



Intelligence Artificielle

Bezzaz Soumia

2025/2026 – Semestre 2

Chapitre 2 : Classification avec apprentissage supervisé

L'apprentissage artificiel (automatique)

- L'apprentissage automatique (machine learning en anglais), un des champs d'étude de l'IA.
- L'objectif de l'apprentissage artificiel est de concevoir des programmes pouvant s'améliorer automatiquement avec l'expérience.

Les Types d'Apprentissage en IA

Il existe essentiellement trois types d'apprentissage artificiel :

L'apprentissage supervisé

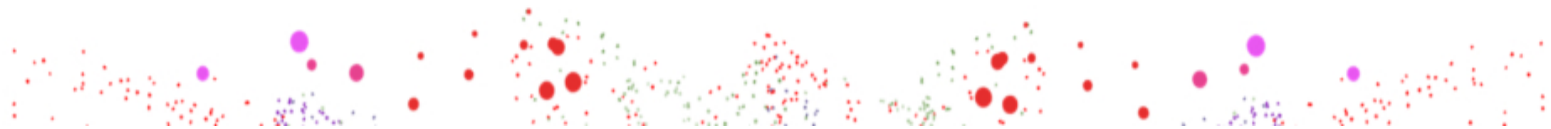
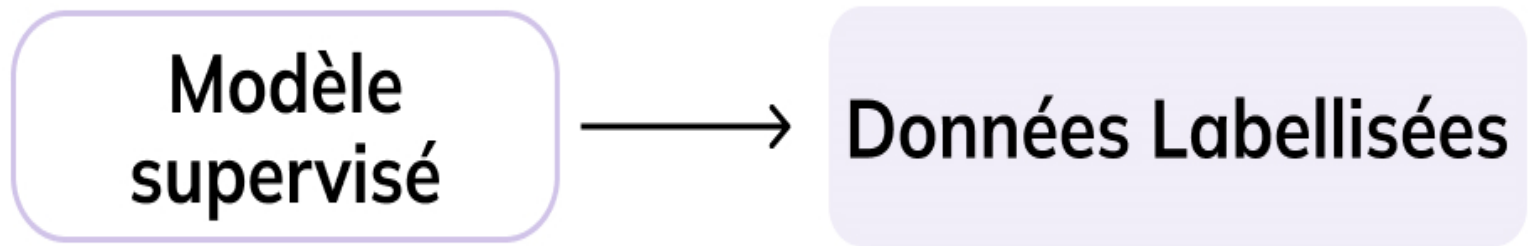
L'apprentissage non supervisé

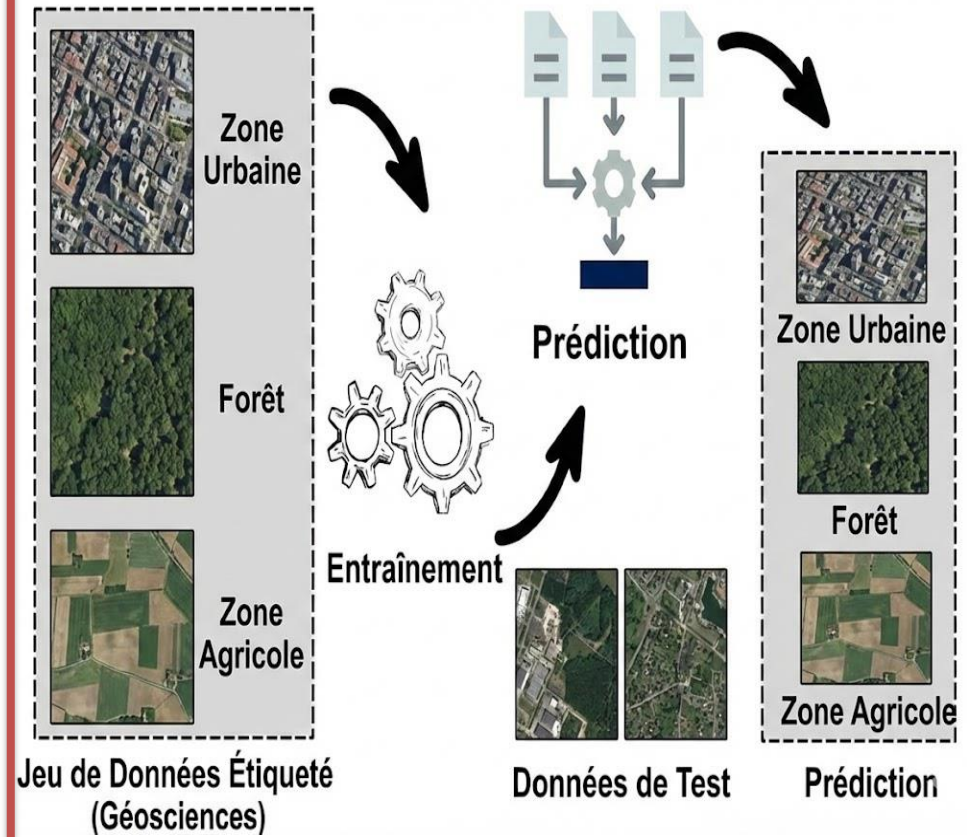
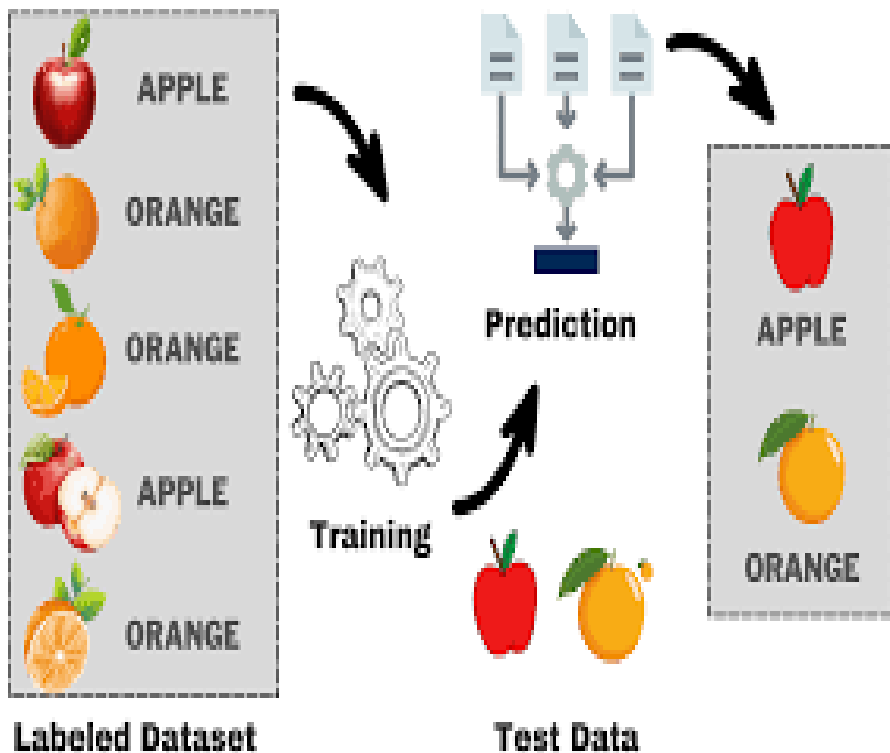
L'apprentissage par renforcement

