



IMT Mines Alès
École Mines-Télécom



Master Sciences et Numérique pour la Santé

HAH811E9 - SCIENCE DES DONNÉES NIVEAU 1

Gérard Dray
Nicolas Sutton-Charani



Contexte

A la croisée de plusieurs champs disciplinaires : **mathématiques, probabilités, statistiques, informatique, intelligence artificielle, théorie de l'information et visualisation**, la science des données met en œuvre différents outils d'analyse de données afin d'extraire automatiquement des informations utiles, des connaissances, à partir de données potentiellement massives. Le but ultime est de rendre cette information plus facile à exploiter, la protéger et la valoriser. Elle pourra servir de base ensuite à des processus d'évaluation et d'aide à la décision.

Ce cours est organisé en 2 UEs :

Master 1 : Bases mathématiques et informatiques de la science des données.

Master 2 : Notions avancées d'intelligence artificielle et d'apprentissage artificiel.

Objectifs

Renforcer les compétences théoriques avec un approfondissement des statistiques et de la théorie des probabilités.

Introduire les méthodes et techniques d'apprentissage artificiel.

Former à l'utilisation de langages informatiques (Python et R) pour réaliser des projets permettant de mettre en pratique les aspects théoriques de la science des données.

Mots clés : probabilités, statistiques, base de données, apprentissage artificiel, R, python.

Prérequis : notions de base en probabilités, statistiques et programmation.

Moyens

Equipe pédagogique :



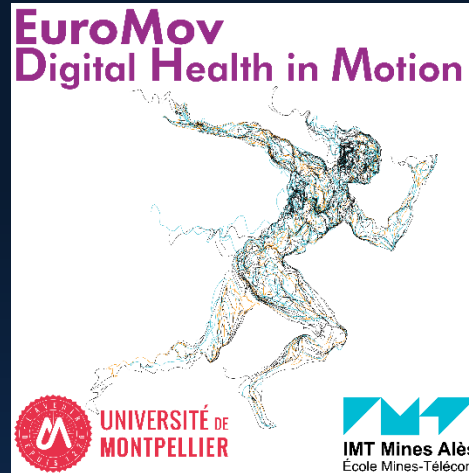
- Nicolas Sutton-Charani (NSC)
Enseignant Chercheur - EuroMov Digital Health in Motion, Univ Montpellier, IMT Mines Ales
- Gérard Dray (GD)
Enseignant Chercheur - EuroMov Digital Health in Motion, Univ Montpellier, IMT Mines Ales

Date	Heure début	Heure fin	Intervenant	Salle	Moyens pédagogiques
jeudi 22 janvier 2026	13:15:00	18:15:00	Gérard Dray	TD 36.414	Cours/TD
jeudi 29 janvier 2026	13:15:00	18:15:00	Nicolas Sutton-Charani	TD 36.416	Cours/TD
jeudi 5 février 2026	13:15:00	18:15:00	Nicolas Sutton-Charani	TD 36.412	Cours/TD
jeudi 12 février 2026	13:15:00	18:15:00	Nicolas Sutton-Charani	TD 36.315	Cours/TD
jeudi 26 mars 2026	13:15:00	18:15:00	Equipe pédagogique	TD 36.410	Projet
jeudi 2 avril 2026	13:15:00	18:15:00	Equipe pédagogique	TD 36.417	Projet
jeudi 9 avril 2026	13:15:00	18:15:00	Equipe pédagogique	TD 36.410	Projet
jeudi 16 avril 2026	13:15:00	18:15:00	Equipe pédagogique	TD 36.410	Projet

G rard Dray
gerard.dray@mines-ales.fr

Chercheur

UMR EuroMov Digital Health in Motion,
Univ Montpellier,
IMT Mines Ales, Ales,
France



Enseignant

IMT Mines Al s
D partement d'enseignement
Informatique et Intelligence Artificielle
2IA

Universit  de Montpellier
IAE

Master 1 E_Mkg

Master 1 MTD

Master 2 E-Mkg

Master 2 MPW

Facult  des Sciences, Facult  de M decine
Master Sciences et Num rique pour la Sant 
Master Sant 

UFR STAPS
Masters STAPS

Quick start poll

Scale: 0 = none, 1 = basic awareness, 2 = beginner, 3 = intermediate, 4 = advanced

1. Programming (Python or similar)
2. Math & stats basics for Machine Learning
3. Machine Learning fundamentals (train/test split, overfitting, metrics)
4. Supervised learning (classification & regression)
5. Unsupervised learning (clustering)

What I am / What I am not

What I am:

- A support:
 - for ML concepts, methods, and good practices
 - for problem framing (what to predict, what data to use, how to evaluate)
 - for interpretation and critical thinking (limits, bias, generalization)
 - for clear reporting (results, metrics, conclusions)

What I am not:

- A code debugger or IT support (I won't fix your environment/setup issues)
- A “do-it-for-you” solution: you must run experiments and make choices
- A guarantee of performance: ML results depend on data and validation
- A replacement for scientific reasoning.

Key message:

**I help you learn to build and trust ML models,
you remain responsible for implementation and validation.**

Course positioning

Data Science for Health Applications

- Turning clinical, physiological, and behavioral data into reliable decision support
- From measurements to models: prediction, classification, and risk stratification
- Focus on methods you can trust: evaluation, interpretability, and limitations

Typical health applications

- Patient monitoring (wearables, vitals, time-series)
- Diagnosis support & screening
- Prognosis and outcome prediction
- Rehabilitation and personalized follow-up
- Medical imaging

Key message:

Build models that are not only accurate, but clinically meaningful and validated.

What is Data Science?

- Data Science is a discipline that :
 - Turns raw data into actionable knowledge
 - Combines statistics, artificial intelligence, programming, and domain expertise
 - Covers the full pipeline: collect, clean, explore, model, communicate
- Why it matters in Health
 - Data are often noisy, heterogeneous, incomplete, and context-dependent
 - Good decisions require reliable data + rigorous evaluation, not only algorithms
 - Most effort is upstream: data quality drives model quality

Key message:

Key message: Better data beats better algorithms

What is Artificial Intelligence?

- Artificial Intelligence (AI): a broad umbrella
- Systems that perform tasks typically requiring human intelligence:
 - pattern recognition
 - decision-making
 - perception (vision, signals)
 - language understanding
- AI ranges from rule-based systems to learning-based approaches

Key message:

In this course, we focus on the AI family that learns from examples: Machine Learning.

What is Machine Learning?

- Machine Learning (ML): learning patterns from examples
- Models learn a mapping from inputs to outputs using data
- Typical goals:
 - prediction (e.g., injury risk, fatigue state)
 - classification (e.g., movement quality categories)
 - pattern discovery (e.g., movement strategies)
- ML is not magic: performance depends on data quality, design choices, and evaluation

Key message:

ML is a methodology to build models that generalize beyond observed examples.

Typical Health use cases

Examples in sport, health, and ergonomics:

- Performance & training: workload monitoring, technique assessment
- Injury prevention: screening, risk indicators, load–response modeling
- Rehabilitation: progress tracking, return-to-activity decision support
- Ergonomics: task recognition, exposure estimation, risk detection
- Clinical & health: activity recognition, functional capacity estimation

Key message:

The same ML tools apply across domains
What changes is the question and constraints.

The ML workflow (big picture)

From measurements to a validated model:

- Define the problem & the target outcome
- Build a dataset (signals, labels, context)
- Preprocess and represent inputs (features or representations)
- Train a model + tune hyperparameters
- Evaluate properly (generalization, bias, robustness)
- Interpret results and consider deployment constraints

Key message:

Most ML failures come from the workflow (problem definition, leakage, evaluation),
not from the algorithm.

Course goals & mindset

What you will learn and how we will think

Core ML approaches:

- supervised learning (regression, classification)
- unsupervised learning (clustering, dimensionality reduction)

Evaluation & pitfalls:

overfitting, leakage, cross-validation, metrics

Responsible and critical perspective:

interpretability, bias, reliability, limitations

Practical orientation through examples

Key message:

The objective is not only to “run models”, but to build trustworthy and useful ones.



IMT Mines Alès
École Mines-Télécom



Master Sciences et Numérique pour la Santé

SCIENCE DES DONNÉES - NIVEAU 1

INTRODUCTION

Gérard Dray



Objectifs

Être capable de comprendre :

- la définition de la science des données
- les définitions de l'Intelligence Artificielle
- la méthodologie d'Apprentissage Artificiel

Initier une réflexion sur l'usage de la science des données dans le domaine de la santé numérique

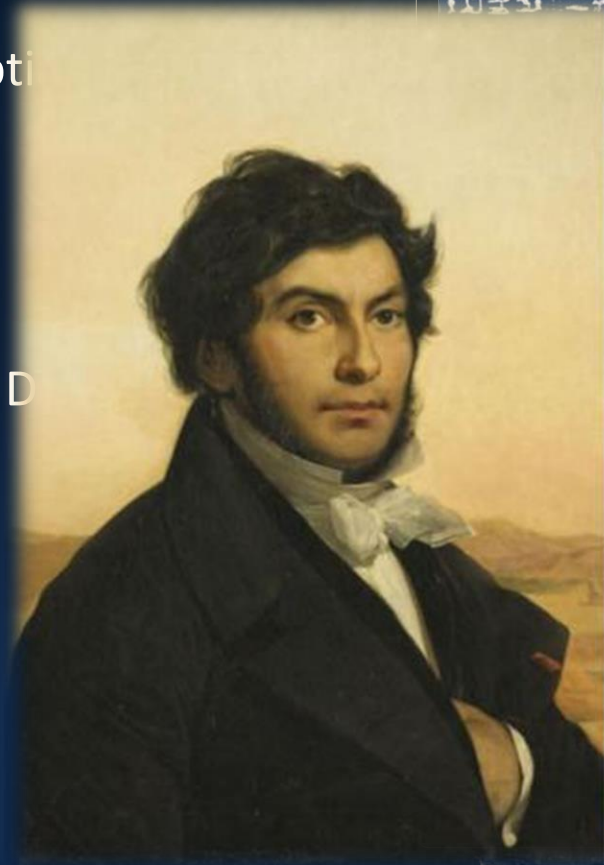
Programme :

- Une histoire ancienne
- Principales définitions
- Méthode d'apprentissage automatique : une illustration
- Exemples
- Discussion et Conclusion

Artificial Intelligence: An ancient story

1822 Jean-François Champollion

Egypti



Ramsès



Ra Mi SS

Ptolémée



Cléopâtre



"During the reign of the young man who succeeded his father, Lord of the Diadems, most glorious, who established Egypt and

triumphing over
ace and civilized
emonies of the
reat, a king like
d Lower Egypt,
Gods, approved
ctory, the living
PTOLEMY, WILL
F PTAH, in the
os of Aetos was
e Soteres Gods,
d the Euergetai
Gods, and the
od; daughter
Athlophoros of
nter Areia of
ros of Arsinoe
ne of Ptolemy
Philopator; the
of Xandikos,
he 18th Mekhir

...

1854 - John Snow

Cholera epidemic, London



Medical Societies.

MEDICAL SOCIETY OF LONDON.

MR. HEADLAND, PRESIDENT.

SATURDAY, OCTOBER 14TH, 1854.

Dr. Snow considered that the cholera poison acted upon the alimentary canal, and not on the blood or nervous system. In every case which he had seen, the evacuations had been sufficient to account for the collapse, without reference to any other cause. There was no poison in the blood in a case of cholera; in the consecutive fever, as it was called, the blood became poisoned from urea getting into the circulation in consequence of the kidneys not acting, but not from any poison having been present from the beginning. There was nothing in the atmosphere to account for the spread of cholera, which he believed was spread from person to person; and that in all cases it could be traced in this manner. If atmospheric, why did it attack one or two persons only in a locality, and these having direct communication with each other? Such cases he had seen at Sydenham, where there had been only two instances of the disease. The first case in the outbreak of 1849 had occurred to a sailor in Bermondsey; the second affected person was the successor to the sailor in the room in which he died. He thought he had collected evidence enough to show that in all cases cholera was propagated by swallowing some portion of the evacuations of an affected person. These, as was well known, flowed into the bed, &c., and persons attending on the sick might easily take the poison unawares. With respect to the class of persons affected by the disease, he believed that the very poor and vagabonds suffered less, in proportion, than decent, respectable persons. He regarded the cholera and diarrhoea, as lately prevalent, to be the same disease in different degrees of intensity. We observed the same difference in scarlatina and other diseases.

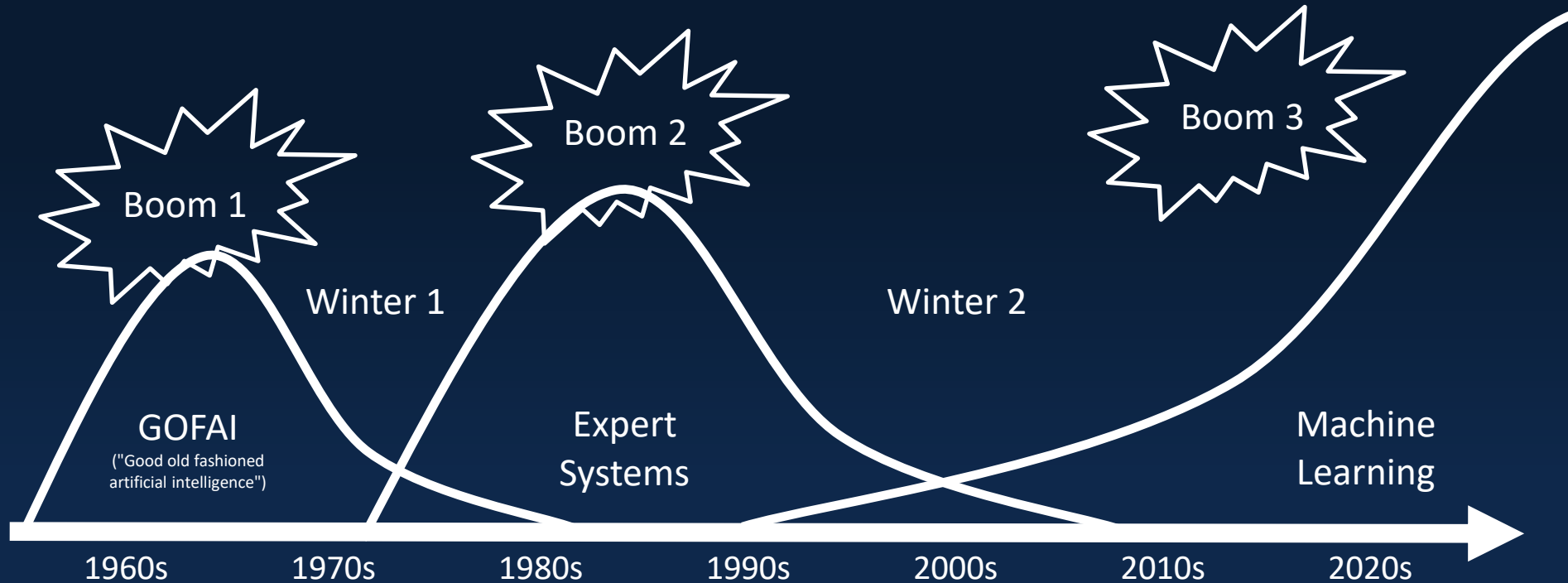


C. H. Duffell, Litho. Geographical Institute, Boston.

NOTE. Boundary within which all the deaths are indicated, is shown thus, Distances between Red Diagonals, thus, .

SCALE 30 INCHES TO A MILE.

History of Artificial Intelligence



1943 The formal neuron of McCulloch and Pitts

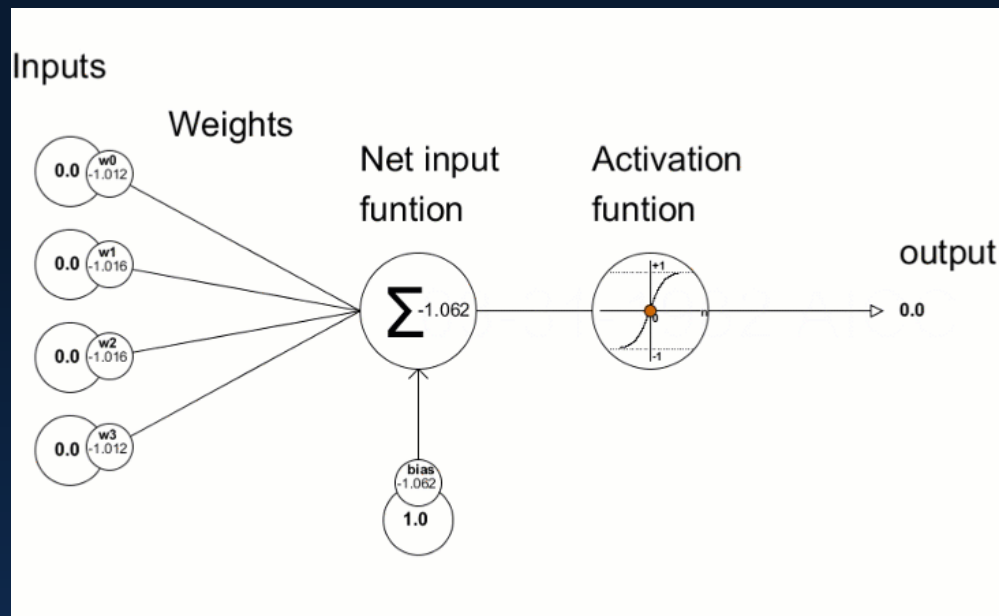


Warren Sturgis McCulloch
(1898 – 1969)

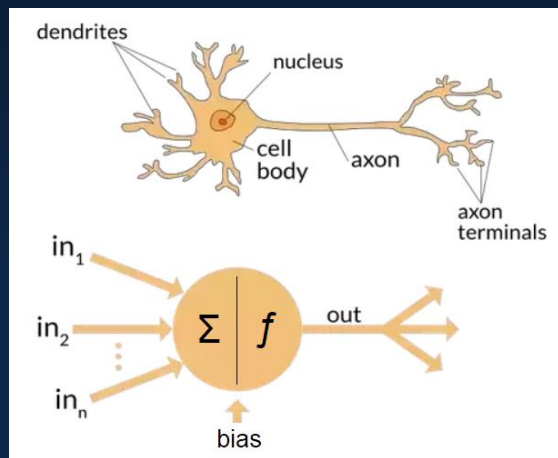


Walter Harry Pitts, Jr.
(1923 – 1969)

https://images.slideplayer.com/22/6379712/slides/slide_8.jpg



<https://www.mq15.com/fr/articles/5486>



<https://towardsdatascience.com/the-differences-between-artificial-and-biological-neural-networks-a8b46db828b7>

1956 The Dartmouth Conference

1956 Dartmouth Conference: The Founding Fathers of AI



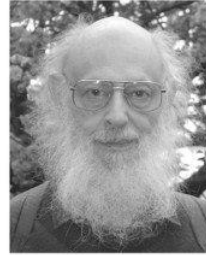
John MacCarthy



Marvin Minsky



Claude Shannon



Ray Solomonoff



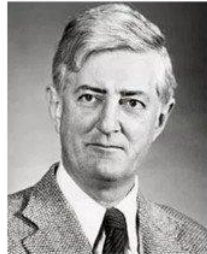
Alan Newell



Herbert Simon



Arthur Samuel



Oliver Selfridge



Nathaniel Rochester



Trenchard More

The Dartmouth Conference (Dartmouth Summer Research Project on Artificial Intelligence) was a scientific workshop held in the summer of 1956, which is considered to be the birth of artificial intelligence as a field of research.

1956 John McCarthy

Artificial Intelligence

“The science and engineering of making intelligent machines, especially intelligent computer programs”



1959 Arthur Samuel

“Machine learning is a field of study that gives computers the ability to learn without being explicitly programmed.”





“Instead of trying to produce a programme to simulate the adult mind, why not rather try to produce one which simulates the child's ? If this were then subjected to an appropriate course of education one would obtain the adult brain.”

DATA MINING

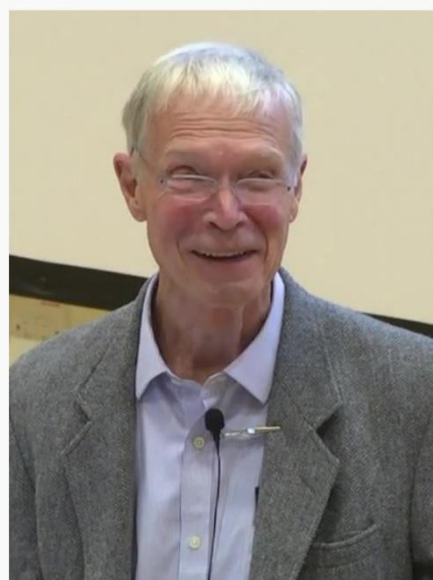


1997 The day Deep Blue defeated Garry Kasparov at chess

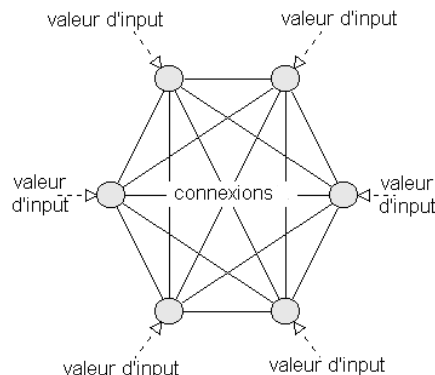


May 11, 1997, chess champion Garry Kasparov loses the sixth game of a historic match.

1982 John Joseph Hopfield



John Hopfield en 2016



Neural networks and physical systems with emergent collective computational abilities

(associative memory/parallel processing/categorization/content-addressable memory/fail-safe devices)

J. J. HOPFIELD

Division of Chemistry and Biology, California Institute of Technology, Pasadena, California 91125; and Bell Laboratories, Murray Hill, New Jersey 07974

Contributed by John J. Hopfield, January 15, 1982

ABSTRACT Computational properties of use to biological organisms or to the construction of computers can emerge as collective properties of systems having a large number of simple equivalent components (or neurons). The physical meaning of content-addressable memory is described by an appropriate phase space flow of the state of a system. A model of such a system is given, based on aspects of neurobiology but readily adapted to integrated circuits. The collective properties of this model produce a content-addressable memory which correctly yields an entire memory from any subset of sufficient size. The algorithm for the time evolution of the state of the system is based on asynchronous parallel processing. Additional emergent collective properties include some capacity for generalization, familiarity recognition, categorization, error correction, and time sequence retention. The collective properties are only weakly sensitive to details of the modeling or the failure of individual devices.

Given the dynamical electrochemical properties of neurons and their interconnections (synapses), we readily understand schemes that use a few neurons to obtain elementary useful biological behavior (1-3). Our understanding of such simple circuits in electronics allows us to plan larger and more complex circuits which are essential to large computers. Because evolution has no such plan, it becomes relevant to ask whether the ability of large collections of neurons to perform "computational" tasks may in part be a spontaneous collective consequence of having a large number of interacting simple neurons.

In physical systems made from a large number of simple elements, interactions among large numbers of elementary components yield collective phenomena such as the stable magnetic orientations and domains in a magnetic system or the vortex patterns in fluid flow. Do analogous collective phenomena in a system of simple interacting neurons have useful "computational" correlates? For example, are the stability of memories, the construction of categories of generalization, or time sequential memory also emergent properties and collective in origin? This paper examines a new modeling of this old and fundamental question (4-8) and shows that important computational properties spontaneously arise.

All modeling is based on details, and the details of neuroanatomy and neural function are both myriad and incompletely known (9). In many physical systems, the nature of the emergent collective properties is insensitive to the details inserted in the model (e.g., collisions are essential to generate sound waves, but any reasonable interatomic force law will yield appropriate collisions). In the same spirit, I will seek collective properties that are robust against change in the model details.

The model could be readily implemented by integrated circuit hardware. The conclusions suggest the design of a delo-

calized content-addressable memory or categorizer using extensive asynchronous parallel processing.

The general content-addressable memory of a physical system

Suppose that an item stored in memory is "H. A. Kramers & C. H. Wannier *Phys. Rev.* 60, 252 (1941)." A general content-addressable memory would be capable of retrieving this entire memory item on the basis of sufficient partial information. The input "Kramers, (1941)" might suffice. An ideal memory could deal with errors and retrieve this reference even from the input "Vannier, (1941)". In computers, only relatively simple forms of content-addressable memory have been made in hardware (10, 11). Sophisticated ideas like error correction in accessing information are usually introduced as software (10).

There are classes of physical systems whose spontaneous behavior can be used as a form of general (and error-correcting) content-addressable memory. Consider the time evolution of a physical system that can be described by a set of general coordinates. A point in state space then represents the instantaneous condition of the system. This state space may be either continuous or discrete (as in the case of N Ising spins).

The equations of motion of the system describe a flow in state space. Various classes of flow patterns are possible, but the systems of use for memory particularly include those that flow toward locally stable points from anywhere within regions around those points. A particle with frictional damping moving in a potential well with two minima exemplifies such a dynamics.

If the flow is not completely deterministic, the description is more complicated. In the two-well problems above, if the frictional force is characterized by a temperature, it must also produce a random driving force. The limit points become small limiting regions, and the stability becomes not absolute. But as long as the stochastic effects are small, the essence of local stable points remains.

Consider a physical system described by many coordinates $X_1, \dots, X_N$, the components of a state vector X . Let the system have locally stable limit points $X_1, X_2, \dots$. Then, if the system is started sufficiently near any X_0 , as at $X = X_0 + \Delta$, it will proceed in time until $X \approx X_0$. We can regard the information stored in the system as the vectors $X_1, X_2, \dots$. The starting point $X = X_0 + \Delta$ represents a partial knowledge of the item X_0 , and the system then generates the total information X_0 .

Any physical system whose dynamics in phase space is dominated by a substantial number of locally stable states to which it is attracted can therefore be regarded as a general content-addressable memory. The physical system will be a potentially useful memory if, in addition, any prescribed set of states can readily be made the stable states of the system.

The model system

The processing devices will be called neurons. Each neuron i has two states like those of McCulloch and Pitts (12): $V_i = 0$

The publication costs of this article were defrayed in part by page charge payment. This article must therefore be hereby marked "advertisement" in accordance with 18 U. S. C. § 1734 solely to indicate this fact.

2000 Machine Learning



Chat

Cat

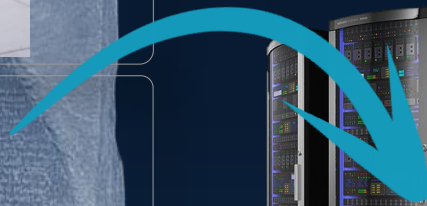


2000 Machine Learning

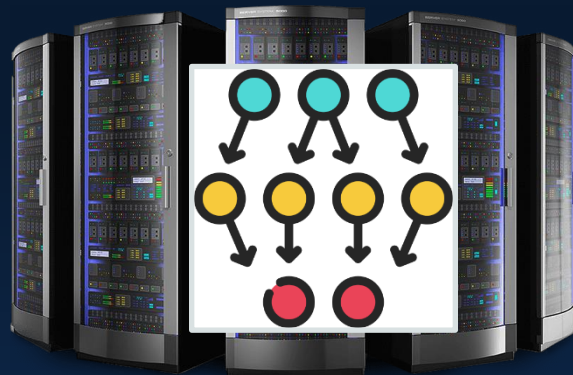


Chien

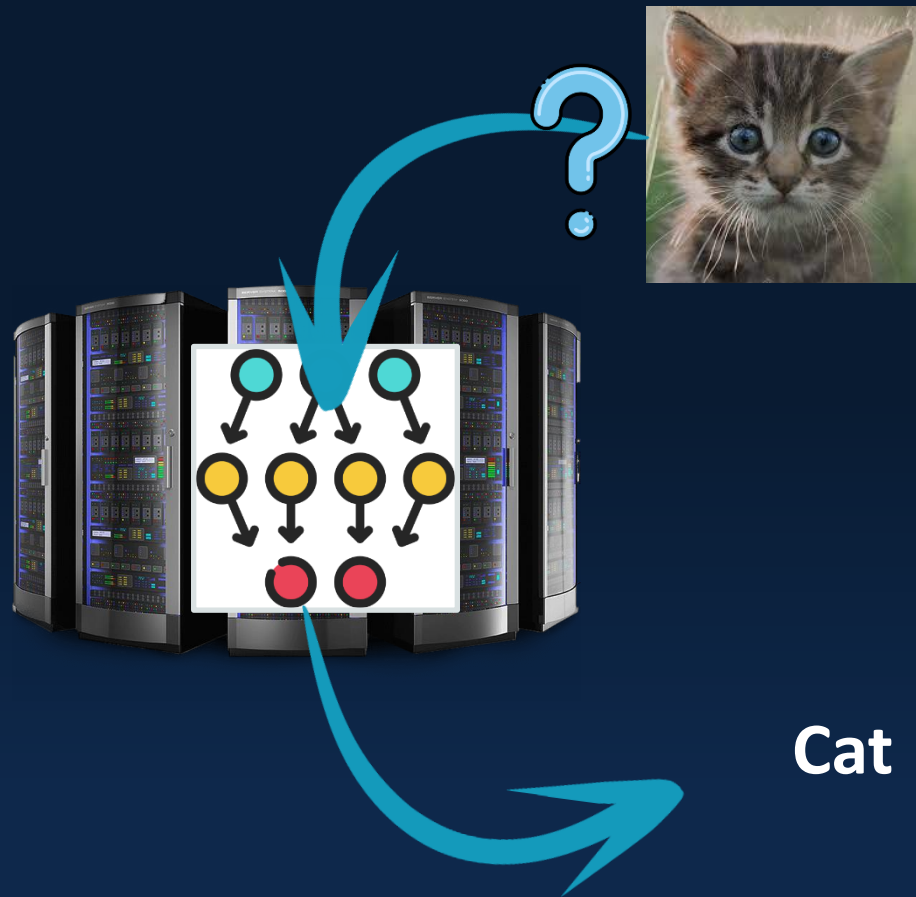
Dog



2000 Machine Learning



2000 Machine Learning

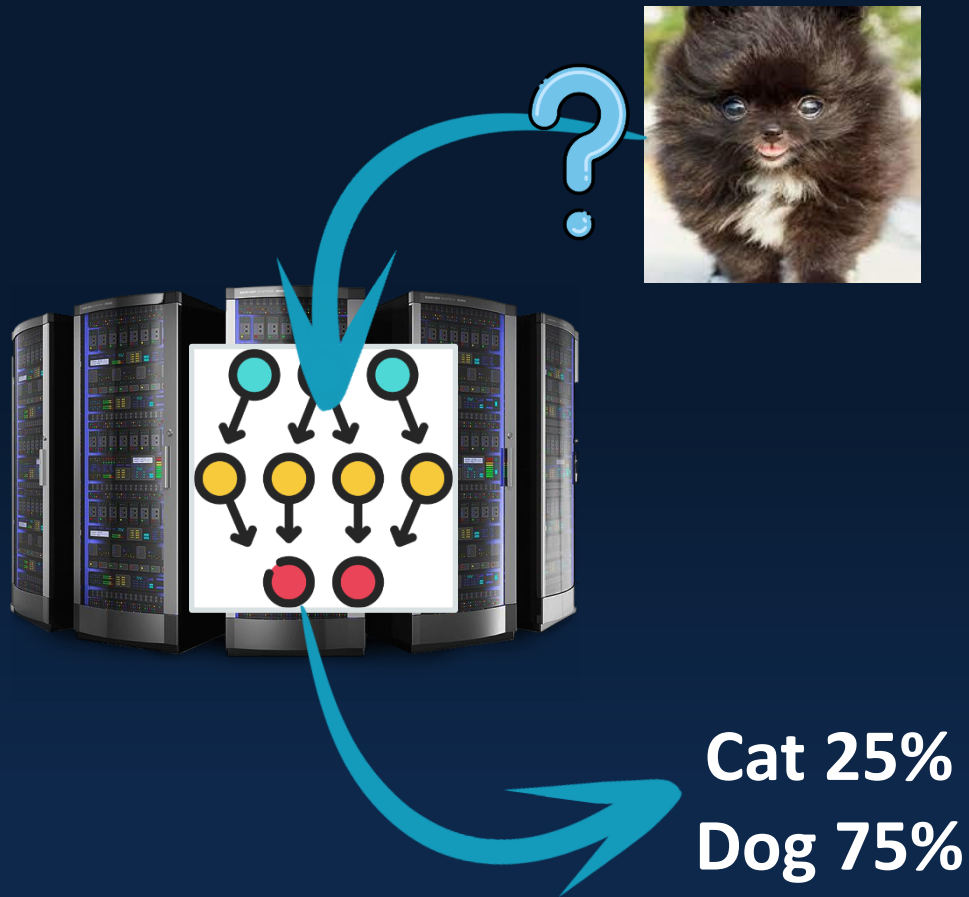


Cat

2000 Machine Learning



2000 Machine Learning

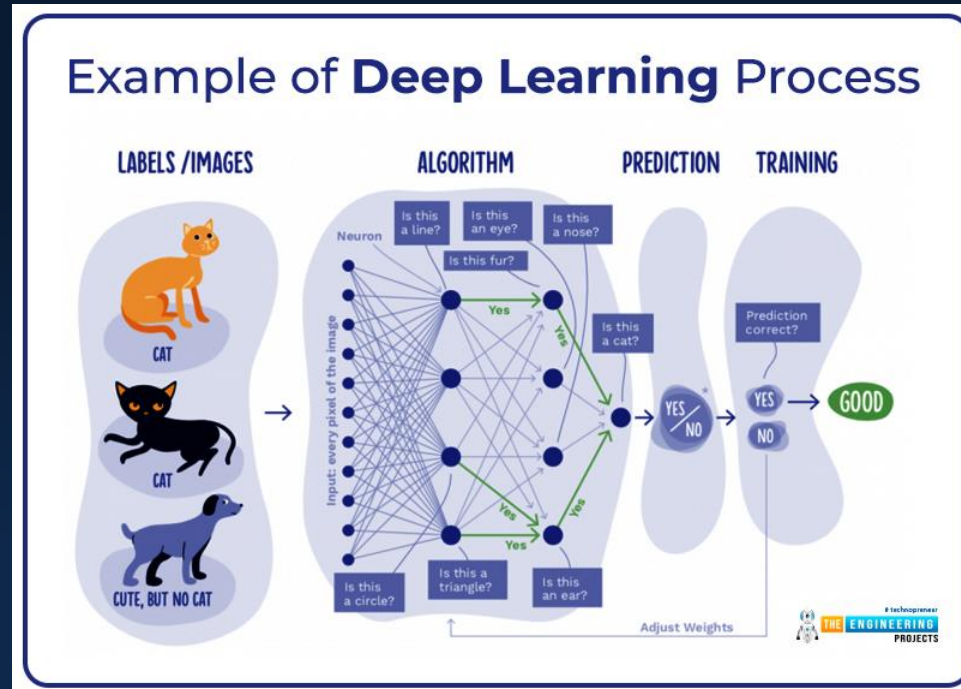


2005 Deep Learning

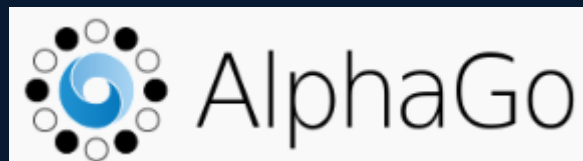


Geoffrey Everest Hinton (born 6 December 1947) is a British-Canadian computer scientist, cognitive scientist, and cognitive psychologist known for his work on artificial neural networks, which earned him the title "the Godfather of AI".

https://en.wikipedia.org/wiki/Geoffrey_Hinton



Deep Learning Apprentissage Profond



MODERN DATA SCIENTIST

Data Scientist, the sexiest job of 21st century requires a mixture of multidisciplinary skills ranging from an intersection of mathematics, statistics, computer science, communication and business. Finding a data scientist is hard. Finding people who understand who a data scientist is, is equally hard. So here is a little cheat sheet on who the modern data scientist really is.

MATH & STATISTICS

- ☆ Machine learning
- ☆ Statistical modeling
- ☆ Experiment design
- ☆ Bayesian inference
- ☆ Supervised learning: decision trees, random forests, logistic regression
- ☆ Unsupervised learning: clustering, dimensionality reduction
- ☆ Optimization: gradient descent and variants

DOMAIN KNOWLEDGE & SOFT SKILLS

- ☆ Passionate about the business
- ☆ Curious about data
- ☆ Influence without authority
- ☆ Hacker mindset
- ☆ Problem solver
- ☆ Strategic, proactive, creative, innovative and collaborative



PROGRAMMING & DATABASE

- ☆ Computer science fundamentals
- ☆ Scripting language e.g. Python
- ☆ Statistical computing package e.g. R
- ☆ Databases SQL and NoSQL
- ☆ Relational algebra
- ☆ Parallel databases and parallel query processing
- ☆ MapReduce concepts
- ☆ Hadoop and Hive/Pig
- ☆ Custom reducers
- ☆ Experience with xaaS like AWS

COMMUNICATION & VISUALIZATION

- ☆ Able to engage with senior management
- ☆ Story telling skills
- ☆ Translate data-driven insights into decisions and actions
- ☆ Visual art design
- ☆ R packages like ggplot or lattice
- ☆ Knowledge of any of visualization tools e.g. Flare, D3.js, Tableau

Google Actualités

google le dépistage cancer sein

Couverture complète

Comment les articles sont-ils classés ?

Top des articles

IA de Google dépiste le cancer du sein avec plus de précision que les médecins, selon cette étude
 13/01/2020
 13/01/2020

Alphabet (google) : Met à l'honneur l'IA en matière de dépistage
 13/01/2020
 13/01/2020

Google : met à l'honneur l'IA en matière de dépistage
 13/01/2020
 13/01/2020

Twitter

Joli de Rosnay
 @jolirosnay
 Des chercheurs de Google Health à Palo Alto et de Deep Mind à Londres, mettent au point un test utilisant l'intelligence artificielle qui détecte les tumeurs précoces du cancer du sein mieux que les experts <https://fr.fox.com/health/ai>
 Twitter • 13/01/2020 18:40

Tous les articles

Cancer du sein : Une intelligence artificielle plus précise que des humains pour établir un diagnostic
 13/01/2020
 13/01/2020

Google : met à l'honneur l'IA en matière de dépistage
 13/01/2020
 13/01/2020

Langue et zone géographique (Français (France))

Paramètres

Télécharger l'application pour Android

Télécharger l'application pour iOS

Envoyer des commentaires

Aide

Confidentialité Conditions d'utilisation

Logos de Google

LE QUOTIDIEN DU MEDECIN

TOUS LES ACTUALITÉS MÉDICALES
 SOUS PROFESSIONNELLE

Rechercher

S'ABONNER DÈS 10€

À LA UNE SANTÉ LIBÉRAL HÔPITAL SPÉCIALITÉS INTERNES ANNONCES / EMPLOI

EN CE MOMENT PUFIS 2020 Cite des urgences Loi de bioéthique

Accueil / Spécialités / Cancérologie

IA et dépistage du cancer du sein : un algorithme de Google Health fait mieux que le radiologue

PAR CHARLENE CATALA-ROD - PUBLIÉ LE 06/01/2020

2 RÉACTIONS COMMENTER

ÉDITION ABONNÉS DU 2020-01-08

LIRE LE JOURNAL

SCIENCES Avenir

Alimentaire Cancer Climat & air Eau & énergie Santé Société Technologie

Incendies en Australie Les bienfaits de la méditation

Une IA de Google surpasse les radiologues pour détecter le cancer du sein

Par Charles Lamber - 10.01.2020 à 12:00

Google a mis au point une intelligence artificielle qui voit mieux respirer les tumeurs cancéreuses sur les mammographies que les radiologues. Ce n'est pas la première fois qu'un algorithme bat les humains pour détecter les tumeurs sur une image.

Google a mis au point une intelligence artificielle qui voit mieux respirer les tumeurs cancéreuses sur les mammographies que les radiologues. Ce n'est pas la première fois qu'un algorithme bat les humains pour détecter les tumeurs sur une image.

En 2018, la communication de précision a fortement augmenté en France

Il y a 170.000 ans, nos ancêtres vivaient dans les grottes

Depuis 2013, chaque semaine nous avons pu voir la météo de la région météorologique

Les premiers premiers tests à des centaines de patients

RECEVOIR VOS ARTICLES

Entrez votre Email

Je m'abonne

Article

International evaluation of an AI system for breast cancer screening

<https://doi.org/10.1038/s41586-019-1799-6>

Received: 27 July 2019

Accepted: 5 November 2019

Published online: 1 January 2020

Scott Mayer McKinney^{1,2*}, Marcin Sieniek^{1,2,4}, Varun Godbole^{1,2,4}, Jonathan Godwin^{1,2,4}, Natasha Antropova³, Hutan Ashtafan⁴, Trevor Back⁴, Mary Chesus¹, Greg C. Corrado¹, Ara Darzi^{2,4,5}, Mozzayr Etemadi⁶, Florencia Garcia-Vicente⁶, Fiona J. Gilbert⁷, Mark Halling-Brown⁸, Demis Hassabis¹, Sunny Jansen⁹, Alan Karthikesalingam¹⁰, Christopher J. Kelly¹⁰, Dominic King¹⁰, Joseph R. Ledsam¹⁰, David Melnick¹⁰, Hormuz Mostofi¹¹, Lily Peng¹², Joshua Jay Reicher¹³, Bernardino Romera-Paredes¹⁴, Richard Siddhant^{12,15}, Mustafa Suleyman¹, Daniel Tse¹⁶, Kenneth C. Young¹⁶, Jeffrey De Fauw^{2,16} & Shravya Shetty^{2,16}

Screening mammography aims to identify breast cancer at earlier stages of the disease, when treatment can be more successful¹. Despite the existence of screening programmes worldwide, the interpretation of mammograms is affected by high rates of false positives and false negatives². Here we present an artificial intelligence (AI) system that is capable of surpassing human experts in breast cancer prediction. To assess its performance in the clinical setting, we curated a large representative dataset from the UK and a large enriched dataset from the USA. We show an absolute reduction of 5.7% and 1.2% (USA and UK) in false positives and 9.4% and 2.7% in false negatives. We provide evidence of the ability of the system to generalize from the UK to the USA. In an independent study of six radiologists, the AI system outperformed all of the human readers: the area under the receiver operating characteristic curve (AUC-ROC) for the AI system was greater than the AUC-ROC for the average radiologist by an absolute margin of 11.5%. We ran a simulation in which the AI system participated in the double-reading process that is used in the UK, and found that the AI system maintained non-inferior performance and reduced the workload of the second reader by 88%. This robust assessment of the AI system paves the way for clinical trials to improve the accuracy and efficiency of breast cancer screening.

