

Flanders AI forum

Antwerpen 7 juni 2022

.AGORIA
Vlaanderen

**VO
KA** Vlaams
netwerk van
ondernemingen



DEPARTEMENT
ECONOMIE
WETENSCHAP &
INNOVATIE

imec
embracing a better life

AGENTSCHAP
INNOVEREN &
ONDERNEMEN



Vlaanderen
is ondernemen



Vlaams AI-
Onderzoeksprogramma

Explaining and Interpreting AI models



Yvan Saeys

yvan.saeys@ugent.be



@Saeyslab

Flanders AI Forum 2022

What do we mean by explainable
/ interpretable ML ?

TRUST

Interpretability

Explainability

Adversarial
examples

Transparency

Usability

Interestingness

Simulatability

Justifiability

Acceptability

Robust ML

Understandability

Comprehensibility

Mental fit

Fairness

Responsible AI

Transparency

Informally, transparency is the opposite of opacity or “black-boxness”. It connotes some sense of understanding the mechanism by which the model works. Transparency is considered here at the level of:

- 1. the entire model (simulatability)*
- 2. the level of individual components such as parameters (decomposability)*
- 3. the level of the training algorithm (algorithmic transparency).*

Interpretability

“Interpretability is the degree to which a human can understand the **cause of a decision**”

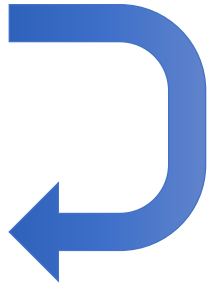
Miller, 2017

- An interpretation is the mapping of an abstract concept (e.g., a predicted class) into a domain that the human can make sense of
- Present some of the properties of an AI model in terms understandable to a human
 - Can we understand what the AI algorithm bases its decision on?
- Note that, in contrast to transparency, to achieve interpretability the data is always involved.

Explainability

- Comes more from the domain of social sciences
- No concise definition yet, as explanations can differ in completeness or the degree of causality
- An explanation is the collection of features of the *interpretable domain*, that have contributed for a given example to produce a decision

Concepts expressed in the understandable terms composing an explanation are self-contained and do not need further explanations



Interpretability versus explainability

- Often, interpretability is related to the whole model
- Explainability is more related to individual predictions of the model
- All in the end should lead to more trustworthy ML

What is an explanation?

An explanation is the answer to a question

- Why is that particular decision made?
 - Why did not the treatment work on the patient?
 - Why was my loan rejected?
- Why is the decision for sample A different to (the nearby) sample B?
- Which features are important for the decision about sample A ?
- Which features are important globally ?
- All depends on the goal, context and background of the ML user

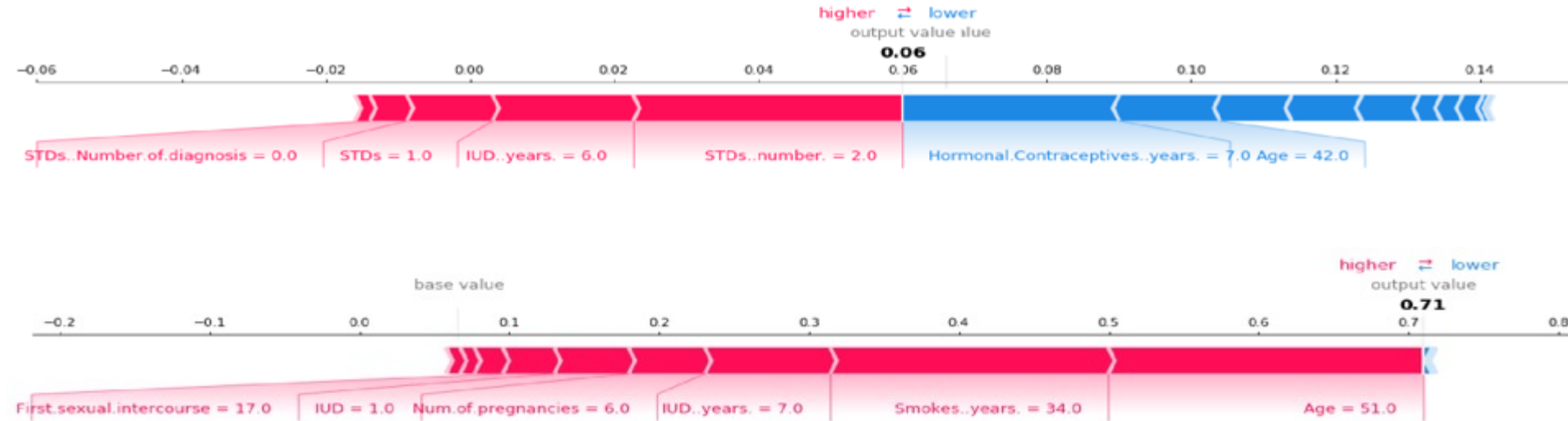
What makes a **good** explanation ?

Lessons from the social sciences

Good explanations are:

- *contrastive*
- *selected*
- *social*
- *truthful (fidelity)*
- *consistent with prior belief*
- *general and probable*

Example: SHAP values to explain the predicted cancer probabilities of two patients



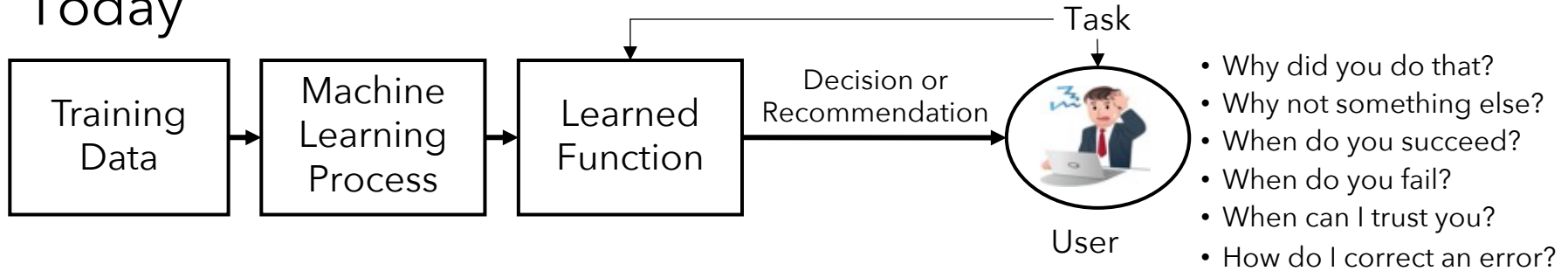
Why do we need explainable /
interpretable ML methods?

Motivations for interpretable ML

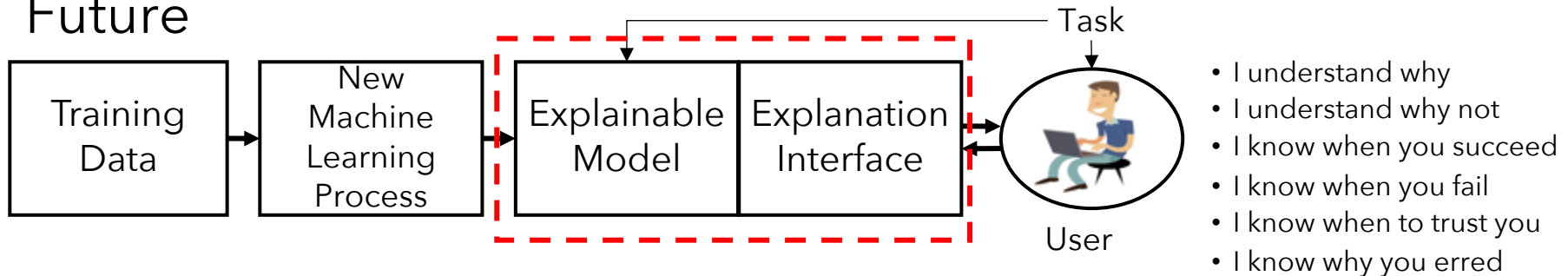
- To justify decisions/build trust
 - To accept a predictive model, medical experts have to understand the intelligence behind the diagnostic
- To validate/verify/understand and improve models
 - Detection of anomalies in the model, study bias/fairness
- To derive new knowledge
 - The model itself is a new source of knowledge
 - Different models provide different insights in the data
- To complement other metrics (typically predictive performance of the model) and thus get a broader view on the model
 - Others include privacy, robustness, causality, ...
- Legal/regulatory obligations

