



U2 MACHINE LEARNING

U2.E2 ALGORITHMS OF MACHINE LEARNING: SUPERVISED, UNSUPERVISED, SEMI-SUPERVISED AND REINFORCEMENT LEARNING

Artificial Intelligence Technician

January 2021, Version 1



Co-funded by the
Erasmus+ Programme
of the European Union

The Development and Research on Innovative Vocational Educational Skills project (DRIVES) is co-funded by the Erasmus+ Programme of the European Union under the agreement 591988-EPP-1-2017-1-CZ-EPPKA2-SSA-B. The European Commission support for the production of this publication does not constitute endorsement of the contents which reflects the views only of the authors, and the Commission cannot be held responsible for any use which may be made of the information contained therein.

The student is able to

AIE.U2.E2.PC1	The student can list the different types of machine learning methods.
AIT.U2.E2.PC2	The student knows how to define and explain the different machine learning methods as well as the conditions of use of each method.

Supervised Learning method trains a function (or algorithm) to compute output variables based on a given data in which both input and output variables are known.

The goal of a learning process is to find a function that minimizes the risk of a prediction error that is expressed as a difference between the actual and the computed output values when tested on a dataset. In such cases, the learning process may be controlled by a predetermined acceptable error threshold.

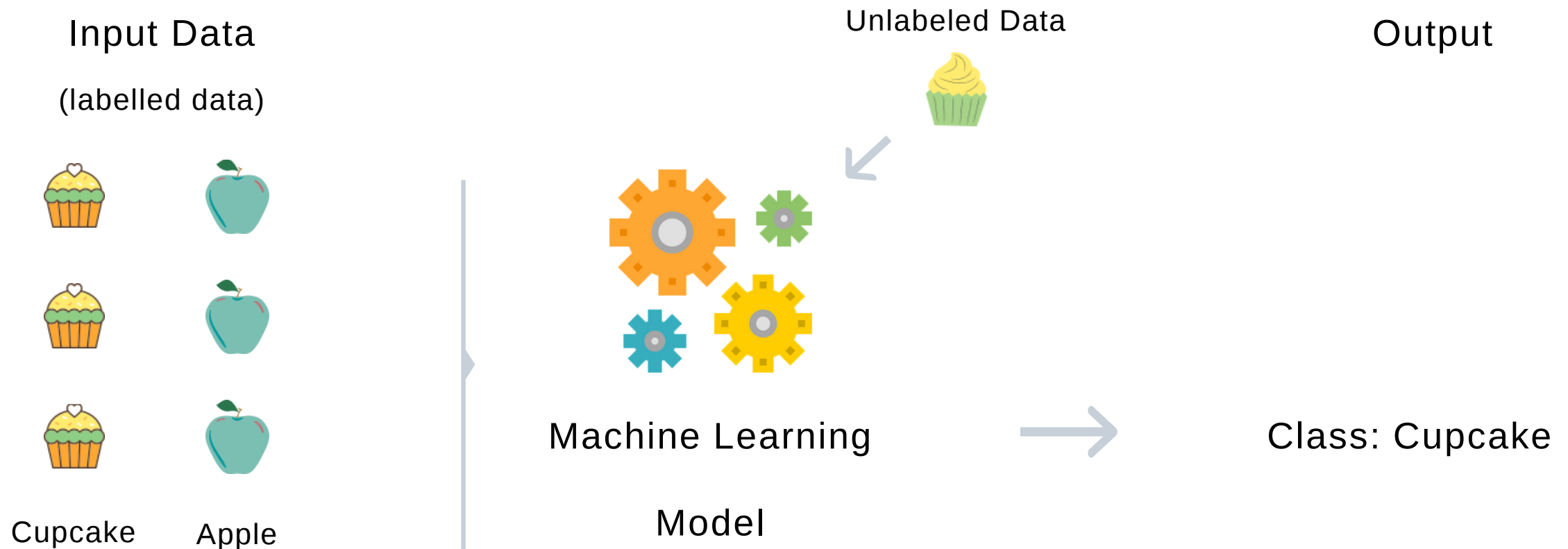


As the name suggests, here the learning occurs under supervision.

The training dataset is
classified

You already know how the
output should be

There is a relationship
between the input and the
output





ANALOGY WITH DRIVING LESSONS

The supervised learning process can be seen as a collection of guidelines provided by a driving instructor to explain what should be done (output variables) in different situations (input variables). These guidelines are adapted by a student driver and turned into a driver behavior.

The predetermined thresholds can be seen as the standards to pass the driving exam. In this case, the student driver knows the standard way to drive (i.e., actual output) and steps to achieve it (i.e., actual inputs) before he/she starts the driving lessons.

For the student driver, it becomes an iterative process to achieve acceptable performance. In each iteration, the student makes mistakes that are corrected by the driving instructor (i.e., training the model). This iterative process ends when the student gets the driving license (i.e., the model achieves a satisfactory performance).



IMPORTANT NOTIONS

Data: a set of labeled data records/instances/cases/examples $\langle x_i, y \rangle$

Training Set

The portion of the dataset from which the ML algorithms discover or learn relationships between the features and the target attribute. The training set is labeled in supervised learning.

Validation Set

Another portion of the dataset to which the ML algorithms are applied in order to check how well they identify relationships between the known outcomes of the target variable and the other features of the dataset.

