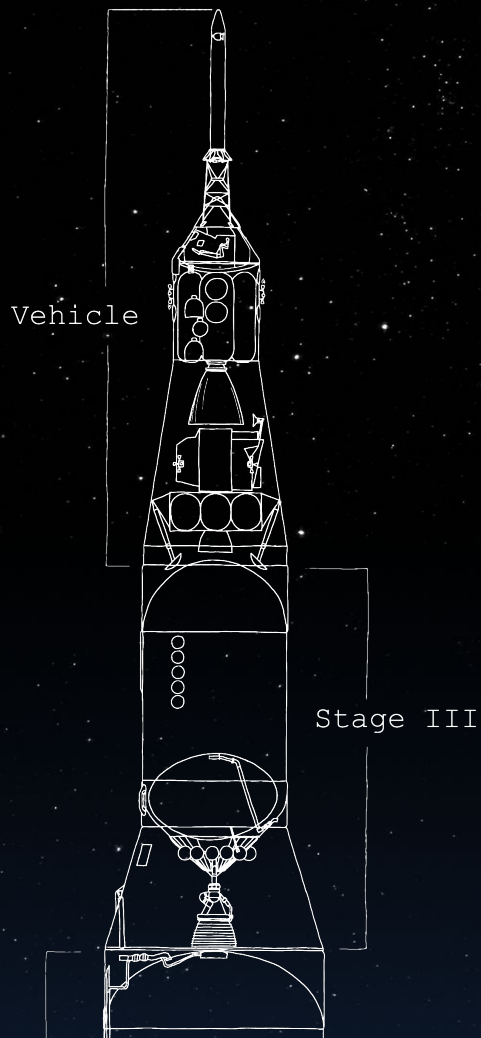




XAI

THE Key to AI operationalization

Clodéric Mars - CTO @ craft ai

17/10/2019 - Séminaire Aristote l'IA est-elle explicable

Explainable Artificial Intelligence

https://www.darpa.mil/program/explainable-artificial-intelligence

DARPA DEFENSE ADVANCED RESEARCH PROJECTS AGENCY

ABOUT US / OUR RESEARCH / NEWS / EVENTS / WORK WITH US

Defense Advanced Research Projects Agency > Program Information

Explainable Artificial Intelligence (XAI)

Mr. David Gunning

The diagram illustrates the need for explainable AI. It shows a flow from an 'AI System' (represented by a neural network icon) to 'DoD and non-DoD Applications' (a box listing Transportation, Security, Medicine, Finance, Legal, and Military), and finally to a 'User' (represented by a person icon). Arrows indicate the direction of the flow.

AI System

- We are entering a new age of AI applications
- Machine learning is the core technology
- Machine learning models are opaque, non-intuitive, and difficult for people to understand

DoD and non-DoD Applications

- Transportation
- Security
- Medicine
- Finance
- Legal
- Military

User

- Why did you do that?
- Why not something else?
- When do you succeed?
- When do you fail?
- When can I trust you?
- How do I correct an error?

RESOURCES

- DARPA-BAA-16-53
- DARPA-BAA-16-53: Proposers Day Slides
- XAI Program Update

Figure 1. The Need for Explainable AI

The current generation of AI systems offer tremendous benefits, but **their effectiveness will be limited by the machine's inability to explain its decisions and actions to users**

- David Gunning - DARPA/I2O XAI Program Update November 2017

Intelligence Artificielle

Les défis actuels et l'action d'Inria



Inria

LIVRE BLANC

N°01

Les systèmes d'IA ont vocation à interagir avec des utilisateurs humains : **ils doivent donc être capables d'expliquer leur comportement**, de justifier d'une certaine manière les décisions qu'ils prennent afin que les utilisateurs humains puissent comprendre leurs actions et leurs motivations.

- Intelligence Artificielle. Les défis actuels et l'action d'Inria



PREPARING FOR THE FUTURE OF ARTIFICIAL INTELLIGENCE

Executive Office of the President
National Science and Technology Council
Committee on Technology

October 2016



Federal agencies [...] should [...] ensure that AI-based products or services purchased with Federal grant funds **produce results in a sufficiently transparent fashion** and are supported by evidence of efficacy and fairness.