Towards explanation interfaces

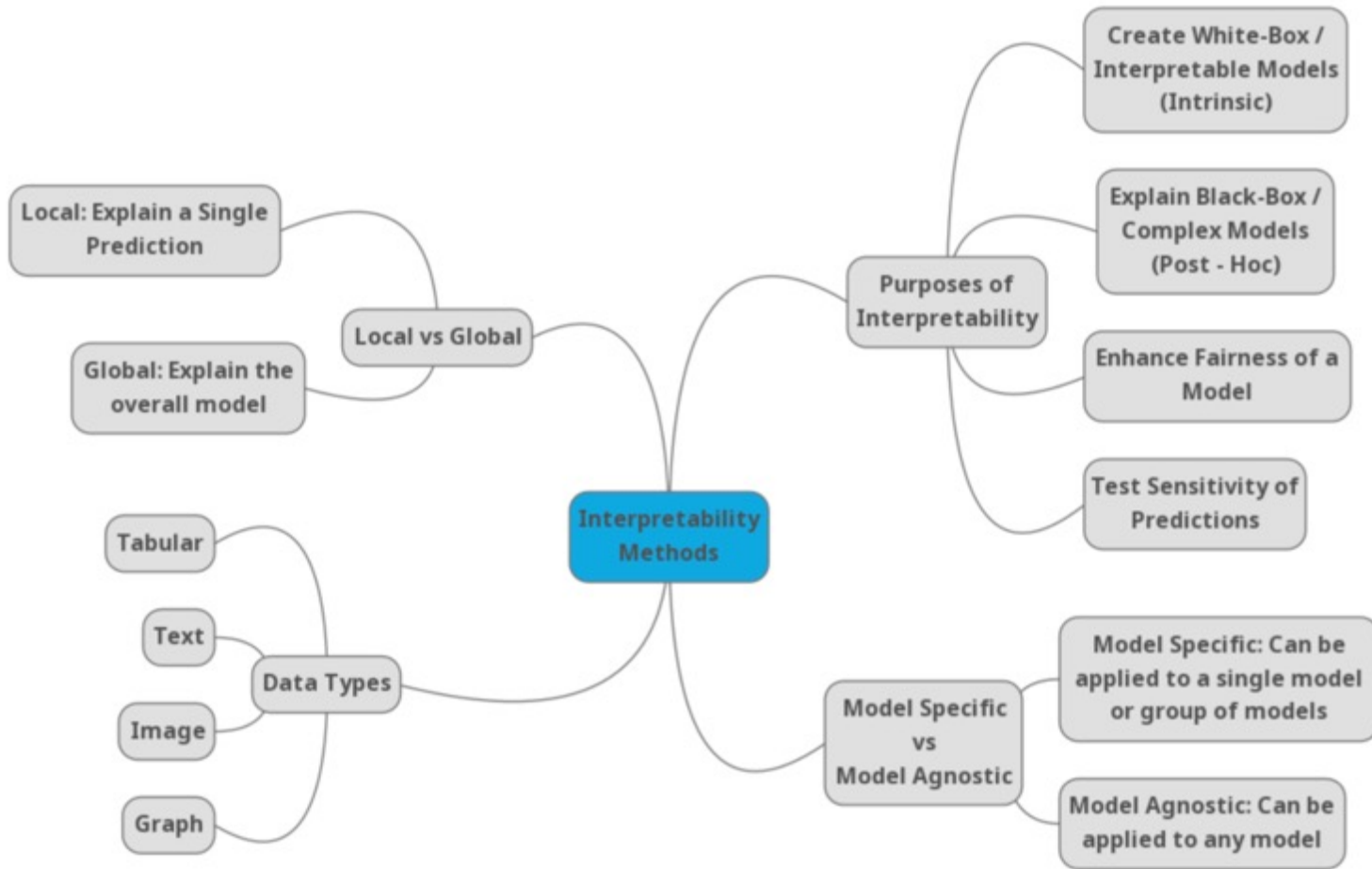
Today



Future

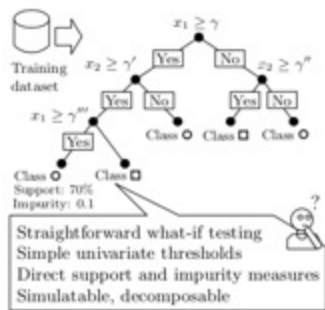


What kind of interpretability methods do exist ?

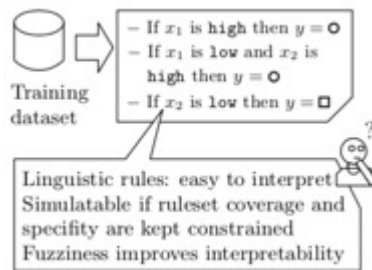


Inherently interpretable models

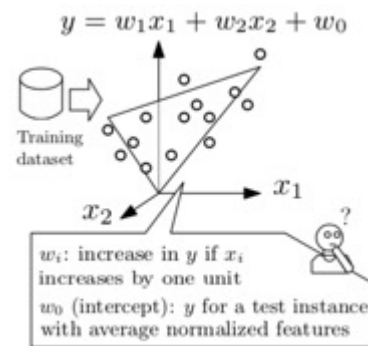
Decision trees



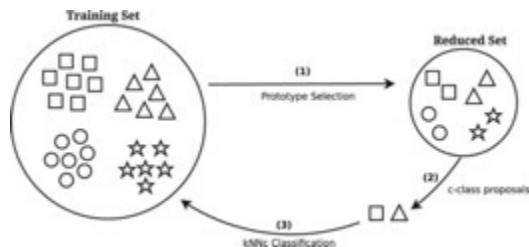
Rule based models



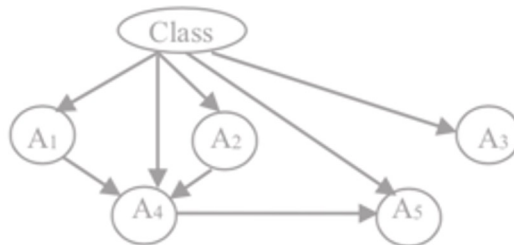
Linear models



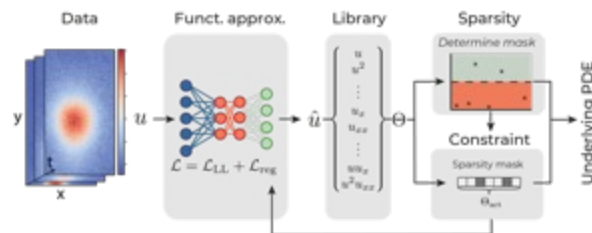
Case based reasoning



Graphical models

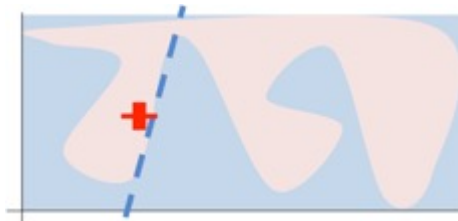


Physics inspired models



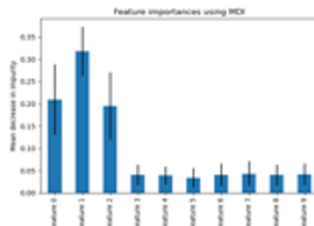
Post-hoc interpretability

Surrogate models



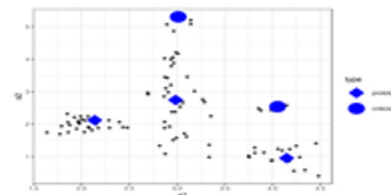
Knowledge distillation, LIME,...

