



ÉCOLE MATHÉMATIQUE AFRICAINE : Bases mathématiques de l'intelligence artificielle

Introduction à l'apprentissage artificiel en santé (cours 1/2)

JEAN-DANIEL ZUCKER



UMMISCO

Unité Mixte Internationale



Jean-Daniel Zucker

UMMISCO : laboratoire de recherche international sous la tutelle de 7 partenaires académiques → explorer et concevoir des méthodes mathématiques et informatiques innovantes pour la modélisation et la simulation de systèmes complexes...



Directeur: Jean-Daniel ZUCKER (IRD)
Adjoint: Hassan HBID (UCA)
Mamadou SY (UGB) Alexis DROGOUL (IRD)



UCA

UGB

UCAD

UY1

72 membres

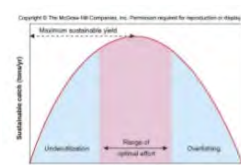
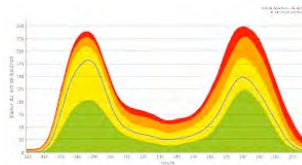
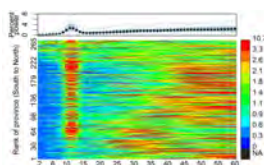
Prof. Maurice Tchuente
Prof. Samuel Bowong



Modèle de partenariat scientifique multi-sites équitable, dans lequel des enseignant-chercheur(e)s de plusieurs pays du Sud et de France collaborent

3
JDZucker

Concevoir des méthodes systémiques dont la finalité est d'être appliquées à des problématiques majeures de développement durable (ODD)

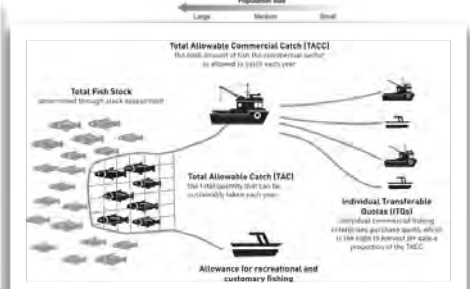


390 million dengue infections/year

Mieux contrôler les épidémies



Réduire la pollution des villes



Atteindre un effort de pêche durable

4
JDZucker

IRD is a public institute that undertakes

65 research units (jointly with other institutions)
2250 agents and a community of 7000 French
researchers

Over 1500 publications/year, 50% in co-
publication with partners

ECOBIO Ecology, biodiversity and
functioning of continental ecosystems

OCEANS Oceans, Climate and Resources

DISCO Internal dynamics and continents surface

SOC Societies and globalisation

SAS Health and Societies



Ce cours:



Plan

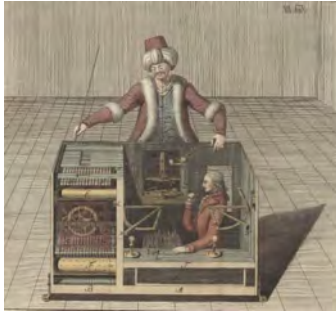
- JDZucker

IA & Santé : Historique



L'Intelligence artificielle (AI) est née au milieu du XXIème siècle

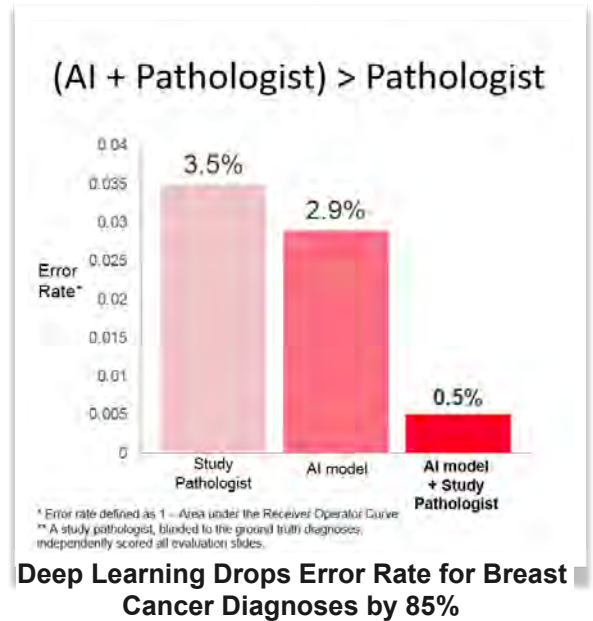
« L'intelligence artificielle (IA) vise à reproduire au mieux, à l'aide de machines, des activités mentales, qu'elles soient de l'ordre de la **compréhension**, de la **perception**, ou de la **décision** » www.larousse.fr



1769 : von Kempelen fabrique le **Turc mécanique** qui joue aux échecs.



1950: le **test de Turing**. Un test d'intelligence artificielle fondée sur la faculté d'une machine à **imiter la conversation humaine**.

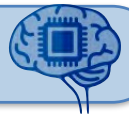


JAMA, vol. 318, no. 22, pp. 2199–2210, Dec. 2017.

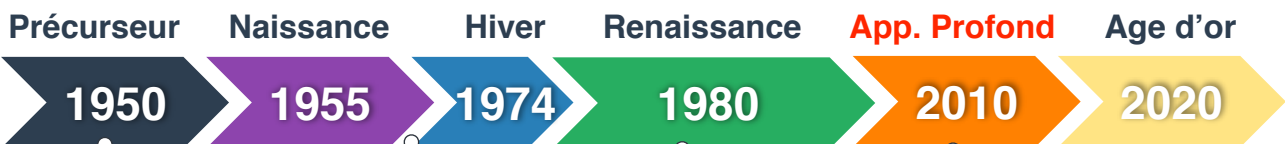
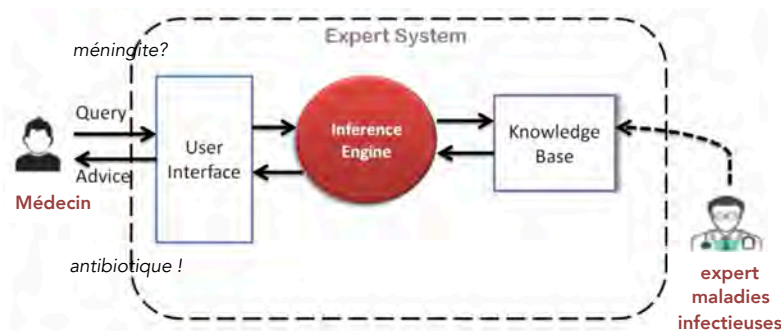
9
JDZucker

IA et médecine... une longue histoire

Intelligence Artificielle : compréhension/perception/décision



Le système expert MYCIN (1973-)



Test de Turing

Systèmes experts

Réseaux de neurones

Réseaux convolutifs

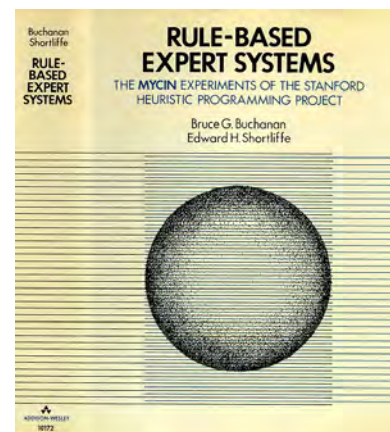
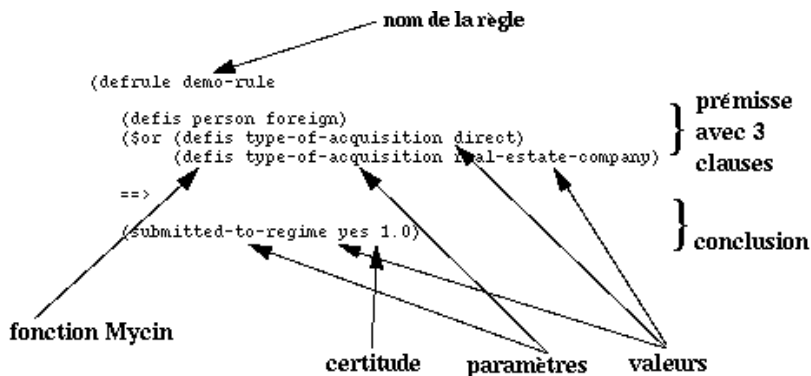
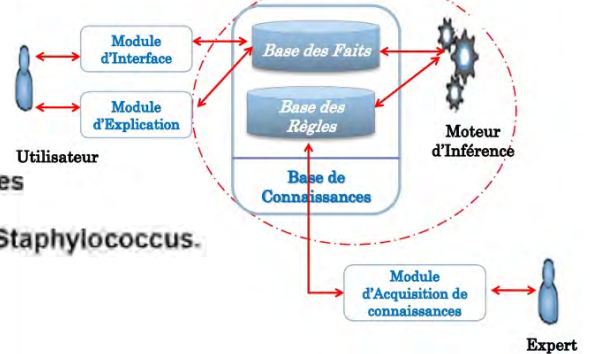
10
JDZucker

Mycin (1973-)

SI

la coloration de l'organisme est GRAM+
et si la morphologie de l'organisme est Cocci
et si le mode de développement de l'organisme est en colonies
ALORS
il existe une évidence (0.7) que l'identité de l'organisme soit Staphylococcus.

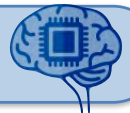
MYCIN : Architecture



11
JDZucker

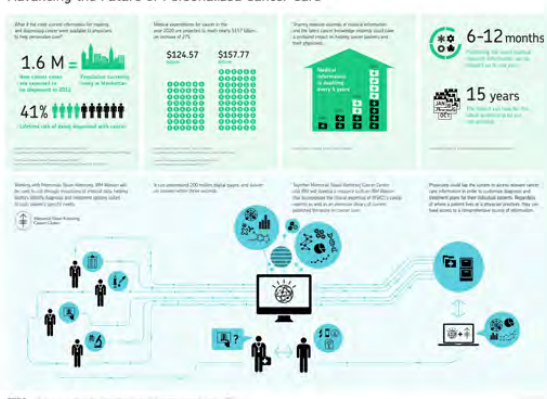
Aujourd'hui IBM a développé Watson

Intelligence Artificielle : compréhension/perception/décision

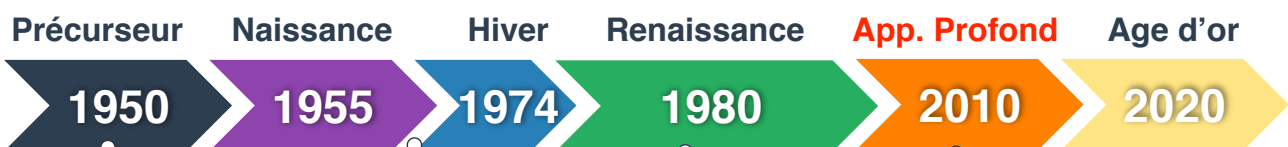


Watson for Oncology (2013)

Memorial Sloan Kettering & IBM Watson:
Advancing the Future of Personalized Cancer Care



<https://www.youtube.com/watch?v=UpFHNGF4F8o&pbjreload=10>



Précurseur Naissance Hiver Renaissance App. Profond Age d'or

1950 1955 1974 1980 2010 2020

Test de Turing Systèmes experts Réseaux de neurones Réseaux convolutifs

12

JDZucker

De nombreuses sociétés s'intéressent à l'IA en santé



13
JDZucker

Aujourd'hui Chatterbot et Medecine



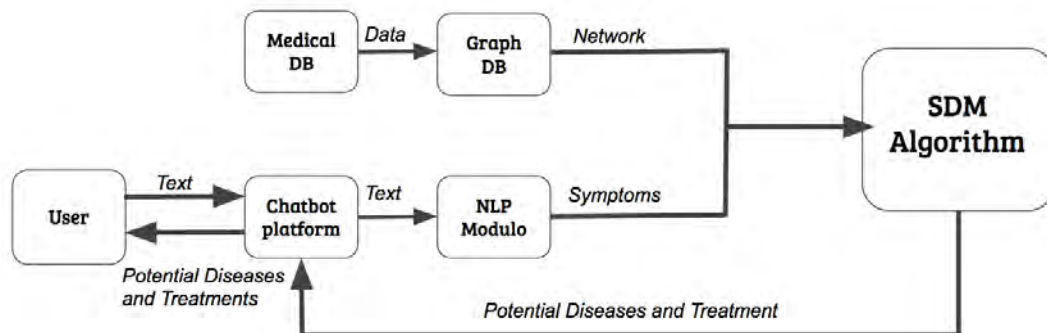
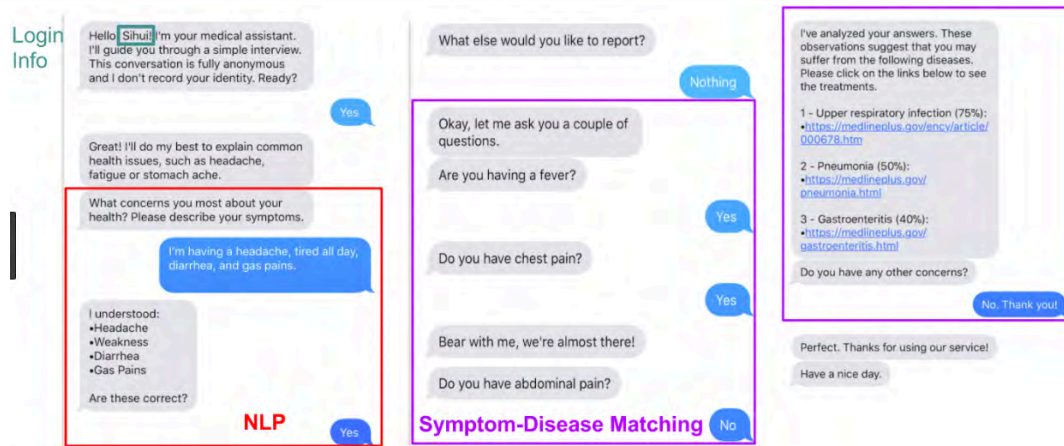
	Your.MD	Sensely	Buoy Health	Infermedica	Florence
Funds raised	\$17.3 MM	\$11.8 MM	\$9 MM	\$5 MM	Unspecified
Year founded	2012	2013	2014	2012	2016
HQ location	London, UK	San Fran., CA	Boston, USA	Wroclaw, POL	Dortmund, DEU
Staff size	49	20	23	26	Unspecified
Target user	Patient, Clinician	Patient, Clinician	Patient	Patient, Clinician	Patient
Types of data processed	Chat (text)	Chat (text), image, video	Chat (text)	Chat (text), image, video	Chat (text)
Est. current users	Unspecified	Unspecified	Unspecified	Unspecified	2,000+?

Data collected from Crunchbase and LinkedIn, March 2018



14
JDZucker

Can a Chat a Day Keep the Doctor Away?



<https://chatbotslife.com/can-a-chat-a-day-keep-the-doctor-away-f7539fea289a>

15
JDZucker

Le rôle croissant de l'apprentissage profond en médecine

Intelligence Artificielle : compréhension/perception/décision

Apprentissage Automatique

Apprentissage Profond

1950

Alan Turing created a test to check if a machine could fool a human being into believing it was talking to a machine.

1957

First neural network for computers (the perceptron) was invented by Frank Rosenblatt, which simulated the thought processes of the human brain.

1979

Students of Stanford University, California, invented the Stanford Cart which could navigate and avoid obstacles on its own.

2002

A software library for Machine Learning, named Torch is first released.

1952

The first computer learning program, a game of checkers, was written by Arthur Samuel.

1967

The Nearest Neighbor Algorithm was written.

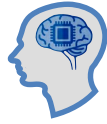
1997

IBM's Deep Blue beats the world champion at Chess.

2016

AlphaGo algorithm developed by Google DeepMind managed to win five games out of five in the Chinese Board Game Go

Dans la médecine, de l'IA venant de 3 grands champs



S. Russell and P. Norvig, The Artificial Intelligence a Modern Approach (AIMA), 3e Preview Edition. Prentice Hall, 2009.

Raisonner/Comprendre



Percevoir/Interagir

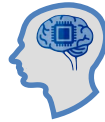


Décider/Agir



17
JDZucker

1er champs de l'IA & Santé:



Raisonner/Comprendre

Systèmes Experts



MÉDECINE PRÉDICTIVE

PRÉDICTION D'UNE MALADIE
ET/OU DE SON ÉVOLUTION

Coach/Application santé

Google Im2Calories project Compter les calories à partir d'une image !

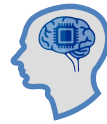


MÉDECINE DE PRÉCISION

RECOMMANDATION DE
TRAITEMENT PERSONNALISÉ

18
JDZucker

2ème champs de l'IA:



Percevoir/Interagir

Chatterbots médicaux



Membres artificiels



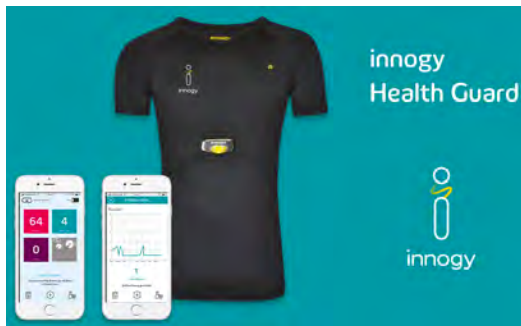
19

Exosquelette



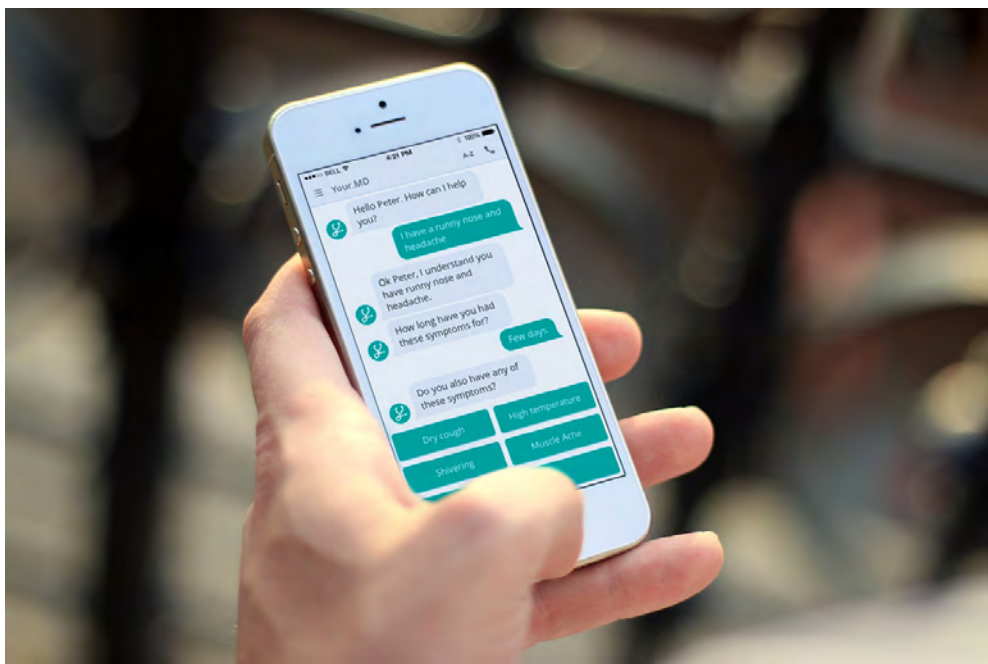
JDZucker

Capteur (e.g. T-Shirt Cardio)



Les techniques qui vont introduire des ruptures : l'IA joue les docteurs grâce aux **Chatterbots**.

Dans (A routine check-up at the digital medical centre. Microbial Biotechnology 2016) les auteurs imaginent dans un futur proche une discussion entre Mme Tristam-Wellness et un robot Dr. Siri von App



20

3ème champs de l'IA:

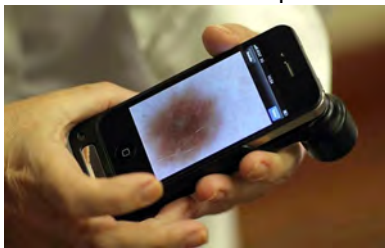
Aide au diagnostic

... de plus de 4 000 maladies génétiques, dont beaucoup sont extrêmement difficiles à identifier



<https://www.face2gene.com/>

... de cancer de la peau



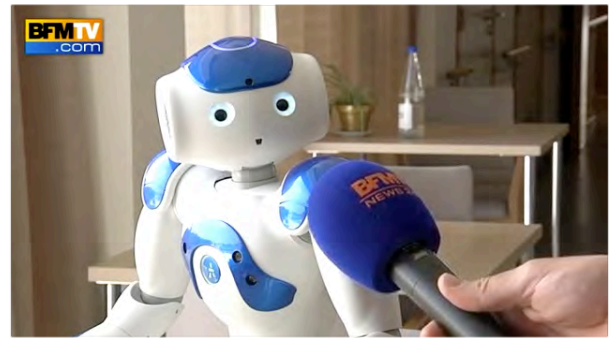
<https://fmh-association.org/dermoscreen-application-diagnostiquer-cancer-peau/>



Décider/Agir

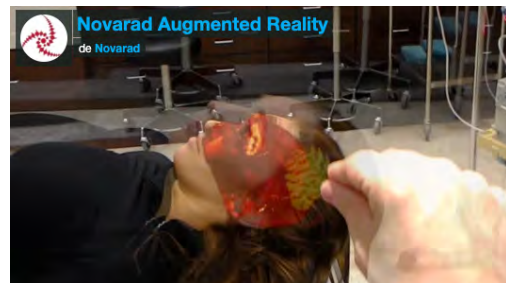


Robots aide-soignant



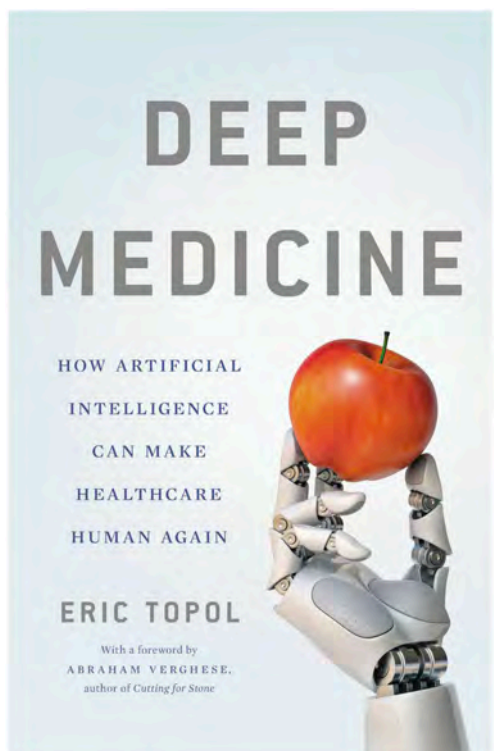
21

Chirurgie assistée par ordinateur



JDZucker

Le rôle croissant de l'apprentissage profond en médecine



chapter one	INTRODUCTION TO DEEP MEDICINE
chapter two	SHALLOW MEDICINE
chapter three	MEDICAL DIAGNOSIS
chapter four	THE SKINNY ON DEEP LEARNING
chapter five	DEEP LIABILITIES
chapter six	DOCTORS AND PATTERNS
chapter seven	CLINICIANS WITHOUT PATTERNS
chapter eight	MENTAL HEALTH
chapter nine	AI AND HEALTH SYSTEMS
chapter ten	DEEP DISCOVERY
chapter eleven	DEEP DIET
chapter twelve	THE VIRTUAL MEDICAL ASSISTANT
chapter thirteen	DEEP EMPATHY

JDZucker

Très prometteuse, en 2019 l'IA en médecine reste superficielle et n'a pas encore tenu ses promesses



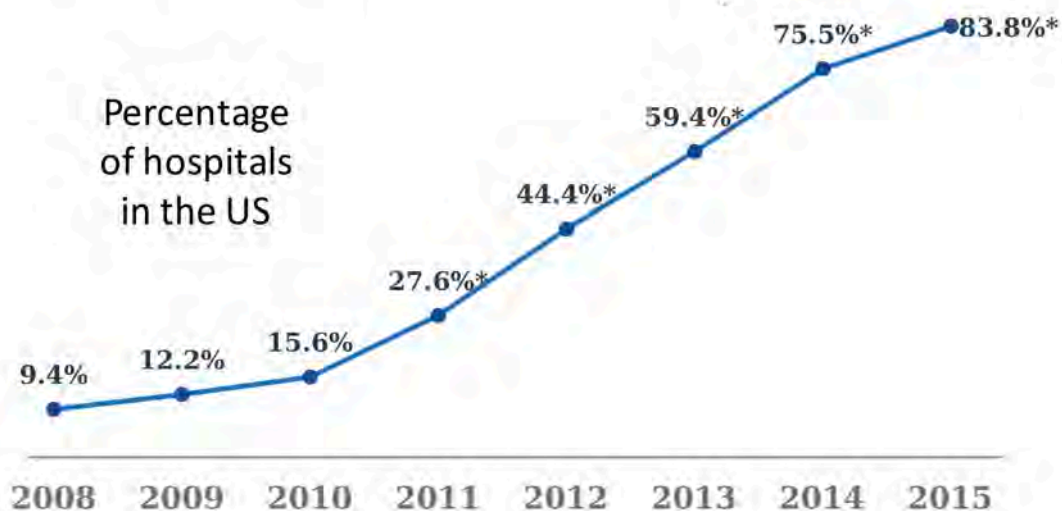
FIGURE 3.3: IBM Watson ad.

L'informatique peut faire une grande différence, mais jusqu'à présent → pas tenu ses promesses.

Difficulté → l'assemblage et l'agrégation des données a été sous-estimée, non seulement par Watson, mais aussi par toutes les entreprises de technologie qui s'intéressent aux soins de santé.

Et la médecine est en retard

Adoption of Electronic Health Records (EHR) has increased 9x since 2008



[Henry et al., ONC Data Brief, May 2016]

Et la médecine est en retard

"Bien que la plupart des secteurs du monde soient bien avancés dans la **quatrième révolution industrielle**, qui est centrée sur l'utilisation de l'IA,

la médecine est encore bloquée dans la première phase de la troisième, qui a vu la première utilisation généralisée des ordinateurs et de l'électronique.

Que les fichiers MP3 soient compatibles avec toutes les marques de lecteurs de musique, par exemple, alors que la médecine n'a encore jamais vu de dossiers médicaux électroniques largement compatibles et conviviaux illustre la lutte du secteur pour le changement. » Eric Topol Deep Medicine 2019

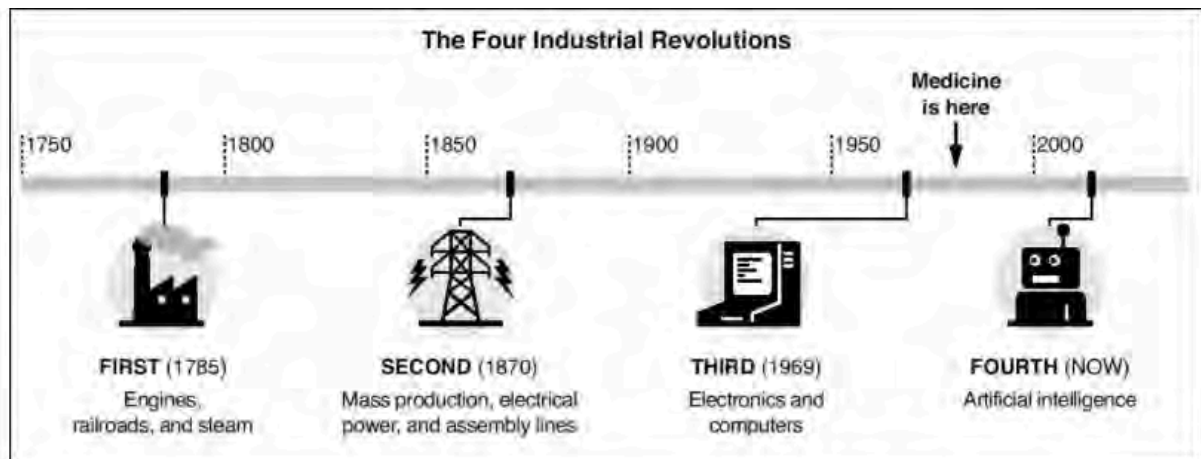
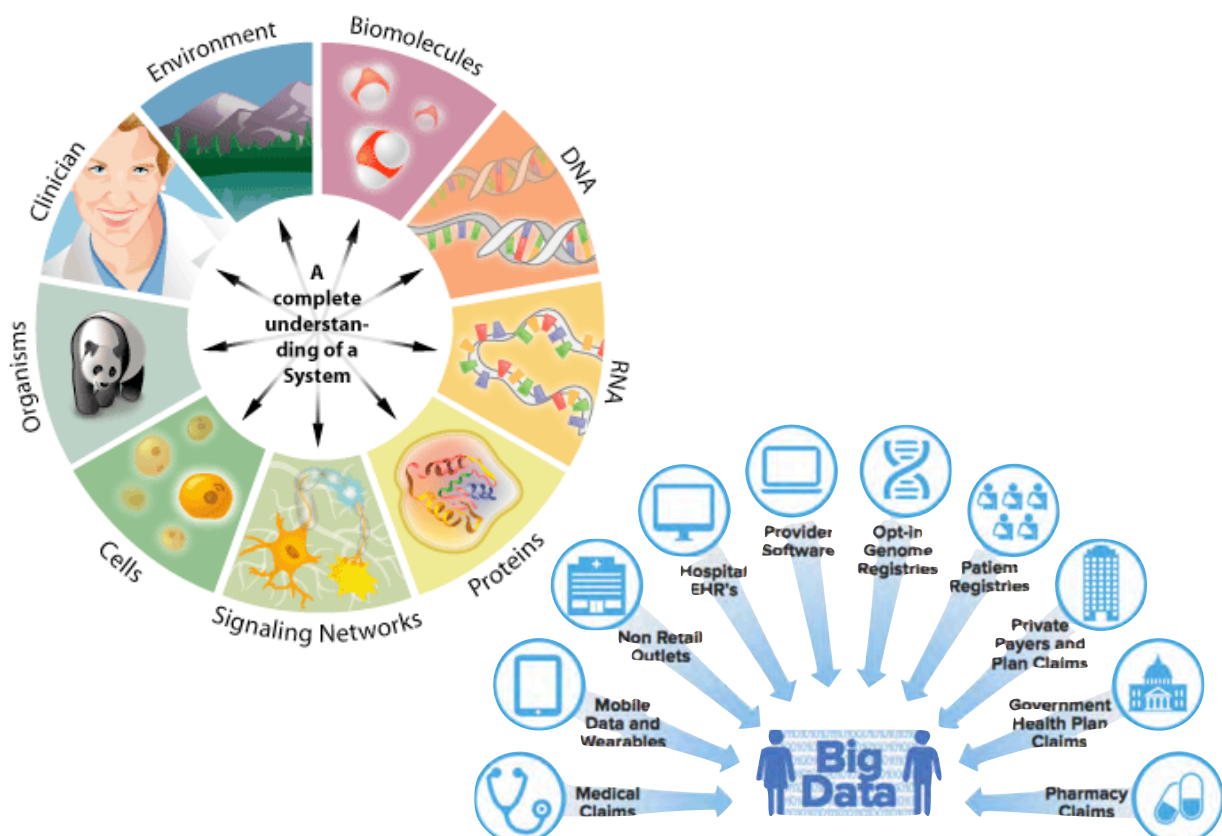


FIGURE 1.4: The four Industrial Revolutions. Source: Adapted from A. Murray, "CEOs: The Revolution Is Coming," *Fortune* (2016): <http://fortune.com/2016/03/08/davos-new-industrial-revolution>.