- Preparing for the future of Artificial Intelligence

CÉDRIC VILLANI

Mathématicien et député de l'Essonne

DONNER UN SENS À L'INTELLIGENCE ARTIFICIELLE

POUR UNE STRATÉGIE
NATIONALE ET EUROPÉENNE

Composition de la mission

Marc Schoenauer Directeur de recherche INRIA • **Yann Bonnet** Secrétaire général du Conseil national du numérique • **Charly Berthet** Responsable juridique et institutionnel du Conseil national du numérique • **Anne-Charlotte Cornut** Rapporteur au Conseil national du numérique • **François Levin** Responsable des affaires économiques et sociales du Conseil national du numérique • **Bertrand Rondepierre** Ingénieur de l'armement, Direction générale de l'armement.

En premier lieu, il faut accroître la transparence et l'auditabilité des systèmes autonomes d'une part, en développant les capacités nécessaires pour observer, comprendre et auditer leur fonctionnement et, d'autre part, en **investissant massivement dans la recherche sur l'explicabilité.**

- Cédric Villani - Donner Un Sens À L'Intelligence Artificielle



LUC JULIA

L'Intelligence artificielle n'existe pas



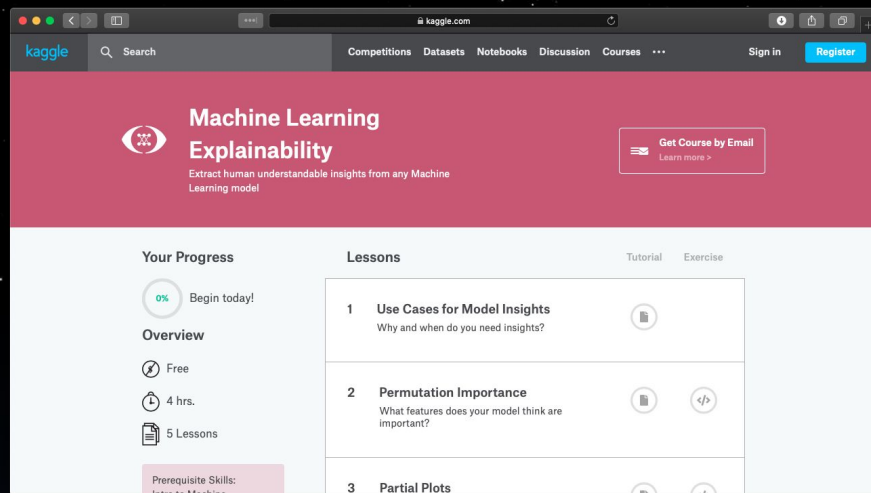
Le cocréateur de *Siri*
déconstruit le mythe de l'IA !

FIRST
EDITIONS



L'explicabilité est importante parce qu'elle apporte la confiance. **Si on est capable d'expliquer pourquoi et comment on fait les choses, ça enlève le côté magique.**

- Luc Julia - L'intelligence artificielle n'existe pas



Many important decisions are made by humans. For these decisions, **insights can be more valuable than predictions.**

In practice, **showing insights [...] will help build trust, even among people with little deep knowledge of data science.**

- Dan Becker - Data Scientist, Instructor @ Kaggle



Explanations are mandatory
when AI empowers humans to
perform complex tasks

- Antoine Buhl - XAI, a game changer for AI in production @ AI Night 2019

When it comes to create AI for
critical systems, **trustability** and
certifiability are mandatory.

- David Sadek - XAI, a game changer for AI in production @ AI Night 2019



Quantmetry

Nos offres

Nos compétences

R&D et Innovation

Qui sommes-nous ?

Carrières

Événements

Blog

Interprétabilité des modèles

De nombreux projets data s'appuient sur la création d'un algorithme dont les performances peuvent être correctes mais dont la mise en production pose question faute de fonctionnement compréhensible. Chez **Quantmetry**, nous sommes convaincus que cette démarche d'intelligibilité, nécessaire pour rendre moins opaque les modèles prédictifs, sera bientôt incontournable pour l'adoption de l'Intelligence Artificielle à grande échelle.

En nous appuyant sur des travaux de R&D internes, nos contributions open source et des échanges réguliers avec le monde de la recherche, nous avons développé une expertise, des convictions et un ensemble de bonnes pratiques sur l'utilisation de techniques favorisant l'interprétation des décisions prises par les modèles prédictifs.

visant à mieux comprendre les facteurs liés à une décision en particulier, nous sommes convaincus que notre expertise pourra vous être utile, comme elle l'a déjà été pour plusieurs de nos clients :

- Extraction de règles logiques pour décrire – et donc mieux comprendre – d'une manière approchée un modèle prédictif complexe.
- Activation des bons leviers d'action au niveau individuel, basé sur le calcul des variables contribuant le plus à une prédiction de churn associé.
- Analyse de l'impact de plusieurs variables issues de données externes, afin de valider que cet impact sur les prédictions était conforme à l'attendu (sens de variation, effet de seuils, etc.).