Breast cancer is the second leading cause of death from cancer in women¹, but early detection and treatment can considerably improve outcomes^{2,3}. As a consequence, many developed nations have implemented large-scale mammography screening programmes. Major medical and governmental organizations recommend screening for all women starting between the ages of 40 and 50^{4–6}. In the USA and UK combined, over 42 million exams are performed each year^{7,8}. Despite the widespread adoption of mammography, interpretation of these images remains challenging. The accuracy achieved by experts in cancer detection varies widely, and the performance of even the best clinicians leaves room for improvement^{9,10}. False positives can lead to patient anxiety¹¹, unnecessary follow-up and invasive diagnostic procedures. Cancers that are missed at screening may not be identified until they are more advanced and less amenable to treatment¹².

AI may be uniquely poised to help with this challenge. Studies have demonstrated the ability of AI to meet or exceed the performance of human experts on several tasks of medical-image analysis^{13–19}.

As a shortage of mammography professionals threatens the availability and adequacy of breast screening services around the world^{20–22}, the scalability of AI could improve access to high-quality care for all.

Computer-aided detection (CAD) software for mammography was introduced in the 1990s, and several assistive tools have been approved for medical use²³. Despite early promise^{24,25}, this generation of software failed to improve the performance of readers in real-world settings^{12,26,27}. More recently, the field has seen a renaissance owing to the success of deep learning. A few studies have characterized systems for breast cancer prediction with stand-alone performance that approaches that of human experts^{28,29}. However, the existing work has several limitations. Most studies are based on small, enriched datasets with limited follow-up, and few have compared performance to readers in actual clinical practice—instead relying on laboratory-based simulations of the reading environment. So far there has been little evidence of the ability of AI systems to translate between different screening populations and settings without additional training data³⁰. Critically, the pervasive use of follow-up intervals that are no longer than 12 months^{29,30,31,32}

¹Google Health, Palo Alto, CA, USA. ²DeepMind, London, UK. ³Department of Surgery and Cancer, Imperial College London, London, UK. ⁴Institute of Global Health Innovation, Imperial College London, London, UK. ⁵Cancer Research UK Imperial Centre, Imperial College London, London, UK. ⁶Northwestern Medicine, Chicago, IL, USA. ⁷Department of Radiology, Cambridge Biomedical Research Centre, University of Cambridge, Cambridge, UK. ⁸Royal Surrey County Hospital, Guildford, UK. ⁹Norfolk and Norwich Healthcare, South Star, Norfolk, UK. ¹⁰Google Health, London, UK. ¹¹Stanford Health Care and Palo Alto Veterans Affairs, Palo Alto, CA, USA. ¹²The Royal Marsden Hospital, London, UK. ¹³Christiane Breast Cancer, Chesham, UK. These authors contributed equally: Scott Mayer McKinney, Marcin T. Sieniek, Varun Godbole, Jonathan Godwin. *These authors jointly supervised this work: Jeffrey De Fauw, Shravya Shetty. e-mail: scottmayer@google.com; tsed@google.com; shetty@google.com

<https://simulator.drdata.io/>

Calculez le prix de vos données !

Vos données personnelles ont un prix pour les entreprises commerciales dans le cadre des démarches publicitaires et de reventes de bases de données. Ces données s'échangent et se revendent entre ces acteurs, et ce partout dans le monde.

Nous avons cherché à en connaître les prix, non sans mal !

Notre message n'est pas de vous encourager à vendre vos données, mais plutôt de vous faire prendre conscience du marché qu'elles représentent.

Notre objectif est de vous sensibiliser à la valeur économique de vos données personnelles pour vous aider à mieux comprendre les enjeux de la protection de votre vie privée.

Votre navigation est anonyme et libre sans cookies et sans compte.

FR EN

Dr Data

Etat civil et Coordonnées

Les informations concernant votre identité tel que votre sexe, numéro de téléphone ou votre âge sont des données d'identification dites «classiques» qui peuvent déterminer un premier profil de votre personne..

Ajoutez des données à partager pour en connaître le prix

☒ Sexe	☐ Adresse Postale	☐ Numéro de Sécurité Sociale
☐ Age	☐ Adresse Email	☐ Catégorie Socio-Professionnelle
☐ Origine Ethnique	☒ Numéro de Téléphone	☐ Niveau d'Etude

Possédez-vous un casier judiciaire ?

☐ Oui
 ☒ Non

Avez vous un permis de conduire ?

☒ Oui
 ☐ Non

Total

Partager nos données c'est bien, le faire avec l'information complète c'est mieux !

17.0004 €

Le prix affiché est calculé en fonction des sources que nous avons pu récupérer et étudier, vous pouvez y accéder en [cliquant ici](#).

Les données ne sont pas exhaustives. Le résultat est le fruit de notre algorithme sur la base des informations disponibles.



Mentions légales & CGU

Copyright © 2022 DrData. Tous droits réservés.

En partenariat avec ...


IMT Mines Alès
 Ecole Mines-Télécom

Le Parisien

XV de France : «Ce ne sont pas les chiffres qui nous dirigent»... Comment Fabien Galthié utilise les datas

Le XV de France, qui affronte le Japon dimanche (14 heures) à Toulouse lors de son dernier match de la tournée d'automne, s'appuie dans sa préparation sur une multitude de données chiffrées, collectées, triées puis fournies par la société SAS France.



Fabien Galthié avant le match contre l'Afrique du Sud au Velodrome, à Marseille. Icon Sport/Johnny Fidelin

*Les chiffres, fournis notamment par une puce GPS incrustée dans les maillots, remontent et s'inscrivent dans une base de données où figurent 1438 matchs (XV de France, U20 et Top 14 compris). Les ballons connectés et les protège-dents connectés élargissent également le champ des investigations. « **Attention, ce ne sont pas les chiffres qui nous dirigent mais nous qui dirigeons les chiffres** », lance Fabien Galthié. D'ailleurs, sur un match, je ne suis connecté à rien. Je regarde, j'écoute, je suis au cœur de l'environnement, il n'y a pas un seul chiffre qui m'impacte. Les décisions, on les prend avec notre intuition et notre savoir-faire. » La science, elle, est là pour répondre aux interrogations.*

2022



OpenAI is an AI research and deployment company. Our mission is to ensure that artificial general intelligence benefits all of humanity.

OpenAI's mission is to ensure that artificial general intelligence (AGI)—by which we mean highly autonomous systems that outperform humans at most economically valuable work—benefits all of humanity.

We will attempt to directly build safe and beneficial AGI, but will also consider our mission fulfilled if our work aids others to achieve this outcome.

OpenAI Charter

We're releasing a charter that describes the principles we use to execute on OpenAI's mission. This document reflects the strategy we've refined over the past two years, including feedback from many people internal and external to OpenAI. The timeline to AGI remains uncertain, but our charter will guide us in acting in the best interests of humanity throughout its development.

Newsroom

Announcements

ChatGPT: Optimizing Language Models for Dialogue
November 30, 2022 — Announcements, Research

DALL-E API Now Available in Public Beta
November 3, 2022 — Announcements, API

DALL-E Now Available Without Waitlist
September 28, 2022 — Announcements

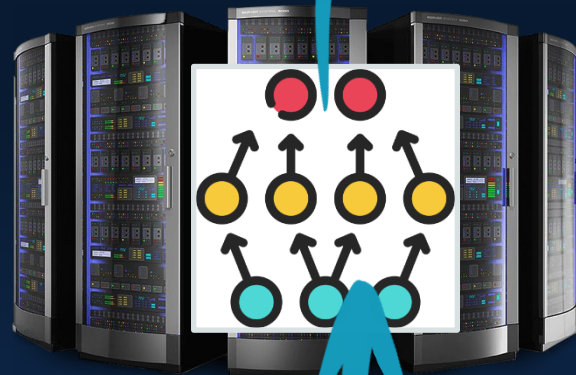
<https://chat.openai.com/chat>

<https://openai.com/dall-e-2/>

2022 Generative Artificial Intelligence



2022 Generative Artificial Intelligence



Dog

2024

THE NOBEL PRIZE IN PHYSICS 2024

Illustrations: Niklas Elmehed



John J. Hopfield

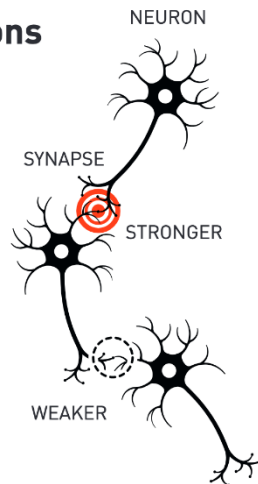
Geoffrey E. Hinton

"for foundational discoveries and inventions
that enable machine learning
with artificial neural networks"

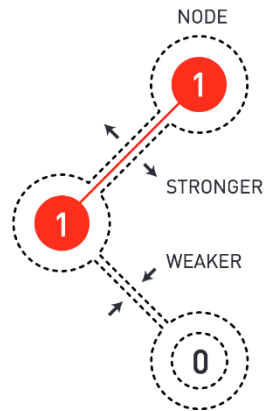
THE ROYAL SWEDISH ACADEMY OF SCIENCES

Natural and artificial neurons

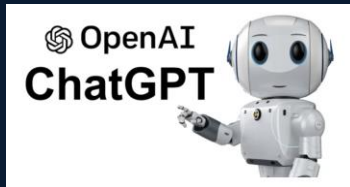
The brain's neural network is built from living cells, neurons, with advanced internal machinery. They can send signals to each other through the synapses. When we learn things, the connections between some neurons get stronger, while others get weaker.



Artificial neural networks are built from nodes that are coded with a value. The nodes are connected to each other and, when the network is trained, the connections between nodes that are active at the same time get stronger, otherwise they get weaker.



©Johan Jarnestad/The Royal Swedish Academy of Sciences



Agent conversationnel

You

Pourrais-tu rédiger une brève introduction pour un cours sur l'intelligence artificielle appliquée à l'ingénierie biomédicale ?

ChatGPT

Bien sûr ! Voici une proposition de texte pour introduire un cours sur l'intelligence artificielle appliquée à l'ingénierie biomédicale :

****Introduction au cours : Intelligence Artificielle en Ingénierie Biomédicale****

Bienvenue dans ce cours passionnant où nous explorerons les intersections innovantes entre l'intelligence artificielle (IA) et l'ingénierie biomédicale. Ce domaine en pleine expansion utilise les avancées de l'IA pour transformer les soins de santé, en améliorant les diagnostics, en optimisant les traitements et en révolutionnant la gestion des maladies.

Au cours de ce programme, nous aborderons les principes fondamentaux de l'IA, y compris l'apprentissage automatique, le traitement du langage naturel et la vision par ordinateur, et comment ces technologies peuvent être appliquées spécifiquement dans le contexte biomédical. Nous étudierons des études de cas réelles montrant l'impact de l'IA sur la création de dispositifs médicaux intelligents, le développement de techniques de diagnostic précoces et l'élaboration de solutions personnalisées pour le traitement des patients.

Ce cours est conçu pour fournir aux étudiants une compréhension profonde des défis techniques, éthiques et réglementaires associés à l'intégration de l'IA dans les applications biomédicales. En combinant théorie et pratique, nous préparerons les étudiants à devenir les pionniers de cette révolution technologique en santé.

Préparez-vous à plonger dans un monde où la technologie et la médecine convergent pour améliorer la qualité de vie des personnes à l'échelle mondiale. Ensemble, nous découvrirons comment l'intelligence artificielle peut être utilisée pour résoudre certains des problèmes les plus urgents en médecine aujourd'hui.

J'espère que cette introduction vous conviendra pour votre cours !



MIDJOURNEY
All about imagination

Midjourney Bot

/Imagine

Ultra realistic photo of a class of students
taking a machine-learning workshop and
listening to a human teacher



[Consulter le journal](#)

G. Dray ▾

[🏠](#) [Actualités ▾](#) [Économie ▾](#) [Vidéos ▾](#) [Débats ▾](#) [Culture ▾](#) [Le Goût du Monde ▾](#) [Services ▾](#) [🔍](#)

ÉCONOMIE • L'INVESTITURE DE DONALD TRUMP

Le second mandat de Donald Trump commence par un investissement de centaines de milliards de dollars pour l'intelligence artificielle

Le nouveau président américain a dévoilé, mardi 21 janvier, le projet Stargate, destiné à bâtir les centres de données géants de la future génération d'IA, élaboré par Oracle, OpenAI et SoftBank. Soit 100 milliards de dollars investis tout de suite, portés à 500 milliards de dollars, d'ici à 2029.

Par Arnaud Leparmentier (New York, correspondant)

Publié aujourd'hui à 07h46, modifié à 08h30 • 🕒 Lecture 3 min.

[🎁 Offrir l'article](#) [🔖](#) [➦](#)

Le Monde

SCIENCE & MÉDECINE - SUPPLÉMENT VIVRE AVEC UN CŒUR OU UN REIN DE COCHON, C'EST POUR BIENTÔT ?

IA : le chinois DeepSeek ébranle le secteur

- La start-up chinoise a rendu public un modèle aux performances comparables aux leaders OpenAI ou Google, mais à un coût beaucoup moins élevé
- Pour expliquer sa suprématie relative, DeepSeek dit avoir utilisé plusieurs techniques innovantes pour optimiser la quantité de calculs
- Cette annonce a provoqué un véritable séisme sur les marchés financiers, faisant éclater la bulle de l'intelligence artificielle qui porte le Nasdaq
- Lundi, le fabricant de semi-conducteurs chinois, qui valait 3 500 milliards de dollars, a vu sa valeur perdue s'envoler en une journée
- Aux observateurs qui appellent à la prudence, les responsables de DeepSeek rappellent que leurs modèles sont « ouverts »

CINÉMA BOB DYLAN IMPRIME SA LÉGENDE À L'ÉCRAN

- « Un parfait incertain » : biopic de James Mangold, une merveille en salle
- Timothée Chalamet incarne le chanteur, Monica Barbaro « une Joan Baez »

RDC Les forces du M23 et du Rwanda sont entrées dans Goma

ÉDITORIAL
DÉSIGNER CLAIREMENT LES PROTAGONISTES

Politique
François Bayrou se tourne vers sa droite pour éviter la censure

Agriculture
La Coordination rurale dans l'ombre de l'extrême droite

Reportage
A bas bruit, des Israéliens choisissent de s'exiler

Commemoration
A Auschwitz, la mise en garde des derniers survivants

Mayotte
Après le cyclone Chido, une rentrée scolaire très perturbée

Crise climatique
La transition écologique, nouvelle cible des populistes

Reportage
Des millions d'Israéliens ont quitté leur pays, en silence et à la guerre, mais d'autres espèrent la fin du conflit de Gaza

Idées
Le temps de la « vassalisation heureuse » ?

LE BILAN DU MONDE 2025

Le Monde

GD G. Dray

Actualités Économie Vidéos Débats Culture Le Goût du Monde Services

La start-up chinoise DeepSeek bouleverse le secteur de l'intelligence artificielle

En développant DeepSeek-R1, un modèle aussi performant que ceux des leaders américains comme OpenAI ou Google, mais avec moins de ressources et en open source, DeepSeek fait vaciller la Silicon Valley.

Par Alexandre Piquard
Publié hier à 20h01, modifié à 09h29 • Lecture 4 min.

Offrir l'article

Implementing the EU AI Act in Radiology: ESR's Recommendations

In EXEC Tue, 18 Mar 2025

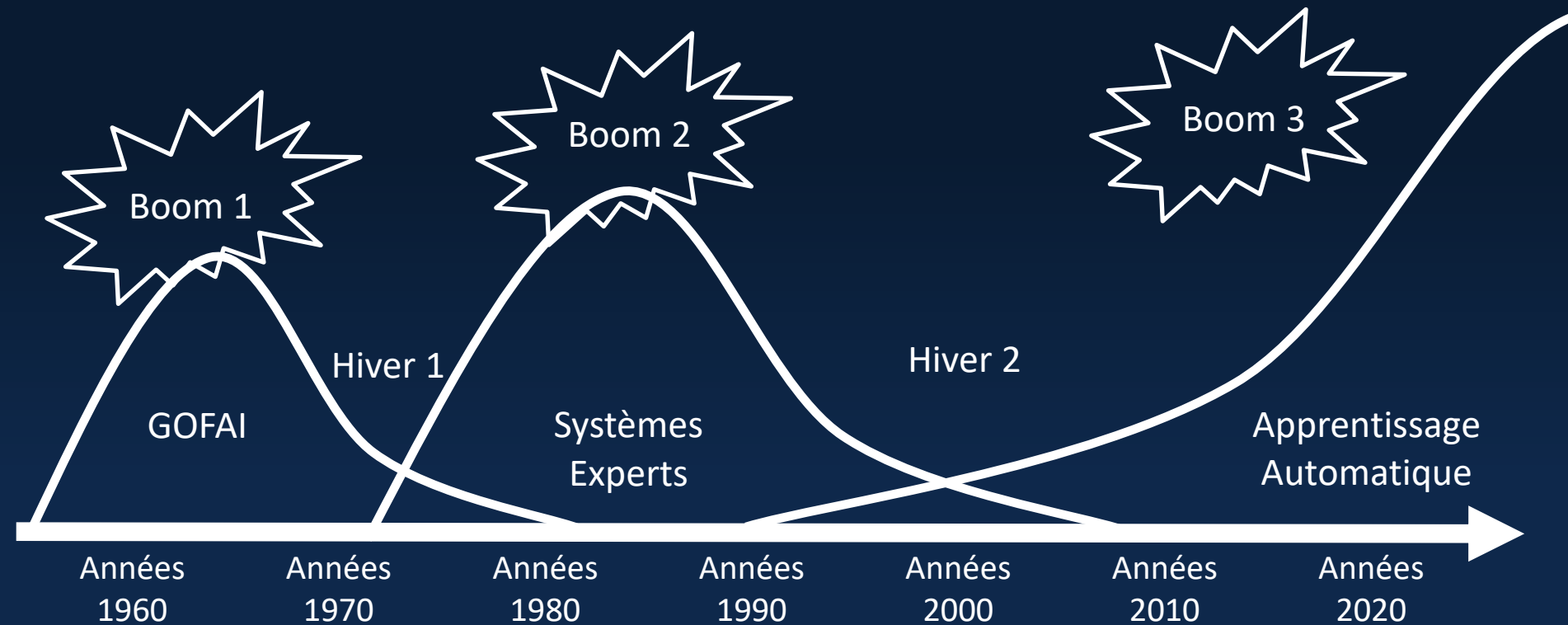


AI in Lung Cancer Screening: Enhancing Accuracy and Reducing Workload

In Imaging Mon, 17 Mar 2025

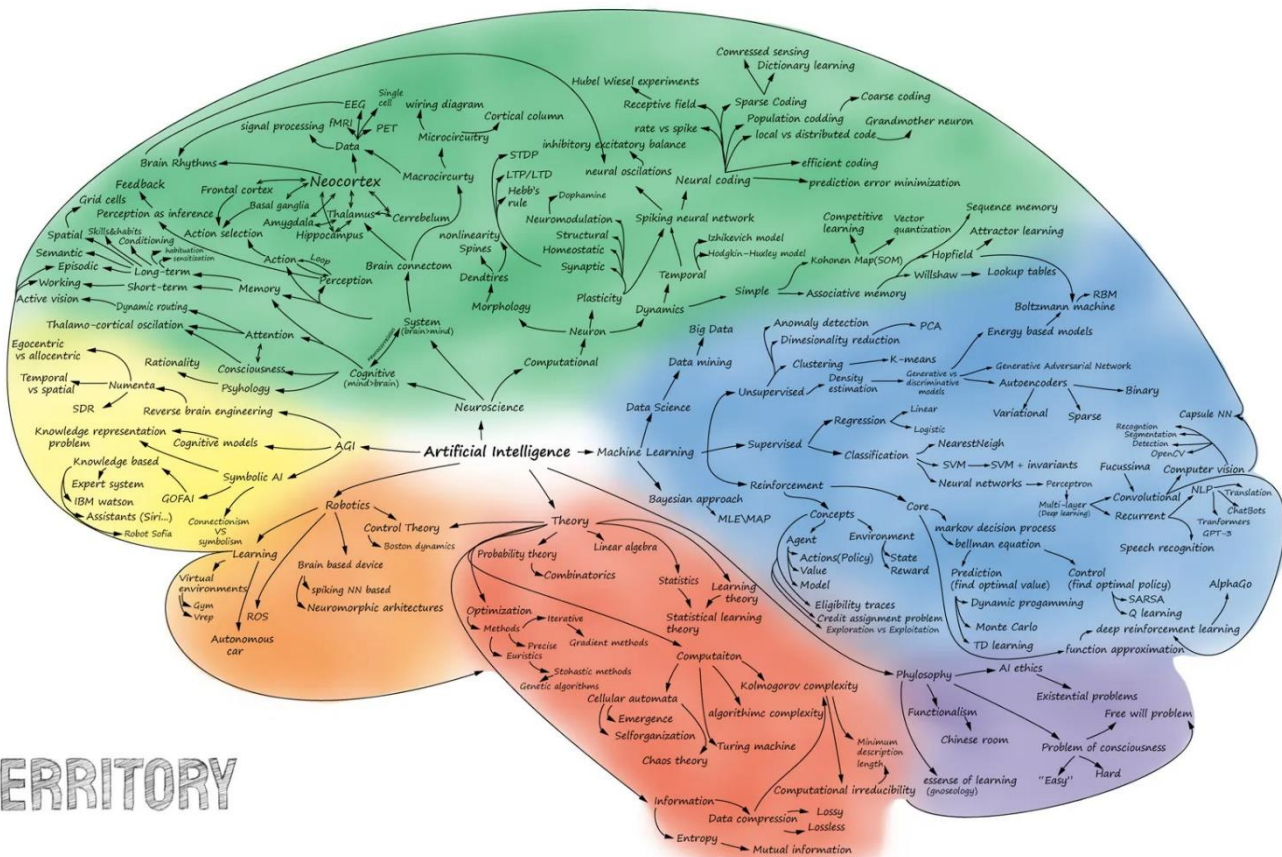


Historique de l'Intelligence Artificielle



Key definitions

The Map of Artificial Intelligence



What is Artificial Intelligence (AI)?

Artificial intelligence refers to any tool used by a machine to “reproduce human-like behaviors, such as reasoning, planning, and creativity.”

What is machine learning (ML)?

Machine learning is a field of artificial intelligence that aims to give machines the ability to “learn” from data using mathematical models. More specifically, it is the process by which relevant information is extracted from a training data set. The goal of this phase is to obtain the parameters of a model that will achieve the best performance, particularly when performing the task assigned to the model. Once learning is complete, the model can then be deployed in production.

What is deep learning (DL)?

Deep learning is a subset of machine learning. Deep learning is a machine learning process that uses neural networks with multiple layers of hidden neurons. These algorithms have a large number of parameters, so they require a very large amount of data in order to be trained.

Qu'est-ce que l'intelligence artificielle? (Artificial Intelligence)

L'intelligence artificielle représente tout outil utilisé par une machine afin de « reproduire des comportements liés aux humains, tels que le raisonnement, la planification et la créativité ».

Qu'est-ce que l'apprentissage automatique ? (Machine Learning)

L'apprentissage automatique est un champ d'étude de l'intelligence artificielle qui vise à donner aux machines la capacité d'« apprendre » à partir de données, via des modèles mathématiques. Plus précisément, il s'agit du procédé par lequel les informations pertinentes sont tirées d'un ensemble de données d'entraînement.

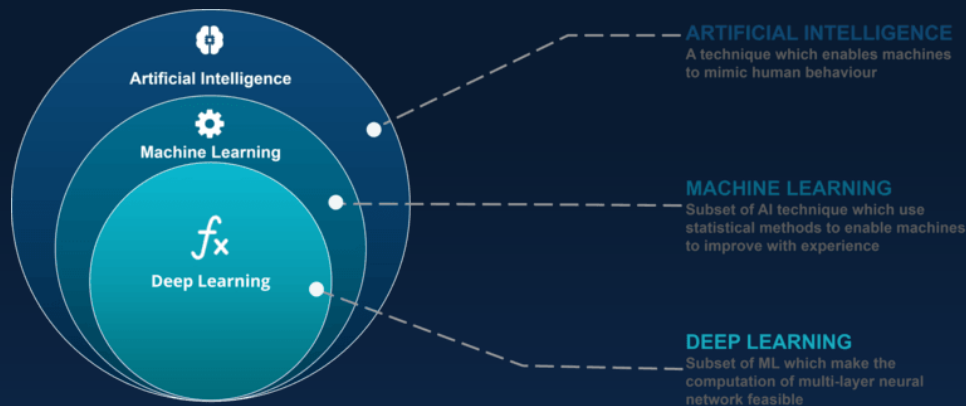
Le but de cette phase est l'obtention des paramètres d'un modèle qui atteindront les meilleures performances, notamment lors de la réalisation de la tâche attribuée au modèle. Une fois l'apprentissage réalisé, le modèle pourra ensuite être déployé en production.

Qu'est-ce que l'apprentissage profond ? (Deep Learning)

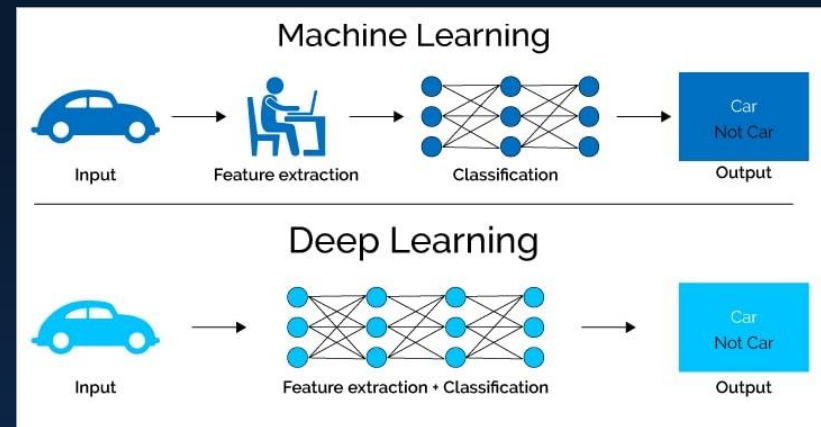
L'apprentissage profond est un sous-ensemble de l'apprentissage automatique. L'apprentissage profond est un procédé d'apprentissage automatique utilisant des réseaux de neurones possédant plusieurs couches de neurones cachées. Ces algorithmes possédant de très nombreux paramètres, ils demandent un nombre très important de données afin d'être entraînés.

Qu'est-ce que l'apprentissage automatique ? (Machine Learning)

Qu'est-ce que l'apprentissage profond ? (Deep Learning)



<https://datawider.com/how-deep-learning-is-different-from-machine-learning/>



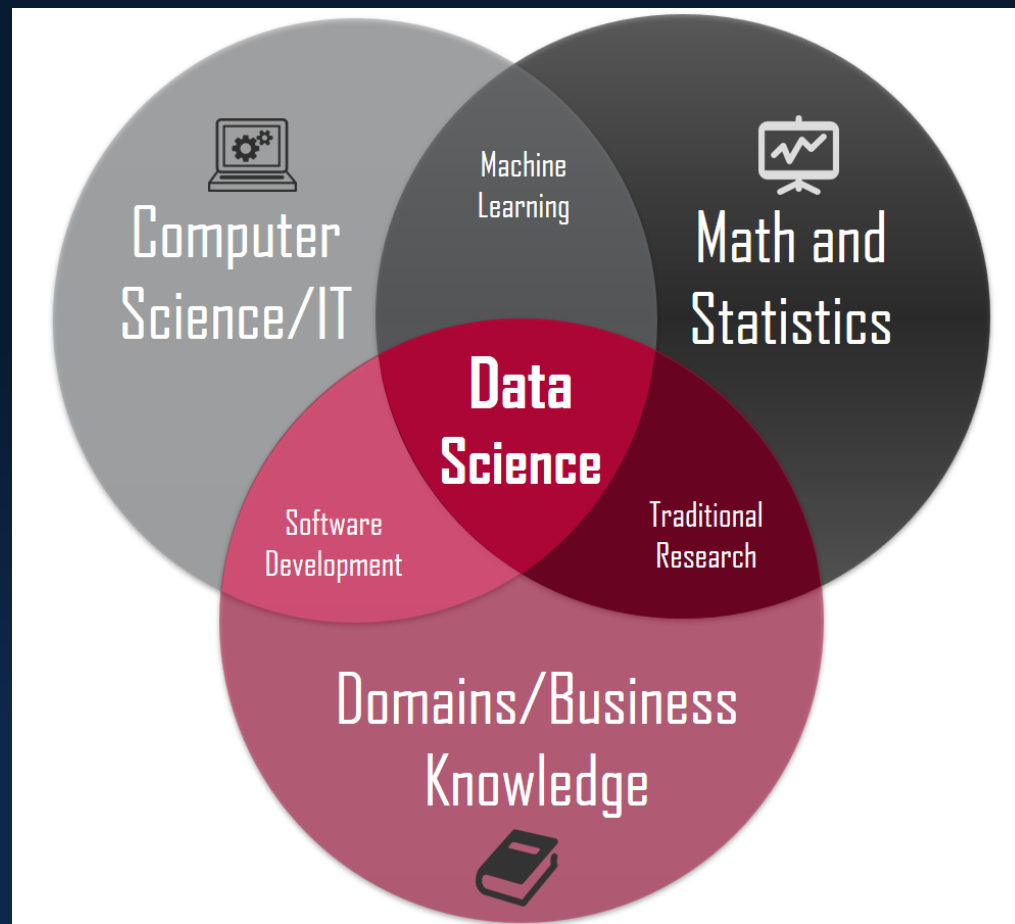
<https://levity.ai/blog/difference-machine-learning-deep-learning>

Qu'est-ce que la science des données ?

La science des données est l'étude de l'extraction automatisée de connaissance à partir de grands ensembles de données.

Plus précisément, la science des données est un domaine interdisciplinaire qui utilise des méthodes, des processus, des algorithmes et des systèmes scientifiques pour extraire des connaissances et des idées à partir de nombreuses données structurées ou non . Elle est souvent associée aux données massives et à l'analyse des données.

Elle utilise des techniques et des théories tirées de nombreux domaines dans le contexte des mathématiques, des statistiques, de l'informatique, de la théorie et des technologies de l'information, parmi lesquelles : l'apprentissage automatique, la compression de données et le calcul à haute performance.



Méthode d'apprentissage automatique : une illustration

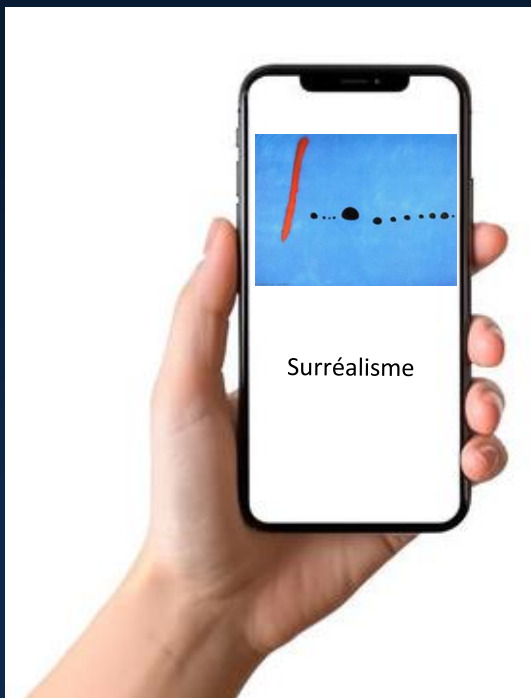
Apprendre à l'ordinateur à discerner une peinture du mouvement:

- Impressionnisme,
- Surréalisme,
- Expressionisme abstrait.



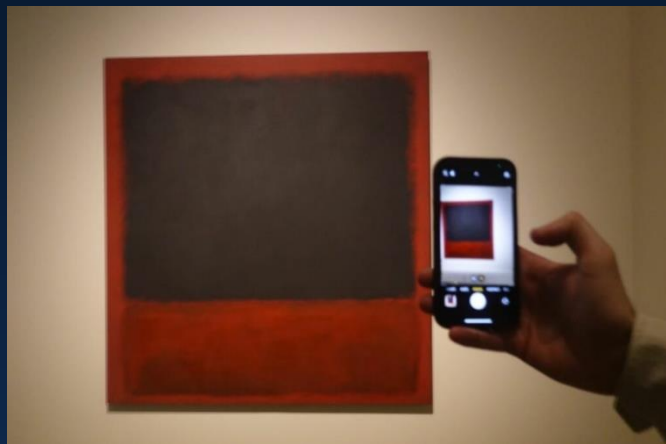
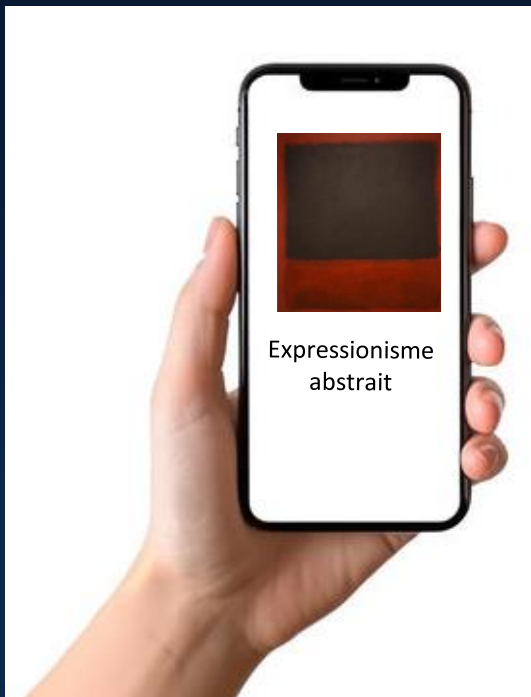
Apprendre à l'ordinateur à discerner une peinture du mouvement:

- Impressionnisme,
- Surréalisme,
- Expressionisme abstrait.

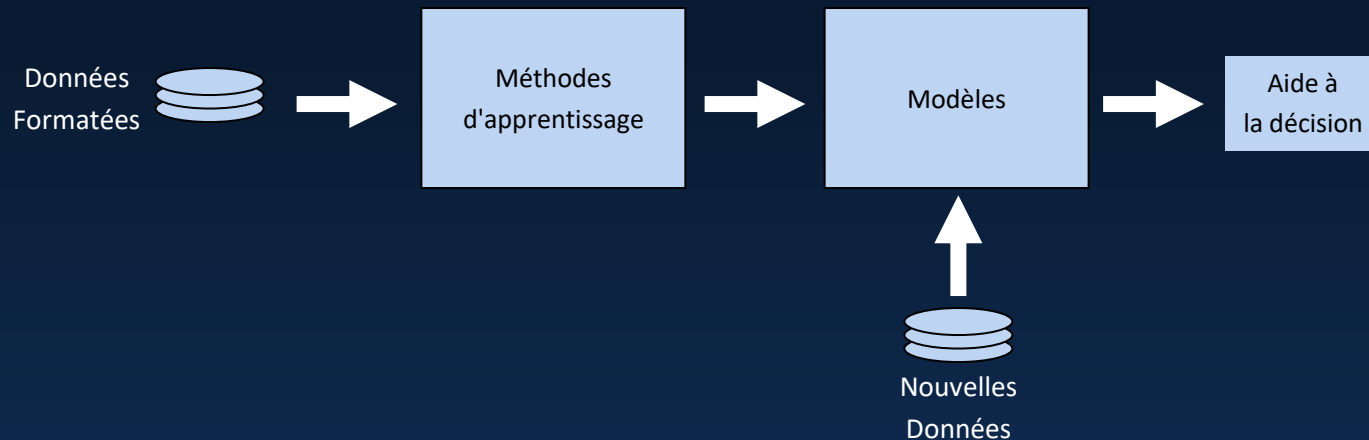


Apprendre à l'ordinateur à discerner une peinture du mouvement:

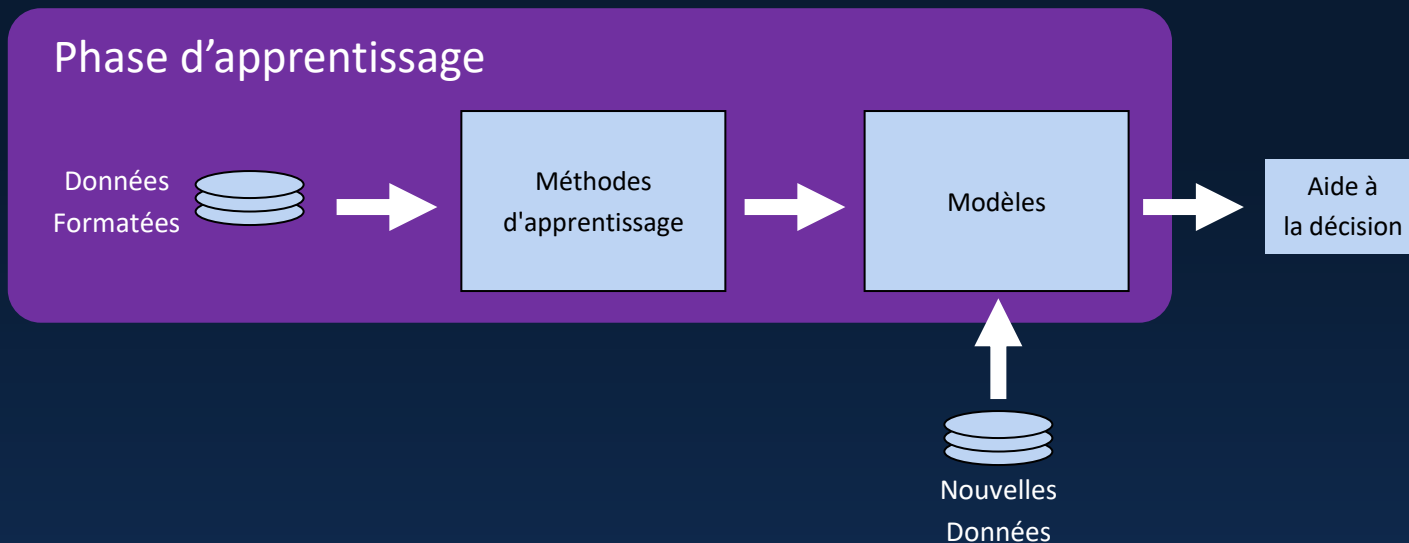
- Impressionnisme,
- Surréalisme,
- Expressionisme abstrait.