25
cker

IA & Santé : Types de données



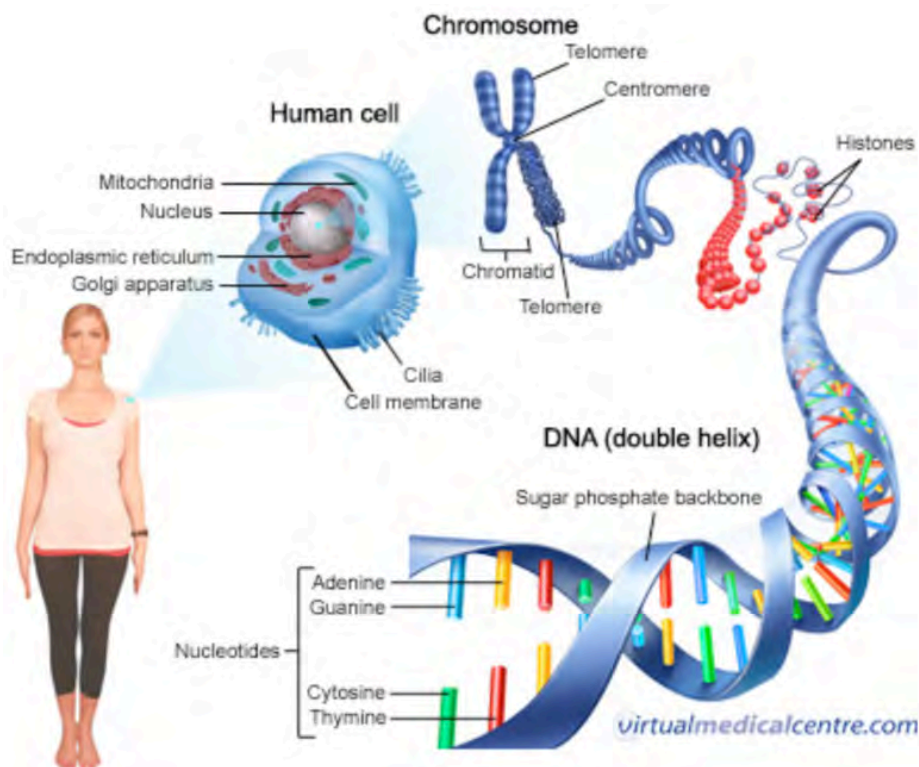
26
JDZucker

Une grande diversité de données



27
JDZucker

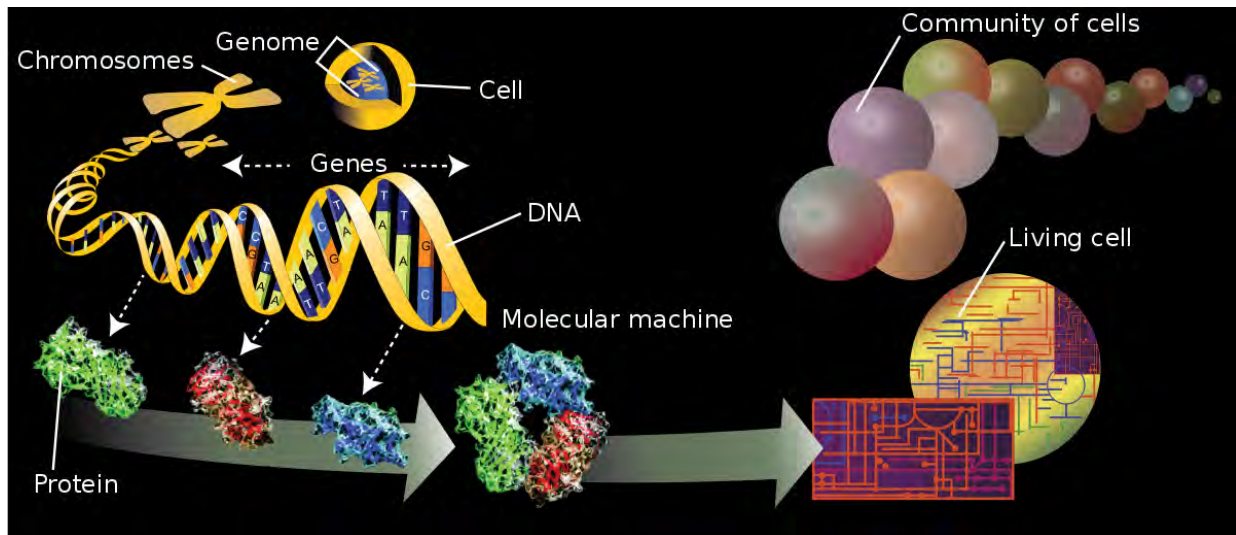
De l'ADN à la vie chez l'humain



1 corps = 10^{14} cellules humaines (et 100 x plus non humaines)
 1 cellule = 6×10^9 ACGT codant pour 20 000+ gènes

28
JDZucker

Quantification des omes → omiques

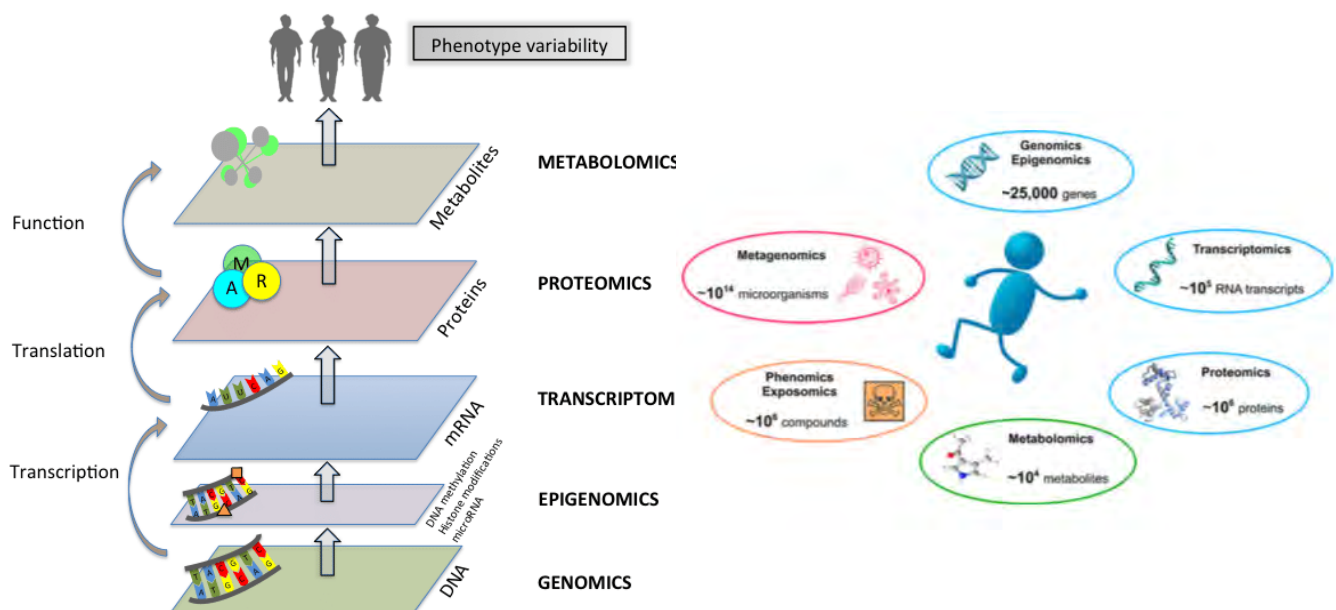


Le néologisme omics de langue anglaise fait référence de façon informelle à un domaine d'étude en biologie se terminant par -omique, comme la génomique, la protéomique ou la métabolomique.

Omics vise la caractérisation et la quantification collectives de pools de molécules biologiques qui se traduisent par la structure, la fonction et la dynamique d'un organisme ou d'organismes. Wikipedia

29
JDZucker

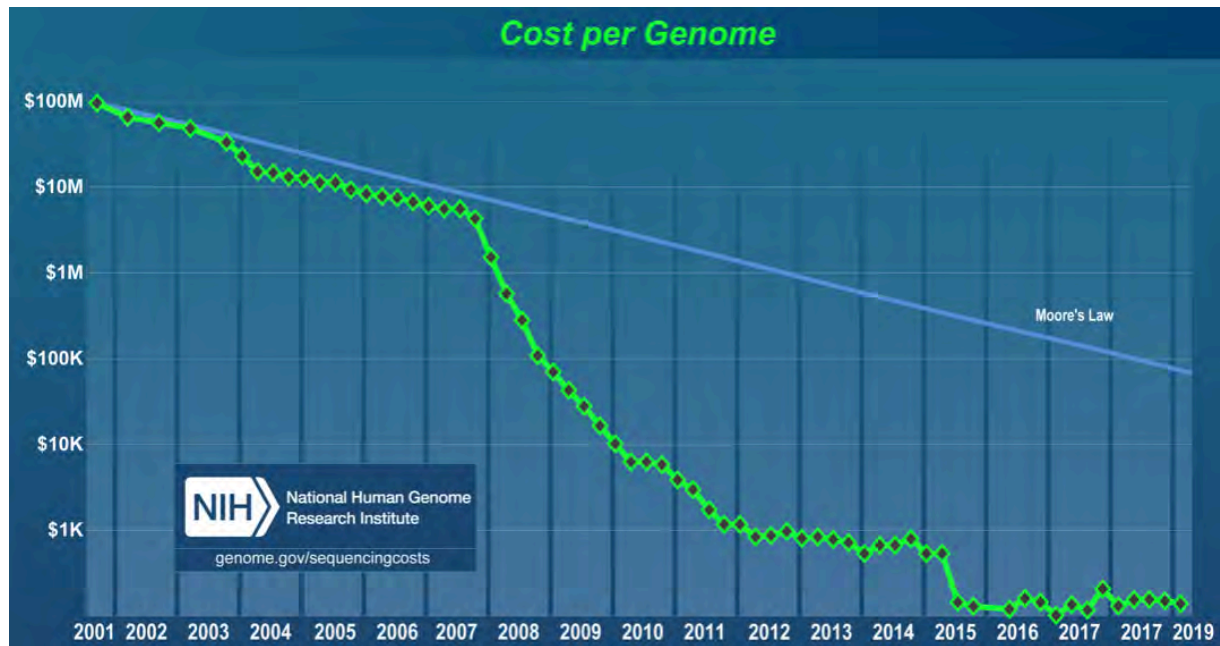
Les données « Omics » au phénotypage profond



L. Yu, K. Li, and X. Zhang, "Next-generation metabolomics in lung cancer diagnosis, treatment and precision medicine: mini review," *Oncotarget*, 8(10) 2017.

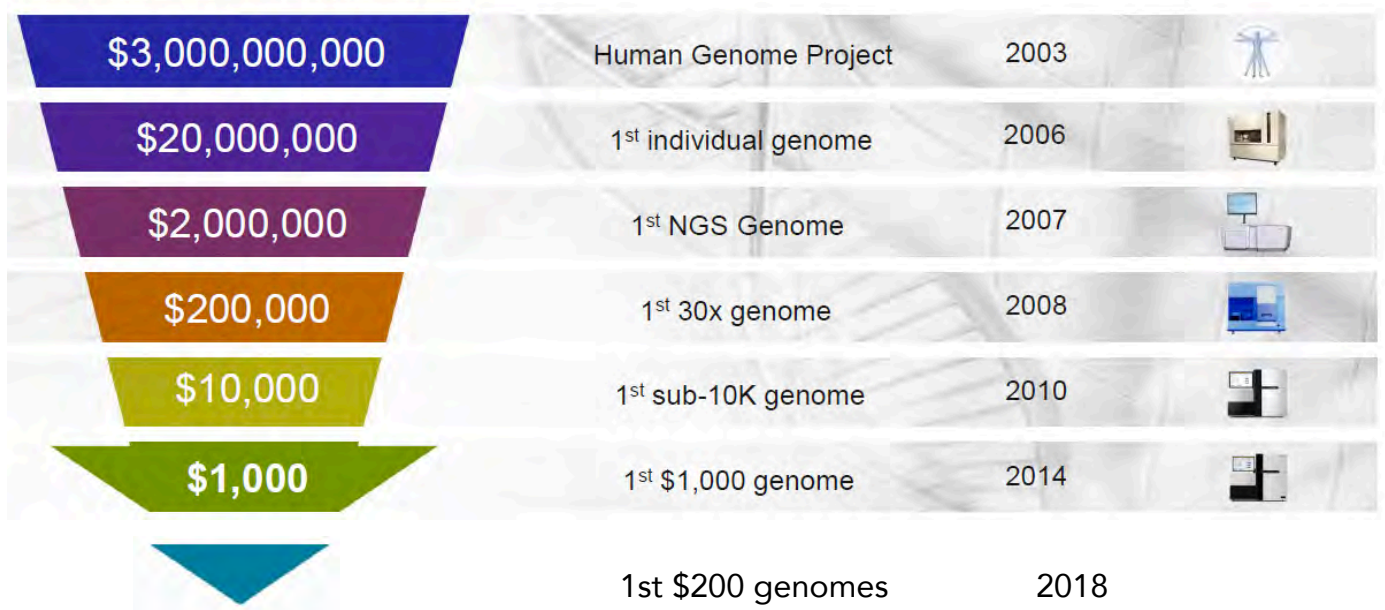
JDZucker

Coût des génomes



31
JDZucker

Cost Per Genome

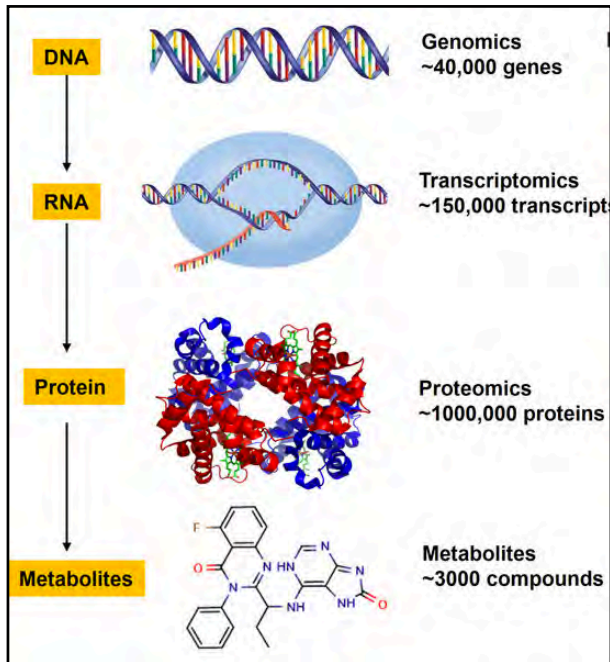


<https://www.cnbc.com/2019/07/01/for-600-veritas-genetics-sequences-6point4-billion-letters-of-your-dna.html>

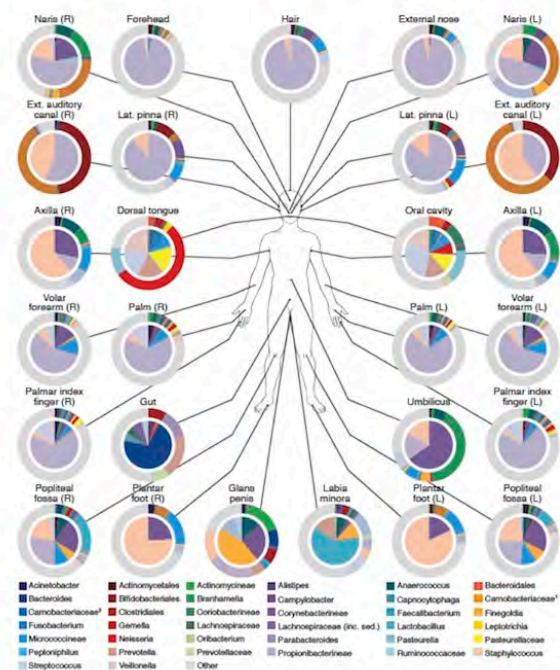
32
JDZucker

Les données « Omics » permettent de nous caractériser très finement, nous et... nos hôtes.

Analyser nos propres cellules



Analyser nos bactéries



F. S. MD, et al., "Translational Radiomics: Defining the Strategy Pipeline and Considerations for Application-Part 1: From Methodology to Clinical Implementation," Journal of the American College of Radiology, 2018

JDZucker 35

Microbiote intestinal humain : un organe oublié

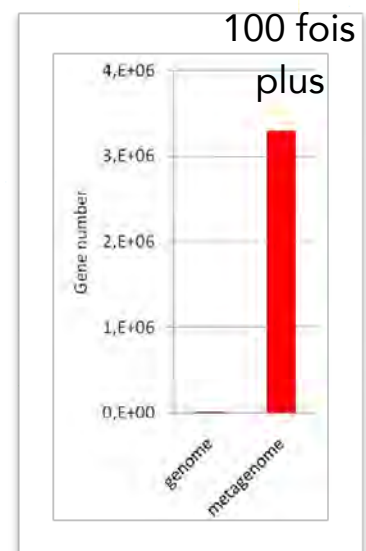
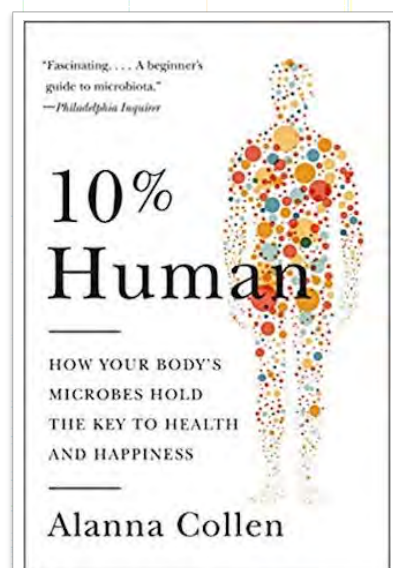
Du bébé "stérile" à la naissance → 2 kg de micro-organismes, sur les 100 billions de cellules du corps humain, seule 1 sur 10 est humaine.

Métabolisme

production de vitamines
dégradation des aliments
extraction énergétique

Système immunitaire

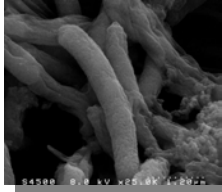
"éducation" des défenses immunitaires innées



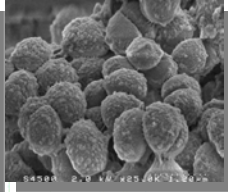
JDZucker 36

Quantification de notre microbiome

La plupart des micro-organismes sont inconnus et non cultivables...



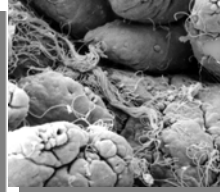
Faecalibacterium prausnitzii
Photos UEPSD



Ruminococcus spp



Clostridium difficile
en caecum souris



Bactéries ancrées dans
une Plaque de Peyer,
Intestin de souris



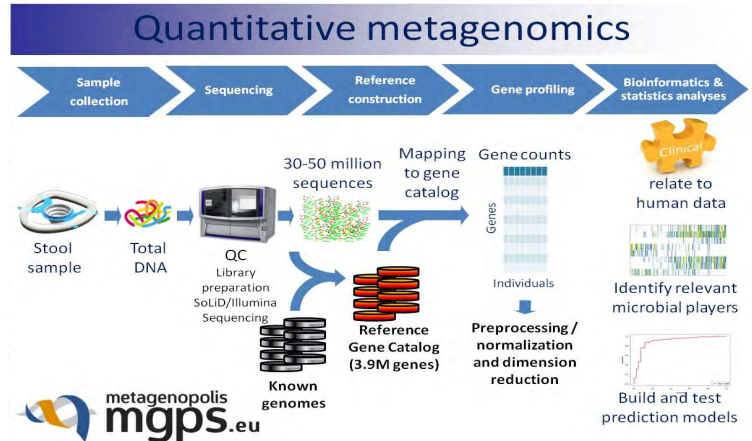
Bacteroides dorei



Escherichia coli



Ehrlich SD ©

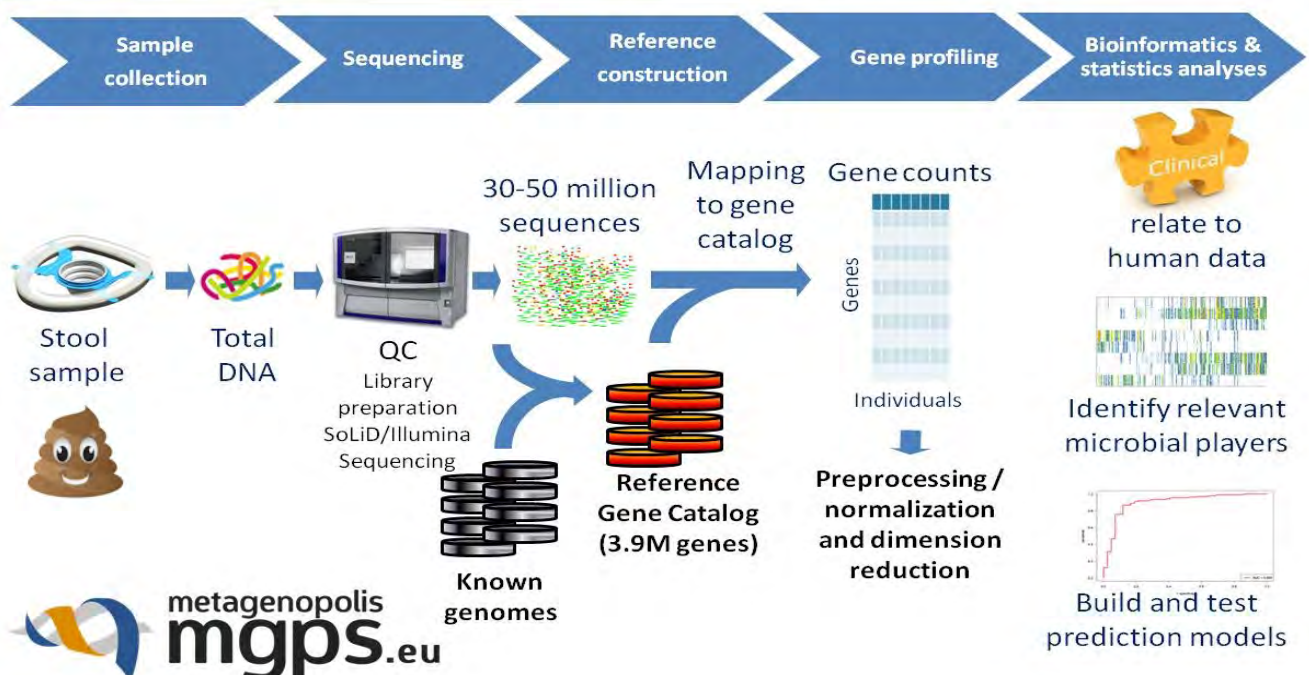


A Powerful Microscope to Scan the neglected organ

Quantification de notre microbiome

La plupart des micro-organismes sont inconnus et non cultivables...

Quantitative metagenomics

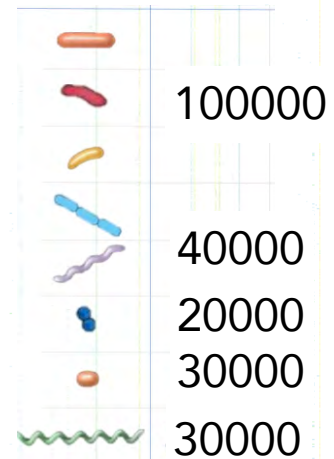


Quantification de notre microbiome (vision simple)



100 €

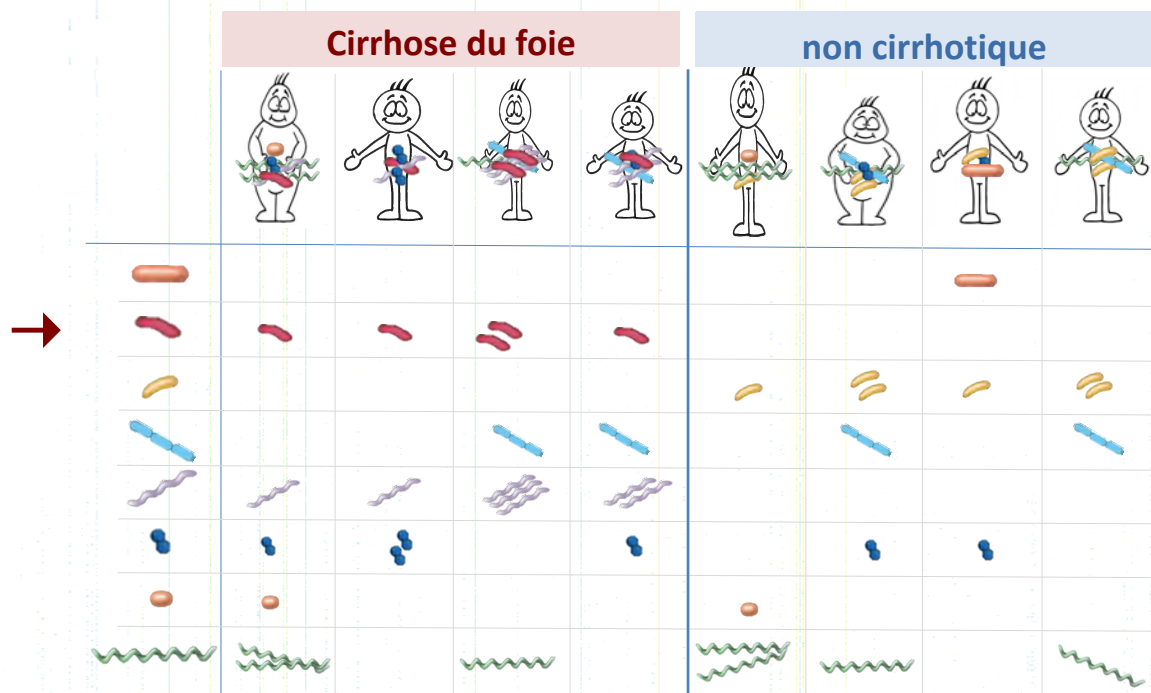
Bactéries Nombre



37

JDZucker

Vers une médecine de précision des maladies intégrant la métagénomique.



Signature géniques impliquant plusieurs gènes bactériens parmi des **millions**...

38

JDZucker

Le microbiote intestinal comme moyen de diagnostic

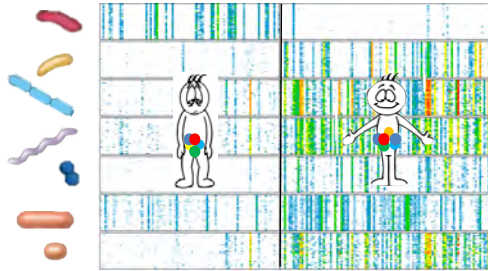
Cirrhose du foie



Diagnostic actuel:

- biopsie 40%
- symptômes cliniques ou imagerie 60%

Apprentissage Automatique

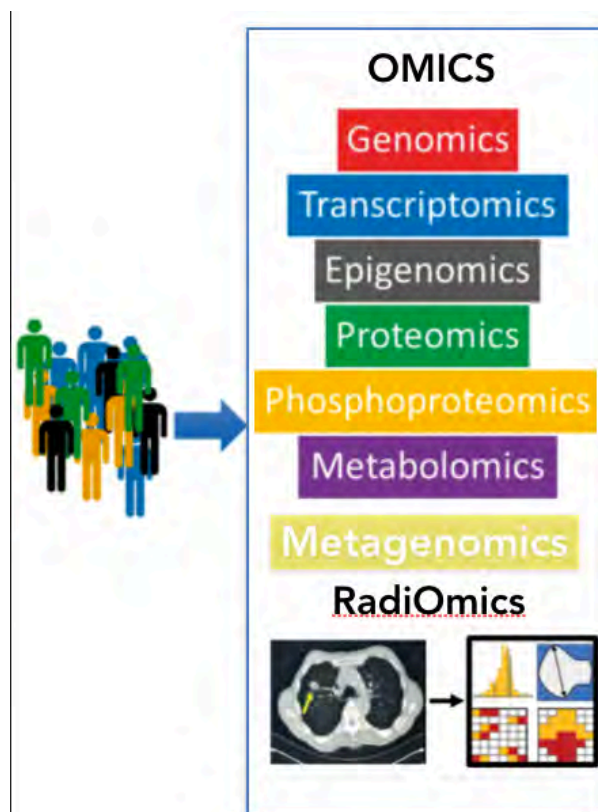


Modèle prédictif à partir de sept espèces

Diagnostic très performant (AUC=94%), non invasif

Qin N. et al. Nature 2014 ³⁹ JDZucker

Des données massives



41

[illegible]

instances/
input
vectors/
observations/
examples/
records/
...

