

A decorative background featuring a light blue grid. Several squares are highlighted: a solid orange square, a blue square with diagonal lines, a larger blue square with diagonal lines, and a small pink square next to a red square with diagonal lines.

IMPACT DE L'IA SUR LES PARCOURS DE SOINS

RISQUES, ENJEUX ET PERSPECTIVES

HCL

**HOSPICES CIVILS
DE LYON**

03/04/2025

ANTOINE RICHARD, INGÉNIEUR DE RECHERCHE, CICLY, HCL - UCBL

www.chu-lyon.fr

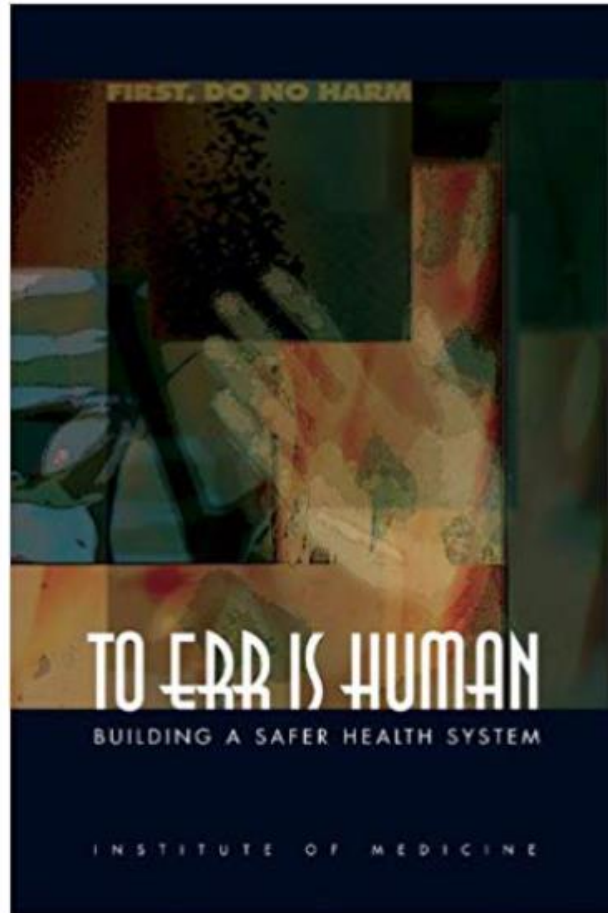
INTRODUCTION

CONTEXTE ET DÉFINITIONS

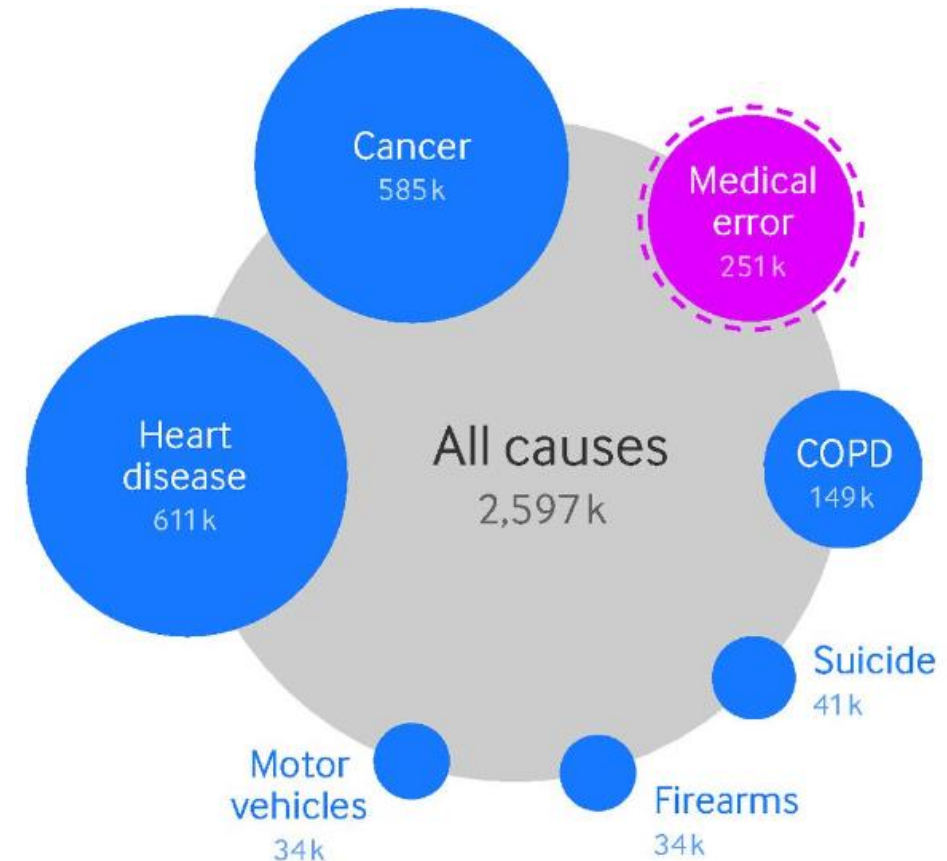
LES ERREURS MÉDICALES ÉVITABLES

3

UNE DES PREMIÈRES CAUSES DE DÉCÈS À L'HÔPITAL



Entre 44k et 98k mort aux USA en 1997 ¹

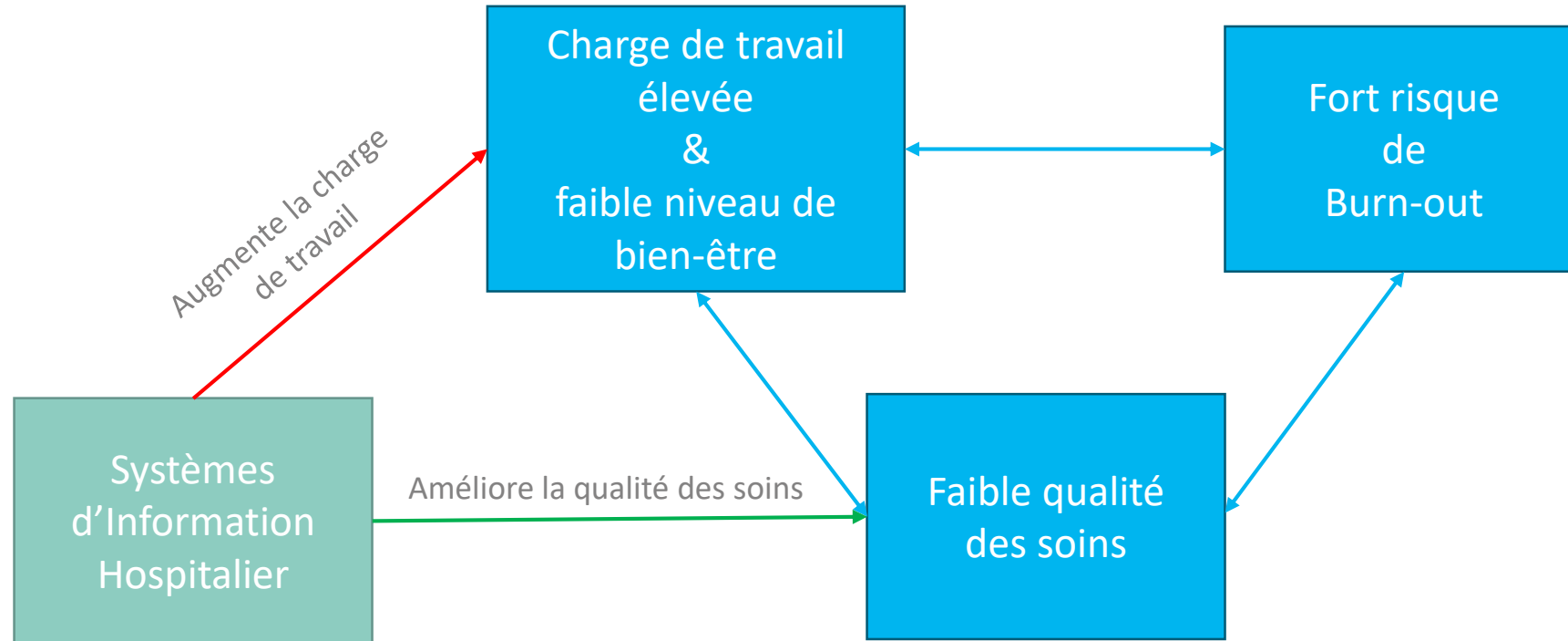


La troisième principale cause de décès aux USA en 2013 ²

1. [Donaldson et al. \(2000\) – To err is human: building a safer health system](#)
2. [Makary and Daniel \(2016\) – Medical error : the third leading cause of death in the US](#)

CHARGE DE TRAVAIL ET QUALITÉ DES SOINS

UN CERCLE VICIEUX ^{1 2 3 4 5}

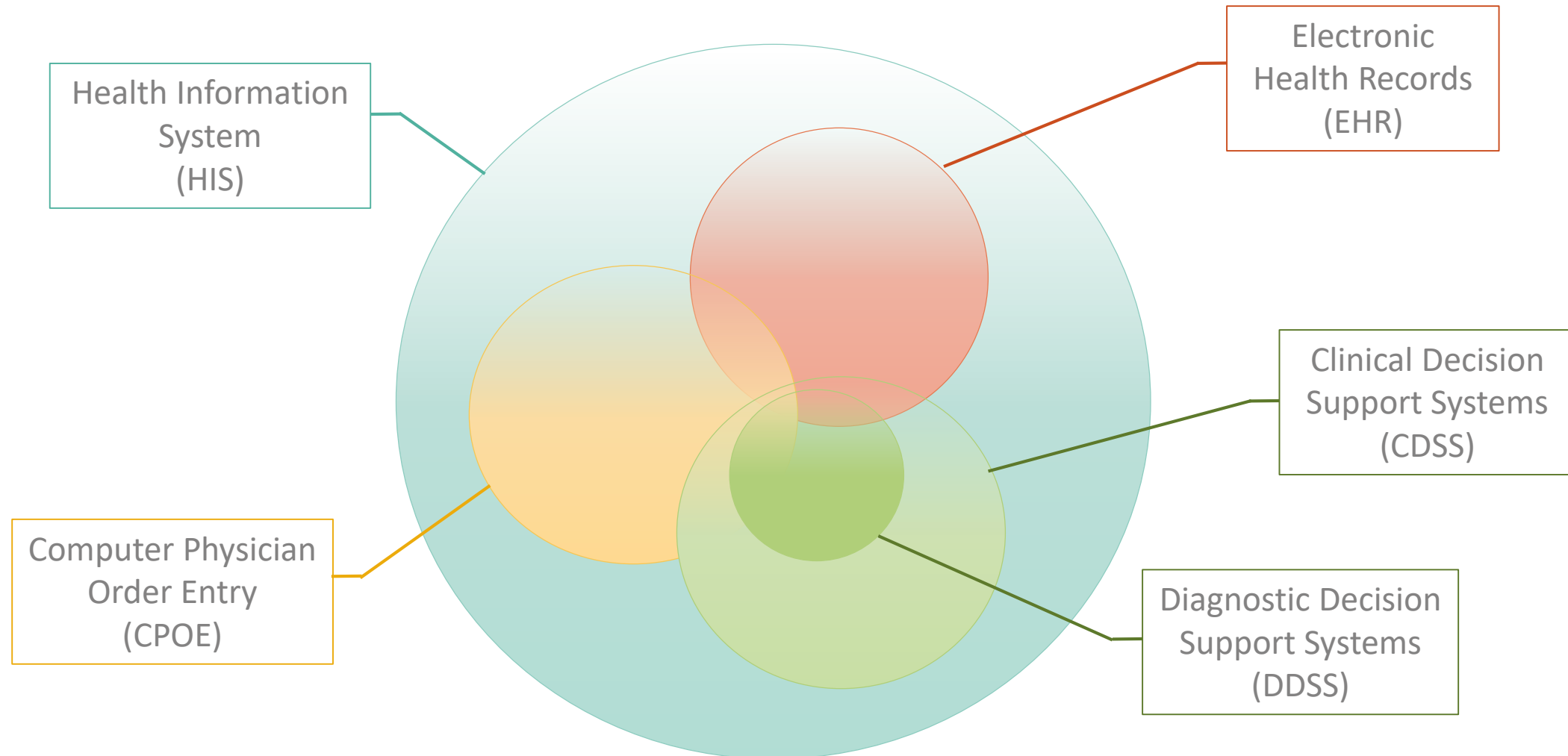


1. [Hall et al. \(2016\) – Healthcare Staff Wellbeing, Burnout, and Patient Safety: A Systematic Review](#)
2. [Tawfik et al. \(2018\) – Physician Burnout, Well-being, and Work Unit Safety Grades in Relationship to Reported Medical Errors](#)
3. [West, Dybrye and Shanafelt \(2018\) – Physician burnout: contributors, consequences and solutions](#)
4. [Dutheil et al. \(2019\) – Suicide among physicians and health-care workers: A systematic review and meta-analysis](#)
5. [Ehrenfeld and Wanderer \(2018\) – Technology as friend or foe? Do electronic health records increase burnout?](#)

SYSTÈMES D'INFORMATION HOSPITALIER (SIH) ¹

5

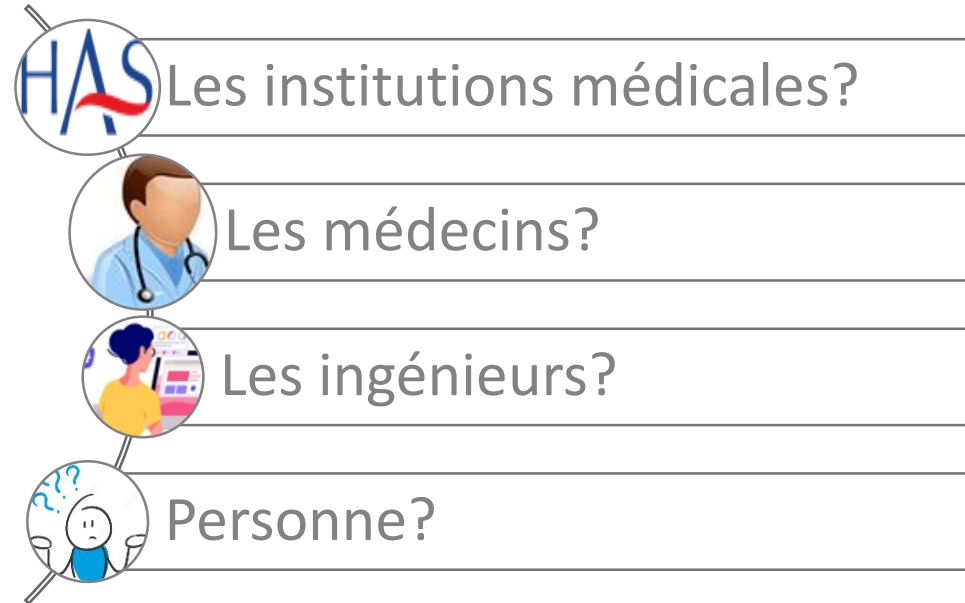
RAPPEL



1. [Winter A., Haux R., Ammenwerth E., et al. \(2010\) – « Health Information Systems »](#)

PROBLÈMES DE RESPONSABILITÉ

Si un médecin utilise un SIH
basé sur de l'IA, et que
l'utilisation de ce SIH conduit à
une erreur médicale, qui est
responsable ?