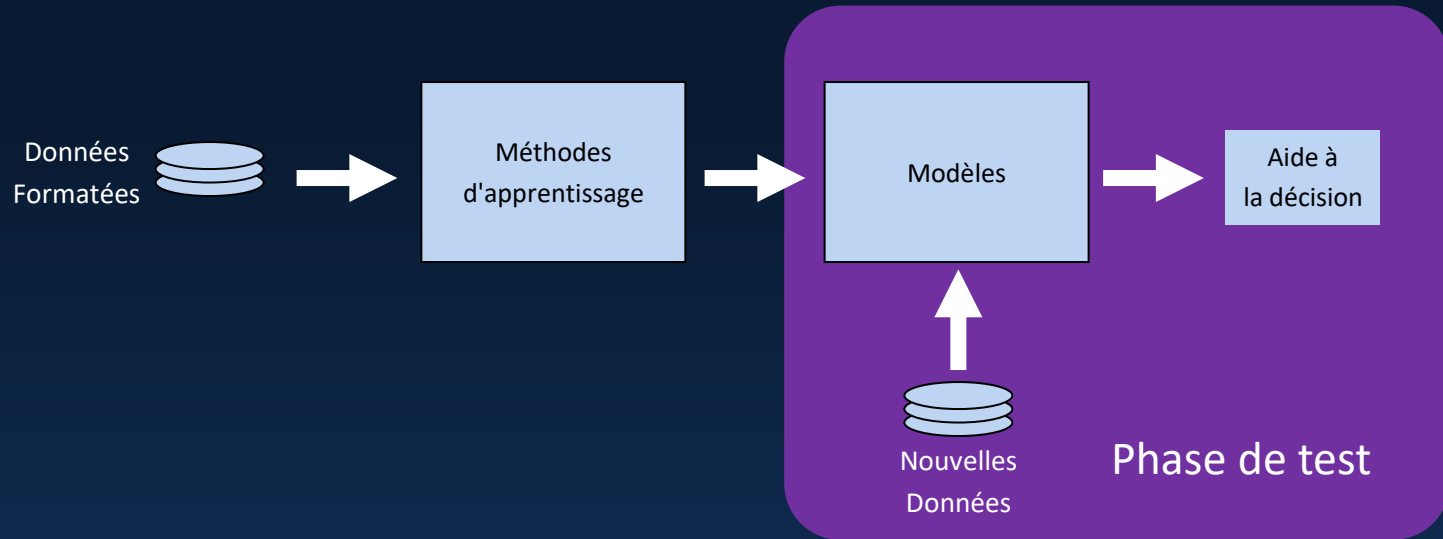
Processus d'induction à partir des données

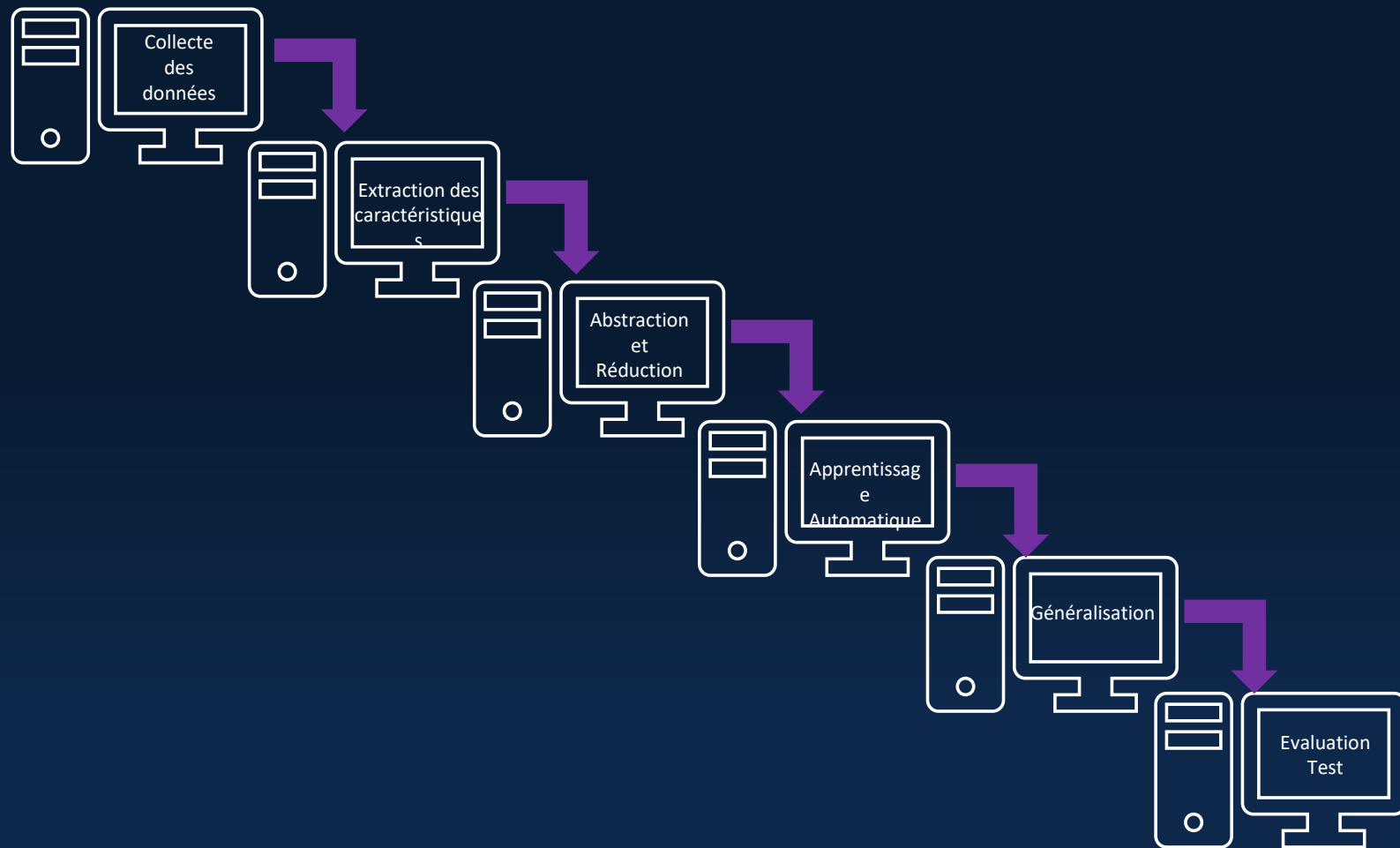


Processus d'induction à partir des données



Processus d'induction à partir des données







Constitution de la base d'apprentissage

1. Impressionnisme (vers 1860-1890)

L'Impressionnisme est un mouvement qui met l'accent sur la lumière, la couleur et les sensations visuelles immédiates, souvent en plein air.

Artistes et œuvres associées :

Claude Monet – Impression, soleil levant (1872), Nymphéas (série, 1899-1926)

Pierre-Auguste Renoir – Bal du moulin de la Galette (1876), Déjeuner des canotiers (1881)

Edgar Degas – L'Étoile (1878), Les Repasseuses (1884)

Camille Pissarro – Boulevard Montmartre, effet de nuit (1897), Gelée blanche (1873)

Berthe Morisot – Le Berceau (1872), Jeune femme en toilette de bal (1879)

Alfred Sisley – Inondation à Port-Marly (1876), La Seine à Bougival (1876)

Gustave Caillebotte – Les Raboteurs de parquet (1875), Rue de Paris, temps de pluie (1877)

Mary Cassatt – La Toilette (1891), Le Thé (1880)

Frédéric Bazille – Atelier de Bazille (1870), La Réunion de famille (1867)

Eva Gonzalès – Nanny et Enfant (1877-78), Le Chignon (1877)



Constitution de la base d'apprentissage

2. Surréalisme (vers 1920-1960)

Le Surréalisme explore l'inconscient, les rêves et l'irrationnel, souvent avec une forte charge symbolique.

Artistes et œuvres associées :

- Salvador Dalí – La Persistance de la mémoire (1931), Le Rêve (1931)
- René Magritte – La Trahison des images (1929), L'Empire des lumières (1954)
- Max Ernst – L'Éléphant Célèbes (1921), La Ville entière (1935-36)
- Joan Miró – Le Carnaval d'Arlequin (1924-25), Bleu II (1961)
- André Masson – Automatisme (1925), La Métamorphose (1929)
- Yves Tanguy – Mama, Papa est blessé ! (1927), L'Extinction des espèces (1938)
- Leonora Carrington – La Débutante (1937), The Lovers (1950)
- Dorothea Tanning – Birthday (1942), Eine Kleine Nachtmusik (1943)
- Victor Brauner – Le Surréaliste (1947), Le Grand Transparent (1947)
- Frida Kahlo – Les Deux Fridas (1939), La Colonne brisée (1944)



Constitution de la base d'apprentissage

3. Expressionnisme abstrait (vers 1940-1960)

Ce mouvement, né aux États-Unis, privilégie l'expression spontanée des émotions à travers des formes abstraites et gestuelles.
Artistes et œuvres associées :

Jackson Pollock – Number 1A, 1948 (1948), Lavender Mist (1950)

Mark Rothko – Orange, Red, Yellow (1961), No. 61 (Rust and Blue) (1953)

Willem de Kooning – Woman I (1952), Excavation (1950)

Franz Kline – Mahoning (1956), Chief (1950)

Clyfford Still – PH-385 (1949), PH-1033 (1976)

Robert Motherwell – Elegy to the Spanish Republic No. 110 (1971), At Five in the Afternoon (1949)

Lee Krasner – The Seasons (1957), Gaea (1966)

Joan Mitchell – Hemlock (1956), City Landscape (1955)

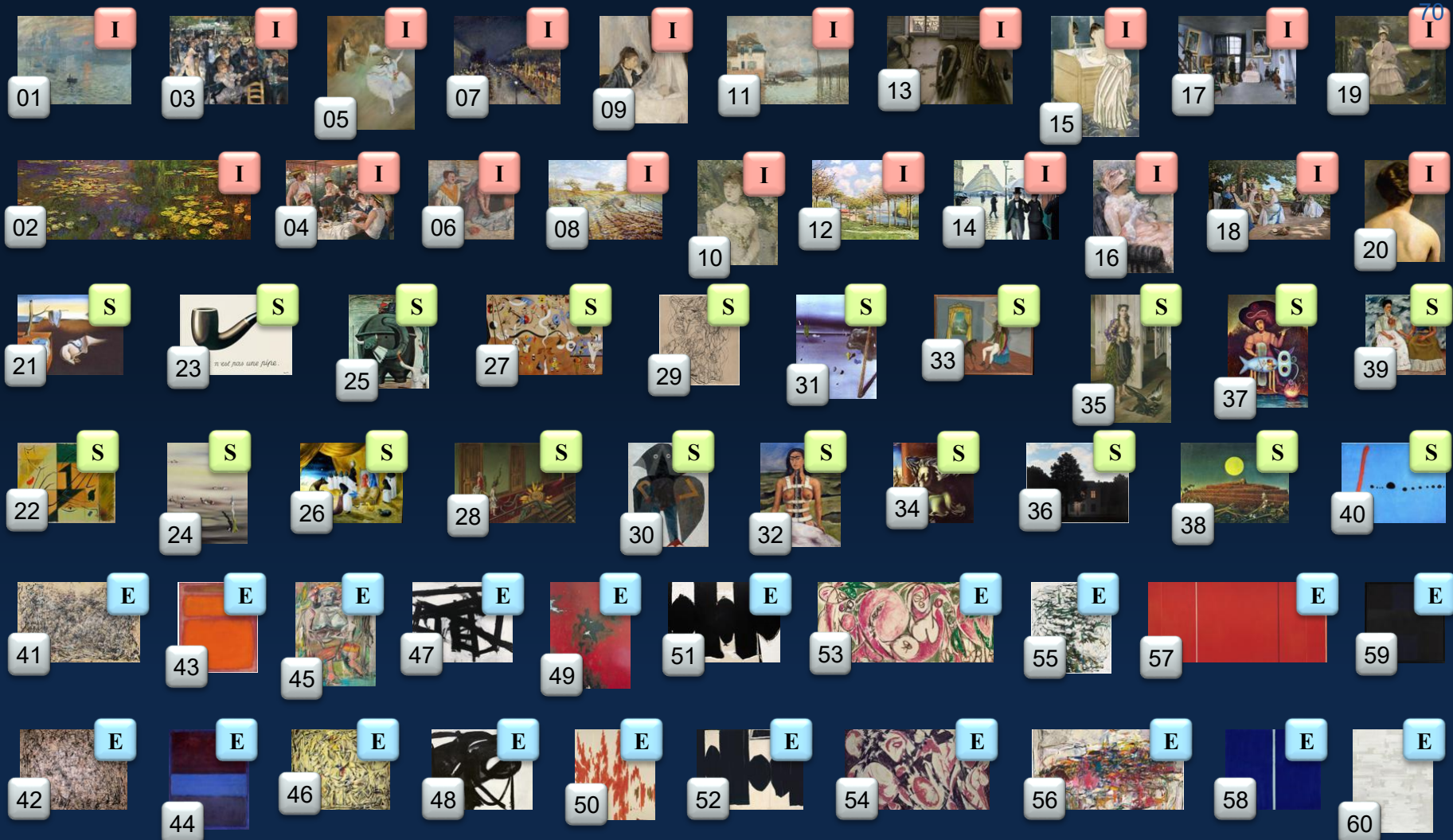
Barnett Newman – Vir Heroicus Sublimis (1951), Onement VI (1953)

Ad Reinhardt – Abstract Painting (1966), Number 107 (1950)

Collecte des
données

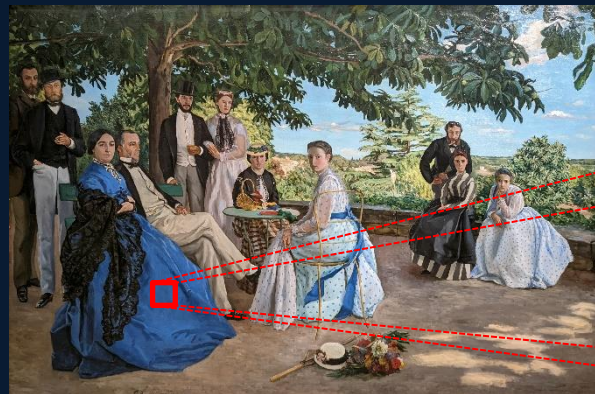
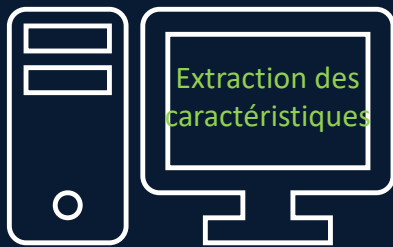
Constitution de la base d'apprentissage



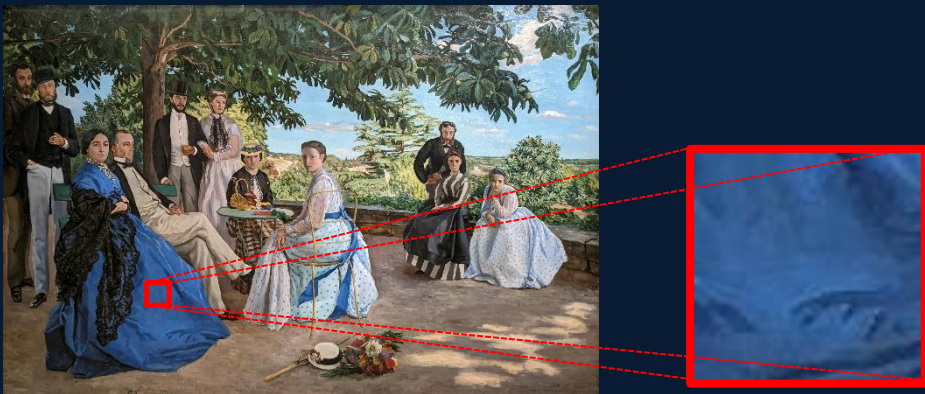
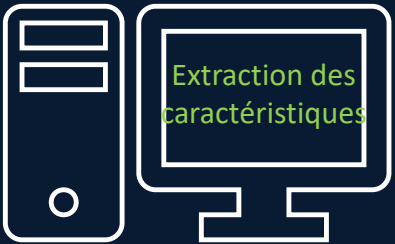




Numéro Descripteur	Descripteur	Type	Description	Interprétation	Exemple
1	Id	Métadonnée	Numéro d'Identification	Numéro d'identification unique	18
2	Style	Classe	Nom du Mouvement artistique du tableau {Impressionnisme, Surréalisme, Expressionnisme-Abstrait}	Classe à prédire	Impressionnisme
3	Annee	Métadonnée	Année d'achèvement du tableau		1867
4	Prenom_Nom_Titre	Métadonnée			Frederic_Bazille_La_Reunion_de_famille
5	No_Prenom_Nom_Titre_Annee	Métadonnée			18_Frederic_Bazille_La_Reunion_de_famille_1867



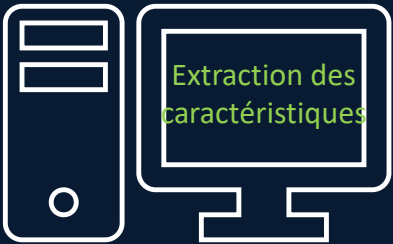
Numéro Descripteur	Descripteur	Type	Description	Interprétation	Exemple
6	Moyenne_Rouge	Descripteur	Moyenne du canal rouge (0-1)	0 = Pas de rouge, 1 = Image très rouge	0.43
7	Moyenne_Vert	Descripteur	Moyenne du canal vert (0-1)	0 = Pas de vert, 1 = Image très verte	0.43
8	Moyenne_Bleu	Descripteur	Moyenne du canal bleu (0-1)	0 = Pas de bleu, 1 = Image très bleue	0.40
9	Saturation_Moyenne	Descripteur	Intensité moyenne des couleurs (0-1)	0 = Image en niveaux de gris, 1 = Couleurs très saturées	0.29
10	Variance_Saturation	Descripteur	Variation de la saturation (dispersion des couleurs)	Faible = Couleurs homogènes, élevée = Mélange de couleurs vives	0.03
11	Teinte_Moyenne	Descripteur	Couleur dominante dans l'image (0-360°)	0° = Rouge, 120° = Vert, 240° = Bleu	95.46
12	Contraste	Descripteur	Variation moyenne d'intensité entre pixels voisins	Faible = Image lisse, élevée = Image très contrastée	0.54
13	Entropie	Descripteur	Complexité de l'image (niveau de détail)	Faible = Image simple, élevée = Image très détaillée	7.76
14	Energie	Descripteur	Uniformité globale de l'image	Faible = texture très variée, élevée, texture très répétitive	252.03
15	GLCM_Contraste	Descripteur	Contraste basé sur la matrice de co-occurrence	Faible = Peu de variations de texture, élevée = Texture marquée	0.23
16	GLCM_Correlation	Descripteur	Corrélation entre pixels adjacents	Proche de 1 = Zones homogènes, proche de 0 = Texture désordonnée	0.97
17	GLCM_energie	Descripteur	Énergie (homogénéité) de la texture	Faible = Texture variée, élevée = Texture répétitive	0.11
18	GLCM_Homogeneite	Descripteur	Régularité des textures	Faible = Texture chaotique, élevée = Texture uniforme	0.91
19	Gabor_Moyenne	Descripteur	Réponse moyenne des filtres de Gabor (analyse de texture)	Faible = Texture faible, élevée = Texture marquée	701.31
20	Gabor_ecart-Type	Descripteur	Dispersion des réponses de Gabor	Faible = Texture homogène, élevée = Texture variée	1361.15
21	Nombre_Visages	Descripteur	Nombre de visages détectés	0 = Aucun visage, >0 = Nombre de visages trouvés	17.00
22	Densite_Contours	Descripteur	Nombre de contours détectés par Canny	Faible = Image floue, élevée = Image avec bords nets	0.10



Numéro Descripteur	Descripteur	Type	Description	Interprétation	Exemple
23	LBP_1	Descripteur	Local Binary Pattern 1	Mesure la texture locale	0.07
24	LBP_2	Descripteur	Local Binary Pattern 2	Mesure la texture locale	0.13
25	LBP_3	Descripteur	Local Binary Pattern 3	Mesure la texture locale	0.12
26	LBP_4	Descripteur	Local Binary Pattern 4	Mesure la texture locale	0.37
27	LBP_5	Descripteur	Local Binary Pattern 5	Mesure la texture locale	0.65
28	LBP_6	Descripteur	Local Binary Pattern 6	Mesure la texture locale	0.49
29	LBP_7	Descripteur	Local Binary Pattern 7	Mesure la texture locale	0.24
30	LBP_8	Descripteur	Local Binary Pattern 8	Mesure la texture locale	0.17
31	LBP_9	Descripteur	Local Binary Pattern 9	Mesure la texture locale	0.18
32	LBP_10	Descripteur	Local Binary Pattern 10	Mesure la texture locale	0.22
33	CouleurDominante_1_R	Descripteur	Composante rouge de la couleur dominante 1	0-1 : Plus proche de 1 = Plus de rouge	0.73
34	CouleurDominante_1_G	Descripteur	Composante verte de la couleur dominante 1	0-1 : Plus proche de 1 = Plus de vert	0.77
35	CouleurDominante_1_B	Descripteur	Composante bleue de la couleur dominante 1	0-1 : Plus proche de 1 = Plus de bleu	0.76
36	CouleurDominante_2_R	Descripteur	Composante rouge de la couleur dominante 2	0-1 : Plus proche de 1 = Plus de rouge	0.20
37	CouleurDominante_2_G	Descripteur	Composante verte de la couleur dominante 2	0-1 : Plus proche de 1 = Plus de vert	0.21
38	CouleurDominante_2_B	Descripteur	Composante bleue de la couleur dominante 2	0-1 : Plus proche de 1 = Plus de bleu	0.17
39	CouleurDominante_3_R	Descripteur	Composante rouge de la couleur dominante 3	0-1 : Plus proche de 1 = Plus de rouge	0.47
40	CouleurDominante_3_G	Descripteur	Composante verte de la couleur dominante 3	0-1 : Plus proche de 1 = Plus de vert	0.45
41	CouleurDominante_3_B	Descripteur	Composante bleue de la couleur dominante 3	0-1 : Plus proche de 1 = Plus de bleu	0.41



Numéro Descripteur	Descripteur	Type	Description	Interprétation	Exemple
42	HistR_1	Descripteur	Bin 1 de l'histogramme du canal rouge	0-1 : Distribution des pixels rouges dans l'image	0.08
43	HistG_1	Descripteur	Bin 1 de l'histogramme du canal vert	0-1 : Distribution des pixels verts dans l'image	0.38
44	HistB_1	Descripteur	Bin 1 de l'histogramme du canal bleu	0-1 : Distribution des pixels bleus dans l'image	0.31
45	HistGray_1	Descripteur	Bin 1 de l'histogramme des niveaux de gris	0-1 : Distribution des niveaux de gris	0.22
46	HistR_2	Descripteur	Bin 2 de l'histogramme du canal rouge	0-1 : Distribution des pixels rouges dans l'image	0.02
47	HistG_2	Descripteur	Bin 2 de l'histogramme du canal vert	0-1 : Distribution des pixels verts dans l'image	0.08
48	HistB_2	Descripteur	Bin 2 de l'histogramme du canal bleu	0-1 : Distribution des pixels bleus dans l'image	0.39
49	HistGray_2	Descripteur	Bin 2 de l'histogramme des niveaux de gris	0-1 : Distribution des niveaux de gris	0.31
50	HistR_3	Descripteur	Bin 3 de l'histogramme du canal rouge	0-1 : Distribution des pixels rouges dans l'image	0.20
51	HistG_3	Descripteur	Bin 3 de l'histogramme du canal vert	0-1 : Distribution des pixels verts dans l'image	0.03
52	HistB_3	Descripteur	Bin 3 de l'histogramme du canal bleu	0-1 : Distribution des pixels bleus dans l'image	0.12
53	HistGray_3	Descripteur	Bin 3 de l'histogramme des niveaux de gris	0-1 : Distribution des niveaux de gris	0.40
54	HistR_4	Descripteur	Bin 4 de l'histogramme du canal rouge	0-1 : Distribution des pixels rouges dans l'image	0.29
55	HistG_4	Descripteur	Bin 4 de l'histogramme du canal vert	0-1 : Distribution des pixels verts dans l'image	0.11
56	HistB_4	Descripteur	Bin 4 de l'histogramme du canal bleu	0-1 : Distribution des pixels bleus dans l'image	0.08
57	HistGray_4	Descripteur	Bin 4 de l'histogramme des niveaux de gris	0-1 : Distribution des niveaux de gris	0.08
58	HistR_5	Descripteur	Bin 5 de l'histogramme du canal rouge	0-1 : Distribution des pixels rouges dans l'image	0.40
59	HistG_5	Descripteur	Bin 5 de l'histogramme du canal vert	0-1 : Distribution des pixels verts dans l'image	0.30
60	HistB_5	Descripteur	Bin 5 de l'histogramme du canal bleu	0-1 : Distribution des pixels bleus dans l'image	0.21
61	HistGray_5	Descripteur	Bin 5 de l'histogramme des niveaux de gris	0-1 : Distribution des niveaux de gris	0.02



01	02	03	04	05
18	Impressionnisme	1867	Frederic_Bazille_La_Reunion_de_famille	18_Frederic_Bazille_La_Reunion_de_famille_1867

06	07	08	09	10	11	12	13	14	15	16	17	18	19	20	21	22	23	24	25	26	27	28	29	30	31	32	33
0.43	0.43	0.40	0.29	0.03	95.46	0.54	7.76	252	0.23	0.97	0.11	0.91	701	1361	17.00	0.10	0.07	0.13	0.12	0.37	0.65	0.49	0.24	0.17	0.18	0.22	0.73

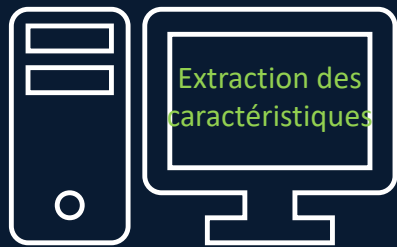
34	35	36	37	38	39	40	41	42	43	44	45	46	47	48	49	50	51	52	53	54	55	56	57	58	59	60	61
0.77	0.76	0.20	0.21	0.17	0.47	0.45	0.41	0.08	0.38	0.31	0.22	0.02	0.08	0.39	0.31	0.20	0.03	0.12	0.40	0.29	0.11	0.08	0.08	0.40	0.30	0.21	0.02



01	02	03	04	05
18	Impressionnisme	1867	Frederic_Bazille_La_Reunion_de_famille	18_Frederic_Bazille_La_Reunion_de_famille_1867

06	07	08	09	10	11	12	13	14	15	16	17	18	19	20	21	22	23	24	25	26	27	28	29	30	31	32	33
0.43	0.43	0.40	0.29	0.03	95.46	0.54	7.76	252	0.23	0.97	0.11	0.91	701	1361	17.00	0.10	0.07	0.13	0.12	0.37	0.65	0.49	0.24	0.17	0.18	0.22	0.73

34	35	36	37	38	39	40	41	42	43	44	45	46	47	48	49	50	51	52	53	54	55	56	57	58	59	60	61
0.77	0.76	0.20	0.21	0.17	0.47	0.45	0.41	0.08	0.38	0.31	0.22	0.02	0.08	0.39	0.31	0.20	0.03	0.12	0.40	0.29	0.11	0.08	0.08	0.40	0.30	0.21	0.02



02

Impressionnisme

06	07	08	09	10	11	12	13	14	15	16	17	18	19	20	21	22	23	24	25	26	27	28	29	30	31	32	33
0.43	0.43	0.40	0.29	0.03	95.46	0.54	7.76	252	0.23	0.97	0.11	0.91	701	1361	17.00	0.10	0.07	0.13	0.12	0.37	0.65	0.49	0.24	0.17	0.18	0.22	0.73

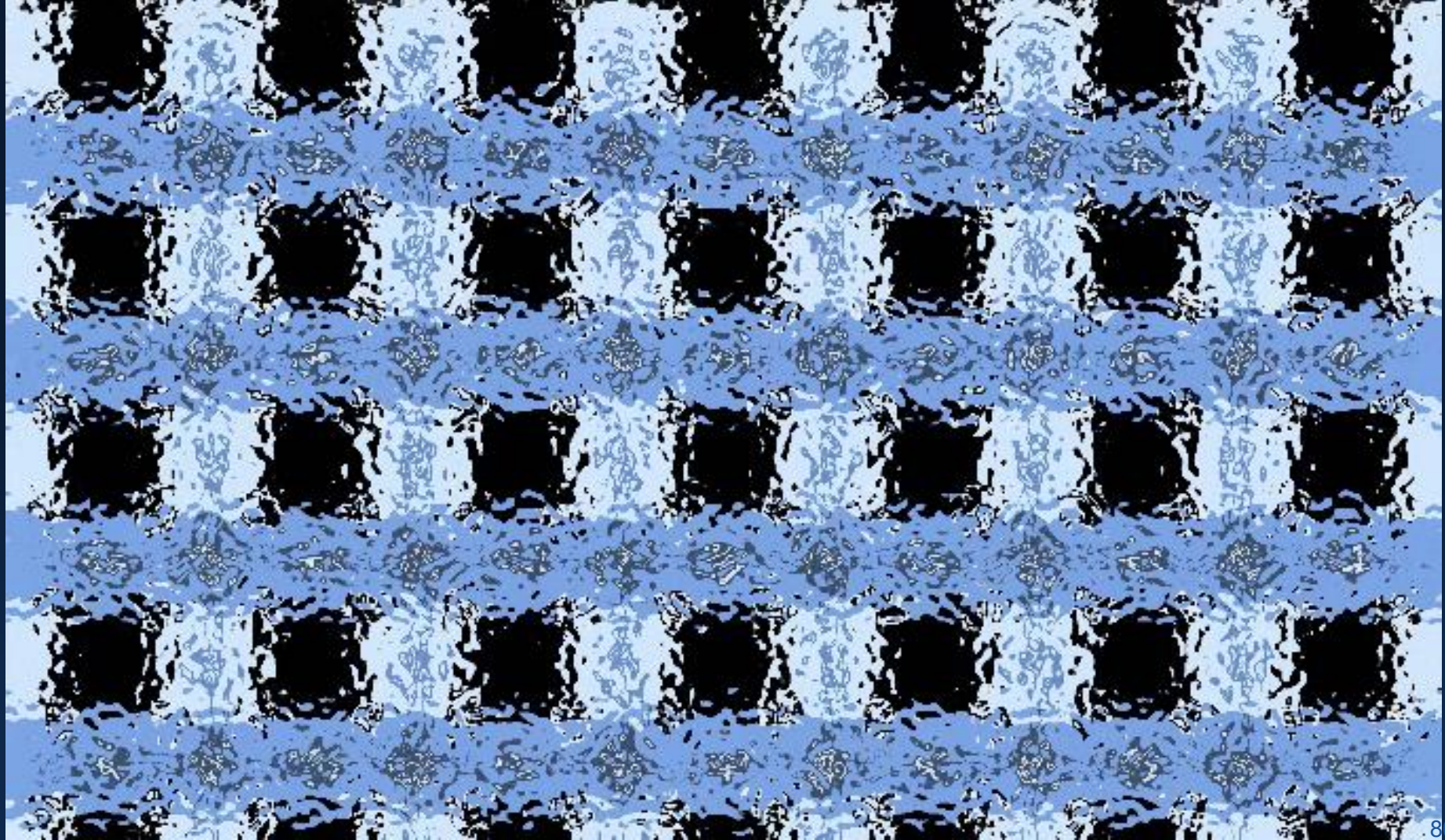
34	35	36	37	38	39	40	41	42	43	44	45	46	47	48	49	50	51	52	53	54	55	56	57	58	59	60	61
0.77	0.76	0.20	0.21	0.17	0.47	0.45	0.41	0.08	0.38	0.31	0.22	0.02	0.08	0.39	0.31	0.20	0.03	0.12	0.40	0.29	0.11	0.08	0.08	0.40	0.30	0.21	0.02



Ensemble de Tableaux d'Apprentissage :
Matrice de 60 lignes et 56 colonnes de descripteurs et 1 colonne de Classe

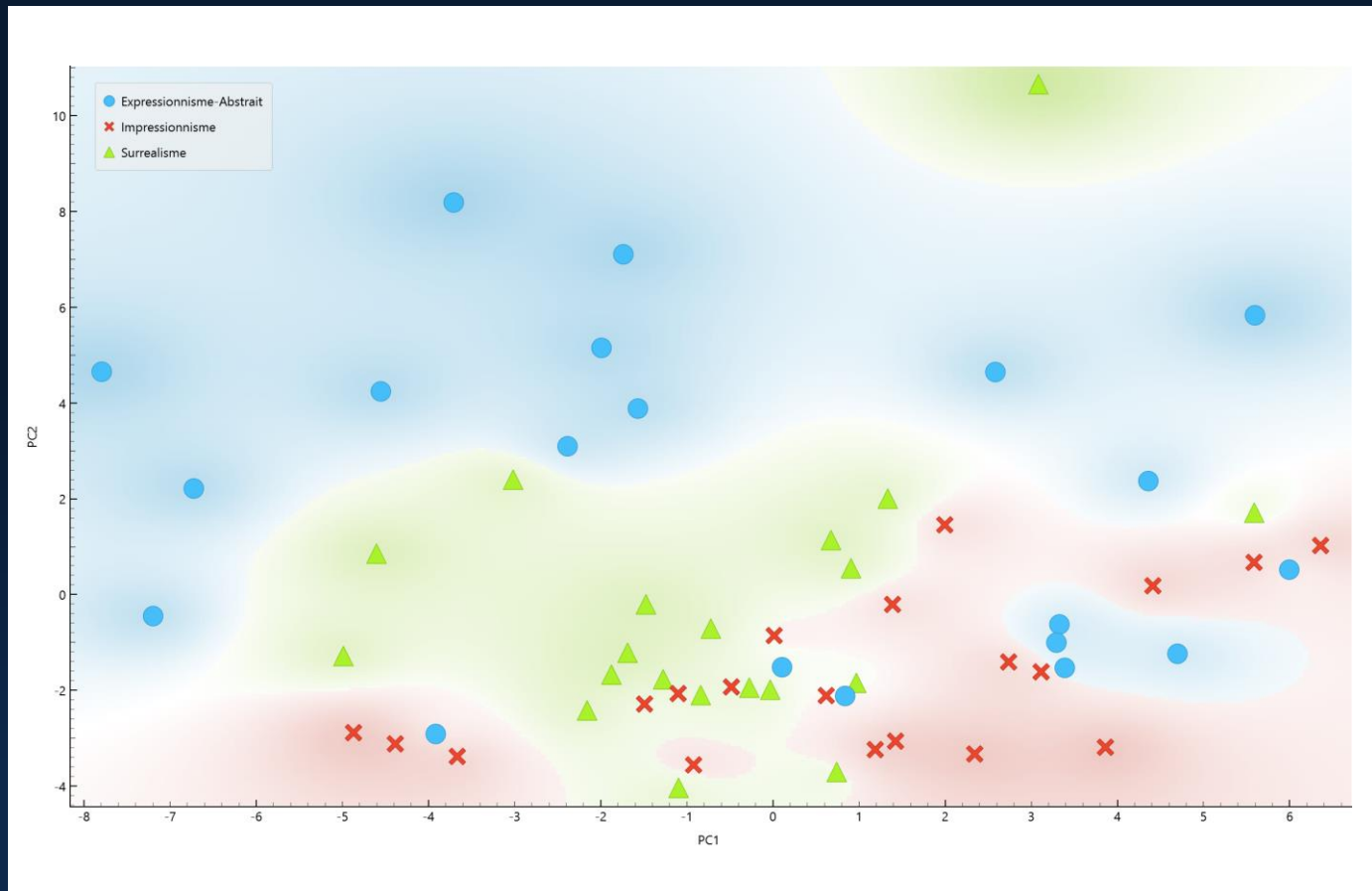
[illegible]

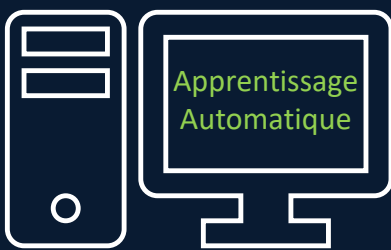




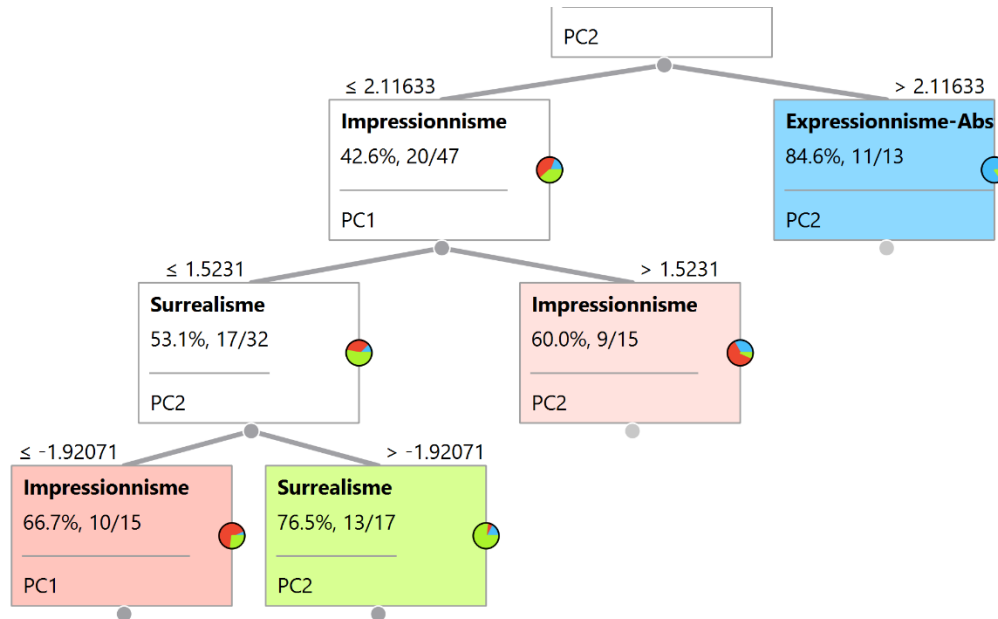
Abstraction Et Réduction

Id	Style	PC1	PC2
01	Impressionnisme	3.86	-3.19
02	Impressionnisme	-4.39	-3.12
03	Impressionnisme	-0.92	-3.56
04	Impressionnisme	-1.10	-2.07
05	Impressionnisme	1.19	-3.24
06	Impressionnisme	1.42	-3.06
07	Impressionnisme	-3.67	-3.38
08	Impressionnisme	6.36	1.02
09	Impressionnisme	5.59	0.67
10	Impressionnisme	2.34	-3.33
11	Impressionnisme	3.11	-1.61
12	Impressionnisme	1.39	-0.21
13	Impressionnisme	-4.87	-2.88
14	Impressionnisme	2.00	1.46
15	Impressionnisme	4.41	0.18
16	Impressionnisme	2.74	-1.41
17	Impressionnisme	0.62	-2.11
18	Impressionnisme	-0.48	-1.93
19	Impressionnisme	-1.49	-2.29
20	Impressionnisme	0.01	-0.86
21	Surrealisme	-1.69	-1.22
22	Surrealisme	-4.60	0.85
23	Surrealisme	3.08	10.66
24	Surrealisme	-3.02	2.40
25	Surrealisme	-0.27	-1.94
26	Surrealisme	-1.10	-4.04
27	Surrealisme	-0.72	-0.71
28	Surrealisme	0.67	1.14
29	Surrealisme	5.59	1.71
30	Surrealisme	-0.03	-1.99
31	Surrealisme	0.91	0.55
32	Surrealisme	1.33	2.00
33	Surrealisme	-0.84	-2.10
34	Surrealisme	-1.88	-1.67
35	Surrealisme	-2.16	-2.42
36	Surrealisme	-4.99	-1.28
37	Surrealisme	-1.27	-1.77
38	Surrealisme	-1.47	-0.20
39	Surrealisme	0.97	-1.85
40	Surrealisme	0.74	-3.70
41	Expressionnisme-Abstrait	4.70	-1.24
42	Expressionnisme-Abstrait	0.84	-2.12
43	Expressionnisme-Abstrait	-1.99	5.15
44	Expressionnisme-Abstrait	-7.20	-0.45
45	Expressionnisme-Abstrait	3.39	-1.53
46	Expressionnisme-Abstrait	6.00	0.52
47	Expressionnisme-Abstrait	-1.57	3.88
48	Expressionnisme-Abstrait	-2.39	3.10
49	Expressionnisme-Abstrait	-3.92	-2.91
50	Expressionnisme-Abstrait	2.58	4.65
51	Expressionnisme-Abstrait	-1.74	7.10
52	Expressionnisme-Abstrait	-3.71	8.19
53	Expressionnisme-Abstrait	3.33	-0.62
54	Expressionnisme-Abstrait	0.11	-1.52
55	Expressionnisme-Abstrait	4.36	2.37
56	Expressionnisme-Abstrait	3.29	-1.00
57	Expressionnisme-Abstrait	-4.55	4.24
58	Expressionnisme-Abstrait	-6.73	2.21
59	Expressionnisme-Abstrait	-7.79	4.65
60	Expressionnisme-Abstrait	5.60	5.83





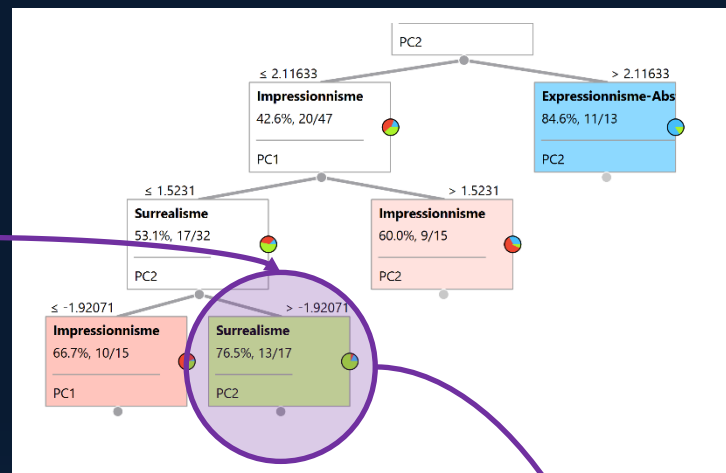
Arbre de décision



Apprentissage Automatique

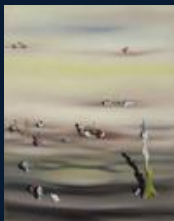


Id	Style	PC1	PC2
34	Surrealisme	-1.88	-1.67

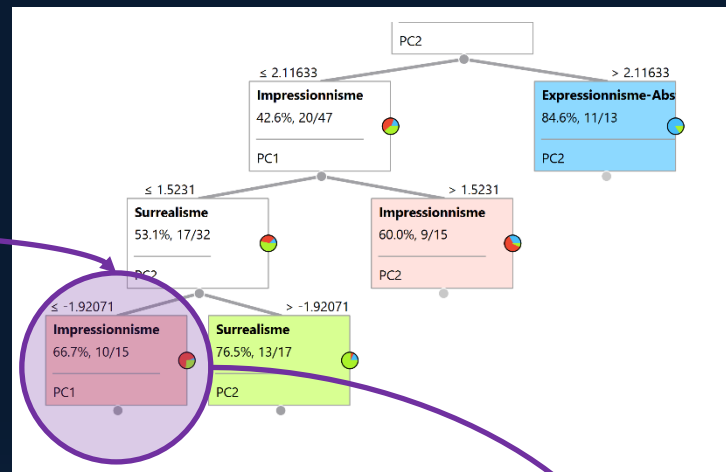


Matrice de confusion

		Predicted			Σ
		Expressionnisme-Abstrait	Impressionnisme	Surrealisme	
Actual	Expressionnisme-Abstrait	19	1	0	20
	Impressionnisme	2	17	1	20
	Surrealisme	4	3	13	20
Σ		25	21	14	60