Vous pouvez retrouver [notre livre blanc "IA, explique-toi !"](#) qui instruit cette problématique de l'intelligibilité des modèles de machine



IA Explique toi ! Quand la performance ne suffit pas

github.com

Why GitHub? Enterprise Explore Marketplace Pricing

Search

Sign inSign up

IBM / AIX360

Watch20Star365Fork57

CodeIssues6Pull requests1Projects0SecurityInsights

Branch: masterAIX360 / README.mdFind fileCopy path

vijay-arya Merge pull request #31 from sadhamanus/masterf26db44 on 16 Sep

5 contributors

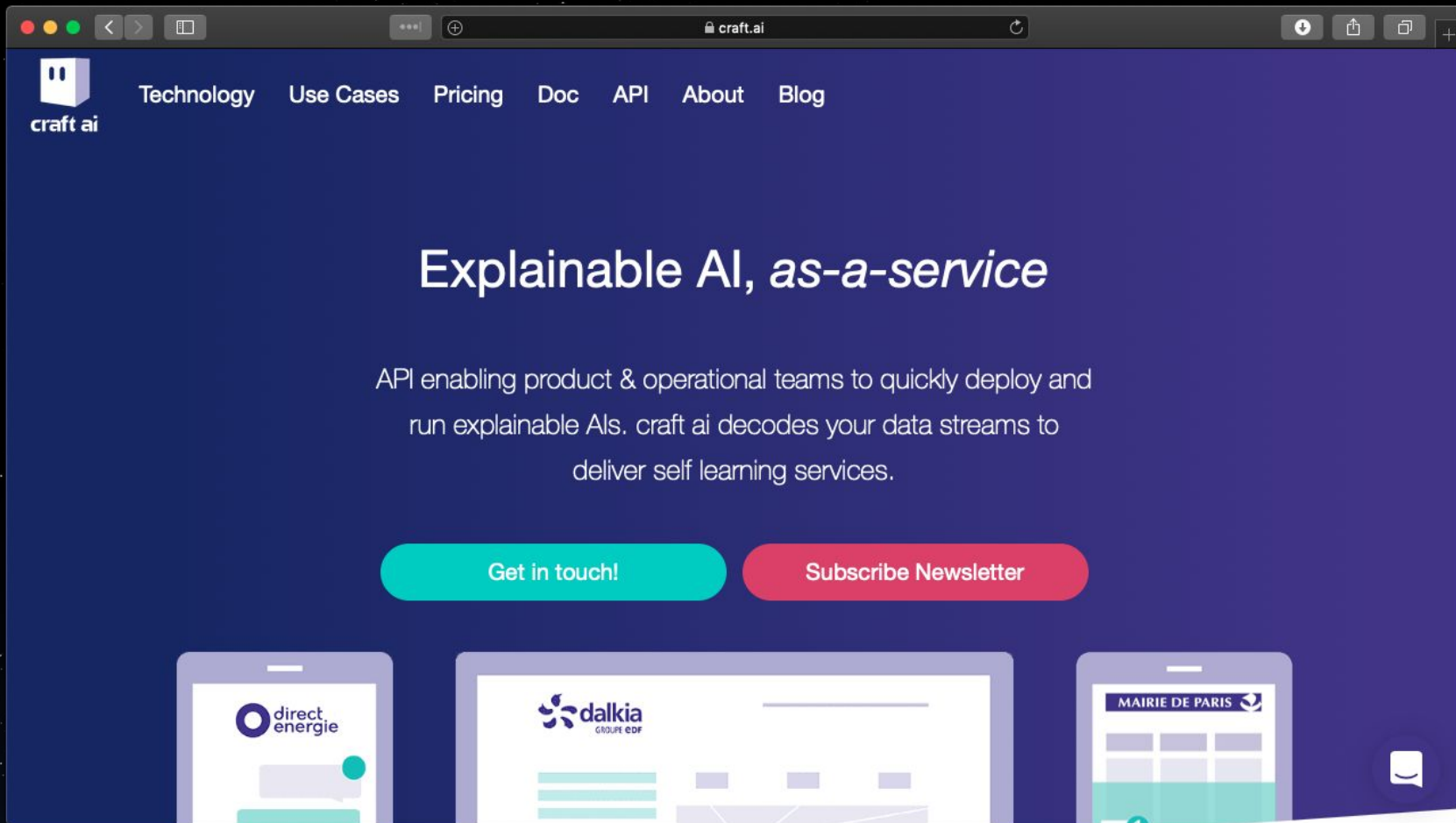
140 lines (92 sloc)6.61 KBRawBlameHistory