Feature importance scores



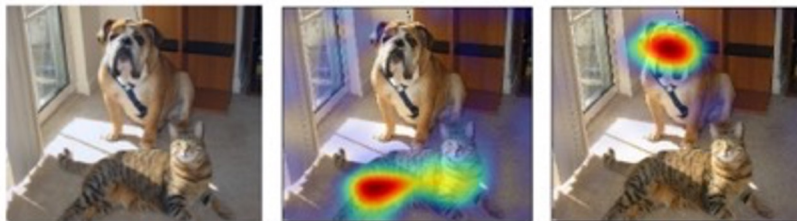
Permutation importance, PDP, ICE, ALE, Shapley values

Explanation by example



Counterfactuals, algorithmic recourse, prototypes and criticisms

Attribution maps / Visualizations



Original Image

Grad-CAM 'Cat'

Grad-CAM 'Dog'

CAM, GRAD-CAM, DeepLift, IG, SmoothGrad, TCAV

Textual explanations

<i>Q: Is this a healthy meal?</i>	Textual Justification	Visual Pointing
	➡ A: No ...because it is a hot dog with a lot of toppings.	
	➡ A: Yes ...because it contains a variety of vegetables on the table.	

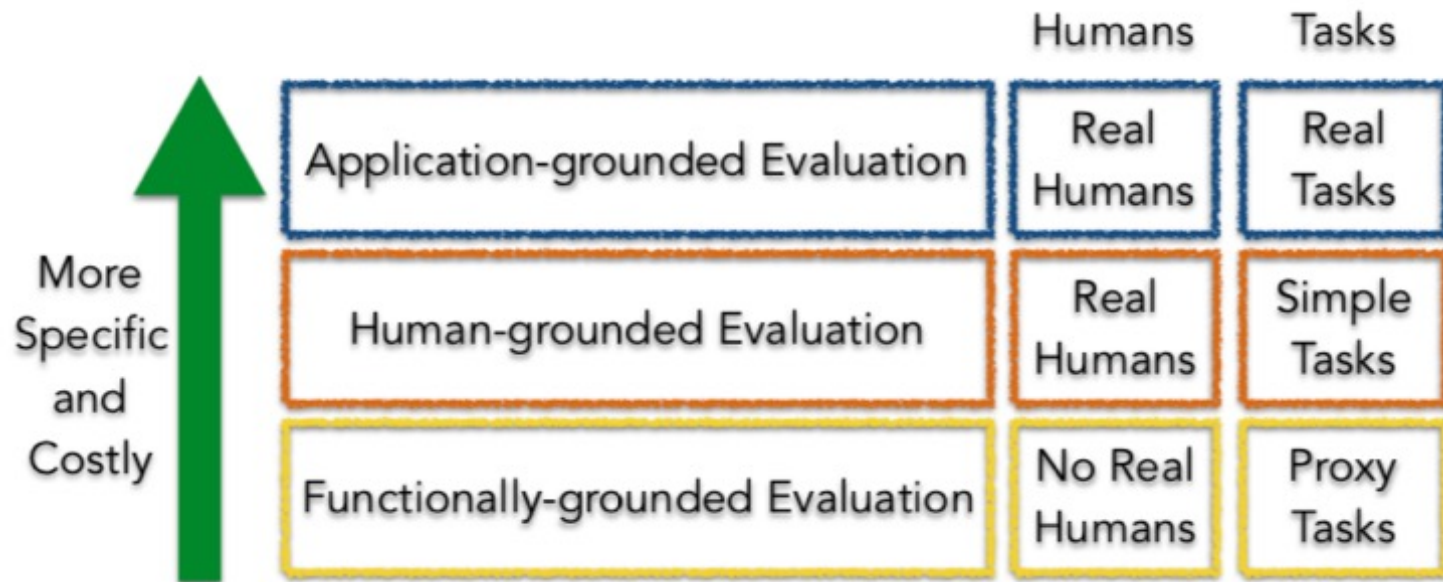
Future prospects

Open problems

Early work in evaluating explanations assumed that there was a single “correct explanation”. Today it is clear that many different explanations can be valid at the same time. This leaves some questions:

- Which types of explanations can be considered “valid”?
- What are the *specific* questions that each explanation provides an answer to?
- How do we evaluate the quality of each type of explanation?

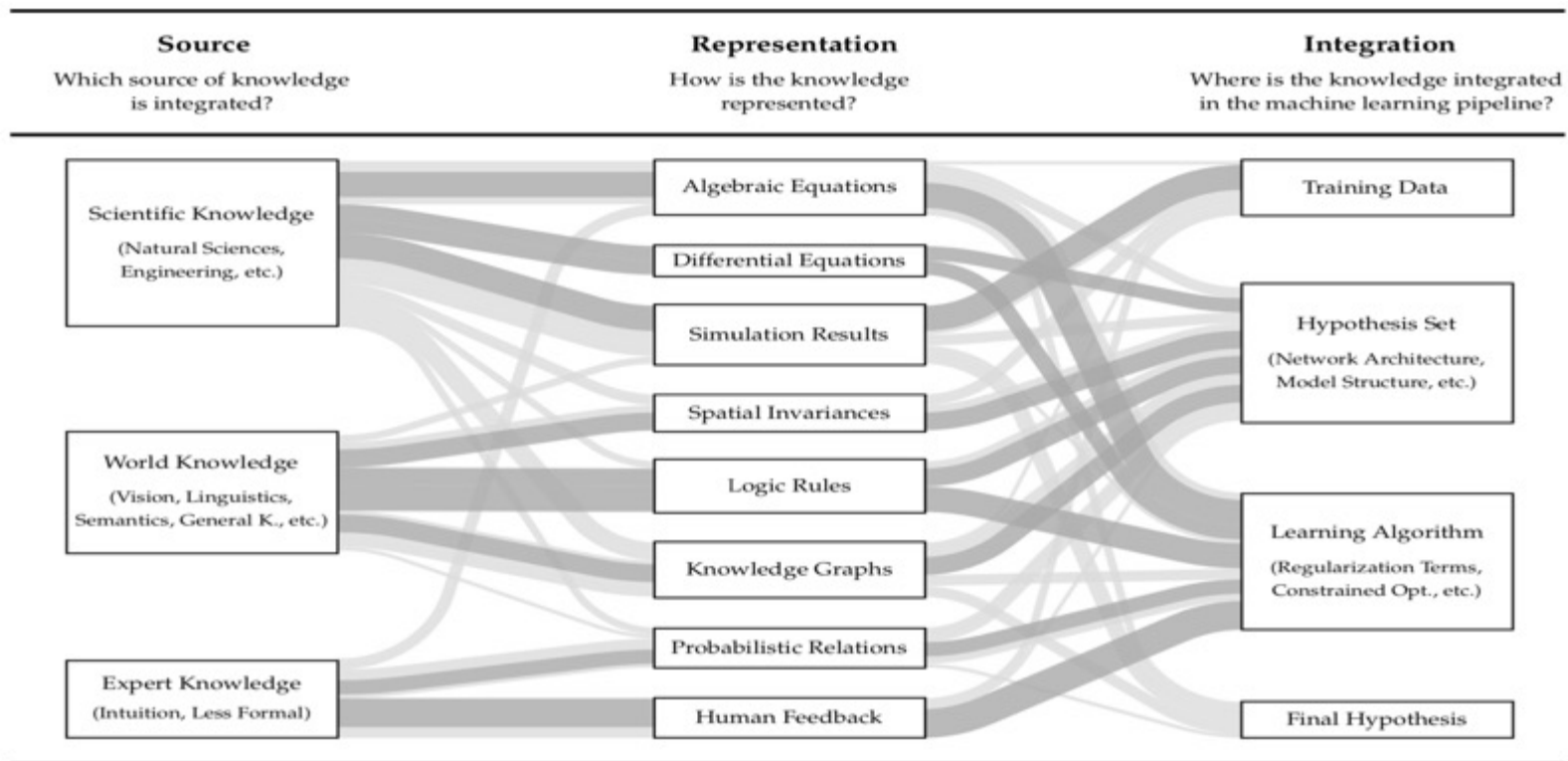
Different levels of evaluating interpretability



Properties of explanations

- **Fidelity:** How well does the explanation represent the “true workings” of the model?
- **Consistency:** How similar are explanations for models that were trained on the same task and produce similar predictions?
- **Stability:** How similar are explanations for similar instances?
- **Comprehensibility:** How well do humans understand the explanation?
 - Note: this depends on the audience
- ...

Future directions: Informed ML



Summary

- Explainable AI methods are a key component in the toolbox of ethical AI
- The field of explainable AI methods is rapidly expanding in multiple directions
 - Few standard definitions exist, multiple points of view
- Methods should be developed in order to align with human values and preferences
 - These values/preferences can be complex and fluid



**GHENT
UNIVERSITY**

ETHIEK IS NIET MOEILIJK *

Prof. Dr. Ir. Tijl De Bie
AIDA-IDLab, Ghent University

WAT IS EEN VERKLARING?