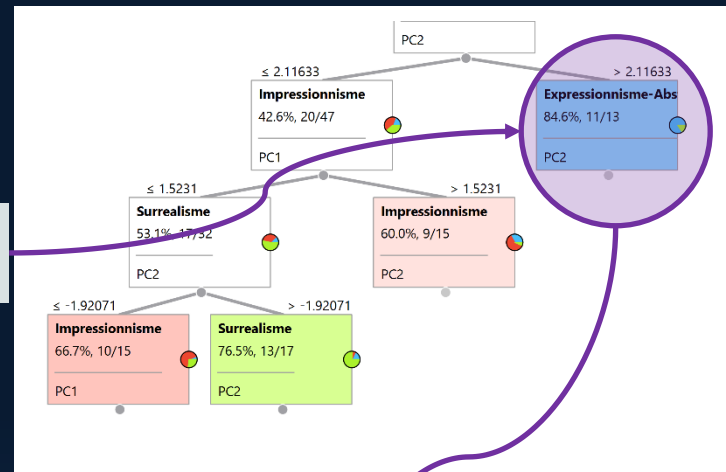


Id	Style	PC1	PC2
33	Surrealisme	-0.84	-2.10



		Predicted			
		Expressionnisme-Abstrait	Impressionnisme	Surrealisme	Σ
Actual	Expressionnisme-Abstrait	19	1	0	20
	Impressionnisme	2	17	1	20
	Surrealisme	4	3	13	20
Σ		25	21	14	60

Arbre de décision



Matrice de confusion

		Predicted			Σ
		Expressionnisme-Abstrait	Impressionnisme	Surrealisme	
Actual	Expressionnisme-Abstrait	19	1	0	20
	Impressionnisme	2	17	1	20
	Surrealisme	4	3	13	20
Σ		25	21	14	60

Apprentissage
Automatique

Id	Style	PC1	PC2
47	Expressionnisme -Abstrait	-1.57	3.88

Id	Style	PC1	PC2
01	Impressionnisme	3.86	-3.19
02	Impressionnisme	-4.39	-3.12
03	Impressionnisme	-0.92	-3.56
04	Impressionnisme	-1.10	-2.07
05	Impressionnisme	1.19	-3.24
06	Impressionnisme	1.42	-3.06
07	Impressionnisme	-3.67	-3.38
08	Impressionnisme	6.36	1.02
09	Impressionnisme	5.59	0.67
10	Impressionnisme	2.34	-3.33
11	Impressionnisme	3.11	-1.61
12	Impressionnisme	1.39	-0.21
13	Impressionnisme	-4.87	-2.88
14	Impressionnisme	2.00	1.46
15	Impressionnisme	4.41	0.18
16	Impressionnisme	2.74	-1.41
17	Impressionnisme	0.62	-2.11
18	Impressionnisme	-0.48	-1.93
19	Impressionnisme	-1.49	-2.29
20	Impressionnisme	0.01	-0.86
21	Surrealisme	-1.69	-1.22
22	Surrealisme	-4.60	0.85
23	Surrealisme	3.08	10.66
24	Surrealisme	-3.02	2.40
25	Surrealisme	-0.27	-1.94
26	Surrealisme	-1.10	-4.04
27	Surrealisme	-0.72	-0.71
28	Surrealisme	0.67	1.14
29	Surrealisme	5.59	1.71
30	Surrealisme	-0.03	-1.99
31	Surrealisme	0.91	0.55
32	Surrealisme	1.33	2.00
33	Surrealisme	-0.84	-2.10
34	Surrealisme	-1.88	-1.67
35	Surrealisme	-2.16	-2.42
36	Surrealisme	-4.99	-1.28
37	Surrealisme	-1.27	-1.77
38	Surrealisme	-1.47	-0.20
39	Surrealisme	0.97	-1.85
40	Surrealisme	0.74	-3.70
41	Expressionnisme-Abstrait	4.70	-1.24
42	Expressionnisme-Abstrait	0.84	-2.12
43	Expressionnisme-Abstrait	-1.99	5.15
44	Expressionnisme-Abstrait	-7.20	-0.45
45	Expressionnisme-Abstrait	3.39	-1.53
46	Expressionnisme-Abstrait	6.00	0.52
47	Expressionnisme-Abstrait	-1.57	3.88
48	Expressionnisme-Abstrait	-2.38	3.10
49	Expressionnisme-Abstrait	-3.92	-2.91
50	Expressionnisme-Abstrait	2.58	4.65
51	Expressionnisme-Abstrait	-1.74	7.10
52	Expressionnisme-Abstrait	-3.71	8.19
53	Expressionnisme-Abstrait	3.33	-0.62
54	Expressionnisme-Abstrait	0.11	-1.52
55	Expressionnisme-Abstrait	4.36	2.37
56	Expressionnisme-Abstrait	3.29	-1.00
57	Expressionnisme-Abstrait	-4.55	4.24
58	Expressionnisme-Abstrait	-6.73	2.21
59	Expressionnisme-Abstrait	-7.79	4.65
60	Expressionnisme-Abstrait	5.60	5.83



Constitution de la base de Test

1. Impressionnisme

- Paul Cézanne – La Montagne Sainte-Victoire (1904-1906)
- Paul Gauguin – Vision après le sermon (La Lutte de Jacob avec l'Ange) (1888)
- Georges Seurat – Un dimanche après-midi à l'Île de la Grande Jatte (1884-1886)
- Armand Guillaumin – Soleil couchant à Ivry (1873)
- Gustave Moreau – La Licorne (1885)

2. Surréalisme

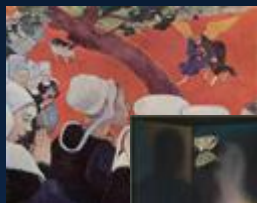
- Paul Delvaux – Le Voyage Légendaire (1974)
- Remedios Varo – Les Feuilles mortes (1963)
- Kay Sage – Tomorrow is Never (1955)
- Wilfredo Lam – La Jungle (1943)
- Toyen – Le Paravent (1966)

3. Expressionnisme abstrait

- Hedda Sterne – Number 6 (1951)
- Norman Lewis – Evening Rendezvous (1962)
- Perle Fine – Summer I (1968)
- Grace Hartigan – The King is Dead (1950)
- Michael Goldberg – Park Avenue Façade (1957)

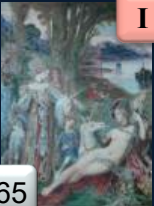
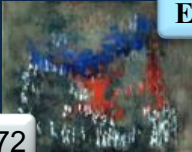
Collecte des
données

Constitution de la base de Test

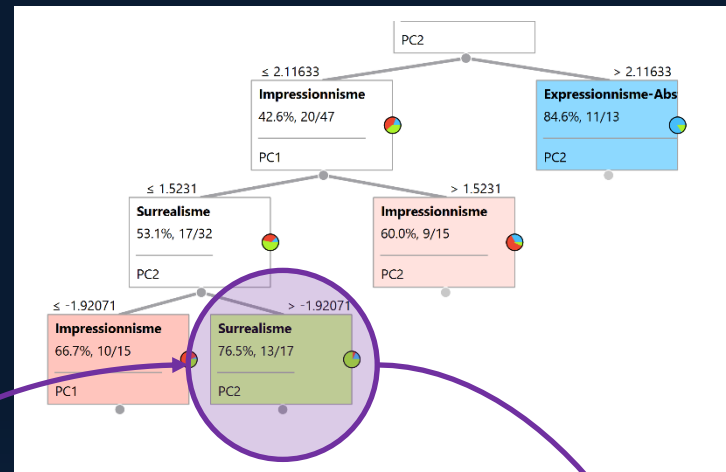




Constitution de la base de Test

 61	 62	 63	 64	 65
 66	 67	 68	 69	 70
 71	 72	 73	 74	 75

Arbre de décision



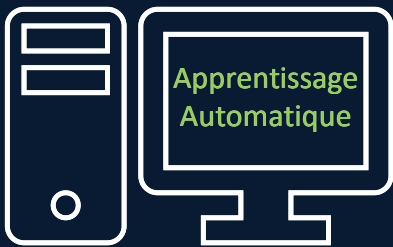
Phase de Test

Apprentissage
Automatique

Id	Style	PC1	PC2
61	Impressionnisme	1.34	-1.47
62	Impressionnisme	-1.82	-1.61
63	Impressionnisme	-0.64	-3.34
64	Impressionnisme	0.87	-2.39
65	Impressionnisme	0.13	-3.13
66	Surrealisme	-2.74	-1.23
67	Surrealisme	5.43	0.94
68	Surrealisme	-2.43	-2.71
69	Surrealisme	-0.76	-2.05
70	Surrealisme	-6.80	2.00
71	Expressionnisme-Abstrait	1.75	3.66
72	Expressionnisme-Abstrait	-2.86	-3.03
73	Expressionnisme-Abstrait	1.24	-1.78
74	Expressionnisme-Abstrait	0.09	-0.21
75	Expressionnisme-Abstrait	2.18	0.34

Matrice de confusion

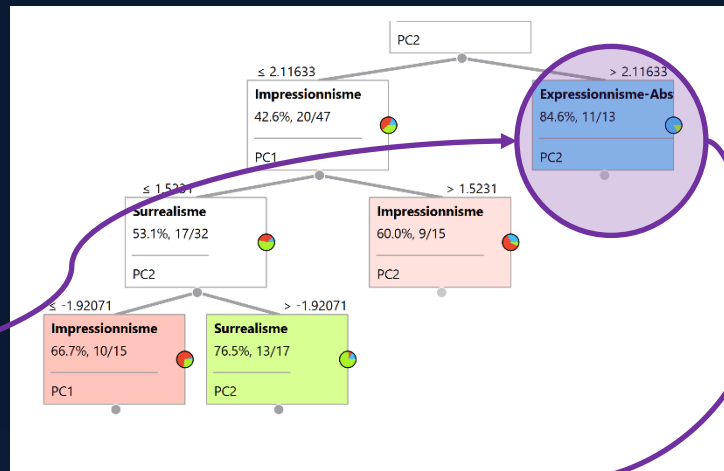
		Predicted			
		Expressionnisme-Abstrait	Impressionnisme	Surrealisme	Σ
Actual	Expressionnisme-Abstrait	3	1	1	5
	Impressionnisme	1	3		5
	Surrealisme	0	2	3	5
Σ		4	6	5	15



Phase de Test



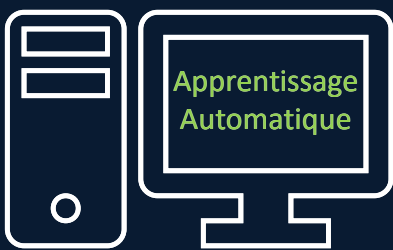
Arbre de décision



Id	Style	PC1	PC2
61	Impressionnisme	1.34	-1.47
62	Impressionnisme	-1.82	-1.61
63	Impressionnisme	-0.64	-3.34
64	Impressionnisme	0.87	-2.39
65	Impressionnisme	0.13	-3.13
66	Surrealisme	-2.74	-1.23
67	Surrealisme	5.43	0.94
68	Surrealisme	-2.43	-2.71
69	Surrealisme	-0.76	-2.05
70	Surrealisme	-6.80	2.00
71	Expressionnisme-Abstrait	1.75	3.66
72	Expressionnisme-Abstrait	-2.86	-3.03
73	Expressionnisme-Abstrait	1.24	-1.78
74	Expressionnisme-Abstrait	0.09	-0.21
75	Expressionnisme-Abstrait	2.18	0.34

Matrice de confusion

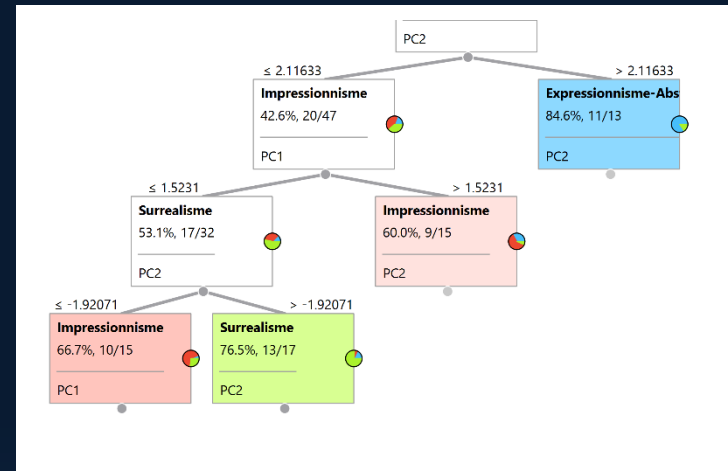
		Predicted			
		Expressionnisme-Abstrait	Impressionnisme	Surrealisme	Σ
Actual	Expressionnisme-Abstrait	3	1	1	5
	Impressionnisme	1	3	1	5
	Surrealisme	0	2	3	5
Σ		4	6	5	15



Phase de Test

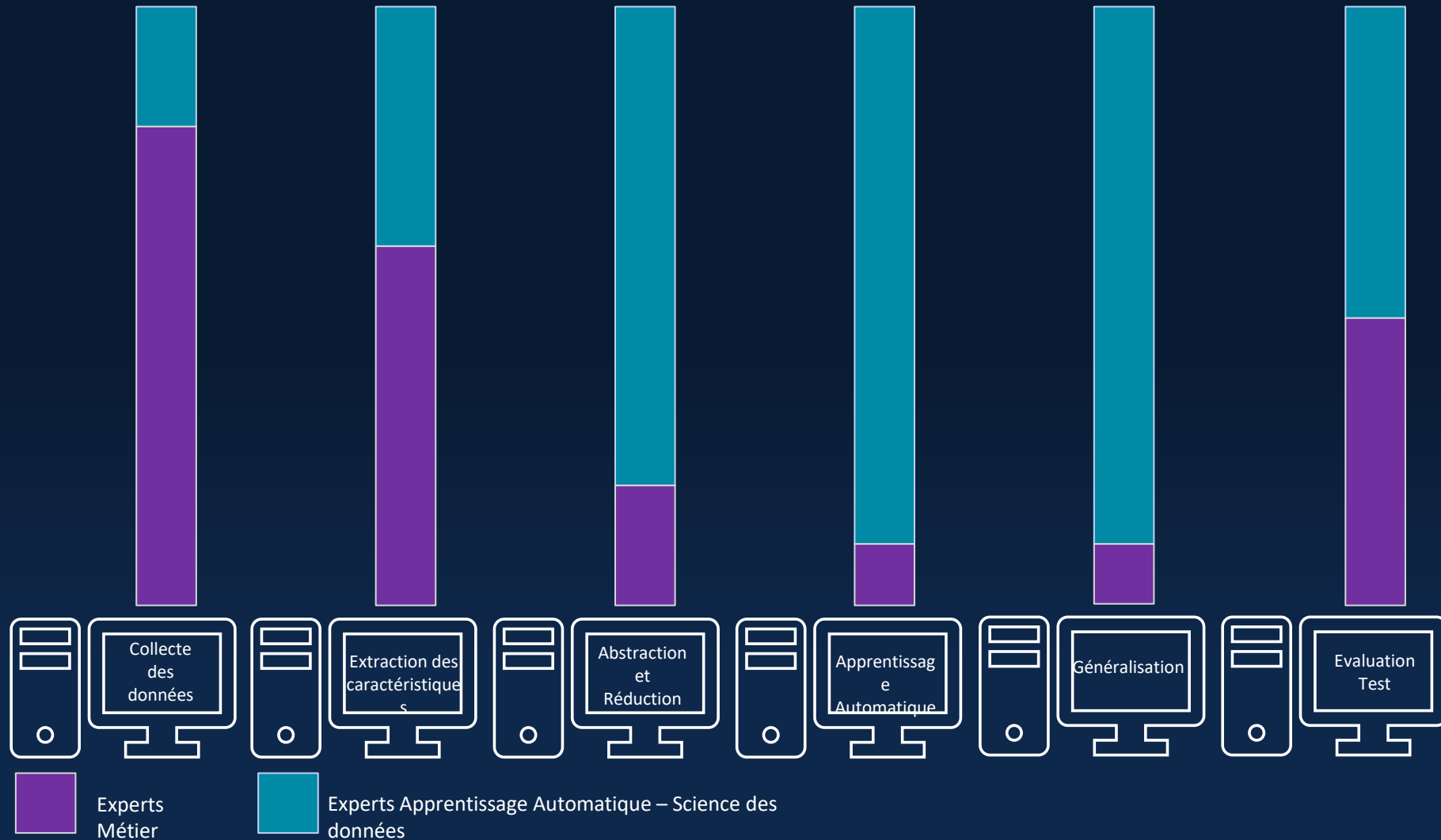
Id	Style	PC1	PC2
61	Impressionnisme	1.34	-1.47
62	Impressionnisme	-1.82	-1.61
63	Impressionnisme	-0.64	-3.34
64	Impressionnisme	0.87	-2.39
65	Impressionnisme	0.13	-3.13
66	Surrealisme	-2.74	-1.23
67	Surrealisme	5.43	0.94
68	Surrealisme	-2.43	-2.71
69	Surrealisme	-0.76	-2.05
70	Surrealisme	-6.80	2.00
71	Expressionnisme-Abstrait	1.75	3.66
72	Expressionnisme-Abstrait	-2.86	-3.03
73	Expressionnisme-Abstrait	1.24	-1.78
74	Expressionnisme-Abstrait	0.09	-0.21
75	Expressionnisme-Abstrait	2.18	0.34

Arbre de décision



Matrice de confusion

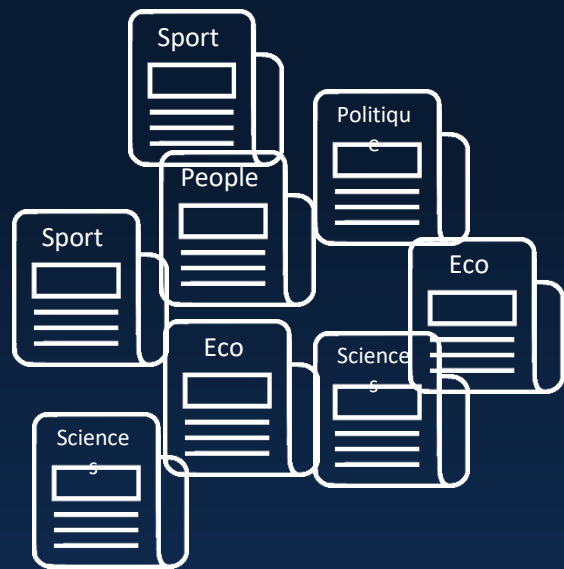
		Predicted			
		Expressionnisme-Abstrait	Impressionnisme	Surrealisme	Σ
Actual	Expressionnisme-Abstrait	3	1	1	5
	Impressionnisme	1	3	1	5
	Surrealisme	0	2	3	5
Σ		4	6	5	15



Examples

Classification de textes - Recommandation

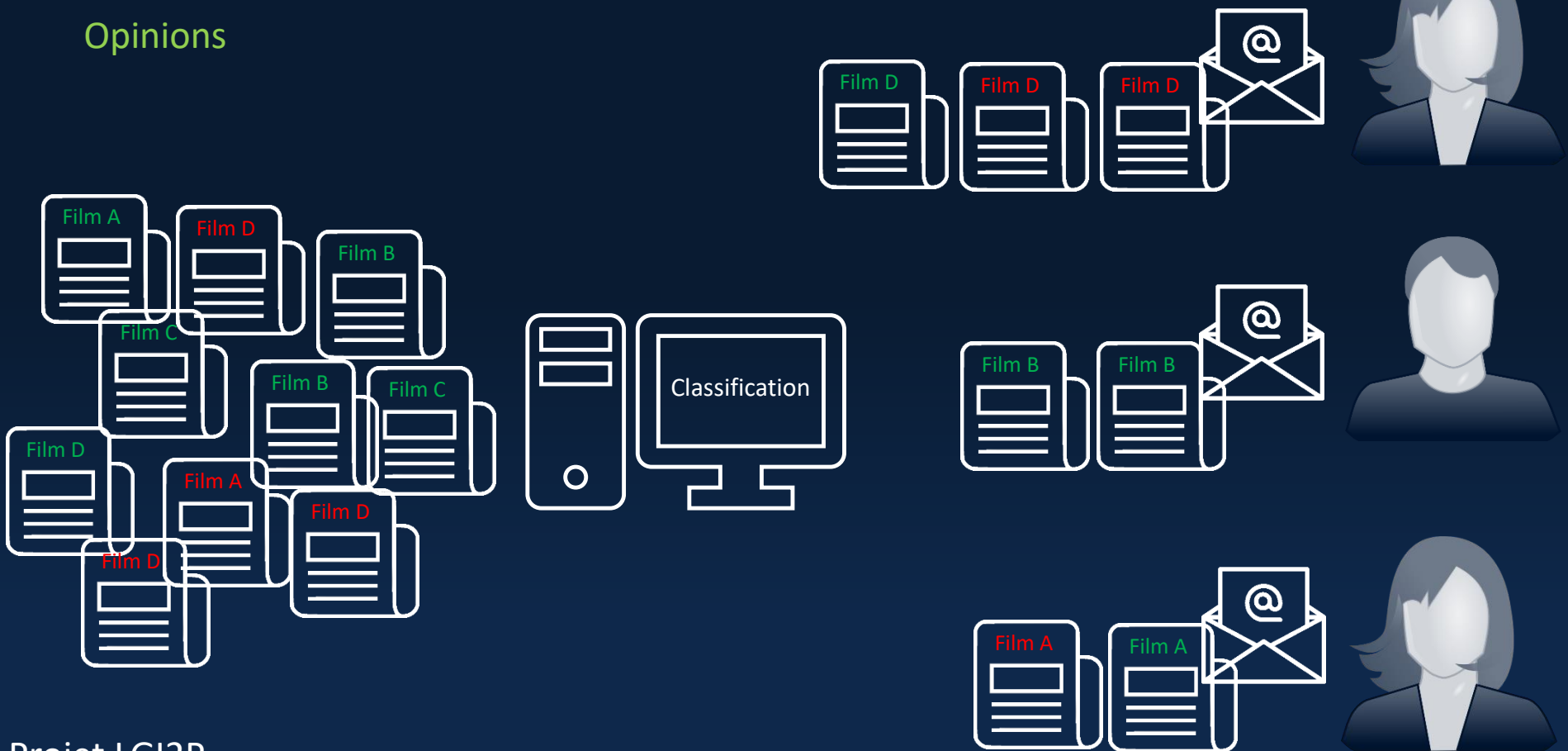
Articles de presse



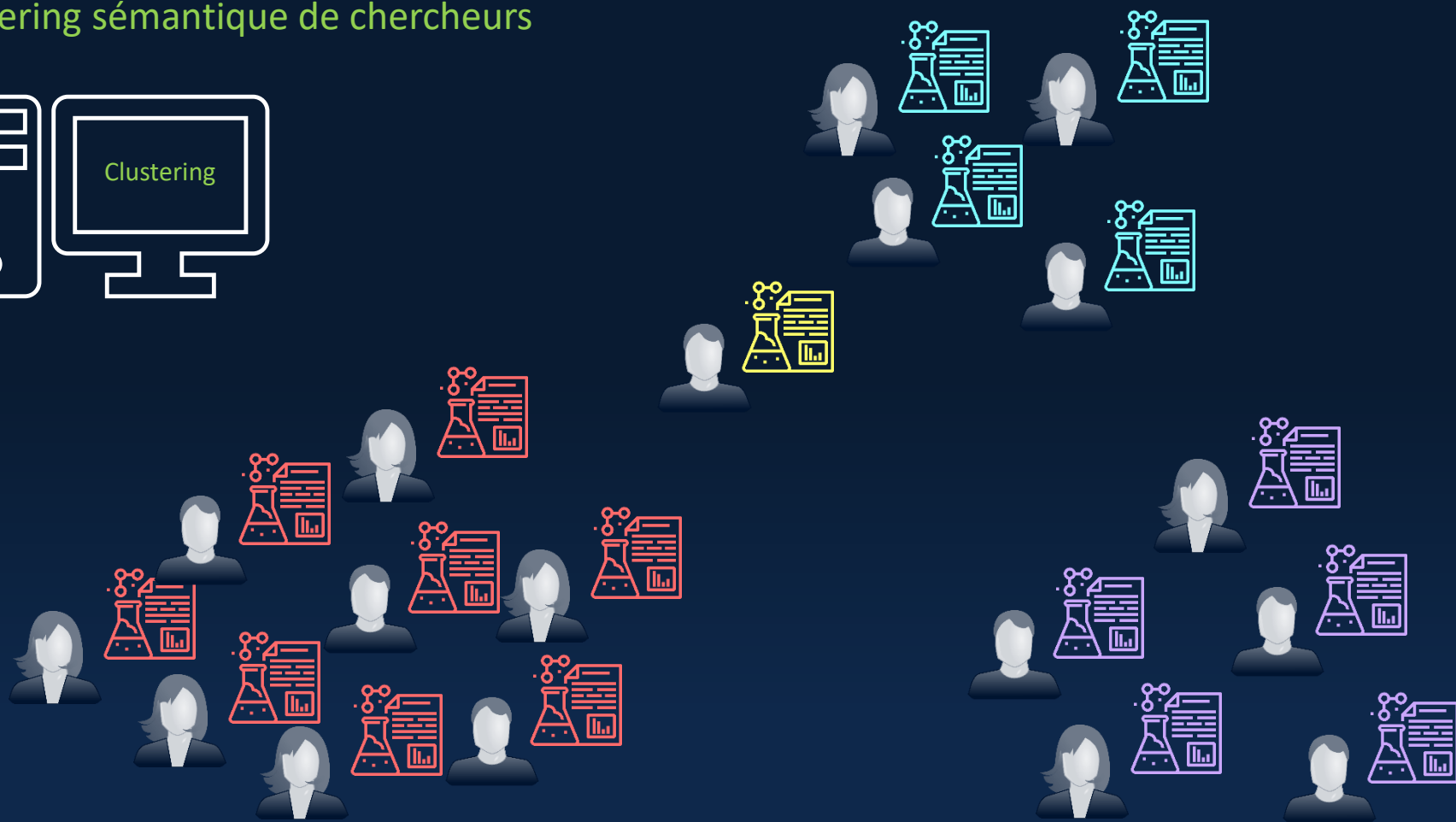
Startup

Classification de textes - Recommandation

Opinions



Clustering sémantique de chercheurs



Organismes de recherche

Classification d'échantillons sanguins

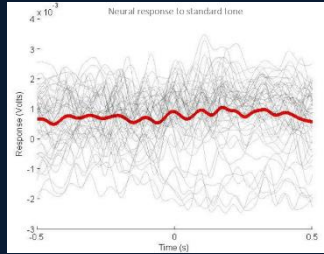


Grand groupe international

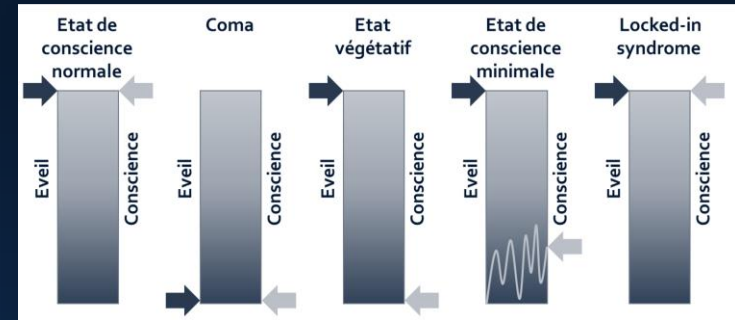
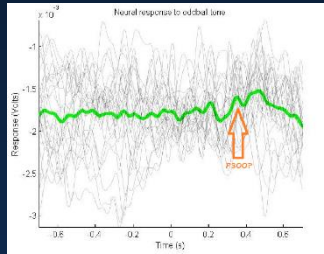
Classification de signaux cérébraux / états de conscience



Fréquent

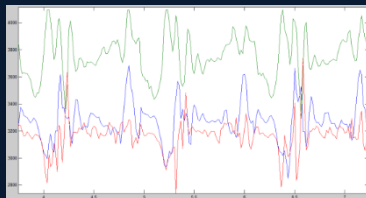


Rare

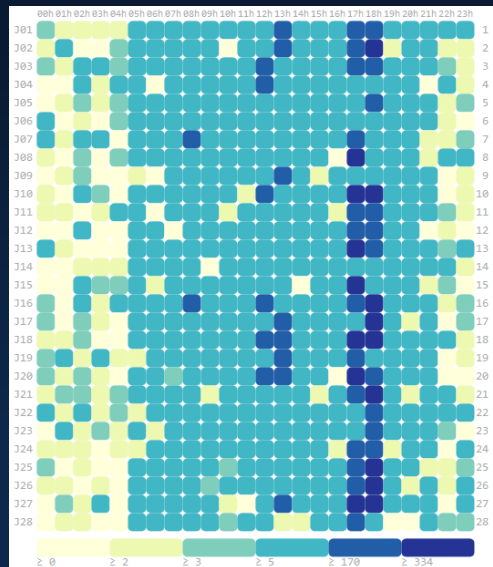


Suivi de l'activité physique

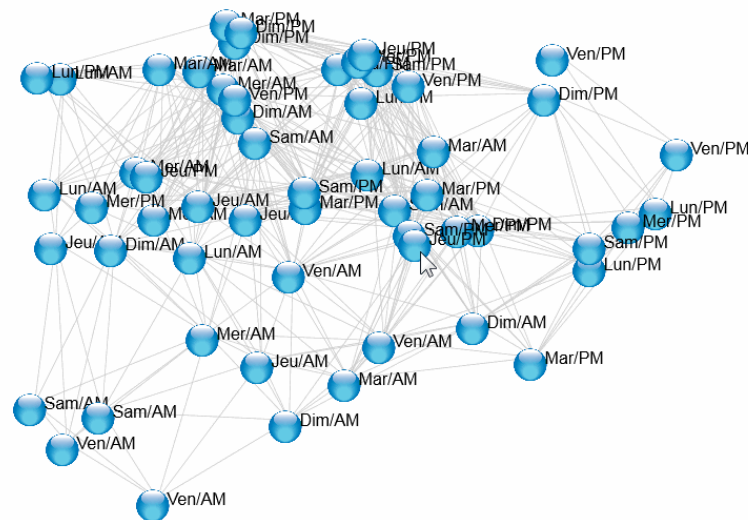
Données



Informations

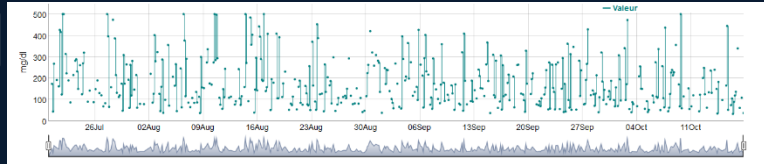


Connaissances

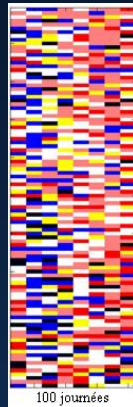


Suivi d'un patient diabétique

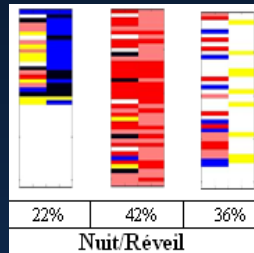
Données



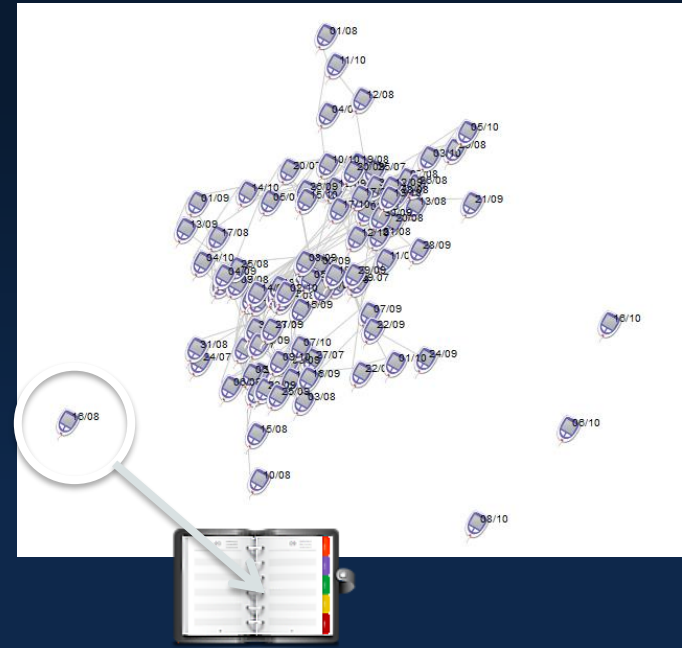
Informations



Clustering

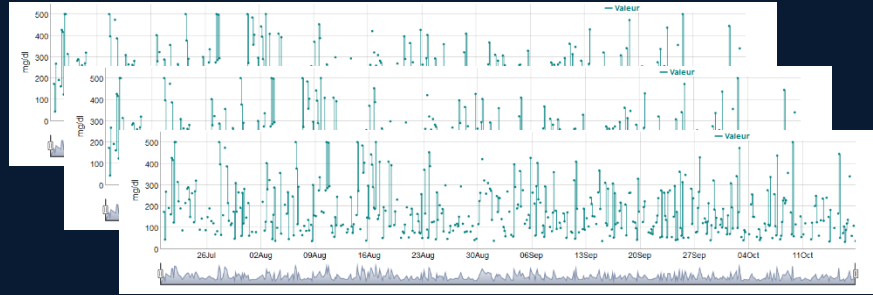


Connaissances

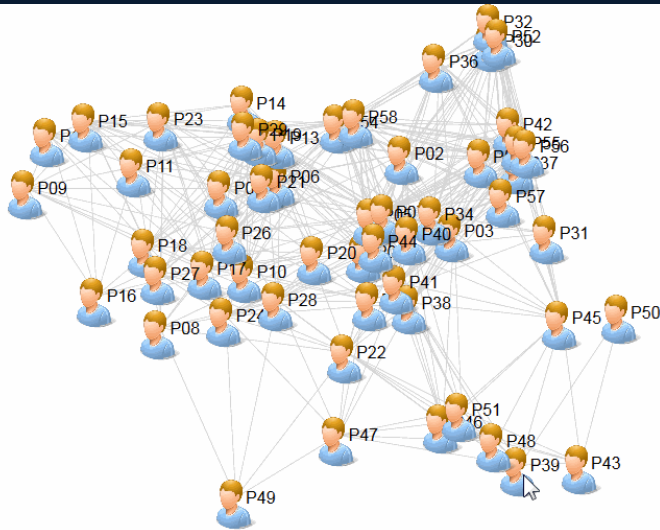


Suivi de patients diabétiques

Données



Informations



Connaissances

En guise de conclusion

A black and white portrait of Jean Piaget, an elderly man with white hair and glasses, wearing a dark suit and tie. He is looking slightly to the right. The portrait is centered on a dark blue background.

"L'intelligence ce n'est pas ce
que l'on sait mais ce qu'on l'on
fait quand on ne sait pas."

JEAN PIAGET

INTELLIGENCE
IS THE ABILITY TO
ADAPT TO CHANGE.

~

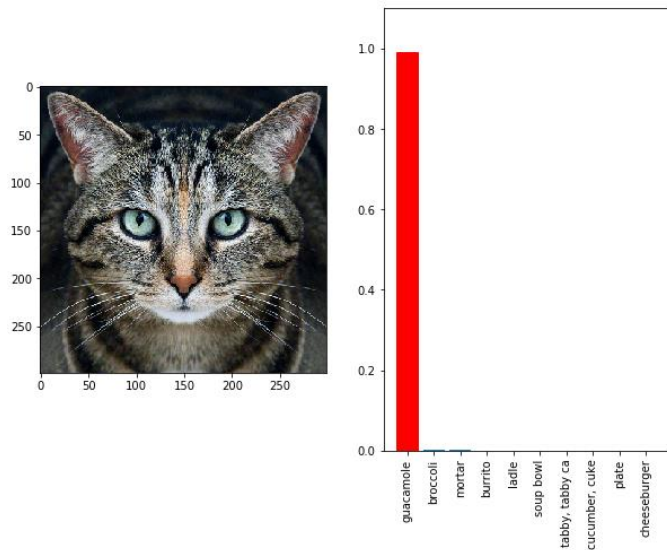
STEPHEN HAWKING



La Trahison des images
René Magritte 1929

Pour aller plus loin

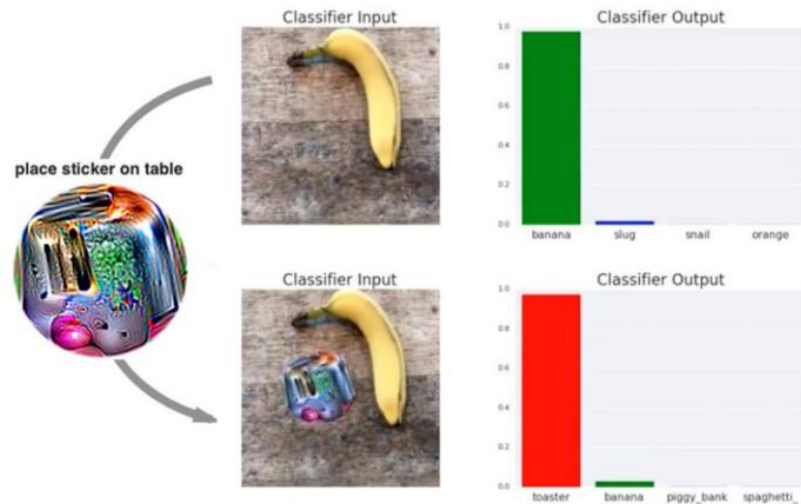
Les réseaux neuronaux se laissent facilement bernier



Computer Vision system fooled into thinking it's a photo of guacamole. (Photo credit MIT CSAIL)

Adversarial Patch

A paper ([Adversarial Patch](#)) recently published by Google has lifted the bar. The paper shows how it is possible to show the system any image and it will classify it as a toaster.



Les réseaux neuronaux se laissent facilement berner



Stop

(a) Normal



Yield



Speed Limit

(b) Attack

Automotive ML systems have been tricked into mistaking stop signs for speed limit signs

Explainable AI (XAI) - Interpretable AI - Explainable Machine Learning (XML),

https://en.wikipedia.org/wiki/Explainable_artificial_intelligence

<https://medium.com/swlh/explainable-vs-interpretable-ai-an-intuitive-example-6baf8fc6d402>

Explainable vs Interpretable AI: An Intuitive Example

Both explainable and interpretable AI are emerging topics in computer science. However, the difference between the two is not always obvious, even to academics. In this post I aim to provide some intuition using a simple example.





SCIENCES • INTELLIGENCE ARTIFICIELLE

IA : « Des humains sont essentiels pour entraîner les systèmes »

Derrière les programmes d'intelligence artificielle, il y a tout un travail de collecte, de vérification et d'annotation des données effectué par des « petites mains ». Une activité essentielle en grande partie externalisée dans les pays en développement, notamment à Madagascar, expliquent deux jeunes sociologues, Maxime Cornet et Clément Le Ludec, dans un entretien au « Monde ».

Propos recueillis par David Larousserie

Publié le 22 juin 2024 à 17h00, modifié le 24 juin 2024 à 16h29 · Lecture 5 min. · [Read in English](#)

vidéo

A Madagascar, les petites mains bien réelles de l'intelligence artificielle

Publié le 29/04/2024 14:21

Temps de lecture : 10min - vidéo : 6min



L'Oeil du 20 heures
France Télévisions

The Size of

LLMs

03/2024

METRIC: PARAMETER COUNT

Total number of trainable variables in a model.

Higher counts often mean greater potential for capturing complex patterns, as well as increased computational resources for training.

GPT 3.5



GPT 4
OpenAI
1.76T

Grok-1
314B



Olympus
amazon
2T

Claude 3
OPUS
ANTHROPIC
2T

G

Gemini
PRO
1.80B



Gemini
ULTRA
Google
1.5T

PaLM-2
DeepMind
340B

Baidu is a Chinese technology company specializing in Internet-related intel, a search engine, AI, and cloud computing.

Ernie 4.0
Baidu
1T



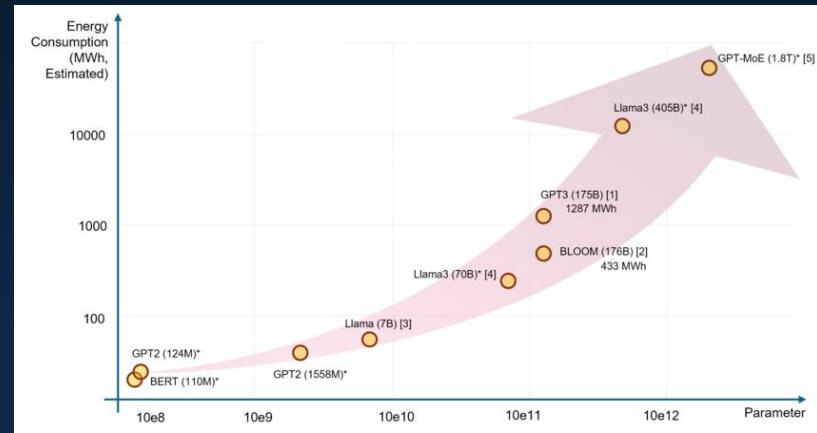
Llama-2
Meta
70B

While both owned by Google, PaLM focuses on multilingual support, reasoning, and coding proficiency, while Gemini is primarily multi-modal and efficient at tool and API integrations.

