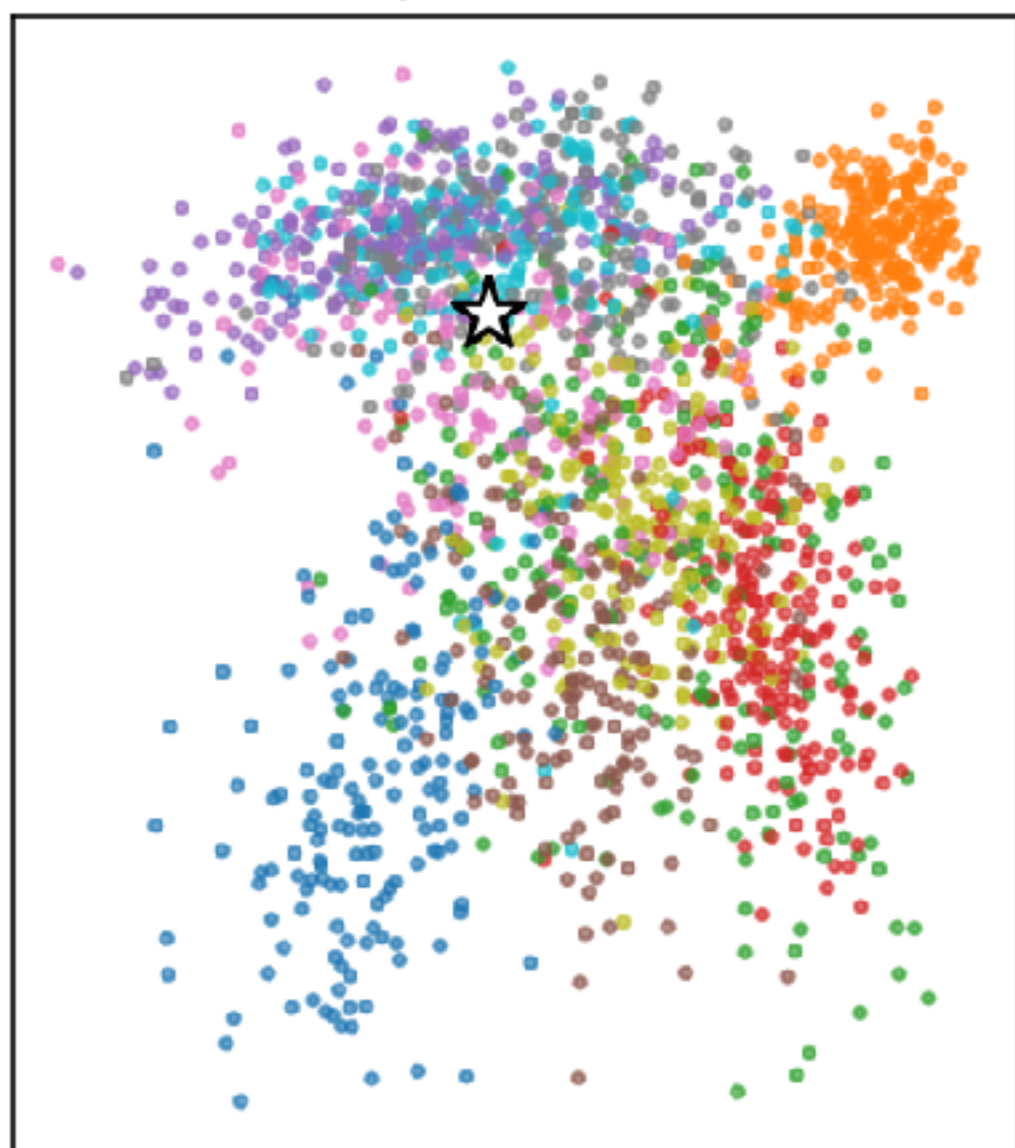


The same 2000 test digits, projected onto the top 2 PCs at each stage (star = the example digit)

