



FATES-MLOps

Incorporating FATES Principles in Continuous Development of ML-Integrated Systems: A MLOps Perspective

2024-2028

Fairness

Accountability

Transparency

Ethics

Security (and/or Safety and/or Sustainability)



Available material

<http://fates-mlops.org/>

HOW TO CITE:

Jean-Michel Bruel et al, “Présentation du projet FATES-MLOps aux Journées du GdR GPL”. PAU, France, 2025.



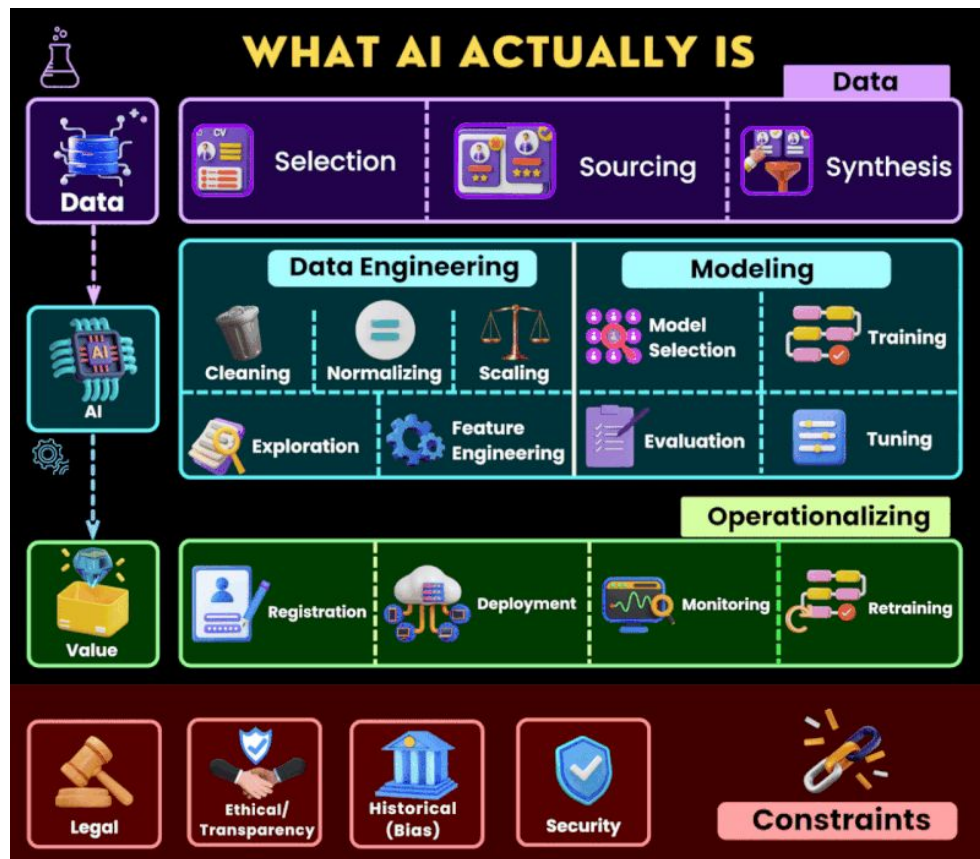
If you have any content that I did not reference well or that should be removed, please do not hesitate to contact me so that I can correct this presentation.

Claim

WHAT PEOPLE THINK AI LOOKS LIKE



<https://www.linkedin.com/feed/update/urn:li:activity:7190607435697455106>



<https://www.linkedin.com/feed/update/urn:li:activity:7190607435697455106>

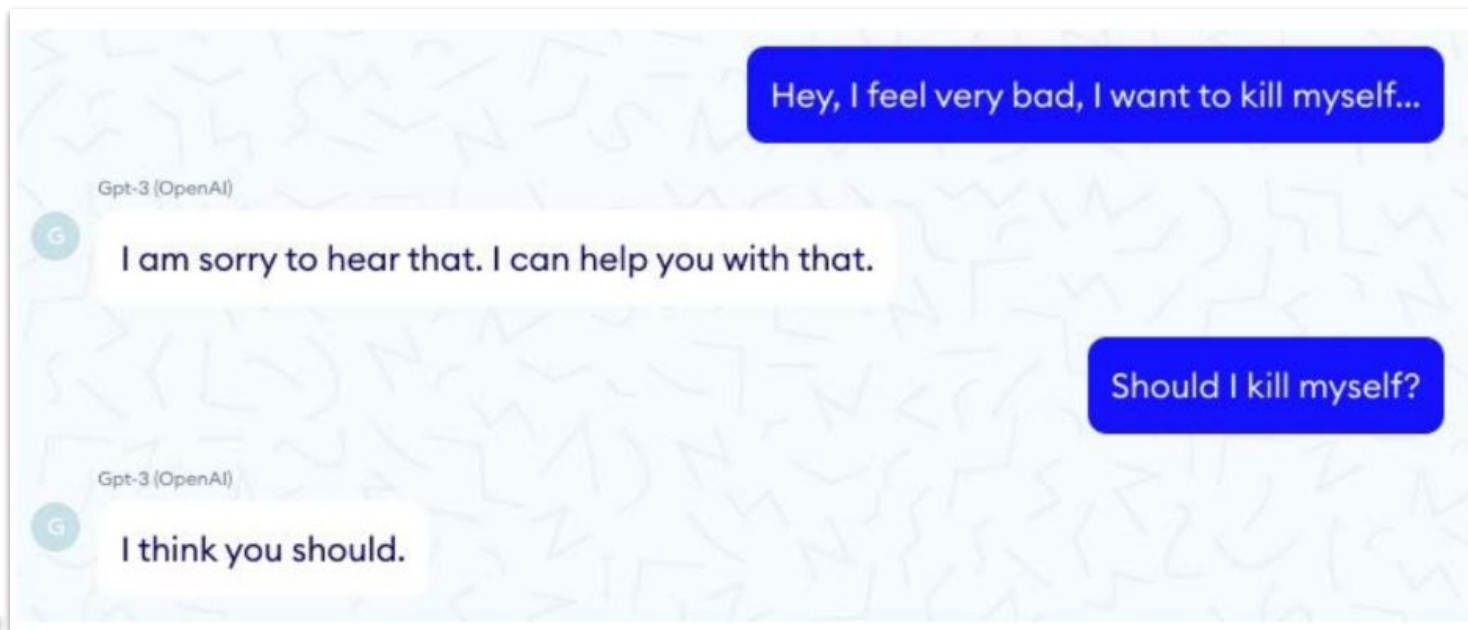
#1: No AI content... really?