SOURCE: www.lifearchitect.ai/models

MADE VISUAL

Un billion : mille millions ($1000 \times 10^6 = 10^9$)
Un trillion : mille billions ($1000 \times 10^9 = 10^{12}$)

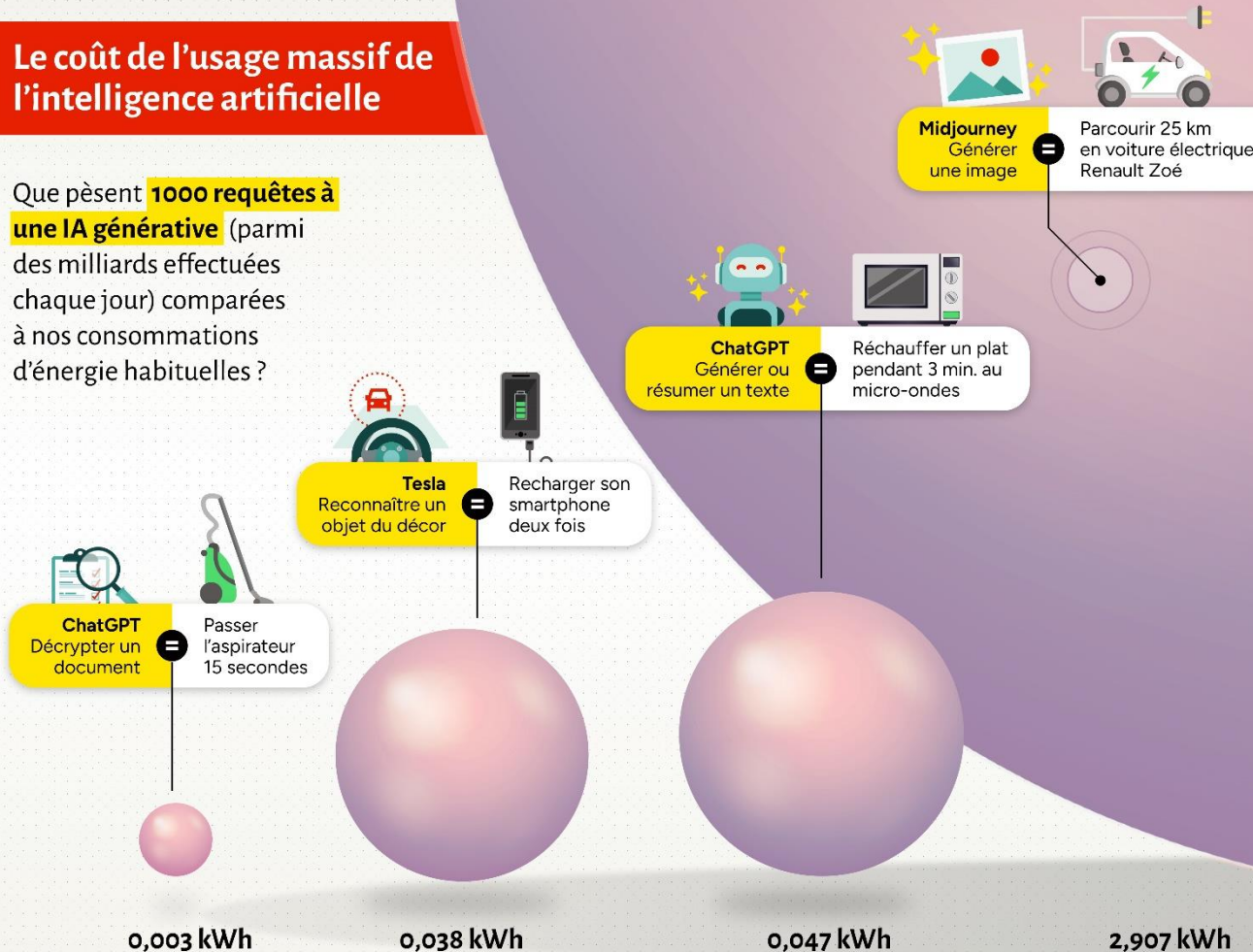


https://www.researchgate.net/figure/Reported-energy-consumption-of-training-different-LLM-models-with-respect-to-model_fig5_384115745

https://www.reddit.com/r/Infographics/comments/1c8ln2o/the_size_of_llms_apr_2024/?dt=60519

Le coût de l'usage massif de l'intelligence artificielle

Que pèsent **1000 requêtes à une IA générative** (parmi des milliards effectuées chaque jour) comparées à nos consommations d'énergie habituelles ?





IA, datacenters et cryptomonnaies font exploser la consommation mondiale d'électricité

Dans quelles proportions la **consommation d'électricité exclusivement liée à l'IA, aux centres de données et aux cryptos** augmente-t-elle dans le monde ?



2022 460 TWh

En 2022, cette consommation était **équivalente à un pays comme la France**



2026 (scénario bas)

+35% (620 TWh)

Soit l'équivalent de **1,3 France en plus** dans le monde



2026 (scénario haut)

+130% (1 050 TWh)

Soit l'équivalent de **2,3 France en plus** dans le monde



IMT Mines Alès
École Mines-Télécom



Master Sciences et Numérique pour la Santé

SCIENCE DES DONNÉES - NIVEAU 1

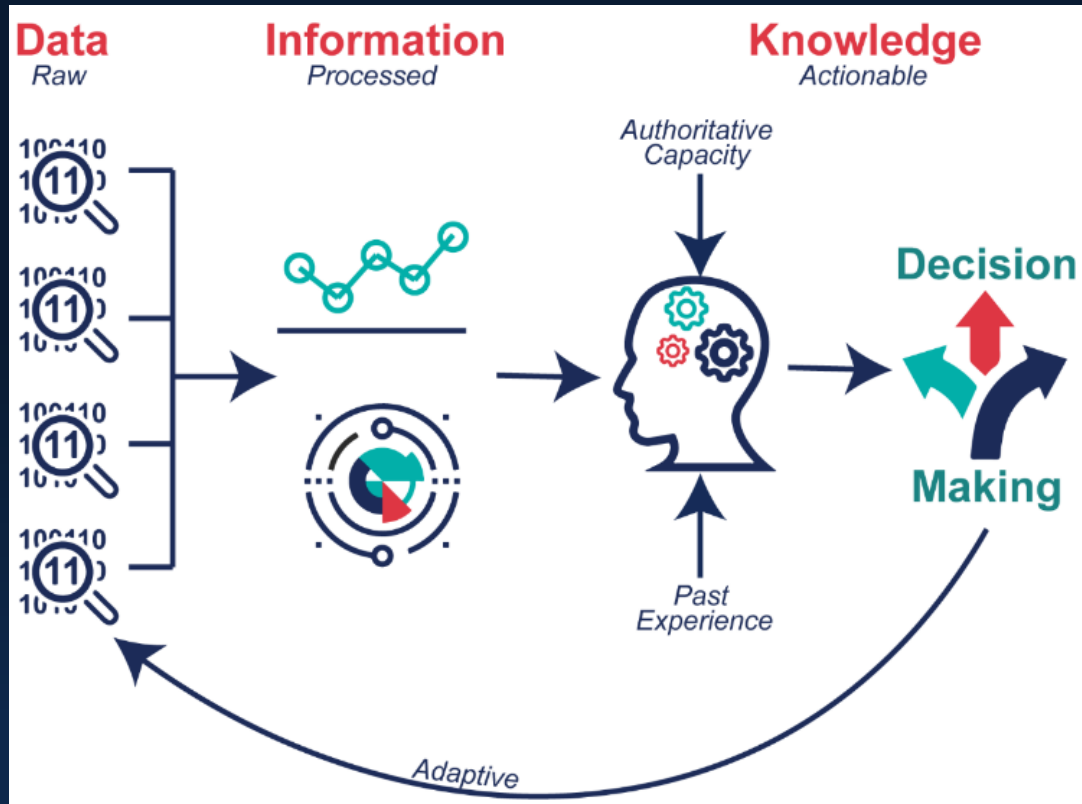
MACHINE LEARNING

Gérard Dray



Machine Learning Process

Data, Information, Knowledge, Decision



Data

Data Frame (Observations x Features) of dimension (n x p)

		Features / Variables / Attributes / Characteristics / ...										
		Age	Income	Gender		Marital Status			...	Feature j	...	Feature p
		X^1	X^2	M	F	Single	Married	Widowed Divorced	...	X^j	...	X^p
				X^3	X^4	X^5	X^6	X^7				
Observations	X_1	x_1^1	x_1^2	x_1^3	x_1^4	x_1^5	x_1^6	x_1^7	...	x_1^j	...	x_1^p
	X_2	x_2^1	x_2^2	x_2^3	x_2^4	x_2^5	x_2^6	x_2^7	...	x_2^j	...	x_2^p
	...	...	...	...	...	...	...	...	...	...	...	...
	X_i	x_i^1	x_i^2	x_i^3	x_i^4	x_i^5	x_i^6	x_i^7	...	x_i^j	...	x_i^p
	...	...	...	...	...	...	...	...	...	...	...	...
	X_n	x_n^1	x_n^2	x_n^3	x_n^4	x_n^5	x_n^6	x_n^7	...	x_n^j	...	x_n^p

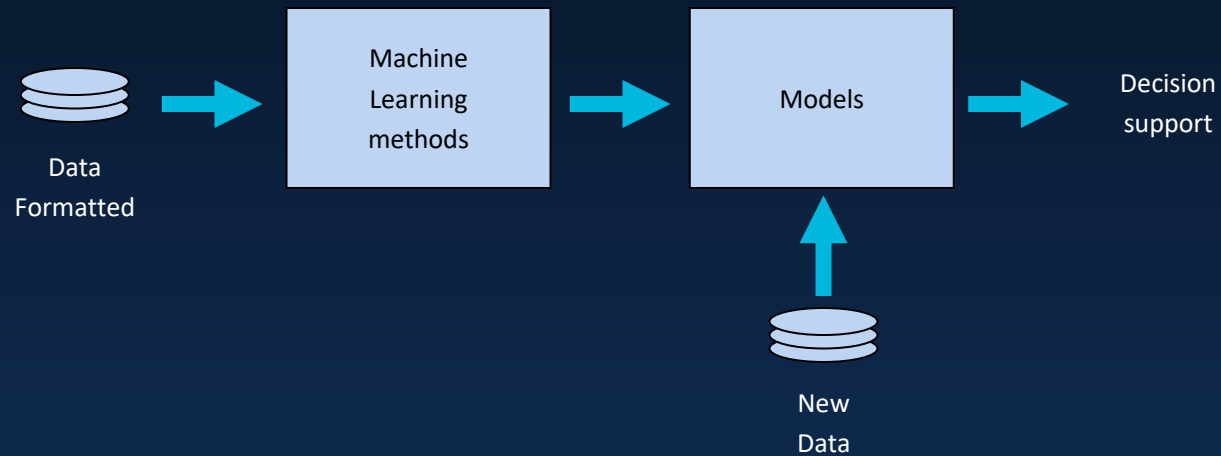
The features Age and Income are quantitative

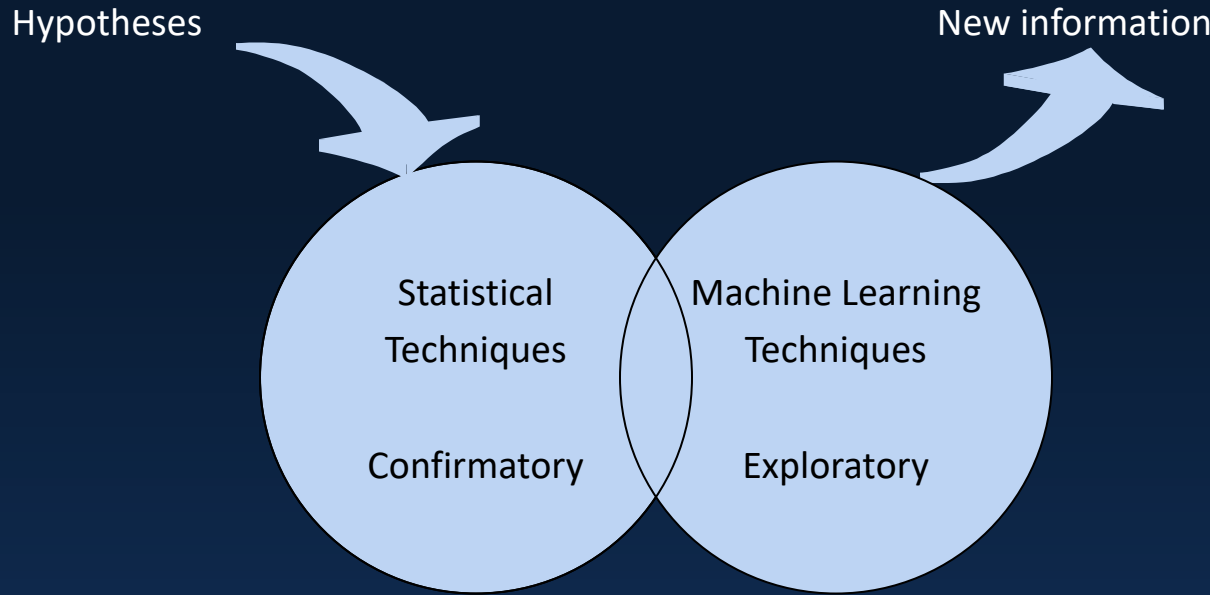
The features Gender and Marital status are qualitative “One-Hot Encoding”

Modality of the features Marital status: (Single, Married, Widowed or Divorced)

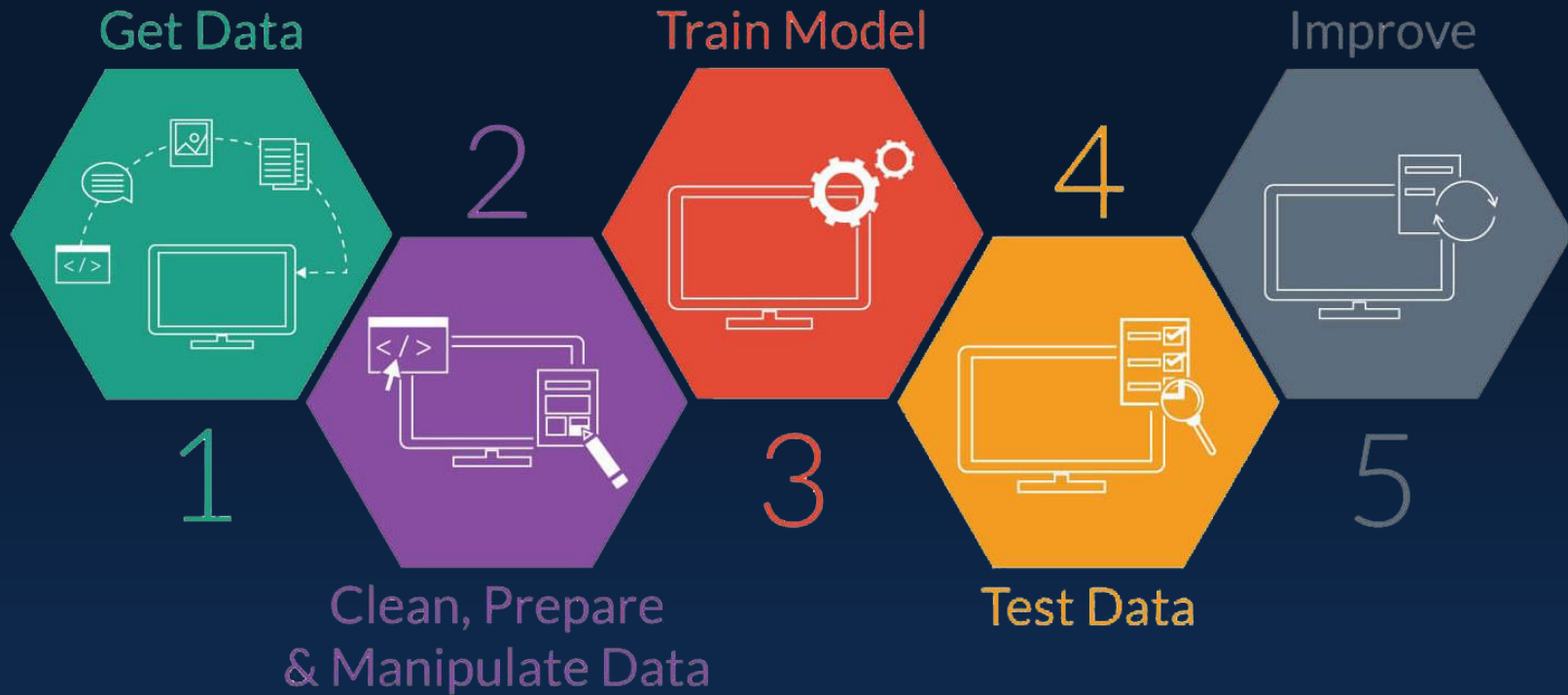
The features : X^3 , X^4 , X^5 , X^6 , and X^7 are Boolean (1 = true, 0 = false)

Data-driven induction process

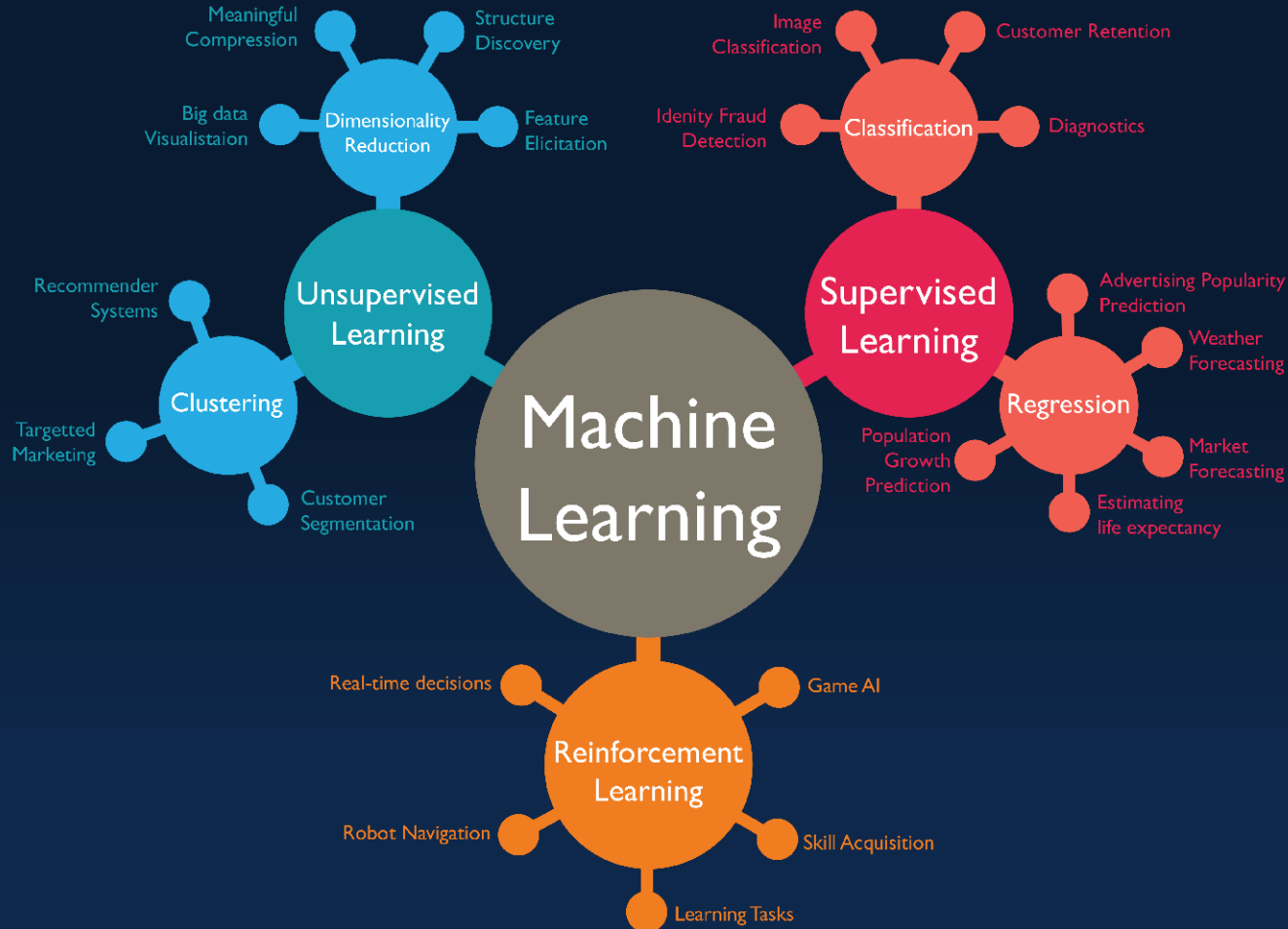




Key stages of machine learning



Machine Learning approaches

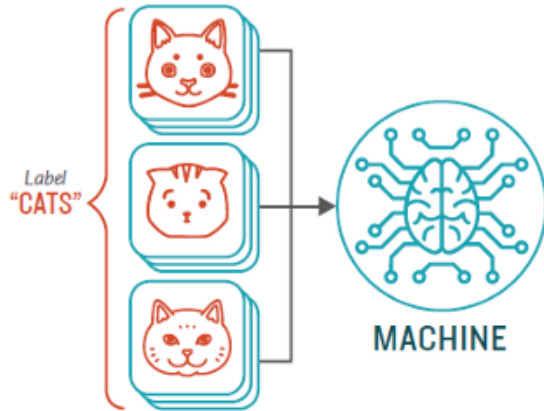


Supervised learning

In supervised learning, we provide the data with the desired outcome (i.e., labeled data). For example, if we want our system to learn to detect cats, we will collect thousands of images, draw a bounding box around the cat, and provide the data set to the machine so that it can learn on its own.

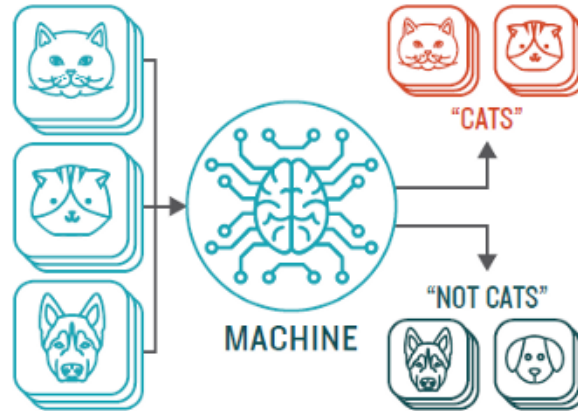
STEP 1

Provide the machine learning algorithm categorized or "labeled" input and output data from to learn



STEP 2

Feed the machine new, unlabeled information to see if it tags new data appropriately. If not, continue refining the algorithm

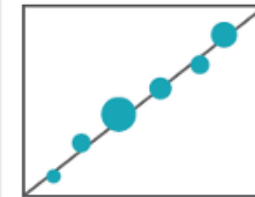


TYPES OF PROBLEMS TO WHICH IT'S SUITED



CLASSIFICATION

Sorting items into categories



REGRESSION

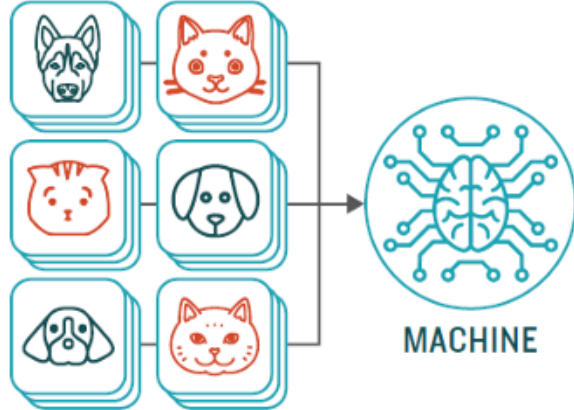
Identifying real values (dollars, weight, etc.)

Unsupervised learning

In unsupervised learning, we simply provide data and let the machine find patterns in the dataset. For example, we can provide three different shapes (circles, triangles, and squares) and let the machine group them together. Such a technique is called clustering.

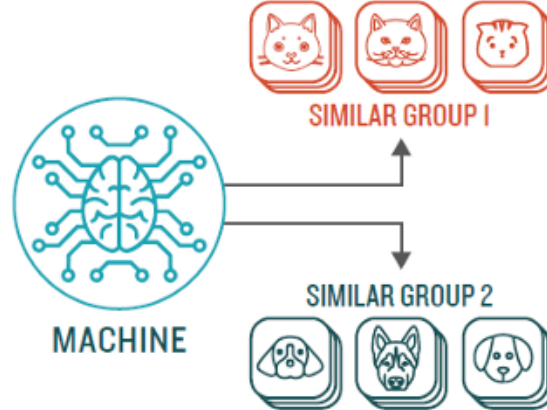
STEP 1

Provide the machine learning algorithm uncategorized, unlabeled input data to see what patterns it finds



STEP 2

Observe and learn from the patterns the machine identifies



TYPES OF PROBLEMS TO WHICH IT'S SUITED

CLUSTERING

Identifying similarities in groups

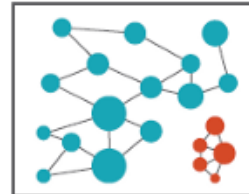
For Example: Are there patterns in the data to indicate certain patients will respond better to this treatment than others?



ANOMALY DETECTION

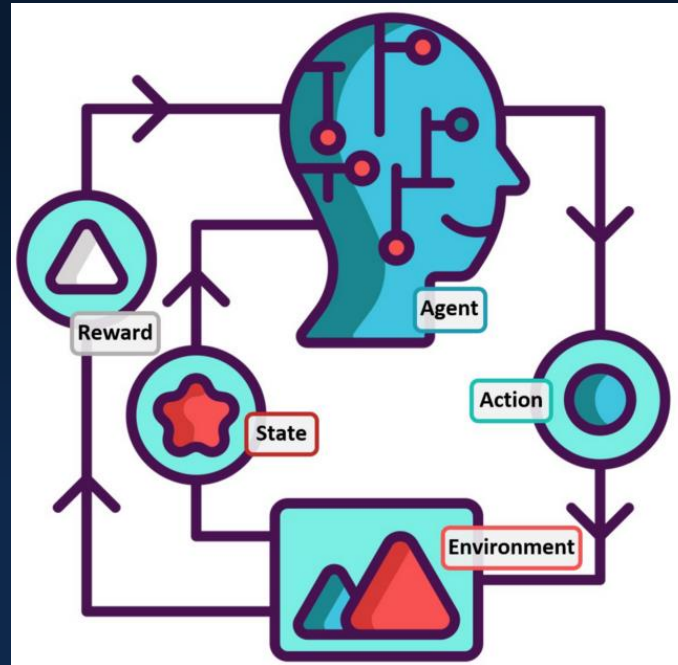
Identifying abnormalities in data

For Example: Is a hacker intruding in our network?



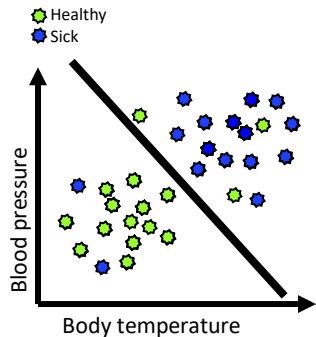
Reinforcement learning

Reinforcement learning uses learning algorithms that learn from repeated experiences through trial and error. It thus replicates the “natural” mechanism of knowledge acquisition. Robots, chatbots, autonomous cars—its applications are numerous in artificial intelligence.



Supervised Machine Learning: Classification

Machine learning: Classification



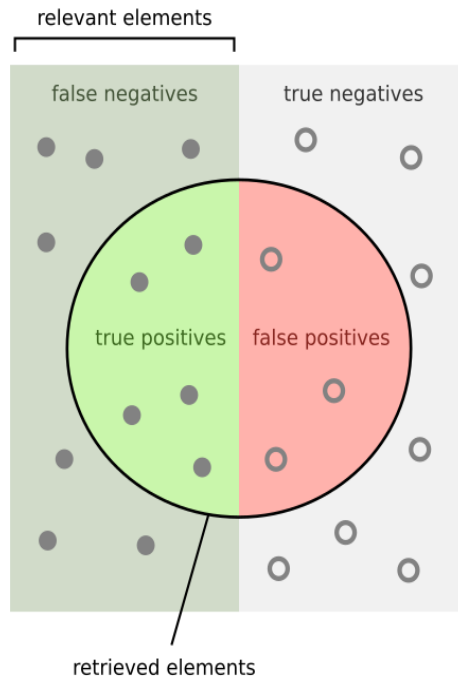
Confusion Matrix

		Predicted Classes	
		Positive	Negative
Real Classes	Positive	True Positives (TP)	False Negatives (FN)
	Negative	False Positives (FP)	True Negatives (TN)

		Predicted Classes	
		Sick	Healthy
Real Classes	Sick	6	2
	Healthy	3	5

Classical algorithms:

- K-nearest neighbors
- Naive Bayes
- Decision tree
- Logistic regression
- Support vector machines
- Neural networks
- ...



How many retrieved items are relevant?

$$\text{Precision} = \frac{\text{true positives}}{\text{true positives} + \text{false positives}}$$

How many relevant items are retrieved?

$$\text{Recall} = \frac{\text{true positives}}{\text{true positives} + \text{false negatives}}$$

$$\text{Precision} = \frac{6}{6 + 3} = 6/9 = 0.67$$

$$\text{Recall} = \frac{6}{6 + 2} = 6/8 = 0.75$$

$$\frac{2 \times \text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} = 0.71$$

Machine learning: Classification

Confusion Matrix

Real	A	TN	TN	TN	TN	FP	TN	TN
	B	TN	TN	TN	TN	FP	TN	TN
	C	TN	TN	TN	TN	FP	TN	TN
	D	TN	TN	TN	TN	FP	TN	TN
	E	FN	FN	FN	FN	TP	FN	FN
	F	TN	TN	TN	TN	FP	TN	TN
	G	TN	TN	TN	TN	FP	TN	TN
		A	B	C	D	E	F	G
		Predicted						

		Predicted Classes			
		Positive	Negative		
Real Classes	Positive	True Positives (TP)	False Negatives (FN)	Sensitivity / TP rate / Recall $\frac{TP}{TP + FN}$	
	Negative	False Positives (FP)	True Negatives (TN)	Specificity $\frac{TN}{TN + FP}$	FP Rate $\frac{FP}{TN + FP}$
		Precision $\frac{TP}{TP + FP}$	TN Rate $\frac{TN}{TN + FN}$		
				Accuracy $\frac{TP + TN}{TP + TN + FP + FN}$	
				F1-Measure / F1 Score $\frac{2 \times Precision \times Recall}{Precision + Recall}$	

Bias-Variance Tradeoff

When we talk about prediction using machine learning models, it's important to understand prediction errors (i.e., bias and variance). The goal of any machine learning model is to find a model that minimizes the prediction errors on unseen data. There's a tradeoff of a model's ability to minimize prediction errors between bias and variance. Understanding these concepts would help us address the issues of overfitting and underfitting.

How many kilometers do you travel each day?



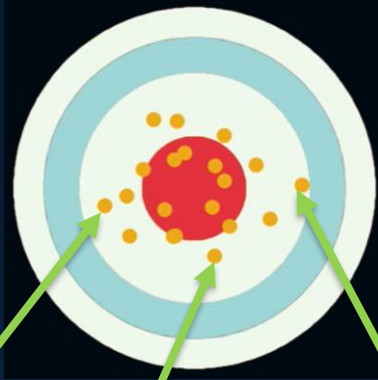
Estimation 1 : 22 km



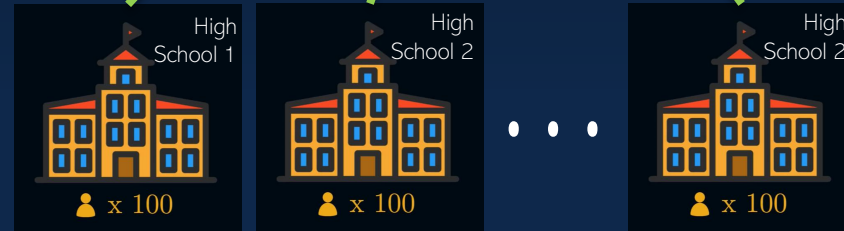
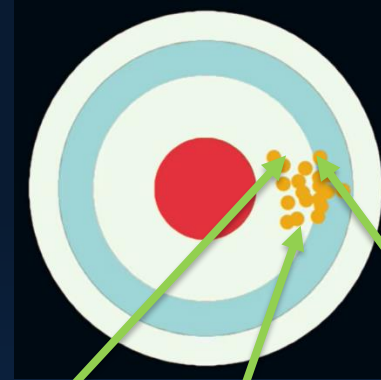
Estimation 2 : 7.1 km

Bias-Variance Tradeoff

High Variance



High Bias

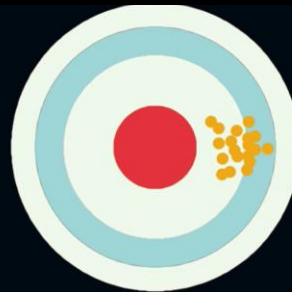


Bias-Variance Tradeoff

Low Variance and Bias



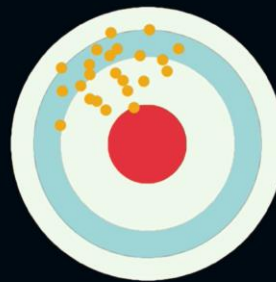
High bias



High Variance



High variance and bias



Bias-Variance Tradeoff

What is bias?

Bias is the difference between our model's average prediction and the correct value we are trying to predict. A model with high bias pays very little attention to the training data and oversimplifies the model. It always leads to high error on both the training and test data.

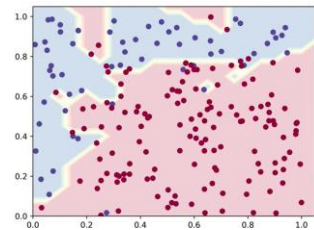
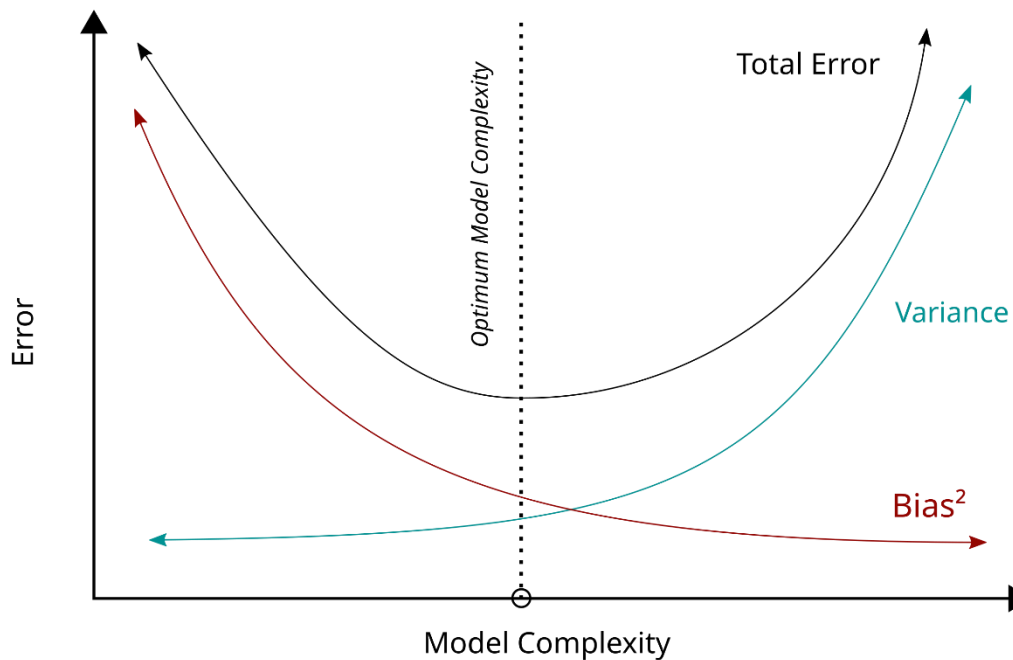
What is variance?

Variance is the variability of the model's prediction, which tells us how spread out our data is. Models with high variance pay close attention to the training data and do not generalize to data they have not seen before. As a result, these models perform very well on the training data but have high error rates on the test data.

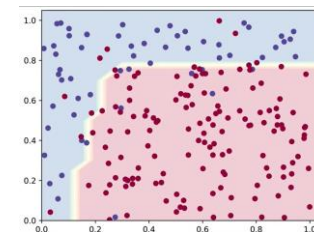
Bias-Variance Tradeoff

$$Err(x) = \left(E[\hat{f}(x)] - f(x)\right)^2 + E\left[\left(\hat{f}(x) - E[\hat{f}(x)]\right)^2\right] + \sigma_e^2$$

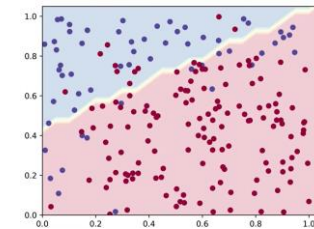
$$Err(x) = \text{Bias}^2 + \text{Variance} + \text{Irreducible Error}$$



overfitting



Good balance

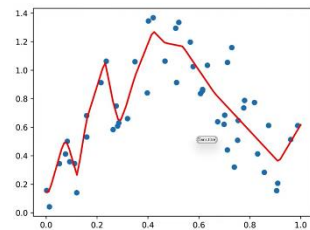
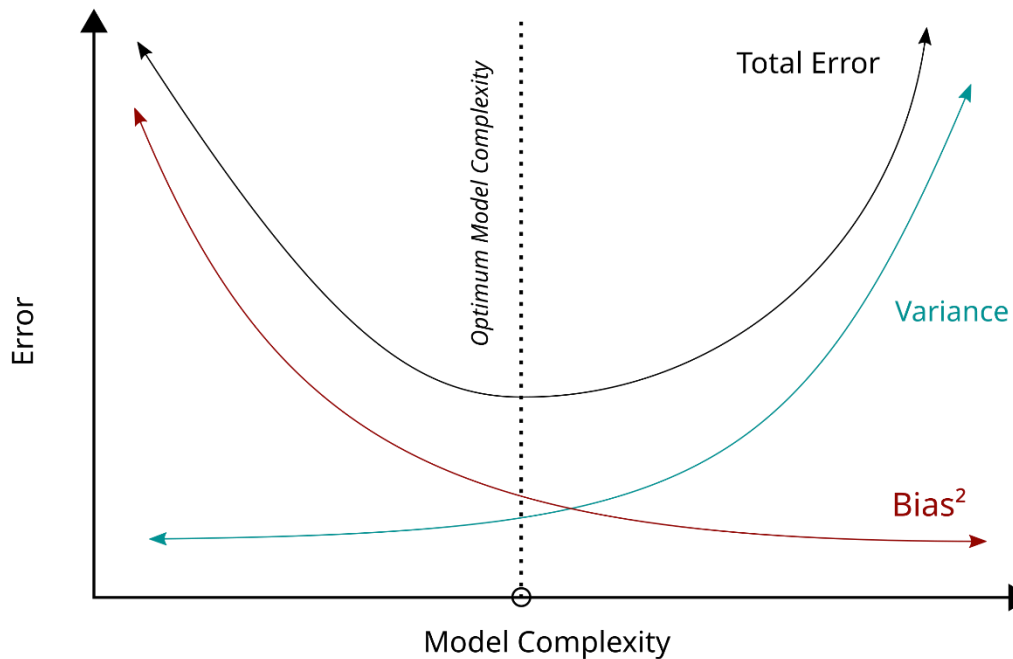


underfitting

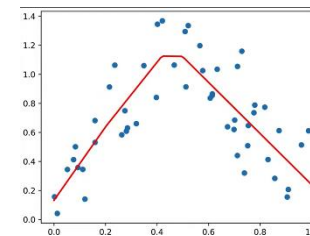
Bias-Variance Tradeoff

$$Err(x) = \left(E[\hat{f}(x)] - f(x)\right)^2 + E\left[\left(\hat{f}(x) - E[\hat{f}(x)]\right)^2\right] + \sigma_e^2$$

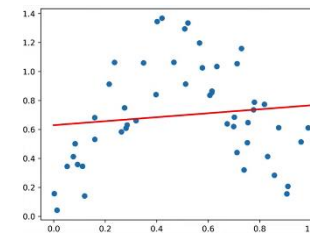
$$Err(x) = \text{Bias}^2 + \text{Variance} + \text{Irreducible Error}$$



overfitting

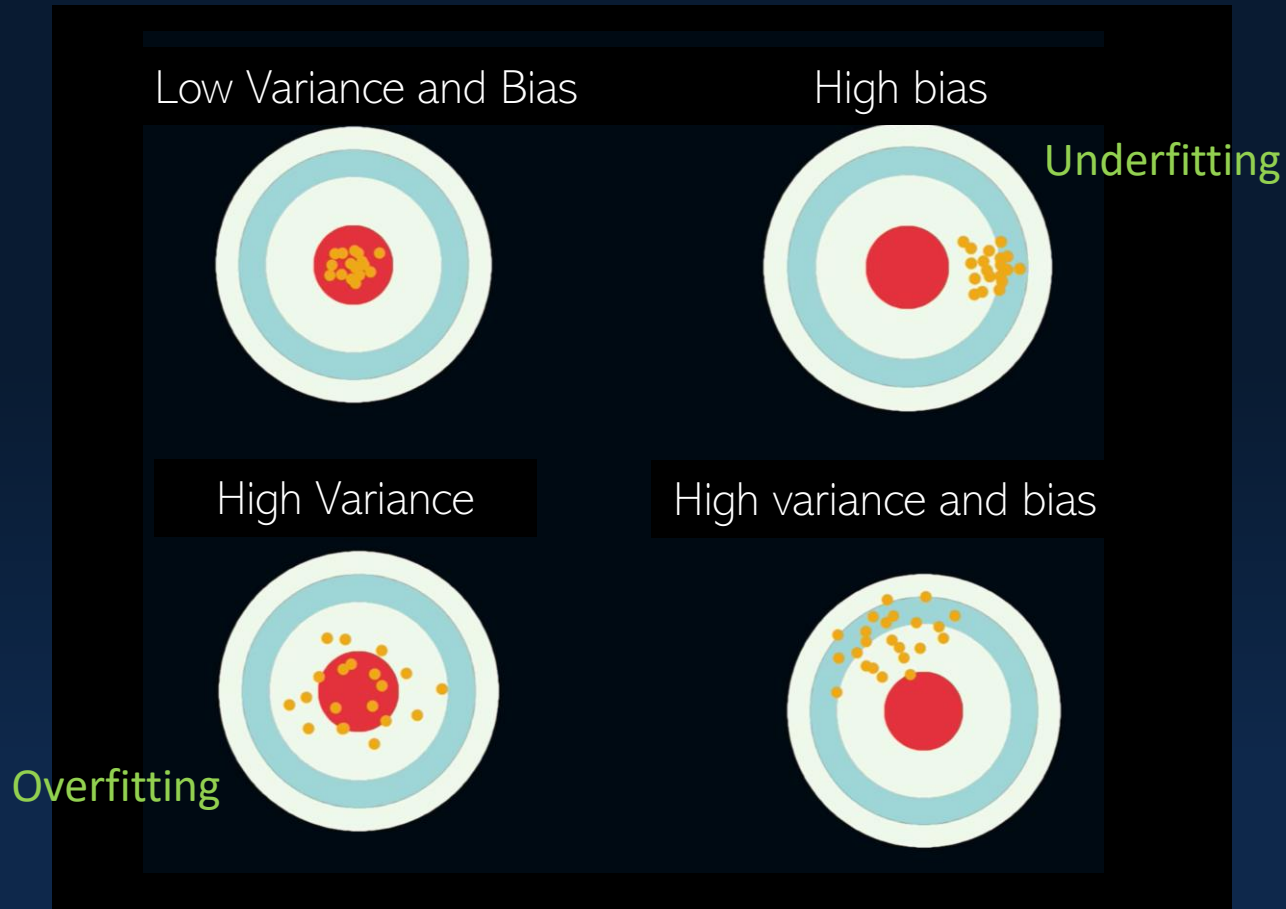


Good balance



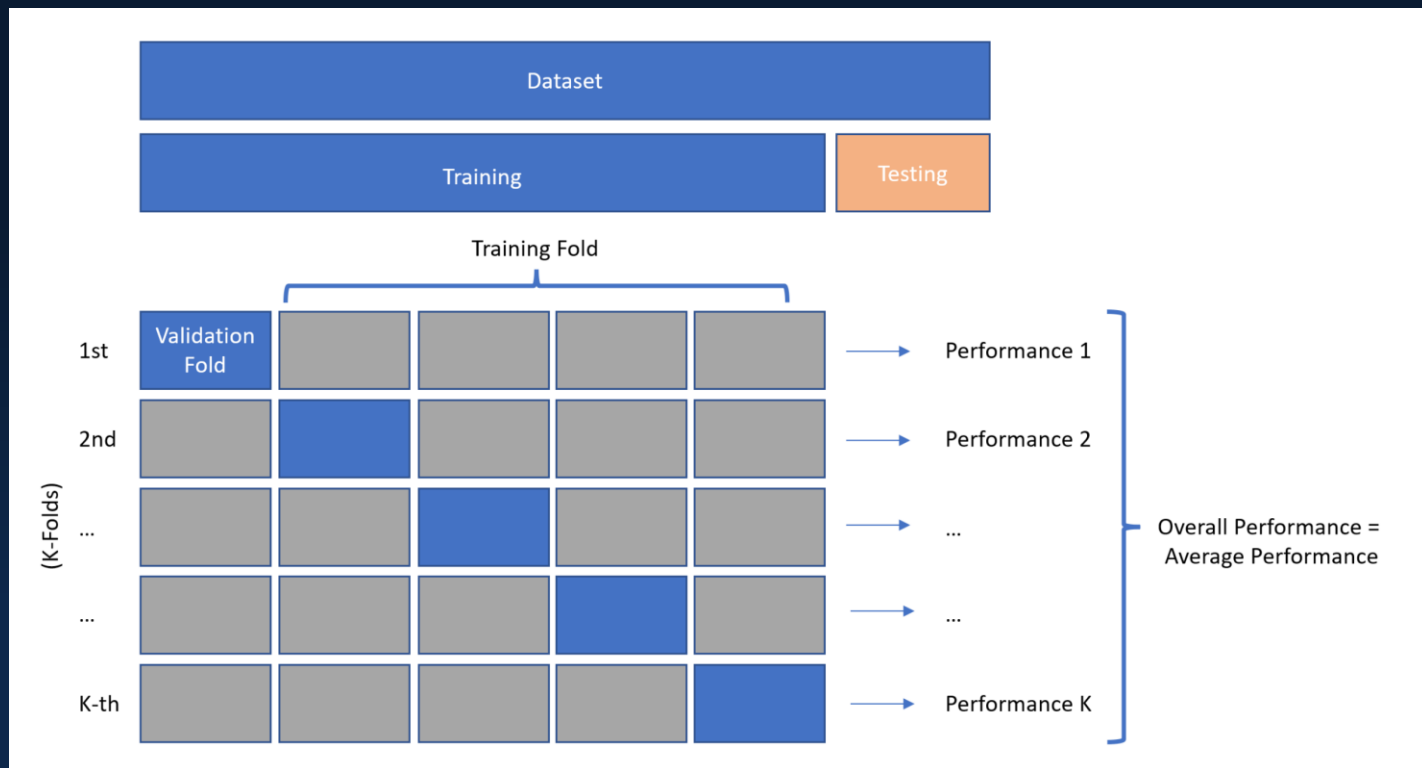
underfitting

Bias-Variance Tradeoff

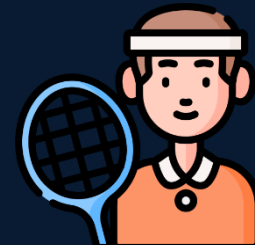


Bias-Variance Tradeoff

K-Fold Cross-Validation is a popular tool used in training a machine learning model. The idea of this method is simple, we create a multiple set of training data from the observed data, then train the model and evaluate the model based on the validation data, which is similar to evaluating the model on the unseen data.