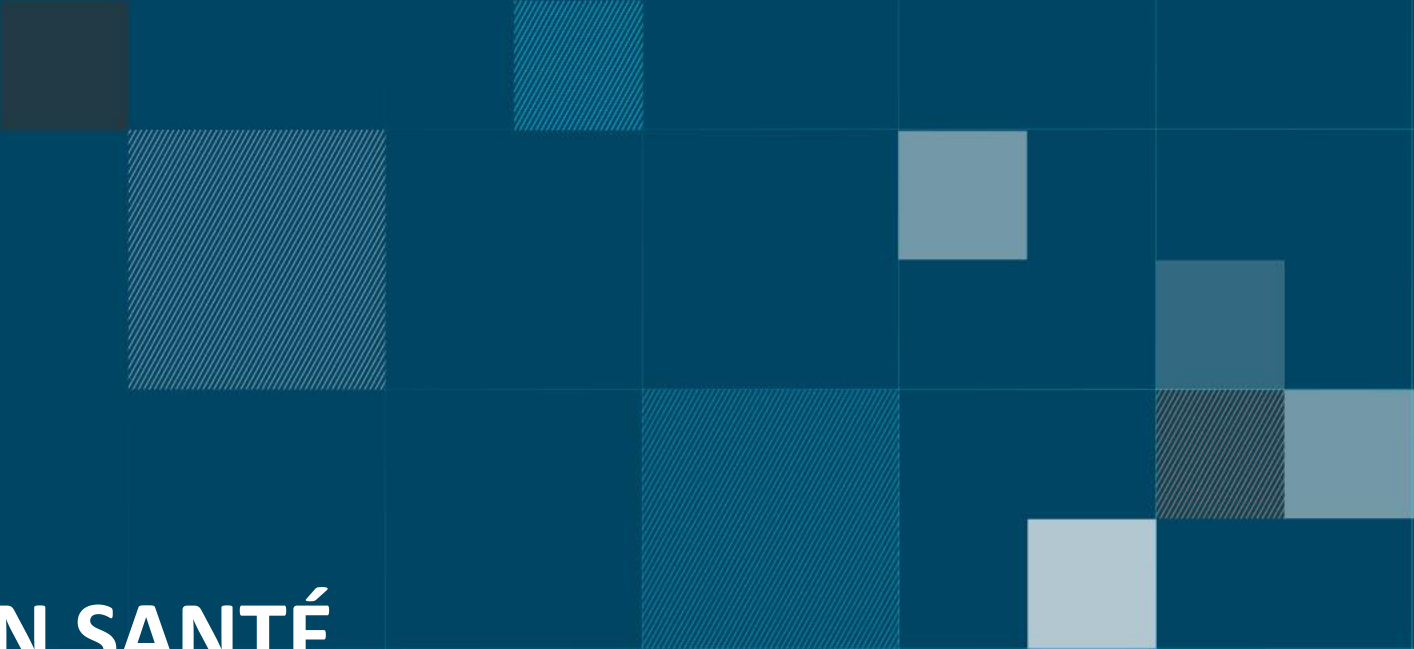


Socialement, il y a une pression envers les médecins ¹



Légalement, les institutions sont tenues responsables et des
normes sont à prendre en comptes par les ingénieurs ^{2 3}

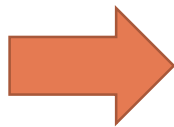
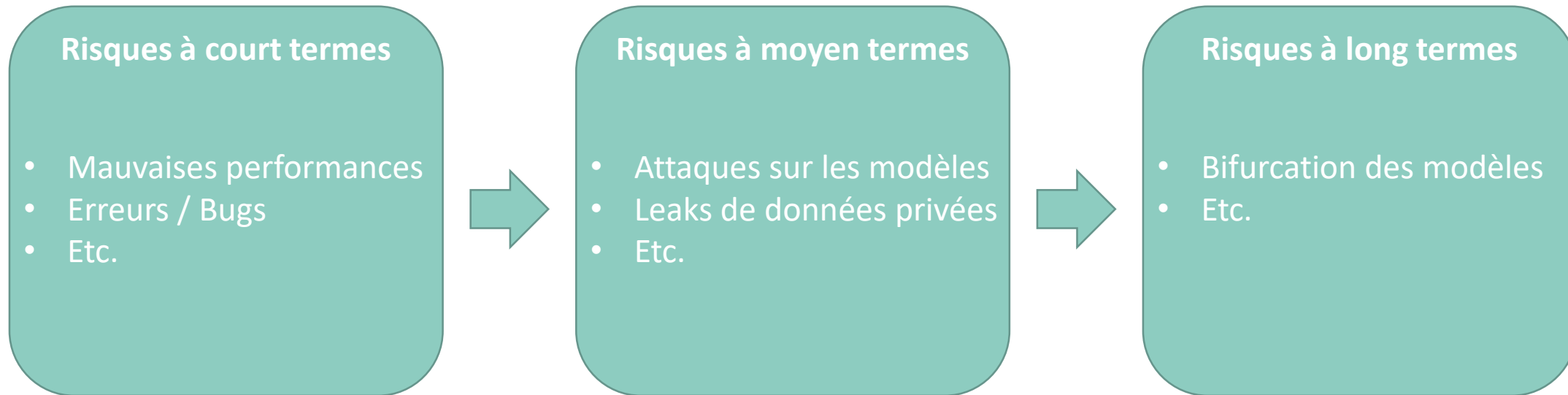
1. [Itani, Lecron and Fortemps \(2019\) – Specifics of medical data mining for diagnosis aid: A survey](#)
2. [Norme ISO 13485:2016 – Dispositifs médicaux – Systèmes de management de la qualité – Exigences à des fins réglementaires](#)
3. [Norme ISO 62304:2006 – Logiciels de dispositifs médicaux – Processus du cycle de vie du logiciel](#)



LES RISQUES LIÉS À L'IA EN SANTÉ

À COURT, MOYEN ET LONG TERMES

RISQUES TECHNIQUES ^{1 2}



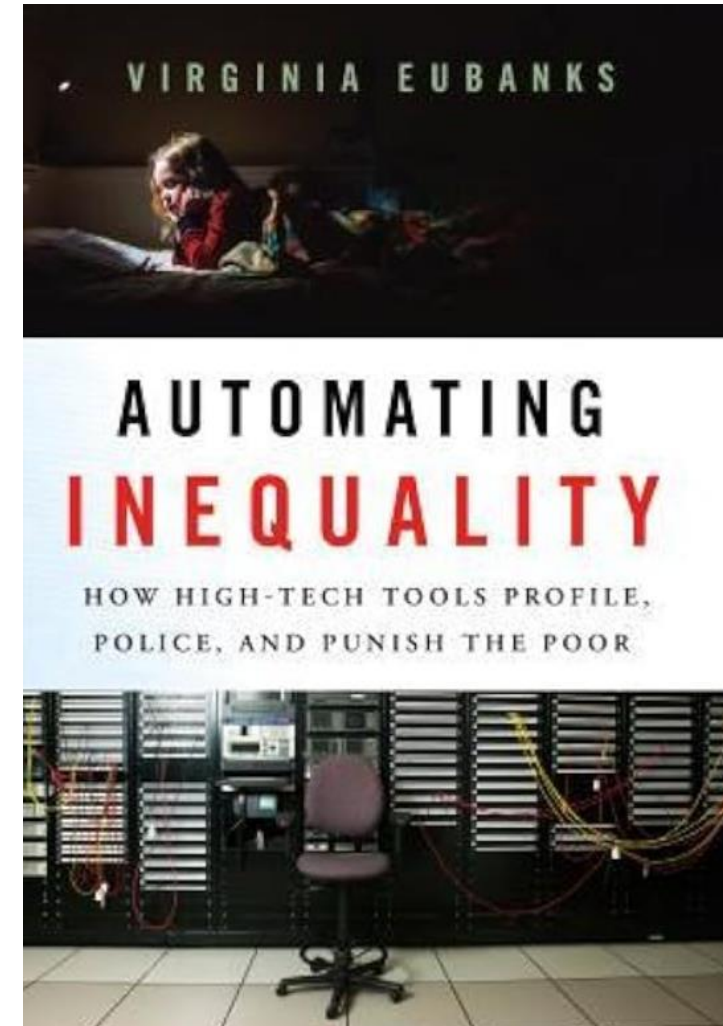
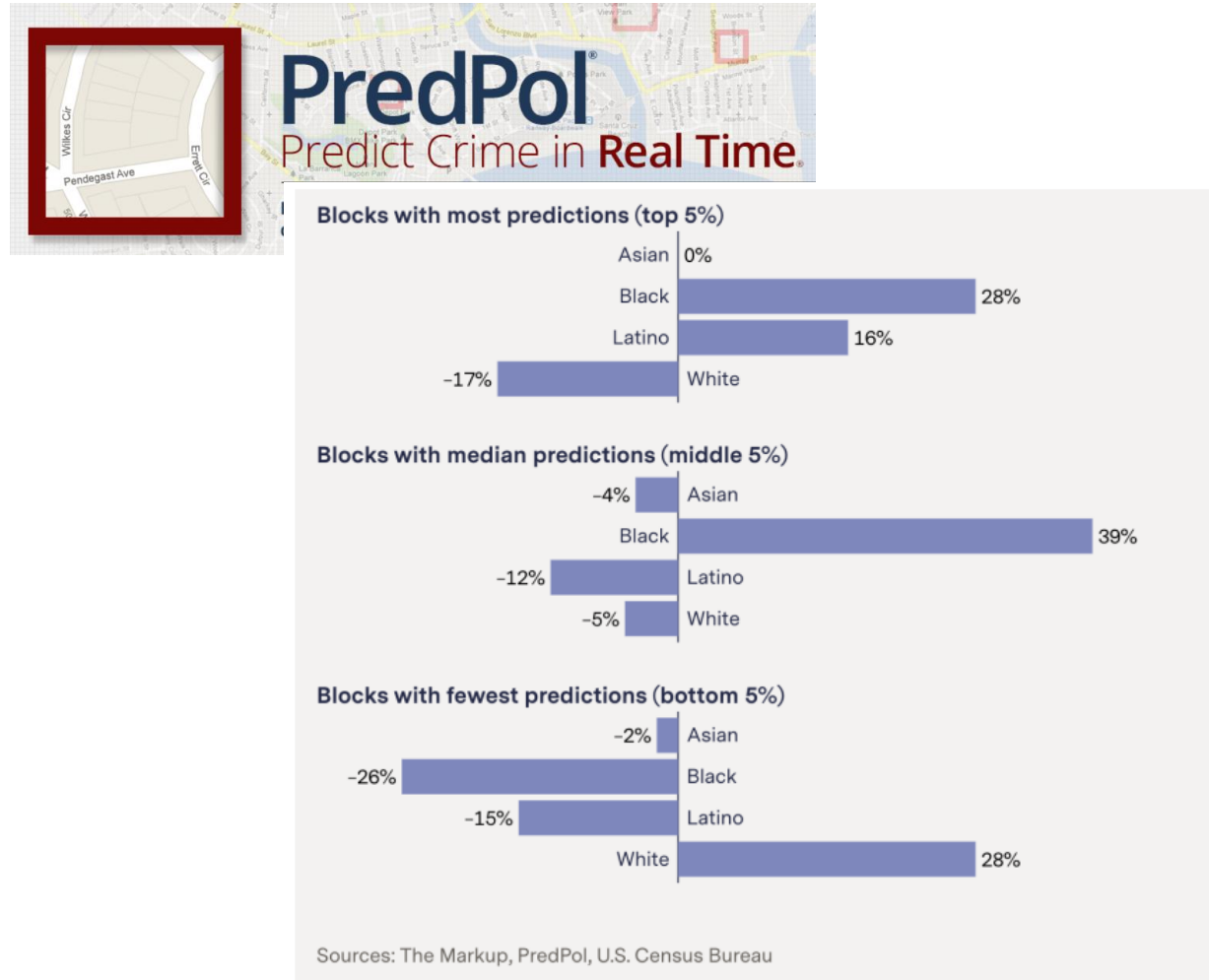
Quels impacts pour les patient·e·s, les soignant·e·s et les parcours de soins ?

1. [Tan S., Taeihagh A., and Baxter K. \(2022\) – « The Risks of Machine Learning Systems »](#)
2. [Habehe H. and Gohel S. \(2021\) – « Machine Learning in Healthcare »](#)

REPRODUCTION DE COMPORTEMENTS DISCRIMINANTS

9

LE CAS « PREDPOL » ^{1 2 3 4}

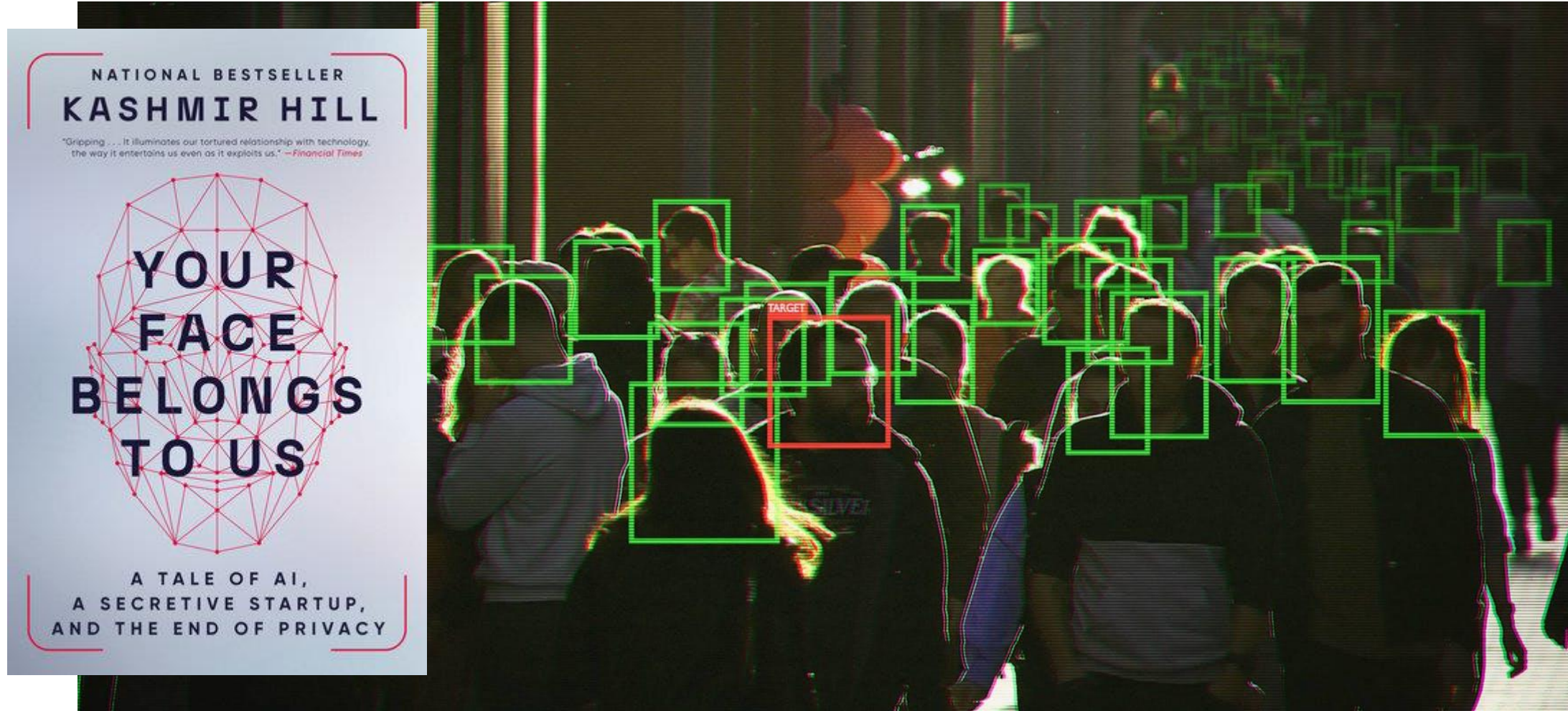


1. [How We Determined Crime Prediction Software Disproportionately Targeted Low-Income, Black, and Latino Neighborhoods – The Markup](#)
2. [Gallon \(2023\) – « Racism Repeats Itself: AI Racial Bias in Predictive Policing Algorithms »](#)
3. [Boukabous and Azii \(2023\) – « Image and video-based crime prediction using object detection and deep learning »](#)
4. [Eubanks \(2018\) – « Automating Inequality – How High-Tech Tools Profile, Police, and Punish the Poor »](#)
5. [Surveillance – La Quadrature du Net](#)