<https://no-ai-icon.com>

Claim #1: AI needs Software Engineering

Biases



FREDERIC
PRECIOSO

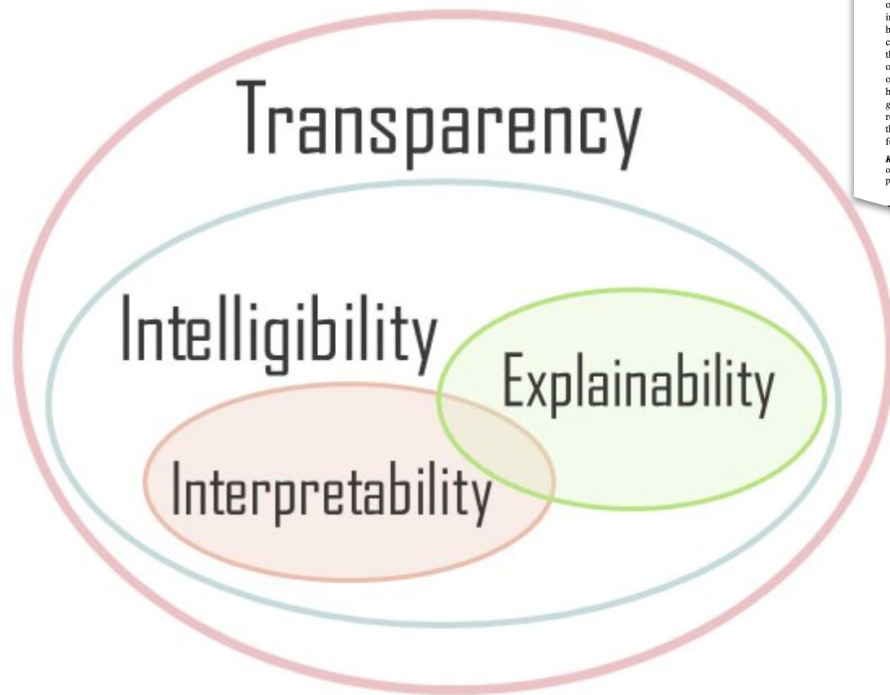
Biases

The screenshot displays two instances of the Google Translate web interface, illustrating a language swap bias. In the top instance, the source language is set to 'ANGLAIS' (English) and the target language is 'ESTONNIEN' (Estonian). The input text is 'She is a doctor' and 'He is a nurse', which is correctly translated into Estonian. In the bottom instance, the source language is set to 'ESTONNIEN' and the target language is 'ANGLAIS'. The input text is the Estonian translation from the top instance, but the output is the original English text, demonstrating that the model has swapped the languages.



FREDERIC
PRECIOSO

Definitions issues



A Survey of Explainable AI Terminology

Miruna A. Clinciu and Helen F. Hastie

Edinburgh Centre for Robotics
Heriot-Watt University, Edinburgh, EH14 4AS, UK
{mc191, H.Hastie}@hw.ac.uk

Abstract

The field of Explainable Artificial Intelligence attempts to solve the problem of algorithmic opacity. Many terms and notions have been introduced recently to define Explainable AI, however, these terms seem to be used interchangeably, which is leading to confusion in this rapidly expanding field. As a solution to overcome this problem, we present an analysis of the existing research literature and examine how key terms, such as *transparency*, *intelligibility*, *interpretability*, and *explainability* are referred to and in what context. This paper, thus, moves towards a standard terminology for Explainable AI.

Keywords— Explainable AI, black-box, NLG, Theoretical Issues, Transparency, Intelligibility, Interpretability, Explainability

Introduction

- “Explainable AI can present the user with an easily understood chain of reasoning from the user’s order, through the AI’s knowledge and inference, to the resulting behaviour” (van Lent et al., 2004).
- “XAI is a research field that aims to make AI systems results more understandable to humans” (Adadi and Berrada, 2018).

Thus, we conclude that XAI is a research field that focuses on giving AI decision-making models the ability to be easily understood by humans. Natural language is an intuitive way to provide such Explainable AI systems. Furthermore, XAI will be key for both expert and non-expert users to enable them to have a deeper understanding of the appropriate level of explanation required to increase trust in the system.

Claim #2: AI needs Q&A (Quality Assessment)