Test/Holdout Set

The subset of the dataset that provides a final estimate of the performance of the ML models after they have been trained and validated. Holdout sets should never be used to make decisions about which algorithms to use or to improve or tune algorithms.

The training set should not be used in testing and the test set should not be used in learning!



An unseen test set provides an unbiased estimation of the model's performance.

➔ IMPORTANT NOTIONS

Features: a measurable property of the object being analysed. In datasets, features appear as columns. Features are also sometimes referred to as “variables” or “attributes”.

# Glucose	# BloodPres...	# SkinThickn...	# Insulin	# BMI	# Age	# Outcome
148	72	35	0	33.6	50	1
85	66	29	0	26.6	31	0
183	64	0	0	23.3	32	1
89	66	23	94	28.1	21	0
137	40	35	168	43.1	33	1
116	74	0	0	25.6	30	0
78	50	32	88	31	26	1
115	0	0	0	35.3	29	0
197	70	45	543	30.5	53	1



IMPORTANT NOTIONS

Features:

The quality of the dataset's features has a huge influence on the quality of the insights that will be achieved when using the dataset for ML.

In addition, different business problems within the same industry do not necessarily require the same features, which is why it is important to have a strong understanding of the business objectives of the data science project.

The quality of dataset's features can be improved with processes such as **feature selection** and **feature engineering**.



IMPORTANT NOTIONS

Features:

Feature Selection

eliminates irrelevant or redundant columns from the dataset without sacrificing accuracy.



1. Reduce the chance of overfitting;
2. Boost the run speed of the algorithm by reducing the CPU, I/O, and RAM load the production system needs to construct and use the model;
3. Increase the interpretability of the model by identifying the most informative factors that drive the results of the model.

Feature Engineering

constructs additional variables to the dataset to improve the performance and accuracy of the ML model.



1. Provide a deeper understanding of the data;
2. Improve the predictive power of the ML model;
3. Deliver more valuable insights;



IMPORTANT NOTIONS

Experimentation cycle:

- Data gathering
- Data preparation
- Learn parameters on the training set
- Hyper-parameter tuning on the validation set
- Compute evaluation metrics on the test set and make predictions



Supervised learning can further be categorized as **Regression** and **Classification** problems.

In the case of a **classification** problem, the aim of the ML algorithm is to categorize or classify the inputs based on the training dataset. The training dataset in a classification problem includes a set of input:output pairs categorized in classes.

- Is a given patient with covid-19 or healthy?
- Is this an image of a dog, a cat, or a horse?

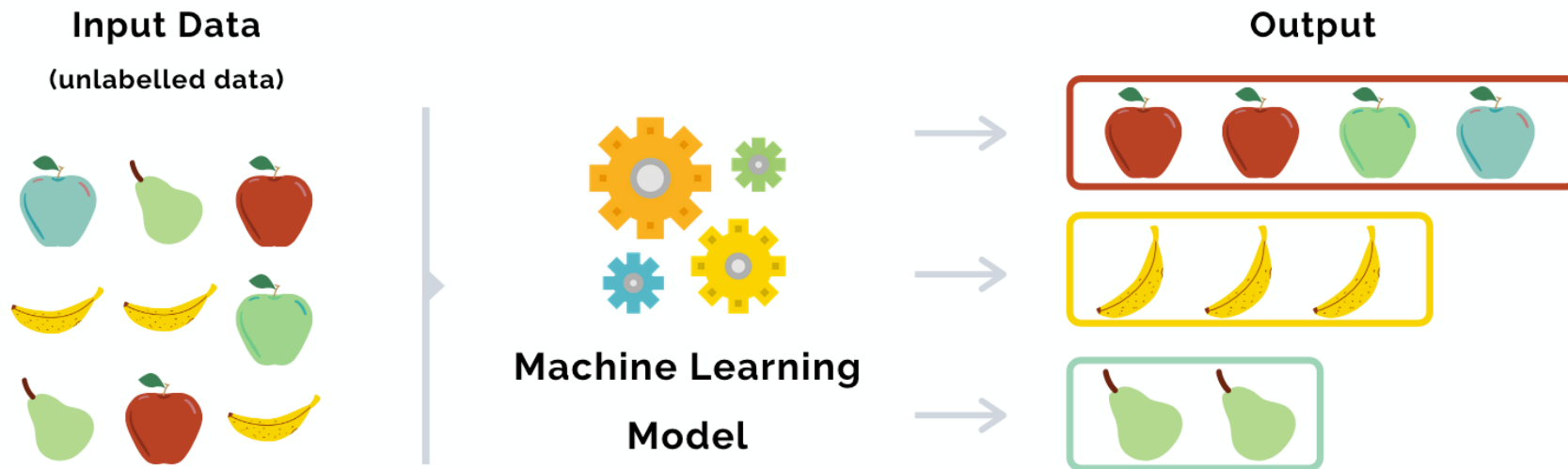
For a **regression** problem, the goal of the ML algorithm is to develop a relationship between outputs and inputs using a continuous function to help machines understand how outputs change for inputs. The relationship between output variables and input variables can be defined by various mathematical functions such as linear, nonlinear, and logistic.

- What is the company's revenue by the end of the year?
- What is the probability that tomorrow will rain?

While classification methods are used when the output is of categorical nature, the regression methods are used for continuous output.