AI Explainability 360 (v0.1.0)

buildpassingdocspassingpypi package0.1.0

The AI Explainability 360 toolkit is an open-source library that supports interpretability and explainability of datasets and machine learning models. The AI Explainability 360 Python package includes a comprehensive set of algorithms that cover different dimensions of explanations along with proxy explainability metrics.

The [AI Explainability 360 interactive experience](#) provides a gentle introduction to the concepts and capabilities by walking through an example use case for different consumer personas. The [tutorials and example notebooks](#) offer a



How XAI makes a
difference?

What's an explanation?

Tim Miller - Explanation in Artificial Intelligence: Insights from the Social Sciences

People look for explanations to improve their understanding of someone or something so that they can derive stable model that can be used for prediction and control

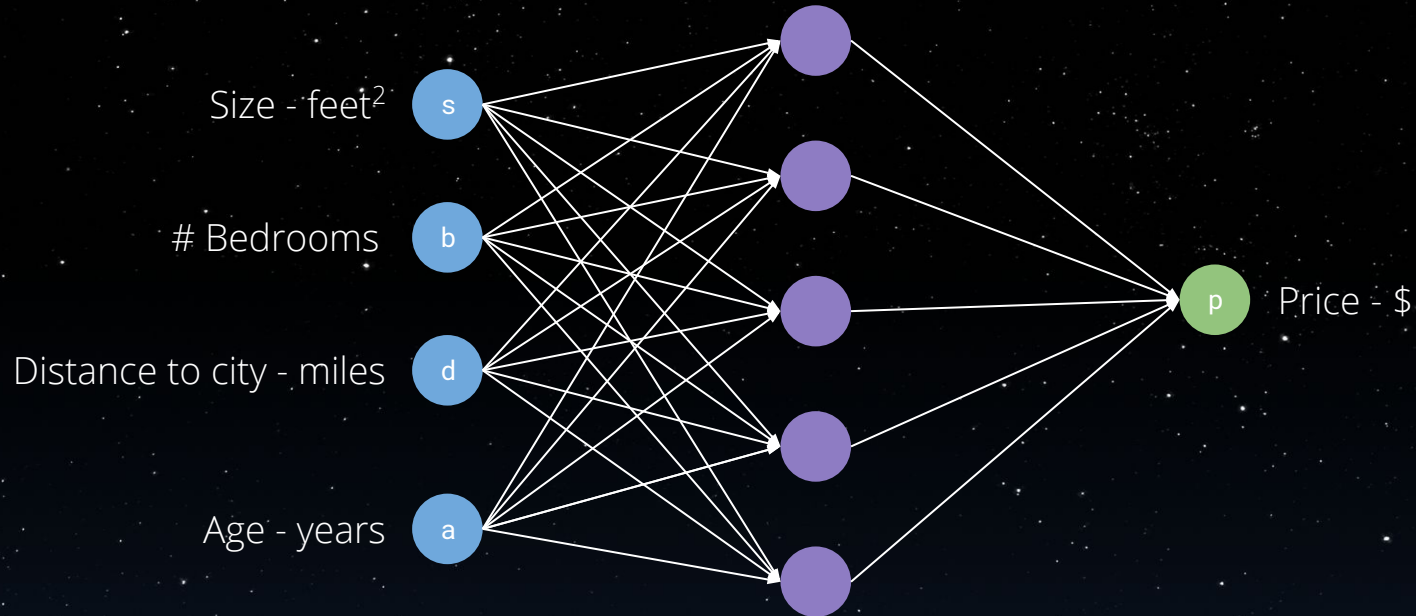
- Fritz Heider, Australian psychologist

1. Answer to "Why?" questions
2. Answer with contrastive explanations
3. Biased towards the explainee

AI, not explainable?