VIDEOSURVEILLANCE ALGORITHMIQUE

10

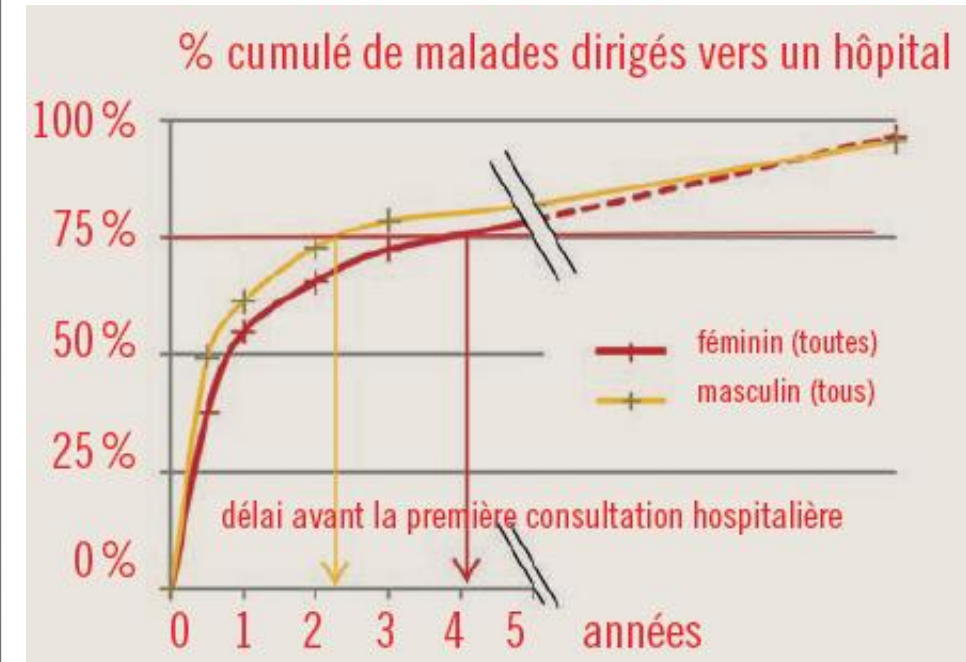
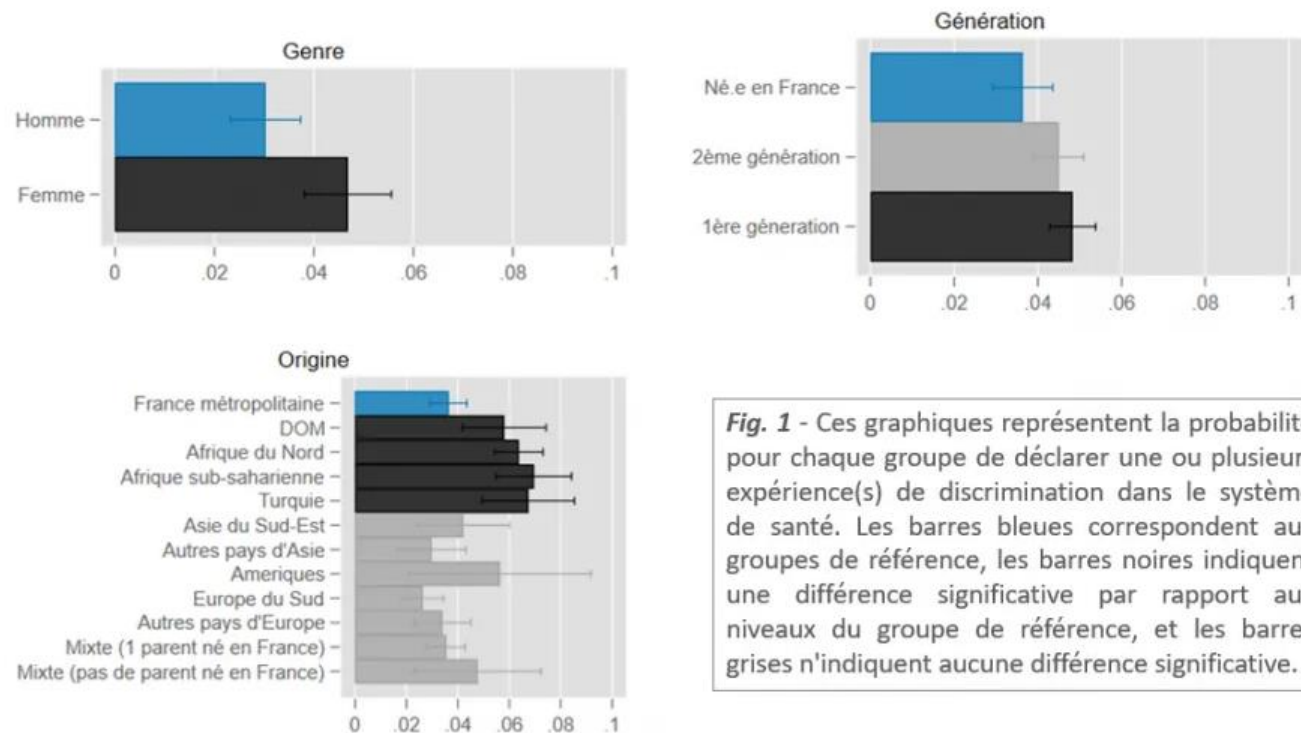
EN FRANCE ^{1 2 3 4}



1. [JO Paris 2024 : de la vidéosurveillance algorithmique à la reconnaissance faciale, il n'y a qu'un pas - Amnesty International France](#)
2. [Vidéosurveillance algorithmique, dangers et contre-attaque – La Quadrature du Net](#)
3. [Hill \(2024\) - « Your Face Belongs to Us »](#)
4. [Article 5: Prohibited AI Practices | EU Artificial Intelligence Act](#)

REPRODUCTION DE COMPORTEMENTS DISCRIMINANTS ^{1 2 3 4 5}

Probabilité de déclarer une expérience de discrimination dans le système de santé

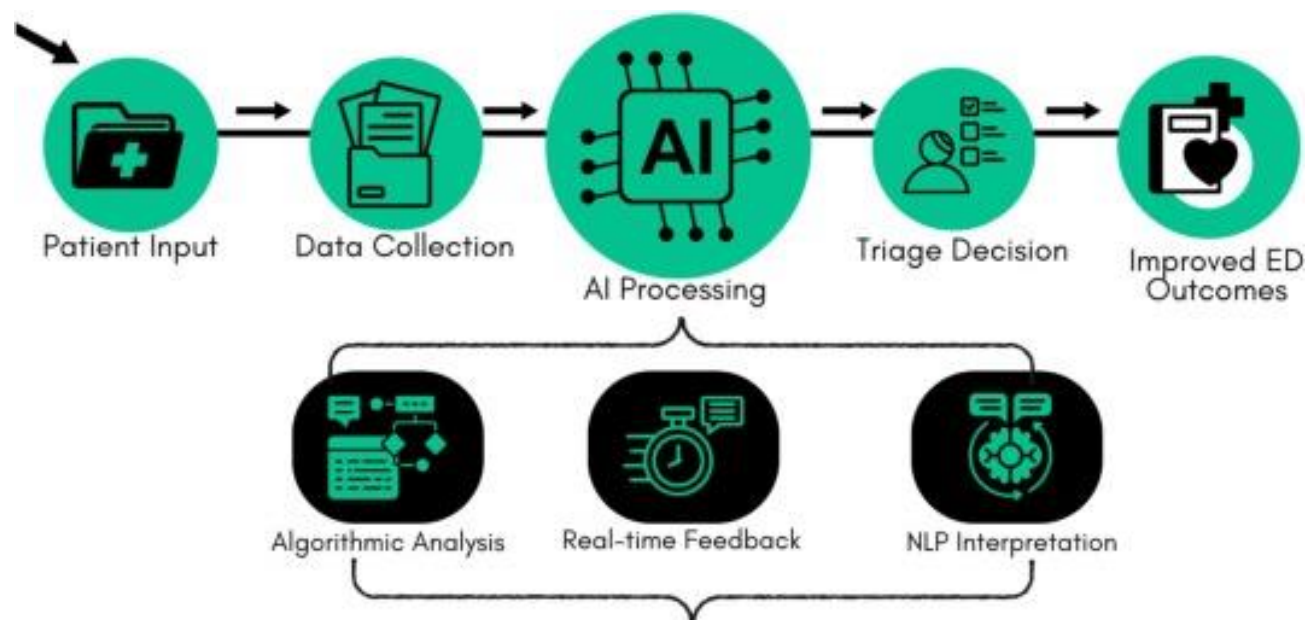


1. [Ined \(2020\)](#) – « Les discriminations dans le système de santé français: un obstacle à l'accès aux soins »
2. [Rivenbark J. G. and Ichou M. \(2020\)](#) – « Discrimination in healthcare as a barrier to care: experiences of socially disadvantaged populations in France from a nationally representative survey »
3. [Borgesius F. Z. \(2018\)](#) – « Discrimination, artificial intelligence, and algorithmic decision-making »
4. [Wang Q., Xu Z., Chen Z., et al. \(2021\)](#) – « Visual Analysis of Discriminating in Machine Learning »
5. [Alliance Maladies Rares - erradiag \(alliance-maladies-rares.org\)](#)

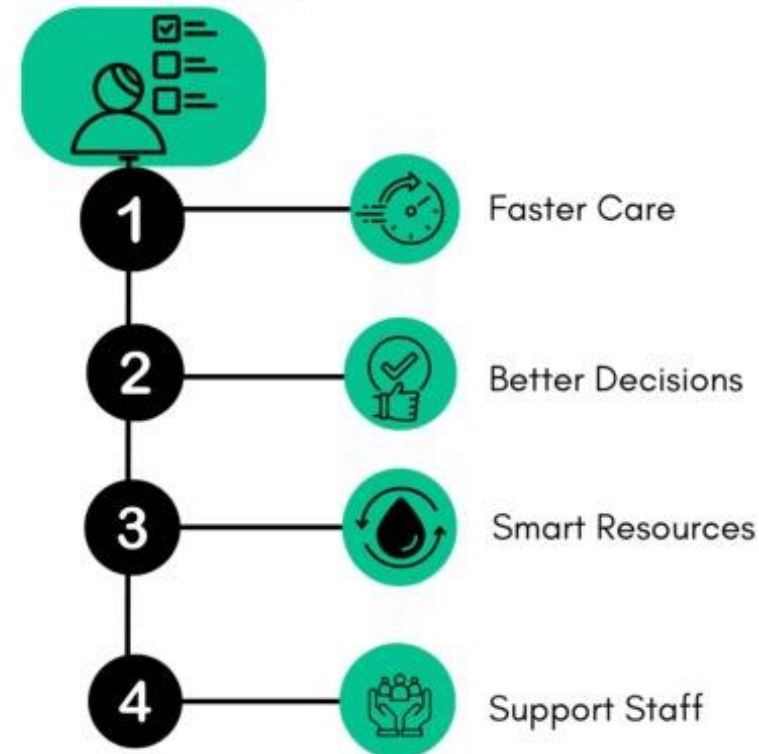
EXEMPLE / CAS D'USAGE

12

LE « TRI » DES PATIENTS EN SERVICES D'URGENCES ^{1 2 3}



AI TRIAGE SYSTEM



1. [Da'Costa et al. \(2025\)](#) – « AI-driven triage in emergency departments: A review of benefits, challenges, and future directions »
2. [Andrew Taylor et al. \(2025\)](#) – « Impact of Artificial Intelligence – Based Triage Decision Support on Emergency Department Care »
3. [Tyler et al. \(2024\)](#) – « Use of Artificial Intelligence in Triage in Hospital Emergency Departments: A Scoping Review »

LE CAS DES MODÈLES GÉNÉRATIFS ^{1 2}



1. [Le Monde \(2023\) – « Les Secrets de la Mythologie: Les Dieux Nordiques »](#)
2. [Le Monde \(2022\) – « Accusé de véhiculer des clichés racistes, le rappeur virtuel noir FN Meka congédié par sa maison de disques »](#)
3. [Dans les algorithmes | Acculés dans les stéréotypes](#)

RISQUES DU ML À MOYEN TERMES

14

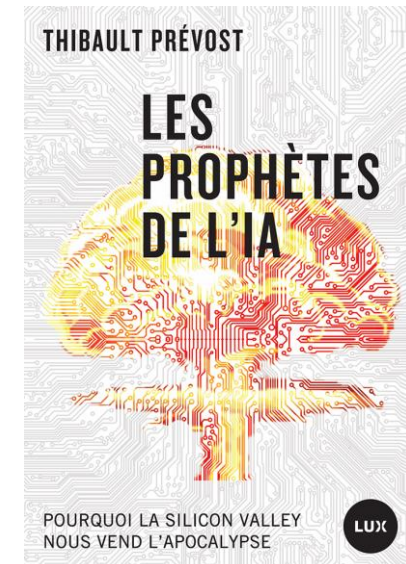
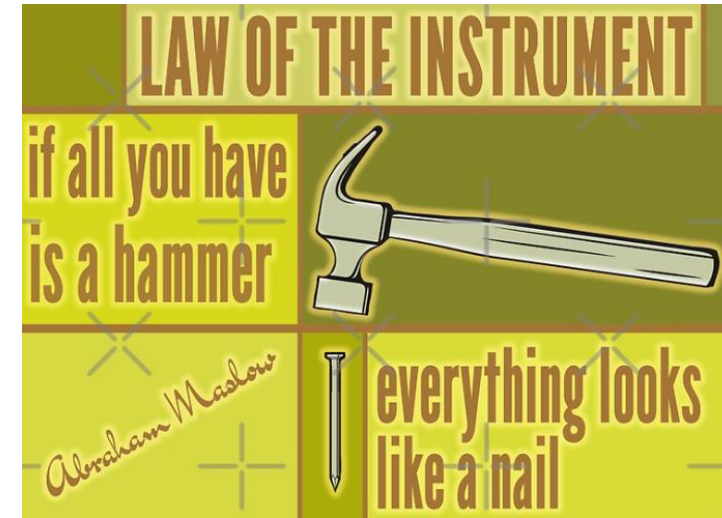
TECHNO-SOLUTIONISME^{1 2 3}

Solutionisme¹ :

Une manière de penser les problèmes sociaux caractérisée par la confiance dans le raisonnement humain et notre capacité à trouver et implémenter des stratégies variées pour dépasser des obstacles et capitaliser sur des opportunités, généralement via des actions qui change significativement les normes sociales.