Aristotelianse 'oorzaken'

Waarom is dit wiel rond?

- Efficiënt (wie maakte het)
- Teleologisch (waarvoor dient het)
- Mechanistisch (hoe is het gemaakt)
- Categorisch (wat voor ding is het)

WAT IS EEN VERKLARING?



Aristotelianse 'oorzaken'

Waarom is dit wiel rond?

- Efficiënt (wie maakte het)
- Teleologisch (waarvoor dient het)
- Mechanistisch (hoe is het gemaakt)
- Categorisch (wat voor ding is het)

Deductief:

Waarom bent u in AI geïnteresseerd?

Klopt voor alle aanwezigen

U bent aanwezig

→ Klopt voor u

WAT IS EEN VERKLARING?



Aristotelianse 'oorzaken'

Waarom is dit wiel rond?

- Efficiënt (wie maakte het)
- Teleologisch (waarvoor dient het)
- Mechanistisch (hoe is het gemaakt)
- Categorisch (wat voor ding is het)



Deductief:

Waarom bent u in AI geïnteresseerd?

Klopt voor alle aanwezigen

U bent aanwezig

→ Klopt voor u

Statistische relevantie:

$P(\text{AI interesse} \mid \text{aanwezig})$

$? = P(\text{AI interesse} \mid \text{aanwezig, stoel})$

WAT IS EEN VERKLARING?



Aristotelianse 'oorzaken'

Waarom is dit wiel rond?

- Efficiënt (wie maakte het)
- Teleologisch (waarvoor dient het)
- Mechanistisch (hoe is het gemaakt)
- Categorisch (wat voor ding is het)



Deductief:

Waarom bent u in AI geïnteresseerd?

Klopt voor alle aanwezigen

U bent aanwezig

→ Klopt voor u

Statistische relevantie:

$P(\text{AI interesse} \mid \text{aanwezig})$

$\neq P(\text{AI interesse} \mid \text{aanwezig, stoel})$

Causaal:

Voorafgaandelijke causale processen en interacties



WAT IS EEN VERKLARING?



Aristotelianse 'oorzaken'

Waarom is dit wiel rond?

- Efficiënt (wie maakte het)
- Teleologisch (waarvoor dient het)
- Mechanistisch (hoe is het gemaakt)
- Categorisch (wat voor ding is het)

Unificerend:

Bvb. natuurwetten

Deductief:

Waarom bent u in AI geïnteresseerd?

Klopt voor alle aanwezigen

U bent aanwezig

→ Klopt voor u

Statistische relevantie:

$P(\text{AI interesse} \mid \text{aanwezig})$

$\neq P(\text{AI interesse} \mid \text{aanwezig, stoel})$

Causaal:

Voorafgaandelijke causale processen en interacties



WAT IS EEN VERKLARING?



Aristoteli

Waarom i

- Efficiën
- Teleolo
- Mechar
- Catego



Deductief:

Waarom bent u in AI geïnteresseerd?

Klopt voor alle aanwezigen

U bent aanwezig

**Complexiteit /
niveau van abstractie
?**

stoel)

Unificerend:

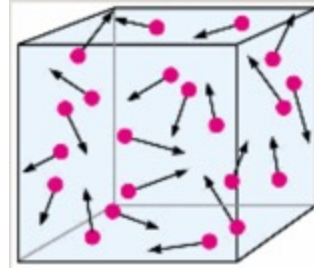
Bvb. natuurwetten

Voorafgaandelijke causale
processen en interacties



HET OOG VAN DE TOESCHOUWER

- Type verklaring
 - Complexiteit / abstractieniveau
 - Complementair met prior kennis
(bv. kennis van natuurwetten)
 - Relevantie
(bv. beoordelen van eerlijkheid / bias)
- ➔ Gebruikerstesten?



VERKLARINGEN VOOR WIE?

Verklaringen



Excuses

Verklaringen waar je mee aan de slag kunt



Verklaringen die je niet vooruit helpen

VERKLARING VAN HET MODEL OF REALITEIT

Belangrijk onderscheid: “Je krijgt geen lening, want...”

“... *volgens het AI-model* betalen mensen zoals jij hun lening vaak niet terug”

Zegt iets over het begrip van
het AI-model van de realiteit

“... *het AI-model* kent die niet toe aan mensen zoals jij”

Zegt iets het AI-model zelf

VERKLARINGEN & TRANSPARANTIE IN AI

- Veelal een pseudo-causale vorm: tegenfeitelijk
- Complexiteit, prior kennis, relevantie, ‘actionability’, eerlijkheid...:
Zelden expliciet meegenomen
- Gebruikerstesten in onderzoeksfase
- Interpreteerbaarheid = minder problematisch, maar verlies aan accuraatheid

Prof. Dr. Ir. Tijl De Bie

AIDA – IDLAB, UGENT

www.ugent.be

tijl.debie@ugent.be

 Universiteit Gent

 @tijldebie



Explainable AI geen louter technische aangelegenheid

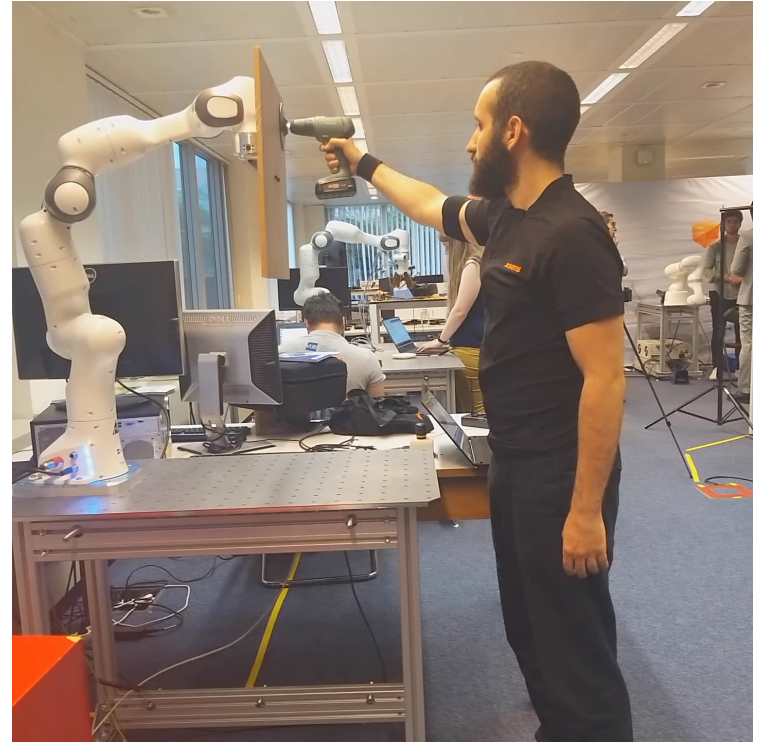
Prof An Jacobs – imec – SMIT, Vrije Universiteit Brussel