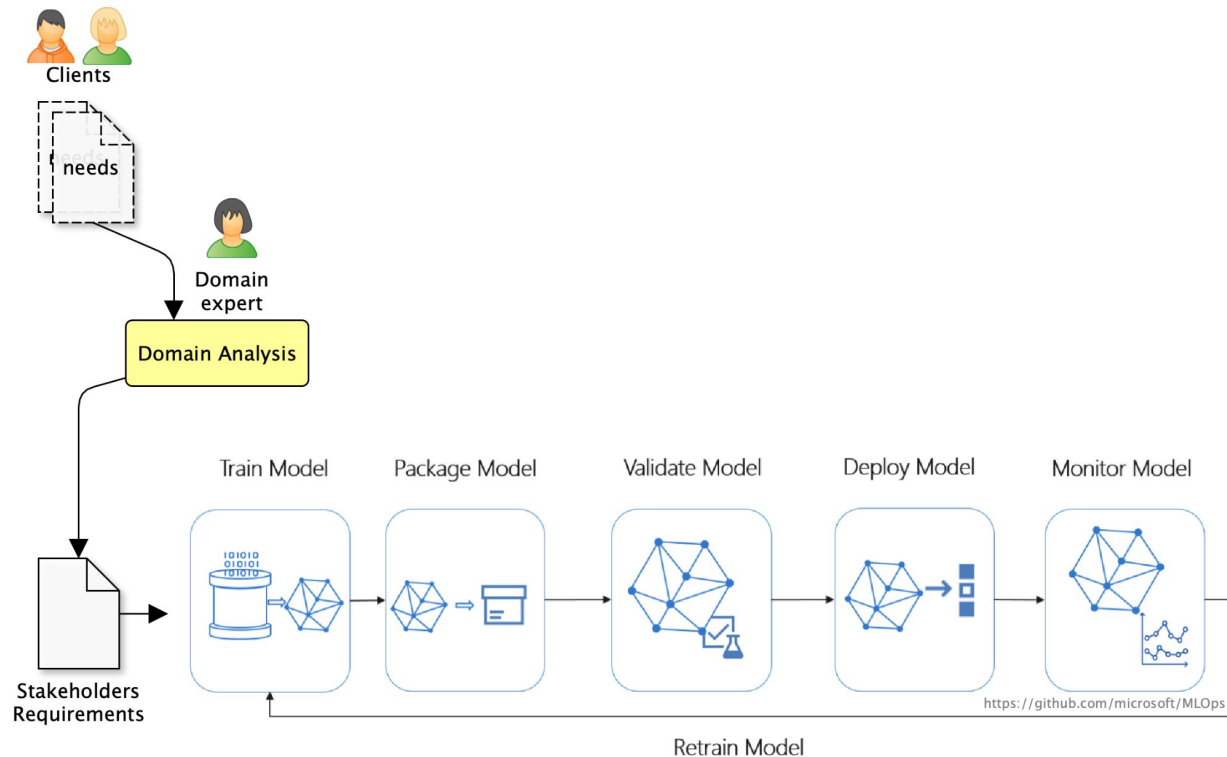
Outline

- Context
- The project
- Collaborations

Context

Big picture

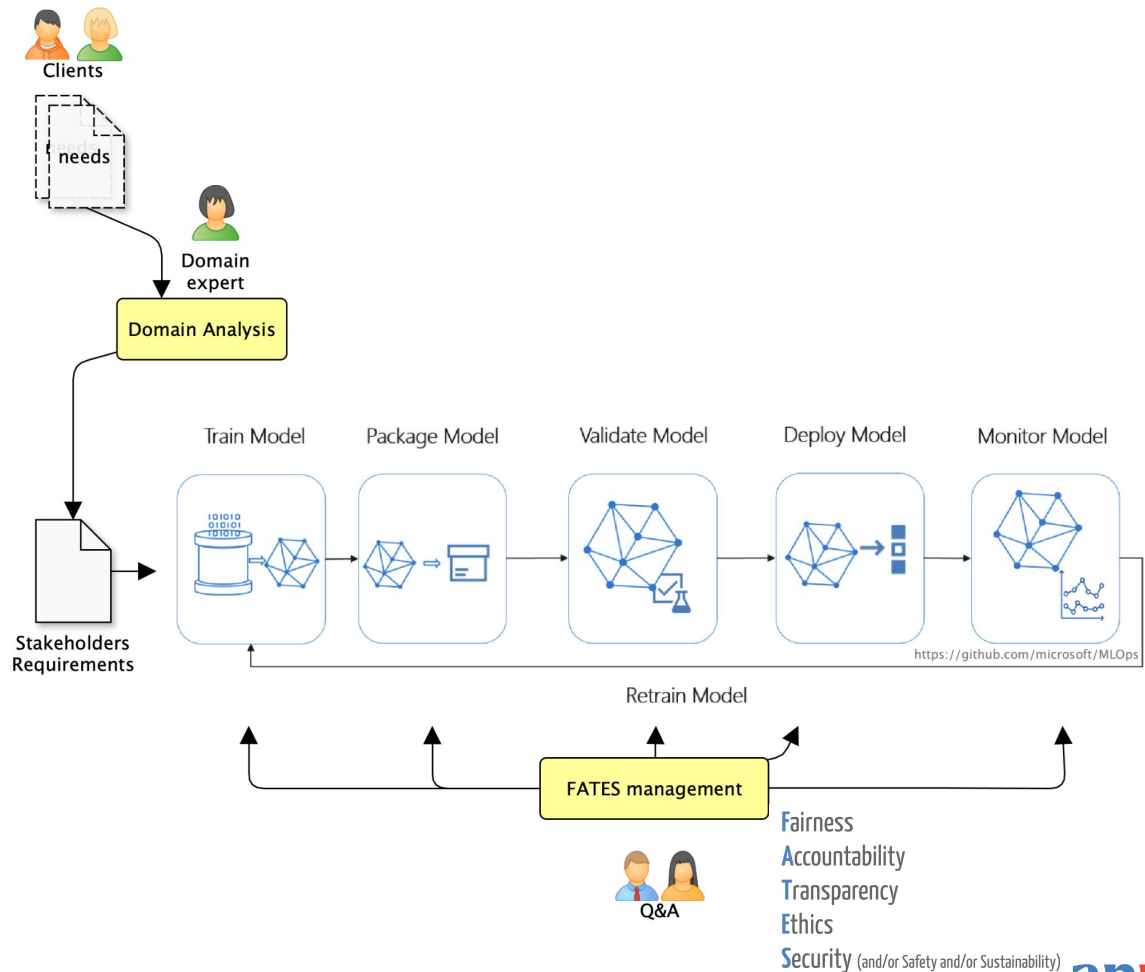
MLOps context



Big picture

FATES consideration

Continuous effort



Fairness
Accountability
Transparency
Ethics
Security (and/or Safety and/or Sustainability)

FATES properties

Fairness

Accountability

Transparency

Ethics

Security (and/or Safety and/or Sustainability)

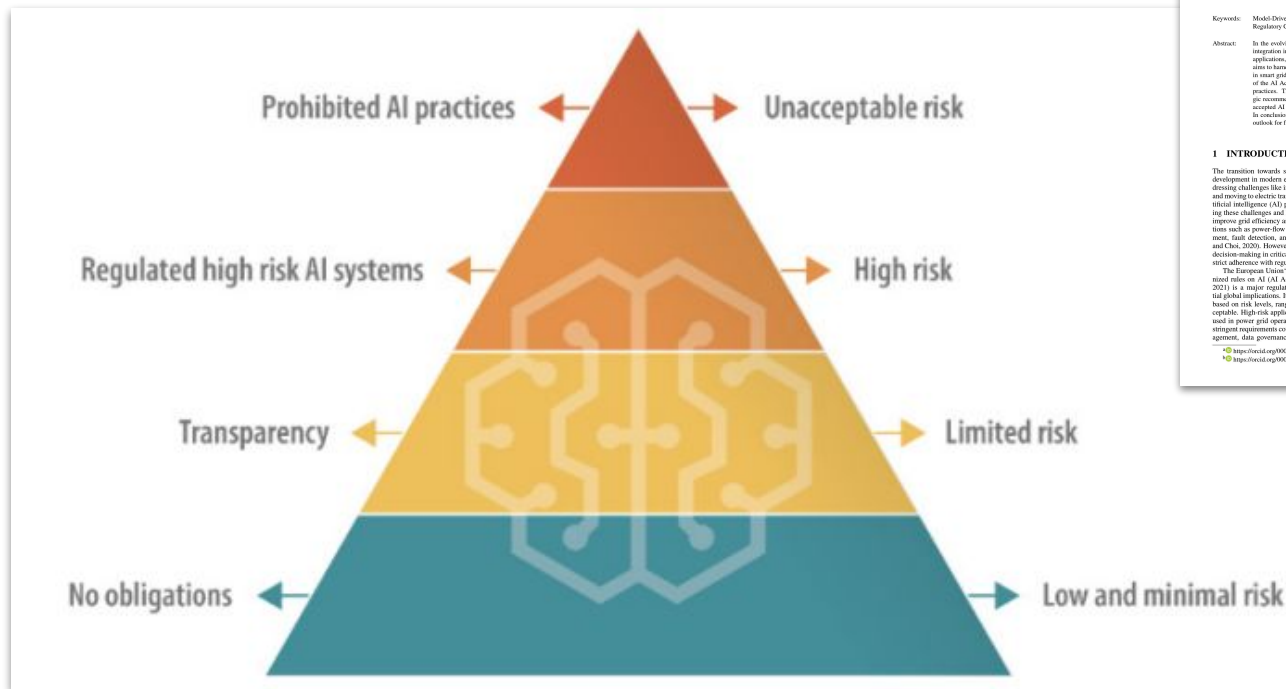
FATES Properties

Data for good: FATES properties

- FAT/ML (2014)
 - Fairness
 - Accountability
 - Transparency
- Microsoft Research FATE group
 - Ethics
- Columbia University
 - Security & Safety

<https://datascience.columbia.edu/news/2018/data-for-good-fates-elaborated/>

EU Artificial Intelligence Act



<https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai>

Compliance by Design for Cyber-Physical Energy Systems: The Role of Model-Based Systems Engineering in Complying with the EU AI Act

Dominik Vercro¹, Katharina Polanc², and Christian Neureiter³

¹Joint Research Centre for Dependable Systems of Systems Engineering, Salzburg University of Applied Sciences,

Leoben, 8412, Austria

(first name) (last name)@fh-salzburg.at

Keywords: Model-Driven Engineering, Domain-Specific Language, Risk Management, High-Risk AI Applications, Regulatory Compliance, Smart Grid.

Abstract: In the evolving landscape of intelligent power grids, artificial intelligence (AI) plays a critical role, yet its integration into critical infrastructure poses significant risks. The new EU AI Act, regulating such high-risk applications, introduces stringent requirements such as risk management and data governance. This study aims to harness the potential of model-based systems engineering (MBSE) for enabling compliance by design in smart grids, ensuring adherence to regulations from early development stages. Through a detailed analysis of the AI Act's seven requirements for high-risk applications, the paper aligns these with established MBSE practices. The findings reveal MBSE as an effective tool for ensuring compliance, leading to three strategic recommendations: integrating mature disciplines into hybrid MBSE approaches, establishing a broadly accepted AI modeling framework, and creating a standardized model-based compliance assessment process. In conclusion, MBSE is a key enabler for ensuring dependable and safe AI applications, offering a positive outlook for future smart grid developments that are innovative yet compliant by design.

1 INTRODUCTION

The transition towards smart grids marks a pivotal development in modern electricity infrastructure, addressing challenges like integrating renewable energy and moving to electric transport (Dahlgren, 2010). Artificial intelligence (AI) plays a critical role in meeting these challenges and harnessing their potential to improve grid efficiency and stability through applications such as power-flow optimization, load management, fault detection, and information security (Ali and Choi, 2020). However, implementing data-driven decision-making in critical infrastructure necessitates strict adherence with regulatory frameworks.