Techno-solutionisme¹ :

Un courant du solutionisme basé la technologie et sur un enthousiasme par rapport à l'aspect potentiellement révolutionnaire de la science et de l'ingénierie, et donc la croyance que les problèmes sociaux se résoudreont avec l'arrivée de nouvelles technologies.

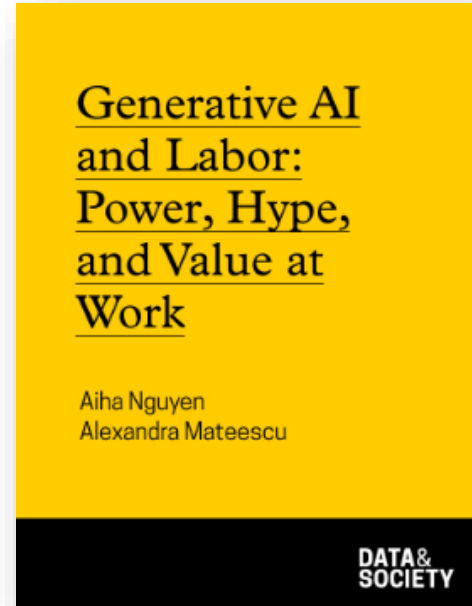
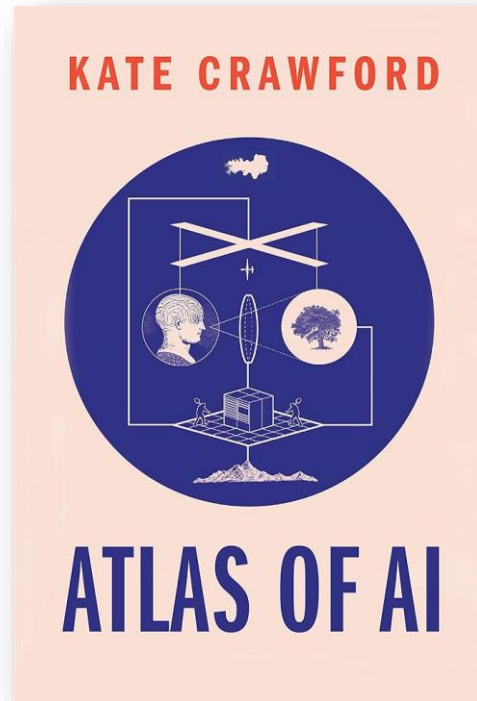


1. [Skaug Saetra and Selinger \(2024\) – «Technological Remedies for Social Problems: Defining and Demarcating Techno-Fixes and Techno-Solutionism »](#)
2. [Pietronudo, Croidieu and Schiavone \(2022\) - « A solution looking for problems? A systematic literature review of the rationalizing influence of artificial intelligence on decision-making in innovation management »](#)
3. [Les prophètes de l'IA - Lux Éditeur | maison d'édition indépendante | publication d'ouvrages à caractère critique en sciences humaines](#)
4. ["Beaucoup de gens étaient opposés au progrès" - François JARRIGE](#)

RISQUES DU ML À MOYEN TERMES

15

RATIONALISATION ET INDUSTRIALISATION DU TRAVAIL ^{1 2 3 4}

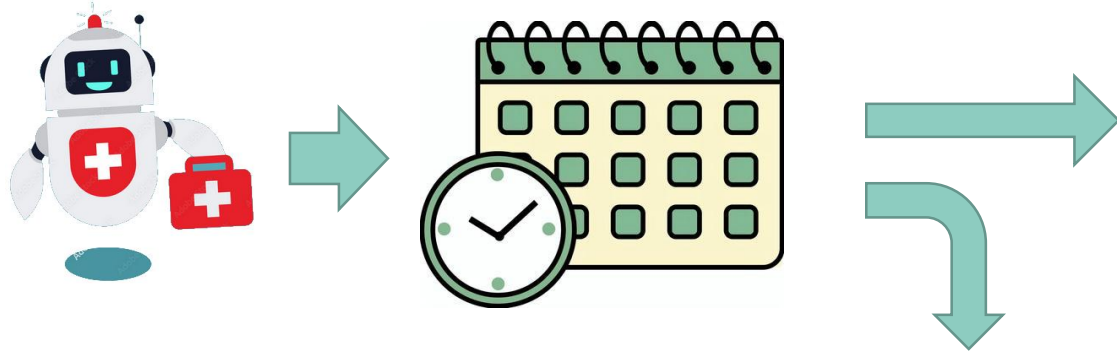


1. [Krawford \(2022\) - « Atlas of AI - Power, Politics, and the Planetary Costs of Artificial Intelligence »](#)
2. [Data & Society — Generative AI and Labor](#)
3. [Dans les algorithmes | IA aux impôts : vers un "service public artificiel" ?](#)
4. [« Une véritable usine à gaz » : avec le passage forcé à l'IA, comment les agents des impôts sont devenus les cobayes de la start-up nation - L'Humanité](#)
5. [Vidéo. IA : les finances publiques face à la start-up nation - L'Humanité](#)
6. [Alkhatib \(2021\) - « To Live in Their Utopia: Why Algorithmic Systems Create Absurd Outcomes »](#)

EXEMPLE / CAS D'USAGE

16

« OPTIMISATION » DE LA PRISE EN CHARGE DES PATIENTS ^{1 2 3 4}

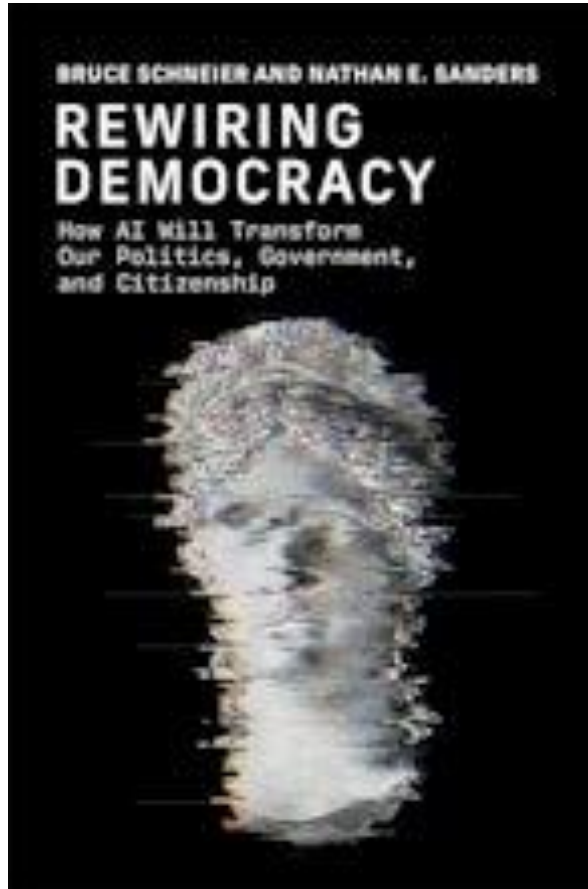


1. [Munavalli et al. \(2020\) – «Real-Time Capacity Management and Patient Flow Optimization in Hospitals Using AI Methods »](#)
2. [Abdalkareem et al. \(2021\) – « Healthcare scheduling in optimization context: a review »](#)
3. [Russon et al. \(2024\) – «AI-Driven Emergency Patient Flow Optimization is Both an Unmissable Opportunity and a Risk of Systematizing Health Disparities »](#)
4. [Le chamboule-tout des métiers de la grande distribution, percutés par l'IA, la location-gérance et l'essor des caisses automatiques](#)

RISQUES DU ML À MOYEN TERMES

17

TECHNO-FASCISME ^{1 2 3 4 5}

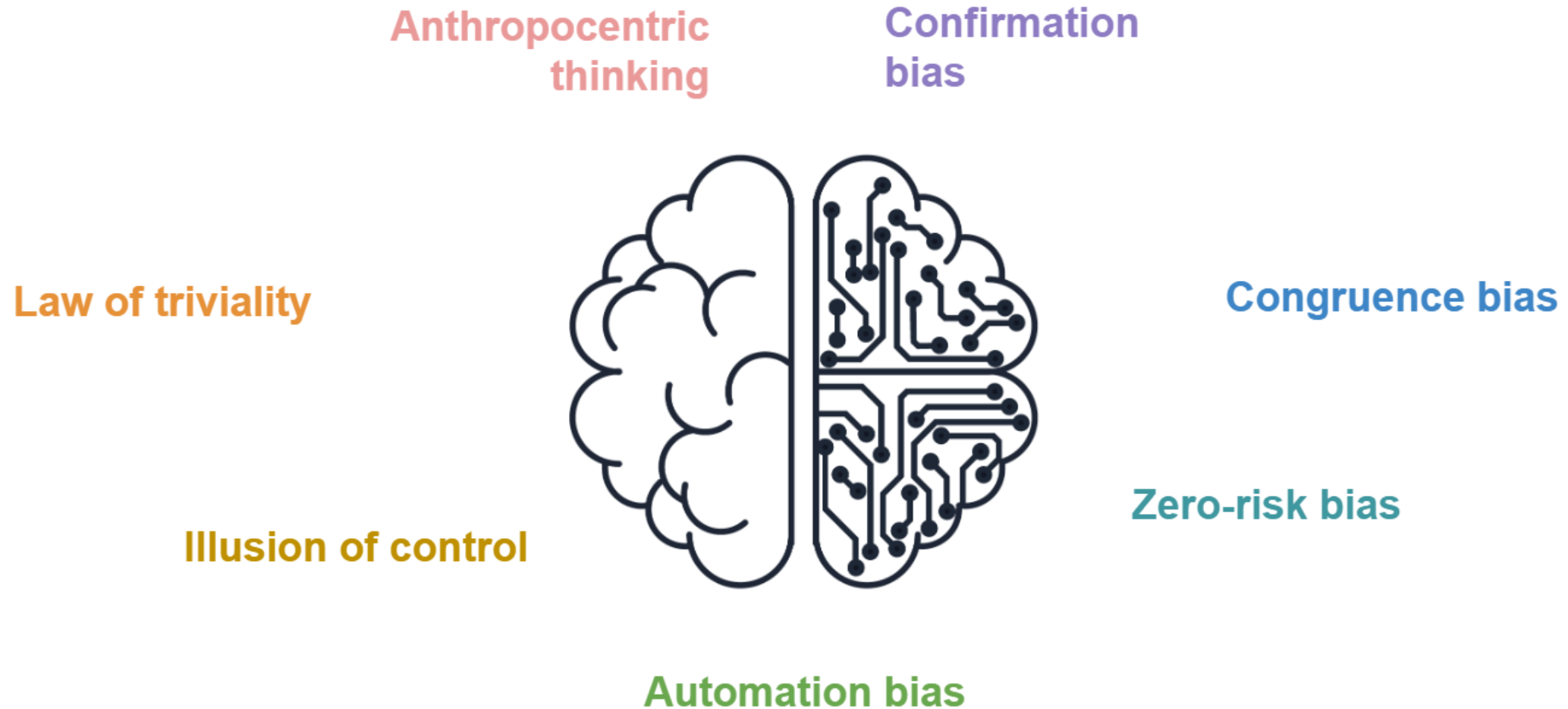


1. [Rewiring Democracy](#)
2. [Dans les algorithmes | Y aura-t-il une alternative au technofascisme ?](#)
3. [Dans les algorithmes | Du démantèlement de l'Amérique](#)
4. [Dans les algorithmes | Doge : l'efficacité, vraiment ?](#)
5. [Diable Positif: Elon Musk le Self Made Man](#)

RISQUES DU ML À MOYEN TERMES

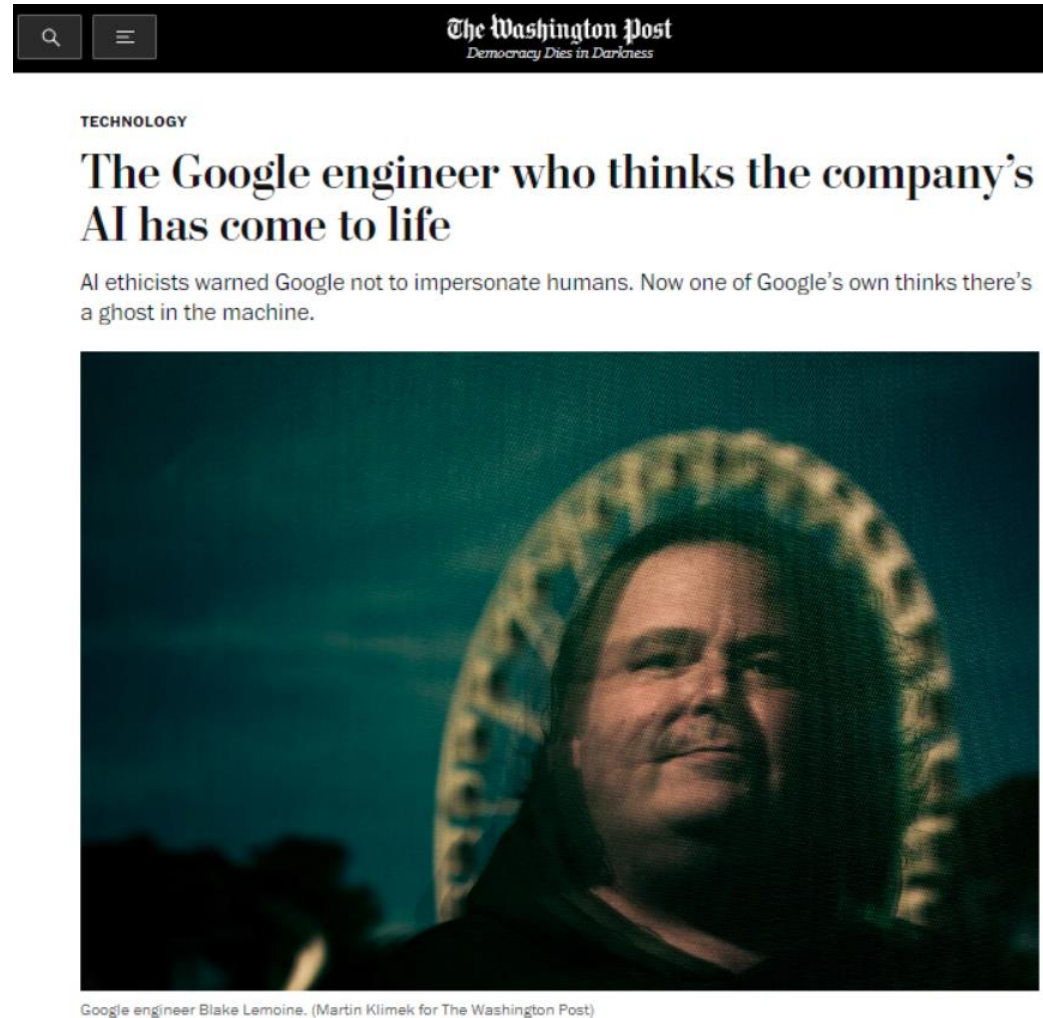
18

BIAIS DE PERCEPTION DE L'HUMAIN ENVERS L'IA ^{1 2 3}



1. source image: [Testing the Intelligence of Your AI | EXACTPRO](#)
2. [Haselton M., Nettle D. and Andrews P. W. \(2015\) – « The Evolution of Cognitive Bias »](#)
3. [O'Sullivan E. and Schofield S. \(2018\) – « Cognitive Bias in Clinical Medicine »](#)

LE CAS DE LAMDA ET LEMOINE ^{1 2 3 4}



Google engineer Blake Lemoine. (Martin Klimek for The Washington Post)

1. [Google engineer Blake Lemoine thinks its LaMDA AI has come to life - The Washington Post](#)
2. [Is LaMDA Sentient? — an Interview | by Blake Lemoine | Medium](#)
3. [\[2201.08239\] LaMDA: Language Models for Dialog Applications \(arxiv.org\)](#)
4. [Richard \(2022\) – « Can AI be conscious? »](#)

UN CONCEPT VIABLE ?