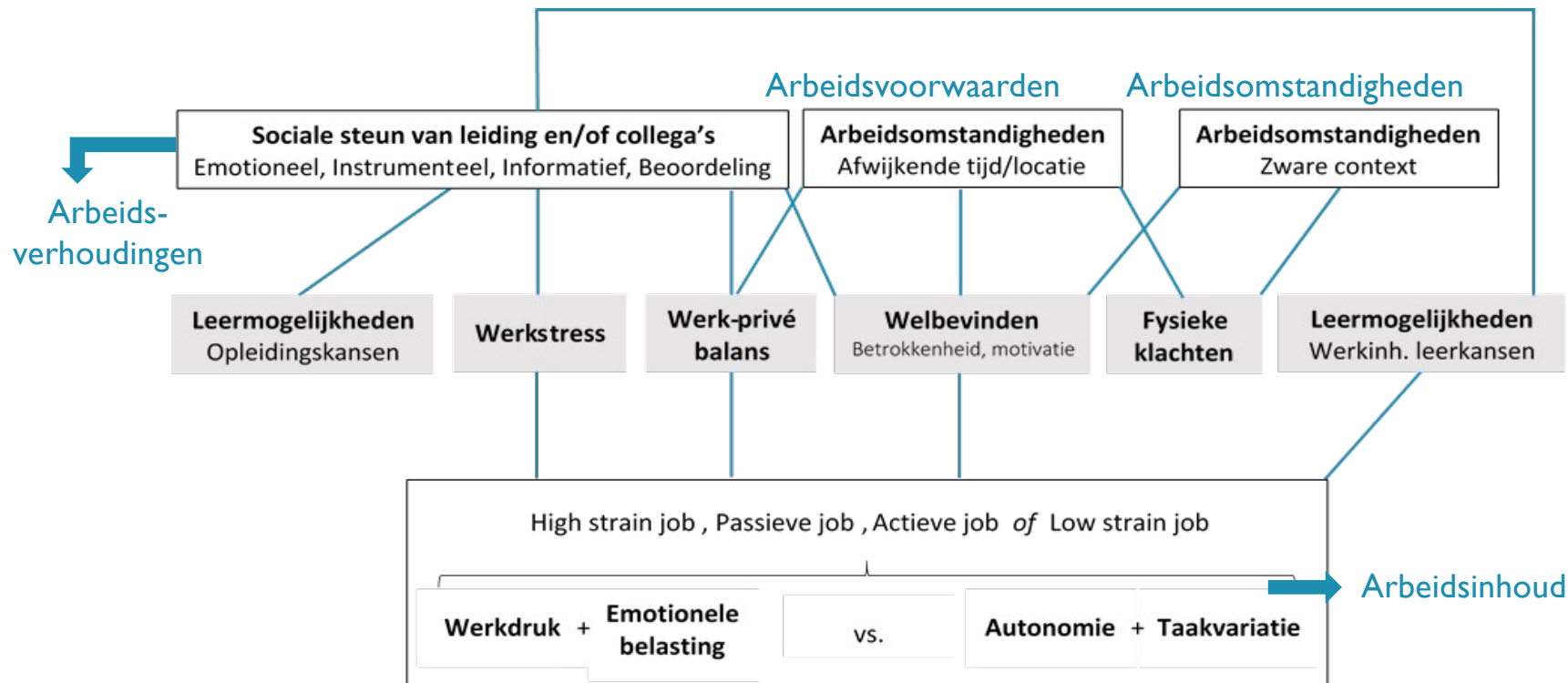
AI potentieel voor werkbaar werk

Zoekende naar ondersteunende
data gedreven technologieën voor meer
werkbaar werk

Bron:
Brubotics demo in EU SOPHIA project
@AIXPVUB



Werkbaar werk



Maar AI oplossingen vandaag leveren
ook volgende uitdagingen op



Bloomberg

Subscribe

Fired by Bot at Amazon: 'It's You Against the Machine'

Contract drivers say algorithms terminate them by email—even when they have done nothing wrong.

By **Spencer Soper**

28 juni 2021 12:00 CEST

Stephen Normandin spent almost four years racing around Phoenix delivering packages as a contract driver for Amazon.com Inc. Then one day, he received an automated email. The algorithms tracking him had decided he wasn't doing his job properly.

 **Copy Link**

Maar AI oplossingen vandaag leveren ook volgende uitdagingen op

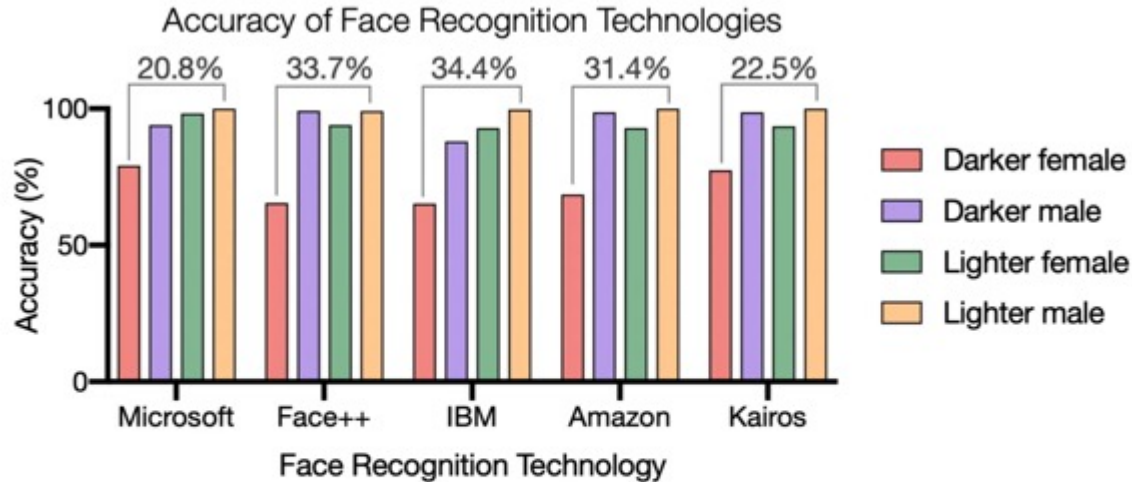


Figure 1: Auditing five face recognition technologies. The **Gender Shades project** revealed **discrepancies** in the classification accuracy of face recognition technologies for different skin tones and sexes. These algorithms consistently demonstrated the poorest accuracy for darker-skinned females and the highest for lighter-skinned males.

Hoe uitdaging aanpakken ?

Trustworthy human centred AI

EU High-Level Expert Group on AI
(AI HLEG)



AI in het bedrijfsleven

C-level analysten en data executives zeggen:

- 73% bedrijven moeite met prioritiseren van AI ethics

65% van de bedrijven kunnen de AI predicties en beslissingen niet verklaren

20% monitors actief models in production on fairness & ethics

Bron:

FICO "The State of Responsible AI", 2020

InformationWeek

 SIGN UP FOR OUR NEWSLETTERS

IT Leadership

DevOps

Security

Cloud

Data Manag

DATA MANAGEMENT

NEWS

6/14/2021
08:00 AM

How and Why Enterprises Must Tackle Ethical AI



Jessica Davis
News

Connect Directly

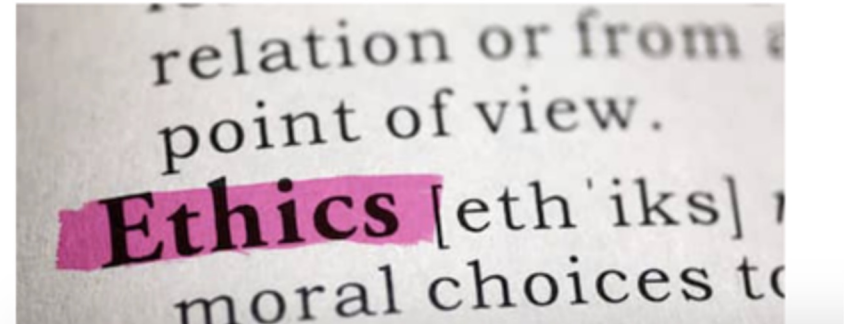


Login



Like

Artificial intelligence is becoming more common in enterprises, but ensuring ethical and responsible AI is not always a priority. Here's how organizations can make sure it is.



Belang ?



International Burn Down a 5G Tower Day

By TorchStone VP, Scott Stewart

May 4, 2020

What is Happening: A group of Internet activists have declared May 3, 2020 to be "International Burn Down a 5G Tower Day." They have been using the hashtags such as #burn5G, #burn5Gtowers, #destroy5G, #destroy5Gtowers, #Disable5Gtowers, #kill5G, #kill5Gtowers and #teardown5G. Twitter, Facebook and YouTube have been working diligently to remove content, forcing it to migrate to other platforms such as Instagram, Telegram, 4Chan and Gab.