Anatomy of a Neural Network

SuperDataScience - Artificial Neural Networks - How do Neural Networks Work?



$$p = f(w_1 f(w_{11}s + w_{21}b + w_{31}d + w_{41}a) + w_2 f(w_{12}s + w_{22}b + w_{32}d + w_{42}a) + \dots)$$

Adversarial attacks

Alexey Kurakin, Ian Goodfellow, Samy Bengio - Adversarial examples in the physical world



(a) Image from dataset



(b) Clean image

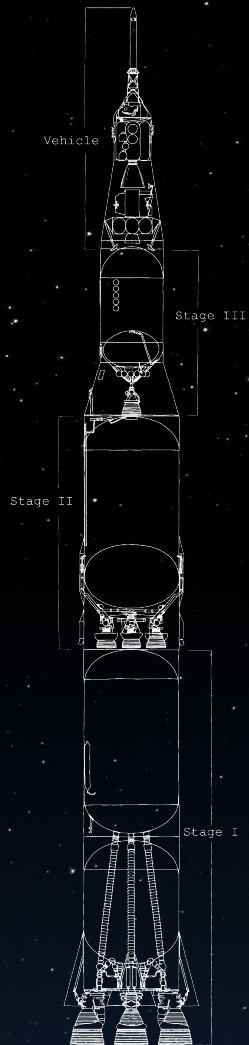


(c) Adv. image, $\epsilon = 4$

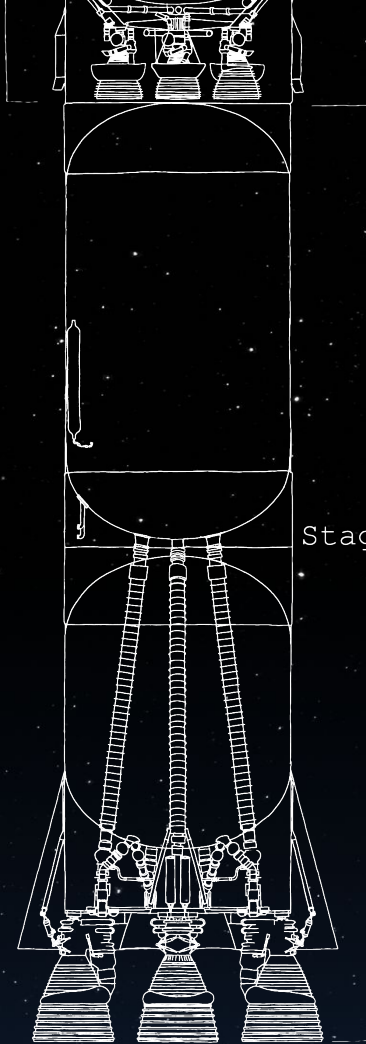


(d) Adv. image, $\epsilon = 8$

How XAI makes a
difference?