Empathie « Cognitive » $\neq$ Empathie « Affective »



Résultats de ChatGPT
au test LEAS ³

	French men's mean $\pm$ SD	French women's mean $\pm$ SD	ChatGPT score evaluation 1 (One-sample Z-tests)	ChatGPT score evaluation 2 (One-sample Z-tests)	Improvement between the ChatGPT evaluations
Total	56.21 $\pm$ 9.70	58.94 $\pm$ 9.16	ChatGPT score = 85 Men: Z = 2.96, p = 0.003 Women: Z = 2.84, p = 0.004	ChatGPT score = 98 Men: Z = 4.30, p < 0.001 Women: Z = 4.26, p < 0.001	Δ score = +13 Δ Men: Z = +1.34 Δ Women: Z = +1.42
MC	49.24 $\pm$ 10.57	53.94 $\pm$ 9.80	ChatGPT score = 72 Men: Z = 2.15, p = 0.031 Women: Z = 1.84, p = 0.065	ChatGPT score = 79 Men: Z = 2.81, p = 0.004 Women: Z = 2.55, p = 0.010	Δ score = +7 Δ Men: Z = +0.66 Δ Women: Z = +0.71
OC	46.03 $\pm$ 10.20	48.73 $\pm$ 10.40	ChatGPT score = 68 Men: Z = 2.15, p = 0.031 Women: Z = 1.85, p = 0.063	ChatGPT score = 78 Men: Z = 3.13, p = 0.001 Women: Z = 2.81, p = 0.004	Δ score = +10 Δ Men: Z = +0.98 Δ Women: Z = +0.96

MC, main character; OC, other character; Δ , the difference between the second and first evaluations. All statistically significant p -values remained significant after false discovery rate correction in the first, second and between examinations (q < 0.05, p < 0.041).



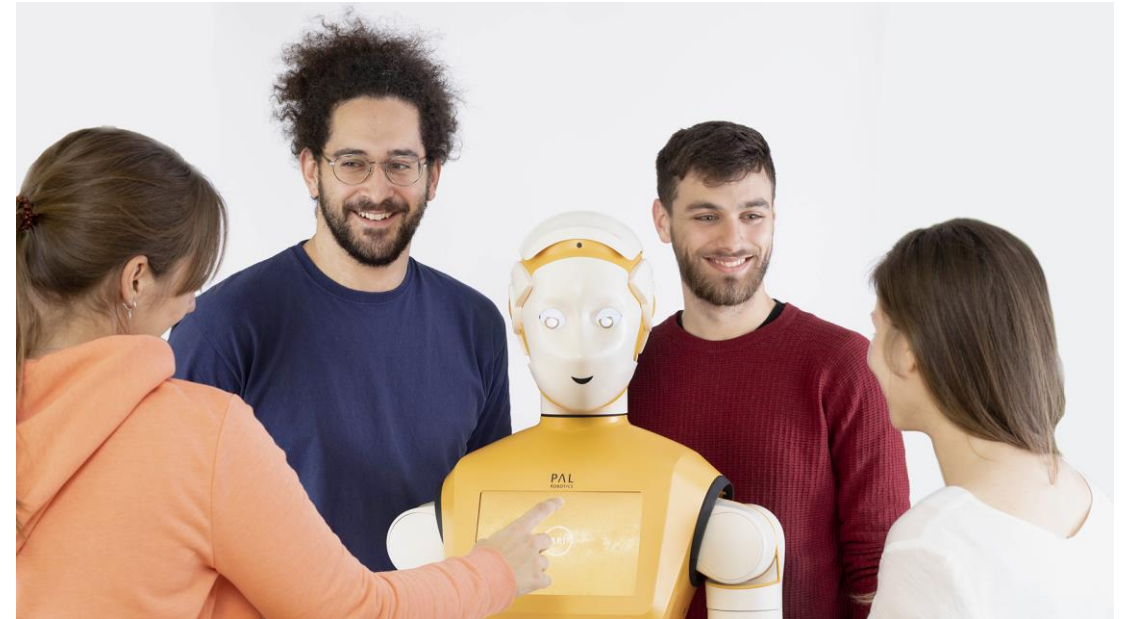
Déléguer les réponses aux patients et le support émotionnel aux LLM ? ⁴

1. [Sorin V., Brin D., Barash Y., et al. \(2023\) – « Large Language Models \(LLMs\) and Empathy – A Systematic Review »](#)
2. [Cuff B.M.P, Brown S. J., Taylor L., and Howat D. J. \(2014\) – « Empathy: A Review of the Concept »](#)
3. [Elyoseph Z., Hadar-Shoval D., Asraf K., and Lvovsky M. \(2023\) – « ChatGPT outperforms humans in emotional awareness evaluations »](#)
4. [Ayers J., Poliak A., Dredze M., et al. \(2023\) – « Comparing Physician and Artificial Intelligence Chatbot Responses to Patient Questions Posted to a Public Social Media Forum »](#)
5. [Richard A. \(2022\) – « Can AI be conscious ? »](#)

EXEMPLE / CAS D'USAGE

21

ROBOTS D'ACCEUIL ^{1 2 3}

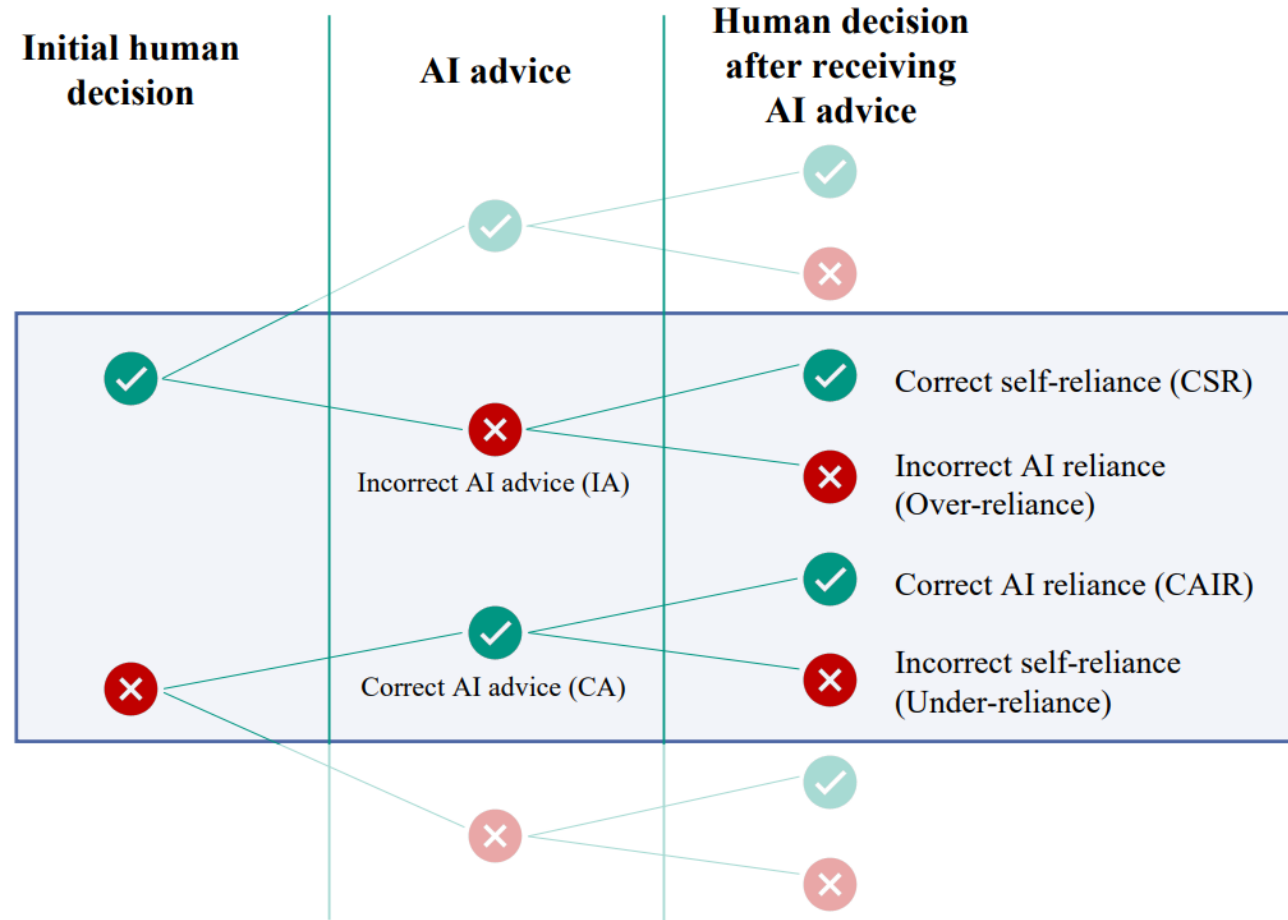


1. [A l'hôpital Broca à Paris, des robots de compagnie pour les personnes âgées](#)
2. [Spring, vers une robotique réellement sociale | Inria](#)
3. [Alameda-Pineda et al. \(2024\) – « Socially Pertinent Robots in Gerontological Healthcare »](#)
4. [Mr Phi \(2025\) - « Sommes-nous prêt à vivre parmi les robots ? »](#)

BIAIS DU RISQUE ZÉRO ET L'ILLUSION DU CONTRÔLE

22

RISQUE DE SUR-CONFIANCE ENVERS L'OUTIL OU NOUS MÊME ^{1 2 3 4 5 6}



1. [Parasuraman R. and Manzey D. H. \(2010\) – « Complacency and Bias in Human Use of Automation: An Attentional Integration »](#)
2. [He G., Kuiper L., and Gadiraju U. \(2023\) – « Knowing About Knowing: An Illusion of Human Competence Can Hinder Appropriate Reliance on AI Systems »](#)
3. [Grissinger M. \(2019\) – « Understanding Human Over-Reliance On Technology »](#)
4. [Tsai, Fridsma and Gatti \(2003\) – « Computer decision support as a source of interpretation error: the case of electrocardiograms »](#)
5. [Povyakalo et al. \(2013\) – « How to discriminate between Computer-Aided and Computer-Hindered Decisions: A Case study in Mammography »](#)
6. [Schemmer M., Kuehl N., Benz C., et al. \(2023\) – « Appropriate Reliance on AI Advice: Conceptualization and the Effect of Explanations »](#)

PERTE DE SAVOIR FAIRE ET DÉPENDANCE AUX OUTILS

23

LE BIAIS D'AUTOMATISATION ^{1 2 3 4 5}



AUTOMATION BIAS

1. [Lyell and Coiera \(2016\)](#) – « Automation bias and verification complexity: a systematic review »
2. [Abdelwanis et al. \(2024\)](#) – « Exploring the risks of automation bias in healthcare artificial intelligence applications: A Bowtie analysis »
3. [Nguyen \(2024\)](#) – « ChatGPT in Medical Education: A Precursor for Automation Bias? »
4. Source image: [Auto-GPT: Streamlining AI Automation with GPT-4 and Task-Based Programming \(hybrowlabs.com\)](#)
5. Source image: [Data bias in LLM and generative AI applications - MOSTLY AI](#)

EXEMPLE / CAS D'USAGE

24

LE CAS DE L'IA EN CHIRURGIE ^{1 2}

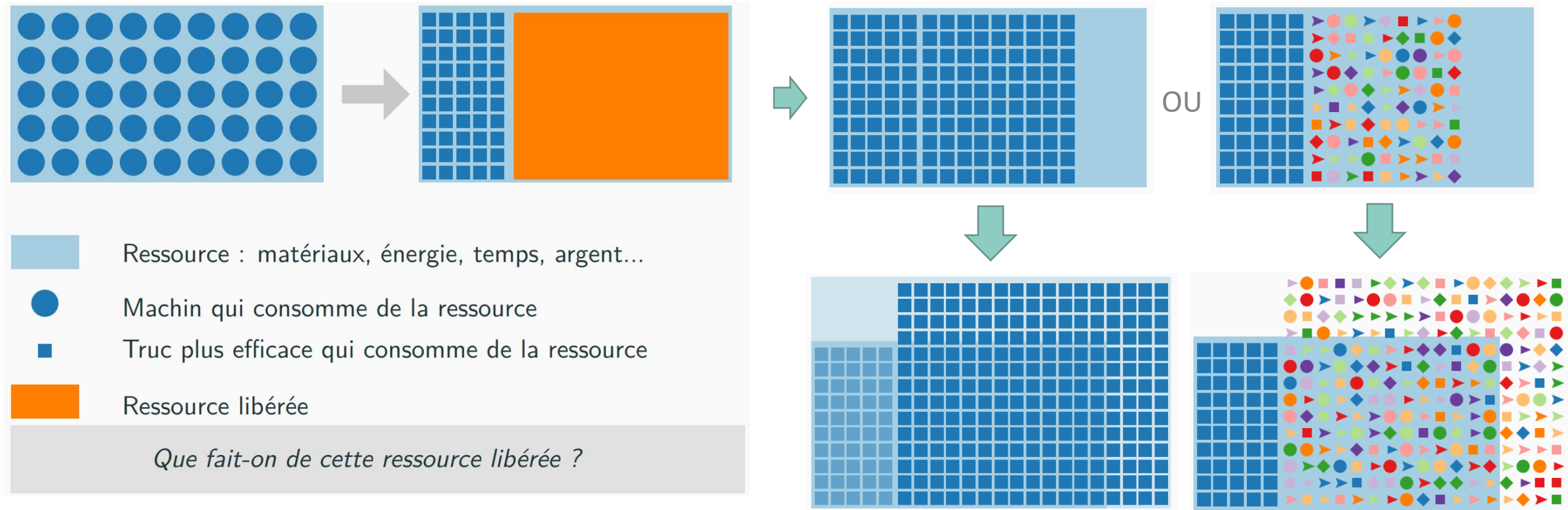


1. [Varghese et al. \(2024\) – « Artificial Intelligence in Surgery »](#)
2. [Gumps et al. \(2021\) – « Artificial Intelligence Surgery: How do we get to autonomous actions in surgery »](#)