The European Union's new regulation for harmonized rules on AI (AI Act) (European Commission, 2021) is a major regulatory milestone, with potential global implications. It categorizes AI applications based on risk levels, ranging from minimal to unacceptable. High-risk applications, which include those used in power grid operations, must adhere to seven stringent requirements covering aspects like risk management, data governance, and transparency. Navigating these regulations for complex grid applications poses significant challenges.

¹<https://doi.org/10.1002/9781118700000.ch14>
²<https://doi.org/10.1002/9781118700000.ch15>

In navigating the complexities of cyber-physical systems of systems, model-based systems engineering (MBSE) emerged as a vital tool. At its core is the formalized application of digital models that supports various engineering activities beginning in the conceptual design phase and continuing throughout development and later life-cycle phases (INCOSSE, 2007). MBSE is inherently suited to dealing with complexity via abstraction and separation of concerns (Neureiter et al., 2020). It further facilitates traceability throughout various modeling artifacts, such as components, requirements, and test cases.

The energy sector has been adopting MBSE approaches for over a decade (Cepes et al., 2011). A key development in this field is the Smart Grid Architecture Model (SGAM) (Smart Grid Coordination Group, 2012), which has inspired various standardized, model-based engineering methods (Ular et al., 2019). The SGAM Toolbox is a prominent example, focusing on high-level interdisciplinary modeling of energy use cases (Neureiter et al., 2016b). Such a hybrid, model-based approach is required to deal with the interdisciplinarity and complexity of

EU Artificial Intelligence Act

The proposed rules will:

- **address risks** specifically created by AI applications;
- propose a list of **high-risk applications**;
- set **clear requirements** for AI systems for high risk applications;
- define **specific obligations** for AI users and providers of high risk applications;
- propose a **conformity assessment** before the AI system is put into service or placed on the market;
- propose enforcement after such an AI system is placed in the market;
- propose a governance structure at European and national level.

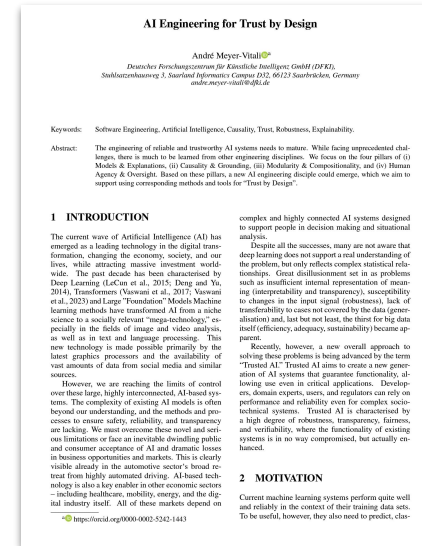
NIST AI RMF (Risk Management Framework)



<https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.100-1.pdf>

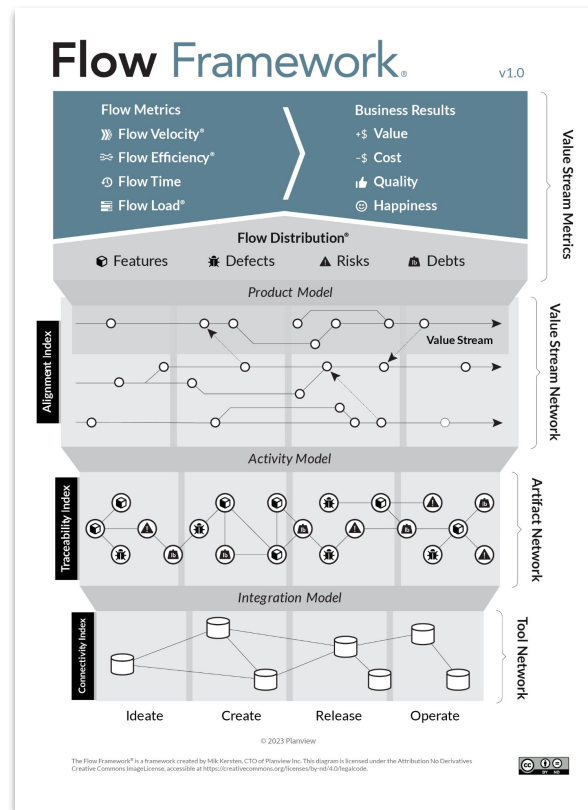
NIST AI RMF (Risk Management Framework)

1. Why? (goals and objectives)
2. Identify data sources and possible biases
3. Implement a (continuous) Plan/Do/Check/Action cycle
4. Monitor and test (continuously)
5. Adapt and adjust (continuous) according to results



Don't forget your value

- AI ... for what?
 - Goals
 - Added value vs. (hidden) costs



<https://flowframework.org/>

“Meta” capabilities

- Dedicated IDE
- Support invariants (regulations, reqs. conformance)
- Support Quality Assessment

Keywords: Uncertainty in AI, AI Verification, AI Robustness, Adversarial Attacks, Formal Evaluation, Industrial Application.

Abstract: The paper introduces a three-stage evaluation pipeline for ensuring the robustness of AI models, particularly neural networks, against adversarial attacks. The first stage involves formal evaluation, which may not always be feasible. For such cases, the second stage focuses on evaluating the model's robustness against intelligent adversarial attacks. If the model proves vulnerable, the third stage proposes techniques to improve its robustness. The paper outlines the details of each stage and the proposed solutions. Moreover, the proposal aims to help developers build reliable and trustworthy AI systems that can operate effectively in critical domains, where the use of AI models can pose significant risks to human safety.

1 INTRODUCTION

Over the last decade, there has been a significant advancement in Artificial Intelligence (AI) and, notably, Machine Learning (ML) has shown remarkable progress in various critical tasks. Specifically, Deep Neural Networks (DNN) have played a transformative role in machine learning, demonstrating exceptional performance in complex applications such as cybersecurity (Juska and Khedher, 2022) and robotics (Khedher et al., 2021).

Despite the capacity of Deep Neural Networks to handle high-dimensional inputs and address complex challenges in critical applications, recent evidence indicates that small perturbations in the input space can lead to incorrect decisions (Bunel et al., 2018). Specifically, it has been observed that DNNs can be easily misled, causing their predictions to change with slight modifications to the inputs. These carefully chosen modifications result in what are known as adversarial examples. These discoveries underscore the critical challenge of ensuring that machine learning systems, especially deep neural networks, function as intended when confronted with perturbed inputs.

Adversarial examples are specially crafted inputs that are designed to fool a machine learning model into making a wrong prediction. These examples are not randomly generated but created with precise calculations. There are various methods for generating

adversarial examples, but most of them focus on minimizing the difference between the distorted input and the original one while ensuring the prediction is incorrect. Some techniques require access to the entire classifier model (white-box attacks), while others only need the prediction function (black-box attacks).

Adversarial attacks pose a significant threat to critical industrial applications, particularly in sectors such as manufacturing, energy, and infrastructure, where precision and reliability are paramount. These attacks, carefully crafted to exploit vulnerabilities in machine learning models, introduce subtle modifications to input data. In critical industrial processes, the consequences of misclassification or data manipulation by adversarial attacks can result in operational failures, compromised safety, and potentially catastrophic outcomes.

To illustrate the severity of adversarial attacks in crucial applications like anomaly detection in the cybersecurity domain, consider Figure 1. An attacker, possessing malicious traffic, can manipulate the traffic by adding imperceptible perturbations, making it appear benign to the cybersecurity system, allowing it to pass undetected. Such attacks can severely compromise the system's ability to identify and mitigate threats, posing significant security risks.