The 3 stages of XAI



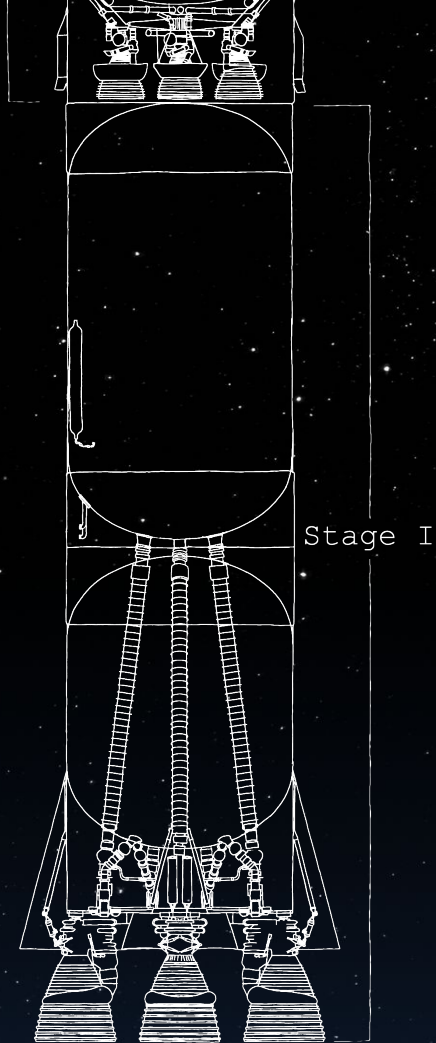
Stage I

Stage I

Explainable building process

Involve Business Experts

Acceptability + Performances



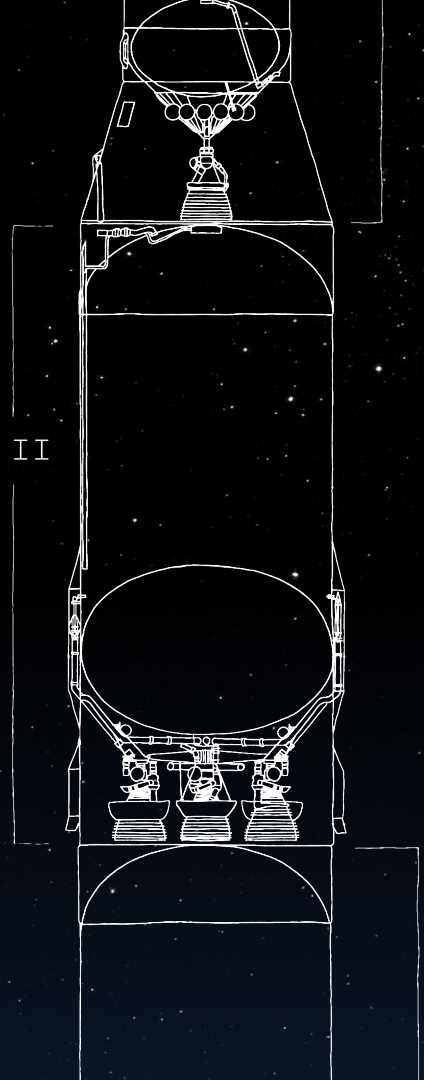
Stage I
The tools

Results Visualization
Model Exploration
Simulation

die längsten reisen fangen an , wenn es auf den straßen dunkel wird .



Translation system debugger & visualizer



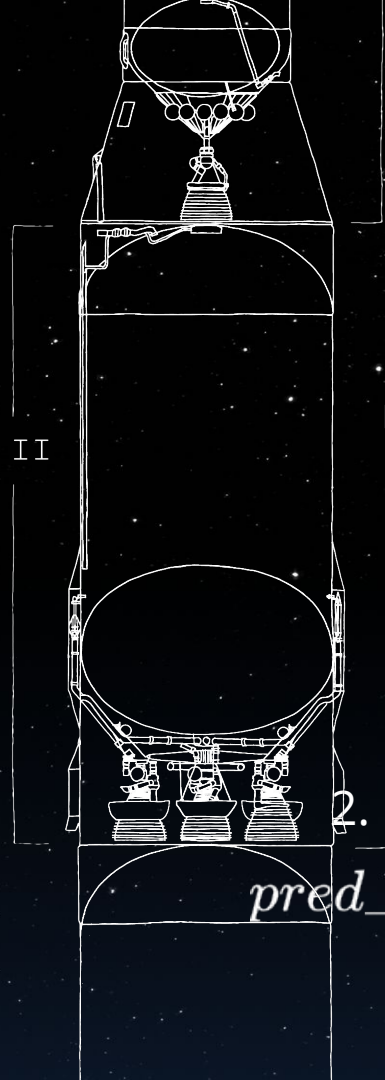
Stage II

Stage II

Explainable decisions

Follow Least Surprise Principle

Trust + Traceability



Stage II

Stage II SHAP

[Scott M. Lundberg, Su-In Lee - A Unified Approach to Interpreting Model Predictions - NeurIPS 2017]

1. Select business understandable features
 $x' = \text{simpl_feat}(x)$

2. Fit a explainable model approximating the actual model

$$\text{pred_model}(x) \approx \text{expl_model}(x') = \phi_0 + \sum_{i=1}^{\dim(x')} \phi_i x'_i$$



BLECKWEN

18/04/2019 - 07:28 | Alert creation : **fraud suspicion**

⚙️ An alert has been created for suspicion of fraud.

RULES FEEDBACK

❌ ml

❌ RULE_MODEL

MODEL_SCORE > 0.5 => **BLOCK**

[See all rules](#)