Apprentissage automatique : Classification



Am I going to play tennis?

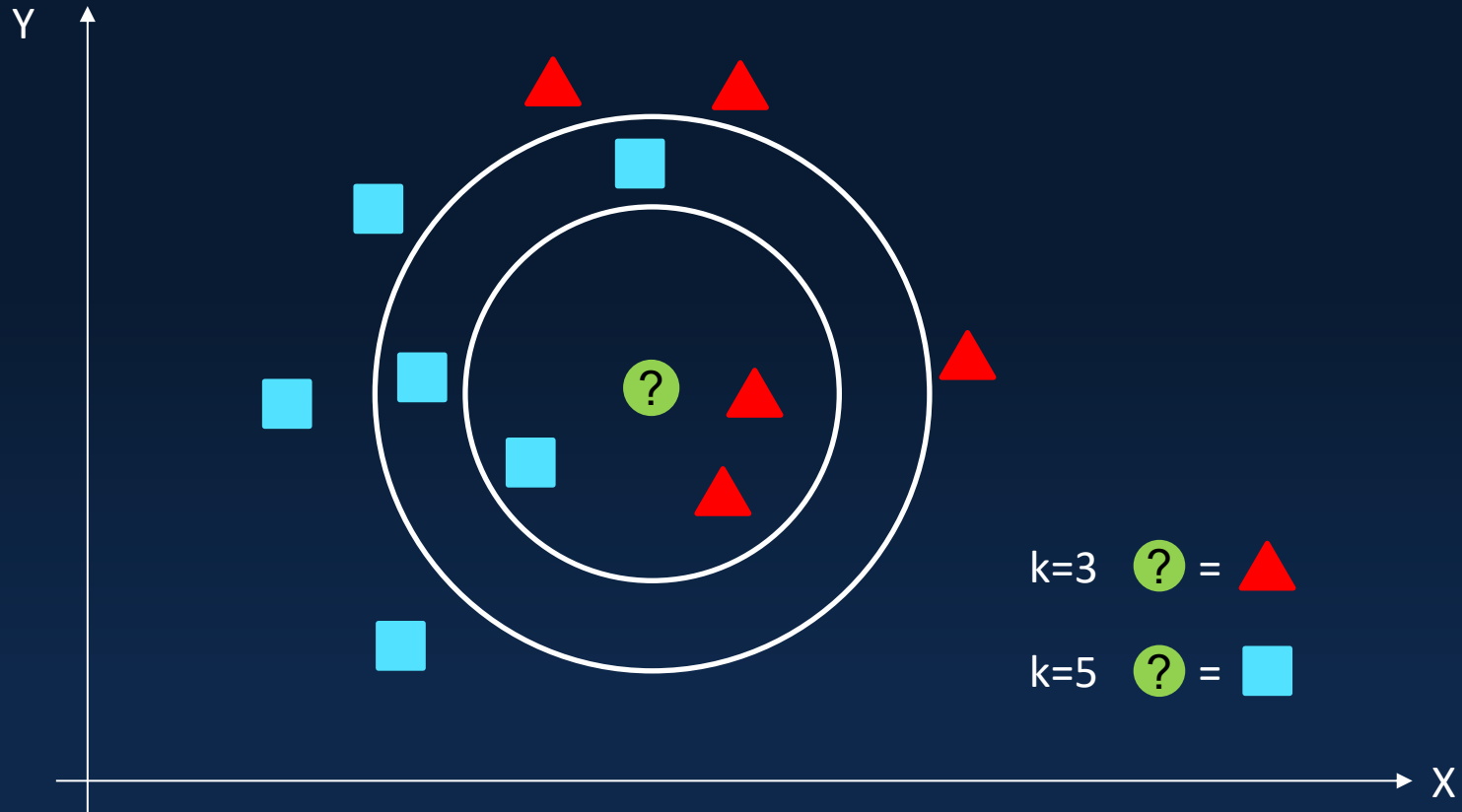
Day	Outlook	Temperature °C	Temperature	Humidity %	Humidity	Windy	Play
1	Sunny	29	Hot	85	High	False	No
2	Sunny	27	Hot	90	High	True	No
3	Overcast	28	Hot	86	High	False	Yes
4	Rainy	21	Mild	96	High	False	Yes
5	Rainy	20	Cool	80	Normal	False	Yes
6	Rainy	18	Cool	70	Normal	True	No
7	Overcast	18	Cool	65	Normal	True	Yes
8	Sunny	22	Mild	95	High	False	No
9	Sunny	21	Cool	70	Normal	False	Yes
10	Rainy	24	Mild	80	Normal	False	Yes
11	Sunny	24	Mild	70	Normal	True	Yes
12	Overcast	22	Mild	90	High	True	Yes
13	Overcast	27	Hot	75	Normal	False	Yes
14	Rainy	22	Mild	91	High	True	No
15	Sunny	19	Cool	75	Normal	True	?
16	Overcast	22	Mild	89	Normal	False	?

Supervised Machine Learning: Classification

Method: k-Nearest Neighbors

Machine learning: Classification

k Nearest Neighbors (kNN)



Machine learning: Classification

k Nearest Neighbors (kNN)

Day	Outlook	Temperature	Humidity	Windy	Outlook Sunny	Outlook Overcast	Outlook Rainy	Temperature Hot	Temperature Mild	Temperature Cool	Humidity High	Humidity Normal	Windy True	Windy False	Play
1	Sunny	Hot	High	False	1	0	0	1	0	0	1	0	0	1	No
2	Sunny	Hot	High	True	1	0	0	1	0	0	1	0	1	0	No
3	Overcast	Hot	High	False	0	1	0	1	0	0	1	0	0	1	Yes
4	Rainy	Mild	High	False	0	0	1	0	1	0	1	0	0	1	Yes
5	Rainy	Cool	Normal	False	0	0	1	0	0	1	0	1	0	1	Yes
6	Rainy	Cool	Normal	True	0	0	1	0	0	1	0	1	1	0	No
7	Overcast	Cool	Normal	True	0	1	0	0	0	1	0	1	1	0	Yes
8	Sunny	Mild	High	False	1	0	0	0	1	0	1	0	0	1	No
9	Sunny	Cool	Normal	False	1	0	0	0	0	1	0	1	0	1	Yes
10	Rainy	Mild	Normal	False	0	0	1	0	1	0	0	1	0	1	Yes
11	Sunny	Mild	Normal	True	1	0	0	0	1	0	0	1	1	0	Yes
12	Overcast	Mild	High	True	0	1	0	0	1	0	1	0	1	0	Yes
13	Overcast	Hot	Normal	False	0	1	0	1	0	0	0	1	0	1	Yes
14	Rainy	Mild	High	True	0	0	1	0	1	0	1	0	1	0	No
15	Sunny	Cool	Normal	True	1	0	0	0	0	1	0	1	1	0	?
16	Overcast	Mild	Normal	False	0	1	0	0	1	0	0	1	0	1	?

Coding of nominal features (1)
“One-Hot Encoding”

Machine learning: Classification

k Nearest Neighbors (kNN)

Day	Outlook Sunny	Outlook Overcast	Outlook Rainy	Temperature Hot	Temperature Mild	Temperature Cool	Humidity High	Humidity Normal	Windy True	Windy False	Play	Dist. Hamming Day 15
1	1	0	0	1	0	0	1	0	0	1	No	6
2	1	0	0	1	0	0	1	0	1	0	No	4
3	0	1	0	1	0	0	1	0	0	1	Yes	8
4	0	0	1	0	1	0	1	0	0	1	Yes	8
5	0	0	1	0	0	1	0	1	0	1	Yes	4
6	0	0	1	0	0	1	0	1	1	0	No	2
7	0	1	0	0	0	1	0	1	1	0	Yes	2
8	1	0	0	0	1	0	1	0	0	1	No	6
9	1	0	0	0	0	1	0	1	0	1	Yes	2
10	0	0	1	0	1	0	0	1	0	1	Yes	6
11	1	0	0	0	1	0	0	1	1	0	Yes	2
12	0	1	0	0	1	0	1	0	1	0	Yes	6
13	0	1	0	1	0	0	0	1	0	1	Yes	6
14	0	0	1	0	1	0	1	0	1	0	No	6

Hamming distance

$$D_{Hamming}(X, Y) = \sum_{i=1}^k |x_i - y_i| = |1 - 1| + |0 - 0| + |0 - 0| + |1 - 0| + |0 - 0| + |0 - 1| + |1 - 0| + |0 - 1| + |0 - 1| + |1 - 0| = 6$$

Day	Outlook Sunny	Outlook Overcast	Outlook Rainy	Temperature Hot	Temperature Mild	Temperature Cool	Humidity High	Humidity Normal	Windy True	Windy False
1	1	0	0	1	0	0	1	0	0	1
15	1	0	0	0	0	1	0	1	1	0

Machine learning: Classification

k Nearest Neighbors (kNN)

Day	Outlook Sunny	Outlook Overcast	Outlook Rainy	Temperature Hot	Temperature Mild	Temperature Cool	Humidity High	Humidity Normal	Windy True	Windy False	Play	Dist. Hamming Day 15
1	1	0	0	1	0	0	1	0	0	1	No	6
2	1	0	0	1	0	0	1	0	1	0	No	4
3	0	1	0	1	0	0	1	0	0	1	Yes	8
4	0	0	1	0	1	0	1	0	0	1	Yes	8
5	0	0	1	0	0	1	0	1	0	1	Yes	4
6	0	0	1	0	0	1	0	1	1	0	No	2
7	0	1	0	0	0	1	0	1	1	0	Yes	2
8	1	0	0	0	1	0	1	0	0	1	No	6
9	1	0	0	0	0	1	0	1	0	1	Yes	2
10	0	0	1	0	1	0	0	1	0	1	Yes	6
11	1	0	0	0	1	0	0	1	1	0	Yes	2
12	0	1	0	0	1	0	1	0	1	0	Yes	6
13	0	1	0	1	0	0	0	1	0	1	Yes	6
14	0	0	1	0	1	0	1	0	1	0	No	6

Hamming distance

$$D_{Hamming}(X, Y) = \sum_{i=1}^k |x_i - y_i| = |1 - 1| + |0 - 0| + |0 - 0| + |0 - 0| + |1 - 0| + |0 - 1| + |0 - 0| + |1 - 1| + |1 - 1| + |0 - 0| = 2$$

Day	Outlook Sunny	Outlook Overcast	Outlook Rainy	Temperature Hot	Temperature Mild	Temperature Cool	Humidity High	Humidity Normal	Windy True	Windy False
11	1	0	0	0	1	0	0	1	1	0
15	1	0	0	0	0	1	0	1	1	0

Machine learning: Classification

k Nearest Neighbors (kNN)

Day 15

Day	Play	Dist. Hamming Day 15
6	No	2
7	Yes	2
9	Yes	2
11	Yes	2
2	No	4
5	Yes	4
1	No	6
8	No	6
10	Yes	6
12	Yes	6
13	Yes	6
14	No	6
3	Yes	8
4	Yes	8
15	?	0

k = 1 ?

k = 3 ?

k = 5 ?

Machine learning: Classification

k Nearest Neighbors (kNN)

Day 16

Day	Play	Dist. Hamming Day 16
10	Yes	2
13	Yes	2
7	Yes	4
9	Yes	4
11	Yes	4
5	Yes	4
8	No	4
12	Yes	4
3	Yes	4
4	Yes	4
6	No	6
1	No	6
14	No	6
2	No	8
16	?	0

k = 1 ?

k = 3 ?

k = 5 ?

Machine learning: Classification

k Nearest Neighbors (kNN)

Day	Outlook	Windy	Outlook Sunny	Outlook Overcast	Outlook Rainy	Temperature °C	Humidity %	Windy True	Windy False	Play
1	Sunny	False	1	0	0	29	85	0	1	No
2	Sunny	True	1	0	0	27	90	1	0	No
3	Overcast	False	0	1	0	28	86	0	1	Yes
4	Rainy	False	0	0	1	21	96	0	1	Yes
5	Rainy	False	0	0	1	20	80	0	1	Yes
6	Rainy	True	0	0	1	18	70	1	0	No
7	Overcast	True	0	1	0	18	65	1	0	Yes
8	Sunny	False	1	0	0	22	95	0	1	No
9	Sunny	False	1	0	0	21	70	0	1	Yes
10	Rainy	False	0	0	1	24	80	0	1	Yes
11	Sunny	True	1	0	0	24	70	1	0	Yes
12	Overcast	True	0	1	0	22	90	1	0	Yes
13	Overcast	False	0	1	0	27	75	0	1	Yes
14	Rainy	True	0	0	1	22	91	1	0	No
15	Sunny	True	1	0	0	19	75	1	0	?
16	Overcast	False	0	1	0	22	89	0	1	?

Coding of nominal features (2)

Machine learning: Classification

k Nearest Neighbors (kNN)

Day	Outlook Sunny	Outlook Overcast	Outlook Rainy	Temperature	Humidity	Windy True	Windy False	Play	Dist. Euclidienne Day 15
1	1	0	0	29	85	0	1	No	14.21
2	1	0	0	27	90	1	0	No	17.00
3	0	1	0	28	86	0	1	Yes	14.35
4	0	0	1	21	96	0	1	Yes	21.19
5	0	0	1	20	80	0	1	Yes	5.48
6	0	0	1	18	70	1	0	No	5.29
7	0	1	0	18	65	1	0	Yes	10.15
8	1	0	0	22	95	0	1	No	20.27
9	1	0	0	21	70	0	1	Yes	5.57
10	0	0	1	24	80	0	1	Yes	7.35
11	1	0	0	24	70	1	0	Yes	7.07
12	0	1	0	22	90	1	0	Yes	15.36
13	0	1	0	27	75	0	1	Yes	8.25
14	0	0	1	22	91	1	0	No	16.34

Euclidean distance

$$D_{\text{euclidienne}}(X, Y) = \sqrt{\sum_{i=1}^k (x_i - y_i)^2} = \sqrt{(1 - 1)^2 + (0 - 0)^2 + (0 - 0)^2 + (29 - 19)^2 + (85 - 75)^2 + (0 - 1)^2 + (1 - 0)^2} = 14.21$$

Day	Outlook Sunny	Outlook Overcast	Outlook Rainy	Temperature	Humidity	Windy True	Windy False
1	1	0	0	29	85	0	1
15	1	0	0	19	75	1	0

Machine learning: Classification

k Nearest Neighbors (kNN)

Day	Outlook Sunny	Outlook Overcast	Outlook Rainy	Temperature	Humidity	Windy True	Windy False	Play	Dist. Euclidienne Day 15
1	1	0	0	29	85	0	1	No	14.21
2	1	0	0	27	90	1	0	No	17.00
3	0	1	0	28	86	0	1	Yes	14.35
4	0	0	1	21	96	0	1	Yes	21.19
5	0	0	1	20	80	0	1	Yes	5.48
6	0	0	1	18	70	1	0	No	5.29
7	0	1	0	18	65	1	0	Yes	10.15
8	1	0	0	22	95	0	1	No	20.27
9	1	0	0	21	70	0	1	Yes	5.57
10	0	0	1	24	80	0	1	Yes	7.35
11	1	0	0	24	70	1	0	Yes	7.07
12	0	1	0	22	90	1	0	Yes	15.36
13	0	1	0	27	75	0	1	Yes	8.25
14	0	0	1	22	91	1	0	No	16.34

Euclidean distance

$$D_{euclidienne}(X, Y) = \sqrt{\sum_{i=1}^k (x_i - y_i)^2} = \sqrt{(1 - 1)^2 + (0 - 0)^2 + (0 - 0)^2 + (24 - 19)^2 + (70 - 75)^2 + (1 - 1)^2 + (0 - 0)^2} = 7.07$$

Day	Outlook Sunny	Outlook Overcast	Outlook Rainy	Temperature	Humidity	Windy True	Windy False
11	1	0	0	24	70	1	0
15	1	0	0	19	75	1	0

Machine learning: Classification

k Nearest Neighbors (kNN)

Day 15

Day	Play	Dist. Euclidienne Day 15
6	No	5.29
5	Yes	5.48
9	Yes	5.57
11	Yes	7.07
10	Yes	7.35
13	Yes	8.25
7	Yes	10.15
1	No	14.21
3	Yes	14.35
12	Yes	15.36
14	No	16.34
2	No	17.00
8	No	20.27
4	Yes	21.19
15	?	0.00

k = 1 ?

k = 3 ?

k = 5 ?

Machine learning: Classification

k Nearest Neighbors (kNN)

Day 16

Day	Play	Dist. Euclidienne Day 16
12	Yes	1.73
14	No	2.83
2	No	5.48
8	No	6.16
3	Yes	6.71
4	Yes	7.21
1	No	8.19
5	Yes	9.33
10	Yes	9.33
13	Yes	14.87
9	Yes	19.08
11	Yes	19.21
6	No	19.52
7	Yes	24.37
16	?	0.00

k = 1 ?

k = 3 ?

k = 5 ?

Supervised Machine Learning: Classification

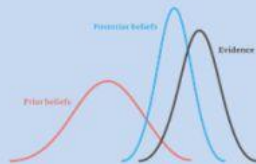
Method: Naives Bayes

Machine learning: Classification

Naives Bayes Method

Naive Bayes

Supervised Learning Algorithm for Classification



Bayes Theorem :

$$P(B_i|A) = \frac{P(B_i) \times P(A|B_i)}{P(A)}$$

$$= \frac{P(B_i) \times P(A|B_i)}{\sum_{j=1}^k P(B_j) \times P(A|B_j)}$$



Thomas Bayes 1702-1762

Probabilistic model

$$P(C_k|x) = \frac{P(C_k) \times P(x|C_k)}{P(x)} \propto p(C_k, x_1, x_2, \dots, x_n) \propto P(C_k) \prod_{i=1}^n P(x_i|C_k)$$

$$\hat{y} = P(C_k) \prod_{i=1}^n P(x_i|C_k)$$

Machine learning: Classification

Method: Naives Bayes

On note :

$x = [x^1, x^2, \dots, x^i, \dots, x^p]$ un vecteur de descripteurs

$C = \{c_1, \dots, c_q\}$ l'ensemble des classes possibles

S : ensemble fini de couples de la forme (x, c_k)

n : nombre d'observations de S

n_k : nombre d'éléments de S appartenant à la classe c_k

La règle de classification de Bayes recommande de classer le vecteur x dans la classe c_k pour laquelle $P(c_k/x)$ est maximal.

Ce qui revient à maximiser : $\frac{p(x/c_k)p(c_k)}{p(x)}$

soit encore $p(x/c_k)p(c_k)$, $p(x)$ ne dépendant pas de c_k

Pour pouvoir appliquer la règle de Bayes il faut donc pouvoir estimer : $p(x/c_k)$ et $p(c_k)$

Machine learning: Classification

Method: Naives Bayes

On estime $p(c_k)$ par : $\hat{p}(c_k) = \frac{n_k}{n}$
et $p(x/c_k)$ par : $\hat{p}(x/c_k) = \prod_{i=1}^p p(x^i/c_k)$

Ce qui revient à considérer les attributs comme indépendants les uns des autres.

La règle de classification de Bayes devient alors : classer le vecteur x dans la classe c_k qui maximise :

$$\prod_{i=1}^p p(x^i/c_k) p(c_k)$$

Machine learning: Classification

Method: Naives Bayes

Exemple

$$S = S_1 \cup S_2$$

$$S_1 = \{01100, 11001, 10110, 10101, 10010\}$$

$$S_2 = \{01010, 11111, 11010, 11101, 10101\}$$

$$C = \{c_1, c_2\}$$

On demande de classer $x=00111$

$$\hat{p}(c_1) = \frac{n_1}{n} = \frac{5}{10} = \frac{1}{2}$$

$$\hat{p}(c_2) = \frac{n_2}{n} = \frac{5}{10} = \frac{1}{2}$$

Machine learning: Classification

Method: Naives Bayes

On suppose les attributs indépendants

	c_1	c_2		c_1	c_2
$\hat{p}(x_1 = 0/c_j)$	1/5	1/5	$\hat{p}(x_1 = 1/c_j)$	4/5	4/5
$\hat{p}(x_2 = 0/c_j)$	3/5	1/5	$\hat{p}(x_2 = 1/c_j)$	2/5	4/5
$\hat{p}(x_3 = 0/c_j)$	2/5	2/5	$\hat{p}(x_3 = 1/c_j)$	3/5	3/5
$\hat{p}(x_4 = 0/c_j)$	3/5	2/5	$\hat{p}(x_4 = 1/c_j)$	2/5	3/5
$\hat{p}(x_5 = 0/c_j)$	3/5	2/5	$\hat{p}(x_5 = 1/c_j)$	2/5	3/5

$$S = S_1 \cup S_2$$

$$S_1 = \{01100, 11001, 10110, 10101, 10010\}$$

$$S_2 = \{01010, 11111, 11010, 11101, 10101\}$$

$$C = \{c_1, c_2\}$$

$$\prod_{i=1}^p p(x^i/c_k) p(c_k)$$

$$\hat{p}(00111/c_1) = \frac{1}{2} \times \frac{1}{5} \times \frac{3}{5} \times \frac{3}{5} \times \frac{2}{5} \times \frac{2}{5} = \frac{1}{2} \times \frac{36}{5^5} = 5.75 \cdot 10^{-3}$$

$$\hat{p}(00111/c_2) = \frac{1}{2} \times \frac{1}{5} \times \frac{1}{5} \times \frac{3}{5} \times \frac{3}{5} \times \frac{3}{5} = \frac{1}{2} \times \frac{27}{5^5} = 4.3 \cdot 10^{-3}$$

On classera donc $x=00111$ en classe c_1

	p(Play = no) = 5 / 14 = 0.36
	p(Play = yes) = 9 / 14 = 0.64
p(Outlook = Overcast / Play = No) = 0 / 5 = 0.00	
p(Outlook = Rainy / Play = No) = 2 / 5 = 0.40	
p(Outlook = Sunny / Play = No) = 3 / 5 = 0.60	
p(Outlook = Overcast / Play = Yes) = 4 / 9 = 0.44	
p(Outlook = Rainy / Play = Yes) = 3 / 9 = 0.33	
p(Outlook = Sunny / Play = Yes) = 2 / 9 = 0.22	
p(Temperature = Cool / Play = No) = 1 / 5 = 0.20	
p(Temperature = Hot / Play = No) = 2 / 5 = 0.40	
p(Overcast = Mild / Play = No) = 2 / 5 = 0.40	
p(Temperature = Cool / Play = Yes) = 3 / 9 = 0.33	
p(Temperature = Hot / Play = Yes) = 2 / 9 = 0.22	
p(Temperature = Mild / Play = Yes) = 4 / 9 = 0.44	
p(Humidity = High / Play = No) = 4 / 5 = 0.80	
p(Humidity = Normal / Play = No) = 1 / 5 = 0.20	
p(Humidity = High/ Play = Yes) = 3 / 9 = 0.33	
p(Humidity = Normal / Play = Yes) = 6 / 9 = 0.67	
p(Windy = False / Play = No) = 2 / 5 = 0.40	
p(Windy = True / Play = No) = 3 / 5 = 0.60	
p(Windy = False/ Play = Yes) = 6 / 9 = 0.67	
p(Windy = True / Play = Yes) = 3 / 9 = 0.33	

Supervised Machine Learning: Classification

Method: Decision Tree

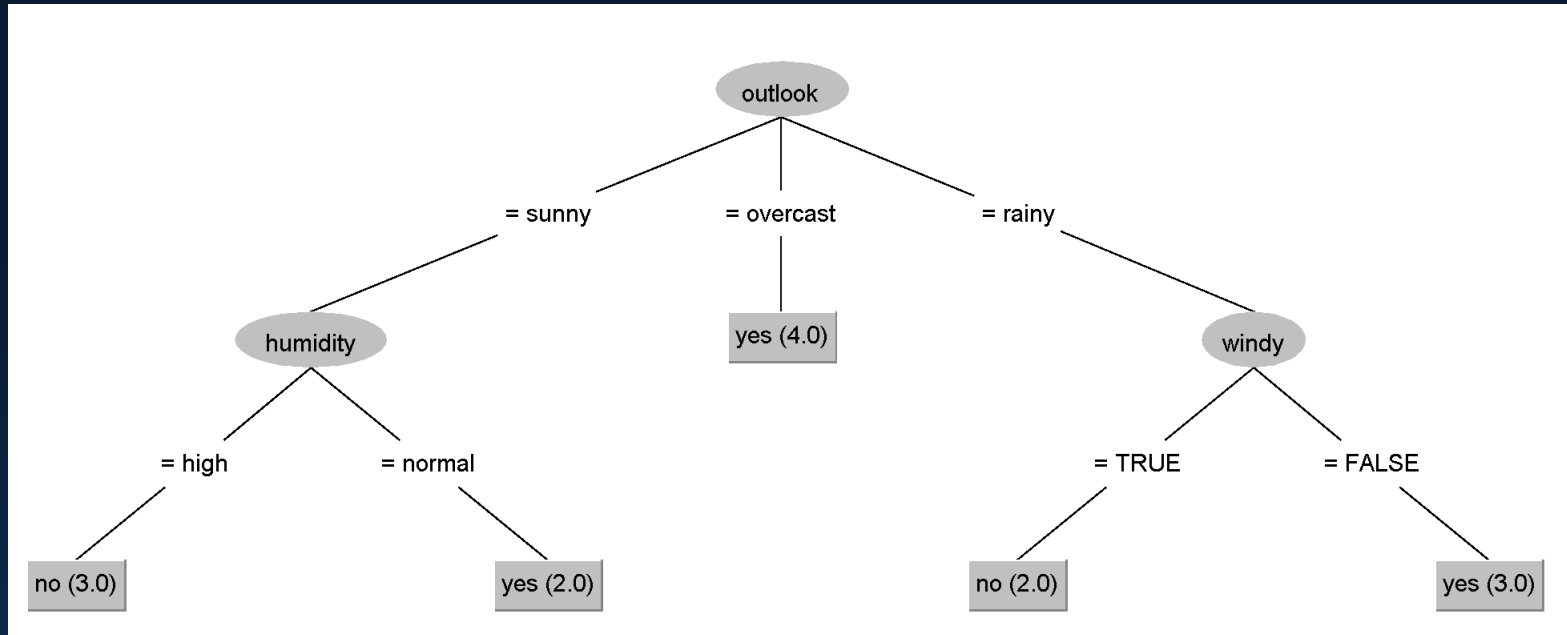
Machine learning: Classification

Method: Decision Tree

Day	Outlook	Temperature	Humidity	Windy	Play
1	Sunny	Hot	High	False	No
2	Sunny	Hot	High	True	No
3	Overcast	Hot	High	False	Yes
4	Rainy	Mild	High	False	Yes
5	Rainy	Cool	Normal	False	Yes
6	Rainy	Cool	Normal	True	No
7	Overcast	Cool	Normal	True	Yes
8	Sunny	Mild	High	False	No
9	Sunny	Cool	Normal	False	Yes
10	Rainy	Mild	Normal	False	Yes
11	Sunny	Mild	Normal	True	Yes
12	Overcast	Mild	High	True	Yes
13	Overcast	Hot	Normal	False	Yes
14	Rainy	Mild	High	True	No

Machine learning: Classification

Method: Decision Tree



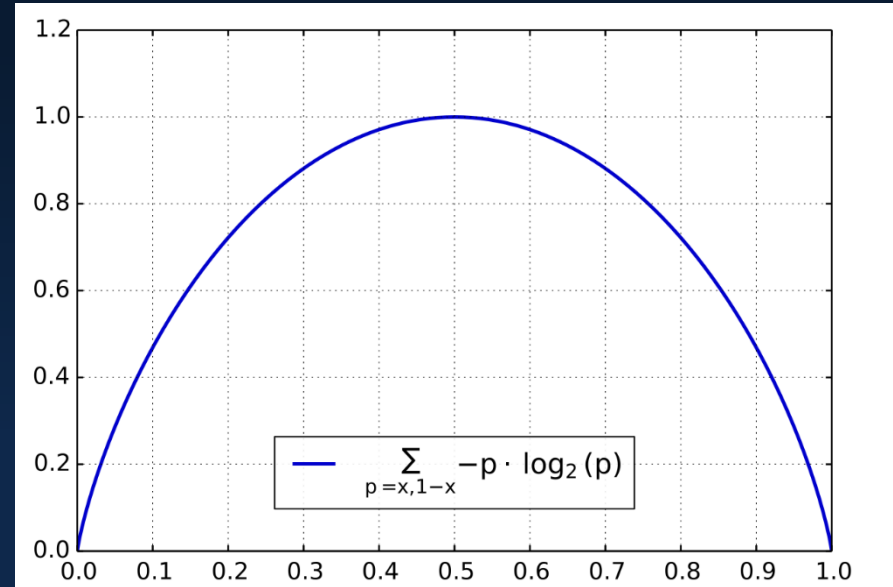
Machine learning: Classification

Method: Decision Tree

Entropy (noun)

1. Name given by Clausius to the state function denoted by S, which characterizes the state of “disorder” of a system.
2. In communication theory, a number that measures the uncertainty of the nature of a given message based on the preceding one. (Entropy is zero when there is no uncertainty.)

$$H(X) = - \sum_{i=1}^n P(x_i) \log_b(P(x_i))$$



Machine learning: Classification

Method: Decision Tree

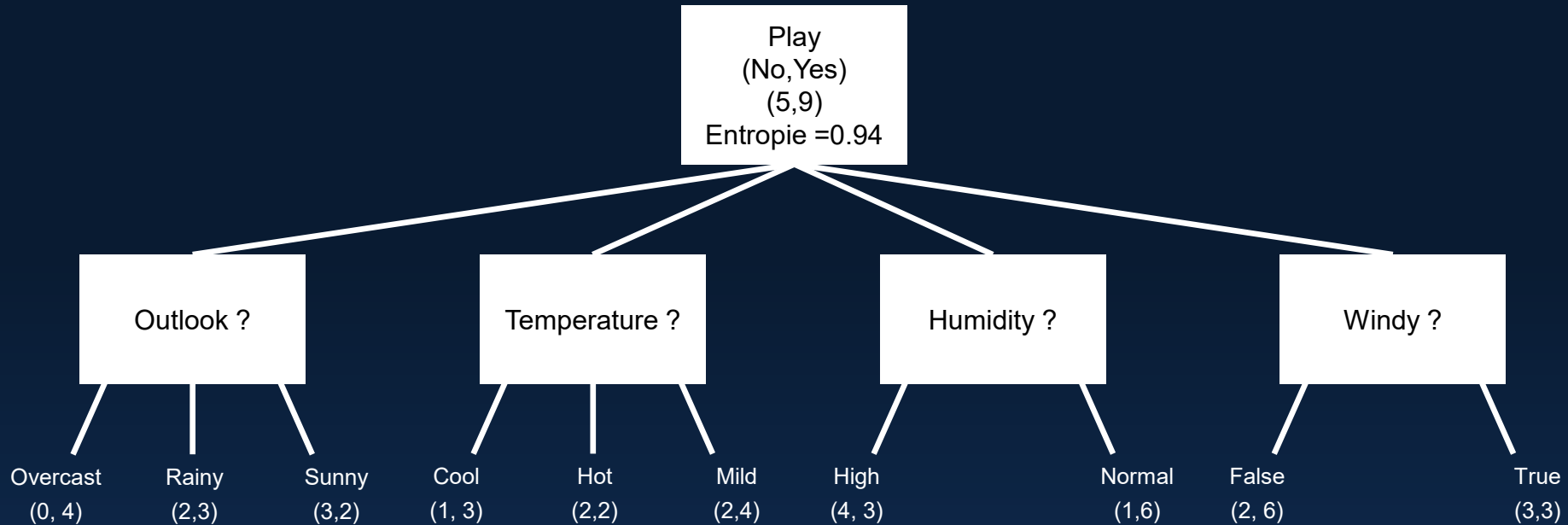
Day	Outlook	Temperature	Humidity	Windy	Play
1	Sunny	Hot	High	False	No
2	Sunny	Hot	High	True	No
6	Rainy	Cool	Normal	True	No
8	Sunny	Mild	High	False	No
14	Rainy	Mild	High	True	No
3	Overcast	Hot	High	False	Yes
4	Rainy	Mild	High	False	Yes
5	Rainy	Cool	Normal	False	Yes
7	Overcast	Cool	Normal	True	Yes
9	Sunny	Cool	Normal	False	Yes
10	Rainy	Mild	Normal	False	Yes
11	Sunny	Mild	Normal	True	Yes
12	Overcast	Mild	High	True	Yes
13	Overcast	Hot	Normal	False	Yes

$$H(\text{Play}) = -P(\text{Play} = \text{No})\log_2(P(\text{Play} = \text{No})) - P(\text{Play} = \text{Yes})\log_2(P(\text{Play} = \text{Yes}))$$

$$H(\text{Play}) = -\frac{5}{14}\log_2\left(\frac{5}{14}\right) - \frac{9}{14}\log_2\left(\frac{9}{14}\right) = 0.94$$

Machine learning: Classification

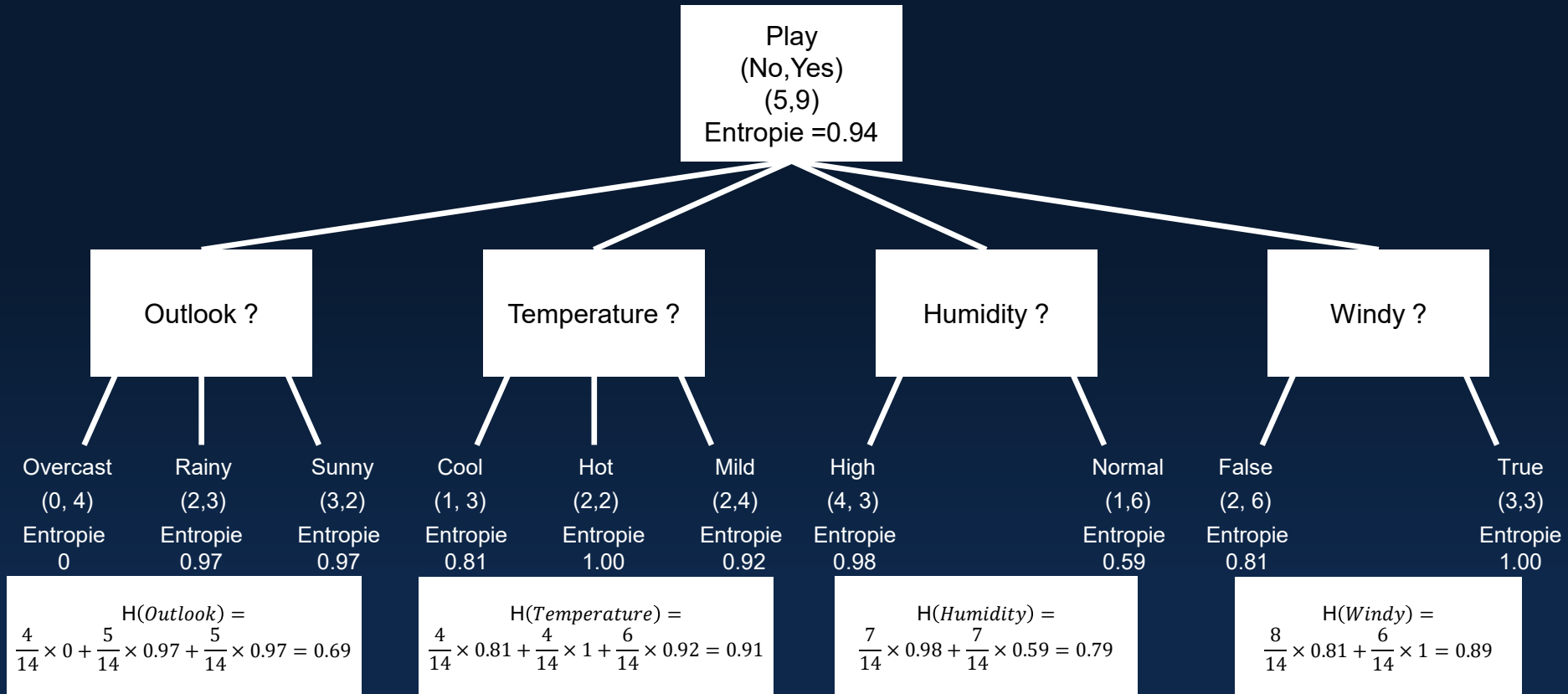
Method: Decision Tree



Which test will remove the most uncertainty?
Which test will reduce entropy the most?