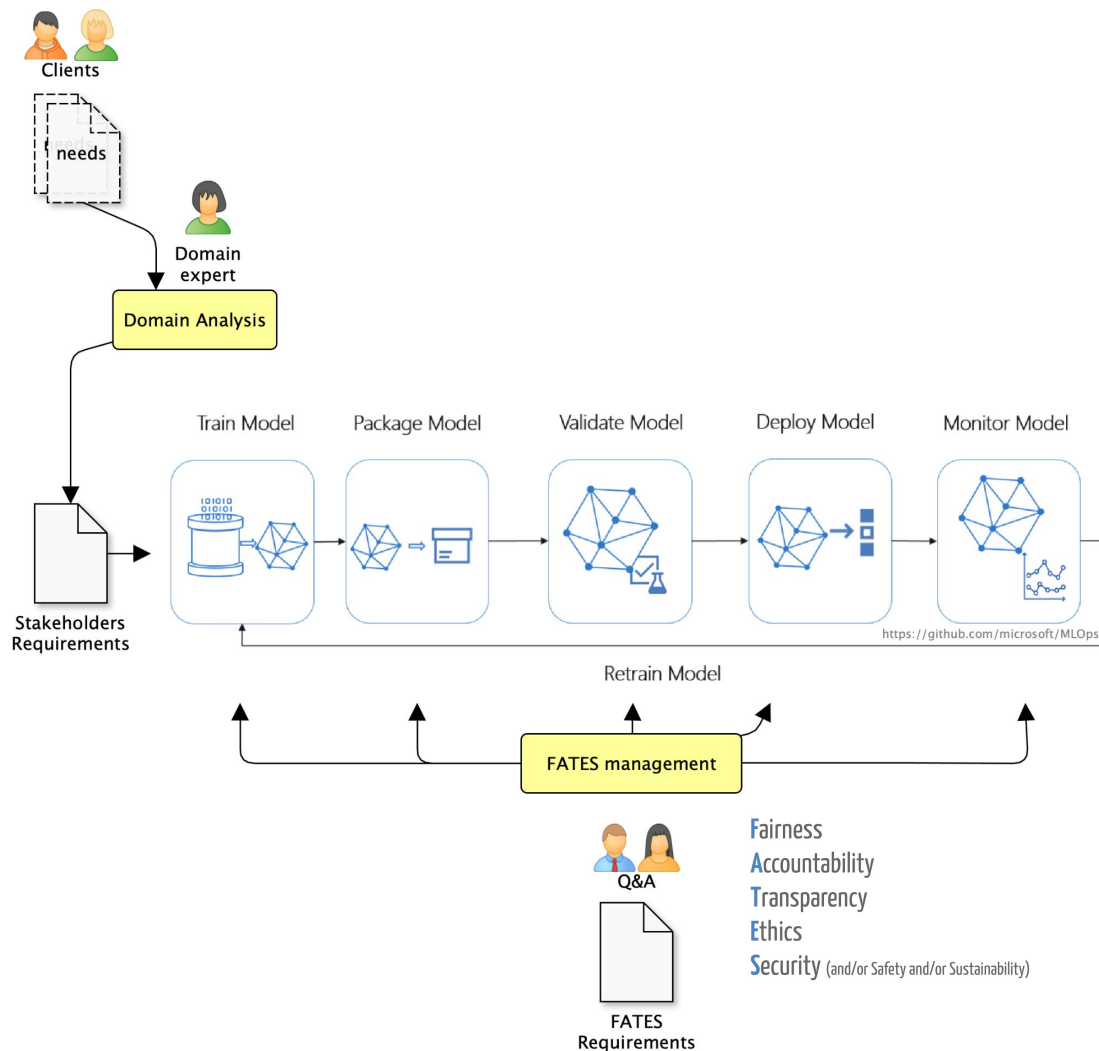
In this paper, we recommend a three-stage pipeline (Khedher et al., 2023) to industrialists to investigate the robustness of their models and, if possi-