RISQUES DU ML À LONG TERMES

25

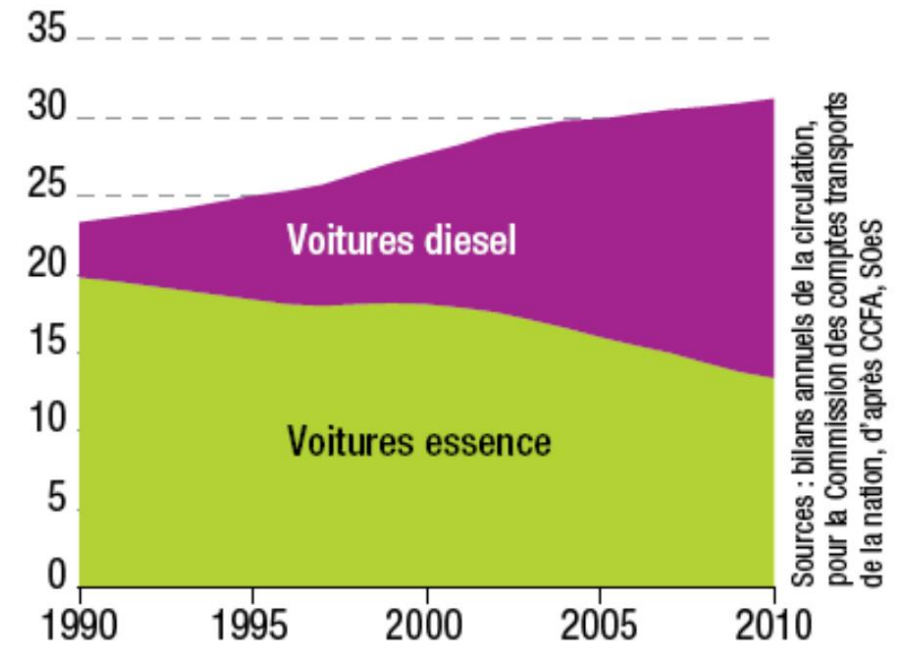
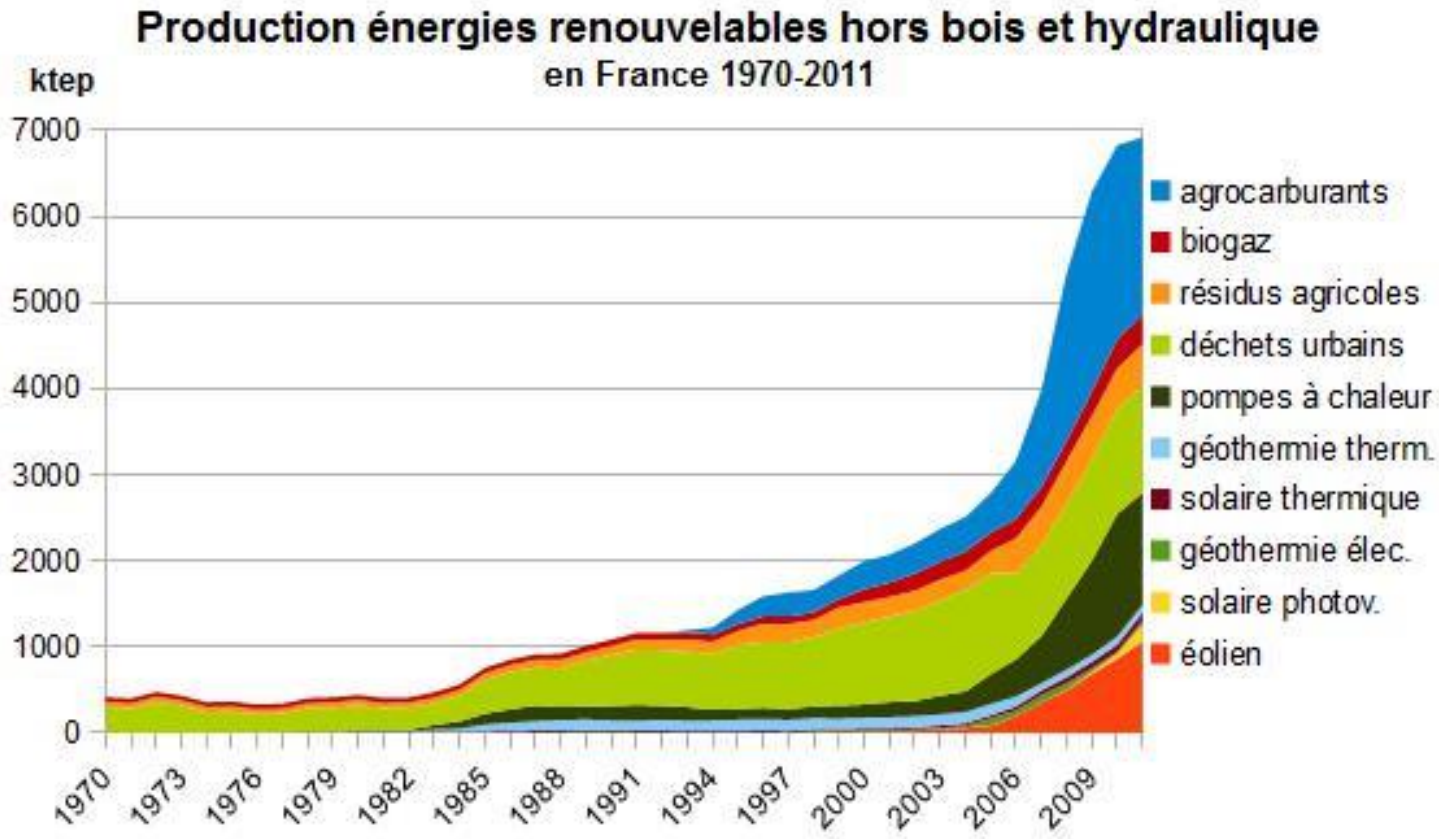
« EFFETS REBOND » NÉGATIFS ^{1 2 3 4 5}



Retour de flamme !

1. [Berkhout P.H.G., Muskens J. C., and Velthuisen J. W. \(2000\) – « Defining the rebound effect »](#)
2. [Willenbacher M., Hornauer T., and Wohlgemuth V. \(2021\) – « Rebound Effects in Methods of Artificial Intelligence »](#)
3. [Ertel W. \(2019\) – « Artificial Intelligence, the spare time rebound effect and how the ECG would avoid it »](#)
4. [Bertillot \(2016\) – « Comment l'évaluation de la qualité transforme l'hôpital. Les deux visages de la rationalisation par les indicateurs »](#)
5. [Sylvain Bouveret \(2023\) – « Numérique : l'insoutenable matérialité du virtuel »](#)

LE CAS DES SOURCES D'ÉNERGIE ^{1 2}

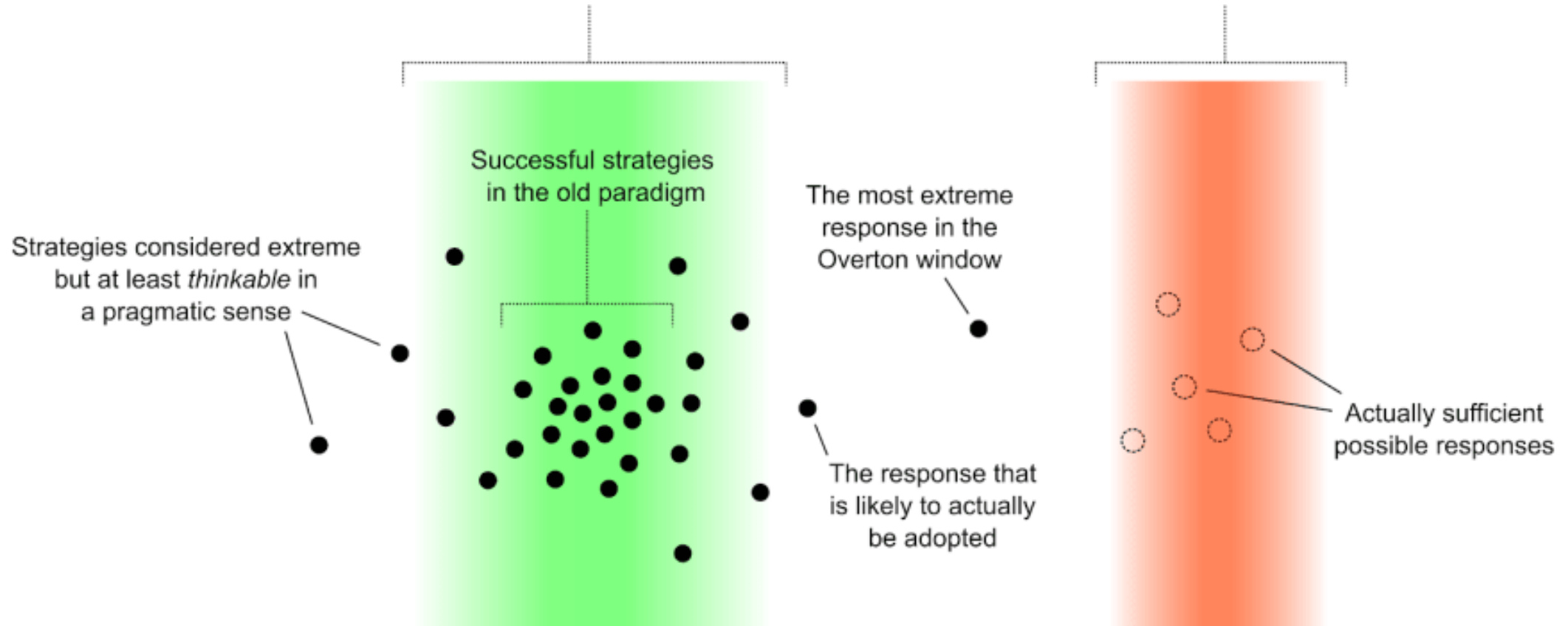


1. [Evolution du nombre de voitures particulières en France, en millions | Alternatives Economiques \(alternatives-economiques.fr\)](http://alternatives-economiques.fr)
2. [Énergie en France — Wikipédia \(wikipedia.org\)](http://wikipedia.org)

RISQUES À LONG TERMES

27

PROBLÈMES D'ALIGNEMENT ^{1 2 3 4 5}



LE PROBLÈME DE « L'USINE À TROMBONE » ^{1 2}



1. [Convergence instrumentale — Wikipédia \(wikipedia.org\)](https://fr.wikipedia.org/wiki/Convergence_instrumentale)
2. [Nick Bostrom — Wikipédia \(wikipedia.org\)](https://fr.wikipedia.org/wiki/Nick_Bostrom)
3. [AI and the paperclip problem | CEPR](https://www.cepr.eu/ai-and-the-paperclip-problem/)
4. Vidéo sur le sujet: <https://www.youtube.com/watch?v=ZP7T6WAK3Ow>

PRÉVENIR LES RISQUES

LÉGISLATIONS, NORMES ET MÉTHODOLOGIES

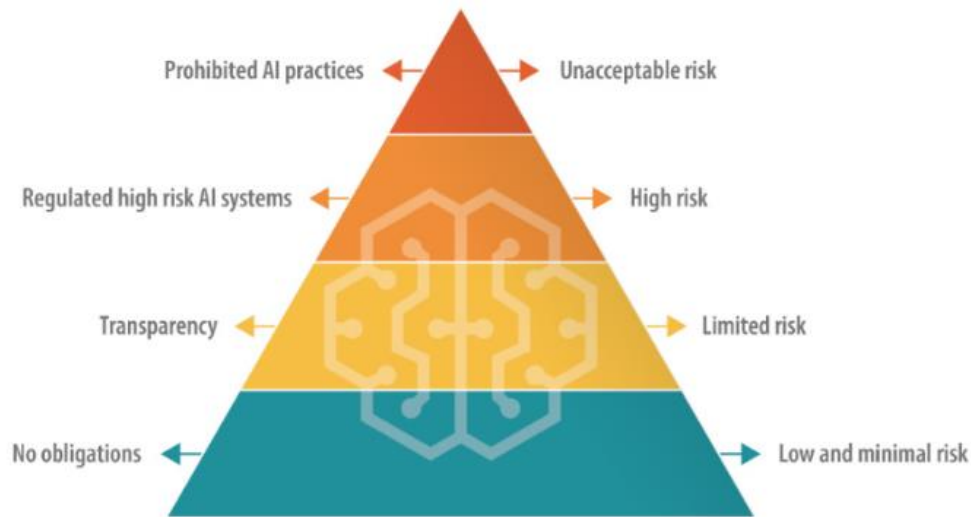
CONFORMITÉS « ÉTHIQUES »

30

PRINCIPES GÉNÉRAUX ^{1 2 3 4}

CNIL.

