

Supervised Learning

Introduction & Framework

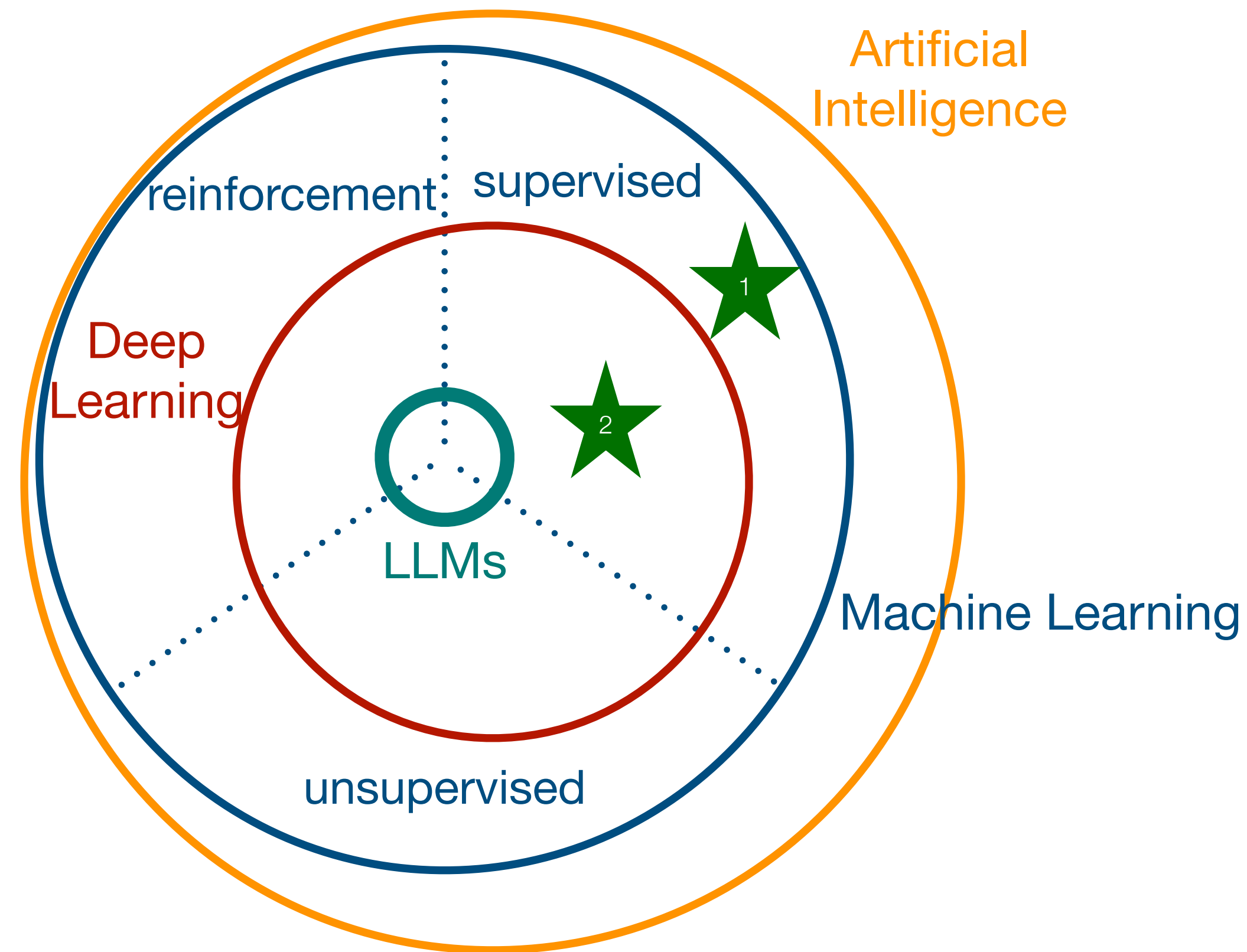
CMPUT 261: Introduction to Artificial Intelligence

P&M §7.1-7.3

Assignments

- **Assignment #2** is now available
 - Due **Oct 21/2025** (two weeks from today) at **11:59pm**

Supervised Learning vs. Machine Learning vs. Deep Learning



Inside an orange circle labeled "Artificial Intelligence" is a blue circle labeled "Machine Learning". The blue circle is divided into three thirds, labeled "reinforcement", "unsupervised", and "supervised". Inside the blue circle intersecting all three thirds is a red circle labelled "deep learning". Inside the red circle, intersecting all three thirds, is a cyan circle labelled "LLMs". Green star #1 is in the "supervised" third of Machine Learning, outside the Deep Learning circle. Green star #2 is in the "supervised" third of Machine Learning, inside the Deep Learning circle.