Unsupervised Learning is a class of Machine Learning algorithms that uses them to analyze and cluster unlabeled datasets

In Unsupervised learning the goal is to learn useful structure without labeled classes, optimization criterion, feedback signal, or any other information beyond the raw data



Prime reasons to use Unsupervised Learning:

- 1 Unsupervised machine learning finds all kind of unknown patterns in data.
- 2 Unsupervised methods help you to find features which can be useful for categorization.
- 3 It is taken place in real time, so all the input data to be analyzed and labeled in the presence of learners.
- 4 It is taken place in real time, so all the input data to be analyzed and labeled in the presence of learners.

Unsupervised Learning can be further classified into two categories:



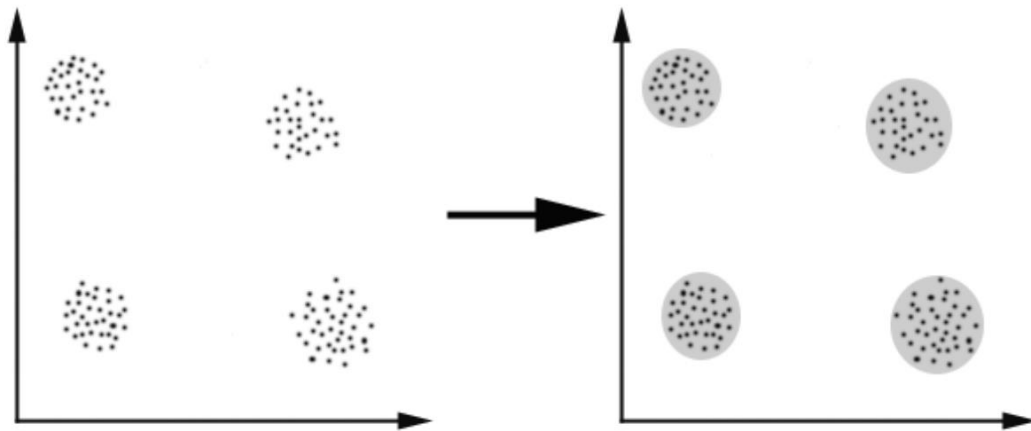
Parametric Unsupervised Learning

- Assumes that sample data comes from a population that follows a probability distribution based on a fixed set of parameters.
- Involves construction of Gaussian Mixture Models and using Expectation-Maximization algorithm to predict the class of the sample in question

Non-Parametric Unsupervised Learning

- The data is grouped into clusters, where each cluster says something about categories and classes present in the data.
- Do not require the modeler to make any assumptions about the distribution of the population, and so are sometimes referred to as a distribution-free method.

A *cluster* is a collection of objects which are “similar” between them and are “dissimilar” to the objects belonging to other clusters.



Distance-based clustering.

Given a set of points, with a notion of distance between points, grouping the points into some number of *clusters*, such that:

- internal (within the cluster) distances should be small
- external (intra-cluster) distances should be large

1 Exclusive Clustering: K-means

Most common type of clustering. Each object belongs to an exclusive cluster. Data point belongs to a definite cluster then it could not be included in another cluster.

2 Overlapping Clustering: Fuzzy C-means

Uses fuzzy sets to cluster data, so that each point may belong to two or more clusters with different degrees of membership.

3 Hierarchical Clustering: Agglomerative clustering, divisive clustering

Is based on the union between the two nearest clusters. The beginning condition is realized by setting every data point as a cluster. After a few iterations it reaches the final clusters wanted

4 Probabilistic Clustering: Mixture of Gaussian models

Data points are clustered based on the likelihood that they belong to a particular distribution



The main objective of the K-Means algorithm is to minimize the sum of the distances between the points and their grouping centroid.



It is an iterative algorithm that tries to partition the data set into distinct K subgroups (clusters) without overlapping

The k-means algorithm works as follows:

- 1 Specify number of K clusters.
- 2 Initialize the centroids by shuffling the data set first and randomly selecting K data points for the centroids without substitution.
- 3 Continue iterating until there are no changes to the centroids. i.e., the assignment of data points to clusters is not changing.

It is similar in process to the K-Means clustering but it works differently:

- 1** Choose a number of clusters (K).
- 2** Assign coefficients randomly to each data point for being in the clusters.
- 3** Repeat until the algorithm has converged.
- 4** Compute the centroid for each cluster.

Therefore this algorithm will not overfit the data for clustering like the k-means algorithm it will mark the data point to multiple clusters instead of the one cluster which will be more helpful than giving the point to the one cluster.

It can be categorized in two ways:



Four different methods are commonly used to measure similarity:

- **Ward's linkage:** States that the distance between two clusters is defined by the increase in the sum of squared after the clusters are merged.
- **Average linkage:** Defined by the mean distance between two points in each cluster
- **Complete (or maximum) linkage:** Defined by the maximum distance between two points in each cluster
- **Single (or minimum) linkage:** Defined by the minimum distance between two points in each cluster



It is the most common type of hierarchical clustering used to group objects in clusters based on their similarity.



It is also known as AGNES, Agglomerative Nesting.

The AGNES algorithm works as follows:

1

Preparing the data. The data should be a numeric matrix with rows(representing observations) and columns (representing variables).

2

Compute similarity information between each pair of objects in the dataset

3

Using linkage function, groups objects into hierarchical cluster trees.