L'apprentissage supervisé

- La machine apprend à partir d'exemples.
- Un expert est employé pour étiqueter correctement des exemples.
- Apprendre une fonction qui associe une entrée à une sortie à partir d'exemples étiquetés.

L'apprentissage supervisé





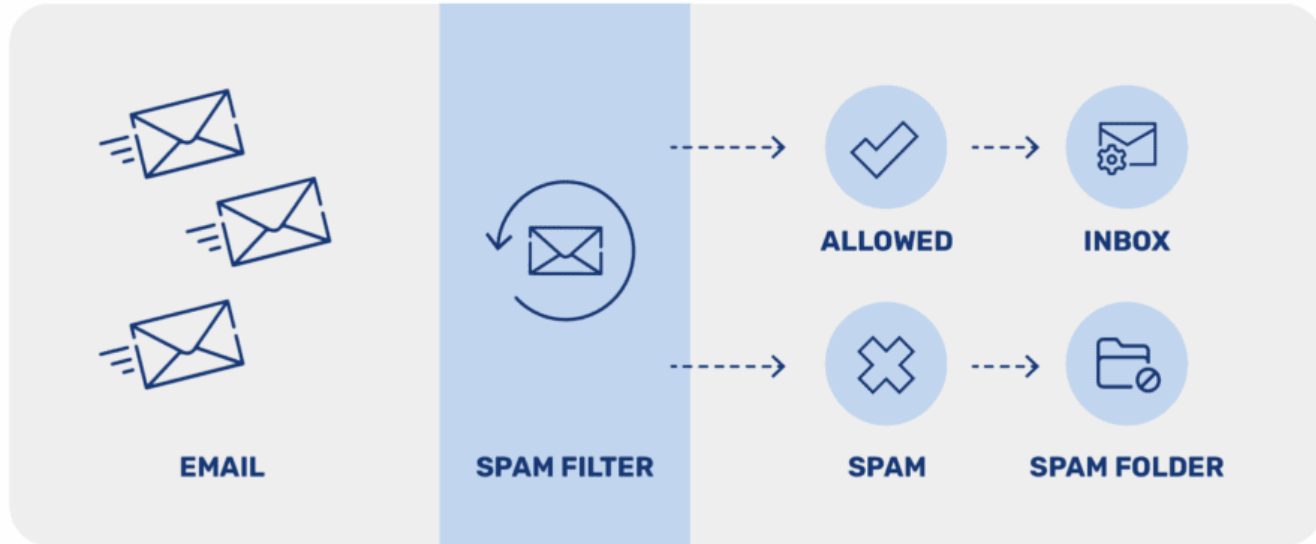
ML
Pipeline

STANDARD MACHINE LEARNING PIPELINE



L'apprentissage supervisé

Par exemple, dans un système de filtrage de spam pour les emails, l'entrée serait un ensemble d'emails, et la sortie serait chaque email marqué comme "spam" ou "non spam".