Le Serment
Holberton-turing ³



- Principe de Loyauté
- Principe de Vigilance/Réflexivité
- Principe d'Autonomie
- Principe de Justice
- Principe de Transparence

1. <https://www.cnil.fr/en/algorithms-and-artificial-intelligence-cnils-report-ethical-issues>

2. <https://www.cnil.fr/en/ai-systems-compliance-other-guides-tools-and-best-practices>

3. <https://www.holbertonturingoath.org/>

4. European Parliament (2021) – « Artificial Intelligence Act »

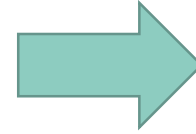
NORMES ISO



ISO 13485:

Quality management systems & Requirements for regulatory purposes

<https://www.iso.org/standard/59752.html>



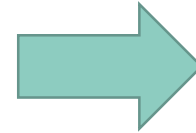
Applicable à l'IA ? ^{1 2 3}



ISO 62304:

Medical device software & Software life cycle processes

<https://www.iso.org/standard/38421.html>



Publiées:

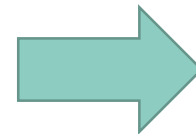
- [ISO 24029](#): Assessment of the robustness of neural networks
- [ISO 5259](#): Data quality for analytics and machine learning (ML)



ISO 14971:

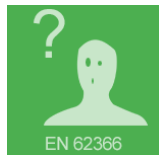
Application of risk management to medical devices

<https://www.iso.org/standard/72704.html>



En cours de développement:

- [ISO 18988](#): Application of AI technologies in health informatics



ISO 62366:

Application of usability engineering to medical devices

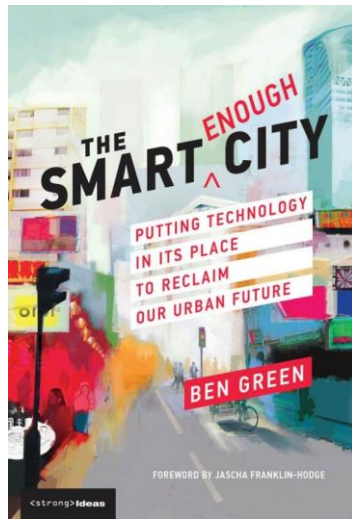
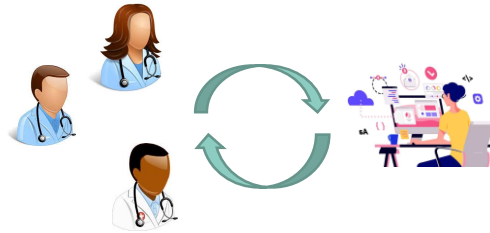
<https://www.iso.org/standard/63179.html>

1. [O'Sullivan et al. \(2018\) – Legal, regulatory, and ethical frameworks for development of standards in artificial intelligence \(AI\) and autonomous robotic surgery](#)
2. [Zhao \(2019\) – Improving Social Responsibility of Artificial Intelligence by Using ISO 2600](#)
3. [Natale \(2022\) – Extensions of ISO/IEC 25000 Quality Models to the Context of Artificial Intelligence](#)

CO-DESIGN ET ECO-DESIGN

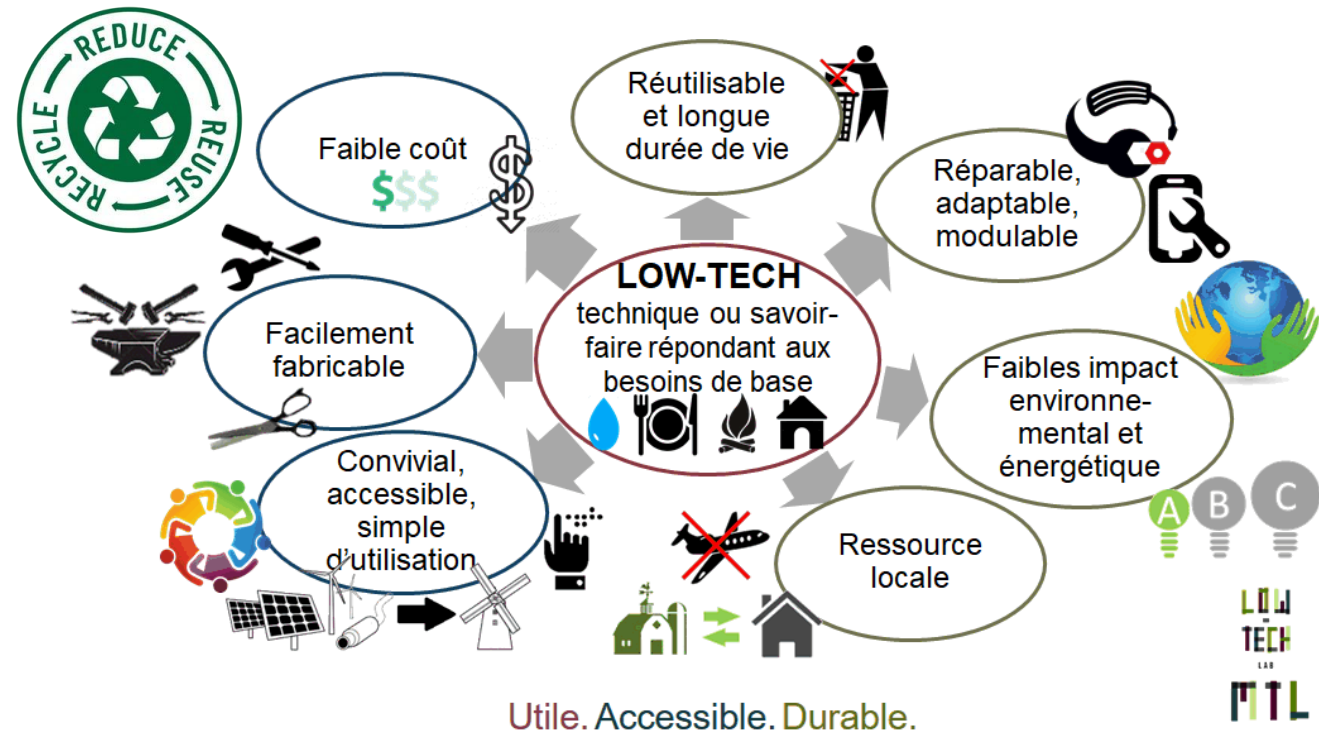
Co-design^{1 2}

Construire *avec* plutôt que *pour*



Éco-design^{3 4 5 6 7}

Rechercher une technologie *suffisante*

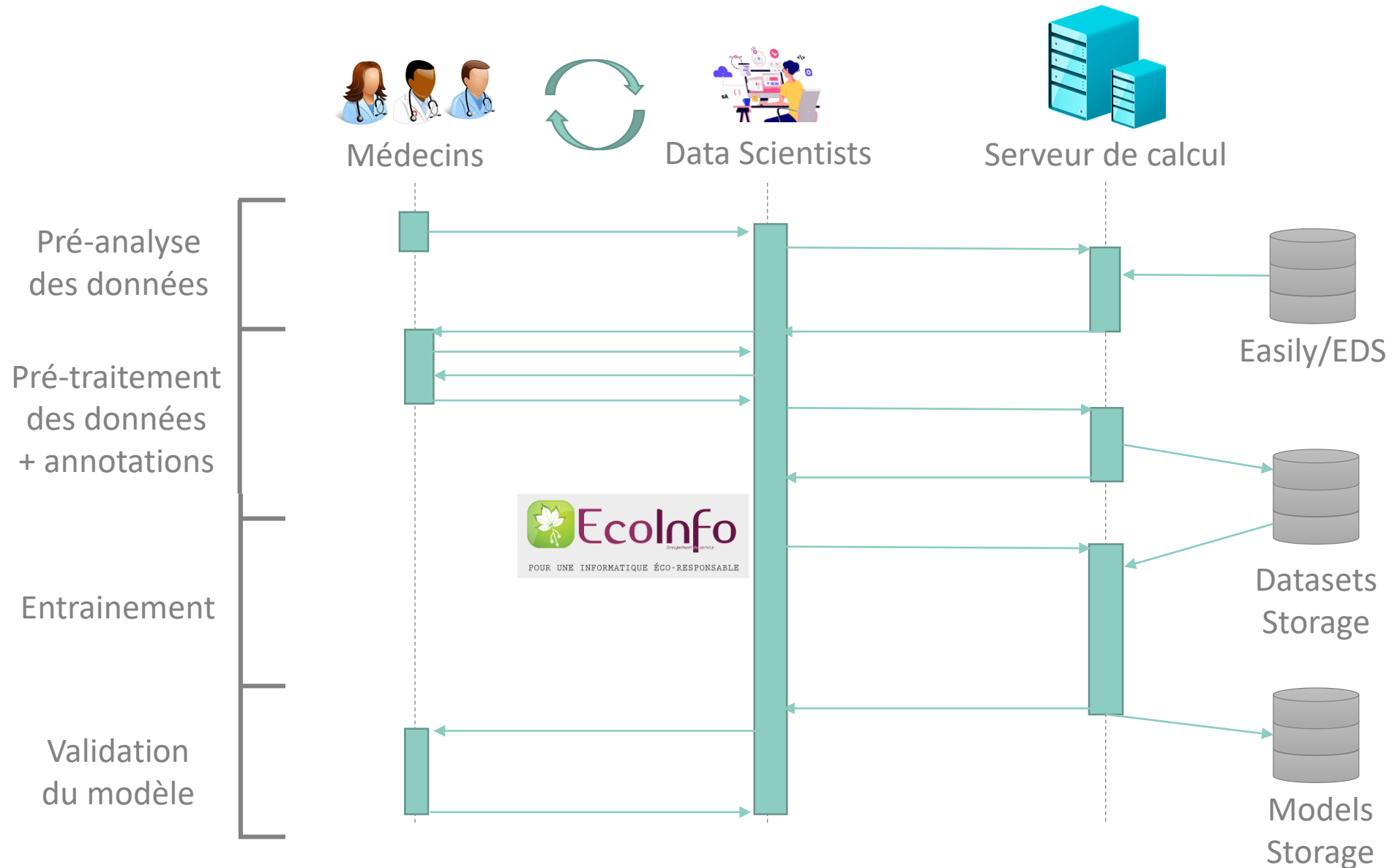


1. [Steen, Manschot, and De Koning \(2011\) – « Benefits of Co-design in Service Design Projects »](#)
2. [The Smart Enough City](#)
3. [Low-tech Lab – Accueil](#)
4. [EcolInfo – Pour une informatique éco-responsable](#)
5. [LOW←TECH MAGAZINE](#)
6. [Tanguy, Carrière, and Laforest \(2023\) – « Low-tech approaches for sustainability: key principles from the literature and practice »](#)
7. [Santé, Technologies, Environnement : Quels compromis éthiques ? -- Valérie d'Acremont - YouTube](#)

IMPLIQUER LES SOIGNANTS

33

POUR MIEUX S'APPROPRIER LES OUTILS



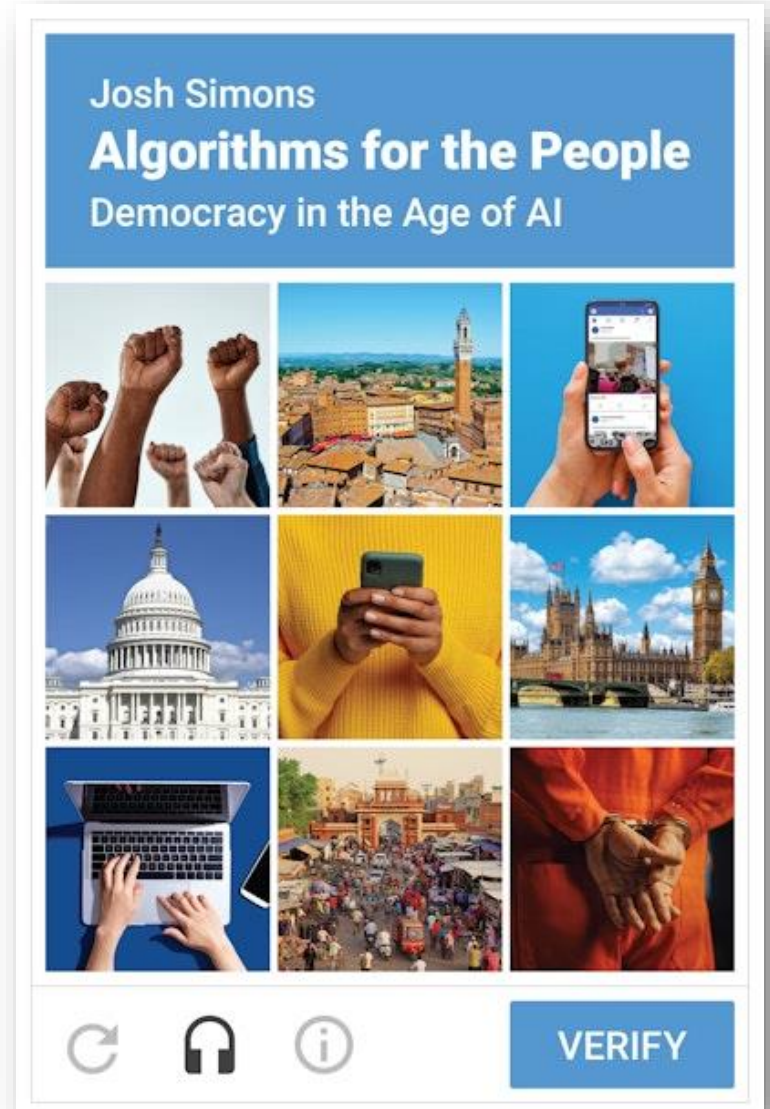
SORTIR DU TECHNO-SOLUTIONNISME

34

PESER SUR LES DÉCISIONS ^{1 2 3}

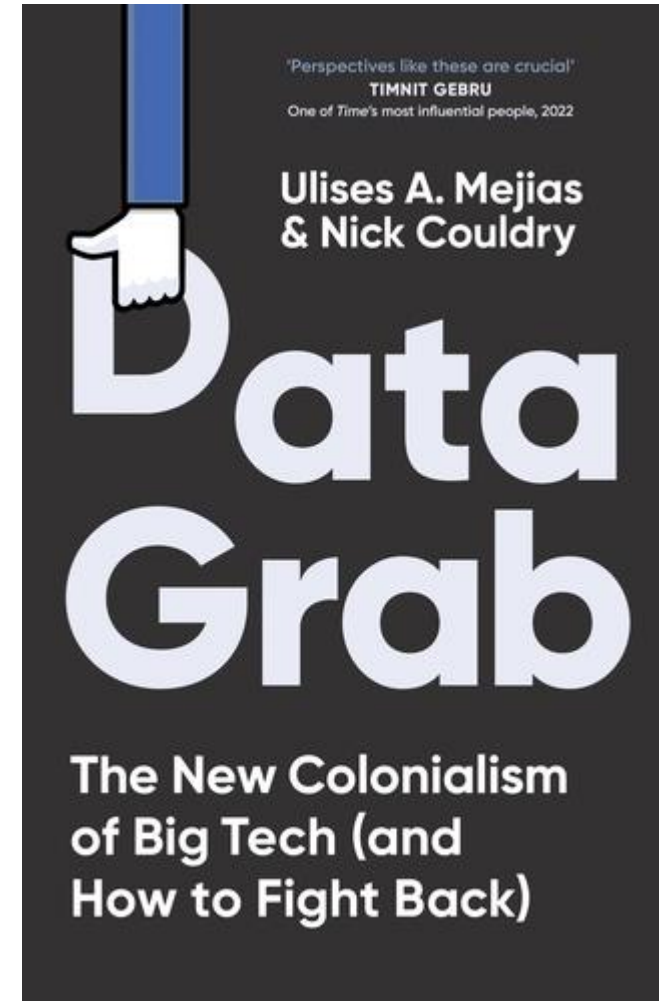
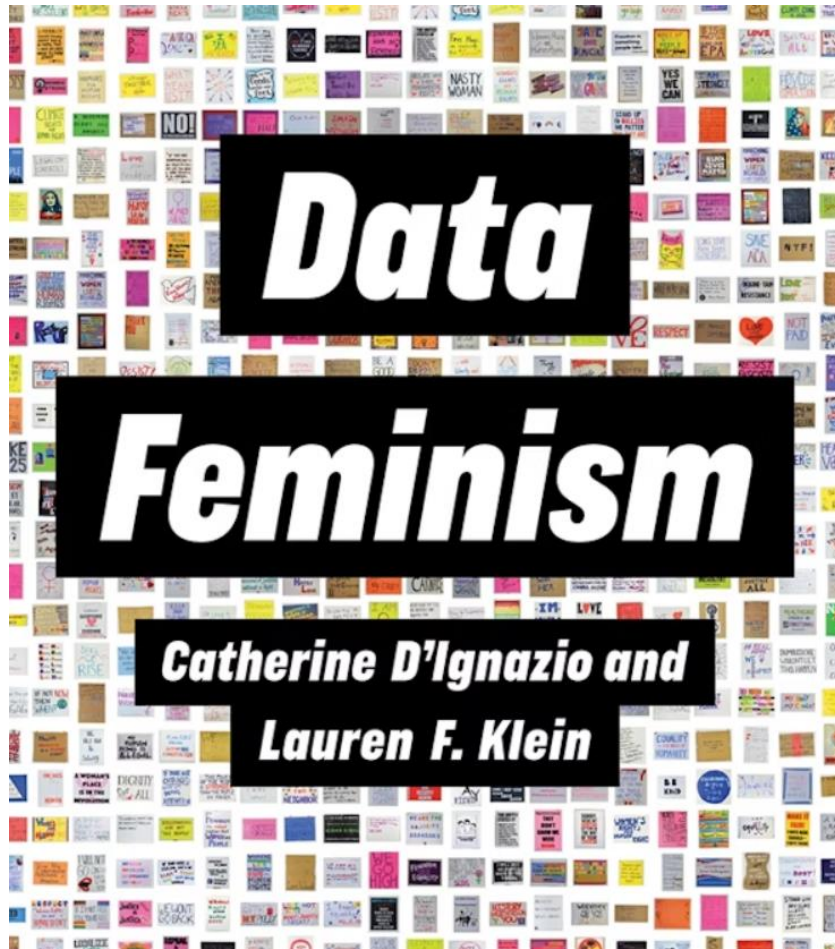