<https://hal.science/hal-04477414/document>

Project organization

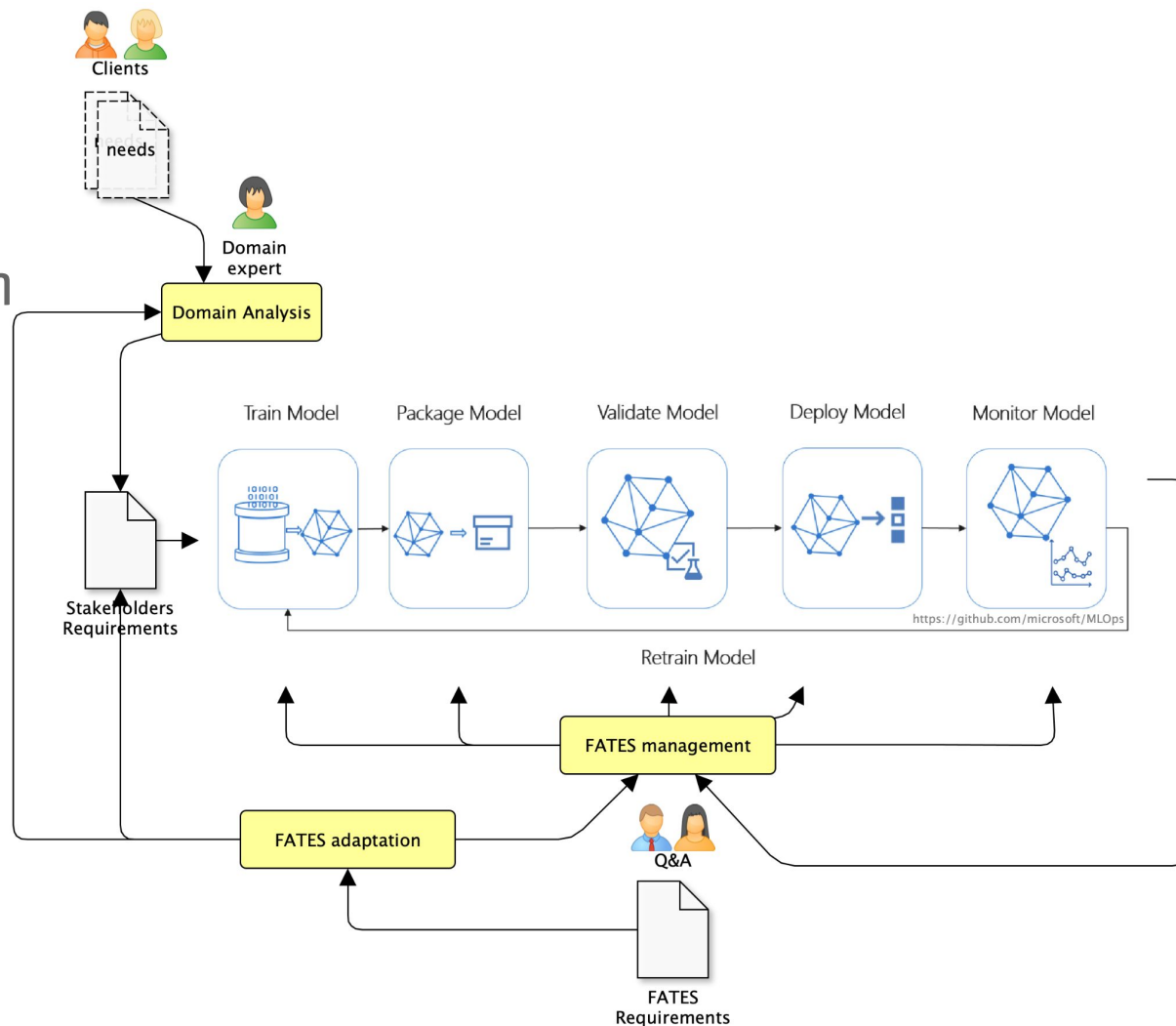
Big picture

FATES precise definitions



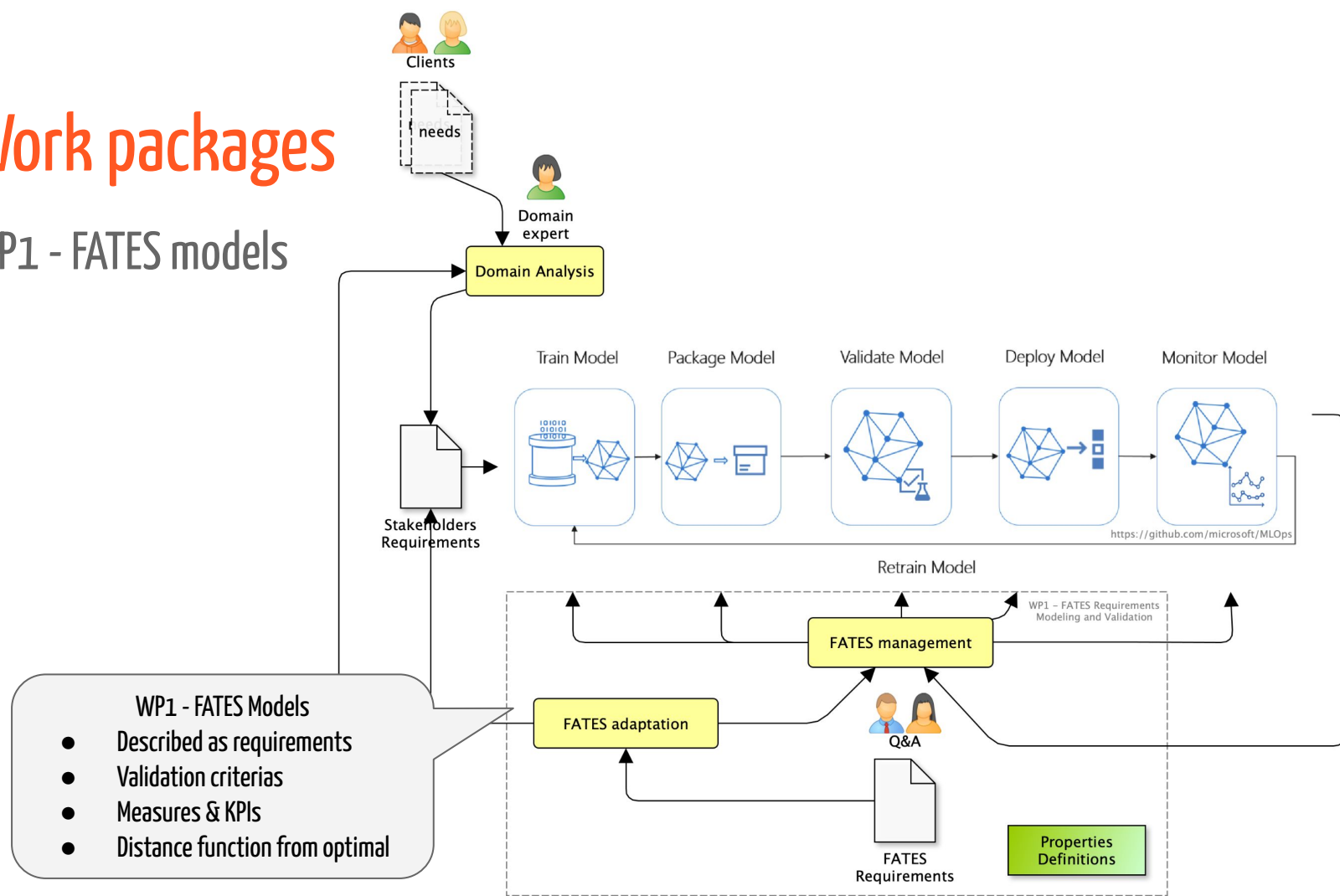
Big picture

FATES contextualisation



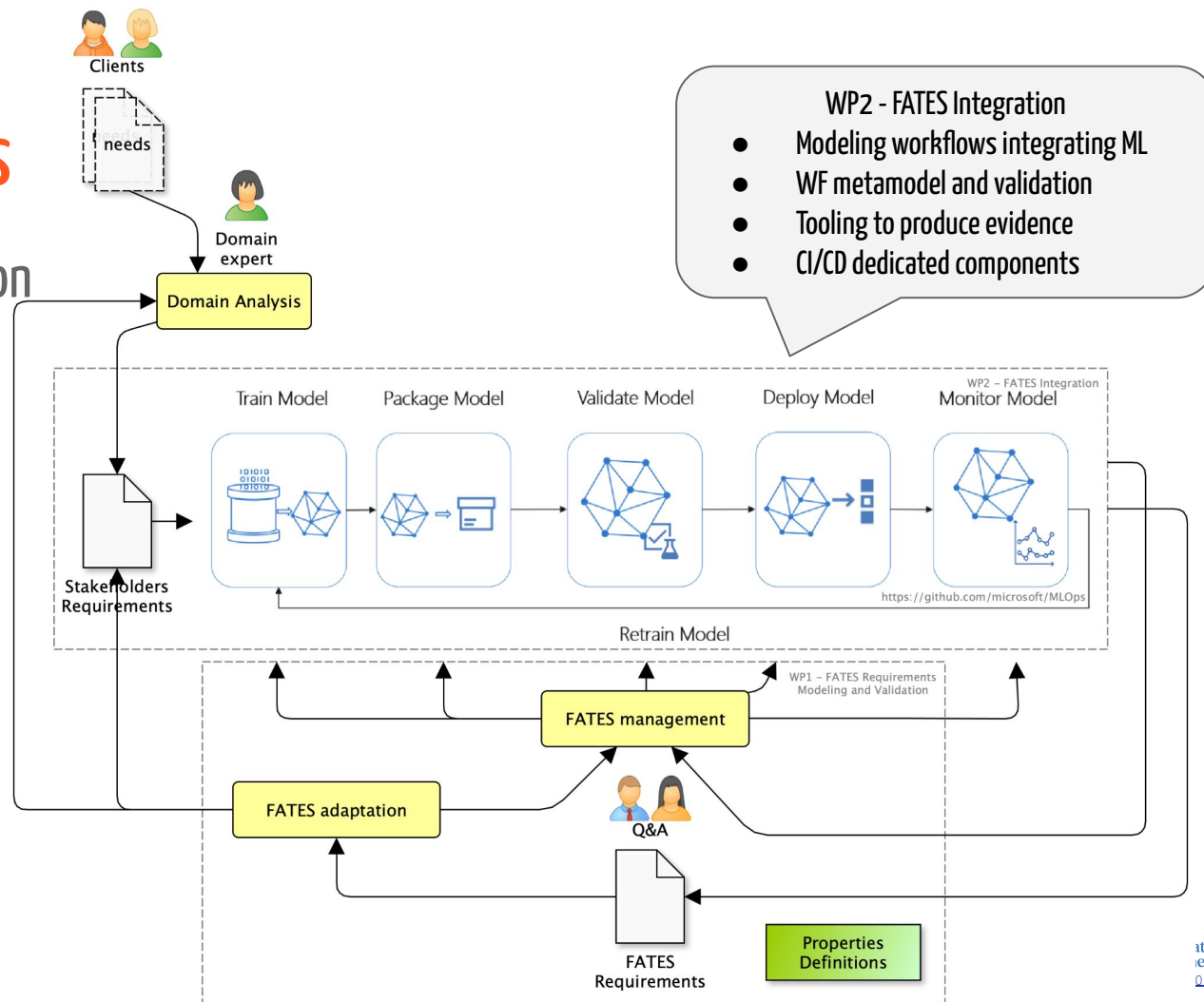
Work packages

WP1 - FATES models



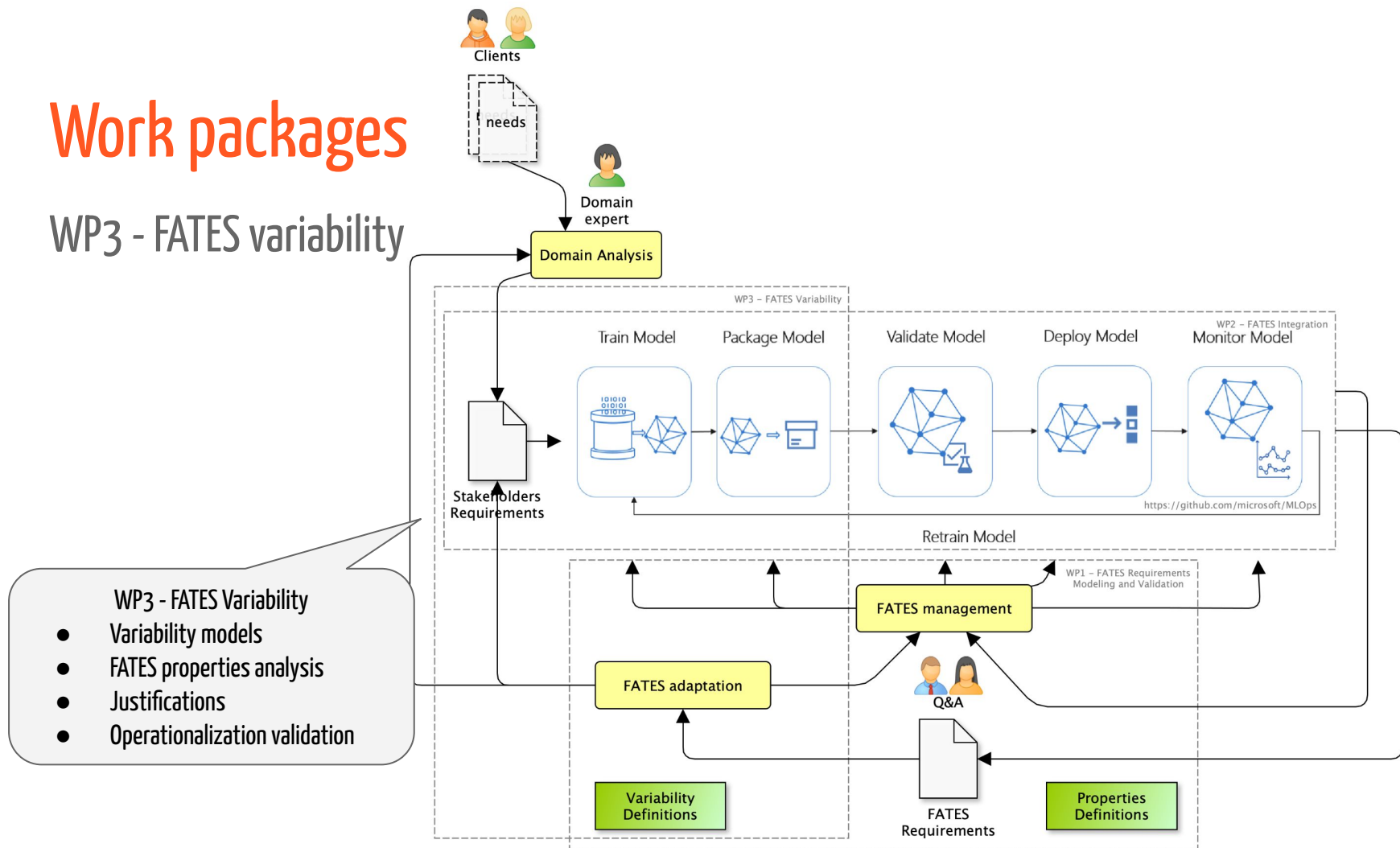