In UCI repository the Median is 42 → [0,0032 , 990114]

42

	Left	Right	More Fun
t1	sNaba8-eb3Y	wYkP68dxEW	left
t2	sNaba8-eb3Y	2h2m5XE-N4	left
t3	fy_FQMqQjok	sNaba8-eb3Y	left
t4	Vr4D8xO2IBY	sNaba8-eb3Y	right
t1	sNaba8-eb3Y	ddtRnstrfE	left
t2	W5mH4E3xxdw	fy_FQMqQjok	right
t3	fy_FQMqQjok	Mo-ZR7-X7es	left
t4	H2PUQQNRvQg	ZruyMxf4JL	right
t1	vX95jGGuDo	wu45U70w7LA	left
t2	sNaba8-eb3Y	2FHf-9teZP0	left
t3	sP3YdG_Uio	fy_FQMqQjok	right
...	...	...	...
!912967	oPDM18jWzG8	Z2cqRpRtQk	right
!912968	6uT5EAP4dAw	t4KvXsgVwMg	right
!912969	W9y6nwBwvQq	t4KvXsgVwMg	right

$N = 1138562$
 $p = 3$
 $N/p = 379,521$

N = 150
p = 4
N/p = 38

N = 72
p = 7729
N/p = 0,009 = 1%



... but **N/p** is another key characteristic.

In UCI repository the Median is 42 → [0,0032 , 990114]

43

Integrated Gene Catalog (IGC)

Li et al. 7/6/2014)

nature
biotechnology

An integrated catalog of reference genes in the human gut microbiome

Junhua Li^{1,2,3,4,5}, Huijie Jia^{1,2,3}, Xianghang Cai^{1,2,3}, Huanzi Zhong^{1,2,3}, Qiang Feng^{1,2,3}, Shinichi Sunagawa⁶, Manimohan Arumugam^{7,8}, Jens Rost Kultima⁹, Edi Prifti¹⁰, Trine Nielsen¹¹, Agnieszka Sierakowska Juncker¹², Chayavanh Manichanh¹³, Bing Chen¹⁴, Wensui Zhang¹⁵, Florence Leventer¹⁶, Juan Wang¹⁷, Xun Xu¹⁸, Liang Xiao¹⁹, Suihua Liang²⁰, Dongya Zhang²¹, Zhaodi Zhang²², Wensong Chen²³, Hailong Zhao²⁴, Jumana Yousef Al-Aama^{25,26}, Sherif Edris^{11,12}, Huanming Yang^{11,12,13}, Jian Wang^{11,13}, Torben Hansen¹¹, Henrik Bjørn Nielsen¹¹, Søren Brunak¹¹, Karsten Kristiansen¹¹, Francisco Guarner¹¹, Oluf Pedersen¹¹, Joel Dore^{11,12}, S Dusko Ehrlich^{14,15}, MetaHIT Consortium¹⁶, Peer Bork^{17,18} & Jun Wang^{11,13,19}

N = 1267

p = 9,879,896

N/p = 0.00013

Co-Abundance Gene Groups (CAGS)

(Nielsen, Almeida & al., 7/6/2014)

nature
biotechnology

Identification and assembly of genomes and genetic elements in complex metagenomic samples without using reference genomes

H Bjørn Nielsen^{1,2,3,4}, Mathieu Almeida^{1,5,6,7}, Agnieszka Sierakowska Juncker^{1,2}, Simon Rasmussen¹, Junhua Li^{8,9}, Shinichi Sunagawa¹⁰, Damien R Plichta¹¹, Laurent Gautier¹², Anders G Pedersen¹³, Emmanuelle Le Chatelier¹⁴, Eric Pelletier^{15,16}, Ida Bonde¹⁷, Trine Nielsen¹⁸, Chayavanh Manichanh¹⁹, Manimohan Arumugam^{20,21}, Jean-Michel Ratto²², Marcelo B Quintanilha dos Santos²³, Nikunj Bhon²⁴, Natalia Borrero²⁵, Kristoffer S Burghoff²⁶, Fouad Boumezeur²⁷, Francine Castellu²⁸, Joël Dore²⁹, Piotr Dworzynski³⁰, Francisco Guarner³¹, Torben Hansen^{32,33}, Falk Hildebrand^{34,35}, Rolf S Kaas³⁶, Sean Kennedy³⁷, Karsten Kristiansen³⁸, Jens Rost Kultima³⁹, Pierre Leonard⁴⁰, Florence Leventer⁴¹, Ole Lund⁴², Bouziane Mounen⁴³, Denis Le Poulley^{44,45}, Nicolas Pons⁴⁶, Oluf Pedersen^{47,48,49}, Edi Prifti⁵⁰, Junjie Qian⁵¹, Jørgen Raas^{52,53,54}, Søren Sørensen⁵⁵, Julien Tap⁵⁶, Sebastian Timm⁵⁷, David W Usery⁵⁸, Takaji Yamada^{59,60}, MetaHIT Consortium⁶¹, Pierre Renaud⁶², Thomas Sicheritz-Ponten⁶³, Peer Bork^{64,65}, Jun Wang^{66,67,68}, Søren Brunak^{1,2} & S Dusko Ehrlich^{64,65}

N = 49

p = 7381

N/p = 0.0066

MetaGenomics Species (MGS)

nature
biotechnology

Identification and assembly of genomes and genetic elements in complex metagenomic samples without using reference genomes

H Bjørn Nielsen^{1,2,3,4}, Mathieu Almeida^{1,5,6,7}, Agnieszka Sierakowska Juncker^{1,2}, Simon Rasmussen¹, Junhua Li^{8,9}, Shinichi Sunagawa¹⁰, Damien R Plichta¹¹, Laurent Gautier¹², Anders G Pedersen¹³, Emmanuelle Le Chatelier¹⁴, Eric Pelletier^{15,16}, Ida Bonde¹⁷, Trine Nielsen¹⁸, Chayavanh Manichanh¹⁹, Manimohan Arumugam^{20,21}, Jean-Michel Ratto²², Marcelo B Quintanilha dos Santos²³, Nikunj Bhon²⁴, Natalia Borrero²⁵, Kristoffer S Burghoff²⁶, Fouad Boumezeur²⁷, Francine Castellu²⁸, Joël Dore²⁹, Piotr Dworzynski³⁰, Francisco Guarner³¹, Torben Hansen^{32,33}, Falk Hildebrand^{34,35}, Rolf S Kaas³⁶, Sean Kennedy³⁷, Karsten Kristiansen³⁸, Jens Rost Kultima³⁹, Pierre Leonard⁴⁰, Florence Leventer⁴¹, Ole Lund⁴², Bouziane Mounen⁴³, Denis Le Poulley^{44,45}, Nicolas Pons⁴⁶, Oluf Pedersen^{47,48,49}, Edi Prifti⁵⁰, Junjie Qian⁵¹, Jørgen Raas^{52,53,54}, Søren Sørensen⁵⁵, Julien Tap⁵⁶, Sebastian Timm⁵⁷, David W Usery⁵⁸, Takaji Yamada^{59,60}, MetaHIT Consortium⁶¹, Pierre Renaud⁶², Thomas Sicheritz-Ponten⁶³, Peer Bork^{64,65}, Jun Wang^{66,67,68}, Søren Brunak^{1,2} & S Dusko Ehrlich^{64,65}

N = 49

p = 741

N/p = 0.067

Precision Medicine from Metagenomics requires scaling up Data Mining algorithms

44

- Learning **prognostics/diagnostics models** ? estimate their performances ?
- reconstructing (causal) networks **between genes or Metagenomic Species** ?



Clinical Data

p=100

p²=10⁶



Transcriptomics Data

p=10,000

p²=10⁸



Metagenomics Data

p=10,000,000

p²=10¹⁴

Des données "massives" pour décrire chaque patient - Big Data



Données Clinique
 $p=100$
 $p^2=10^6$



Données Transcriptomiques
 $p=10,000$
 $p^2=10^8$



Données Métagénomiques
 $p=10,000,000$
 $p^2=10^{14}$

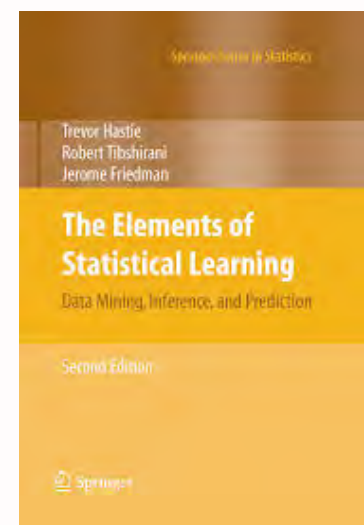
Also referred as High-Dimensional Problems: $p \gg N$

46

New chapter in the second edition of The Elements of Statistical Learning
 [Hastie et al., 2009]

High-Dimensional Problems: $p \gg N$

18 High-Dimensional Problems: $p \gg N$	649
18.1 When p is Much Bigger than N	649
18.2 Diagonal Linear Discriminant Analysis and Nearest Shrunken Centroids	651
18.3 Linear Classifiers with Quadratic Regularization	654
18.4 Linear Classifiers with L_1 Regularization	661



● **Hard problems** because the required number N is about the VCDim of the concept space

● If concepts are linear then $\text{VCDim} = \text{\#attributes} = p$ and we need $N \geq p$ ($N/p \geq 1$) to learn...

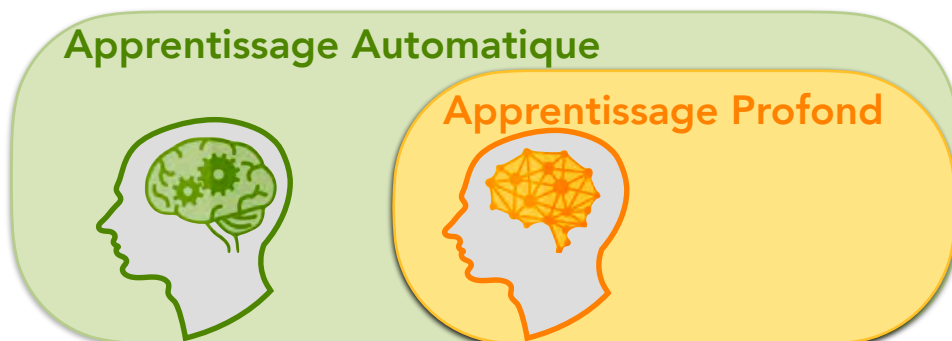
Problèmes liés à l'apprentissage en santé

- ✓ la quantité de données
- ✓ l'intégration de données
- ✓ l'anonymisation des données

- ✓ la compréhension des biais des algorithmes (fairness)
- ✓ la robustesse des algorithmes
- ✓ l'efficacité et les garanties des algorithmes sur des problèmes $n \ll p$

- ✓ l'interprétabilité des modèles
- ✓ la prise en compte du coût dans les modèles...

IA & Santé : l'importance de l'apprentissage



Décliner les trois domaines du Machine Learning en santé

Les 3 grandes tâches du Machine Learning



Supervisé



Apprendre à partir
d'exemples

Non
Supervisé



Apprendre à
organiser des
observations

Par
Renforcement



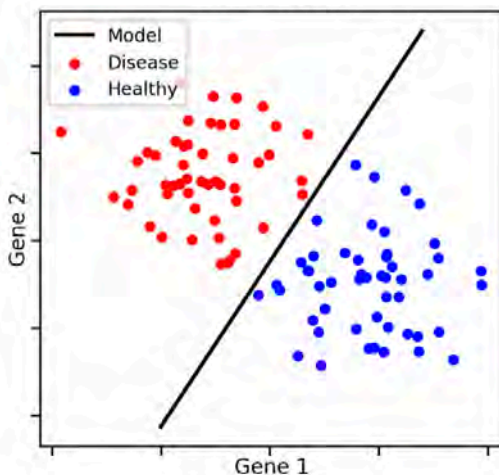
Apprendre à
décider par
essais/erreur

49
JDZucker

Apprentissage supervisé : à partir de données géniques

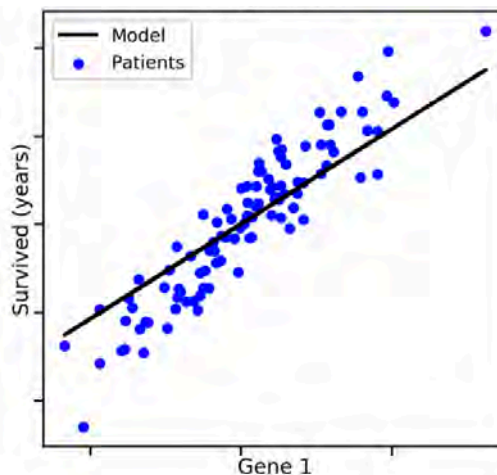


Classer des patients selon
l'expression de leur gènes



Séparateur linéaire
SVM Gradient Boosting RF

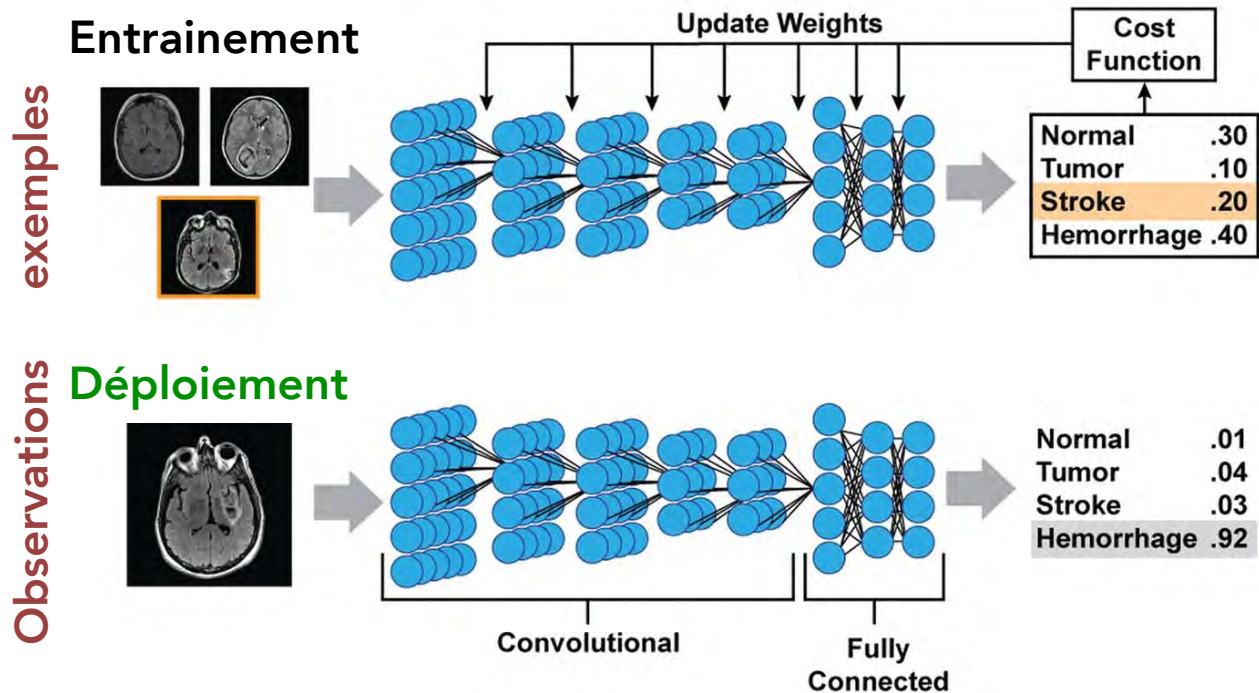
Prédire la survie des patients en
fonction de l'expression d'un gène



Régression linéaire
Régression logistique
Régression polytonique

50
JDZucker

Apprentissage supervisé : à partir d'images



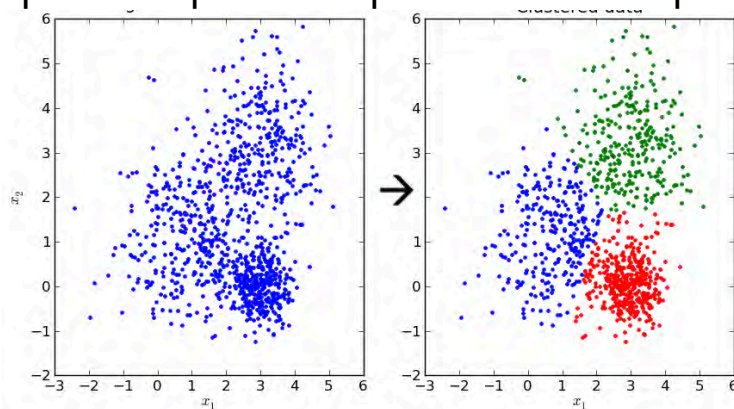
Zaharchuk et al. "Deep Learning in Neuroradiology," AJNR Am J Neuroradiol, (39)10. 2018.

JDZucker

Apprentissage non supervisé: clinique, omique ou images

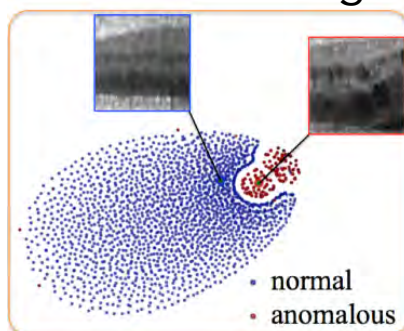


Regrouper des patients proches: exemple k-means



k ?

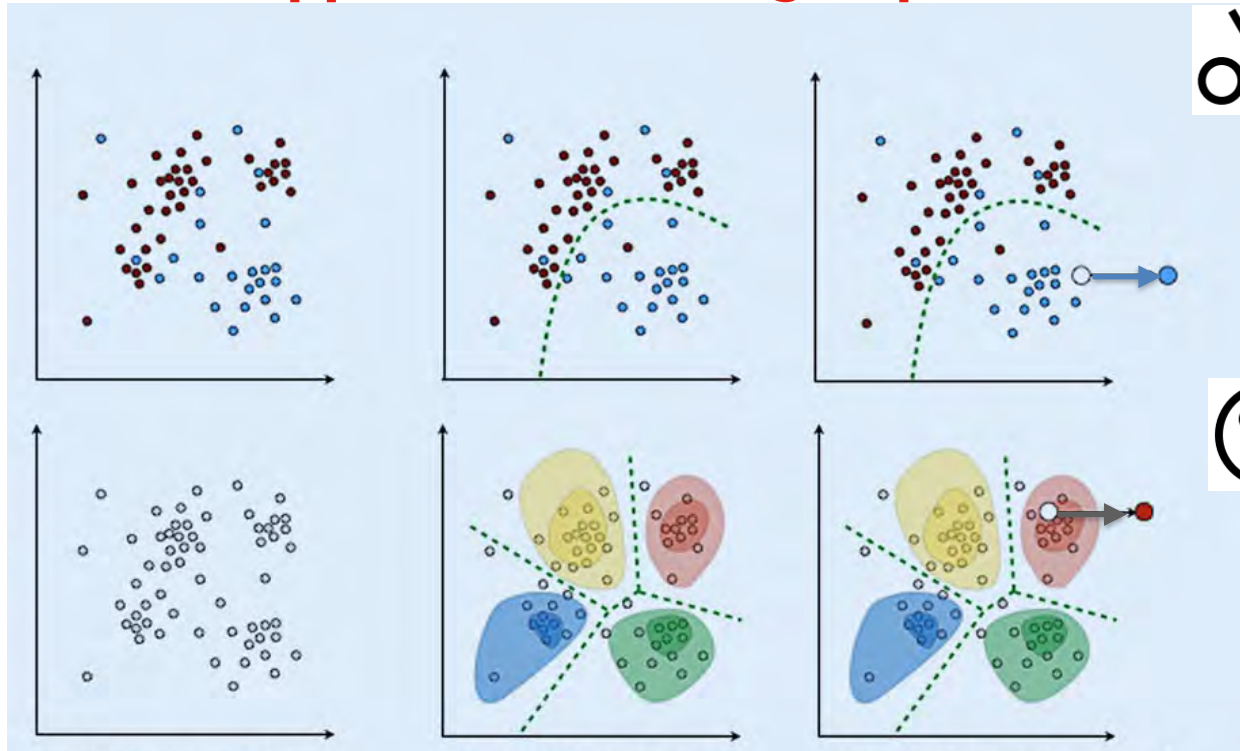
détecter des anomalies d'images d'examens



seuil ?

52
JDZucker

Utilisation des modèles pour prédire une classe ou l'appartenance à un groupe



Données patients

Modèles

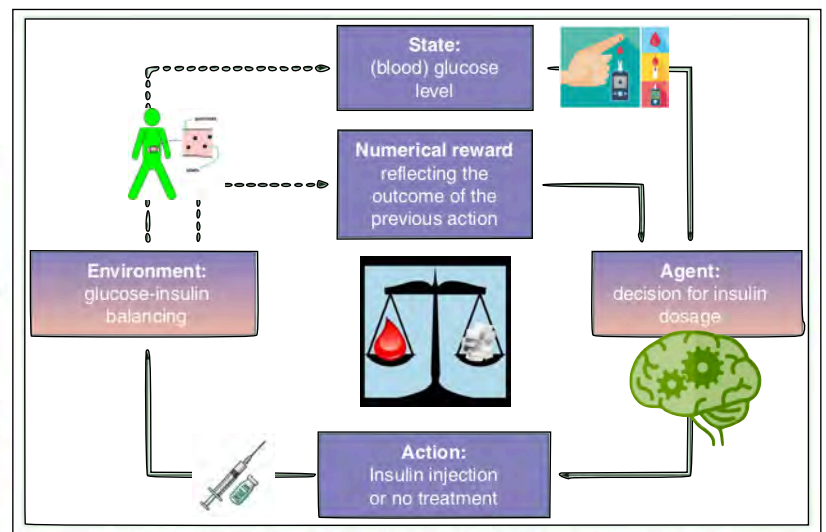
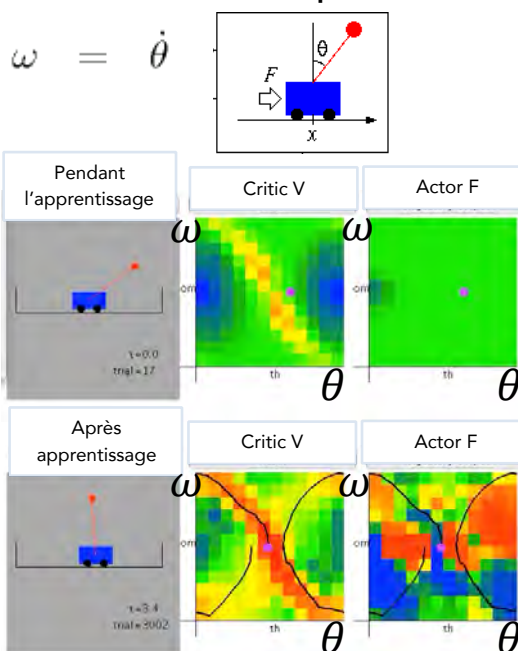
Utilisation du modèle sur nouveau patient

S. Röhrich, et al. "Machine learning: from radiomics to discovery and routine," 2018.

53
JDZucker

Apprentissage par renforcement : programmer le contrôle de l'injection optimale d'insuline

Quelle force F pour maintenir l'équilibre ?



Bothe et al., "The use of reinforcement learning algorithms to meet the challenges of an artificial pancreas," Expert Review of Medical Devices, (10)5,2014.

<https://brain.cc.kogakuin.ac.jp/~kanamaru/NN/CPRL/>

blue: negative, green: 0, red: positive.

<https://towardsdatascience.com/a-review-of-recent-reinforcement-learning-applications-to-healthcare-1f8357600407>

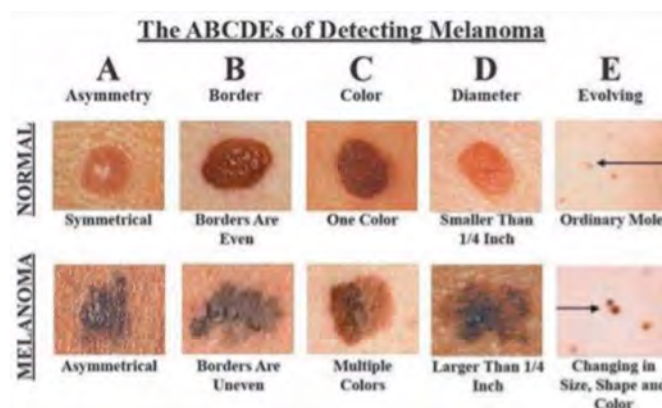
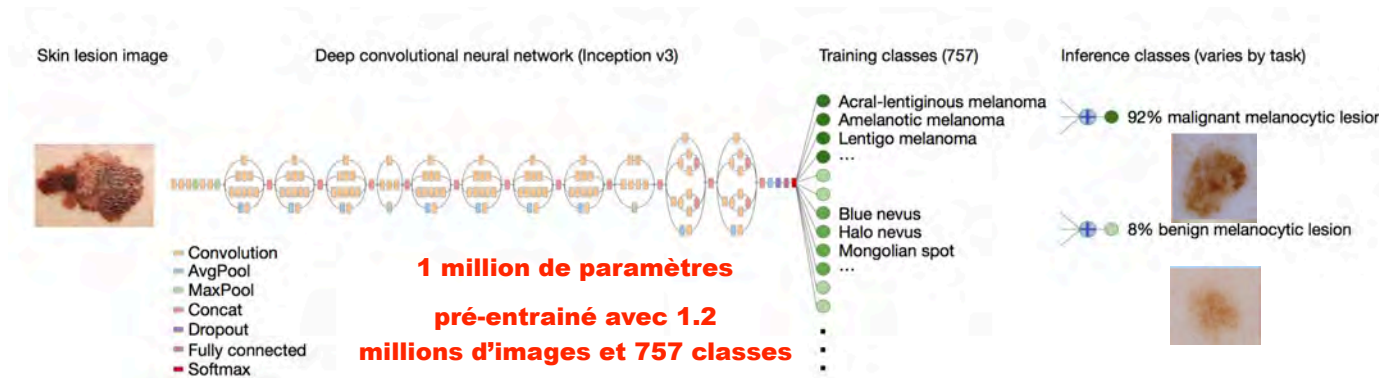
54
JDZucker

3 predictions in **2016** on Machine Learning as a disruptive technology for Medicine in the next few years.

Obermeyer, Z. & Emanuel, E.J., **2016**. Predicting the Future — Big Data, Machine Learning, and Clinical Medicine. New England Journal of Medicine, 375(13), pp.1216–1219.

- First, ML will **dramatically improve prognosis**. We can precisely identify large patient subgroups with mortality rates approaching 100% and others with rates as low as 10%.
prediction → come into use in the next 5 years.
- Second, ML will **displace much of the work of radiologists and anatomical pathologists**. Algorithms will also monitor and interpret streaming physiological data, replacing aspects of anesthesiology and critical care.
prediction → disruptions is within next years, not decades.
- Third, ML will **improve diagnostic accuracy**.
Obstacles: a) gold standard for diagnosis unclear → harder to train algorithms. b) high-value EHR data are often stored in unstructured formats c) models need to be built and validated individually for each diagnosis.
prediction → to develop, over the next decade.

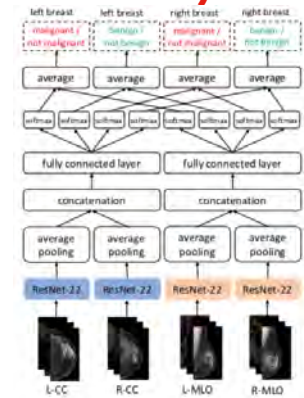
Classification des cancer de la peau du niveau d'un expert dermatologue (Nature,2017)



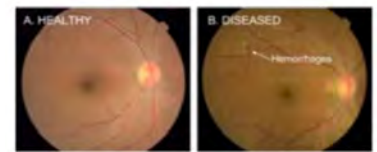
A. Esteva, et al., "Dermatologist-level classification of skin cancer with deep neural networks," Nature, 2017.

Les performances des « IA » dépassent régulièrement celles des radiologues et anatho-pathologistes (2016-2019)

- diagnostiquer les cancers du sein mieux que les radiologues (Nan Wu, et al. 2019). Trained and evaluated on over 200,000 exams (over **1,000,000 images**). AUC of 0.895/



- diagnostiquer la rétinopathie diabétique comme les **ophtalmologistes** (Gulshan, JAMA, 2016.) **128 000** images



Caveat

- To avoid « **disillusionment** » a stronger appreciation of the **technology's capabilities** and limitations is needed.
- Combining** machine-learning software with the **best human clinician "hardware"** will permit delivery of care that outperforms what either can do alone.