4

Determines where to cut the hierarchical trees into clusters. Creates partitions of the data

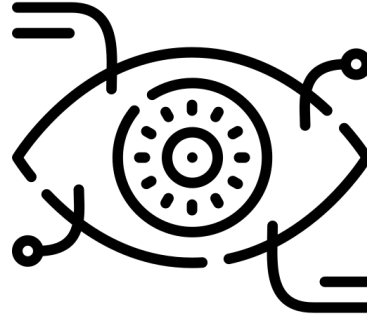
Top-down clustering requires a method for splitting a cluster that contains the whole data and proceeds by splitting clusters recursively until individual data have been split into singleton cluster.

- 1 Considers the entire data as one group
- 2 Iteratively splits the data into subgroups
- 3 If the number of a hierarchical clustering algorithm is known, then the process of division stops once the number of clusters is achieved.
- 4 Else, the process stops when the data can be no more split



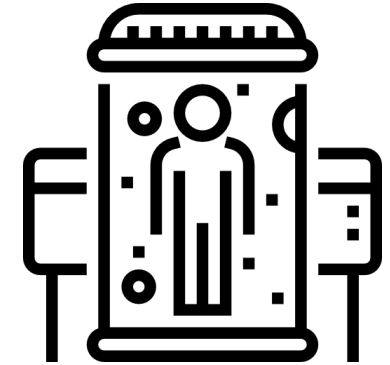
News Sections

Google News uses unsupervised learning to categorize articles on the same story from various online news outlets.



Computer vision

Unsupervised learning algorithms are used for visual perception tasks, such as object recognition.



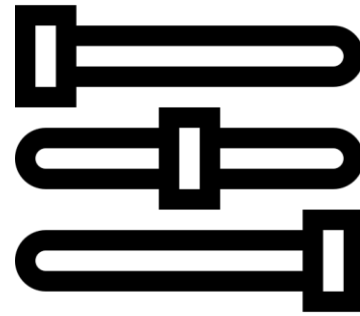
Medical imaging

Unsupervised machine learning provides essential features to medical imaging devices, such as image detection, classification and segmentation



Anomaly detection

Unsupervised learning models can comb through large amounts of data and discover atypical data points within a dataset.



Customer personas

Unsupervised learning allows businesses to build better buyer persona profiles, enabling organizations to align their product messaging more appropriately.



Recommendation Engines

Unsupervised learning can help to discover data trends that can be used to develop more effective.

Is it possible to improve the quality of learning by combining labeled and unlabeled data?

- ➡ There is usually a lot more unlabeled data available than labeled.
- ➡ Semi-Supervised Learning (SSL) is a mixture between supervised and unsupervised approaches.
- ➡ In addition to unlabeled data, the algorithm is provided with some supervision information – but not necessarily for all examples.

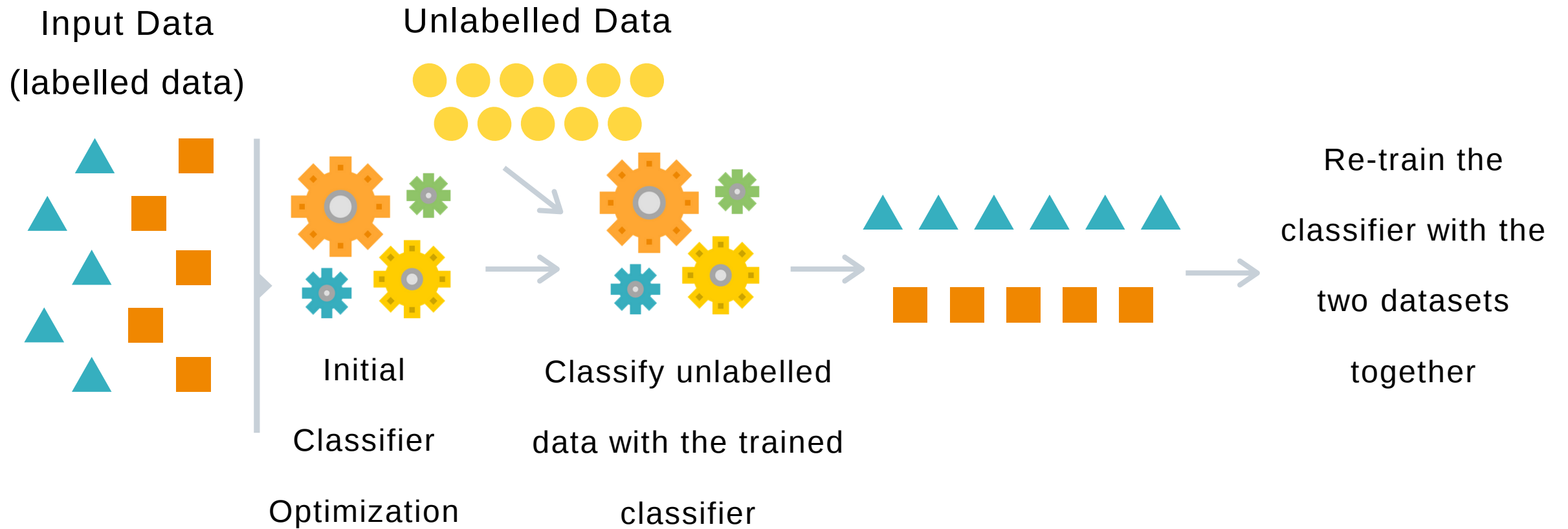


As the name suggests, SSL is halfway between supervised and unsupervised learning.

