

Title: Machine Learning (2021/2022)

Speakers: Lauren Hayward

Collection: Machine Learning (2021/2022)

Date: April 11, 2022 - 11:30 AM

URL: <https://pirsa.org/22040067>

Abstract: This course is designed to introduce modern machine learning techniques for studying classical and quantum many-body problems encountered in condensed matter, quantum information, and related fields of physics. Lectures will focus on introducing machine learning algorithms and discussing how they can be applied to solve problem in statistical physics. Tutorials and homework assignments will concentrate on developing programming skills to study the problems presented in lecture.

## Outline for today

- k-NN algorithm
- Intro to supervised learning (SL)  
with neural networks (NN)
  - Network architecture and propagation of information
    - ↳ weights
    - ↳ biases
    - ↳ activation functions
  - Basics of training ("learning")

Recall the goal of SL: Given a dataset  $\mathcal{D} = \{(\vec{x}, \vec{y})\}$ ,  
fit  $\vec{f}(\vec{x})$  to  $\vec{y}$

- $\vec{x} = (x_1, x_2, \dots, x_{d_x})$

- $\vec{y} = (y_1, y_2, \dots, y_{d_y})$

- $\vec{f}: \mathbb{R}^{d_x} \rightarrow \mathbb{R}^{d_y}$

We have studied an SL algorithm: linear regression

## Other SL algorithms:

- k-Nearest Neighbours (kNN)
- Decision trees
- Logistic regression
- Support Vector machines
- Neural networks (NNs)

⋮



## kNN

For a dataset  $\mathcal{D} = \{(\vec{x}^{(i)}, y^{(i)})\}_{i=1}^M$

To determine  $f(\vec{x})$  for kNN:

↳ Step 1: Calculate the distance

$$d(\vec{x}, \vec{x}^{(i)}) = \|\vec{x} - \vec{x}^{(i)}\|_2 \text{ for all } \vec{x}^{(i)} \text{ in } \mathcal{D}$$

↳ Find the  $k$  closest points to  $\vec{x}$ . Store these as  $\vec{x}^{*(j)}$  and their labels as  $y^{*(j)}$  ( $1 \leq j \leq k$ )

## kNN

For a dataset  $\mathcal{D} = \{(\vec{x}^{(i)}, y^{(i)})\}_{i=1}^M$

To determine  $f(\vec{x})$  for kNN.

↳ Step 1: Calculate the distance

$$d(\vec{x}, \vec{x}^{(i)}) = \|\vec{x} - \vec{x}^{(i)}\|_2 \text{ for all } \vec{x}^{(i)} \text{ in } \mathcal{D}$$

↳ Find the  $k$  closest points to  $\vec{x}$ . Store these as  $\vec{x}^{*(j)}$  and their labels as  $y^{*(j)}$  ( $1 \leq j \leq k$ )