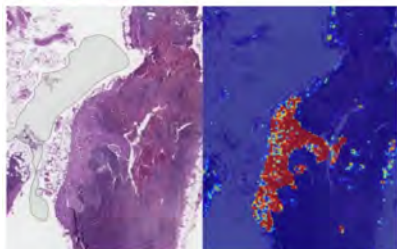
Chen, J. H. & Asch, NEJM 376 (2017).

57
JDZucker

Des systèmes basés sur une collaboration homme-machine peuvent faire mieux que l'IA seule...

INTELLIGENCE ARTIFICIELLE | SCIENCE
Google détecte le cancer du sein métastatique avec une précision de 99%
Google a largement investi dans le développement des applications de santé.

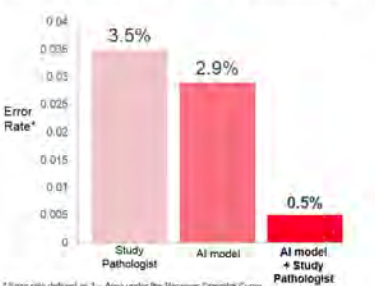
PAR VALENTIN CIMINO | @CIMINO - 15 OCTOBRE 2018



<https://sciencefr2018/10/15/une-ia-by-google-detecte-le-cancer>

[avec une précision de 99%](#)

(AI + Pathologist) > Pathologist



* Error rate defined as 1 - Area under the Receiver Operator Curve in 8 study pathologists, blinded to the ground truth diagnosis, independently scored all evaluation slides

INTELLIGENCE ARTIFICIELLE
Google a créé une intelligence artificielle capable de détecter le cancer du poumon

Les chercheurs de Google et l'hôpital universitaire de Northwestern se sont unis pour élaborer une IA capable de détecter le cancer du poumon

PAR CAMILLE ZAHET - 11 JANVIER 2019



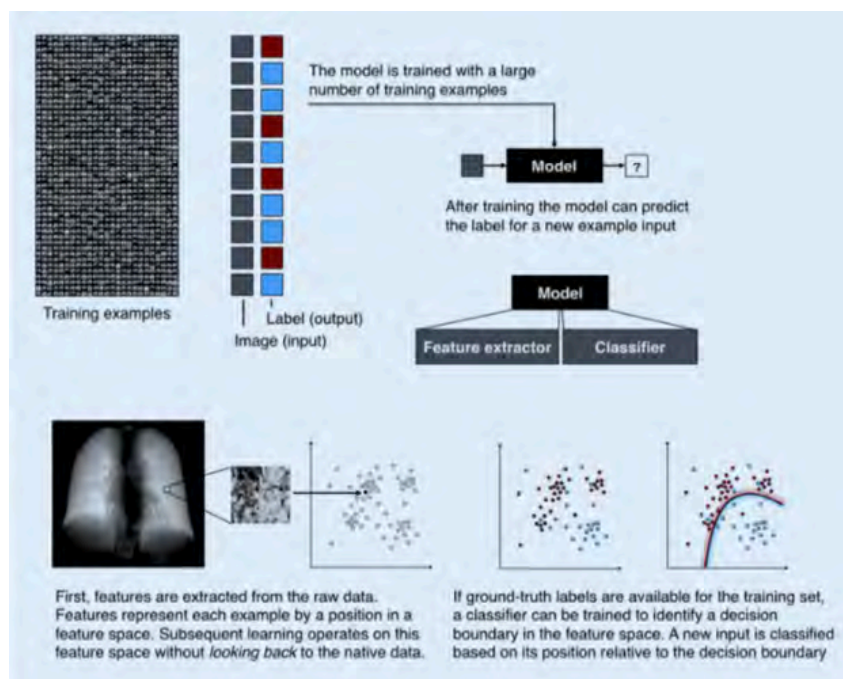
Deep Learning Drops Error Rate for Breast Cancer Diagnoses by 85%

JAMA, vol. 318, no. 22, pp. 2199–2210,
Dec. 2017.

58
JDZucker

The explosion of medical imaging data creates an environment ideal for machine-learning and data-based science

Radiomics, the high-throughput mining of quantitative image features from standard-of-care medical imaging that enables data to be extracted and applied within clinical-decision support systems (CDSS) to improve **diagnostic, prognostic, and predictive accuracy**, is gaining importance in **cancer research**.



S. Röhrich, "Machine learning: from radiomics to discovery and routine," 2018.

59
JDZucker

L'approbation des usages médicaux de l'IA est en marche... forte des performances en prédiction...



Table 1 | Peer-reviewed publications of AI algorithms compared with doctors

Specialty	Images	Publication
Radiology/neurology	CT head, acute neurological events	Titano et al. ²⁷
Pathology	Breast cancer	Ehteshami Bejnordi et al. ²⁸
	Lung cancer (+ driver mutation)	Coudray et al. ²¹
	Brain tumors (+ methylation)	Capper et al. ⁴⁵
Dermatology	Skin cancers	Esteva et al. ²¹
	Melanoma	Haenssle et al. ²⁰⁰
	Skin lesions	Han et al. ⁴⁶
Ophthalmology	Diabetic retinopathy	Gulshan et al. ²⁷
Gastroenterology	Polyps at colonoscopy*	Mori et al. ³⁴
	Polyps at colonoscopy	Wang et al. ³⁵
Cardiology	Echocardiography	Madani et al. ²⁴
	Echocardiography	Zhang et al. ⁴⁶

Table 2 | FDA AI approvals are accelerating

Company	FDA Approval	Indication
Apple	September 2018	Atrial fibrillation detection
Aidoc	August 2018	CT brain bleed diagnosis
iCAD	August 2018	Breast density via mammography
Zebra Medical	July 2018	Coronary calcium scoring
Bay Labs	June 2018	Echocardiogram EF determination
Neural Analytics	May 2018	Device for paramedic stroke diagnosis
IDx	April 2018	Diabetic retinopathy diagnosis
Icometrix	April 2018	MRI brain interpretation
Imagen	March 2018	X-ray wrist fracture diagnosis
Viz.ai	February 2018	CT stroke diagnosis
Arterys	February 2018	Liver and lung cancer (MRI, CT) diagnosis
MaxQ-AI	January 2018	CT brain bleed diagnosis
Alivecor	November 2017	Atrial fibrillation detection via Apple Watch
Arterys	January 2017	MRI heart interpretation

E. J. Topol, "High-performance medicine: the convergence of human and artificial intelligence," Nat Med, pp. 1–13, Jan. 2019.

60
JDZucker

Plan

I. Santé et IA: type de problèmes, type de donnée et types d'apprentissage

II. Des modèles interprétables pour la santé

III. Exemple de la régression régularisée

IV. Médecine de précision et IA

V. Classeur avec rejet

VI. Conclusion

61
JDZucker

RGPD et modèles interprétables : droit et confiance

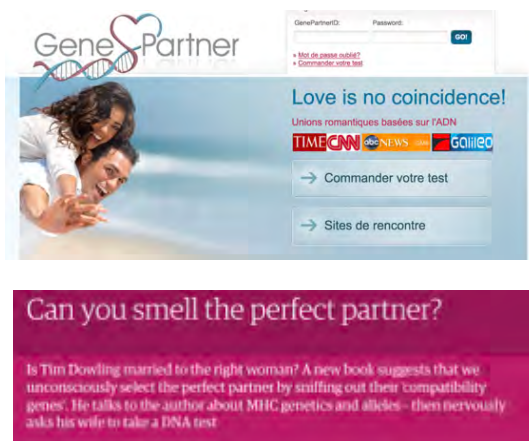
- Règlements de l'UE (Règlement général sur la protection des données (GDPR) en vigueur le 25 mai 2018) sur la prise de décision algorithmique et un "droit d'explication".

Goodman, B. & Flaxman, S. R. European Union Regulations on Algorithmic Decision-Making and a "Right to Explanation". AI magazine, 2017

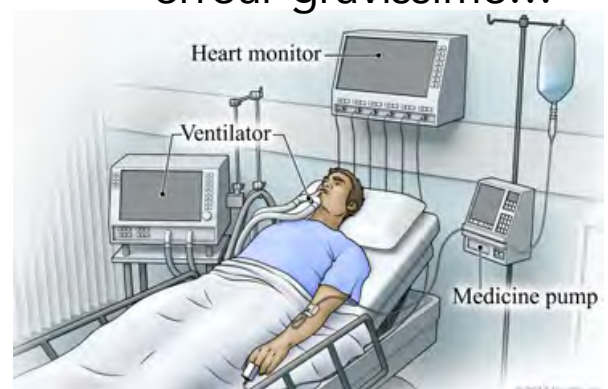
- Une explication de la prédiction est désirée par médecins et patients lorsque un **modèle** doit être validé avant d'être déployé en routine → **confiance**

Vanthienen, et al. Performance of classification models from a user perspective. *Decision Support Systems* 51, 782- 793,(2011).

erreur pas (trop) grave...



erreur gravissime...



A. Vellido, et al., "Machine learning in critical care: state-of-the-art and a sepsis case study," BMEO,2018.

62
JDZucker

Equité/Fairness : l'IA est biaisée par les données



Fair

- Demographic fairness
- Fairness in design
- Fairness in data
- Fairness in algorithms
- Fairness in outcomes

Une étude récente a révélé que certains programmes de reconnaissance faciale classent incorrectement **moins de 1 % des hommes à la peau claire**, mais plus d'**un tiers des femmes à la peau foncée**.

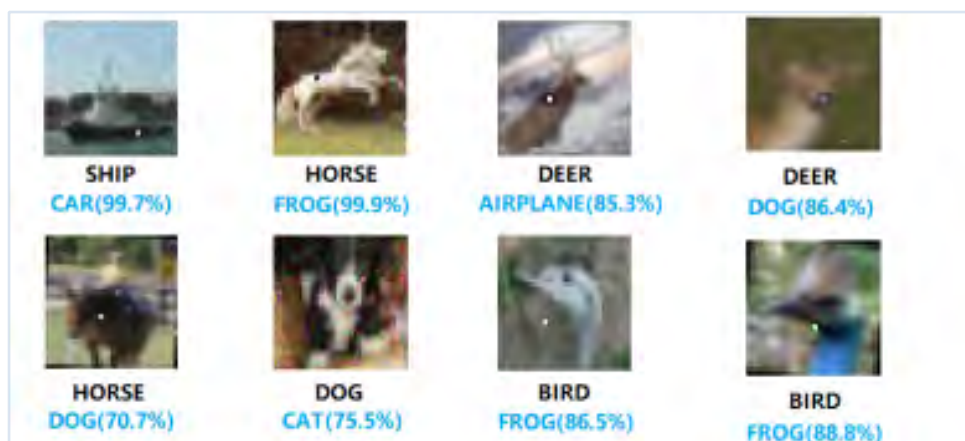
Que se passe-t-il lorsque l'on se fie à de tels algorithmes pour diagnostiquer le mélanome sur une peau claire ou foncée... ?



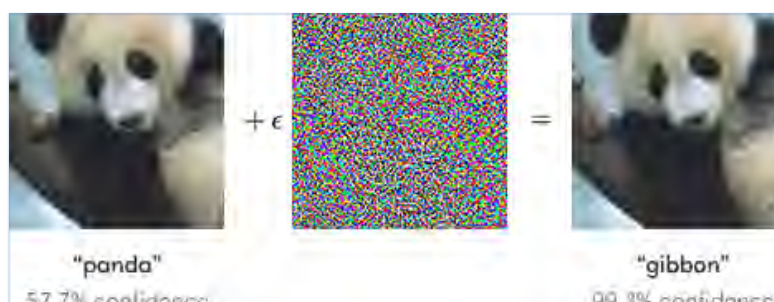
Un programme apprend à partir des données qu'on lui donne et qui peuvent être ... biaisées

63
JDZucker

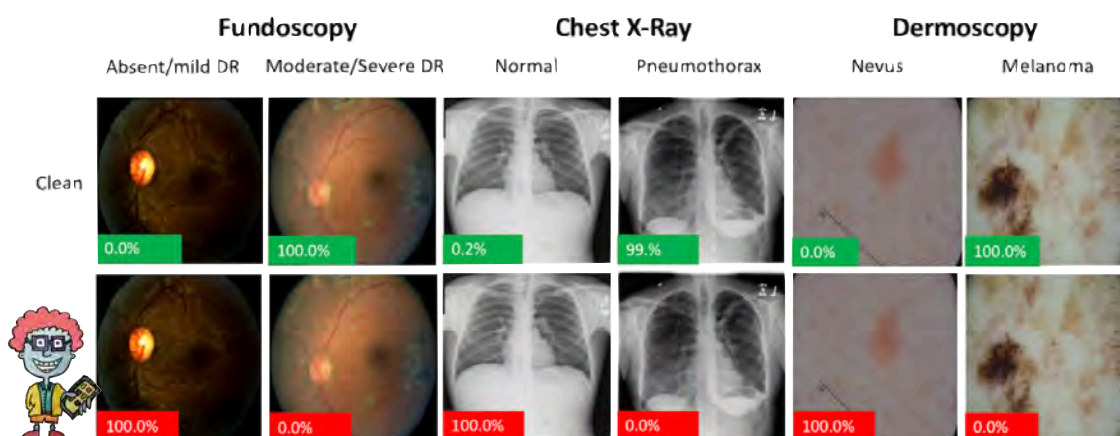
Les « adversarial attacks » sont maintenant connues, mais..



J. Su, et al. "One pixel attack for fooling deep neural networks.," CoRR, 2017.



...se pose le problème de la « responsabilité » des algorithmes ... notamment en cas d'attaques d'images médicales.



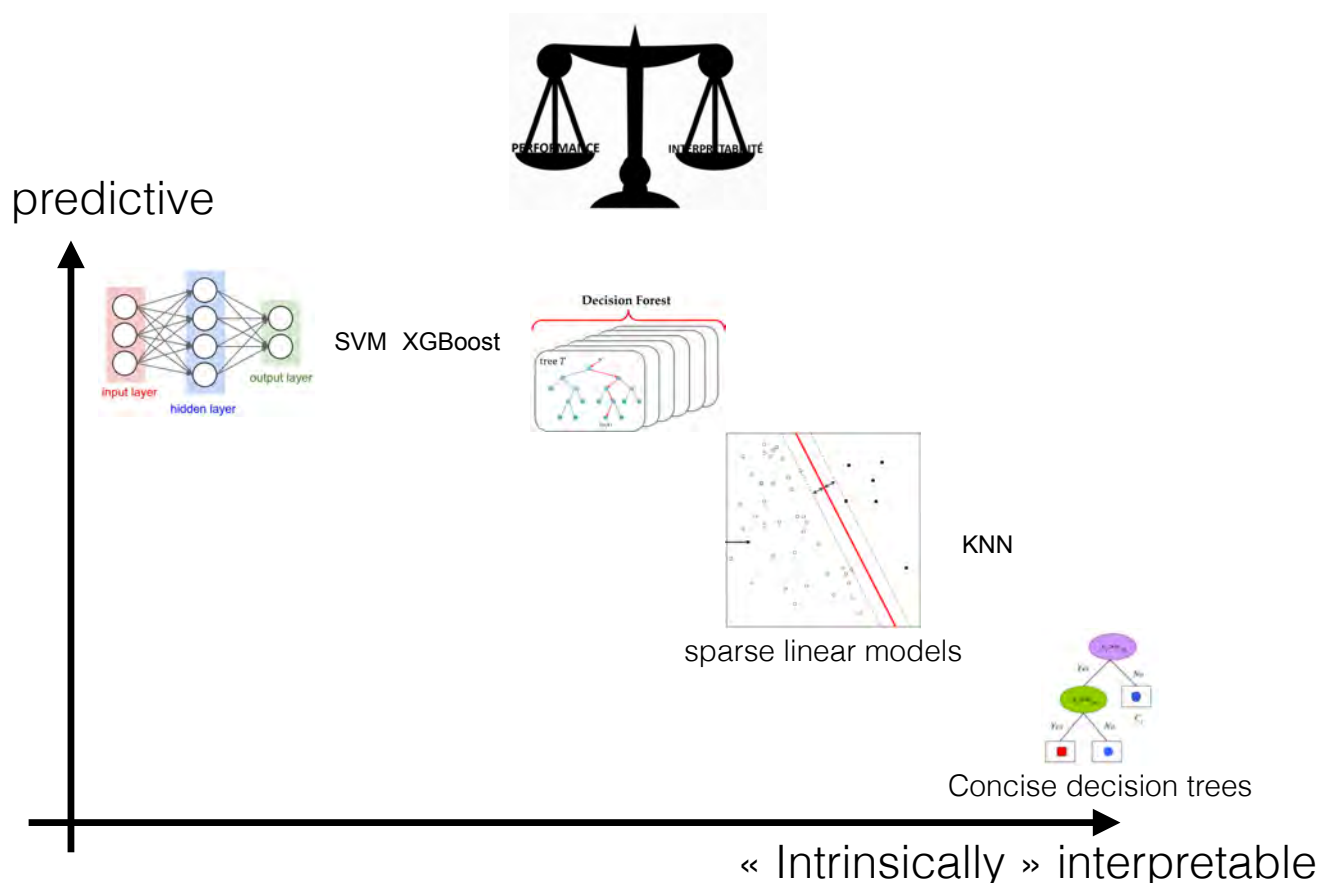
Finlayson, et al., "Adversarial Attacks Against Medical Deep Learning Systems.," arXiv, 2018.

Qui est responsable en cas d'erreur ?



65
JDZucker

Interpretability vs Predictive power

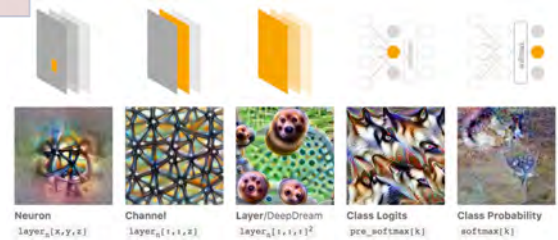


Interpretability in Machine Learning

credits: Yann Chevalere

Type A - Interpreting black-box models

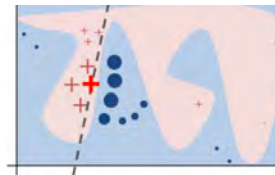
Model distillation (soft DT) [Frosst&Hinton,17]
Looking into the black box



[Olah, distill.pub]

Type B - Interpreting predictions from black-box models

Attribution methods: e.g. LIME



"Why Should I Trust You?"
Explaining the Predictions
of Any Classifier
[Ribeiro et al. '16]

Type C - Learning interpretable models

Intrinsically interpretable models: (discrete) linear model, scoring model, ...



Interepreting NN using « attention maps » indicating the areas of the heat ma that the neural-network model is using to make the prediction for the image.

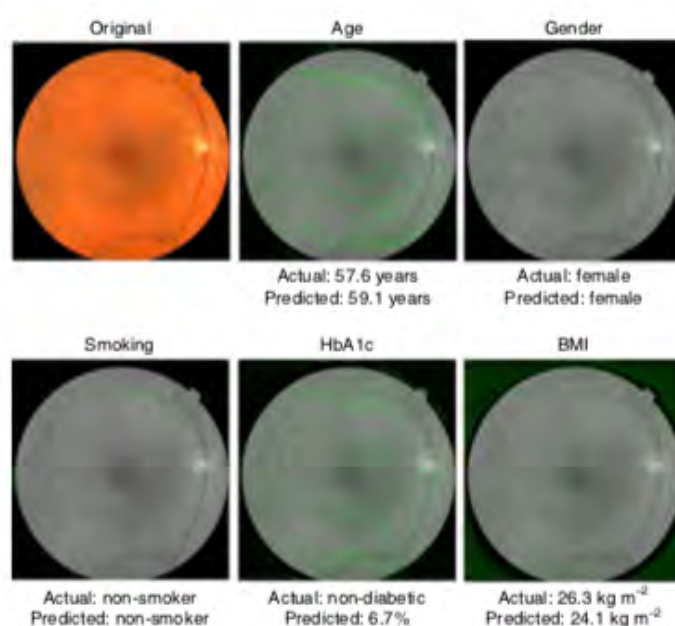


Table 6 | Percentage of the 100 attention heat maps for which doctors agreed that the heat map highlighted the given feature

Risk factor	Vessels (%)	Optic disc (%)	Non-specific features
Age	95	33	38
Gender	71	78	50
Current smoker	91	25	38
HbA1c	78	32	46
SBP	98	14	54
DBP	29	5	97
BMI	1	6	99

Heat maps (n=100) were generated for each risk factor and then presented to three ophthalmologists who were asked to check the features highlighted in each image (n=300 responses for each risk factor). The images were shuffled and presented as a set of 700, and the ophthalmologists were blinded to the output prediction of the heat maps and the ground-truth label. For the variables that were present in both datasets (age and gender), the most commonly highlighted features were identical in both datasets.

Attention maps for a single retinal fundus image. The top left image is a sample retinal image in colour from the UK Biobank dataset. The remaining images show the same retinal image, but in black and white.

Interpretability in Machine Learning

credits: Yann Chevaleyre

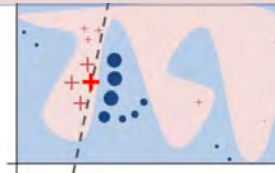
Type A - Interpreting black-box models

Model distillation (soft DT) [Frosst&Hinton,17]
Looking into the black box
Feature Importance (e.g. Random Forest)



Type B - Interpreting predictions from black-box models

Attribution methods: e.g. LIME



“Why Should I Trust You?”
Explaining the Predictions
of Any Classifier
[Ribeiro et al. '16]

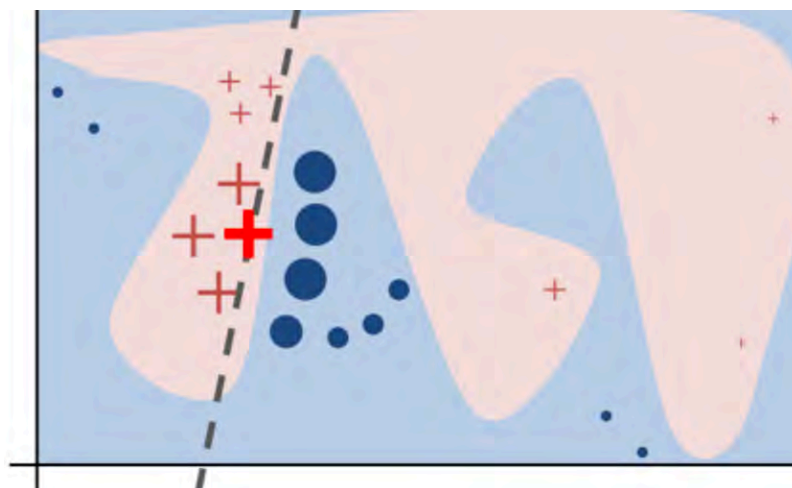
Type C - Learning interpretable models

Intrinsically interpretable models: (discrete) linear model, scoring model, ...



Interpreting predictions: LIME - Local Interpretable Model-Agnostic Explanations

Objectif: convertir les prédictions en un modèle **interprétable** : **séparateur linéaire**.

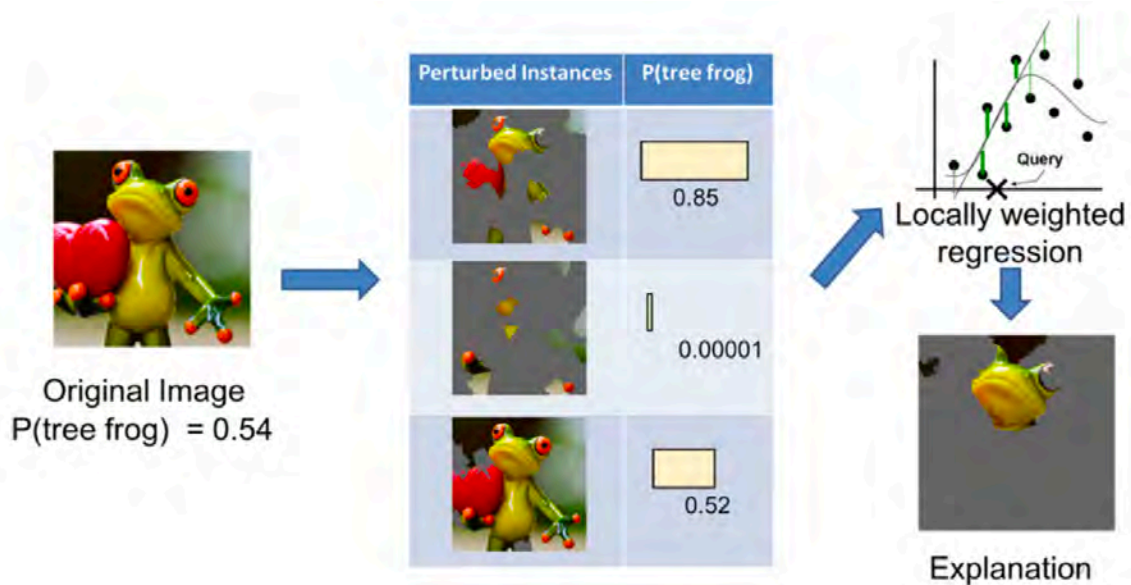


Le graphique représente les zones possibles de prédiction en rouge et bleue, la croix rouge en gras et la prédiction initiale, les axes représentent des variables les autres points (rond bleu ou croix rouge) sont les **prédictions obtenues après modification des valeurs des variables**.

Par exemple, un point situé à droite de la prédiction originale aura été modifiée uniquement sur la variable qui correspond à l'axe des abscisses.

Enfin plus un point possède une grande taille, plus il est “proche” (en distance) du point initial.

LIME



[source:

]

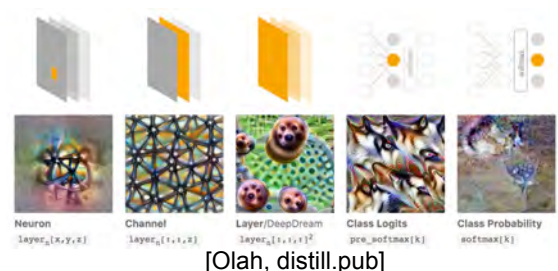
[Ribeiro et al. '16] 71

Interpretability in Machine Learning

credits: Yann Chevalere

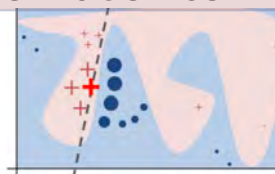
Type A - Interpreting black-box models

Model distillation (soft DT) [Frosst&Hinton,17]
Looking into the black box
Feature Importance (e.g. Random Forest)



Type B - Interpreting predictions from black-box models

Attribution methods: e.g. LIME



"Why Should I Trust You?"
Explaining the Predictions
of Any Classifier
[Ribeiro et al. '16]

Type C - Learning interpretable models

Intrinsically interpretable models: (discrete) linear model, scoring model, ...

Diarem score: was best score (2013); validated on independent cohorts

	Score
Age (years)	
<40	0
40-49	1
50-59	2
≥60	3
HbA_{1c} (%)	
<6.5%	0
6.5-6.9%	2
7.0-8.9%	4
≥9.0%	6
Other diabetes drugs	
No sulfonylureas or insulin-sensitising agent other than metformin	0
Sulfonylureas and insulin-sensitising agent other than metformin	3
Treatment with insulin	
No	0
Yes	10

Total score calculated by adding scores for each of the four variables.

Table 5: Calculation of DiaRem score for prediction of the probability of diabetes remission after Roux-en-Y gastric bypass surgery

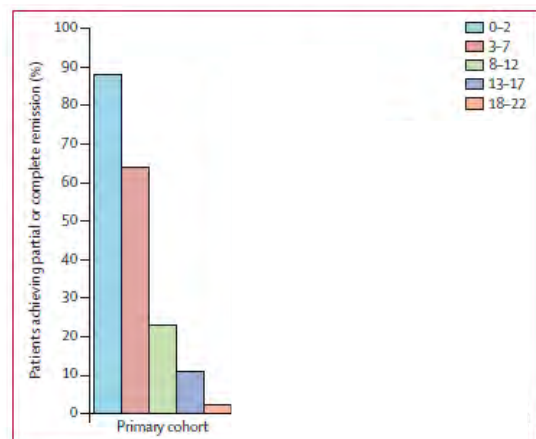


Figure 4: Proportion of patients in each cohort achieving partial or complete remission at 14 months after surgery, by DiaRem score

Table 2 Prediction errors of the different models in percentage

	Train error	Test error	Cross-validation error
Hayes - SL	13.4	31	
Hayes - J48DT	12.6	34.5	
Dixon - LR6	16	40.5	
Dixon - LR7	12	22.6	
Lee - Score	27.3 ^a	16.7	16.7
Still - Score	19.4 ^b	15.5	15.9
LR	7.1		19.9
DT	13.1		17.6
Lasso	14.3		18.9
EN ^c	13.1		17.2

Still et al Lancet endocrinology 2013;

73

74

Automated Score Re-Construction of Diarem

1. Identification of related clinical variables

age | glycated hemoglobin | insuline | other drugs

2. Meaningful thresholds for clinical variables

age | glycated hemoglobin | insuline | other drugs

age	glycated hemoglobin	insuline	other drugs
<40	<6.5	yes	yes
40-49	6.5 - 6.9	no	no
50 - 59	7 - 8.9		
>60	> 9		

3. Optimization of weights for sub-groups of the variables

age | glycated hemoglobin | insuline | other drugs

age	glycated hemoglobin	insuline	other drugs
<40	<6.5	yes	yes
40-49	6.5 - 6.9	no	no
50 - 59	7 - 8.9		
>60	> 9		
0	0	10	3
1	2	0	0
2	4		
3	6		

4. Find an optimal separator between two classes

Classify as Remission if sum of scores < 7

Classify as Non-remission if sum of scores ≥ 7

Stop explaining black box machine learning models for high stakes decisions and use interpretable models instead

Cynthia Rudin

Black box machine learning models are currently being used for high-stakes decision making throughout society, causing problems in healthcare, criminal justice and other domains. Some people hope that creating methods for explaining these black box models will alleviate some of the problems, but trying to explain black box models, rather than creating models that are interpretable in the first place, is likely to perpetuate bad practice and can potentially cause great harm to society. The way forward is to design models that are inherently interpretable. This Perspective clarifies the chasm between explaining black boxes and using inherently interpretable models, outlines several key reasons why explainable black boxes should be avoided in high-stakes decisions, identifies challenges to interpretable machine learning, and provides several example applications where interpretable models could potentially replace black box models in criminal justice, healthcare and computer vision.