The dataset consists of a combination of labelled and unlabelled data

First, similar data is grouped using an unsupervised learning algorithm and only then the supervised learning algorithms are applied

It can refer to either transductive or inductive learning



➡ *Labeled vs. Unlabeled Data*

0 1 2 3 4 5 6 7 8 9

"0" "1" "2" "3" "4" "5" "6" "7" "8" "9"



➡ Human Expert / Special Equipment ➡

"dog" "cat" "dog"

Unlabeled data, X_i

Labeled data, Y_i

Cheap and Abundant

Expensive and Scarce

Feature Space X **Label Space Y**

GOAL: Construct a predictor $f: X \rightarrow Y$ to minimize $R(f) \equiv \mathbb{E}_{XY}[\text{loss}(Y, f(X))]$

An optimal predictor (Bayes Rule) depends on unknown P_{XY} . So, instead *learn* a good prediction rule from training data $\{(X_i, Y_i)\}_{i=1}^n \stackrel{\text{iid}}{\sim} P_{XY}(\text{unknown})$

Training Data

 $\{(X_i, Y_i)\}_{i=1}^n$ *labeled*

Learning Algorithm



Prediction Rule

 $\hat{f}_n$

Training Data

$$\{(X_i, Y_i)\}_{i=1}^n$$
$$\{X_i\}_{i=1}^m$$



Learning Algorithm



Prediction Rule

$$\hat{f}_{n,m}$$

Supervised Learning

Labeled Data $\{(X_i, Y_i)\}_{i=1}^n$



X_i

“dog”

Y_i

Semi-Supervised Learning

Labeled Data $\{(X_i, Y_i)\}_{i=1}^n$ and Unlabeled Data $\{X_i\}_{i=1}^m$

$$m \gg n$$

Goal: Learn a better prediction rule than based on labeled data alone.

➡ COGNITIVE SCIENCE

Computational modal of how humans learn from labeled and unlabeled data:

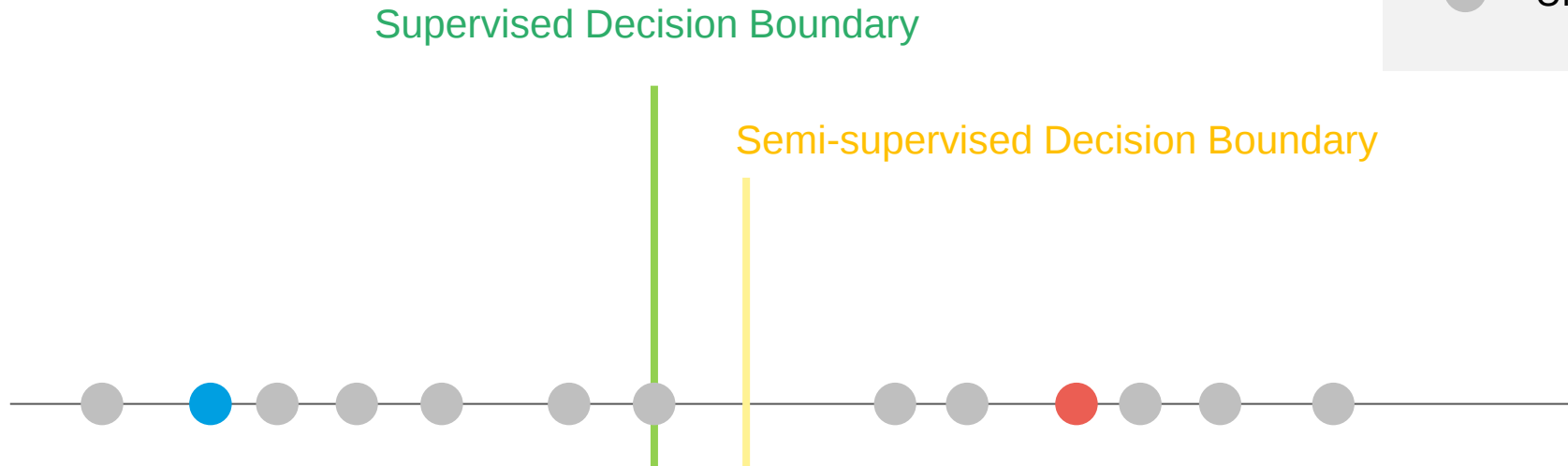
- Concept learning in children: x =animal, y =concept (e.g. cat)
- Dad/Mum points to an animal and says “this is a cat”
- Children also observe animals by themselves



CAN UNLABELED DATA HELP?

“Similar” data points have “similar” labels

- Positive labeled data
- Negative labeled data
- Unlabeled data



- Assume each class is a coherent group (e.g. Gaussian)
- Then, unlabeled data can help identify the boundary more accurately

➡ WHY SHOULD WE USE SSL?

- ➡ Labeling data is usually expensive, specially when a significant number of cases are being addressed;

I have a good idea, but I can't afford to label lots of data!

- ➡ Even when labeled data are available in significant amounts, it is useful to use more data to introduce as much variability as possible in the studies carried out;

I have lots of labeled data, but I have even more unlabeled data available!

- ➡ Domain adaptation.

I have labeled data from a domain, but I want a model for a different domain!