What is the difference between Supervised Learning, Machine Learning, and Deep Learning?

Lecture Outline

1. Recap & Logistics
2. Supervised Learning Problem
3. Measuring Prediction Quality

After this lecture, you should be able to:

- define supervised learning task, classification, regression, loss function
- represent categorical target values in multiple ways (indicator variables, indexes)
- define generalization performance
- identify an appropriate loss function for different tasks
- explain why a separate test set estimates generalization performance
- define 0/1 error, absolute error, (log-)likelihood loss, mean squared error, worst-case error

Regression Example

- Aim is to predict the value of **target** Y based on **features** X
- Both X and Y are **real-valued**
 - **Exact values** of both targets and features may not have been in the training set
 - Input 8 is an **interpolation** problem: X is **within the range** of the training examples' values
 - Input 9 is an **extrapolation** problem: X is **outside** the range of the training examples' values

i	$x^{(i)}$	$y^{(i)}$
1	0.7	1.7
2	1.1	2.4
3	1.3	2.5
4	1.9	1.7
5	2.6	2.1
6	3.1	2.3
7	3.9	7

8	2.9	?
9	5.0	?

Classification Example: Holiday Preferences

- An agent wants to learn a person's preference for the **length** of holidays
- Holiday can be for 1,2,3,4,5, or 6 days
- Two possible representations:

i	$y^{(i)}$
1	1
2	6
3	6
4	2
5	1

i	$y^{(i)}_1$	$y^{(i)}_2$	$y^{(i)}_3$	$y^{(i)}_4$	$y^{(i)}_5$	$y^{(i)}_6$
1	1	0	0	0	0	0
2	0	0	0	0	0	1
3	0	0	0	0	0	1
4	0	1	0	0	0	0
5	1	0	0	0	0	0

Question:

What are the advantages/
disadvantages of
each representation?

Generalization Example

Example: Consider binary two classifiers, **P** and **N**

- **P** classifies all the **positive examples** from the training data as **true**, and all others as **false**
- **N** classifies all of the **negative examples** from the training data as **false**, and all others as **true**

Question: Which classifier performs better on the **training data**?

Question: Which classifier **generalizes** better?

Bias

- The **hypothesis** is the function $h(X)$ that we learn
- The **hypothesis space** is the set of **possible hypotheses**
 - "Training a model" =
"Choosing a hypothesis from the hypothesis space based on data"
- A preference for one hypothesis over another is called **bias**
 - Bias is not a bad thing in this context!
 - Preference for "simple" models is a bias
 - Which bias works best for **generalization** is an **empirical** question

Measuring Prediction Error

- We choose our hypothesis partly by measuring its **performance** on training data
 - **Question:** What is the other consideration?
- This is usually described as **minimizing** some quantitative measurement of **error** (or **loss**)
 - **Question:** What might error mean?

0/1 Error

Definition:

The **0/1 error** for a dataset of n examples and hypothesis h is the number of examples for which the prediction was not correct:

$$\sum_{i=1}^n 1[y^{(i)} \neq h(\mathbf{x}^{(i)})]$$

- Not appropriate for **real-valued** target features (**why?**)
- Does not take into account **how wrong** the answer is
 - e.g., $1[2 \neq 1] = 1[6 \neq 1]$
- Most appropriate for **binary** or **categorical** target features

$1[\cdot]$ is **indicator function**:
value is 1 if the expression
in brackets is TRUE, else 0

Absolute Error

Definition:

The **absolute error** for a dataset of n examples and hypothesis h is the sum of absolute distances between the predicted target value and the actual target value:

$$\sum_{i=1}^n \left| y^{(i)} - h(\mathbf{x}^{(i)}) \right|$$

- Meaningless for **categorical** variables
- Takes account of **how wrong** the predictions are
- Most appropriate for **cardinal** or *possibly* **ordinal** values