NATURE MACHINE INTELLIGENCE | VOL 1 | MAY 2019 | 206–215 | www.nature.com/natmachintell

Prise de décisions à enjeux élevés **santé**, justice pénale, etc.

La voie à suivre est de concevoir des modèles qui sont intrinsèquement interprétables

Nous ne pouvons pas utiliser boîte noire AI pour trouver le lien de causalité, ou de compréhension → l'intelligence humaine et l'IA interprétable. Acceptons les Boîtes noires pour fournir de la valeur, optimiser des résultats et susciter l'inspiration.

ARTIFICIAL INTELLIGENCE

In defense of the black box

Black box algorithms can be useful in science and engineering

By Elizabeth A. Holm

The science fiction writer Douglas Adams imagined the greatest computer ever built, Deep Thought, programmed to answer the deepest question ever asked: the Great Question of Life, the Universe, and Everything. After 7.5 million years of processing, Deep Thought revealed its answer: Forty-two (7). As artificial intelligence (AI) systems enter every sector of human endeavor—including science, engineering, and health—humanity is confronted by the same conundrum that Adams encapsulated so succinctly: What good is knowing the answer when it is unclear why it is the answer? What good is a black box?

In an informal survey of my colleagues in the physical sciences and engineering, the top reason for not using AI methods such as deep learning, voiced by a substantial majority, was that they did not know how to interpret the results. This is an important objection, with implications that range from practical to ethical to legal (2). The goal of scientists and the responsibility of engineers is not just to predict what happens but to understand why it happens. Both an engineer and an AI system may learn to predict whether a bridge will collapse. But only the engineer can explain that decision in terms of physical models that can be communicated to and evaluated by others. Whose bridge would you rather cross?

5 APRIL 2019 • VOL 364 ISSUE 6435
sciencemag.org **SCIENCE**

75
JDZucker



ÉCOLE MATHÉMATIQUE AFRICAINE : Bases
mathématiques de l'intelligence artificielle

Introduction à l'apprentissage artificiel en santé (cours 2/2)

JEAN-DANIEL ZUCKER



JDZucker

Compléments de cours suite aux questions à la fin du premier cours de la veille

1. Qu'est-ce qu'un Gibbon ?
2. Quelle est la différence entre Machine Learning et Statistical Modeling ?
3. Risque empirique, validation croisée et LOO.
4. Courbes ROC
5. Forêts aléatoires

77
JDZucker

1) A gibbon



Panda



Gibbon

Ian J. Goodfellow et al. 2015

Seen as a gibbon



x
"panda"
57.7% confidence

$+ .007 \times$



$\text{sign}(\nabla_x J(\theta, x, y))$
"nematode"
8.2% confidence

$=$



$x + \epsilon \text{sign}(\nabla_x J(\theta, x, y))$
"gibbon"
99.3 % confidence

78
JDZucker

2) What is the difference between Machine Learning and Statistical modeling?

The data modeling culture (Generative modeling)

- Assume a **stochastic model** (5-50 parameters).
- Estimate model parameters.
- Interpret model and make predictions.
- Estimated population: **98%** of statisticians

The algorithmic modeling culture (Predictive modeling)

- Assume relationship between predictor vars and response vars has a **functional form** (10^2 -- 10^6 parameters).
- Search (efficiently) for the best prediction function.
- Make predictions [and interpret model]
- Interpretation / causation - mostly an after-thought.
- Estimated population: **2%** of statisticians (many in other fields).

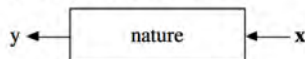
79
JDZucker

Leo Breiman Statistical Modeling / the two cultures

Statistical Science, 2001

1. INTRODUCTION

Statistics starts with data. Think of the data as being generated by a black box in which a vector of input variables $\mathbf{x}$ (independent variables) go in one side, and on the other side the response variables $\mathbf{y}$ come out. Inside the black box, nature functions to associate the predictor variables with the response variables, so the picture is like this:



There are two goals in analyzing the data:

Prediction. To be able to predict what the responses are going to be to future input variables;

Information. To extract some information about how nature is associating the response variables to the input variables.

There are two different approaches toward these goals:

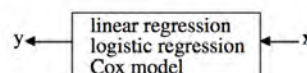
The Data Modeling Culture

The analysis in this culture starts with assuming a stochastic data model for the inside of the black box. For example, a common data model is that data are generated by independent draws from

response variables = $f(\text{predictor variables, random noise, parameters})$

Leo Breiman is Professor, Department of Statistics, University of California, Berkeley, California 94720-4735 (e-mail: leo@stat.berkeley.edu).

The values of the parameters are estimated from the data and the model then used for information and/or prediction. Thus the black box is filled in like this:

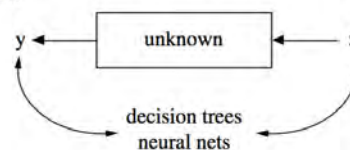


Model validation. Yes—no using goodness-of-fit tests and residual examination.

Estimated culture population. 98% of all statisticians.

The Algorithmic Modeling Culture

The analysis in this culture considers the inside of the box complex and unknown. Their approach is to find a function $f(\mathbf{x})$ —an algorithm that operates on $\mathbf{x}$ to predict the responses $\mathbf{y}$. Their black box looks like this:



Model validation. Measured by predictive accuracy.

Estimated culture population. 2% of statisticians, many in other fields.

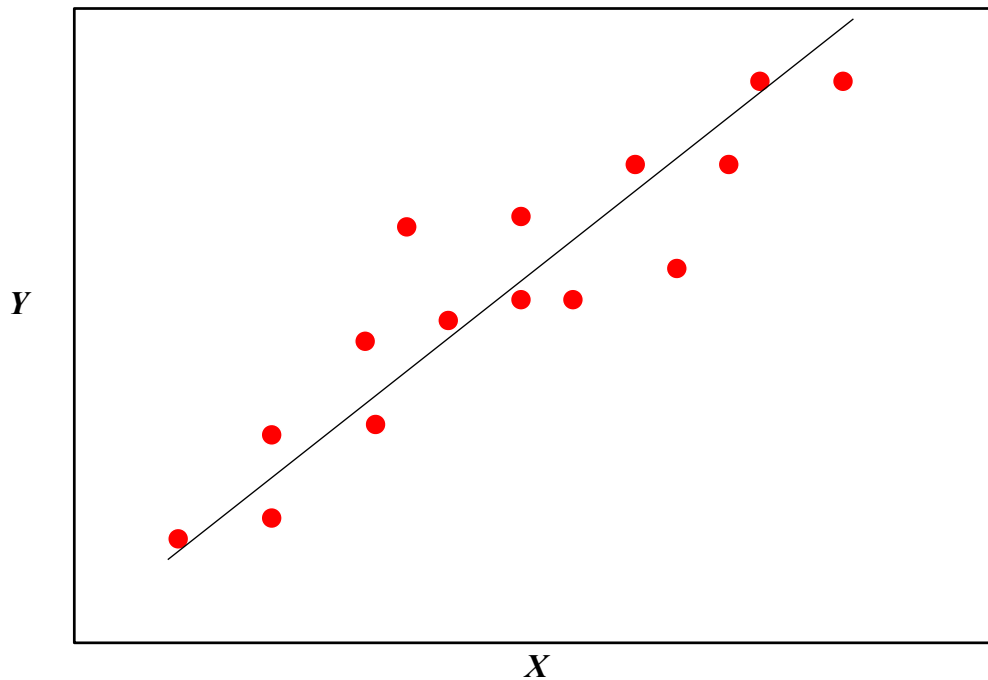
In this paper I will argue that the focus in the statistical community on data models has:

- Led to irrelevant theory and questionable scientific conclusions;

80
JDZucker

The difference between a data model and an algorithm.

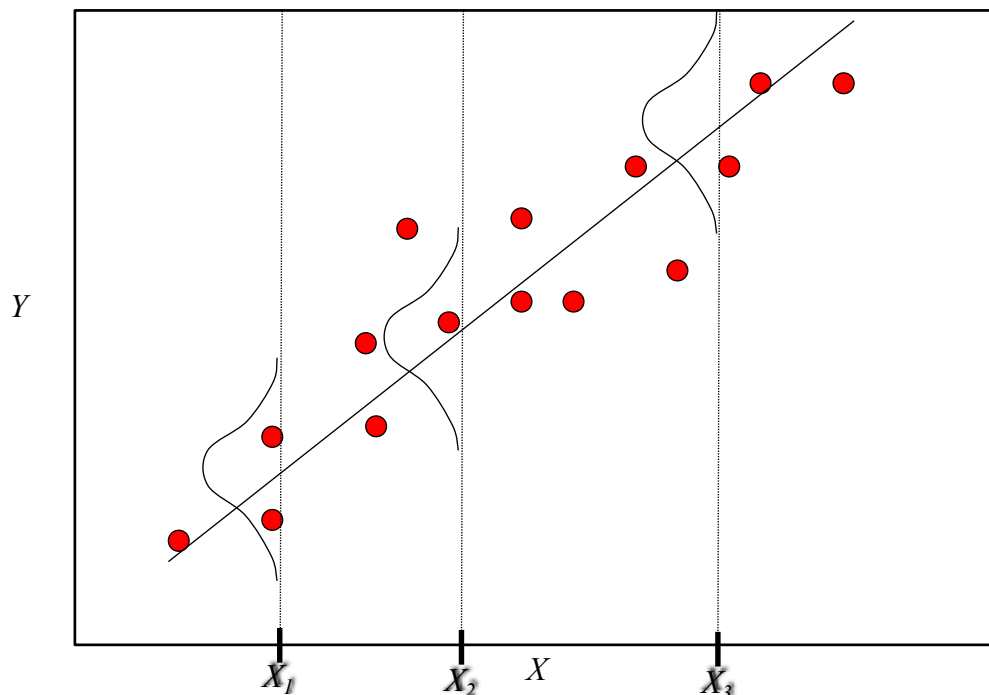
Illustrative example: simple linear regression



81
JDZucker

Least squares fit based on probabilistic assumptions

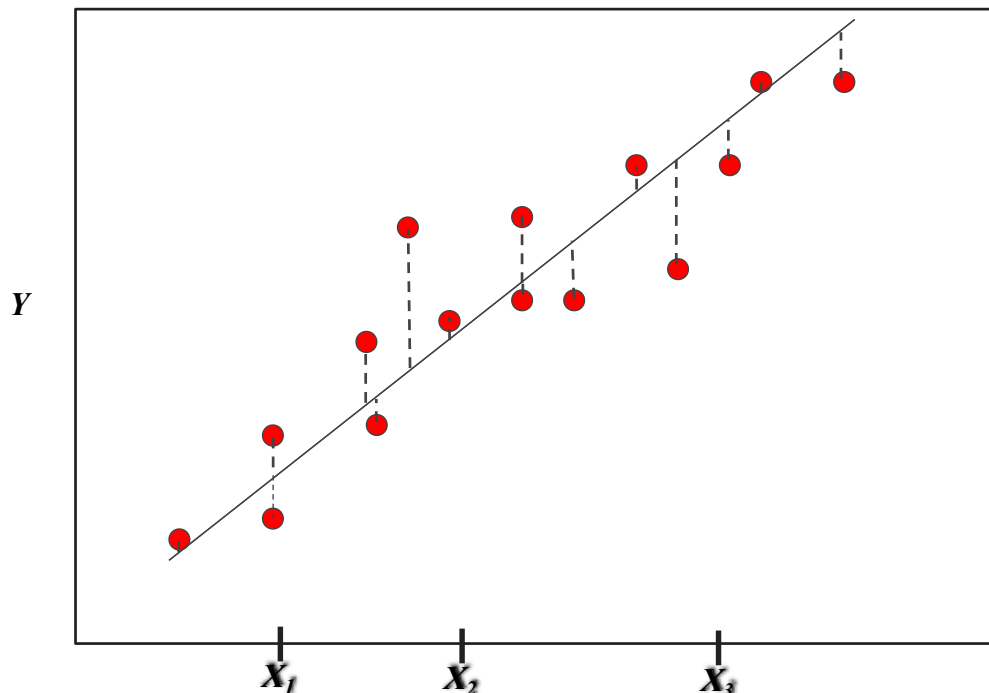
Chaque Y observé est supposé avoir une distribution de probabilité (p. ex. normale) pour un X donné.



82
JDZucker

Least squares fit base on geometrical assumptions

Il n'y a pas de rôle explicite pour la probabilité. Les points (X,Y) sont simplement approximés par une ligne droite. Une considération géométrique (non probabiliste) (minimiser une distance) conduit également à l'ajustement des moindres carrés.



Les 2 approches donnent tous deux des ajustements des moindres carrés. Mais ne poursuivent⁸³ pas même objectif. Pour le modèle de données, la ligne est une estimation de la caractéristique probabiliste ; pour le modèle algorithmique, la ligne est un dispositif approximatif. JDZucker

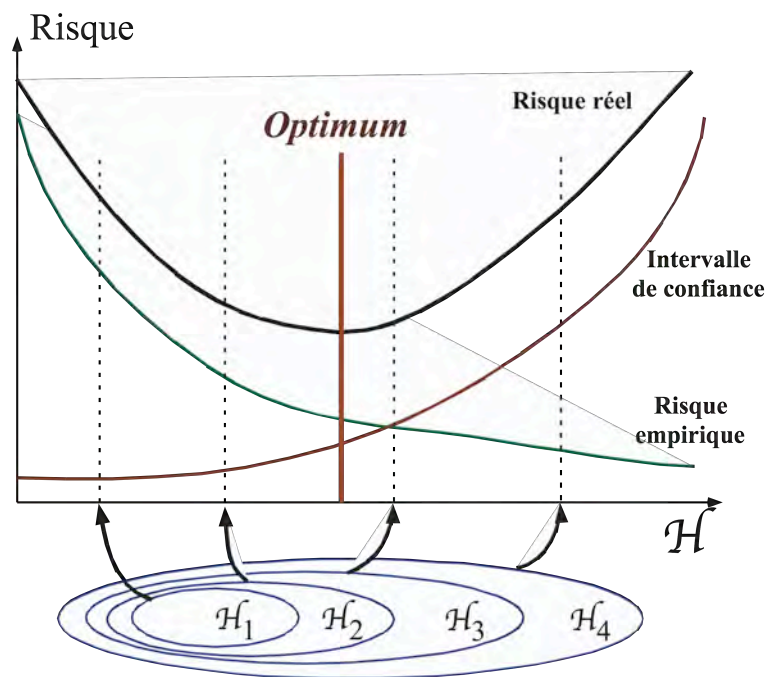
ML & SL de plus en plus similaire

<https://www.r-bloggers.com/whats-the-difference-between-machine-learning-statistics-and-data-mining/>

Machine learning	Statistics
network, graphs	model
weights	parameters
learning	fitting
generalization	test set performance
supervised learning	regression/classification
unsupervised learning	density estimation, clustering
large grant = \$1,000,000	large grant = \$50,000
nice place to have a meeting: Snowbird, Utah, French Alps	nice place to have a meeting: Las Vegas in August

It should seem clear by now that machine learning, statistics, and data mining are all fundamentally similar.

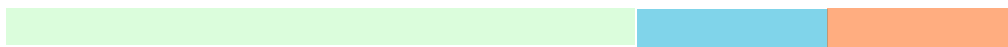
3) Risque Empirique et estimation de l'erreur



85
JDZucker

Risque réel, empirique et validation

Toutes les données disponibles



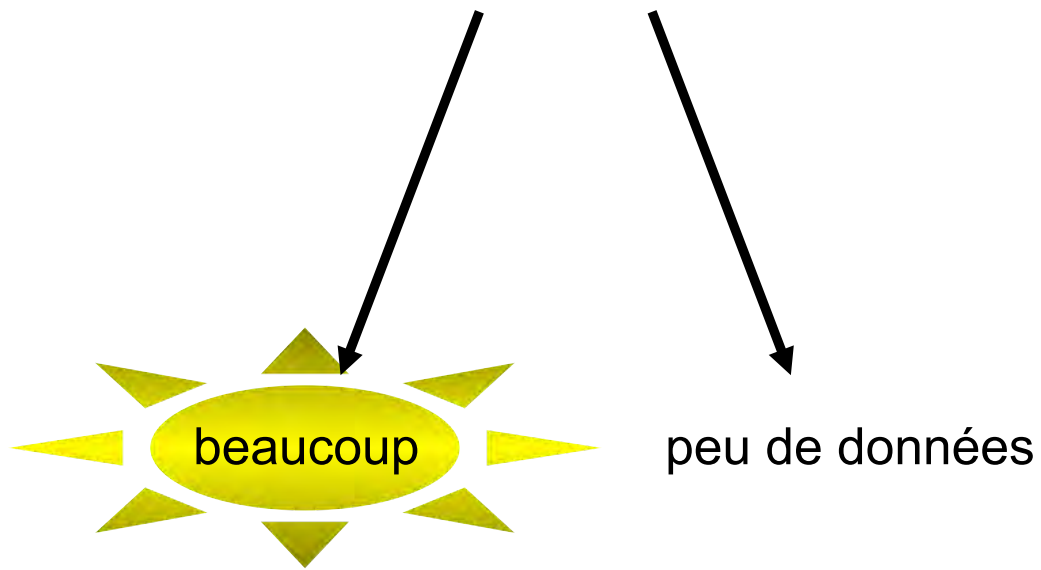
*Ensemble
d'apprentissage*

*Ensemble
de test*

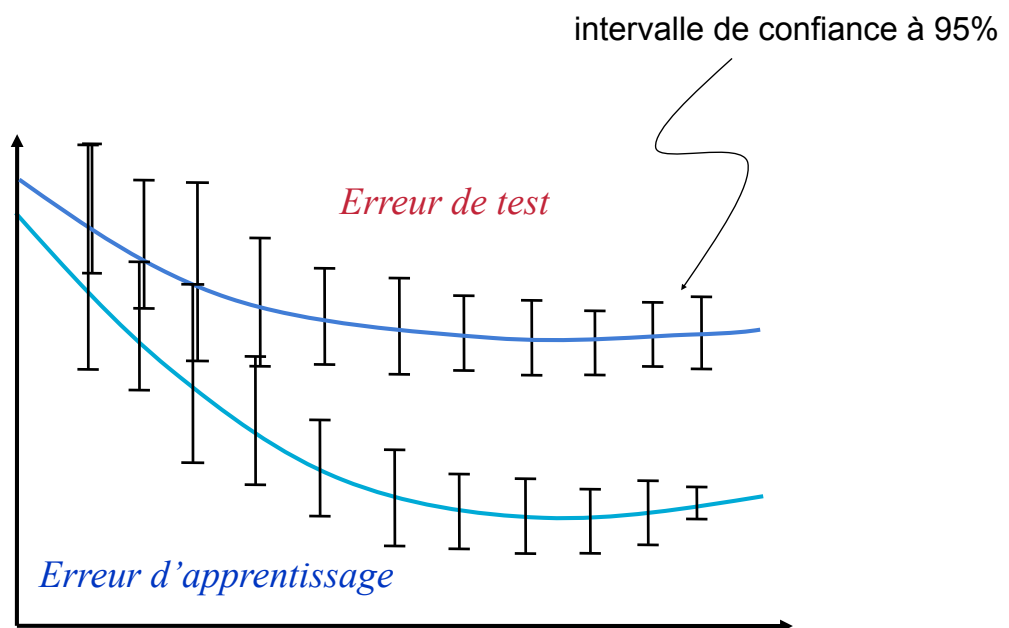
*Ensemble
de validation*

86
JDZucker

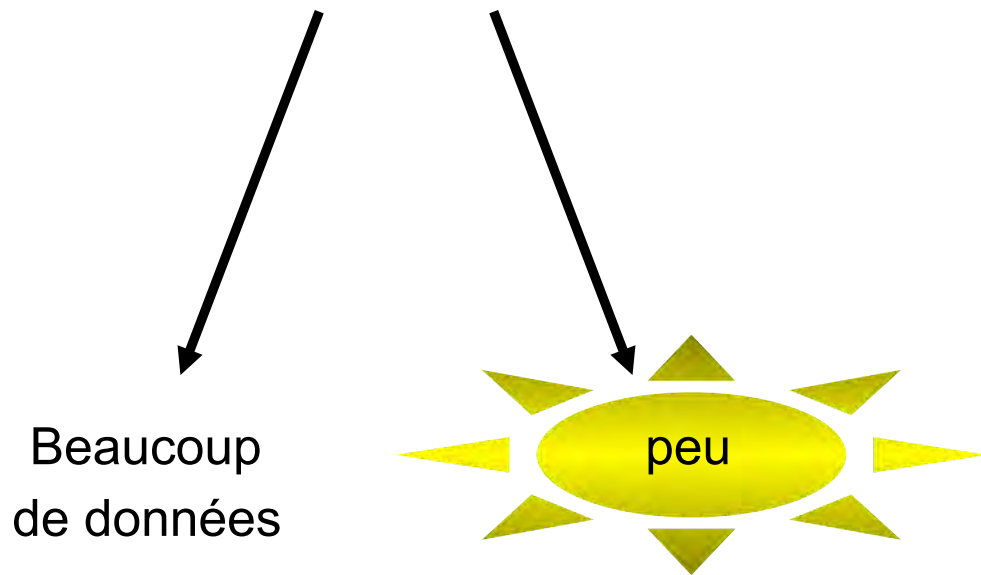
Évaluation des hypothèses produites



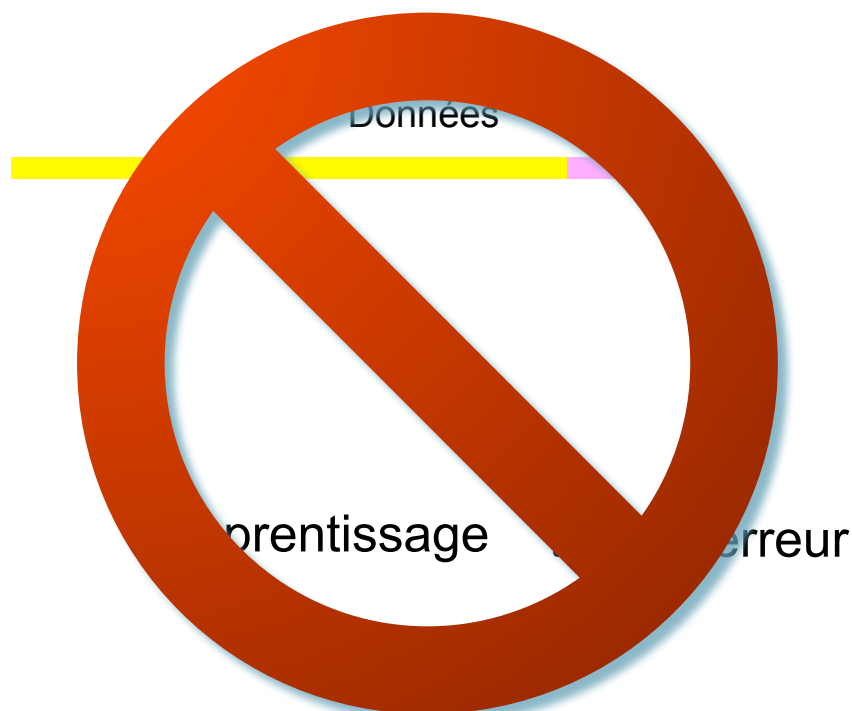
Courbes de performance



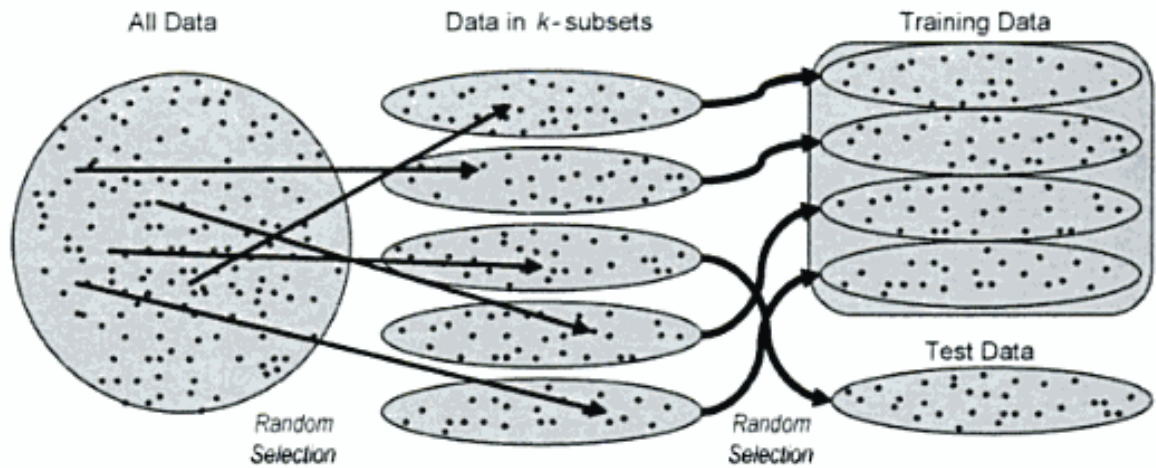
Évaluation des hypothèses produites



Différents ensembles

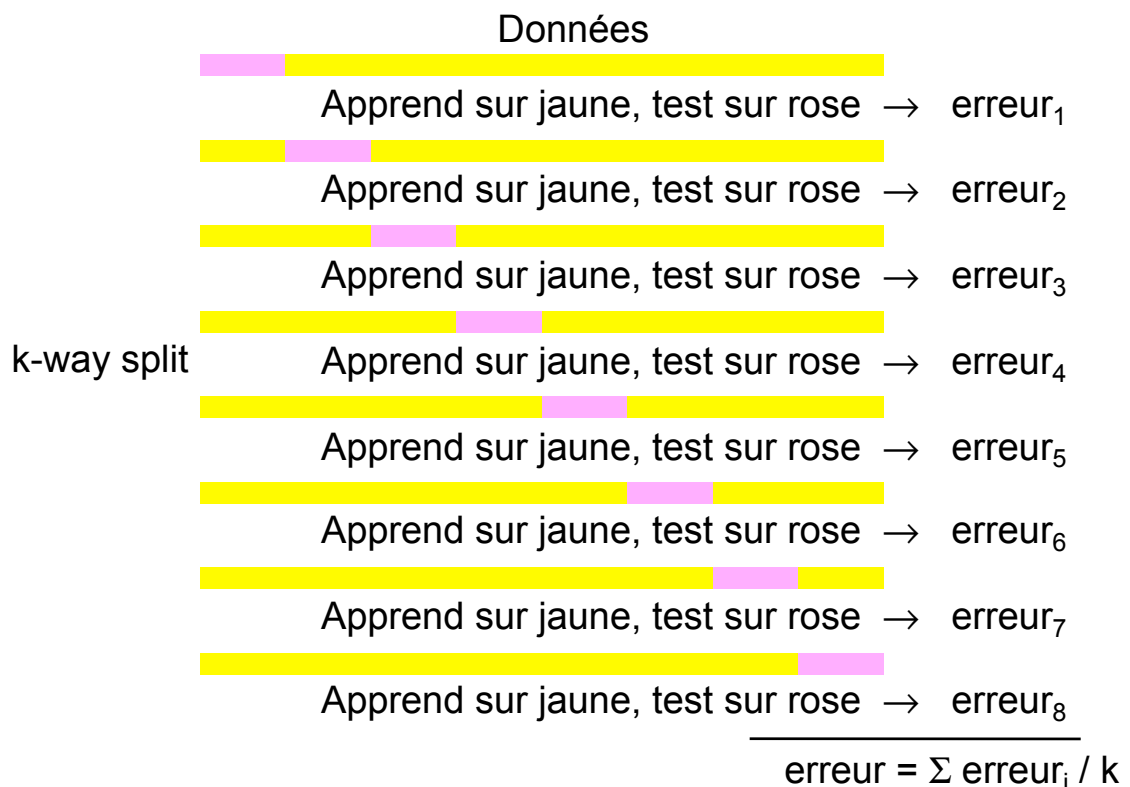


Validation croisée à k plis (k -fold)



91
JDZucker

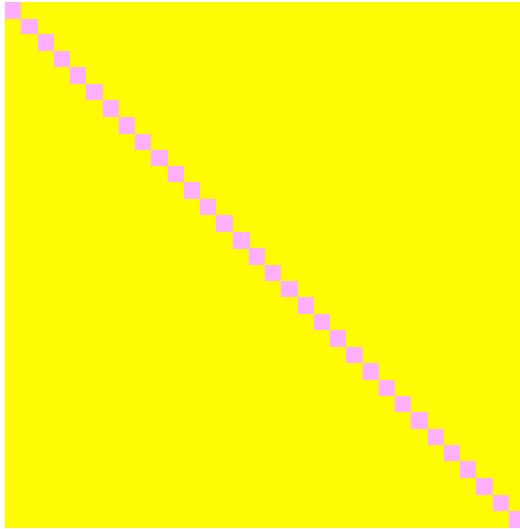
Validation croisée à k plis (k -fold)



92
JDZucker

Procédure "leave-one-out" (LOO)

Données

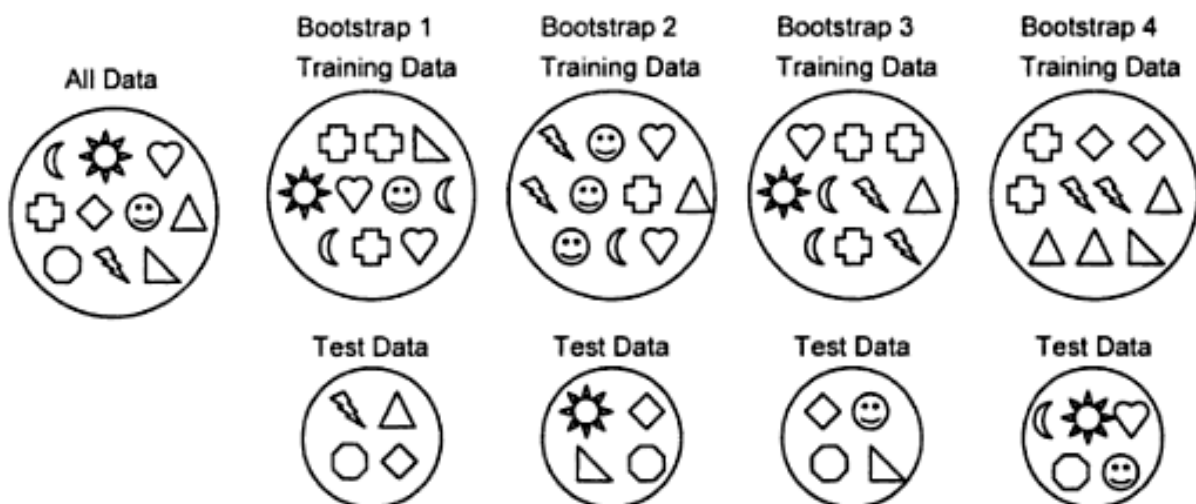


- Faible biais
- Haute variance
- Tend à sous-estimer l'erreur si les données ne sont pas vraiment i.i.d.