L'apprentissage supervisé

Les filtres anti-spam analysent les emails entrants en utilisant des méthodes, qui évalue la probabilité qu'un email soit du spam selon ses caractéristiques (mots-clés, expéditeur, structure).

Ils s'appuient sur des règles prédéfinies, des listes noires et l'apprentissage automatique pour classer les messages.

Gmail, par exemple, ne permet pas de désactiver cette analyse mais permet de personnaliser les filtres via les paramètres Spam.

1. Analyse par mots-clés

Le filtre recherche des mots ou expressions suspects fréquemment utilisés dans les spams.

Exemples de mots-clés suspects	Contexte
"Gratuit", "Free"	Offres trop belles pour être vraies
"Urgent", "Agissez maintenant"	Pression psychologique
"Gagnant", "Félicitations"	Fausses loteries
"Viagra", "Casino"	Produits douteux
"Cliquez ici"	Incitation au clic

Principe : Plus un mail contient de mots-clés suspects, plus son **score de spam** augmente, et il finit dans les indésirables.

2. Analyse de la structure du message

Le filtre examine comment le mail est construit :

Format HTML suspect : code mal formé, caché ou excessif

Ratio texte/images : un mail composé uniquement d'images (pour contourner l'analyse de mots-clés) est suspect

Pièces jointes douteuses : fichiers .exe, .zip, .scr

Liens hypertextes : URL masquées ou raccourcies pointant vers des sites malveillants

Mise en forme agressive : MAJUSCULES PARTOUT, couleurs criardes, beaucoup de points d'exclamation !!!

En-têtes techniques (headers) : incohérences dans les métadonnées du mail (chemin de routage, horodatage...)

Principe : *Un mail légitime a généralement une structure propre et professionnelle.*

3. Analyse de l'expéditeur

Le filtre vérifie qui envoie le message :

Vérification	Détail
Adresse e-mail	banque@xyz123.ru ≠ adresse officielle d'une banque
Domaine	Le domaine est-il connu, fiable, récent ?
Réputation	L'IP de l'expéditeur est-elle sur une liste noire (blacklist) ?
Authentification SPF/DKIM/DMARC	Protocoles qui vérifient que l'expéditeur est bien autorisé à envoyer au nom de ce domaine
Usurpation (spoofing)	Le nom affiché dit "Amazon" mais l'adresse réelle est hacker@faux-site.com

Principe : Si l'expéditeur n'est pas authentifié ou est déjà signalé, le mail est bloqué.

L'apprentissage supervisé

Chaque critère attribue un score, et si le total dépasse un certain seuil, le mail est classé comme spam.

Un seul critère ne suffit pas toujours, mais la combinaison des trois rend le filtrage très efficace.

L'apprentissage supervisé

Il existe de nombreuses méthodes de classification supervisée :

- k plus proches voisins
- analyse (factorielle) discriminante
- régression logistique
- arbres de décision
- forêts aléatoires (random forest)
- réseaux de neurones
- SVM . . .

Nous verrons la méthode SVM (Support Vector Machine), les arbres de décision

Exercice

Expliquer brièvement :

- La différence entre données d'entraînement et données de test.
- Ce qu'est le surapprentissage (overfitting).
- Pourquoi le prétraitement des données est important en géographie et en aménagement urbain.

1. Données d'entraînement vs données de test

Données d'entraînement : C'est l'ensemble de données utilisé pour apprendre les patterns et construire le modèle. Le modèle ajuste ses paramètres en fonction de ces données.

Données de test : C'est un ensemble séparé, que le modèle n'a jamais vu, utilisé pour évaluer ses performances réelles. Cela permet de vérifier si le modèle généralise bien à de nouvelles données.

Analogie : C'est comme étudier avec des exercices (entraînement) puis passer un examen avec de nouvelles questions (test).

2. Le surapprentissage (overfitting)

Le surapprentissage se produit quand un modèle mémorise trop les détails spécifiques des données d'entraînement, incluant le bruit et les particularités, au lieu d'apprendre les patterns généraux.

Conséquence : Le modèle performe très bien sur les données d'entraînement mais échoue sur de nouvelles données.

Exemple urbain : Un modèle prédit parfaitement les prix immobiliers de votre quartier d'entraînement mais se trompe complètement dans un autre quartier similaire.

3. Importance du prétraitement en géographie/aménagement urbain

Le prétraitement est crucial car les données géospatiales présentent des défis spécifiques :

Échelles variables : Coordonnées GPS, densités de population, distances nécessitent une normalisation

Données manquantes : Zones non couvertes, capteurs défaillants

Projections cartographiques : Harmoniser les systèmes de coordonnées

Données hétérogènes : Combiner images satellitaires, recensements, réseaux de transport

Bruit spatial : Erreurs GPS, imprécisions cadastrales