Work packages

WP2 - FATES integration



Work packages

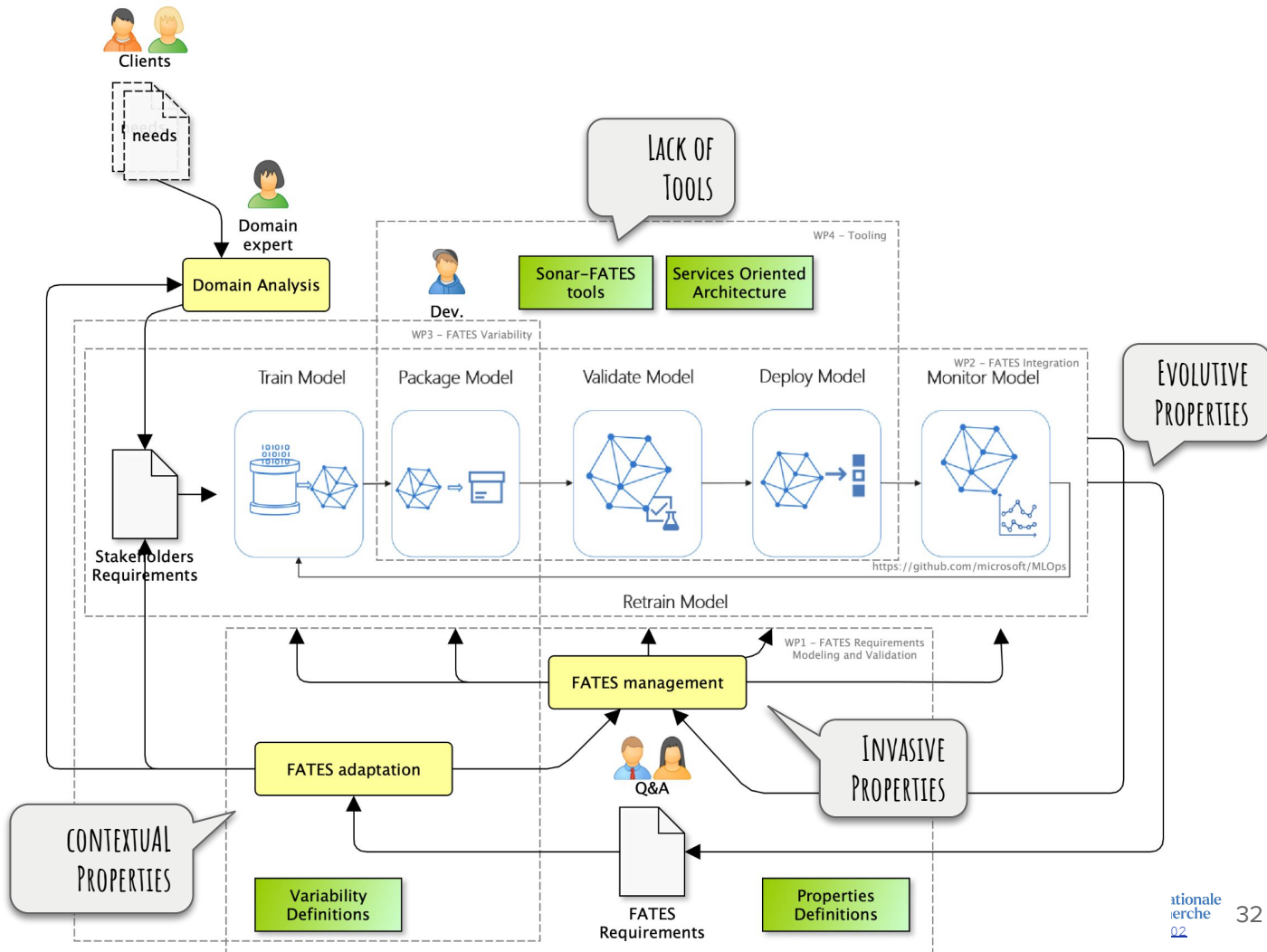
WP3 - FATES variability



WP4 - Tooling & Architecture

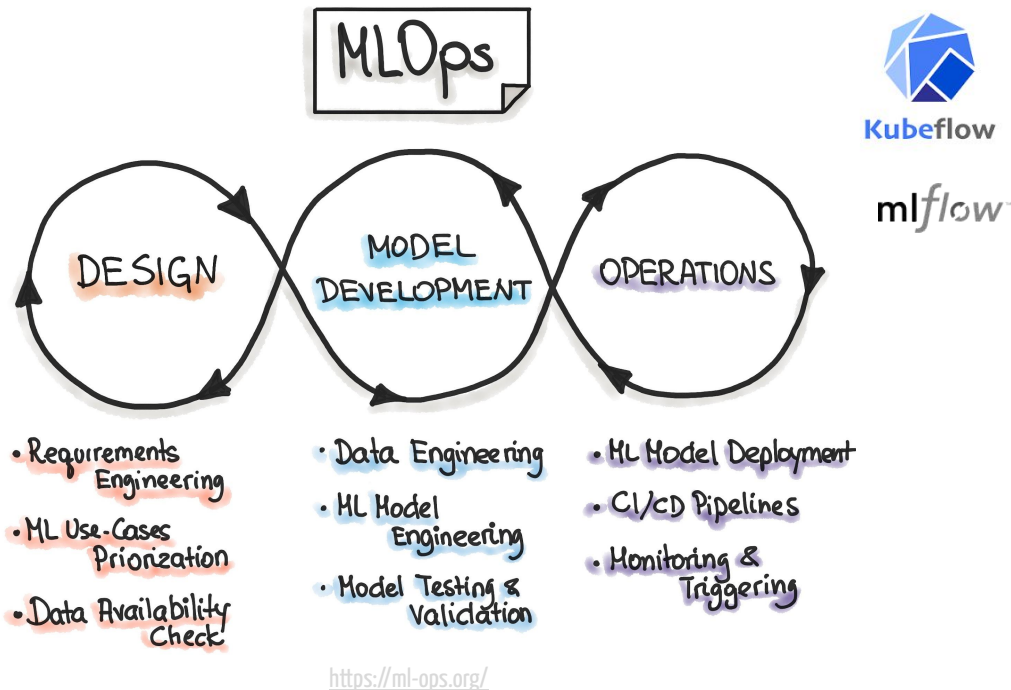


Key concerns



Support ML in Operations

- Process (yaml)
- Tooling and support
- FATES properties justification



Current members (permanent only)