[Guyon & Elisseeff, jMLR, 03]

93
JDZucker

Le Bootstrap



Le bootstrap est biaisé

Le bootstrap est biaisé, car certaines observations **sont utilisées à la fois dans l'échantillon pour construire le modèle et dans l'échantillon pour le valider.**

Le bootstrap "hors du sac" (out-of-the-bag) et le bootstrap .632 tentent de corriger ce biais.

94
JDZucker

Le Bootstrap

Out-of-the-bag bootstrap

Le bootstrap "hors du sac" consiste à ne pas utiliser toutes les observations pour valider le modèle mais uniquement celles qui ne figurent pas déjà dans l'échantillon ayant servi à le construire (c'est d'ailleurs ce qu'on faisait pour la validation croisée).

Bootstrap .632

En fait, le bootstrap "out-of-the-bag" est quand-même biaisé, mais dans l'autre sens. Pour tenter de corriger ce biais, on peut faire une moyenne pondérée du bootstrap initial et du bootstrap oob.

$$.368 * (\text{biais estimé par le bootstrap}) + .632 * (\text{biais estimé par le bootstrap oob})$$

(le coefficient .632 s'interprète ainsi : pour n grand, les échantillons de bootstrap contiennent en moyenne 63,2% des observations initiales).

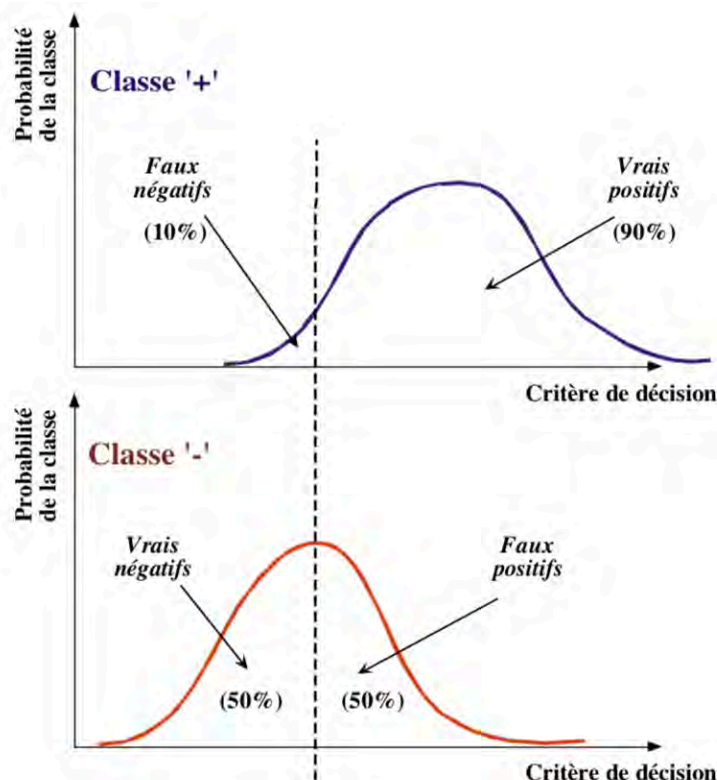
La probabilité de ne pas être choisi après n fois est

$$(1 - 1/n)^n = e^{-1} = 0.368$$

95
JDZucker

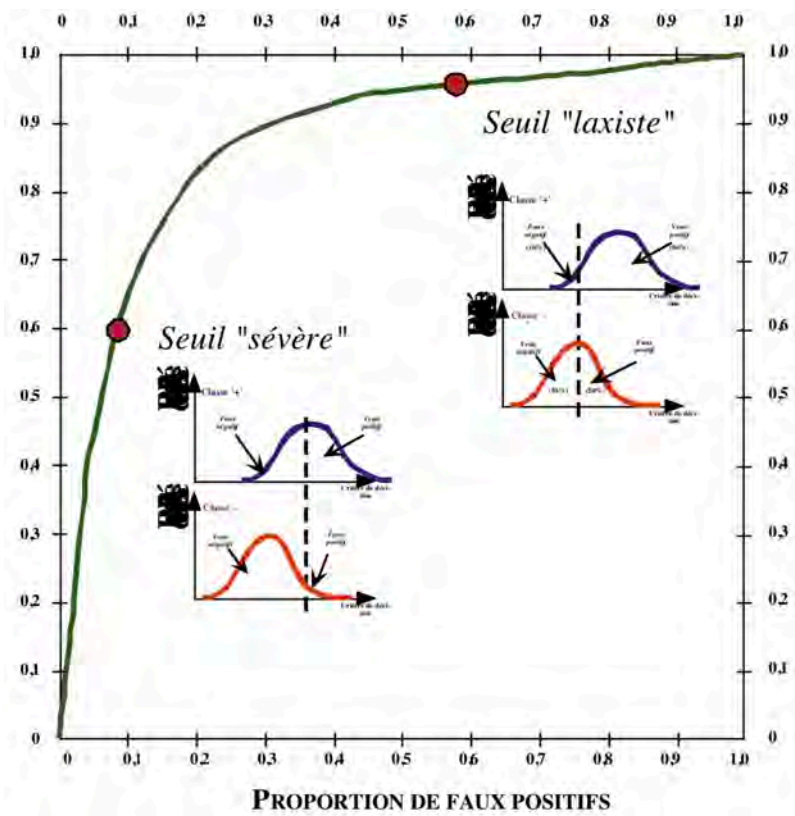
4) Courbe ROC

Receiver Operating Characteristics

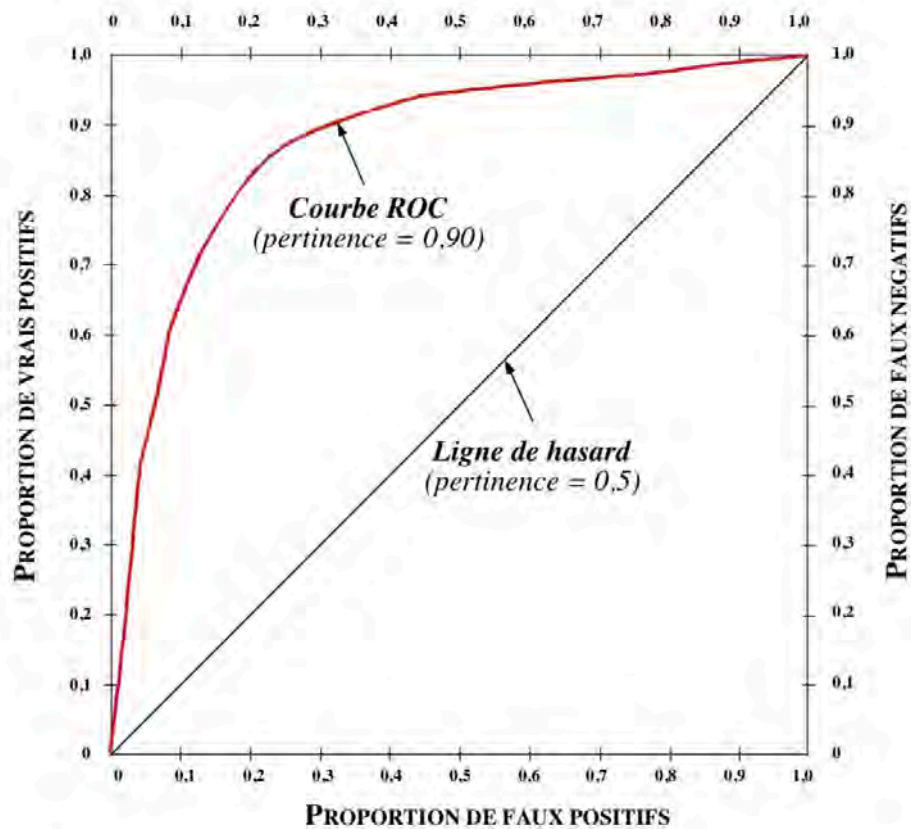


96
JDZucker

proportion
de vrais
positifs



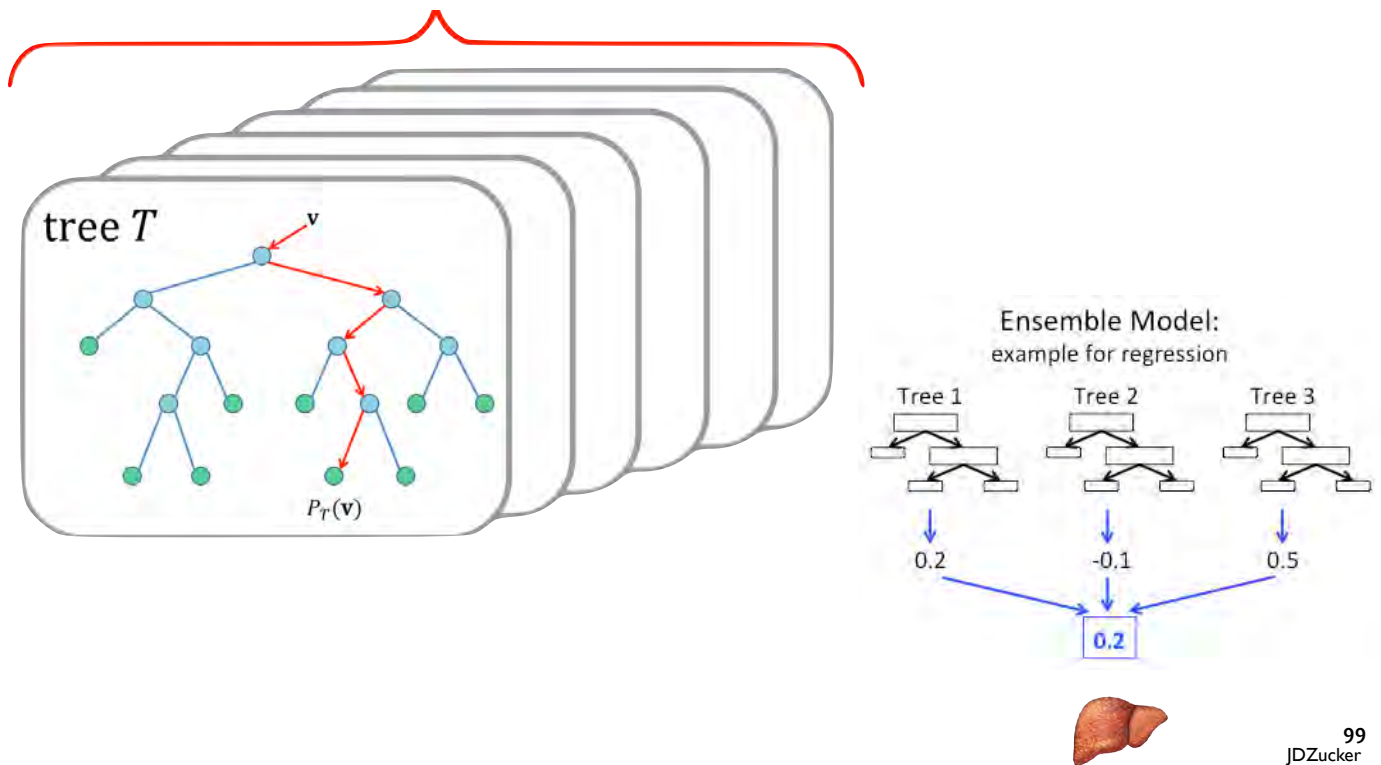
97
JDZucker



98
JDZucker

5) Random Forests: bootstrap on observations and features ($m < p$) for each tree

Random Forest



99
JDZucker

6) Can I export Orange workflow as a Python script?

"Malheureusement, non. Nous en avons débattu longuement et durement. Une fonctionnalité (limitée) de ce type serait probablement possible, mais il s'agirait d'un très gros projet à mener à bien et qui coûterait plus cher que les avantages escomptés. »

<https://orange.biolab.si/faq/>

Je veux utiliser Orange sur de gros volumes de données.

Orange fournit le widget SQL, qui échantillonne les instances de données et permet l'analyse exploratoire des données sur de grands ensembles de données. Vous avez besoin du module **psycopg2** installé pour pouvoir voir le widget.

Installez-le directement en Python avec

pip install psycopg2

100
JDZucker

Plan

I. Santé et IA: type de problèmes, type de donnée et types d'apprentissage

II. Des modèles interprétables pour la santé

III. Exemple de la régression régularisée

IV. Médecine de précision et IA

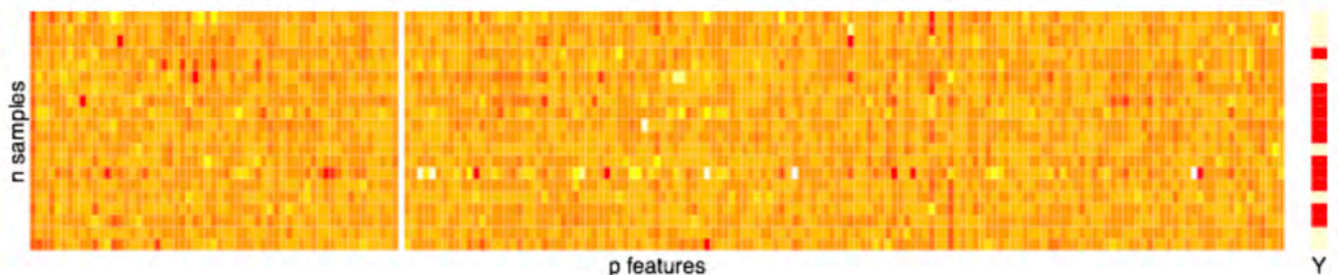
V. Classeur avec rejet

VI. Conclusion

101
JDZucker

Problème de régression en haute dimension

Gene expression



- n samples (patients), p features (genes)
- $X \in \mathbb{R}^{n \times p}$ gene expression profile of each patient
- $Y \in \mathcal{Y}^n$ survival information of each patient
- Fit a linear model for a sample $x \in \mathbb{R}^p$:

$$f(x) = \beta^\top x = \sum_{i=1}^p \beta_i x_i$$

- Standard methods (least squares or logistic regression) **won't work** because $n < p$

crédit: JP. Vert

102
JDZucker

I. La régression logistique simple

Variable dépendante : $Y = 0 / 1$

Variable indépendante : X

Objectif : Modéliser

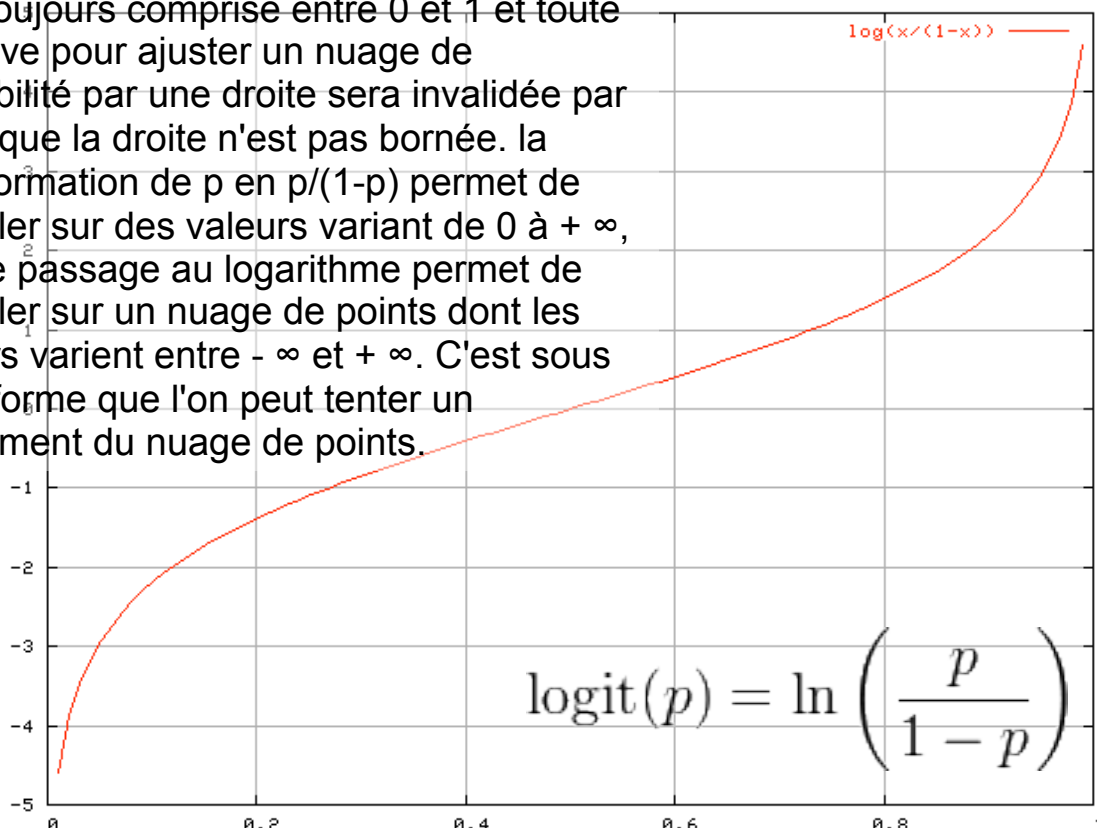
$$\pi(x) = \text{Prob}(Y = 1/X = x)$$

- Le modèle linéaire $\pi(x) = \beta_0 + \beta_1 x$ convient mal lorsque X est continue.
- Le modèle logistique est plus naturel.

103
JDZucker

Logit(x)

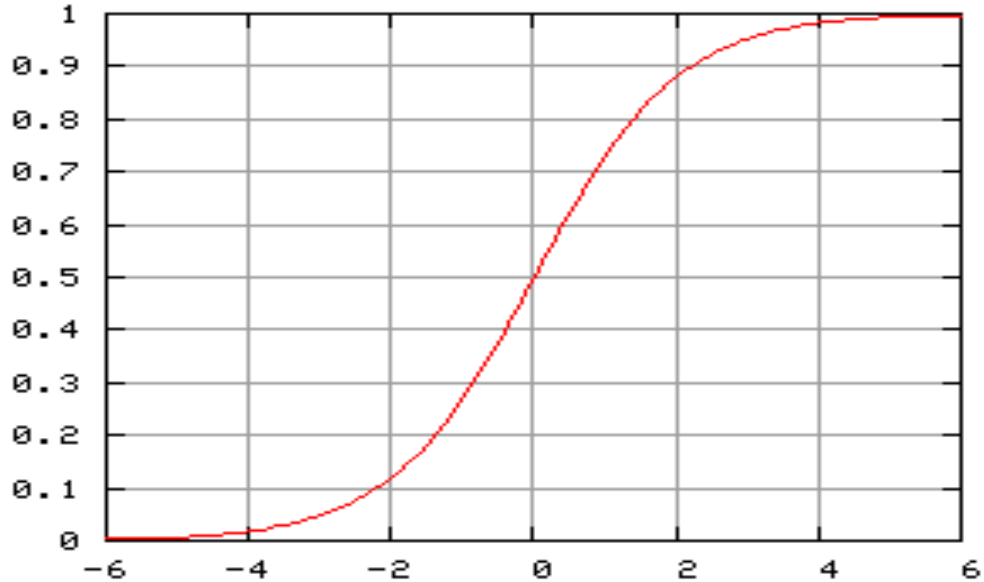
Si p est une probabilité, cette probabilité sera toujours comprise entre 0 et 1 et toute tentative pour ajuster un nuage de probabilité par une droite sera invalidée par le fait que la droite n'est pas bornée. la transformation de p en $p/(1-p)$ permet de travailler sur des valeurs variant de 0 à $+\infty$, puis le passage au logarithme permet de travailler sur un nuage de points dont les valeurs varient entre $-\infty$ et $+\infty$. C'est sous cette forme que l'on peut tenter un ajustement du nuage de points.



104
JDZucker

Odd ratio

$$\text{LOGIT}(p) = \ln\left(\frac{p}{(1-p)}\right) = z \Leftrightarrow p = \frac{\exp(z)}{1 + \exp(z)}$$



105
JDZucker

Estimation des paramètres

Les données

X	Y
x_1	y_1
$\vdots$	$\vdots$
x_i	y_i
$\vdots$	$\vdots$
x_n	y_n

$y_i = 1$ si caractère présent,
0 sinon

Le modèle

$$\begin{aligned}\pi(x_i) &= P(Y = 1 / X = x_i) \\ &= \frac{e^{\beta_0 + \beta_1 x_i}}{1 + e^{\beta_0 + \beta_1 x_i}}\end{aligned}$$

106
JDZucker

Le modèle logistique

Logistic Regression:

$$\ln \left(\frac{P(Y)}{1-P(Y)} \right) = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_K X_K$$

Linear Regression:

$$Y = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_K X_K + \varepsilon$$

Relationship between Odds & Probability

$$\text{Odds}(\text{event}) = \frac{\text{Probability}(\text{event})}{1 - \text{Probability}(\text{event})}$$

$$\text{Probability}(\text{event}) = \frac{\text{Odds}(\text{event})}{1 + \text{Odds}(\text{event})}$$

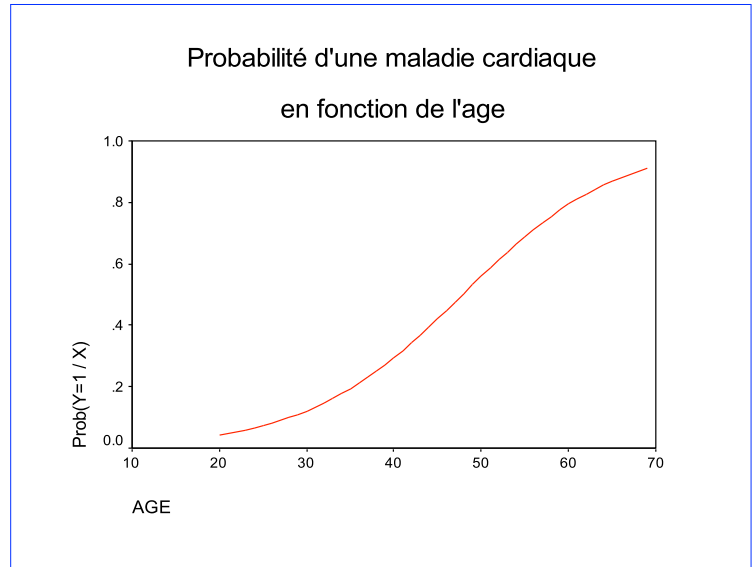
Le modèle logistique

$$\pi(x) = \frac{e^{\beta_0 + \beta_1 x}}{1 + e^{\beta_0 + \beta_1 x}}$$

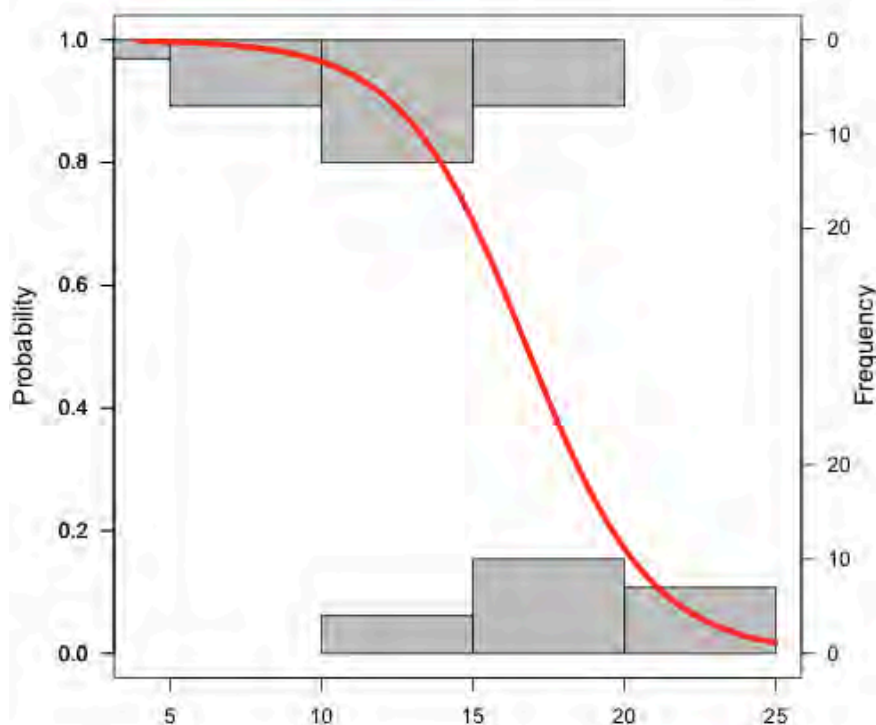
ou

$$\text{Log}\left(\frac{\pi(x)}{1 - \pi(x)}\right) = \beta_0 + \beta_1 x$$

Fonction de lien : Logit



Exemple CARS : voiture Safe (sûre)



Regression Logistique

```
myData = cars
myData$safe = ifelse(dist < mean(dist),1,0)
attach(myData)

quartz(title="speed vs. safe")
# creates a quartz window with title

plot(speed,safe,xlab="speed",ylab="Proba of safe")

g=glm(safe~speed,family=binomial,mydata)
# run a logistic regression model
curve(exp(8-0.5*x)/(1+exp(8-0.5*x)),add = TRUE, col = "blue")
curve(predict(g,data.frame(speed=x),type="resp"),add=TRUE)
# draws a curve based on prediction from logistic regression model

points(speed,fitted(g),pch=20) # optional: you could skip this draws an
invisible set of points of body size survival based on a 'fit' to glm model.
pch= changes type of dots.

### or ####
library(popbio)
logi.hist.plot(speed,safe,boxp=FALSE,type="hist",col="gray")
```

111
JDZucker

Regression Logistique Prédiction

```
myData = cars
myData$safe = ifelse(dist < mean(dist),1,0)
attach(myData)
Newdata = data.frame(speed = c(13.5))

quartz(title="speed vs. safe")
plot(speed,safe,xlab="speed",ylab="Proba of safe")

g=glm(safe~speed,family=binomial,myData)
# run a logistic regression model
print(predict.glm(g,Newdata,type="response"))
```

112
JDZucker

Régression régularisée

Ridge - Lasso - Elasticnet

D'après le cours de Ricco
Rakotomalala

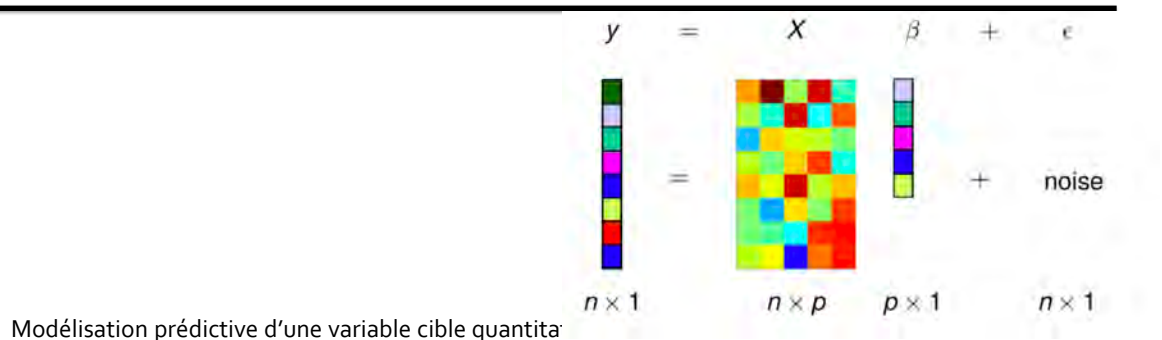
Université Lumière Lyon 2

Ricco Rakotomalala

Tutoriels Tanagra - <http://tutoriels-data-mining.blogspot.fr/>

113
JDZucker

1. Régression Linéaire Multiple



$$y = a_0 + a_1x_1 + \dots + a_px_p + \varepsilon$$

Résume l'information de y que
l'on n'arrive pas à capter avec les
($x_1, \dots, x_p$) variables prédictives.