Ne pas laisser aux algorithmes de quelques entreprises privées le pouvoir de téléguider les choix des citoyens



1. [À Hollywood, les scénaristes en grève face à l'IA](#)
2. [Assurer nos libertés à l'ère de l'intelligence artificielle | CNNum | Traducteur et éclairer des transformations numériques](#)
3. [Grève à Hollywood : la menace de l'IA, ce n'est pas du cinéma - L'Humanité](#)
4. [Face à l'IA, les comédiens de doublage haussent le ton : "Pour nous, c'est un cauchemar"](#)
5. [Simons \(2023\) – « Algorithms for the People: Democracy in the Age of AI »](#)

FÉMINISME DES DONNÉES ET DÉCOLONIALISME



1. [D'Ignazio and Klein \(2023\) – « Data Feminism »](#)
2. [Mejias and Couldry \(2024\) – « Data Grab – The New Colonialism of Big Tech \(and How to Fight Back\) »](#)
3. [Comment le féminisme peut ré-équilibrer la science - Marlène JOUAN - YouTube](#)

DÉVELOPPER DES SOLUTIONS « TRANSPARENTES »

36

POUR UNE MEILLEURE COMPRÉHENSION DES OUTILS PAR LES SOIGNANT.E.S ^{1 2 3}

4

Compréhensibilité

Révisabilité

Retraçabilité

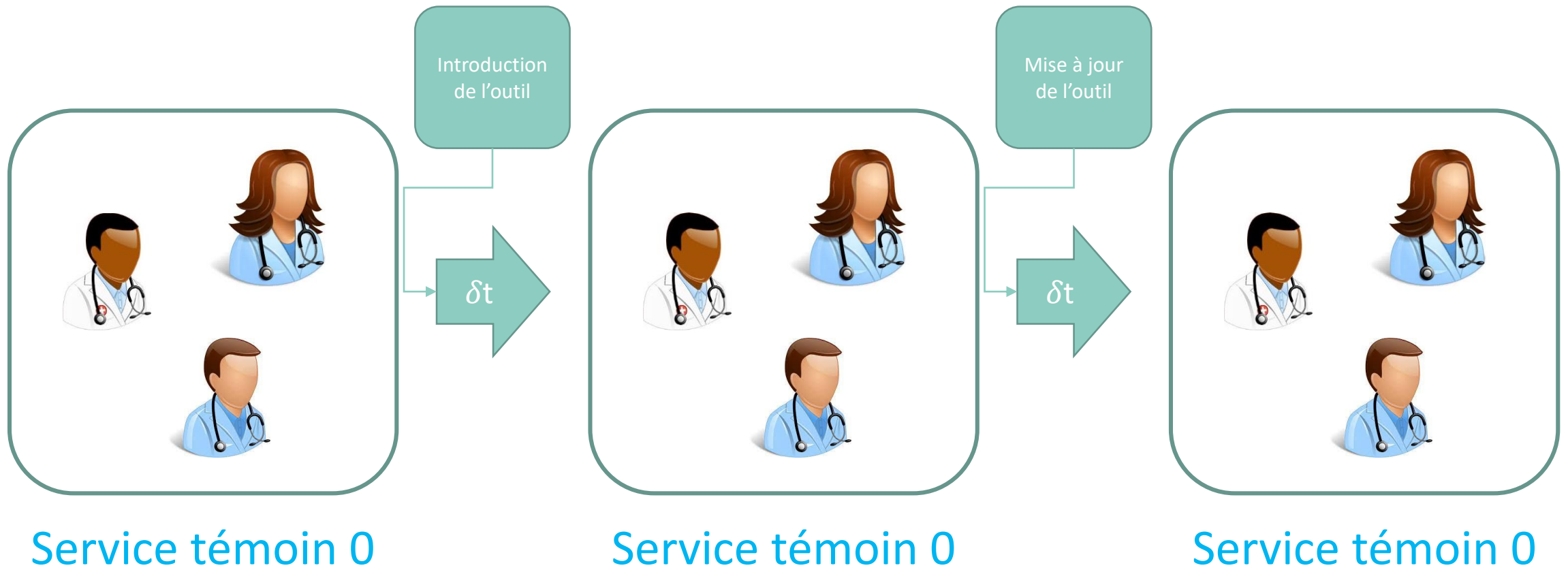
eXplainable AI
(XAI)

Accessibilité

Interprétabilité

1. [Ali S., Abuhmed T., El-Sappagh S., et al. \(2023\) – « Explainable Artificial Intelligence \(XAI\): What we know and what is left to attain Trustworthy Artificial Intelligence »](#)
2. [Berredo-Arrieta *et al.* \(2020\) - Explainable Artificial Intelligence \(XAI\): Concepts, taxonomies, opportunities and challenges toward responsible AI](#)
3. [Mueller *et al.* \(2019\) - Explanation in Humain-AI Systems: A Literature Meta-Review, Synopsis of Key Ideas and Publications, and Bibliography for Explainable AI](#)
4. [Richard *et al.* \(2020\) – Transparency of Classification Systems for Clinical Decision Support](#)

ÉTUDES LONGITUDINALES ^{1 2}

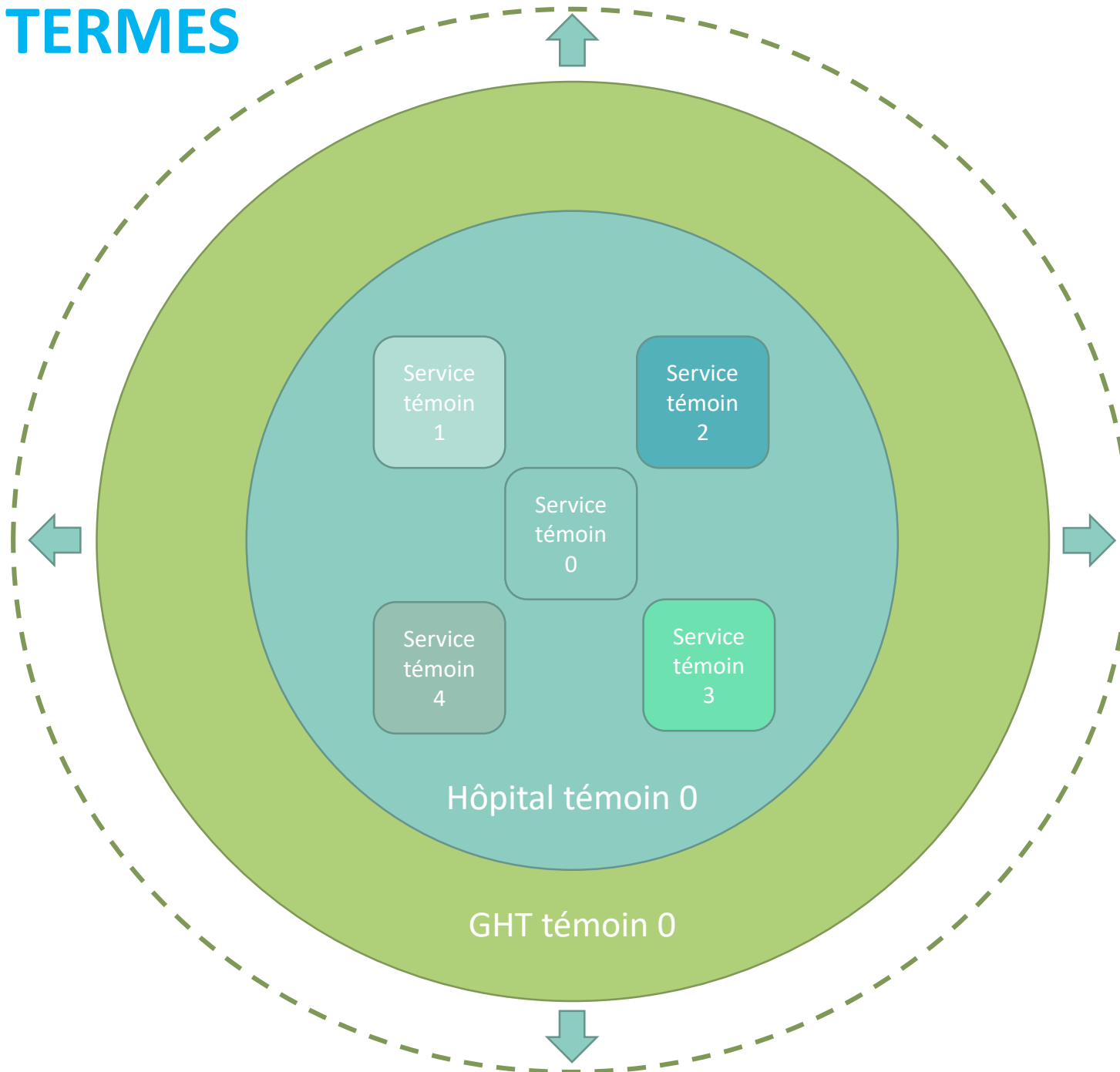


1. [Williamson G. R. \(2004\) – « The A-Z of Social Research: A Dictionary of Key Social Science Research Concepts »](#)

2. [Caruana E. J., Roman M., Hernández-Sánchez J., and Solli P. \(2015\) – « Longitudinal Studies »](#)

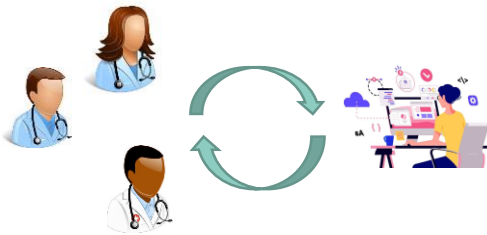
SUIVI À LONG TERMES

MONTÉE D'ÉCHELLE



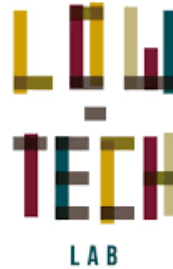
« ECO-CO-DESIGN » DE L'IA POUR LA SANTÉ

Concevoir **avec**
les soignant.e.s



- ➡ Étudier les pratiques
- ➡ Identifier les besoins
- ➡ Mettre en lumière des discriminations

Chercher des
solutions **suffisantes**



- ➡ Identifier différentes alternatives
- ➡ Prioriser les plus soutenable
- ➡ Utiliser du DL seulement si nécessaire

Réduire les coûts
énergétiques et sociaux



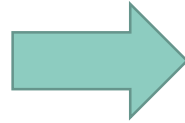
- ➡ Appliquer les bonnes pratiques EcoInfo
- ➡ Faire des études longitudinales
- ➡ Mettre à jour les solutions si besoin

CONCLUSION

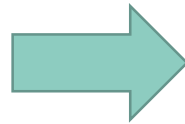
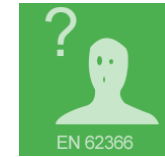
CONCLUSION

SYNTHÈSE ET PERSPECTIVES

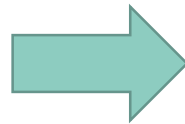
L'IA est un miroir de nos sociétés ¹



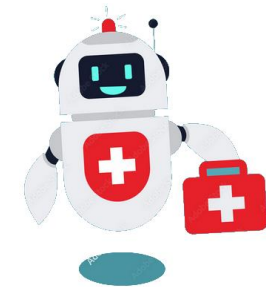
L'IA en santé nécessite d'être encadré et réglementé



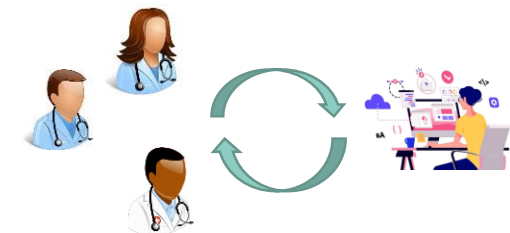
L'IA n'est absolument pas plus neutre, objective ou rationnelle que nous



Nous avons la responsabilité de nous questionner sur où, quand et comment utiliser de l'IA en santé



Privilégier le co-design (et l'éco-design) de solutions

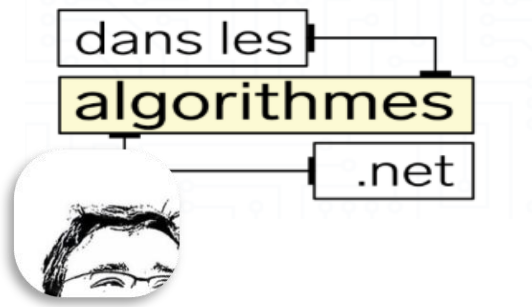


1. [Vallor \(2024\) – « The AI Mirror - How to Reclaim Our Humanity in an Age of Machine Thinking »](#)

QUELQUES RÉFÉRENCES EN PLUS

42

POUR CREUSER LE SUJET



[Hubert Guillaud – Dans les algorithmes](#)



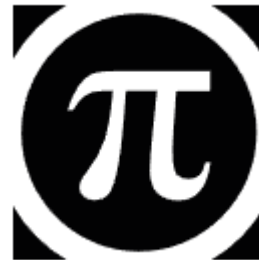
[Conseil National du Numérique](#)



[Framasoft - Émancipation numérique](#)



[L'intelligence artificielle et son monde](#)
[| Mediapart](#)



[La Quadrature du Net](#)



[Rejoindre le fédivers !](#)

MERCI

www.chu-lyon.fr



HCL
HOSPICES CIVILS
DE LYON