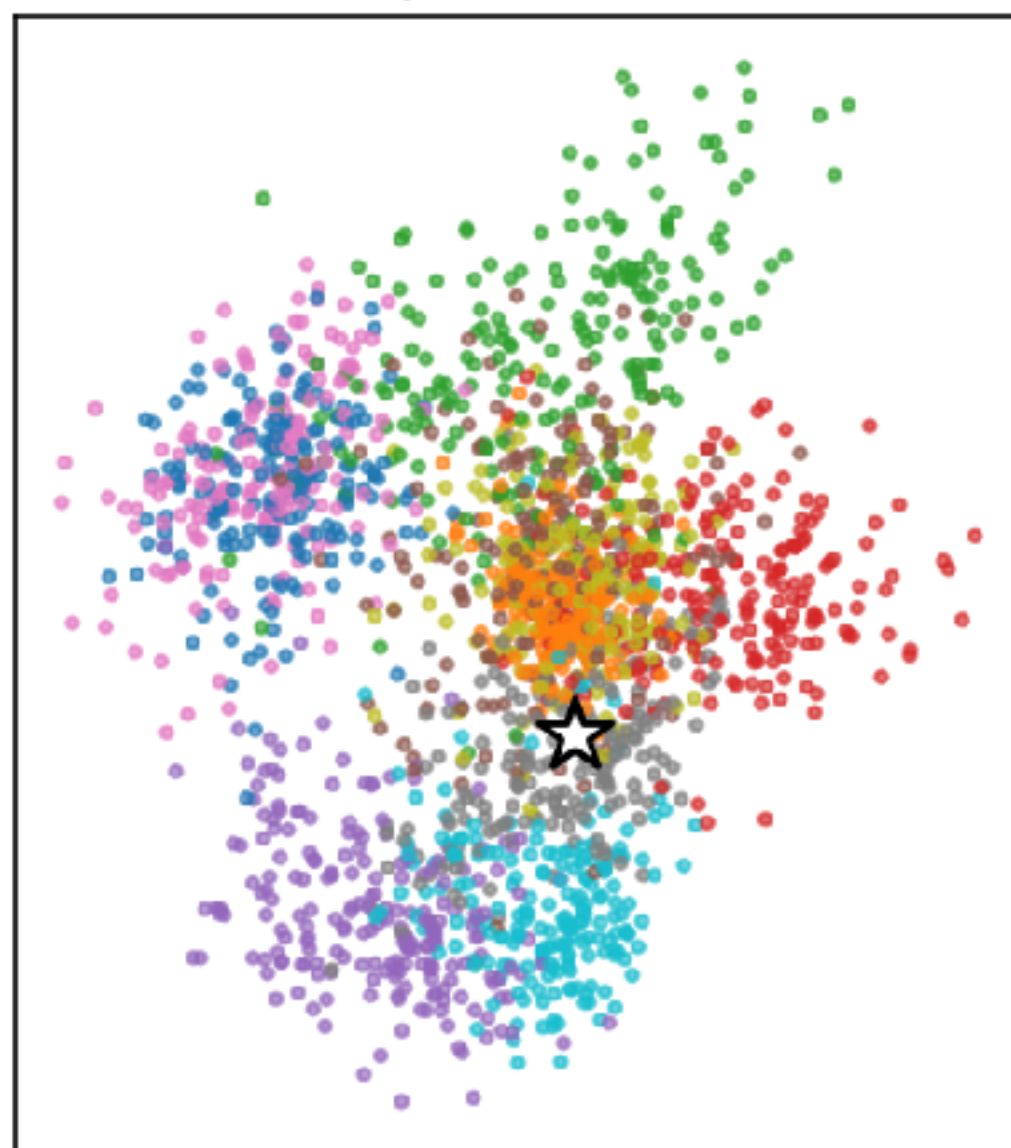
input (784)
PC1+PC2 explain 17% of variance



hidden 1 (128)
PC1+PC2 explain 28% of variance



hidden 2 (64)
PC1+PC2 explain 40% of variance



logits (10)
PC1+PC2 explain 54% of variance