MACHINE LEARNING FEEDBACK

78 %



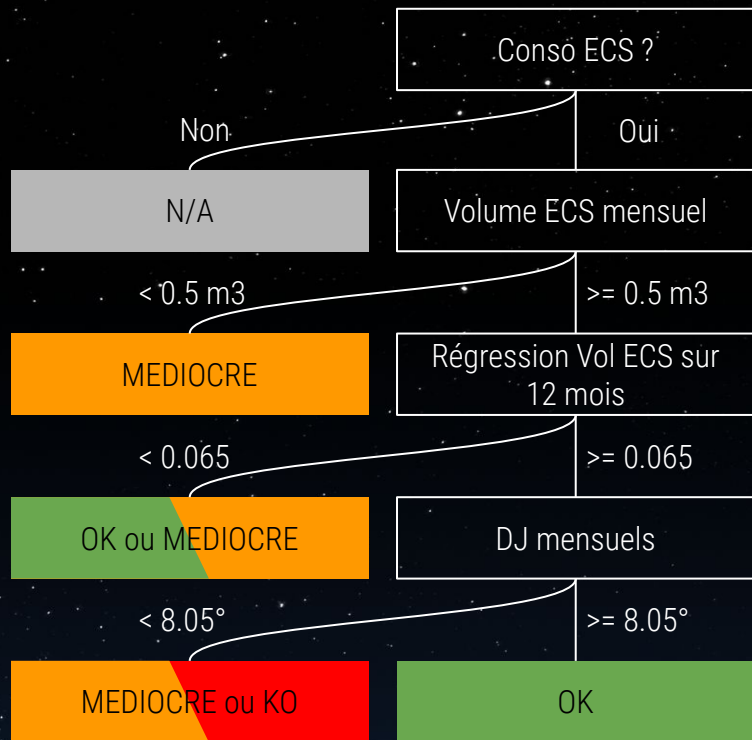
[See all the features](#)

Stage II

SHAP Usage example

Fraud pattern recognition

Explainability = Insights + Traceability

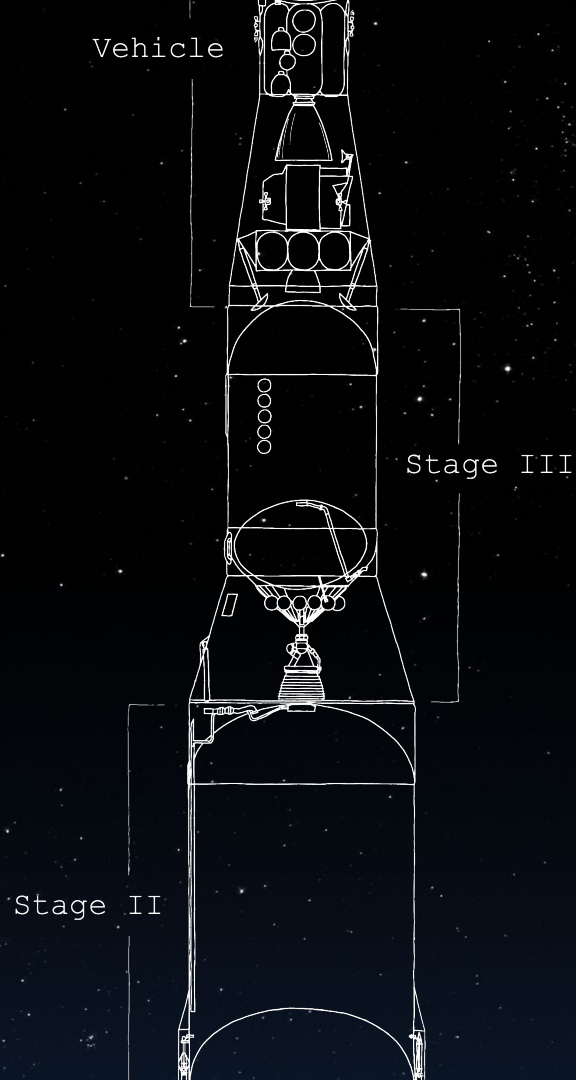


Stage II

Decision Tree to generate
contrastive explanations

Energy manager assistant

Explainability = Productivity



Stage III

Explainable decision process

Enable programmatic inspection of
the decision process and underlying
concepts

Automation + Certifiability

Stage III RuleFit

Generate a minimal set of rules from a trained Random Forest

Predictive Learning via Rule Ensembles

Jerome H. Friedman* Bogdan E. Popescu†

October 5, 2005

Abstract

General regression and classification models are constructed as linear combinations of simple rules derived from the data. Each rule consists of a conjunction of a small number of simple statements concerning the values of individual input variables. These rule ensembles are shown to produce predictive accuracy comparable to the best methods. However their principal advantage lies in interpretation. Because of its simple form, each rule is easy to understand, as is its influence on individual predictions, selected subsets of predictions, or globally over the entire space of joint input variable values. Similarly, the degree of relevance of the respective input variables can be assessed globally, locally in different regions of the input space, or at individual prediction points. Techniques are presented for automatically identifying those variables that are involved in interactions with other variables, the strength and degree of those interactions, as well as the identities of the other variables with which they interact. Graphical representations are used to visualize both main and interaction effects.