➡ SSL ALGORITHMS

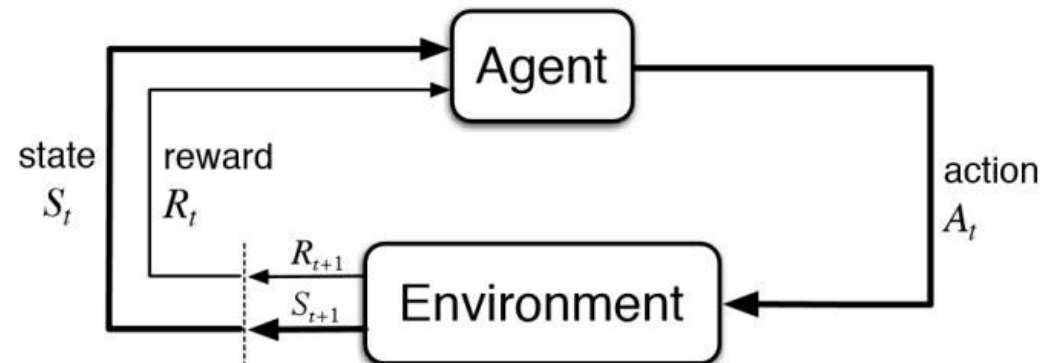
- Self-Training
- Generative models
- Mixture models
- Graph-based models
- Co-Training
- Semi-supervised SVM or Transductive Support Vector Machine (TSVM)
- ...

- Most approaches make strong model assumptions (guesses).
- Some commonly used assumptions:
 - Clusters of data are from the same class
 - Data can be represented as a mixture of parameterized distributions
 - Decision boundaries should go through non-dense areas of the data
 - Model should be as simple as possible

Does not require labeled input/output pairs to be submitted

It is a paradigm that is concerned with how software agents should act in an environment in order to maximize the notion of cumulative reward.

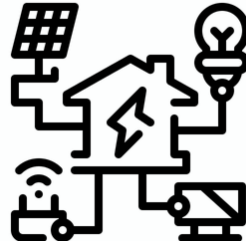
It is a type of dynamic programming that forms algorithms using a reward and punishment system





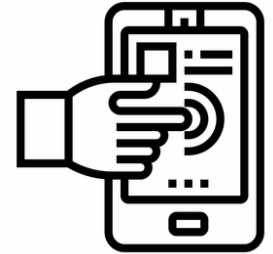
Agent

the learner and the
decision maker.



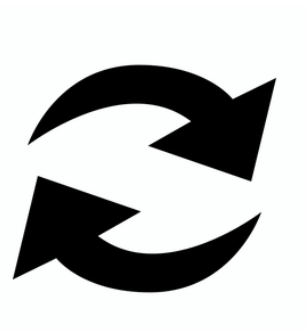
Environment

Where the agent
learns and decides
what actions to
perform.



Action

A set of actions
which the agent can
perform.



State

The state of the agent in the environment.



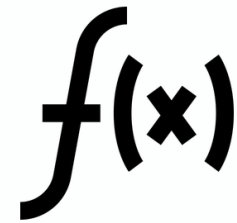
Reward

For each action selected by the agent the environment provides a reward. Usually a scalar value.



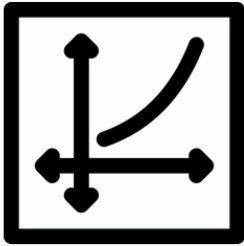
Policy

The decision-making function (control strategy) of the agent, which represents a mapping from situations to actions.



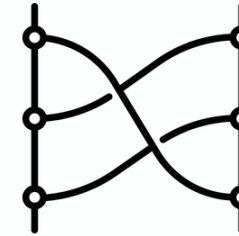
Value function

Mapping from states to real numbers, where the value of a state represents the long-term reward achieved starting from that state, and executing a particular policy.



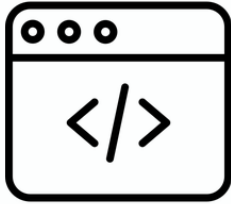
Function approximator

Refers to the problem of inducing a function from training examples. Standard approximators include decision trees, neural networks, and nearest-neighbor methods



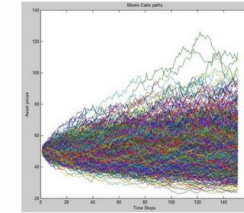
Markov decision process (MDP)

A probabilistic model of a sequential decision problem, where states can be accurately perceived, and the current state and action chosen to determine the probability distribution of future states. Essentially, the outcome of applying an action to a state depends only on the current action and state (and not on previous actions or states).



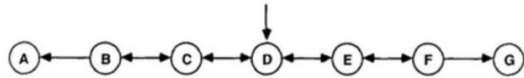
Dynamic programming (DP)

is a class of solution methods for solving sequential decision problems with a compositional cost structure. Richard Bellman was one of the principal founders of this approach.



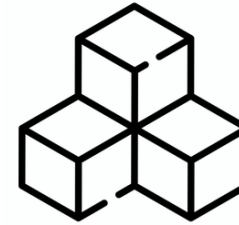
Monte Carlo methods

A class of methods for learning of value functions, which estimates the value of a state by running many trials starting at that state, then averages the total rewards received on those trials.