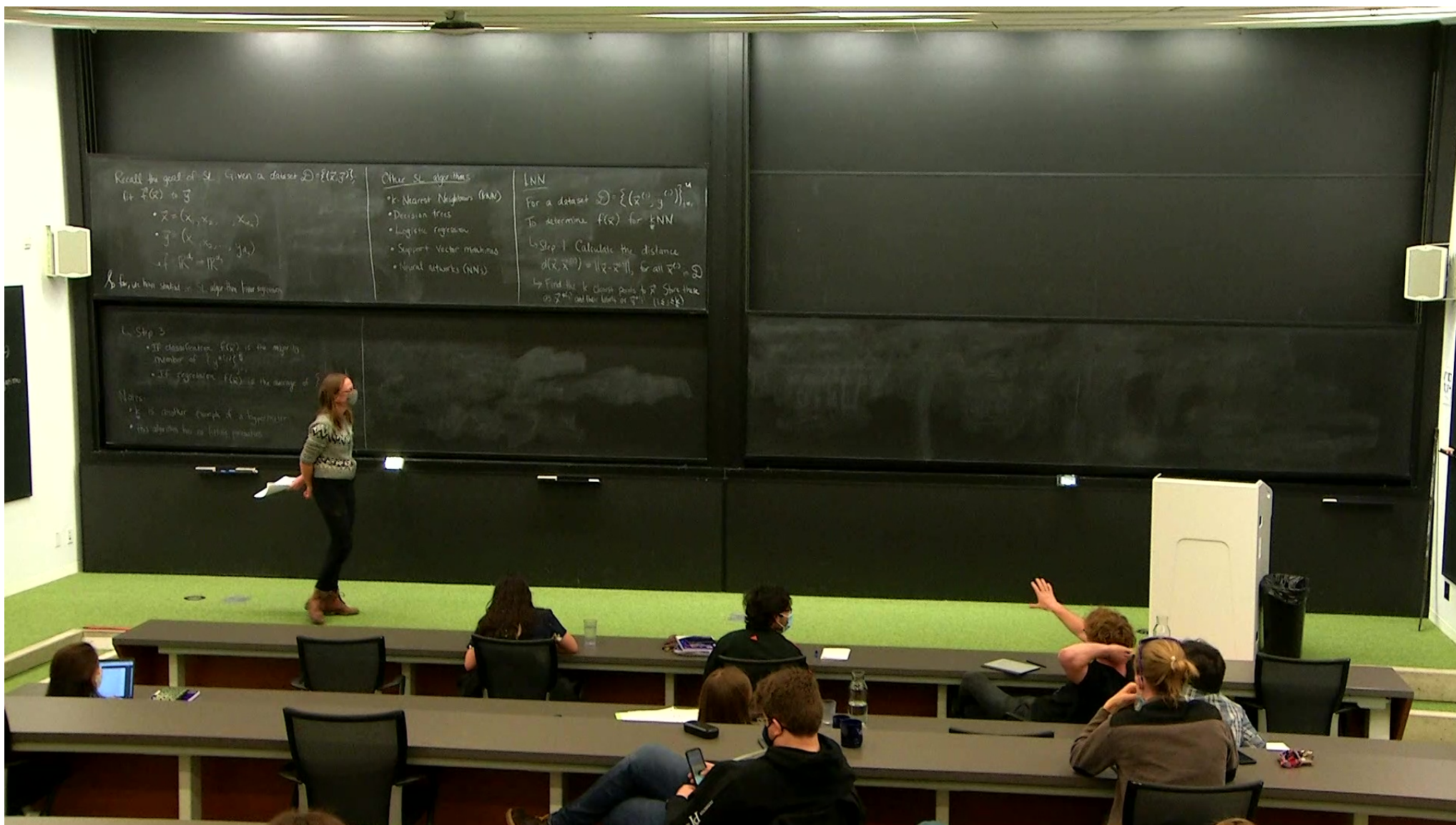
↳ Step 3:

- If classification:  $f(\bar{x})$  is the majority member of  $\{y^{*(j)}\}_{j=1}^k$
- If regression:  $f(\bar{x})$  is the average of  $\{y^{*(j)}\}_{j=1}^k$

Notes:

- $k$  is another example of a hypermeter

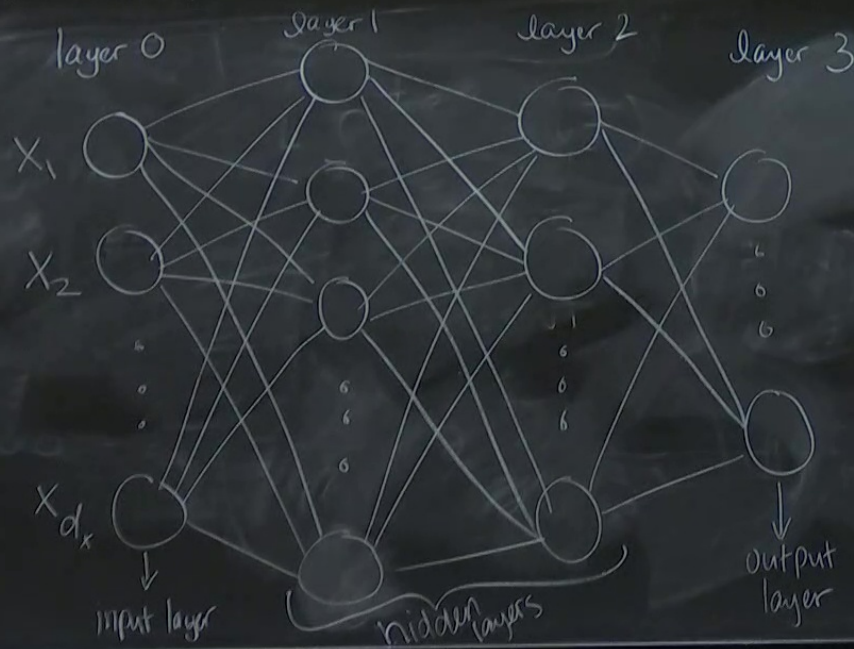




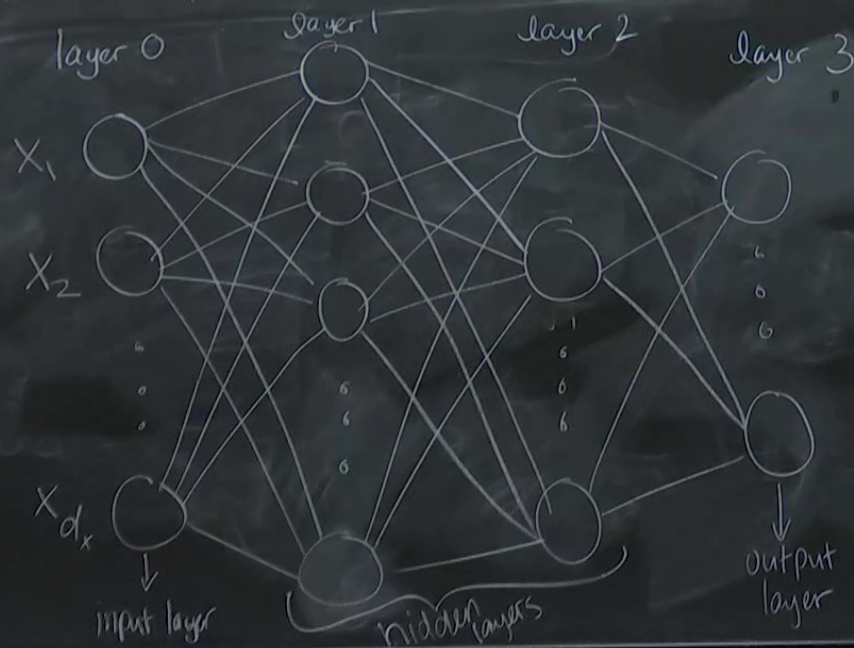


→ Find the  $k$  closest points to  $\vec{x}$ . Store these as  $\vec{x}^{*(j)}$  and their labels as  $\vec{y}^{*(j)}$  ( $1 \leq j \leq k$ )

## Feedforward Neural Networks



# Feedforward Neural Networks



Each  $\bigcirc$  is a "neurons"

Notation:

- $n_l \equiv \#$  of neurons in layer  $l$
- $L \equiv$  largest value of  $l$   
(eg:  $L=3$  here)
- $a_j^{(l)} \equiv$  output from the  $j^{\text{th}}$  neuron  
in layer  $l$

- this algorithm has no fitting parameters

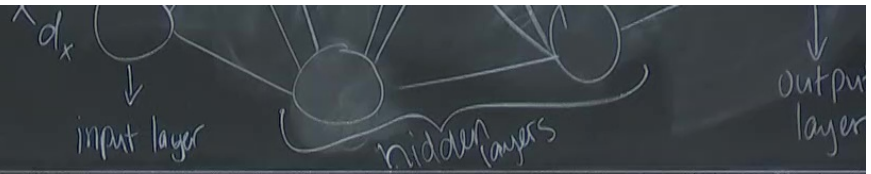
Then:

- $n_0 = d_x$
- $n_L = d_y$
- $a_i^{(0)} = x_i \quad (1 \leq i \leq d_x)$

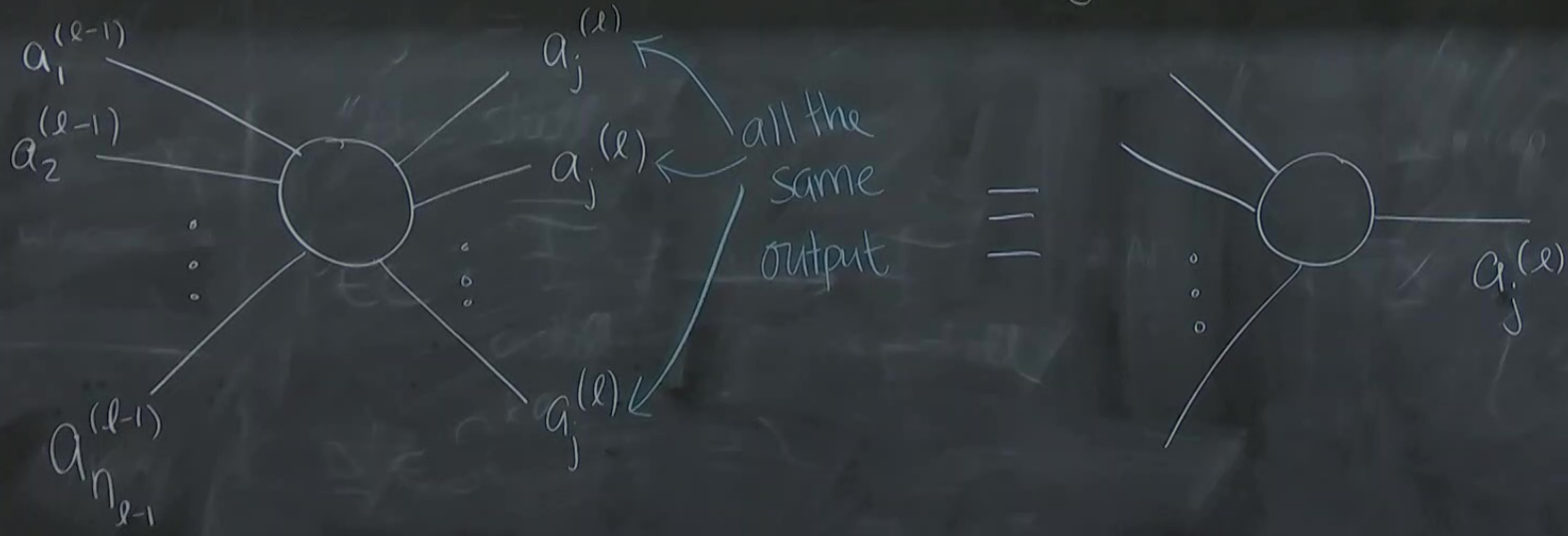
Let's zoom in on the  $j^{\text{th}}$  neuron in



ng parameters



Let's zoom in on the  $j^{\text{th}}$  neuron in layer  $l$



hidden layers

↓  
output layer

$a_j$  - output from the  $j$  neuron  
in layer  $l$

The output from this neuron is:

$$a_j^{(l)} = g_l \left( \underbrace{\sum_{i=1}^{n_{l-1}} a_i^{(l-1)} W_{ij}^{(l)} + b_j^{(l)}}_{\equiv z_j^{(l)}} \right)$$

$a_j^{(l)}$

where:

- each  $g_\ell$  is a nonlinear "activation function" that we choose
- each link has a "weight"  $W_{ij}^{(\ell)}$
- each neuron has a "bias"  $b_j^{(\ell)}$

We will see that the weights and biases are adjusted as the network "learns".



More notation:

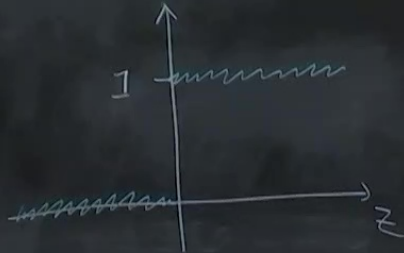
$$\vec{a}^{(l)} \equiv (a_1^{(l)}, a_2^{(l)}, \dots, a_{n_l}^{(l)})$$

we will see that the weights and biases are adjusted as the network "learns".