Disposant d'un échantillon de taille n ,
nous cherchons les paramètres estimés
 $\hat{a} = (\hat{a}_0, \hat{a}_1, \dots, \hat{a}_p)$ qui minimise l'erreur
quadratique moyenne (MSE : *mean squared error*).

$$\min_{a_0, a_1, \dots, a_p} \frac{1}{n} \sum_{i=1}^n \left(y_i - \left(a_0 + \sum_{j=1}^p a_j x_{ij} \right) \right)^2$$

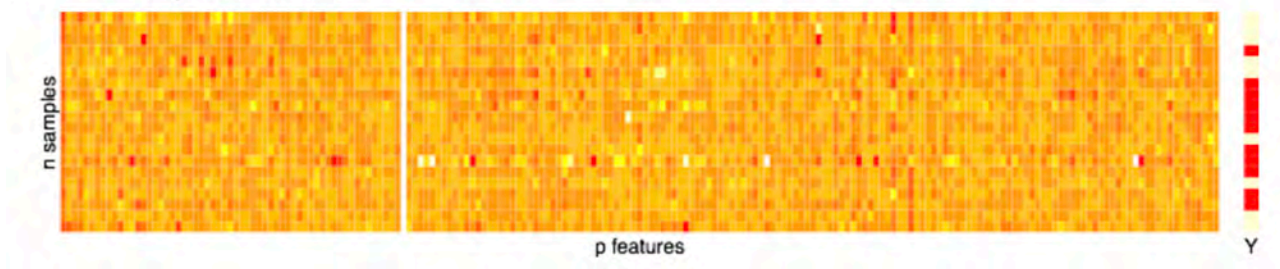
Estimateur des moindres carrés
ordinaires (écriture matricielle)

$$\hat{a}_{MCO} = (X'X)^{-1}X'Y$$

114
JDZucker

Exemple expression génique

Gene expression



Questions: Sélectionner les gènes significatifs, trouver un algorithme avec une bonne prédiction de la réponse ?

Linear model? Difficult to imagine : $p > n$!

$$\hat{a}_{MCO} = (X^T X)^{-1} X^T Y$$

- $X^T X$ is an $p \times p$ matrix, but its rank is lower than n . If $n \ll p$, then

$$rk(X^T X) \leq n \ll p.$$

- Consequence : the Gram matrix $X^T X$ is not invertible and even very ill-conditioned (most of the eigenvalues are 0 !)

Le modèle linéaire échoue

115
JDZucker

Regression Ridge

Ajouter une contrainte sur les coefficients lors de la modélisation pour maîtriser l'amplitude de leurs valeurs (« pour éviter qu'elles partent dans tous les sens »)

$$\min_{\beta_1, \dots, \beta_p} \sum_{i=1}^n \left(y_i - \sum_{j=1}^p \beta_j z_{ij} \right)^2$$

Sous la contrainte $\sum_{j=1}^p \beta_j^2 \leq \tau$ $\left\{ \begin{array}{l} \text{Norme L2} \quad \|\beta\|_2^2 = \sum_{j=1}^p \beta_j^2 \\ \tau (\tau \geq 0) \text{ est un paramètre à fixer} \end{array} \right.$

Remarques :

- On parle de « **shrinkage** » (rétrécissement) : on rétrécit les plages de valeurs que peuvent prendre les paramètres estimés.
- Les variables x_j doivent être centrées et réduites (z_j) pour éviter que les variables à forte variance aient trop d'influence
- La variable cible y doit être centrée pour évacuer la constante de la régression (qui ne doit pas être pénalisée), la cible y peut être éventuellement réduite aussi : nous travaillerons alors sur les paramètres β_j
- $(\tau \rightarrow 0) \Rightarrow \beta_j \rightarrow 0$: les variances des coefficients estimés sont nulles
- $(\tau \rightarrow +\infty) \Rightarrow \beta_{\text{Ridge}} = \beta_{\text{MCO}}$

116
JDZucker

Regression Ridge: Fonction de pénalité

La régression ridge peut être écrite, de manière totalement équivalente : forme lagrangienne

Somme des carrés des résidus pénalisés :

$$\min_{\beta_1, \dots, \beta_p} \sum_{i=1}^n \left(y_i - \sum_{j=1}^p \beta_j z_{ij} \right)^2 + \lambda \sum_{j=1}^p \beta_j^2$$

λ ($\lambda \geq 0$) est un paramètre (coefficient de pénalité) qui permet de contrôler l'impact de la pénalité : à fixer

Fonction de pénalité

Remarques :

- Il y a une équivalence directe entre τ et λ
- $\tau \rightarrow 0 \Leftrightarrow \lambda \rightarrow +\infty : \beta_j \rightarrow 0$ (tous), variances des coefficients nulles
- $\tau \rightarrow +\infty \Leftrightarrow \lambda \rightarrow 0 : \beta_{\text{Ridge}} = \beta_{\text{MCO}}$

Question: comment choisir lambda

117
JDZucker

Utiliser la performance prédictive pour déterminer lambda

S'appuyer sur un critère de performance pure

Principe :

- Fixer une plage de valeurs de λ
- Construire le modèle sur un échantillon d'apprentissage Ω_{App}
- L'évaluer sur un échantillon test Ω_{Test}
- Choisir λ qui minimise un critère d'erreur en test

Dérivé du PRESS (predicted residual error sum of squares) :

$$ERR = \frac{1}{n_{\text{Test}}} \sum_{i \in \Omega_{\text{Test}}} (y_i - \hat{y}_i)^2$$

Les individus Ω_{Test} ne doivent pas faire partie de ceux qui ont permis de construire le modèle. Sinon, c'est $\lambda = 0$ sera la solution systématique.

Validation croisée :

Si la base est de taille restreinte, la partition apprentissage-test n'est pas judicieuse. Mieux vaut passer par la *K-fold* validation croisée (cf.

118
JDZucker

Contrainte sur la norme L1 des coefficients
Least Absolute Shrinkage and Selection Operation

RÉGRESSION LASSO

JDZucker 1

Lasso: contrainte sur la norme L1

Norme L1 : $\|\beta\|_1 = \sum_{j=1}^p |\beta_j|$

Très similaire à Ridge

$$\min_{\beta_1, \dots, \beta_p} \sum_{i=1}^n \left(y_i - \sum_{j=1}^p \beta_j z_{ij} \right)^2 \quad \text{s.c.} \quad \sum_{j=1}^p |\beta_j| \leq \tau$$

Avec la fonction de pénalité

$$\min_{\beta_1, \dots, \beta_p} \sum_{i=1}^n \left(y_i - \sum_{j=1}^p \beta_j z_{ij} \right)^2 + \lambda \sum_{j=1}^p |\beta_j|$$

λ ($\lambda \geq 0$) est un paramètre (coefficient de pénalité) qui permet de contrôler l'impact de la pénalité : à fixer

Idem Ridge : $\left\{ \begin{array}{l} \bullet x_j \text{ doivent être centrées et réduites, } y \text{ au moins centré (mais peut être réduit aussi)} \\ \bullet \text{ Relation inverse entre } \tau \text{ et } \lambda ; (\lambda = 0 \Leftrightarrow \tau = +\infty) \Rightarrow \text{MCO} ; \text{ plus } \lambda \text{ élevé, plus la régularisation est forte (} \tau \text{ faible, les coefficients sont rétrécis)} \end{array} \right.$

Quel intérêt par rapport à Ridge ? LASSO peut faire office de dispositif de sélection de variables en annulant certains coefficients β_j : les variables associées à ($\beta_j = 0$) sont de facto exclues du modèle prédictif.



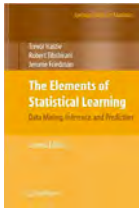
120
JDZucker

Algorithme d'apprentissage-LARS (least-angle regression)

Il n'y a pas de calcul direct pour LASSO. Passer par la régression LARS. C'est une démarche itérative où l'on commence avec tous les ($\beta_j=0$). On fait évoluer les coefficients sélectivement. On a ainsi des scénarios de solutions (LASSO path) à mesure que $\|\beta\|_1$ augmente.

Algorithme LARS

« Elements of Statistical Learning »,
Hastie, Tibshirani, Friedman, Corrected
12th, Jan. 2017 ; page 74.



Légère modification
permettant d'obtenir LASSO

Algorithm 3.2 Least Angle Regression.

1. Standardize the predictors to have mean zero and unit norm. Start with the residual $\mathbf{r} = \mathbf{y} - \bar{\mathbf{y}}$, $\beta_1, \beta_2, \dots, \beta_p = 0$.
2. Find the predictor $\mathbf{x}_j$ most correlated with $\mathbf{r}$.
3. Move β_j from 0 towards its least-squares coefficient $\langle \mathbf{x}_j, \mathbf{r} \rangle$, until some other competitor $\mathbf{x}_k$ has as much correlation with the current residual as does $\mathbf{x}_j$.
4. Move β_j and β_k in the direction defined by their joint least squares coefficient of the current residual on $(\mathbf{x}_j, \mathbf{x}_k)$, until some other competitor $\mathbf{x}_l$ has as much correlation with the current residual.
5. Continue in this way until all p predictors have been entered. After $\min(N-1, p)$ steps, we arrive at the full least-squares solution.

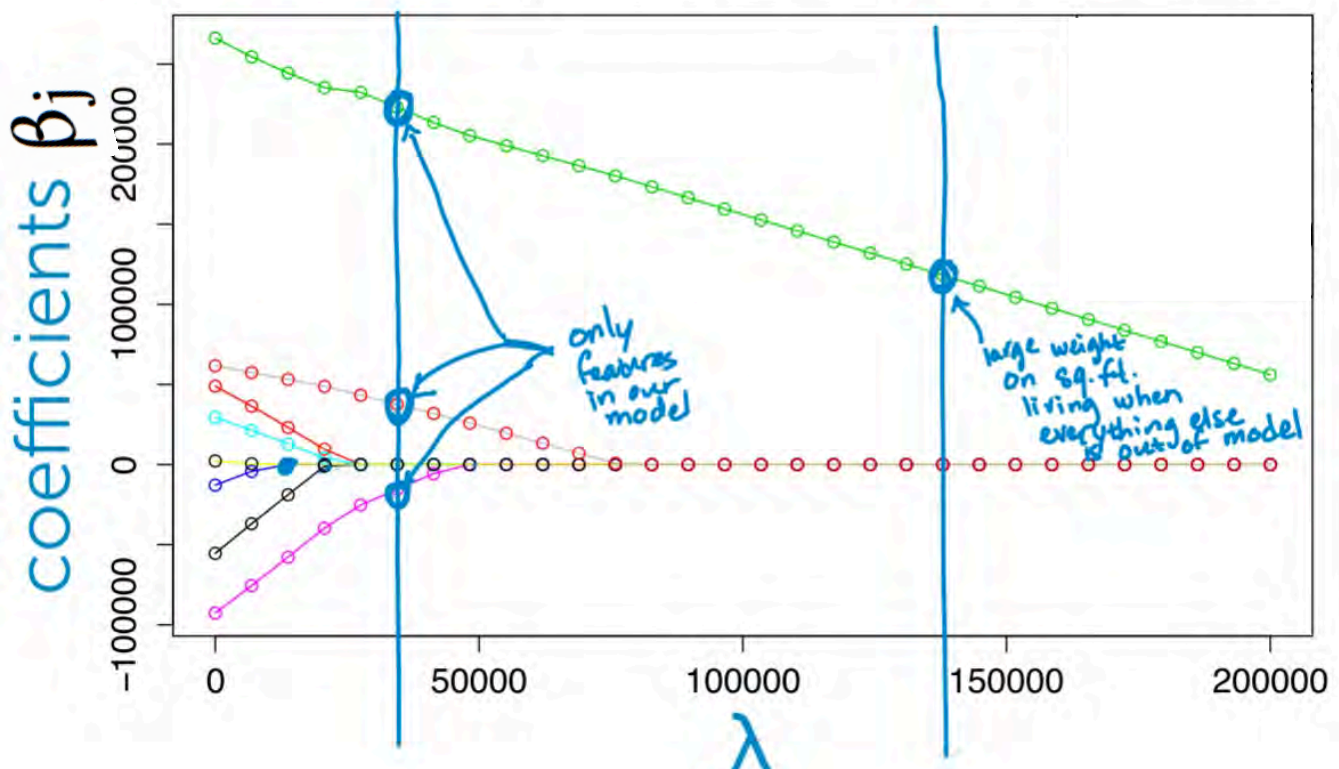
Algorithm 3.2a Least Angle Regression: Lasso Modification.

- 4a. If a non-zero coefficient hits zero, drop its variable from the active set of variables and recompute the current joint least squares direction.

121
JDZucker

Lasso: Least Absolute Shrinkage and

Coefficient path – lasso



Choix du paramètre λ

Dans l'optique prédictive, la minimisation de l'erreur de prédiction en balayant les valeurs de λ reste la méthode la plus efficace. On procède souvent par validation croisée.

Puisque c'est en validation croisée, on peut même avoir un écart-type et donc un intervalle de confiance $[MSE + s.e.(MSE)]$

```
#package R
library(glmnet)

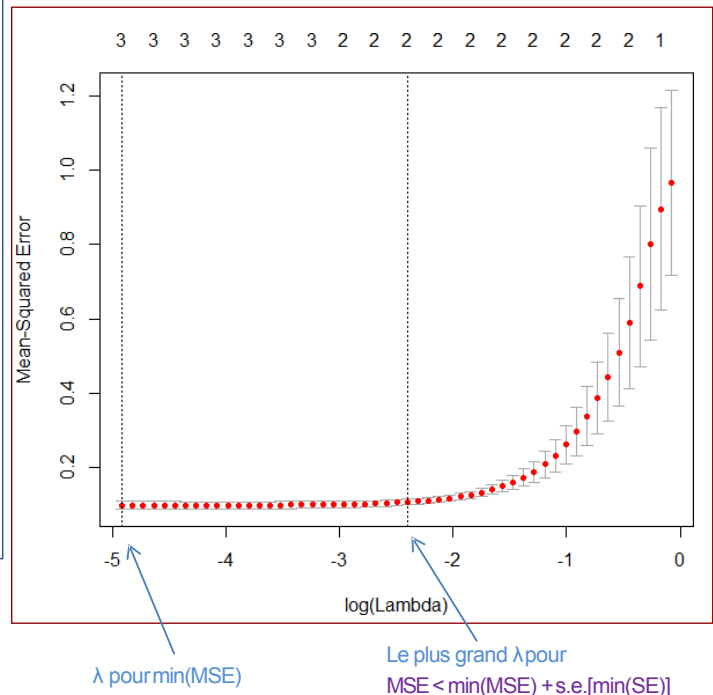
#3-fold cross-validation
cvfit <- cv.glmnet(as.matrix(S[1:3]),S$CONSO,nfolds=3)

#courbe log(Lambda) vs. MSE
plot(cvfit)

#valeur min de MSE(en validation croisée)
print(min(cvfit$cvm)) #0.0987

#lambda corresp.
print(cvfit$lambda.min) #0.007347

#lambda le plus élevé dont le MSE est inf.
#à la borne haute de l'intervalle de min(MSE)
cvfit$lambda.1se #0.09057664 (log[0.09057] = -2.4)
```



123
JDZucker

Bilan LASSO

Avantage

Capacité à sélectionner les variables en acceptant les **coefficients nuls**

Inconvénients

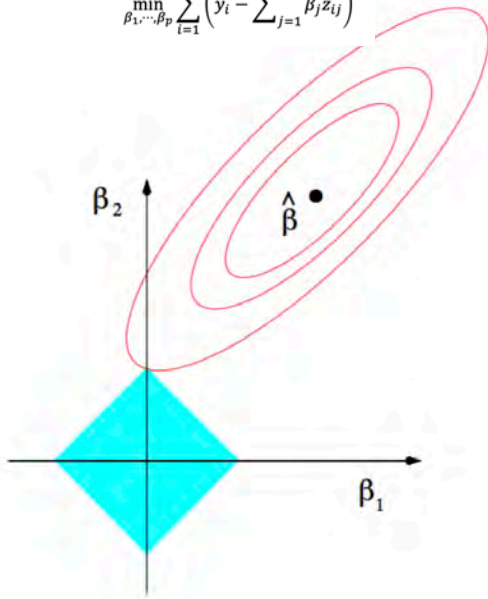
Dans les problèmes à très grandes dimensions ($p \gg n$), LASSO **ne sélectionne que n variables prédictives au maximum**, mécaniquement. C'est une vraie limitation de l'algorithme.

Parmi un groupe de variables corrélées, LASSO en choisit une, celle qui est la plus liée à la cible; souvent, masquant l'influence des autres. Cet inconvénient est inhérent aux techniques intégrant un mécanisme de sélection de variables (*ex. arbres de décision, quoique...*).

124
JDZucker

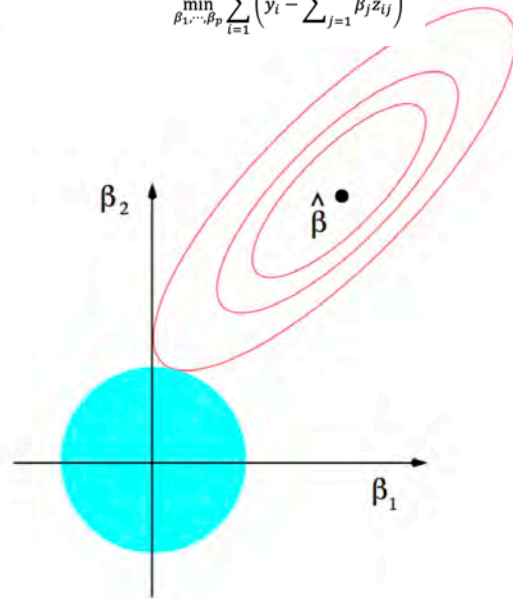
Interprétation géométrique

$$\min_{\beta_1, \dots, \beta_p} \sum_{i=1}^n \left(y_i - \sum_{j=1}^p \beta_j z_{ij} \right)^2$$



Norme L1 : $\|\beta\|_1 = \sum_{j=1}^p |\beta_j|$

$$\min_{\beta_1, \dots, \beta_p} \sum_{i=1}^n \left(y_i - \sum_{j=1}^p \beta_j z_{ij} \right)^2$$



Norme L2 : $\|\beta\|_2^2 = \sum_{j=1}^p \beta_j^2$

125
JDZucker

Mix de Ridge et Lasso

ELASTICNET

Elasticnet : combiner les avantages de ridge et lasso

Formulation sous la forme
d'un problème d'optimisation

$$\min_{\beta_1, \dots, \beta_p} \sum_{i=1}^n \left(y_i - \sum_{j=1}^p \beta_j z_{ij} \right)^2 + \lambda_1 \sum_{j=1}^p |\beta_j| + \lambda_2 \sum_{j=1}^p \beta_j^2$$

λ_1 et λ_2 (≥ 0) sont des paramètres positifs à fixer. ($\lambda_1 = 0$ et $\lambda_2 > 0$) : Ridge ; ($\lambda_1 > 0$ et $\lambda_2 = 0$) : Lasso ; ($\lambda_1 = 0$ et $\lambda_2 = 0$) : MCO.

Formulation alternative
de l'optimisation

$$\min_{\beta_1, \dots, \beta_p} \sum_{i=1}^n \left(y_i - \sum_{j=1}^p \beta_j z_{ij} \right)^2 + \lambda \left[\alpha \sum_{j=1}^p |\beta_j| + (1 - \alpha) \sum_{j=1}^p \beta_j^2 \right]$$

$R(\beta)$ = Fonction de pénalité de la régression elasticnet. ($\alpha = 0$) : Ridge ; ($\alpha = 1$) : Lasso.

Autre formulation :
optimisation sous
contrainte

$$\min_{\beta_1, \dots, \beta_p} \sum_{i=1}^n \left(y_i - \sum_{j=1}^p \beta_j z_{ij} \right)^2$$

s.c. $\alpha \sum_{j=1}^p |\beta_j| + (1 - \alpha) \sum_{j=1}^p \beta_j^2 \leq \tau$

Avec toujours la
relation inverse
entre λ et τ

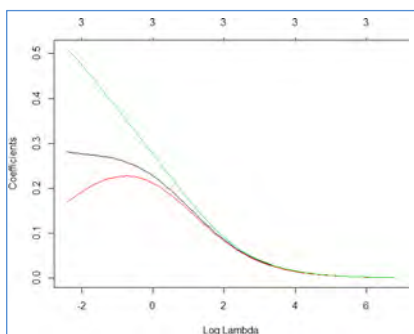
127
JDZucker

Quel intérêt?

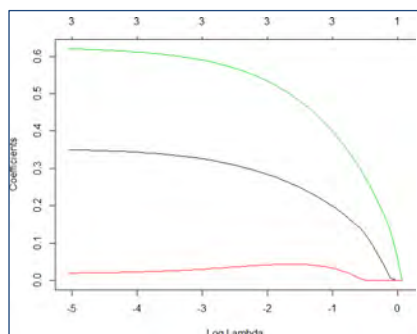
Double avantage :

- Capacité de sélection de variables du LASSO conservée (coefficients nuls): exclusion des variables non pertinentes
- Groupe de variables prédictives corrélées, partage des poids (comme Ridge) et non plus sélection arbitraire

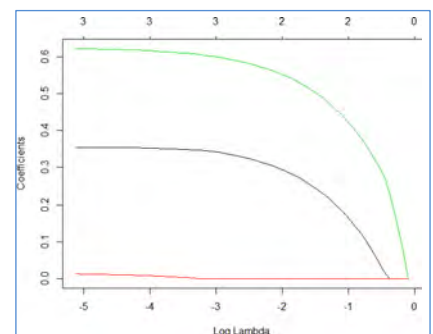
Le « coefficient-path » selon λ dépend de la valeur de α cette fois-ci.



$\alpha = 0$, Ridge
Pas de sélection
de variables



$\alpha = 0.85$
Plus de nuances
dans la sélection



$\alpha = 1$, Lasso
Sélection drastique

128
JDZucker

Conclusion

- ➡ La dimensionalité et la colinéarité sont des problèmes récurrents en régression.
La régularisation via des fonctions de pénalités sur les coefficients est une solution viable.
- ➡ Notamment grâce à sa capacité à sélectionner les variables et à gérer les groupes de variables corrélées.
- ➡ Applicable à la régression linéaire et logistique, mais aussi à d'autres méthodes dès lors que l'on a des combinaisons linéaires de variables avec des coefficients à estimer (ex. réseaux de neurones, perceptron, SVM linéaire, etc.).
- ➡ **Ne jamais oublier de ramener les explicatives sur la même échelle (centrer et réduire est le plus couramment utilisé) pour éviter les inégalités de traitements entre les coefficients.**
- ➡ Ridge, Lasso et Elasticnet sont proposés dans de nombreux packages pour R et Python, dont certains ont servi dans ce support (MASS, glmnet, elasticnet, lars, penalizedSVM, scikit-learn... mais aussi ceux spécialisés dans les réseaux de neurones telles que tensorflow / [keras](#)).

Overfitting and regularization orange3

TP

<https://blog.biolab.si/tag/regularization/>

https://docs.biolab.si/2/_downloads/bio-tutorial.pdf

Plan

I. Santé et IA: type de problèmes, type de donnée et types d'apprentissage

II. Des modèles interprétables pour la santé

III. Exemple de la régression régularisée

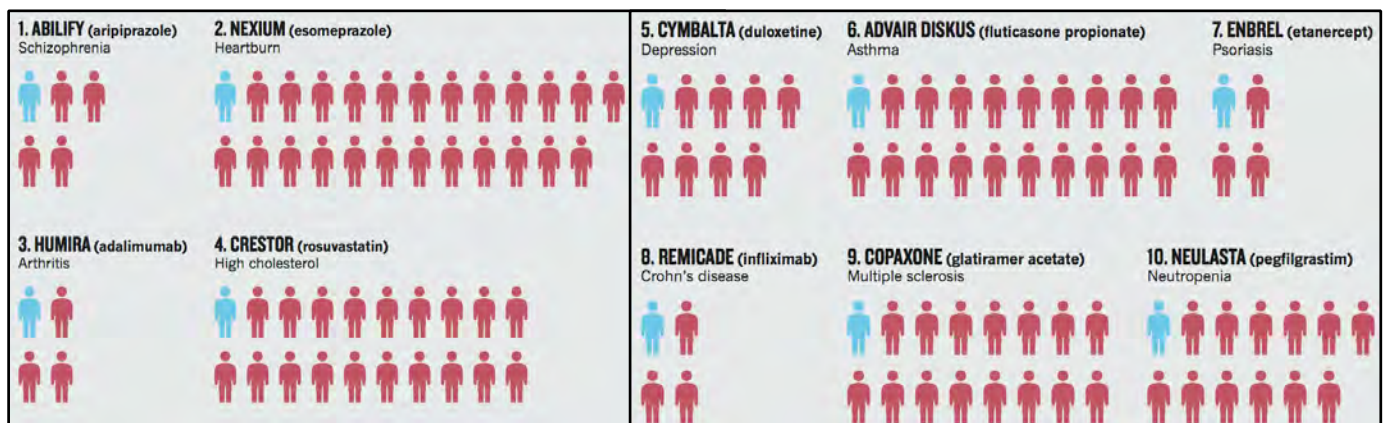
IV. Médecine de précision et IA

V. Classeur avec rejet

JDZucker

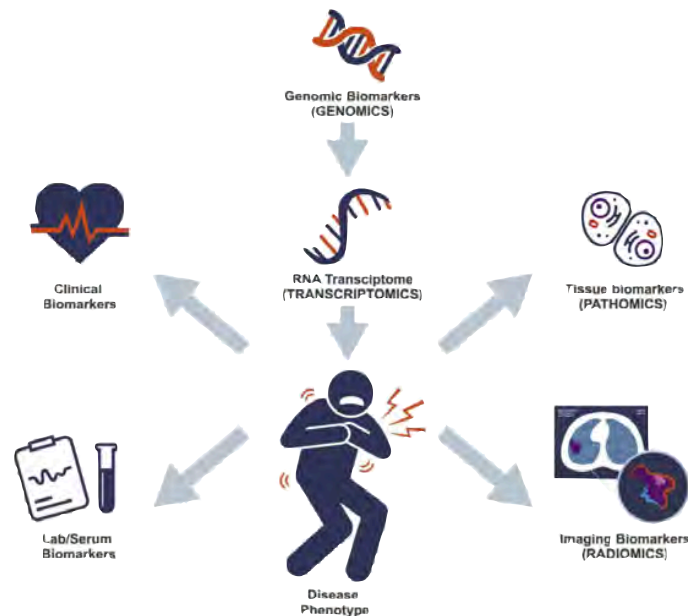
Prédire la réponse aux traitements est un enjeu majeur

Pour **chaque** personne qu'ils **aident (bleu)**, les **dix médicaments** les plus lucratifs aux États-Unis ne parviennent pas à améliorer les conditions d'entre **3** et **24** personnes (rouge).



Emergence de la médecine de précision...

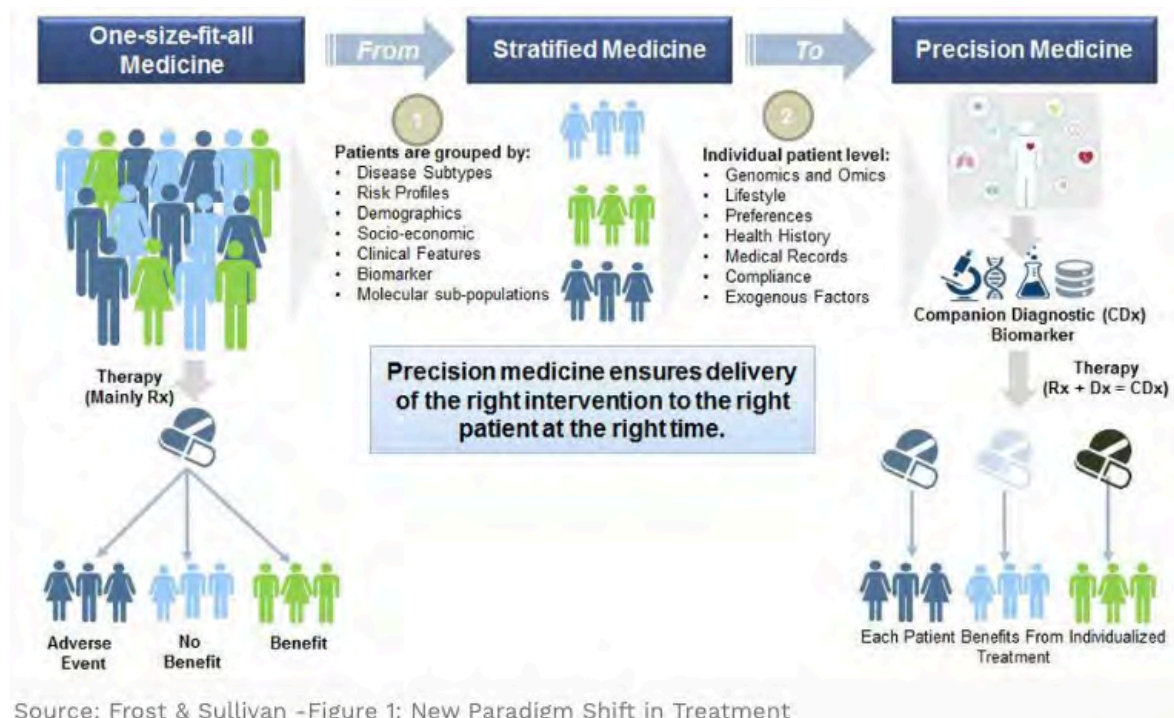
→ fournir les meilleurs soins disponibles à chaque patient, sur la base d'une stratification en **sous-classes de maladies** présentant une base biologique commune.