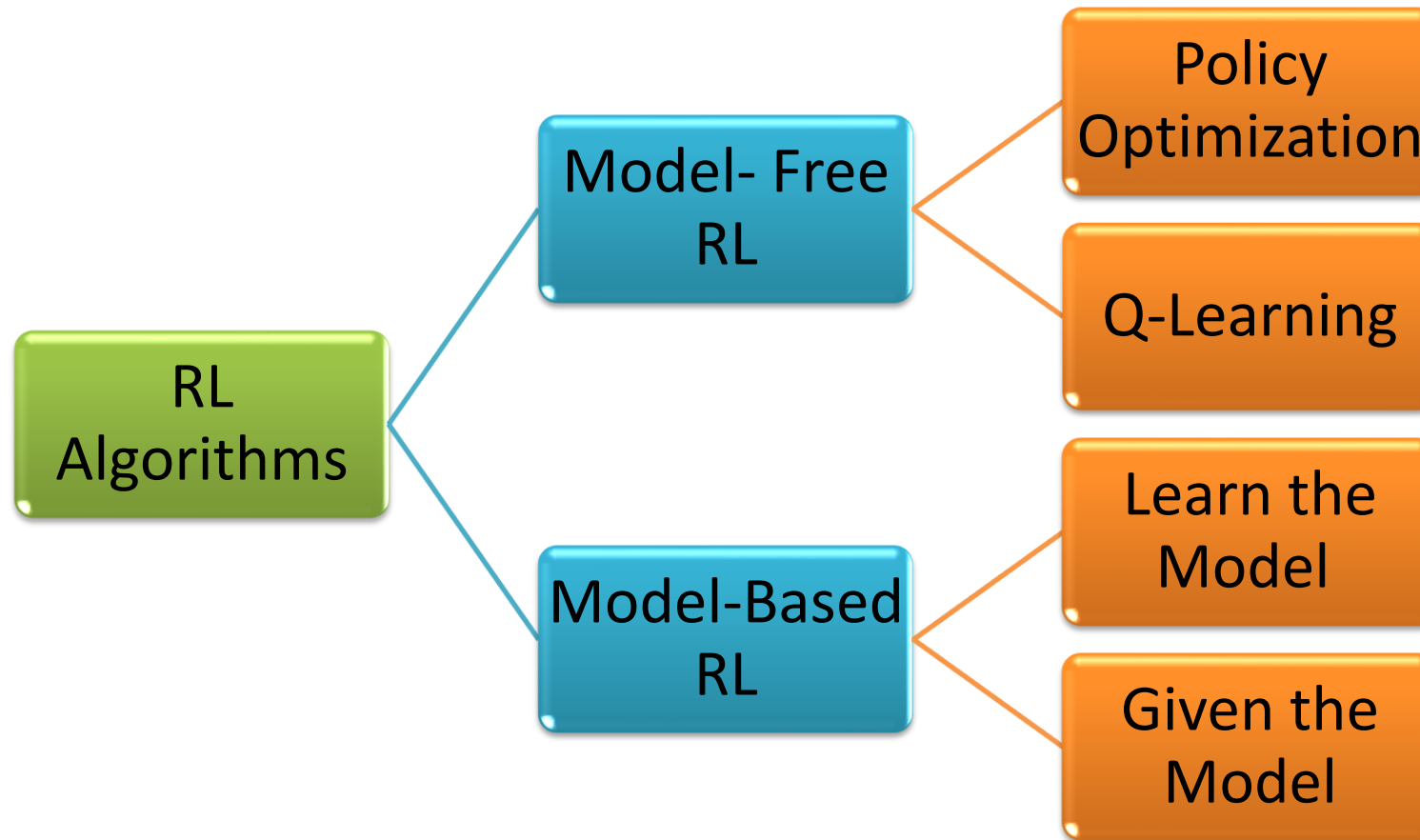
Temporal Difference (TD) algorithms

A class of learning methods, based on the idea of comparing temporally successive predictions. Possibly the single most fundamental idea in all of reinforcement learning.



Model

The agent's view of the environment, which maps state-action pairs to probability distributions over states. Note that not every reinforcement learning agent uses a model of its environment



- Can be used to solve very complex problems that cannot be solved by conventional techniques;
- This technique is preferred to achieve long-term results, which are very difficult to achieve.
- It is very similar to the learning of human beings. Therefore, it is close to achieving perfection.
- Can correct the errors that occurred during the training process.
- Once an error is corrected by the model, the chances of occurring the same error are very less.
- Robots can implement reinforcement learning algorithms to learn how to walk.
- Reinforcement learning models can outperform humans in many tasks
- Reinforcement learning is intended to achieve the ideal behavior of a model within a specific context, to maximize its performance.
- It can be useful when the only way to collect information about the environment is to interact with it.

- Reinforcement learning as a framework is wrong in many different ways, but it is precisely this quality that makes it useful.
- Too much reinforcement learning can lead to an overload of states, which can diminish the results.
- It is not preferable to use for solving simple problems.
- Needs a lot of data and a lot of computation.
- Assumes the world is Markovian, which it is not.

Resources management in computer clusters

Designing algorithms to allocate limited resources to different tasks is challenging and requires human-generated heuristics.

Self-Driving Cars

Various papers have proposed Deep Reinforcement Learning for autonomous driving. In self-driving cars, there are various aspects to consider, such as speed limits at various places, drivable zones, avoiding collisions

Industry Automation

In industry reinforcement, learning-based **robots** are used to perform various tasks. Apart from the fact that these robots are more efficient than human beings, they can also perform tasks that would be dangerous for people.