XAI - Explainable AI

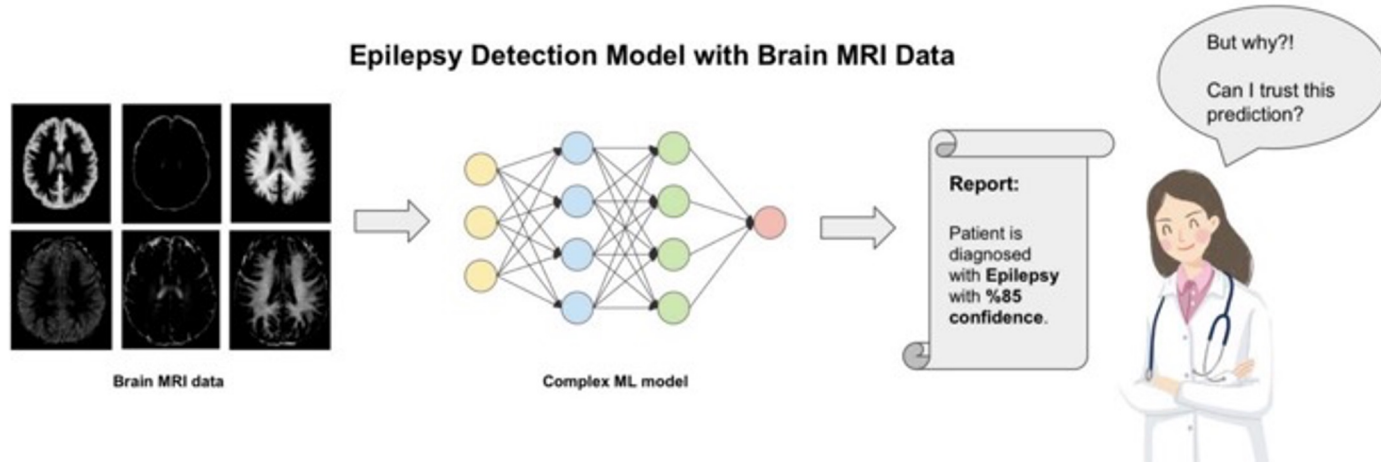
Voor data specialisten :
kunnen controleren **van**
performantie en **robustheid**
van het model in bepaalde
context en de transfereerbaarheid
naar andere situaties

Voor eind-gebruikers en andere stakeholders:

Uitleg, duidelijkheid nodig voordat
"duister" magisch verhaal het gebrek aan
transparantie invult

BASIC REQUIREMENT TO TRUSTWORTHY AI SOLUTIONS

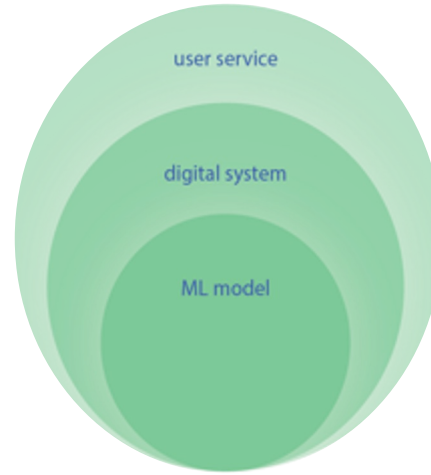
MORE DIVERSE NON-TECHNICAL “EXPLAINABILITY”



“If you can’t explain it to me, I don’t trust it”

N.N - UZ Jette

Waar te beginnen?



Organisatie &
maatschappelijke
context

<https://data-en-maatschappij.ai/en/publications/xai-decision-tool-2021>



understand & align

WHAT
is the system



WHO
needs an explanation



WHY
do they need one



explore & create

WHEN
do they need it



IN WHAT WAY
should the explanation be



CHOOSE
the right technique



TRANSLATE
to a user interface



reeds meeste
aandacht gekregen



evaluate & implement

ML MODEL
performance



XAI TECHNIQUE
performance



USER
performance



IMPACT
assessment



XAI alignment workshop : co-created format

Doel: hoe verschillende betrokken partijen & verschillende “explainability” noden tov system expliciteren en alligneren



The screenshot shows the website for the MyColruyt-app. At the top, there's a navigation bar with links: 'Over de MyColruyt-app', 'Op jouw maat', 'Inspiratie', 'Collect&Go', and 'Druktemeter'. The main heading reads 'Nóg makkelijker boodschappen voorbereiden met MyColruyt'. Below this, a sub-heading says 'Je boodschappenlijstje zó klaar'. A list of bullet points describes the app's features: adding products from the assortment or latest receipt, sorting the list by store layout, and getting price estimates. To the right of the text is a smartphone displaying the app's interface, which shows a shopping list with items like 'Aardappelen' and 'Kippenfilet'.



We asked ourselves, is it possible to create a self-learning algorithm to make suggestions based on peoples preferences, search behaviour on our website and their profile

XAI alignment workshop :



facilitatie face-to-face of online via MIRO

Wat? **Explainability (uitlegbaarheid)** van een AI systeem = uitleggen hoe het systeem werkt en waar de resultaten op zijn gebaseerd

Waarom? Om te weten of een (autonoom) systeem de juiste beslissing maakt, om te weten of er fouten (bugs) zijn en om te weten of wij als mens een andere beslissing zouden willen maken (overruling). Al deze aspecten helpen bij het verbeteren van **de betrouwbaarheid van het systeem**.

Doel van deze workshop: om gezamenlijk over de wat, wie en waarom van explainability na te denken en uit te stippelen waar je als (innovatie) team naartoe wil gaan werken vanaf het begin van het ontwikkelingsproces. Op die manier wordt explainability een onderdeel van het systeem vanaf het begin en geeft het sturing aan de keuzes die worden gemaakt voor de technieken en ontwerpen van het systeem.

Deze workshop is gebaseerd op de volgende tools:



Artificial Intelligence Impact Assessment
Klik hier voor meer informatie



WORKSHOP DEEL I Wat en wie ? (2 uur + voorgesprek)
WORKSHOP DEEL II Waarom en hoe ? (2 uur)

XAI alignment workshop :

Best bij start van project wanneer eerste idee over use case gevormd is

**Wat en wie ? (2 uur +
voorgesprek)**

Wat wil je dat de toepassing doet,
welke data input, welke output, ...

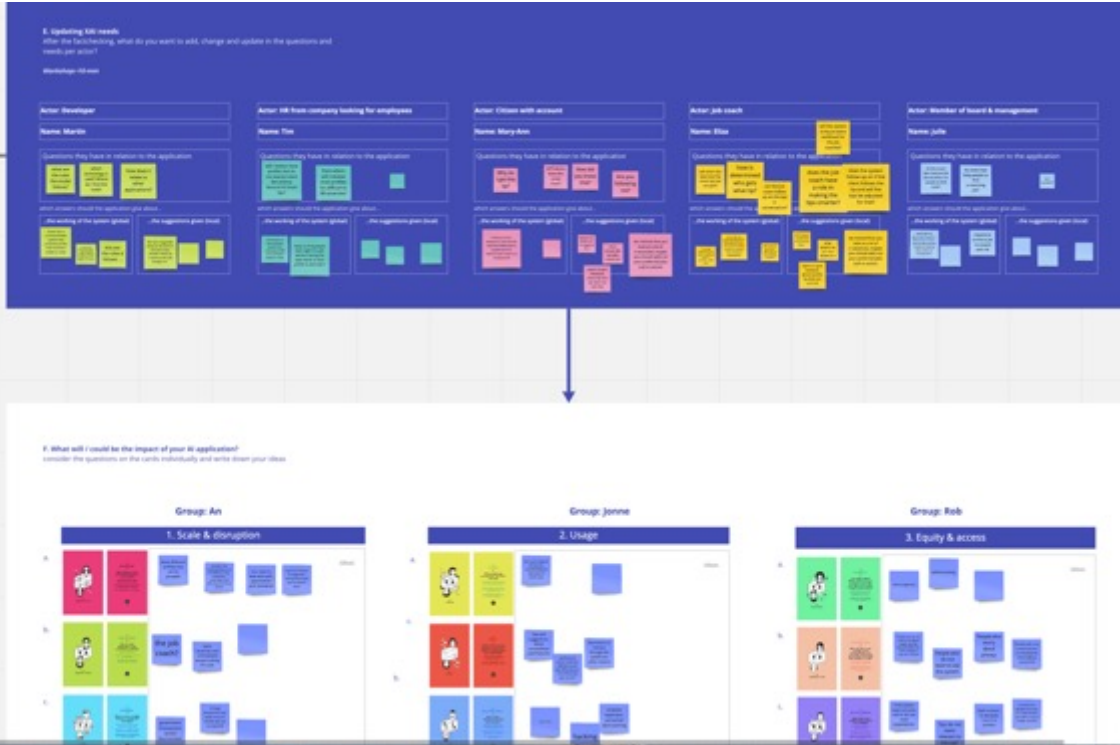
Wie intern en extern speelt een rol
bij toepassing, wat is hun doel en
wat is voor hen de meerwaarde

Ontwikkelen van persona kaarten
voor 5 belangrijkste actoren



XAI alignment workshop :

Best bij start van project wanneer eerste idee over use case gevormd is



Waarom en hoe ? (2 uur)

Verdiepen van persona's explainability noden en pijnpunten,

Welke globale en locale info wil men verkrijgen van systeem, waar en hoe ?