Machine learning: Classification

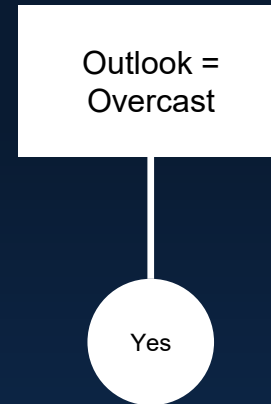
Method: Decision Tree



Machine learning: Classification

Method: Decision Tree

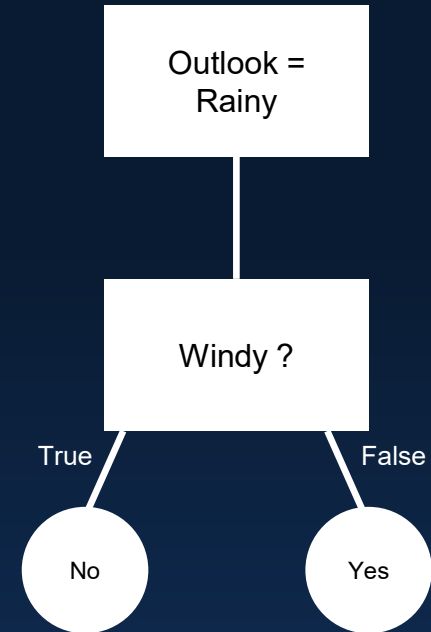
Day	Outlook	Temperature	Humidity	Windy	Play
3	Overcast	Hot	High	False	Yes
7	Overcast	Cool	Normal	True	Yes
12	Overcast	Mild	High	True	Yes
13	Overcast	Hot	Normal	False	Yes



Machine learning: Classification

Method: Decision Tree

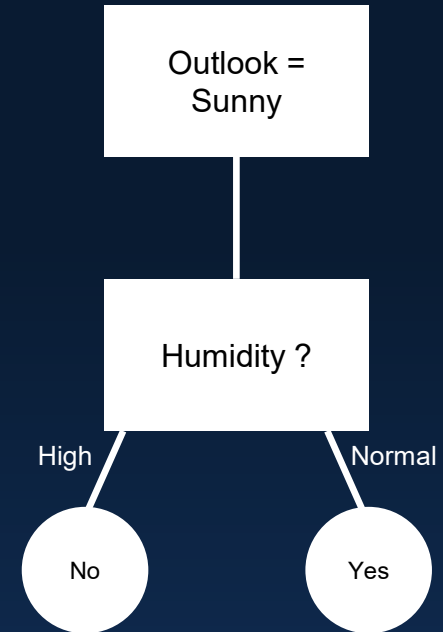
Day	Outlook	Temperature	Humidity	Windy	Play
6	Rainy	Cool	Normal	True	No
14	Rainy	Mild	High	True	No
4	Rainy	Mild	High	False	Yes
5	Rainy	Cool	Normal	False	Yes
10	Rainy	Mild	Normal	False	Yes



Machine learning: Classification

Method: Decision Tree

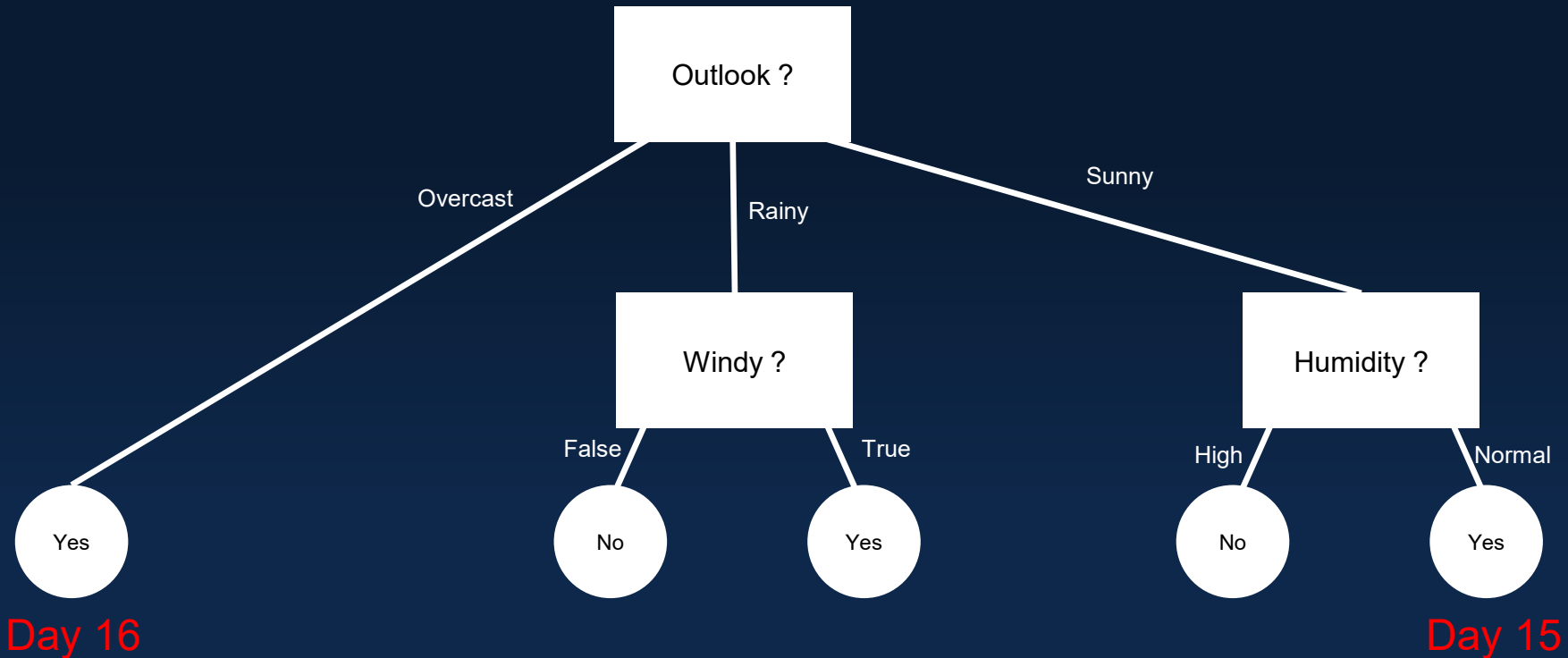
Day	Outlook	Temperature	Humidity	Windy	Play
1	Sunny	Hot	High	False	No
2	Sunny	Hot	High	True	No
8	Sunny	Mild	High	False	No
9	Sunny	Cool	Normal	False	Yes
11	Sunny	Mild	Normal	True	Yes



Machine learning: Classification

Method: Decision Tree

Day	Outlook	Temperature	Humidity	Windy	Play
15	Sunny	Cool	Normal	True	?
16	Overcast	Mild	Normal	False	?



Classification mise en pratique :

 « Orange Data Mining »
étape par étape

Classification mise en pratique : « Orange Data Mining » étape par étape

Étape 1 : Poser le problème

Étape 2 : Rechercher les données

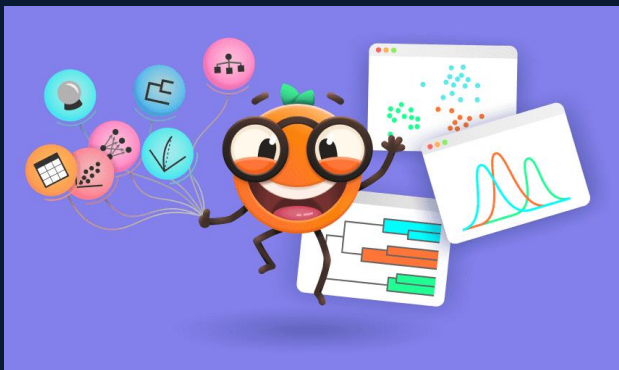
Étape 3 : Nettoyer les données

Étape 4 : Transformer les données

Étape 5 : Sélectionner les données pertinentes

Étape 6 : Rechercher les modèles

Étape 7 : Évaluer et valider les résultats



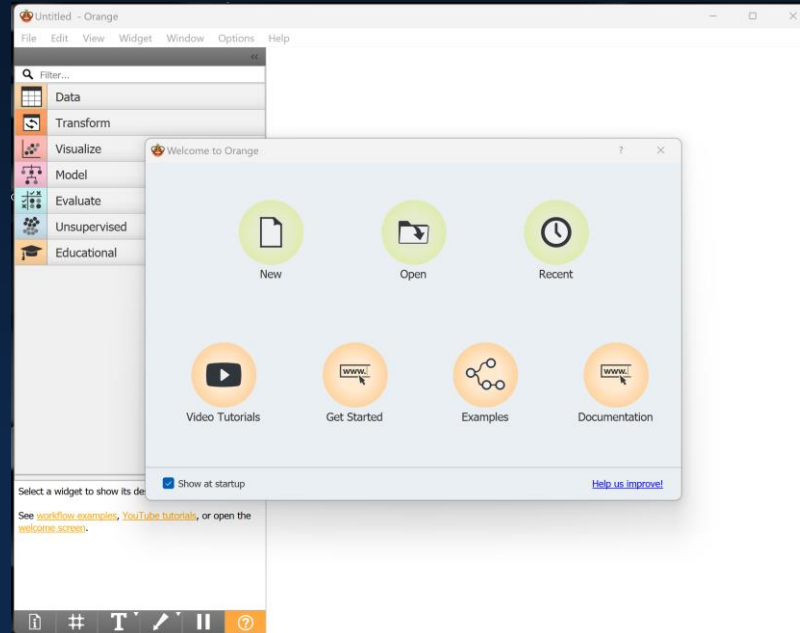
<https://orangedatamining.com/>

Data Mining Fruitful and Fun

Open source machine learning and data
visualization.

Download Orange 3.36.2

<https://orangedatamining.com/download/>



Étape 1: Poser le problème

Un exemple : Classification supervisée de fleurs d'iris

Connaissant :

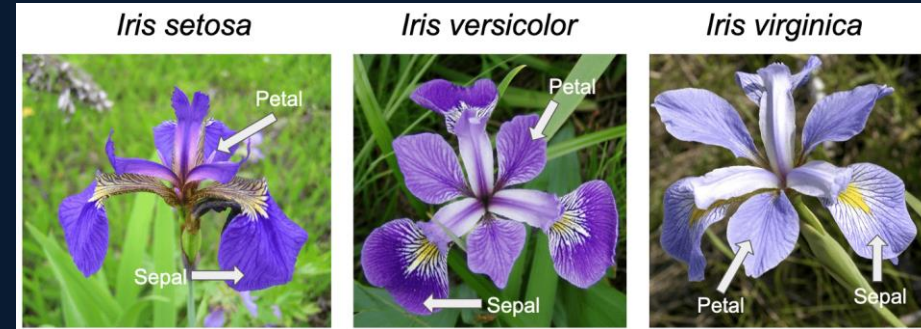
- la longueur des sépales d'une fleur d'iris
- la largeur des sépales d'une fleur d'iris
- la longueur des pétales d'une fleur d'iris
- la largeur des pétales d'une fleur d'iris

Est-il possible de classer les fleurs dans les catégories suivantes ?

- Iris Setosa - classe 1
- Iris Versicolour - classe 2
- Iris Virginica - classe 3

L'objectif est de fournir un modèle de classification supervisée qui accepte en entrée les 4 caractéristiques des sépales et des pétales des fleurs d'iris et qui fournit en sortie le numéro de la classe de la fleur (1, 2 ou 3).

Les résultats de la classification seront évalués sur une base de données de test représentant 25% des données collectées sous la forme d'une matrice de confusion.



Étape 2 : Rechercher les données



Datasets

Contribute Dataset

About Us

Search datasets...

Login



Iris

Donated on 6/30/1988

A small classic dataset from Fisher, 1936. One of the earliest known datasets used for evaluating classification methods.

Dataset Characteristics	Subject Area	Associated Tasks
Tabular	Biology	Classification
Feature Type	# Instances	# Features
Real	150	4

Dataset Information

What do the instances in this dataset represent?
Each instance is a plant

Additional Information

This is one of the earliest datasets used in the literature on classification methods and widely used in statistics and machine learning. The data set contains 3 classes of 50 instances each, where each class refers to a type of iris plant. One class is linearly separable from the other 2; the latter are not linearly separable from each other....

SHOW MORE

Has Missing Values?

No

Introductory Paper

[The Iris data set: In search of the source of virginica](#)

By A. Unwin, K. Kleinman. 2021

Published in Significance, 2021

Variables Table

DOWNLOAD

IMPORT IN PYTHON

CITE

352 citations

791028 views

Keywords

ecology

Creators

R. A. Fisher

DOI

10.24432/C56C76

License

This dataset is licensed under a [Creative Commons Attribution 4.0 International](#) (CC BY 4.0) license.

This allows for the sharing and adaptation of the datasets for any purpose, provided that the appropriate credit is given.

Étape 2 : Rechercher les données

Un exemple : Classification de fleurs d'iris

Title: Iris Plants Database Updated Sept 21 2000 by C.Blake

File : iris.dat

Sources:

- (a) Creator: R.A. Fisher
- (b) Donor: Michael Marshall
- (c) Date: July, 1988

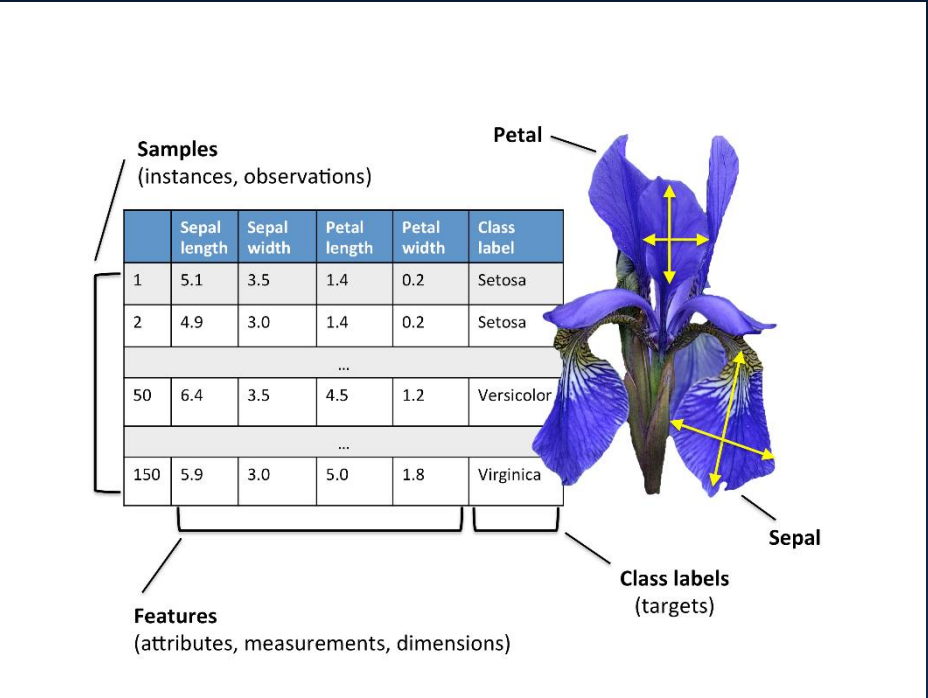
Number of Instances: 150 (50 in each of three classes)

Number of Attributes: 4 numeric attributes and the class

Attribute Information:

- 1. sepal length in cm
- 2. sepal width in cm
- 3. petal length in cm
- 4. petal width in cm
- 5. class:
 - Iris Setosa
 - Iris Versicolour
 - Iris Virginica

Missing Attribute Values: None



Étape 2 : Recherche les données

File Edit View Widget Window Options Help

Filter...

Group by Pivot Table

Preprocess Impute Continuity Discretize

Randomize Purge Domain Split Formula

Create Class Create Instance Python Script

Visualize

Tree Viewer Box Plot Violin Plot Distributions

Scatter Plot Line Plot Bar Plot Flow Diagram

Music Display Interact Linear Projection Radviz

Heat Map Venn Diagram Pythagorean Tree

Pythagorean Forest CH2 Rule Viewer Nomogram

Model

Constant CH2 Rule Induction Calibrated Learner M3N

Tree Random Forest Gradient Boosting SVM

CSV File Import Data Select Columns Data Table Selected Data -- Data Distributions

File Edit View Window Help

Info

150 instances (no missing data)

4 features

Target with 3 values

No meta attributes.

Variables

☒ Show variable labels (if present)

☒ Visualize numeric values

☒ Color by instance classes

Selection

☒ Select full rows

	species	sepal_length	sepal_width	petal_length	petal_width
41	setosa	5	3.5	1.3	0.3
42	setosa	4.5	2.3	1.3	0.3
43	setosa	4.4	3.2	1.3	0.2
44	setosa	5	3.5	1.6	0.6
45	setosa	5.1	3.8	1.9	0.4
46	setosa	4.8	3	1.4	0.3
47	setosa	5.1	3.8	1.6	0.2
48	setosa	4.6	3.2	1.4	0.2
49	setosa	5.3	3.7	1.5	0.2
50	setosa	5	3.3	1.4	0.2
51	versicolor	7	3.2	4.7	1.4
52	versicolor	6.4	3.2	4.5	1.5
53	versicolor	6.9	3.1	4.9	1.5
54	versicolor	5.5	2.3	4	1.3
55	versicolor	6.5	2.8	4.6	1.5
56	versicolor	5.7	2.8	4.5	1.3
57	versicolor	6.3	3.3	4.7	1.6
58	versicolor	4.9	2.4	3.3	1
59	versicolor	6.6	2.9	4.6	1.3
60	versicolor	5.2	2.7	3.9	1.4
61	versicolor	5	2	3.5	1
62	versicolor	5.9	3	4.2	1.5
63	versicolor	6	2.2	4	1
64	versicolor	6.1	2.9	4.7	1.4
65	versicolor	5.6	2.9	3.6	1.3
66	versicolor	6.7	3.1	4.4	1.4
67	versicolor	5.6	3	4.5	1.5
68	versicolor	5.8	2.7	4.1	1
69	versicolor	6.2	2.2	4.5	1.5
70	versicolor	5.6	2.5	3.9	1.1
71	versicolor	5.9	3.2	4.8	1.8
72	versicolor	6.1	2.8	4	1.3
73	versicolor	6.3	2.5	4.9	1.5
74	versicolor	6.1	2.8	4.9	1.2
75	versicolor	6.4	2.9	4.3	1.3
76	versicolor	6.6	3	4.4	1.4
77	versicolor	6.8	2.8	4.8	1.4
78	versicolor	6.7	3	5	1.7
79	versicolor	6	2.9	4.5	1.5
80	versicolor	5.7	2.6	3.5	1
81	versicolor	5.5	2.4	3.8	1.1
82	versicolor	5.5	2.4	3.7	1
83	versicolor	5.8	2.7	3.9	1.2
84	versicolor	6	2.7	5.1	1.6
85	versicolor	5.4	3	4.5	1.5
86	versicolor	6	3.4	4.5	1.6

Distributions - Orange

File Edit View Window Help

Variable

Filter...

species

sepal_length

sepal_width

petal_length

petal_width

☐ Sort categories by frequency

Distribution

Fitted distribution None

Bin width 10

Smoothing 10

Hide bars

Columns

Split by species

☒ Stack columns

☐ Show probabilities

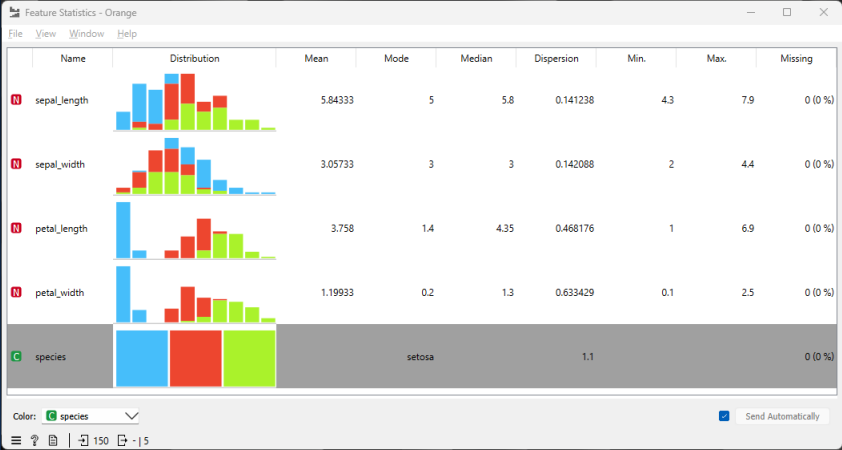
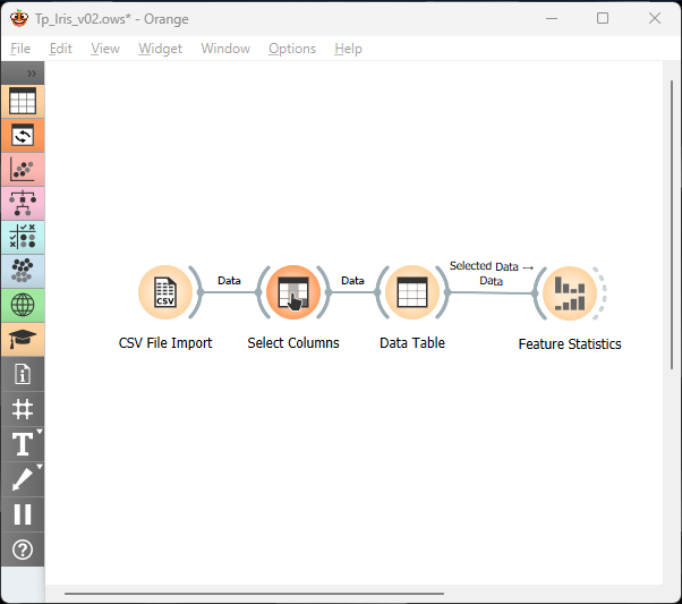
☐ Show cumulative distribution

Apply Automatically

Restore Original Order

☒ Send Automatically

Étape 2 : Recherche les données



Étape 3 : Nettoyer les données

Vérifier l'origine des données

Traiter les valeurs aberrantes

Traiter les valeurs manquantes

Traiter les valeurs nulles

Estimer la qualité des données

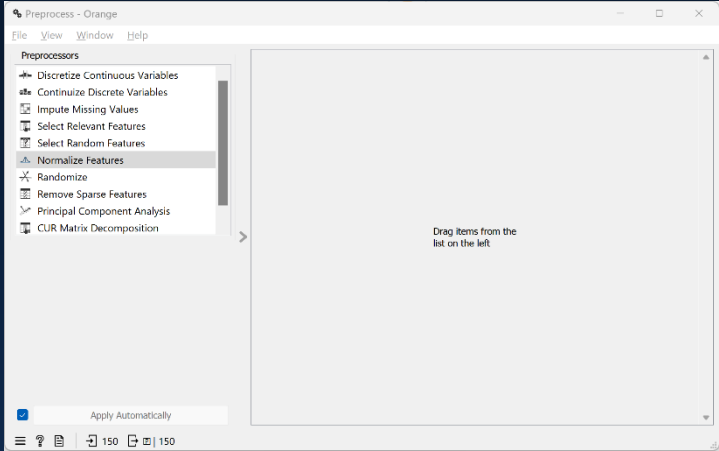
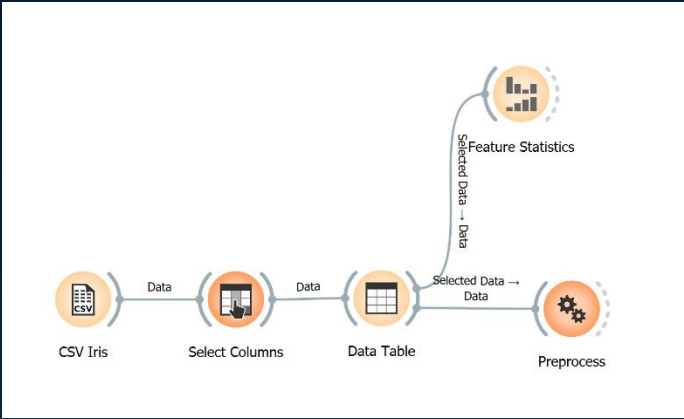
Étape 4 : Transformer les données

Cette étape peut consister à :

- Normaliser les données
- Coder les informations qualitatives
- Coder en ratios (pourcentages)
- Transformer les dates en durées
- Transcoder les données,
 exemple : code postal en coordonnées géographiques
- Exprimer des fréquences
- Exprimer des tendances
- Réaliser des combinaisons de variables

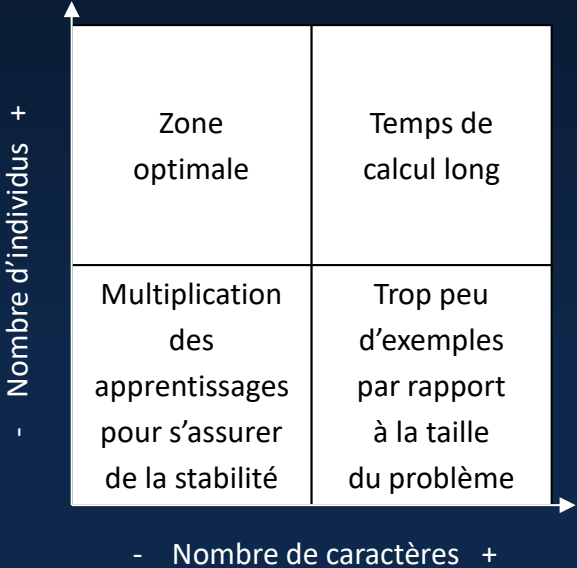
...

Étape 4 : Transformer les données



Étape 5 : Sélectionner les données pertinentes

		Caractères										
		Age	Revenu	Sexe		Situation Matrimoniale			...	Caractère j	...	Caractère p
		X^1	X^2	M	F	Marié	Célibataire	Veuf Divorcé	...	X^j	...	X^p
				X^3	X^4	X^5	X^6	X^7	...			
Individus	X_1	x_1^1	x_1^2	x_1^3	x_1^4	x_1^5	x_1^6	x_1^7	...	x_1^j	...	x_1^p
	X_2	x_2^1	x_2^2	x_2^3	x_2^4	x_2^5	x_2^6	x_2^7	...	x_2^j	...	x_2^p
	...	...	...	...	...	...	...	...	...	...	...	...
	X_i	x_i^1	x_i^2	x_i^3	x_i^4	x_i^5	x_i^6	x_i^7	...	x_i^j	...	x_i^p
	...	...	...	...	...	...	...	...	...	...	...	...
	X_n	x_n^1	x_n^2	x_n^3	x_n^4	x_n^5	x_n^6	x_n^7	...	x_n^j	...	x_n^p



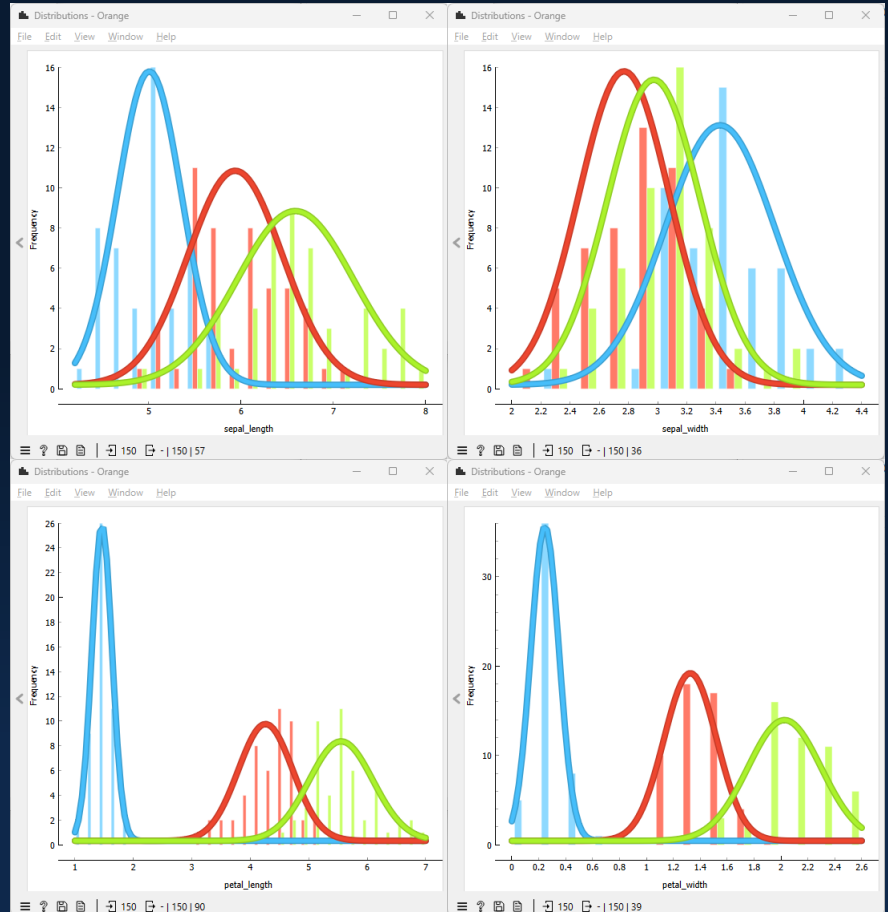
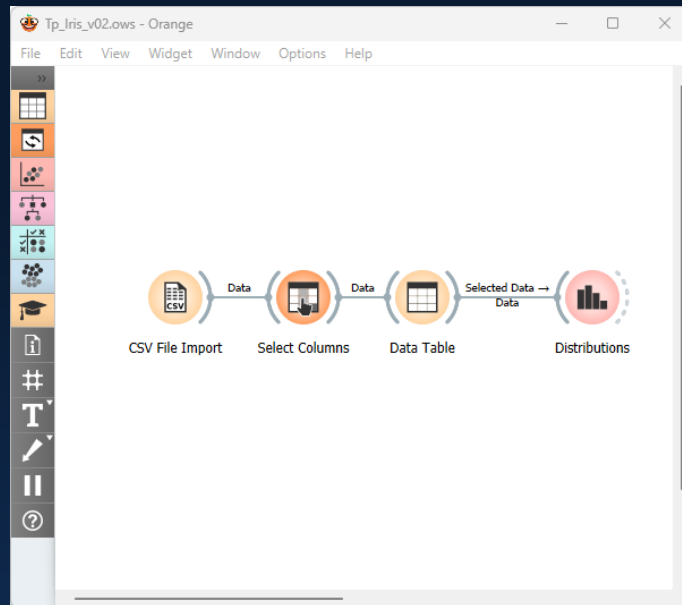
Réduire les dimensions

- Expertise humaine
- Analyses graphiques
- Information Mutuelle
- Analyses de corrélation
- Analyse en composantes principales
- ...

Étape 5 : Sélectionner les données pertinentes

Réduire les dimensions

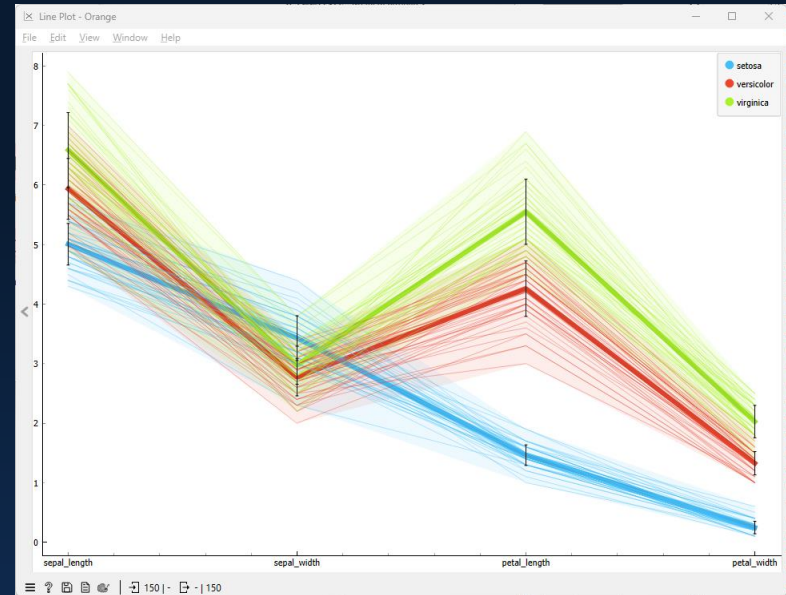
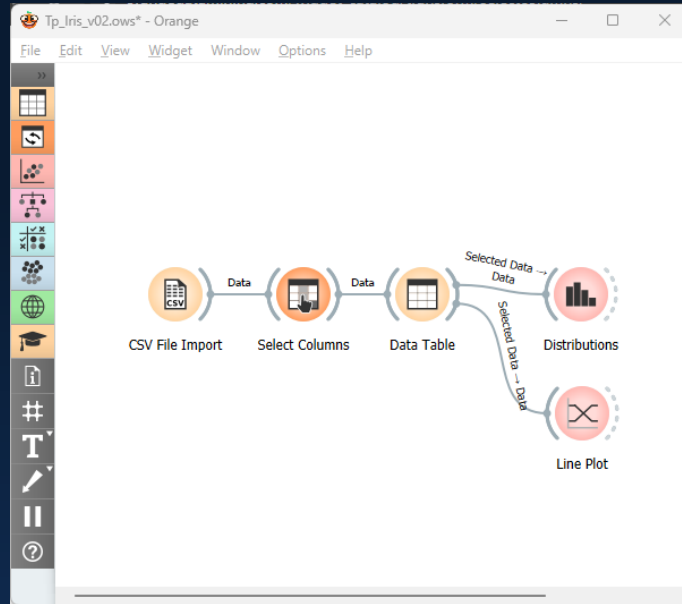
- Expertise humaine



Étape 5 : Sélectionner les données pertinentes

Réduire les dimensions

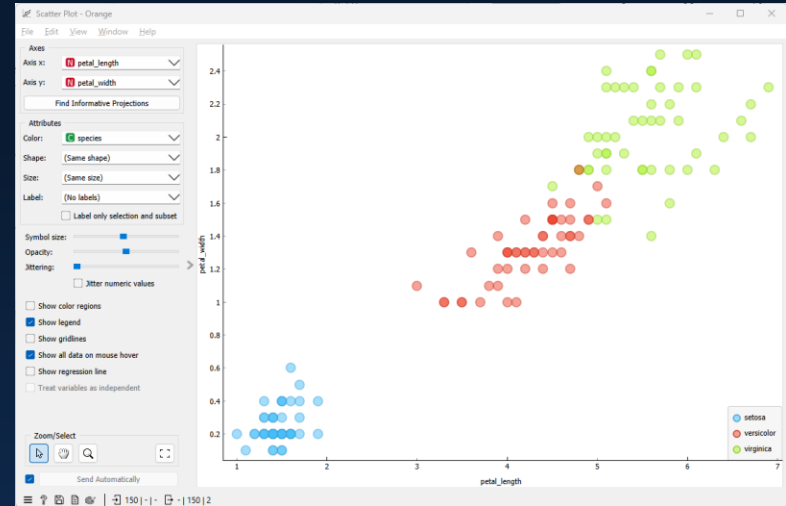
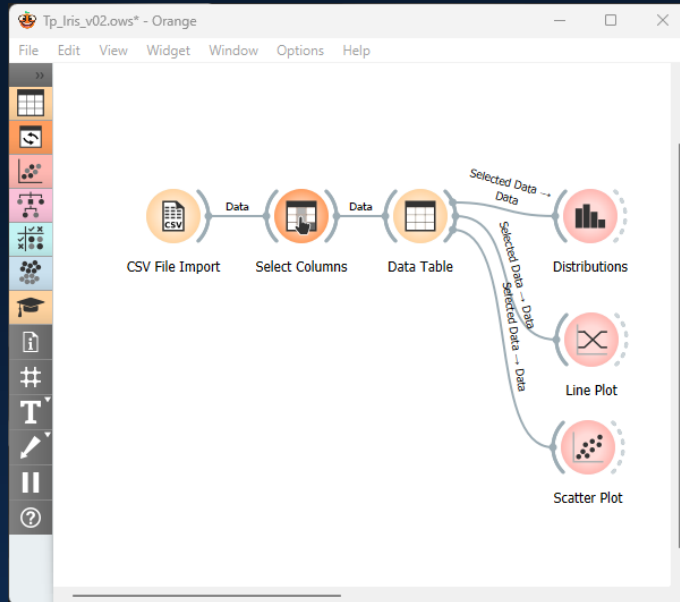
- Expertise humaine



Étape 5 : Sélectionner les données pertinentes

Réduire les dimensions

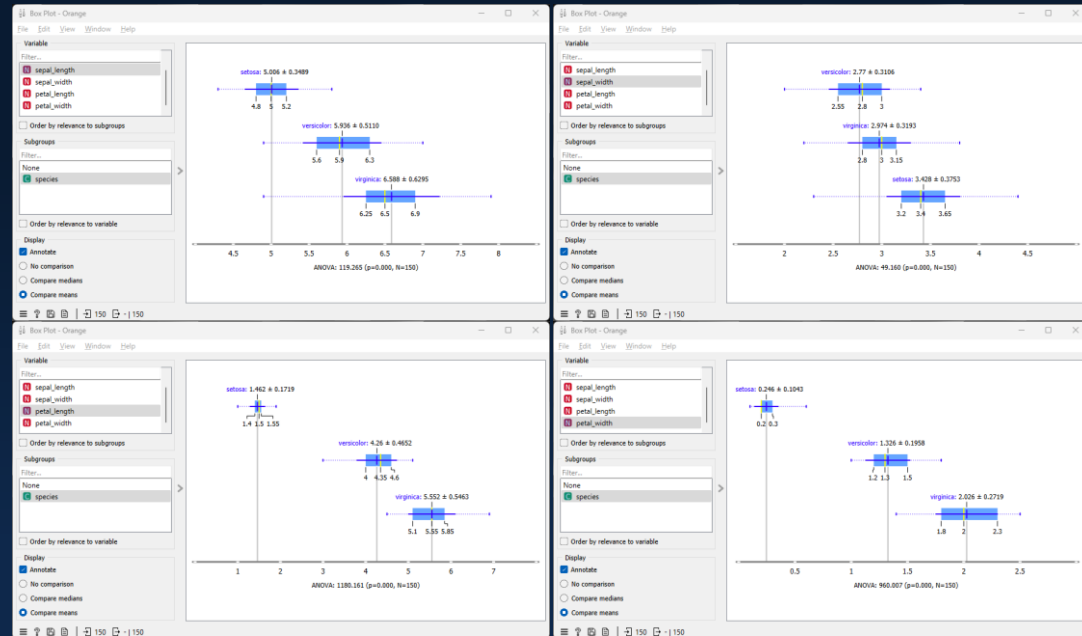
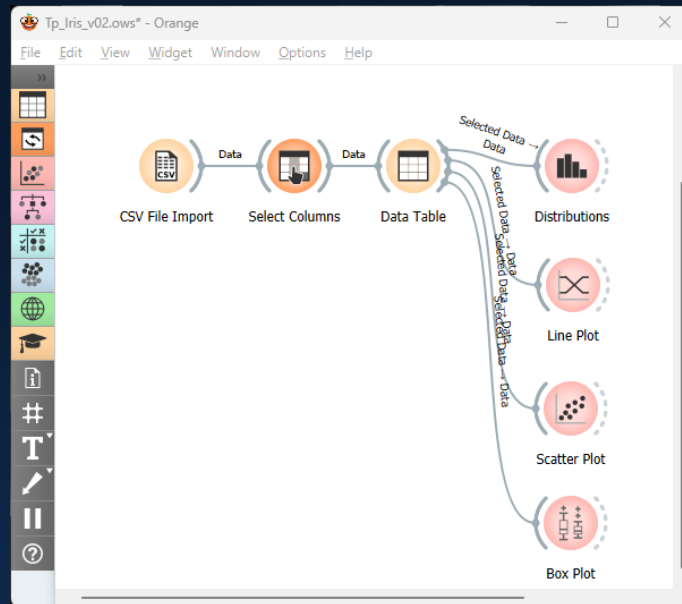
- Expertise humaine



Étape 5 : Sélectionner les données pertinentes

Réduire les dimensions

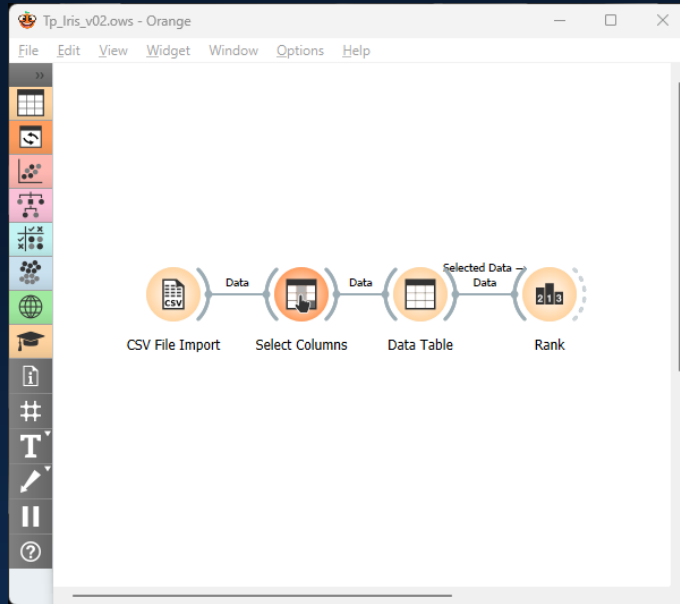
- Expertise humaine



Étape 5 : Sélectionner les données pertinentes

Réduire les dimensions

- Information Mutuelle



The 'Rank' widget interface shows the 'Scoring Methods' section with 'Information Gain' selected. Below this, a table displays the 'Info. gain' for various attributes. The attributes are ranked from highest to lowest information gain.