Shaikh et al., "Translational Radiomics: Defining the Strategy Pipeline and Considerations for Application-Part 1: From Methodology to Clinical Implementation," JACR, 2018.

133
JDZucker

La fin de la médecine « one-size-fit all » annonce celle de la médecine de précision...



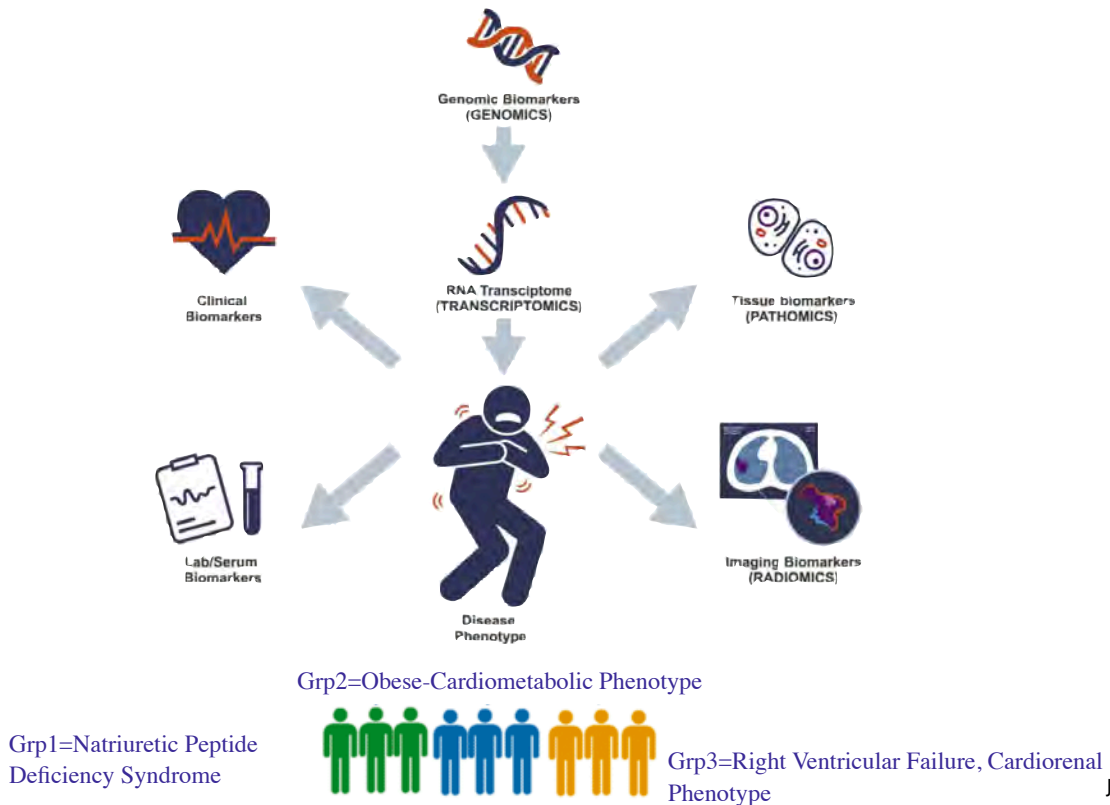
Source: Frost & Sullivan -Figure 1: New Paradigm Shift in Treatment

→ la bonne intervention au bon patient au bon moment.

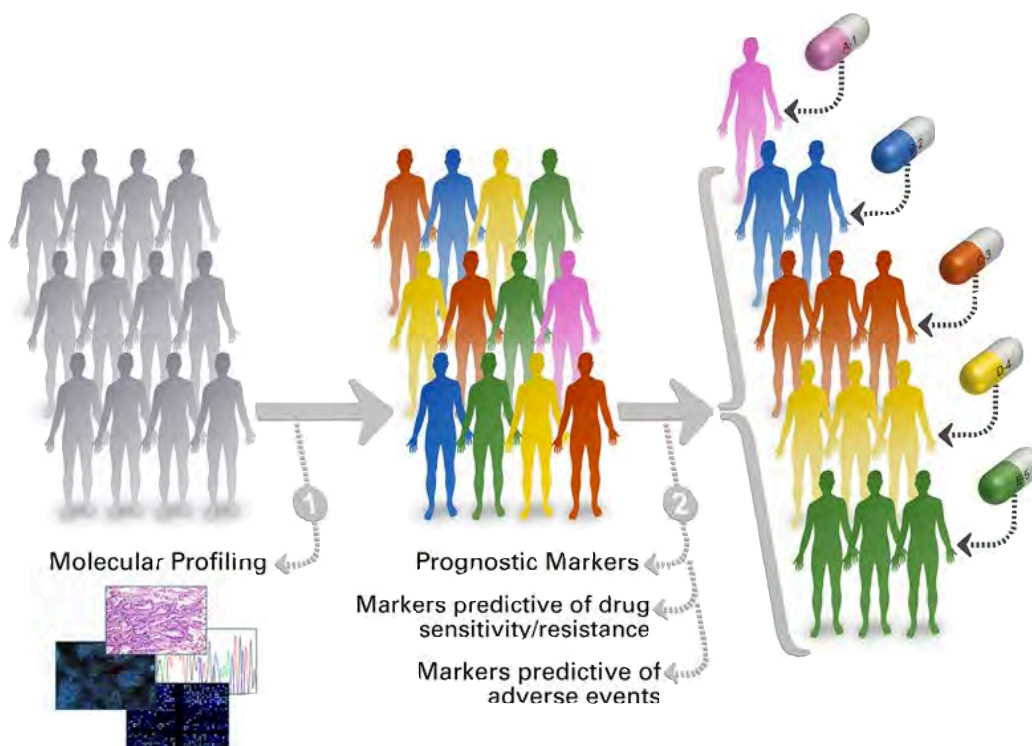
134
JDZucker

Emergence de la médecine de précision...

→ fournir les meilleurs soins disponibles à chaque patient, sur la base d'une stratification en **sous-classes de maladies** présentant une base biologique commune.



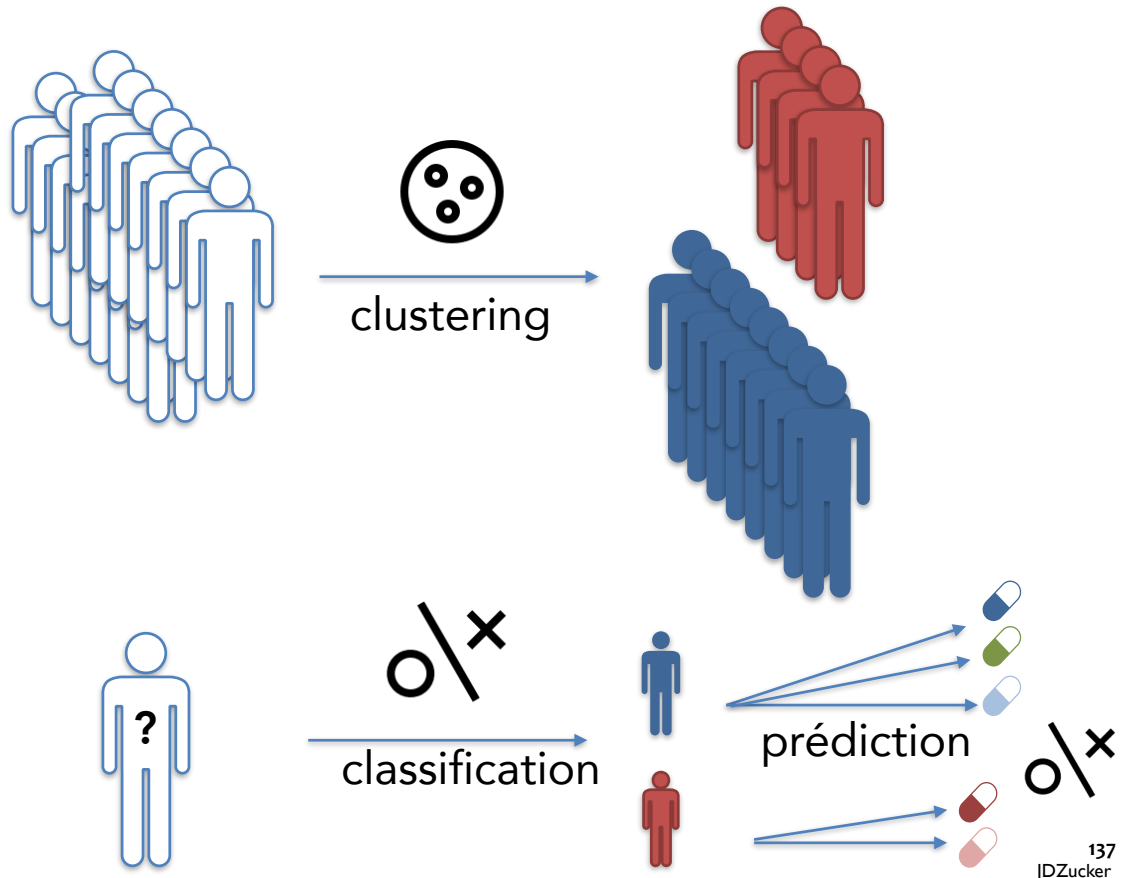
Médecine de précision pour le cancer



→ biomarqueurs tumoraux associés au pronostic du patient et / réponse tumorale au traitement.

Médecine de précision et apprentissage automatique:

Clustering pour stratifier et Classifier/prédire pour traiter

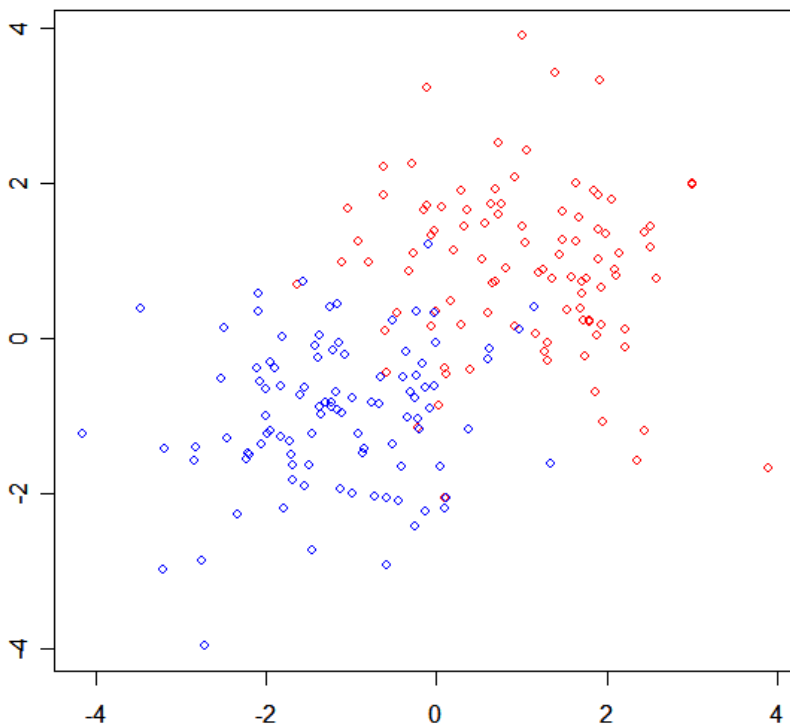


Plan

- I. Santé et IA: type de problèmes, type de donnée et types d'apprentissage
- II. Des modèles interprétables pour la santé
- III. Exemple de la régression régularisée
- IV. Médecine de précision et IA
- V. Classeur avec rejet
- VI. Conclusion

Principle

Crédit: Blaise Hanczar

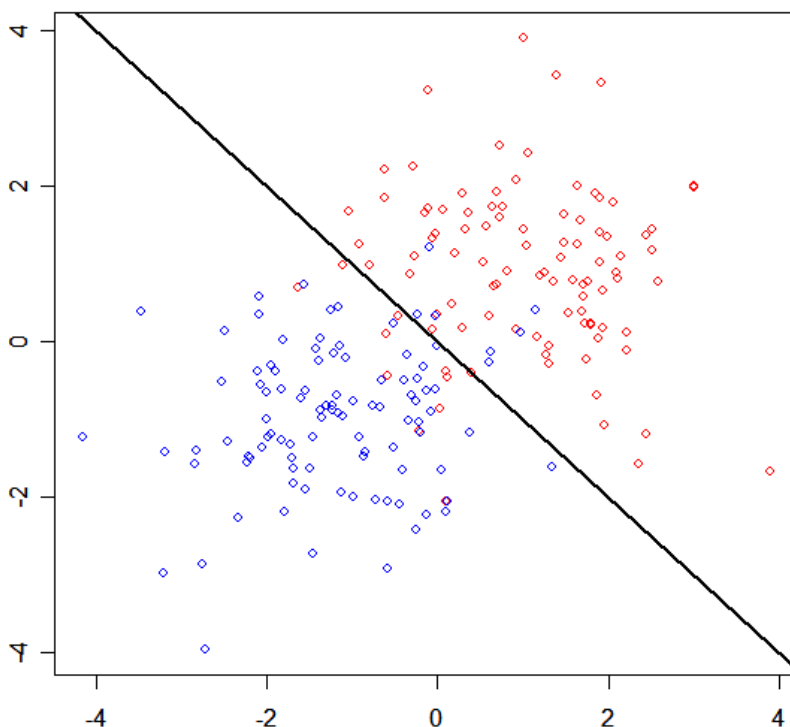


2-class classification problem

		Actual class	
Predicted class		P	N
Predicted class	P		
	N		

139
JDZucker

Principle



2-class classification problem

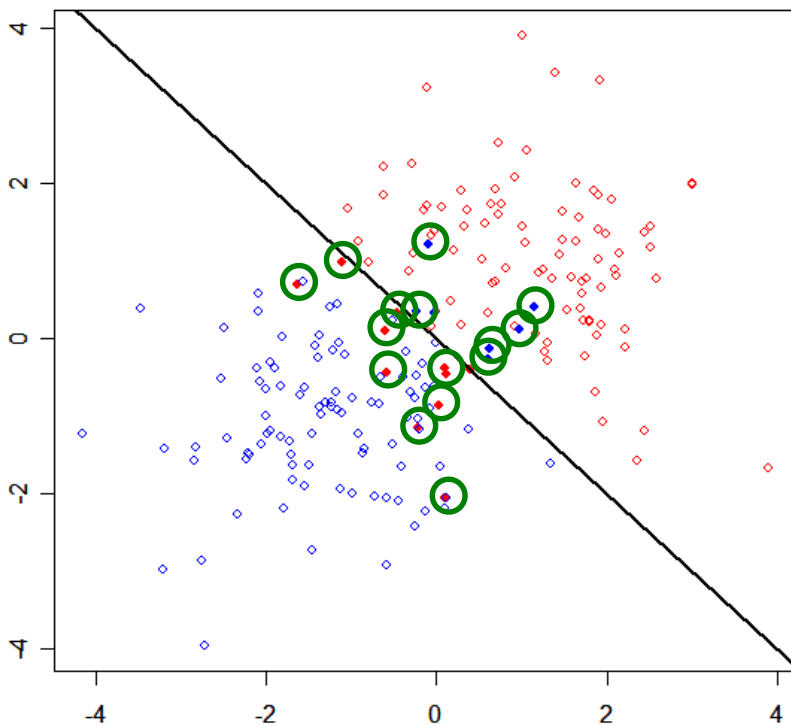
		Actual class	
Predicted class		P	N
Predicted class	P		
	N		

Classifier

$$\psi(x): \mathbb{R}^d \rightarrow \{P, N\}$$

140
JDZucker

Principle



2-class classification problem

Actual class		
Predicted class	P	N
P		
N		

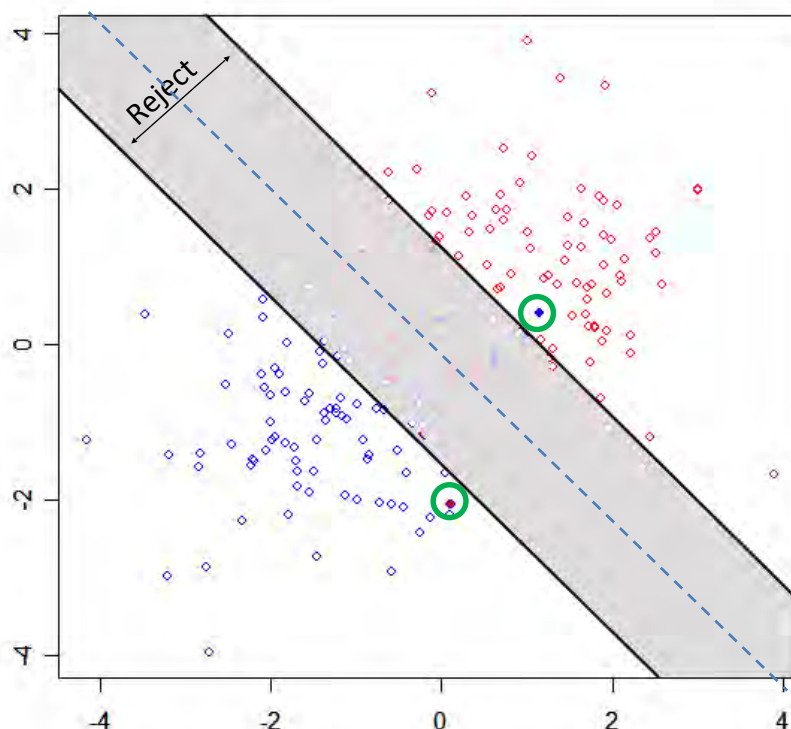
Classifier

$$\psi(x): \mathbb{R}^d \rightarrow \{P, N\}$$

16 misclassifications

All errors are close to the separator

Principle



2-class classification problem

Actual class		
Predicted class	P	N
P		
N		
R		

Classifier with rejection option

$$\psi(x): \mathbb{R}^d \rightarrow \{P, N, R\}$$

2 misclassifications

41 rejections

Plug-in methods

Classifier computes a score for each example :

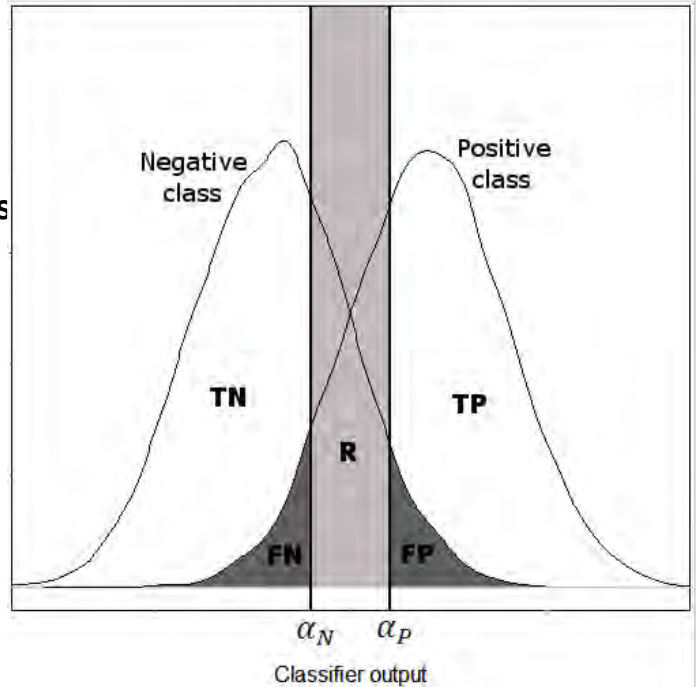
$$\begin{aligned}\psi: \mathbb{R}^d &\rightarrow \mathbb{R} \rightarrow \{P, N, R\} \\ x &\mapsto \omega(x) \mapsto \psi(x)\end{aligned}$$

Two thresholds are applied on this score to do the class assignment :

$$\psi_{\alpha_P, \alpha_N}(x) = \begin{cases} P & \text{if } \omega(x) \geq \alpha_P \\ N & \text{if } \omega(x) \leq \alpha_N \\ R & \text{if } \alpha_N < \omega(x) < \alpha_P \end{cases}$$

The best classifier minimizes :

$$L(\psi_{\alpha_P, \alpha_N}) = \pi_P(\lambda_{TP}TPR + \lambda_{FN}FNR + \lambda_R RPR) + \pi_N(\lambda_{TN}TNR + \lambda_{FP}FPR + \lambda_R RNR)$$



Plug-in methods

Classifier with reject option can also be viewed as a combination of two no-reject classifiers :

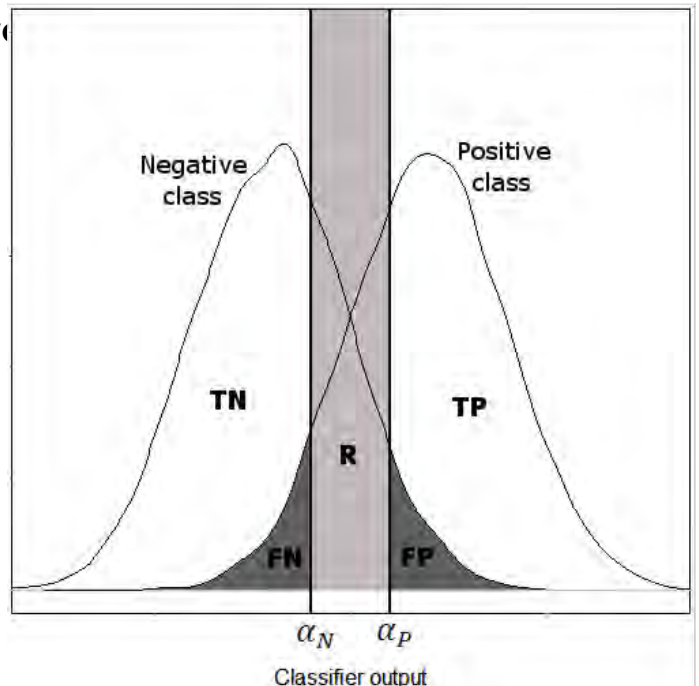
ψ_{α_P} is focused on Positive class

ψ_{α_N} is focused on Negative class

$$\psi_{\alpha_P, \alpha_N}(x) = \begin{cases} P & \text{if } \psi_{\alpha_P}(x) = \psi_{\alpha_N}(x) = P \\ N & \text{if } \psi_{\alpha_P}(x) = \psi_{\alpha_N}(x) = N \\ R & \text{if } \psi_{\alpha_P}(x) \neq \psi_{\alpha_N}(x) \end{cases}$$

The best classifier minimizes :

$$L(\psi_{\alpha_P, \alpha_N}) = \pi_P(\lambda_{TP}TPR + \lambda_{FN}FNR + \lambda_R RPR) + \pi_N(\lambda_{TN}TNR + \lambda_{FP}FPR + \lambda_R RNR)$$



Plan

I. Santé et IA: type de problèmes, type de donnée et types d'apprentissage

II. Des modèles interprétables pour la santé

III. Exemple de la régression régularisée

IV. Médecine de précision et IA

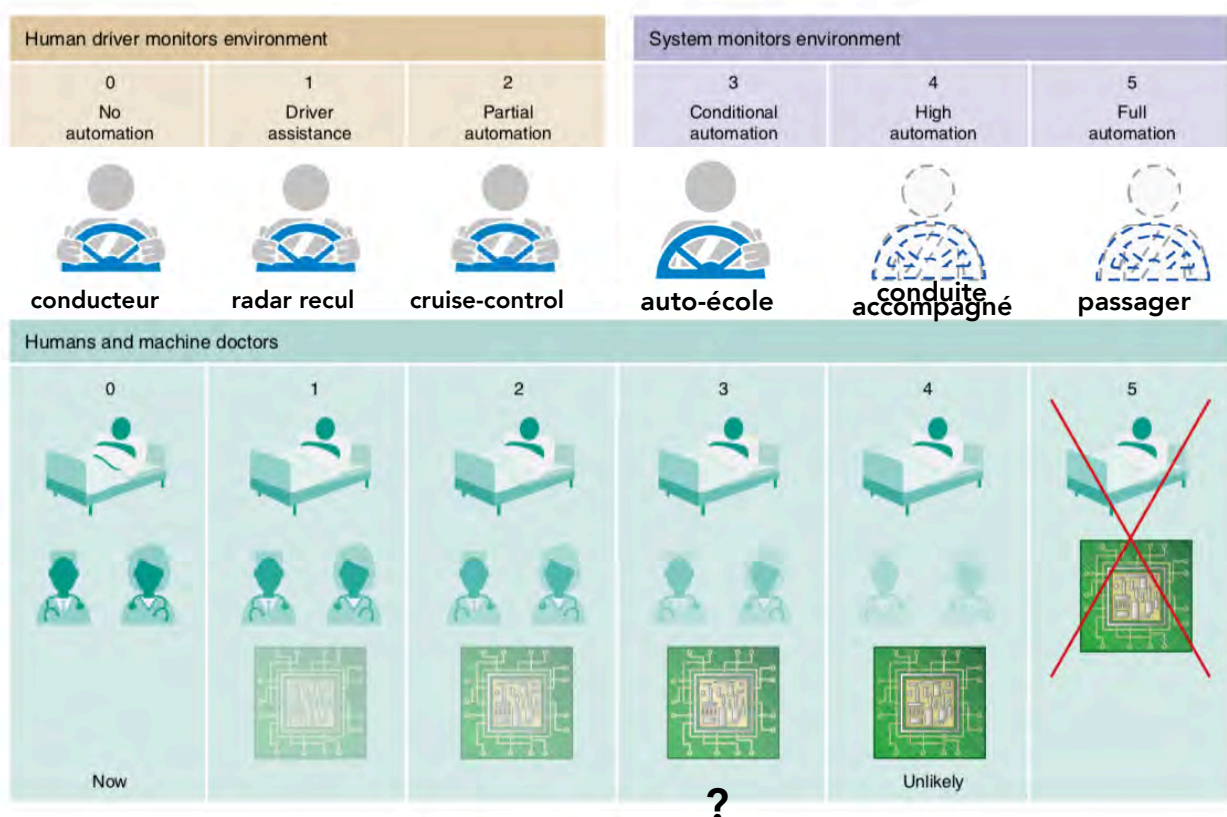
V. Classeur avec rejet

VI. Conclusion

145
JDZucker

Jusqu'à quel point l'IA décidera de notre santé ?

Le but est une synergie, afin que le clinicien reste au commande....



Conclusion

- 👤 **La médecine de précision** annonce un bouleversement dans la prise en charge des patients, leur parcours de soin et leur suivi grâce à l'IA.
- 👤 **Nouveaux diagnostics moléculaires** (omiques) et **d'imagerie**
 - **stratification** des maladies → meilleurs diagnostic,
 - aide au **prognostic** → meilleurs choix des traitements,
 - **désert médicaux** → tri des patients les plus à risques.
- 👤 Progrès de l'IA et du **Deep Learning** posent des **questions éthiques** sur son adoption en médecine : équité/confiance/transparence/**interprétabilité**
- 👤 **L'IA doit aider les cliniciens** (pas se substituer) à être plus efficace mais l'**interprétabilité** est utile pour éviter les erreurs et contribuer à la recherche de l'étiologie ...

147
JDZucker

Milou, l'interprétabilité et sa consommation de Deep Learning



Les images extraites de l'œuvre de Hergé sont la propriété exclusive de MOULINSART SA.
© Hergé-Moulinsart 2018.

Jean-Daniel Zucker

148
JDZucker

Merci

Centre FRANCE
Centre 202 C de France, Bâtiment
Sorbonne Université, Paris, France
12 chercheurs, 2 postdoctorats administratifs,
20 doctorants, 2 post-doctorats, 4 masters
40 articles publiés et en cours
Participation à 22 cours et séminaires de Master 2
Organisation de "Journées Open-Source" et "Ateliers de
Modélisation des Systèmes Complexes"
Le séminaire "Littérature d'usage à la gestion des
Espaces Urbains par les communs"
Ecole Nord-Est REGACI "Missions de l'opérateur"

Centre ASIE DU SUD-EST Vietnam
Université des Sciences et Technologies, Hanoi, Vietnam
12 chercheurs,
14 doctorants, 14 masters / en
20 articles publiés et en cours
Participation à 8 modules et séminaires de Master 2
Plateforme de modélisation et simulation GAMA
Modèles d'aide à la décision environnementale et urbaine
Fouilles de données complexes, bioinformatique
Modélisation hybride informatique / mathématique

Centre MEDITERRANEE Maroc
Université Cadi Ayyad, Marrakech, Maroc
14 chercheurs,
25 doctorants, 20 masters / en
25 articles publiés et en cours
Participation à 12 cours et séminaires de Master 2
Modélisation et analyse mathématiques des
systèmes complexes
Modèle intégré de gestion de la pollution urbaine

Centre AFRIQUE CENTRALE Cameroun
Université Nguinsu 1, Yaoundé, Cameroun
13 chercheurs,
32 doctorants, 11 masters
11 articles publiés et en cours
Participation à 10 modules de Master 2
Plateforme EPICAM de surveillance géodésique
déployée sur 12 centres de surveillance la télédétection
Organisation annuelle de la Conférence de
Recherche en Informatique du Cameroun
Modélisation mathématique pour l'aide à la
prise de décision des villes

Centre AFRIQUE DE L'OUEST Sénégal
Université Cheikh Anta Diop, Dakar, Sénégal
10 chercheurs, 20 doctorants, 10 masters
20 articles publiés et en cours
21 doctorants, 20 masters
20 articles publiés et en cours
Participation à 13 cours et séminaires de Master 2
Colloque de Master 2 "Systèmes Complexes" de l'UMISCO
Modélisation des dynamiques de l'écologie de pêche et de
la dynamique démographique de la ressource
Conception d'applications intelligentes Big Data, IoT,
pour l'aide à la décision des villes et du territoire

UMMISCO

149
JDZucker

Bibliographie