Bredere reflectie over impact van toepassing

Kiezen van prioriteiten qua verbeter punten

Testimonials na de XAI alignment workshops

- Bespreekbaar maken en definiëren van ethische/explainability vereisten op een gestructureerde wijze

"It (This workshop) allowed for an interesting discussion on how non-tangible objectives are tranfomed into metrics which is always difficult and often based on beliefs. When uninteded bias creeps in it would be here. This was a new insight for me "

***data scientist COLRUYT,
2020***

"I have learnt that explainability is no one-size-fits-all story.

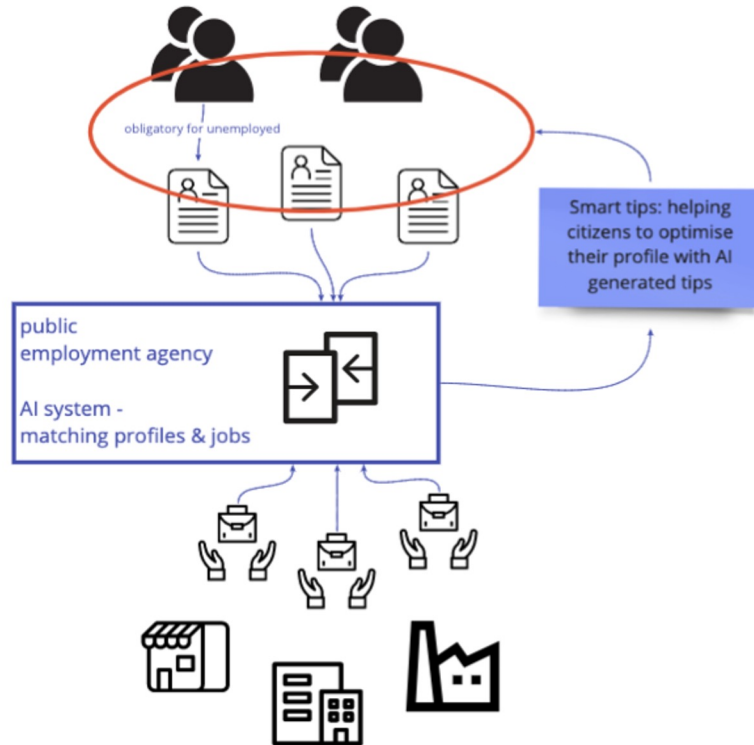
There are different actors who all have different needs.

These types of exercises are not easy, but definitely necessary for bringing different perspectives together.

If you'd ask me what is explainability, as a developer I would give a technical answer trying to make algorithms explainable with several techniques out of academic literature, but there is actually much more to think about."

data scientist VDAB, 2021

Er is ook train-the-trainer versie: met vaste use case (2 uur)



The case: Smart Tips

A **public employment agency** has a job bank where citizens can create a **profile** and look for a job. Since it's a public agency, citizens who are **officially unemployed** and who receive an allowance from the government, are **obliged to create a profile** in this job bank.

The agency would like to improve their recommendation and **matching system** by helping citizens to **optimise their profile**. For this, they would like to develop a '**smart tip**' **application** which offers AI-generated, personal tips to citizens about how to improve their profile. The smart tip engine will work together with existing models in the system to **improve the 'chance of work'** score and '**matchability**' score of the profiles. The final aim is to help citizens **to find a fitting job or**

WORKSHOP

Human-centered
explainable AI



Workshop: Human-centered Explainable AI

Over transparantie en uitlegbaarheid van AI-systemen

workshop ethiek vertrouwen uitlegbaarheid
explainability trustworthy AI transparantie

DIY WORKSHOP

AI Blindspots in
een zorgcontext
identificeren



DIY workshop: AI Blindspots in een zorgcontext identificeren

Over het identificeren van ethische, juridische en maatschappelijke uitdagingen in zorgtoepassingen

workshop AI Blindspots ethiek maatschappij
juridisch healthcare ethische valkuilen

ONLINE VORMING

Hoe rekening houden met
ethiek en regelgeving in
een projectvoorstel over AI?



Webinar: Algemene Verordening Gegevensbescherming

Over het toepassen van de AVG in KMO's

lezing vorming recht AVG/GDPR



understand & align

WHAT
is the system



WHO
needs an explanation



WHY
do they need one



explore & create

WHEN
do they need it



IN WHAT WAY
should the explanation be



CHOOSE
the right technique



TRANSLATE
to a user interface



evaluate & implement

ML MODEL
performance



XAI TECHNIQUE
performance



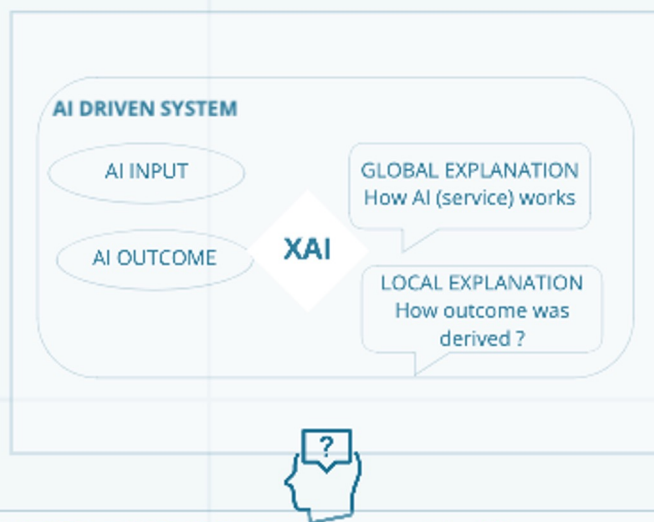
USER
performance



IMPACT
assessment



MEASUREMENT INSTRUMENTS



EVALUATION DIMENSIONS to assess with different user types

meaning
explanation

trust in
explanation

utility of
explanation

quality of
explanation

ease of use
of explanation

acceptance
of outcome

quality of
of outcome

...



embracing a better life

— Ethiek in de ontwikkeling van AI-toepassingen



— Voorstelling

Dr. Tomas Folens

Lector en onderzoek ethiek en coaching aan VIVES Hogeschool

Executive coach (ICF)

Wetenschappelijk medewerker KULeuven, Faculteit TRW, OE Theologische Ethiek

Contact

tomas.folens@vives.be

<https://digital-ethics.blog/>

www.conectar.be

— startvraag

Hoe ethiek vertalen naar de werkvloer?

Hard law – verplichtend karakter, voor alle bedrijven hetzelfde (Ricoeur: gestolde ethiek)

Soft law – normen, deontologische codes..., minder dwingend karakter

Ethiek – waarden, wat zou moeten gebeuren, voordeel: aanpasbaarder

Europese commissie: tussen soft law & ethiek (principalism), wat betekent dat de aanpasbaarheid aan concrete situaties veel moeilijker is – gevolg: vaak kloof tussen abstracte ethische kaders en de concrete werkvloer.

Uitdaging om de kloof te overwinnen: stappenplan

— Concrete aanpak in stappen



— **Stap 1: beginsituatie**

Business Understanding

What does the business need?



'Modelleren' van de
beginsituatie
(bedrijfsprocessen)

Bepalen stakeholders &
bepalen van significante
rollen

Ethische beginsituatie
Welke waarden en normen
spelen in ons bedrijf? (CSR)
Welke waarden en normen
worden bedreigd?