## Activation functions

Examples:

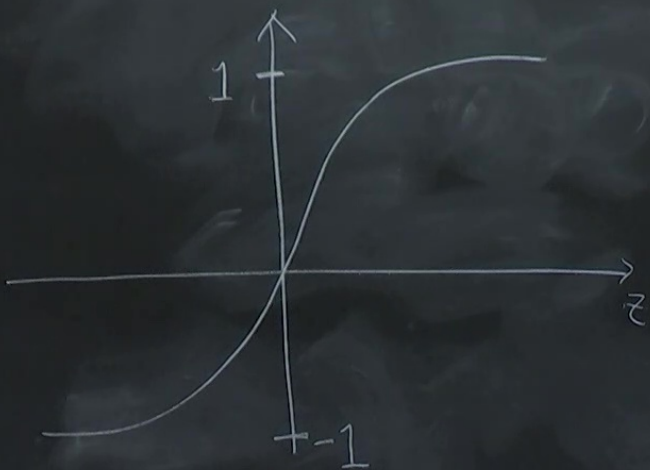
① Perceptron  $\textcircled{\text{---}}(z)$   
(step function)



learning is difficult  
because most "learning  
algorithms" take the  
gradient of the neurons' output

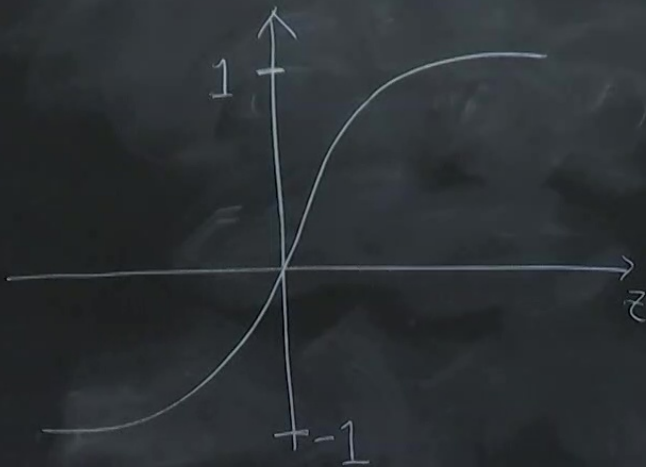
gradient of  $\tanh$  function output

$$\textcircled{3} \tanh(z) = \frac{e^z - e^{-z}}{e^z + e^{-z}}$$

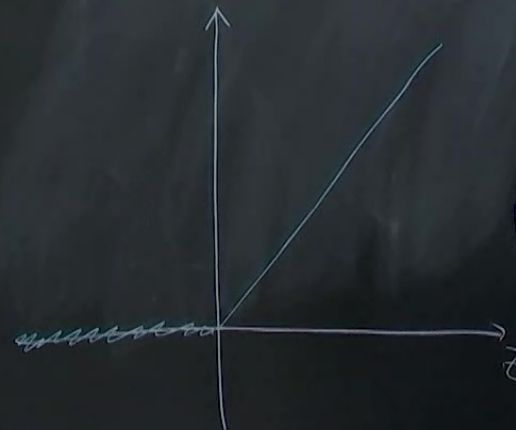




③  $\tanh(z) = \frac{e^z - e^{-z}}{e^z + e^{-z}}$



④ Rectified linear unit  
 $\text{ReLU}(z) = \max(0, z)$



⑤ Exponential linear unit

$$\text{ELU}(z) = \begin{cases} \alpha(e^z - 1) & \text{if } z < 0 \\ z & \text{if } z \geq 0 \end{cases}$$

( $\alpha > 0$  is a hyperparameter)



For classification with  $d_y > 1$ , it is common to apply the softmax activation function to the output layer

$$\left[ \text{softmax}(z_1, z_2, \dots, z_{n_e}) \right]_i = \frac{e^{z_i}}{\sum_{j=1}^{n_e} e^{z_j}} \quad (1 \leq i \leq n_e)$$

Feedfor

$x_1$

$x_2$

$x_d$



For classification with  $d_y > 1$ , it is common to apply the softmax activation function to the output layer

$$\left[ \text{softmax}(z_1, z_2, \dots, z_{n_c}) \right]_i = \frac{e^{z_i}}{\sum_{j=1}^{n_c} e^{z_j}} \quad (1 \leq i \leq n_c)$$

Outputs can be interpreted as probabilities