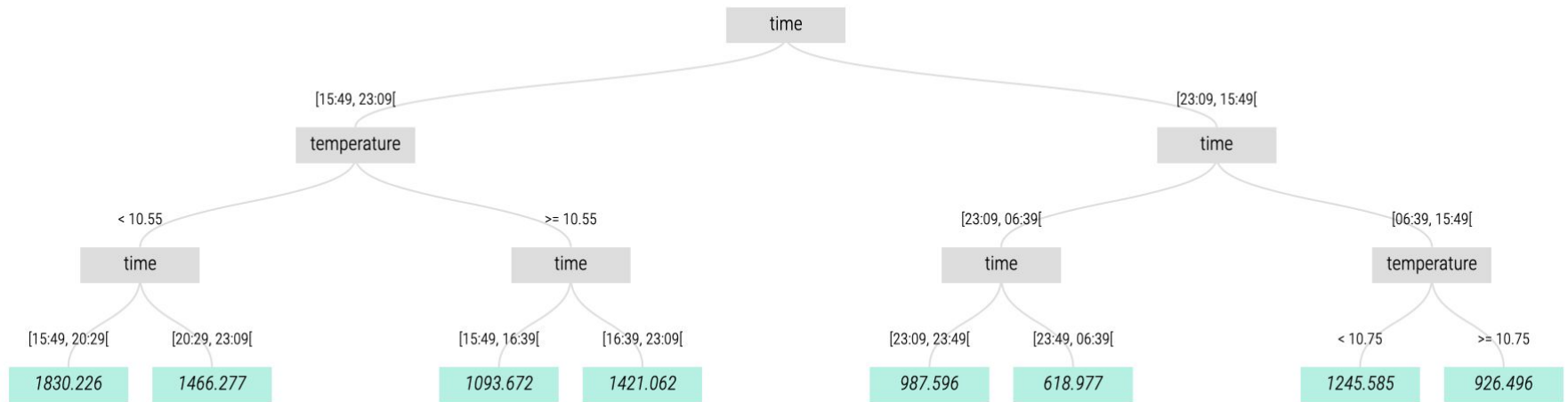
Key words and phrases: regression, classification, learning ensembles, rules, interaction effects, variable importance, machine learning, data mining.

1 Introduction

Predictive learning is a common application in data mining, machine learning and pattern recognition. The purpose is to predict the unknown value of an attribute y of a system under study, using the known joint values of other attributes $\mathbf{x} = (x_1, x_2, \dots, x_n)$ associated with that system. The prediction takes the form $\hat{y} = F(\mathbf{x})$, where the function $F(\mathbf{x})$ maps a set of joint values

Stage III

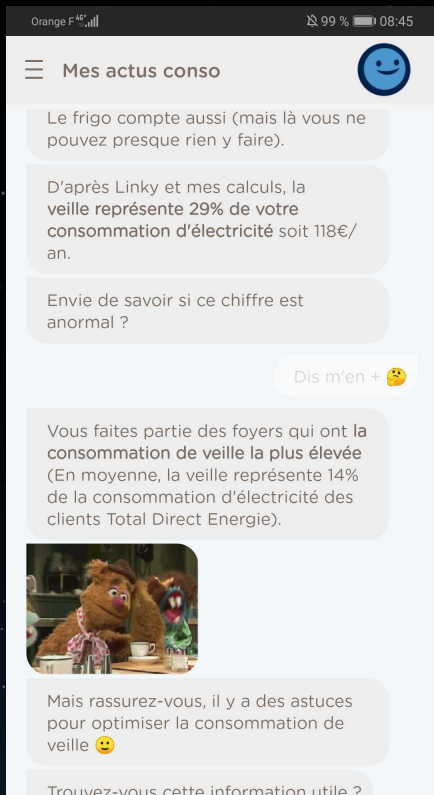
Decision Trees



Energy coach

Household consumption predictive model used as a knowledge base

Explainability = Interface with symbolic systems



Challenges

| Not because they are
easy, but because they
are hard

- John F. Kennedy

How can we evaluate explainability?

Challenges

| Not because they are
easy, but because they
are hard

- John F. Kennedy

Improve the performances of XAI

Improve Data Engineering

New Machine Learning approaches

Hybrid models

Takeaways

We set sail on this new sea because there is new knowledge to be gained, and new rights to be won, and they must be won and used for the progress of all people.

- John F. Kennedy

Explainability is an opportunity

Regulation is coming

No system in production below stage II