Trading And Finance

Supervised time series models can be used for predicting future sales as well as predicting stock prices. However, these models don't determine the action to take at a particular stock price.

Healthcare

In healthcare, patients can **receive treatment** from policies learned from RL systems. RL can find optimal policies using previous experiences without the need for previous information on the mathematical model of biological systems.

News Recommendation

User preferences can change frequently, therefore **recommending news** to users based on reviews and likes could become obsolete quickly. With reinforcement learning, the RL system can track the reader's return behaviors.

- There are 4 big types of machine learning: Supervised Learning, Unsupervised Learning, Semi-Supervised Learning and Reinforcement Learning;
- Supervised Learning method trains a function (or algorithm) to compute output variables.
- In Supervised Learning there are 3 important notions: training set, validation set , test set and features
- Supervised learning can further be categorized as Regression and Classification problems.
- In Unsupervised learning the goal is to learn useful structure without labeled classes
- Unsupervised Learning can be further classified into two categories: Parametric or Non-parametric supervised learning

- There are 4 types of clustering algorithms: Exclusive Clustering (Kmeans), Hierarchical Clustering (Agglomerative and Divisive Clustering), Overlapping Clustering (Fuzzy c-means), Probabilistic Clustering (Mixture of Gaussian models)
- Applications of unsupervised learning: News Sections, Computer vision, Medical imaging, Anomaly detection, Customer personas, Recommendation Engines
- Semi-Supervised Learning is a mixture between supervised and unsupervised approaches.
- Reinforcement Learning is a type of dynamic programming that forms algorithms using a reward and punishment system
- Reinforcement Learning Algorithms are divided between model based and model free approaches
- Reinforcement Learning can be applied to Resources management in computer clusters, Self-Driving Cars, Industry Automation, among others

- Zhao, Z., & Liu, H. (2007, June). Spectral feature selection for supervised and unsupervised learning. In *Proceedings of the 24th international conference on Machine learning* (pp. 1151-1157). Barlow, H. B. (1989). Unsupervised learning. *Neural computation*, 1(3), 295-311.
- Celebi, M. E., & Aydin, K. (Eds.). (2016). *Unsupervised learning algorithms*. Berlin: Springer International Publishing.
- Hastie, T., Tibshirani, R., & Friedman, J. (2009). Unsupervised learning. In *The elements of statistical learning* (pp. 485-585). Springer, New York, NY.
- Ashfaq, R. A. R., Wang, X. Z., Huang, J. Z., Abbas, H., & He, Y. L. (2017). Fuzziness based semi-supervised learning approach for intrusion detection system. *Information Sciences*, 378, 484-497.
- Sathya, R., & Abraham, A. (2013). Comparison of supervised and unsupervised learning algorithms for pattern classification. *International Journal of Advanced Research in Artificial Intelligence*, 2(2), 34-38.
- Barlow, H. B. (1989). Unsupervised learning. *Neural computation*, 1(3), 295-311.
- Zhu, X., & Goldberg, A. B. (2009). Introduction to semi-supervised learning. *Synthesis lectures on artificial intelligence and machine learning*, 3(1), 1-130.
- Hastie, T., Tibshirani, R., & Friedman, J. (2009). Overview of supervised learning. In *The elements of statistical learning* (pp. 9-41). Springer, New York, NY.



Regina Sousa

- PhD student
in Biomedical Engineering
- Research Collaborator of the
Algoritmi Research Center

 [0000-0002-2988-196X](https://orcid.org/0000-0002-2988-196X)



Diana Ferreira

- PhD student
in Biomedical Engineering
- Research Collaborator of the
Algoritmi Research Center

 [0000-0003-2326-2153](https://orcid.org/0000-0003-2326-2153)



José Machado

- Associate Professor with
Habilitation at the University of
Minho
- Integrated Researcher
of the Algoritmi Research Center

 [0000-0003-4121-6169](https://orcid.org/0000-0003-4121-6169)



António Abelha

- Assistant Professor at the University of Minho
- Integrated Researcher of the Algoritmi Research Center

 [0000-0001-6457-0756](https://orcid.org/0000-0001-6457-0756)



Victor Alves

- Assistant Professor at the University of Minho
- Integrated Researcher of the Algoritmi Research Center

 [0000-0003-1819-7051](https://orcid.org/0000-0003-1819-7051)

This Training Material has been certified according to the rules of **ECQA – European Certification and Qualification Association**.

The Training Material was developed within the international job role committee “**Artificial Intelligence Technician**”:

UMINHO – University of Minho (<https://www.uminho.pt/PT>)

The development of the training material was partly funded by the EU under Blueprint Project DRIVES.



Thank you for your attention

DRIVES project is project under **The Blueprint for Sectoral Cooperation on Skills in Automotive Sector**, as part of New Skills Agenda.

The aim of the Blueprint is **to support an overall sectoral strategy and to develop concrete actions to address short and medium term skills needs.**

Follow DRIVES project at:



More information at:

www.project-drives.eu



Co-funded by the
Erasmus+ Programme
of the European Union

The Development and Research on Innovative Vocational Educational Skills project (DRIVES) is co-funded by the Erasmus+ Programme of the European Union under the agreement 591988-EPP-1-2017-1-CZ-EPPKA2-SSA-B. The European Commission support for the production of this publication does not constitute endorsement of the contents which reflects the views only of the authors, and the Commission cannot be held responsible for any use which may be made of the information contained therein.