— Stap 2: Analyse

Analyse



- Analyse
- Functioneel – technisch-ethisch
- selecteren: data

Dialoog met de stakeholders

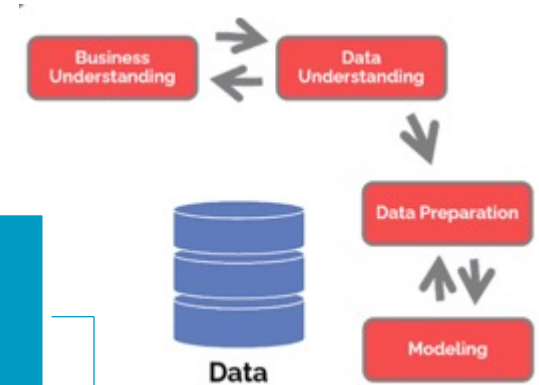


- Nood aan gesprek over ethiek
- Wat is voor jullie belangrijk in deze casus? Welke waarden realiseren?
- Voordeel: innovatie & implementatie

Bepalen van opties



- Opties naar beschikbare data toe
- Opties op basis van technische criteria
- Opties op basis van ethische criteria



Beschikbare data?
Bias?

Ethics by design
Vertaling > technische criteria

— **Stap 3: kiezen model en algoritme**

Bij kiezen algoritme:

Ethische keuzes

Modeling

Data Preparation

Bv. Positief of negatief getest op Corona (hoeveel false positive & false negative?)

Selecteren technieken – FB-mechanismes

— Stap 4 – Evaluatie en implementatie



Checken met stakeholders & ethische evaluatie na proefdraaien van algoritme (impact algoritme, gevolgen)

Implementatie

**Ethiek is niet moeilijk:
Waar(om), wat, hoe?
Rob Heyman**



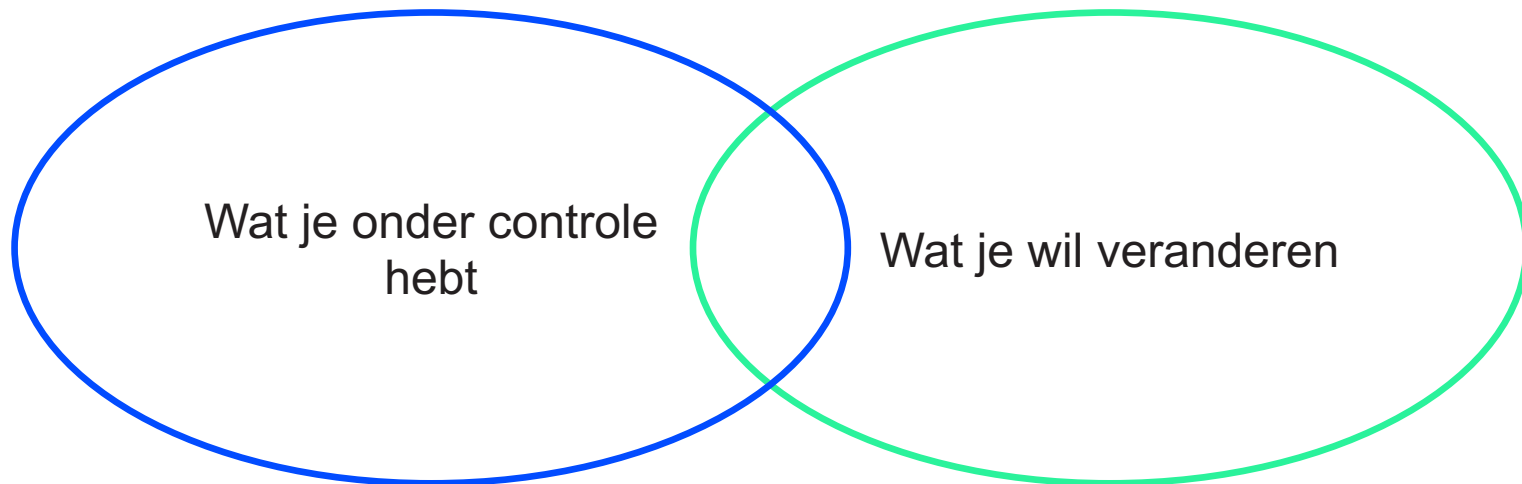
Waar denk je aan als je het woord ethiek hoort?



HIC SUNT DRACONES

- Visie
- Goedkeuring van een project
- Iets voelt uitdagend aan (data, context, stakeholders, ...)
- Investeerder vraagt het
- Het moet

Wat?



Hoe? wel

Als onderdeel van een project, met middelen en tijd

INDEPENDENT
HIGH-LEVEL EXPERT GROUP ON
ARTIFICIAL INTELLIGENCE
SET UP BY THE EUROPEAN COMMISSION



THE ASSESSMENT LIST FOR
TRUSTWORTHY ARTIFICIAL
INTELLIGENCE (ALTAI)
for self assessment

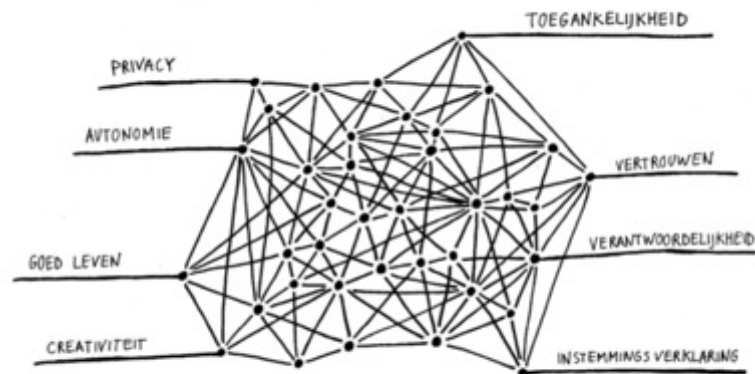
EU Projecten

- **Step1: Register to the [ALTAI](#)**
- **Step 2: Log in to start using the online tool**
- **Step 3: Read the instructions**
 - ALTAI is best completed involving a multidisciplinary team of people from within or outside your organisation with specific competences or expertise on each of the 7 requirements.

Data Ethics Decision Aid

Deze [tool](#) is gemaakt door de [Utrecht Data School](#) in samenwerking met data-analisten van de [Gemeente Utrecht](#). Ze bestaat uit verschillende methodieken die je kan inzetten bij het begin van een project. Het doel is om de ethische kwesties en beslissingen daaromtrent te documenteren, bijvoorbeeld ter verantwoording naar stakeholders.

-



The page features decorative geometric shapes on the left and right sides. On the left, there is a large, multi-colored hexagonal structure composed of several smaller hexagons in shades of red, yellow, green, blue, and orange. On the right, there is a smaller, red and orange geometric shape. The background is a light blue gradient.

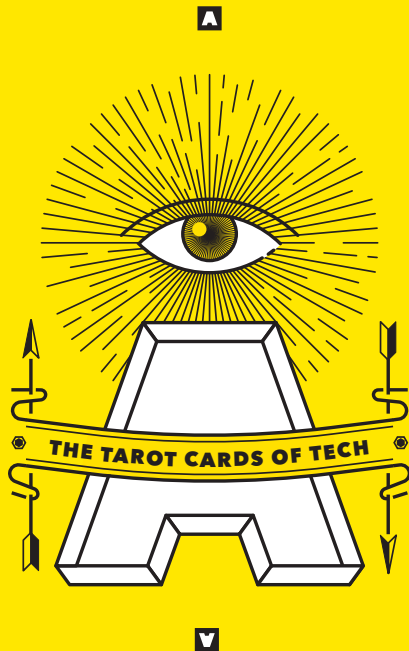
SDG Impact Assessment Tool

GUIDE 1.0

We recommend using the SDG Impact Assessment Tool in an iterative five-step process.



Follow the five-step process to make the best use of the tool.



**Artefact created
The Tarot Cards of Tech
to help creators of all
kinds consider the impact
of technology.**

Each card contains provocations that will not only help you foresee unintended consequences, but also reveal opportunities for creating positive change. Take The Tarot Cards of Tech to your next brainstorm or team meeting to gaze into the future and better understand the potential impact of your products.

If you like these cards, let us know at info@artefactgroup.com.

From workshops to design thinking, Artefact would love to learn how best we can help you and your organization.

www.artefactgroup.com

TOOLS

FILTERS

Ethische principes ▾

Procesfase ▾

Moeilijkheid ▾

Doelgroep ▾

Types ▾

Prijs ▾

Onderwerp ▾

[Hoe gebruik ik de filters?](#)



Tool: AI Ethics cards



Tool: Cards for Humanity



Tool: amai!-kaartspel