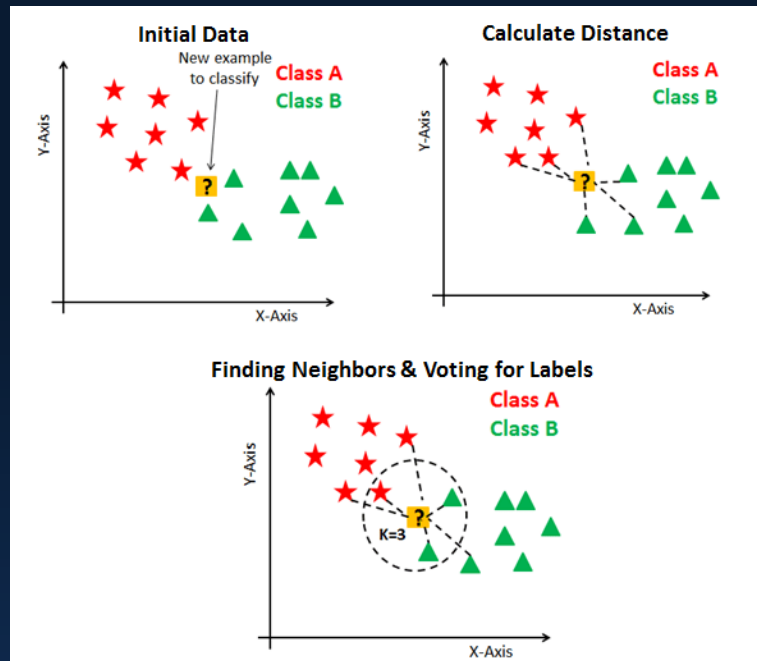
#	Attribute	Info. gain
1	petal_length	1.086
2	petal_width	1.059
3	sepal_length	0.624
4	sepal_width	0.369

Additional options visible in the interface include 'Information Gain Ratio', 'Gini Decrease', 'ANOVA', χ^2 , 'ReliefF', 'FCBF', 'Select Attributes' (None, All, Manual), and 'Best ranked: 5'. The 'Send Automatically' checkbox is also checked.

Étape 6 : Recherche les modèles

Choix d'une méthode ou d'une technique :

- K plus proches voisins
- Naive Bayes
- Arbres de décision
- Réseaux de neurones
- Régression
- Règles
- Analyses factorielles
- Logique Floue
- Algorithmes génétiques
- Le raisonnement à base de cas
- Les réseaux Bayésiens
- ...



Étape 6 : Rechercher les modèles

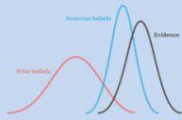
Choix d'une méthode ou d'une technique :

- K plus proches voisins
- Naive Bayes
- Arbres de décision
- Réseaux de neurones
- Régression
- Règles
- Analyses factorielles
- Logique Floue
- Algorithmes génétiques
- Le raisonnement à base de cas
- Les réseaux Bayésiens
- ...



Naive Bayes

Supervised Learning Algorithm for Classification



Bayes Theorem :

$$P(B_i|A) = \frac{P(B_i) \times P(A|B_i)}{P(A)}$$
$$= \frac{P(B_i) \times P(A|B_i)}{\sum_{j=1}^k P(B_j) \times P(A|B_j)}$$


Thomas Bayes 1702-1762

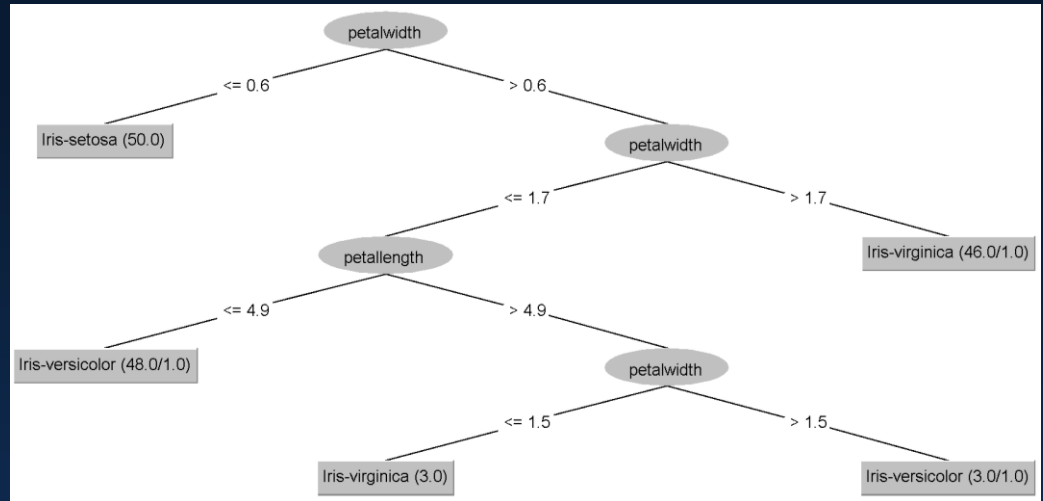
Probabilistic model

$$P(C_k|x) = \frac{P(C_k) \times P(x|C_k)}{P(x)} = \propto p(C_k, x_1, x_2, \dots, x_n) \propto P(C_k) \prod_{i=1}^n P(x_i|C_k)$$
$$\hat{y} = P(C_k) \prod_{i=1}^n P(x_i|C_k)$$

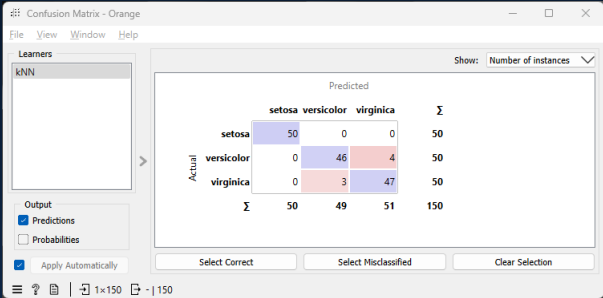
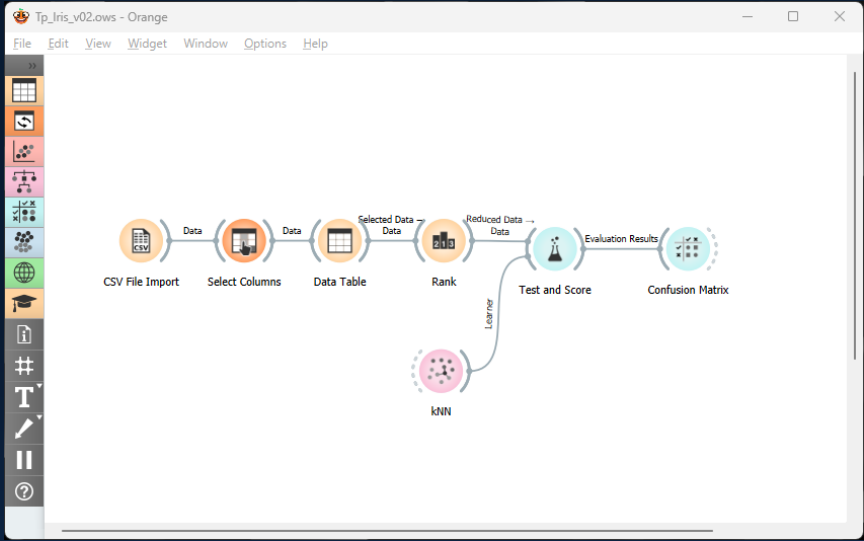
Étape 6 : Recherche les modèles

Choix d'une méthode ou d'une technique :

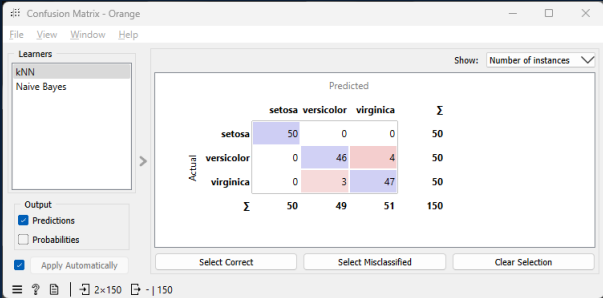
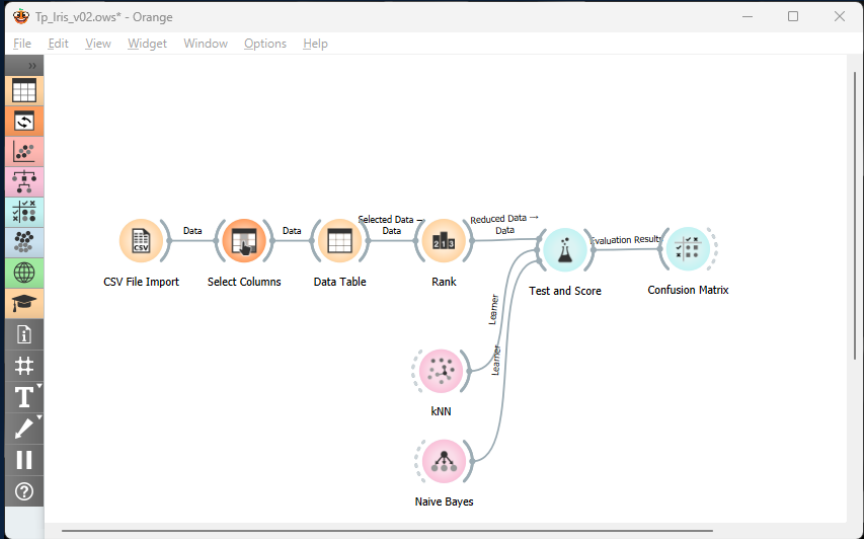
- K plus proches voisins
- Naive Bayes
- Arbres de décision
- Réseaux de neurones
- Régression
- Règles
- Analyses factorielles
- Logique Floue
- Algorithmes génétiques
- Le raisonnement à base de cas
- Les réseaux Bayésiens
- ...



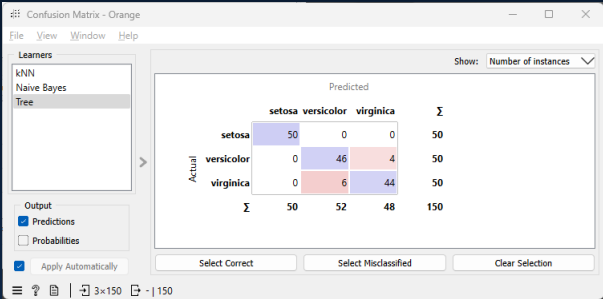
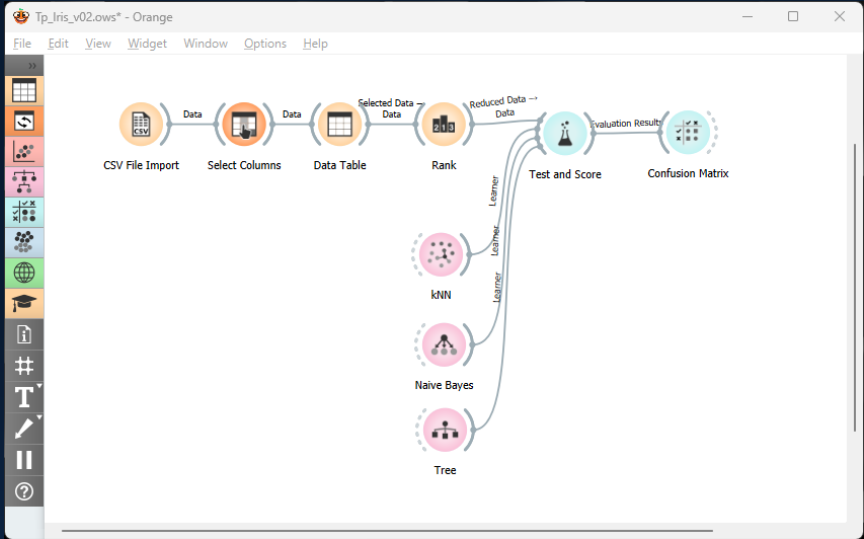
Étape 7 : Évaluer et valider les résultats



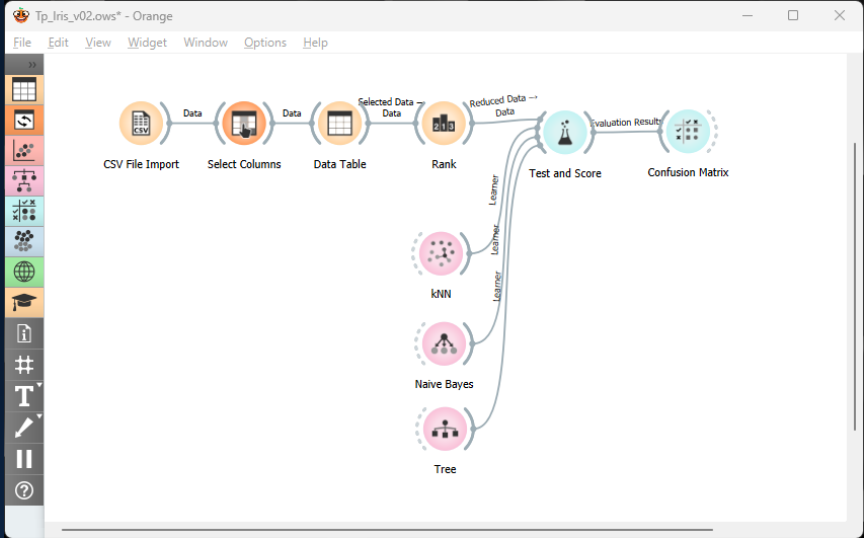
Étape 7 : Évaluer et valider les résultats



Étape 7 : Évaluer et valider les résultats



Étape 7 : Évaluer et valider les résultats



Test and Score - Orange

File Edit View Window Help

☒ Cross validation
Number of folds: 10
☒ Stratified

☐ Cross validation by feature

☐ Random sampling
Repeat train/test: 10
Training set size: 66 %
☒ Stratified

☐ Leave one out
☐ Test on train data
☐ Test on test data

Evaluation results for target (None, show average over classes) ▾

Model	AUC	CA	F1	Prec	Recall	MCC
kNN	0.988	0.953	0.953	0.953	0.953	0.930
Naive Bayes	0.981	0.893	0.893	0.894	0.893	0.840
Tree	0.963	0.933	0.933	0.934	0.933	0.900

Compare models by: Area under ROC curve ▾ ☐ Negligible dff.: 0.1

	kNN	Naive Bayes	Tree
kNN		0.863	0.953
Naive Bayes	0.137		0.902
Tree	0.047	0.098	

Table shows probabilities that the score for the model in the row is higher than that of the model in the column. Small numbers show the probability that the difference is negligible.

150 | 150 | 3x150

Machine Learning: Feature Engineering

Machine learning: Feature Engineering

Feature engineering is a preprocessing step in supervised machine learning and statistical modeling which transforms raw data into a more effective set of inputs. Each input comprises several attributes, known as features. By providing models with relevant information, feature engineering significantly enhances their predictive accuracy and decision-making capability.



		Features / Variables / Attributes / Characteristics / ...										
		Age	Income	Gender		Marital Status			...	Feature j	...	Feature p
		X^1	X^2	M	F	Single	Married	Widowed Divorced	...	X^j	...	X^p
				X^3	X^4	X^5	X^6	X^7				
Observations	X_1	x_1^1	x_1^2	x_1^3	x_1^4	x_1^5	x_1^6	x_1^7	...	x_1^j	...	x_1^p
	X_2	x_2^1	x_2^2	x_2^3	x_2^4	x_2^5	x_2^6	x_2^7	...	x_2^j	...	x_2^p
	...	...	...	...	...	...	...	...	...	...	...	...
	X_i	x_i^1	x_i^2	x_i^3	x_i^4	x_i^5	x_i^6	x_i^7	...	x_i^j	...	x_i^p
	...	...	...	...	...	...	...	...	...	...	...	...
	X_n	x_n^1	x_n^2	x_n^3	x_n^4	x_n^5	x_n^6	x_n^7	...	x_n^j	...	x_n^p

Feature engineering techniques:

- Feature imputation
- Feature encoding
- Feature scaling
- Feature transformation
- Feature selection
- ...

Machine learning: Feature Engineering

Feature Imputation

Missing data can skew your analysis, reduce model accuracy, and lead to biased results. Imputation fills these gaps, allowing you to use the complete dataset for analysis and modeling.

Imputation Techniques				
Technique	Description	Pros	Cons	Use Case
Mean/Median/Mode Imputation	Replace missing values with mean, median, or mode	Simple to implement	May distort original data distribution	Useful when data is normally distributed
KNN Imputation	Use nearest neighbors to fill missing values	More accurate imputation	Computationally expensive	When dataset has relationships between features
Forward/Backward Fill	Use next/previous value to fill missing values	Maintains trend in data	Can propagate errors	Time series data
Interpolation	Estimate missing values based on other values	Handles continuous data well	May not be accurate for all types of data	Continuous data with trends

Machine learning: Feature Engineering

Feature Encoding

Machine learning models require numerical input. Encoding transforms categorical data into a numerical format, enabling models to process and learn from it.

Encoding Techniques

Technique	Description	Pros	Cons	Use Case
One-Hot Encoding	Converts categories to binary columns	Preserves all information	Can create high-dimensional data	Categorical data with no ordinal relationship
Label Encoding	Assigns unique number to each category	Simple and efficient	Implies ordinal relationship	Ordinal categorical data
Ordinal Encoding	Assigns numbers following order	Preserves ordinal relationship	Assumes order is meaningful	Ordered categorical data
Frequency Encoding	Uses frequency of categories	Handles high cardinality well	Loses category information	High cardinality categorical data

Machine learning: Feature Engineering

Feature Scaling

Feature scaling ensures that all features contribute equally to the model. It can significantly improve the performance of algorithms that rely on distance measurements or gradient-based optimization.

Scaling Techniques				
Technique	Description	Pros	Cons	Use Case
Standardization	Transforms data to have mean of 0 and std dev of 1	Handles outliers well	Assumes normal distribution	Data with outliers
Min-Max Scaling	Scales data to a fixed range (0 to 1)	Preserves relationships	Sensitive to outliers	Data without outliers
Robust Scaling	Uses median and IQR	Not sensitive to outliers	May not scale all features equally	Data with many outliers

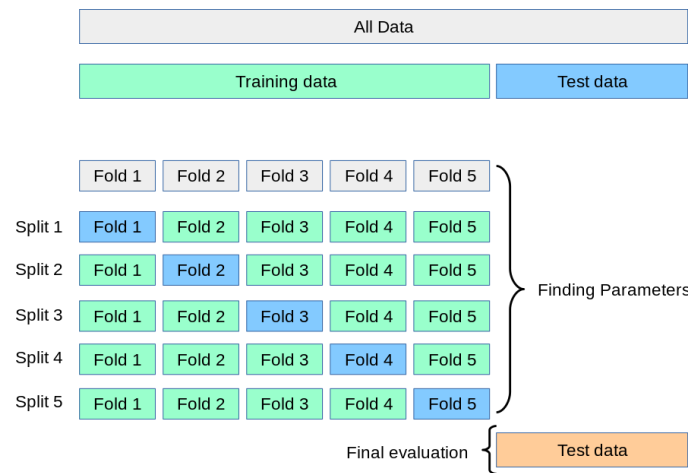
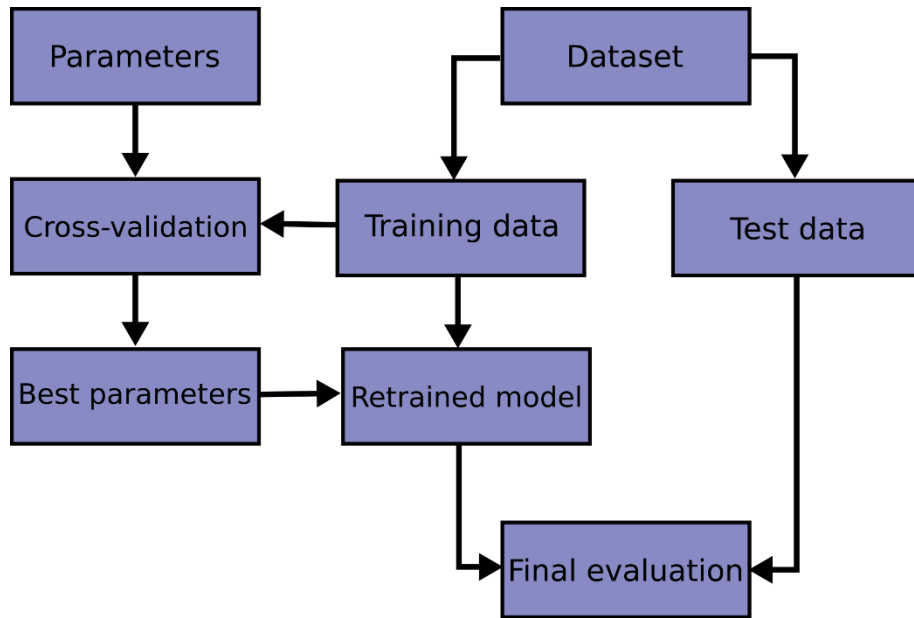
Machine learning: Feature Engineering

Feature Transformation

Feature transformation modifies the data into a more suitable format for modeling, helping to improve model performance and interpretability.

Feature Transformation Techniques

Technique	Description	Pros	Cons	Use Case
Log Transformation	Reduces skewness in data	Handles skewed data well	Not applicable for negative values	Skewed data
Box-Cox Transformation	Transforms data to be more normal	Flexible and powerful	Requires positive data	Data requiring normalization
Polynomial Features	Adds polynomial terms	Captures non-linear relationships	Can create very large feature space	Non-linear relationships in data



Machine learning: Feature Engineering

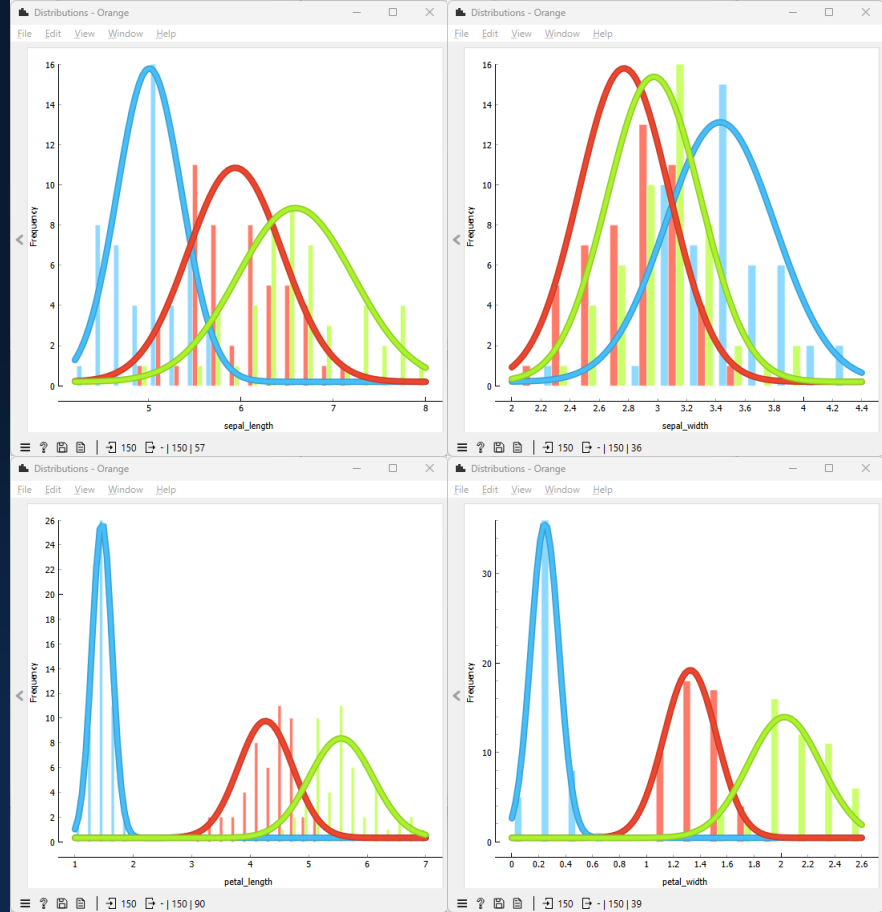
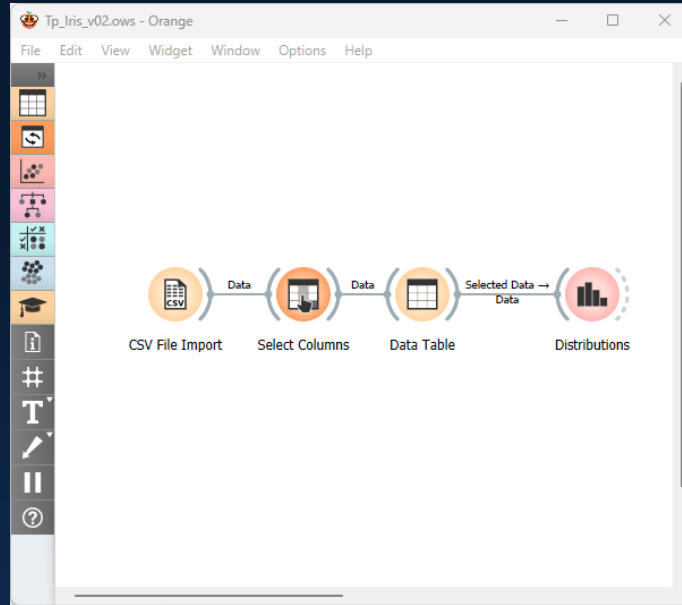
Feature selection

Feature selection

- Human expertise
- Graphical analyses
- Mutual information
- Correlation analyses
- Principal component analysis ...

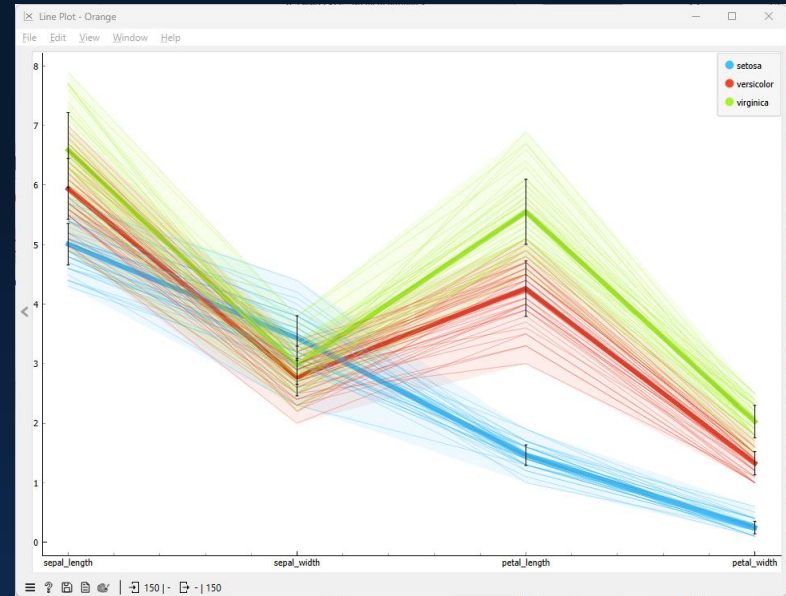
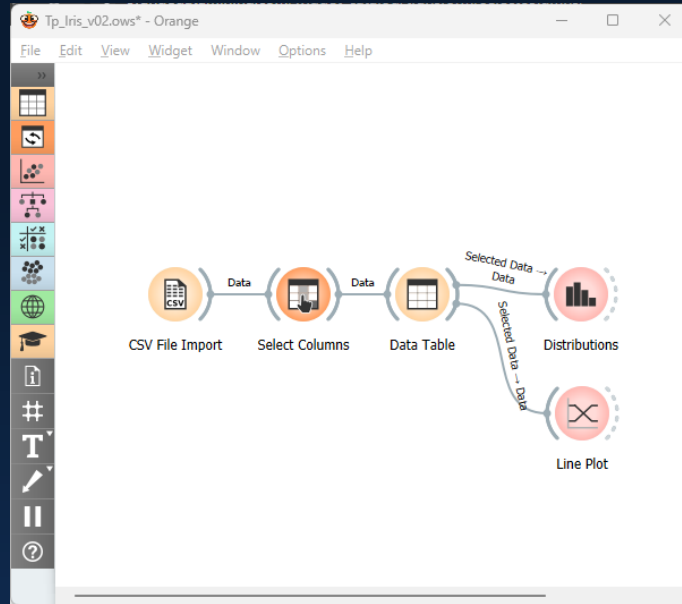
Machine learning: Feature Engineering

Feature selection



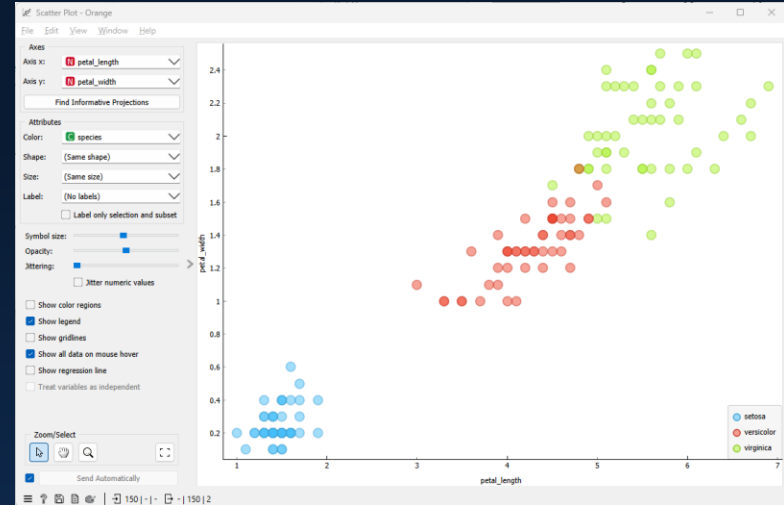
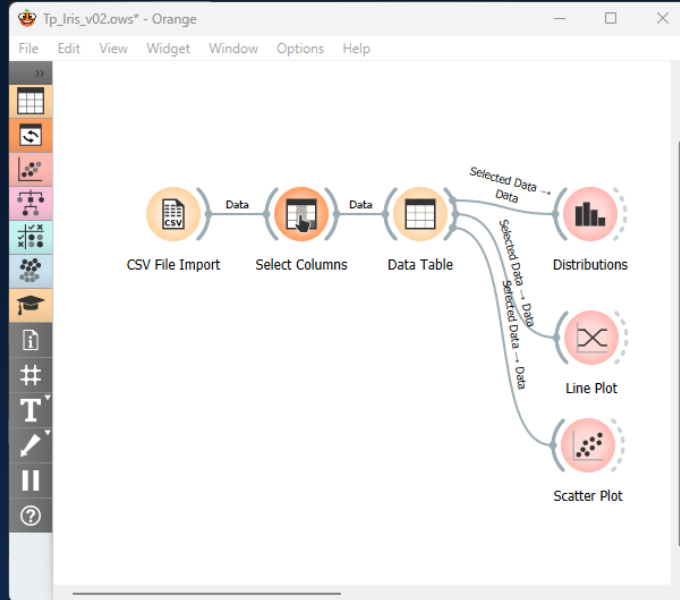
Machine learning: Feature Engineering

Feature selection



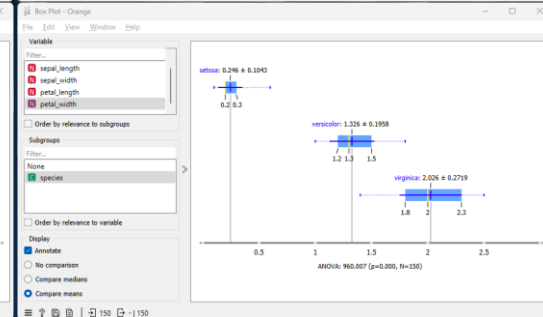
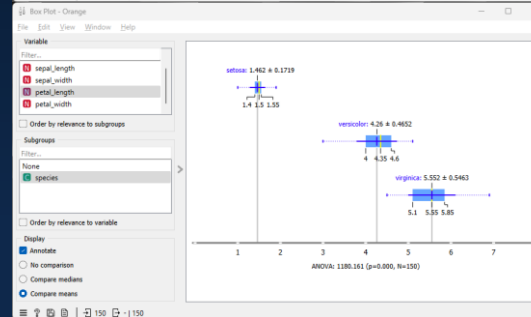
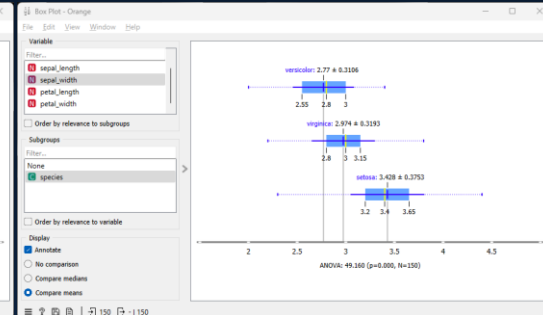
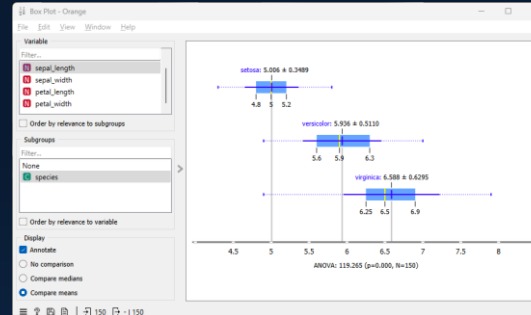
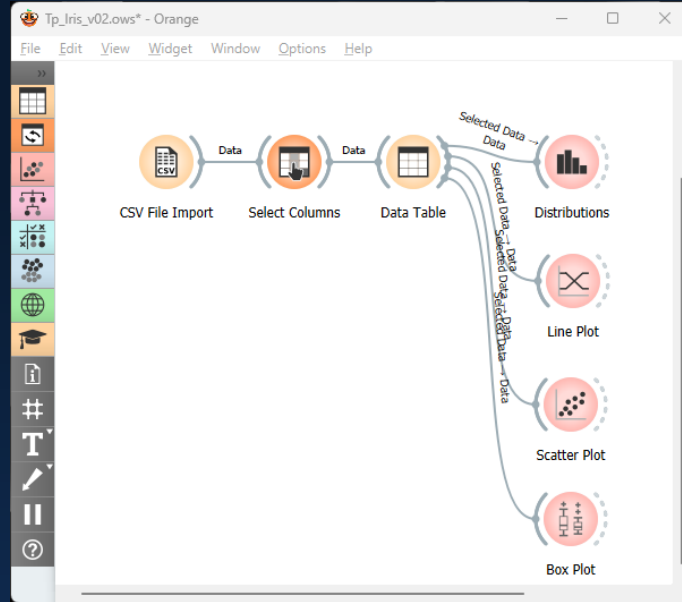
Machine learning: Feature Engineering

Feature selection



Machine learning: Feature Engineering

Feature selection



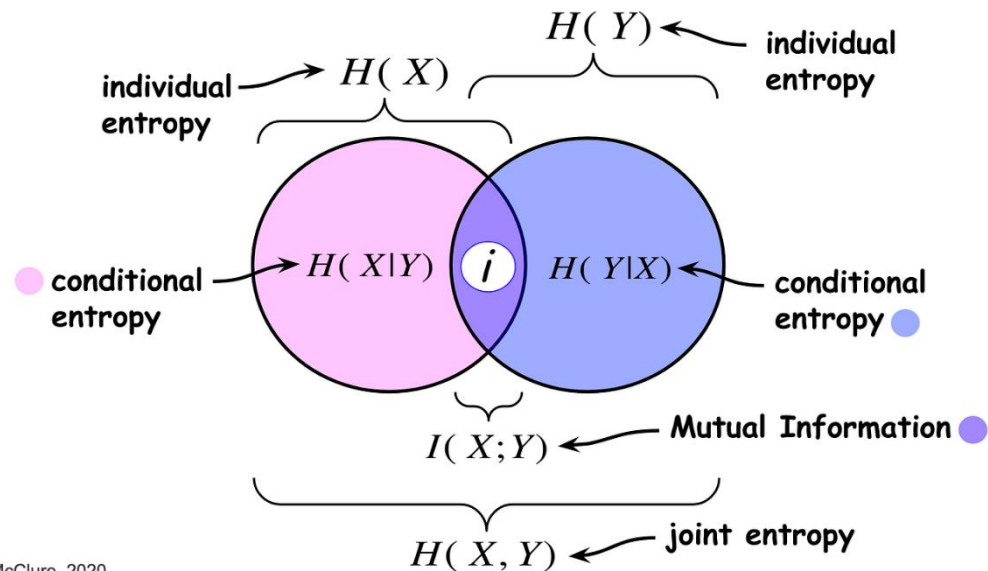
Machine learning: Feature Engineering

Feature selection

Y is a feature

X is the class

$$\begin{aligned} I(X; Y) &\equiv H(X) - H(X | Y) \\ &\equiv H(Y) - H(Y | X) \end{aligned}$$



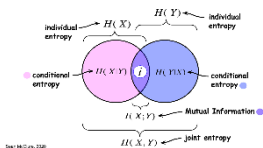
Sean McClure, 2020

$$H[x] = E[h(x)] = - \sum_x p(x) h(x) = - \sum_x p(x) \log_2 p(x)$$

$$H[x|y] = - \sum_x \sum_y p(x, y) \log_2 p(x|y) = - \sum_x \sum_y p(x, y) \log_2 \frac{p(x, y)}{p(y)}$$

Machine learning: Feature Engineering

Feature selection



$$H[x] = E[h(x)] = - \sum_x p(x) h(x) = - \sum_x p(x) \log_2 p(x)$$

$$H[x|y] = - \sum_x \sum_y p(x,y) \log_2 p(x|y) = - \sum_x \sum_y p(x,y) \log_2 \frac{p(x,y)}{p(y)}$$

$$I(X; Y) \equiv H(X) - H(X | Y)$$

[illegible]

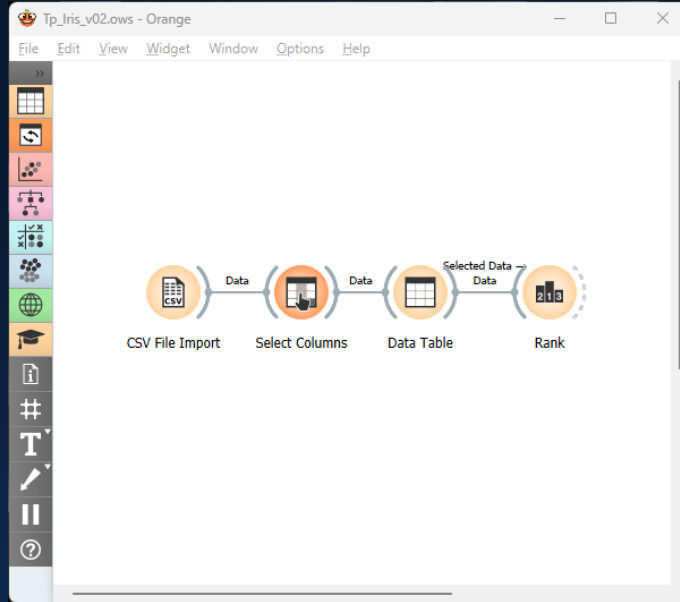
		X = Play				X = Play															
		No	Yes			No	Yes	Marginal p(x)			No	Yes									
Y = Temperature	Cool	1	3	Cool	0.07	0.21	0.29	Cool	0.14	0.09											
	Hot	2	2	Hot	0.14	0.14	0.29	Hot	0.14	0.14											
	Mild	2	4	Mild	0.14	0.29	0.43	Mild	0.23	0.17	H(X=Play)= 0.94		H(Play Temperature)= 0.91		I(Play,Temperature)= 0.03						
				Marginal p(y)	0.36	0.64	1.00														

		X = Play				X = Play				X = Play										2
		No	Yes			No	Yes	Marginal p(x)				No	Yes							
Y = Humidity	High	4	3	High	0.29	0.21	0.50	High	0.23	0.26										
	Normal	1	6	Normal	0.07	0.43	0.50	Normal	0.20	0.10			H(X=Play) = 0.94		H(Play Humidity) = 0.79		I(Play,Humidity) = 0.15			
				Marginal p(y)		0.36	0.64	1.00												

		X = Play				X = Play		Marginal p(x)				X = Play								3
		No	Yes			No	Yes	0.00				No	Yes							
Y = Windy	False		2	6	False		0.14	0.43	0.57	False		0.29	0.18							
	True		3	3	True		0.21	0.21	0.43	True		0.21	0.21	H(X=Play)= 0.94		H(Play Windy)= 0.89		I(Play,Windy)= 0.05		
				Marginal p(y)			0.36	0.64	1.00											

Machine learning: Feature Engineering

Feature selection



Rank - Orange

File View Window Help

Scoring Methods

- ☒ Information Gain
- ☐ Information Gain Ratio
- ☐ Gini Decrease
- ☐ ANOVA
- ☐ χ^2
- ☐ ReliefF
- ☐ FCBF

Select Attributes

- ☐ None
- ☒ All
- ☐ Manual
- ☐ Best ranked: 5

☒ Send Automatically

	#	Info. gain
1	petal_length	1.086
2	petal_width	1.059
3	sepal_length	0.624
4	sepal_width	0.369