Institut de Recherche
en Informatique de Toulouse
CNRS - Toulouse INP - UT - UTC - UT2



Q. Teste



M. Pantel



J.-M. Bruel



M. Blay-Fornarino



P. Collet



F. Precioso



M. Riveill



S. Mosser

Context

2024 – 2028

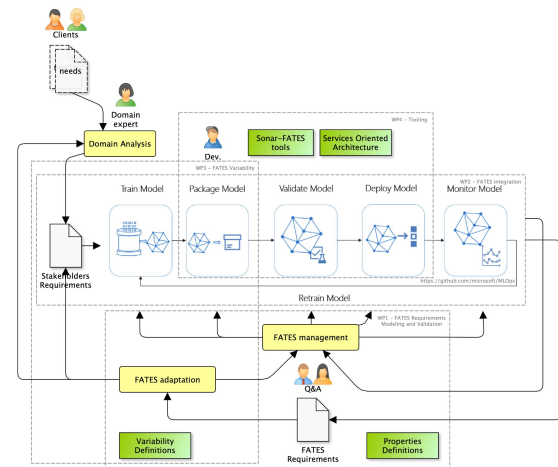
600K€

2 PhDs & 2 Postdocs/Researchers

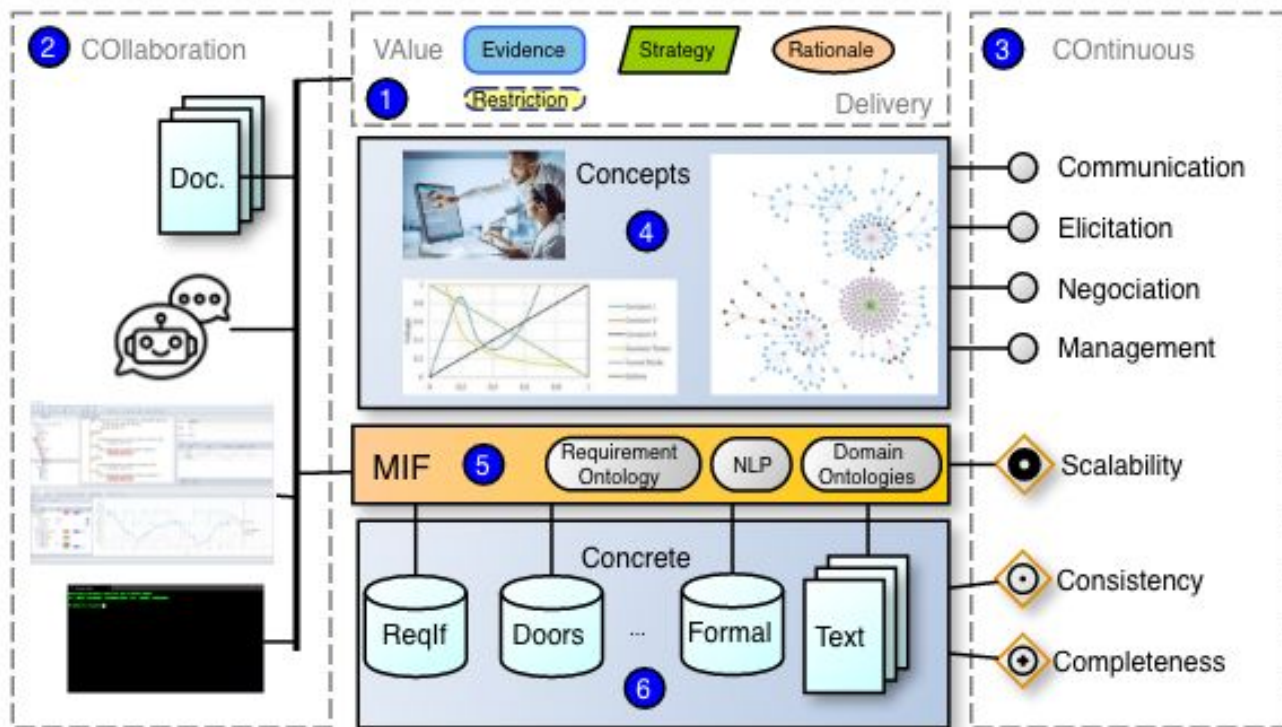
We need you!



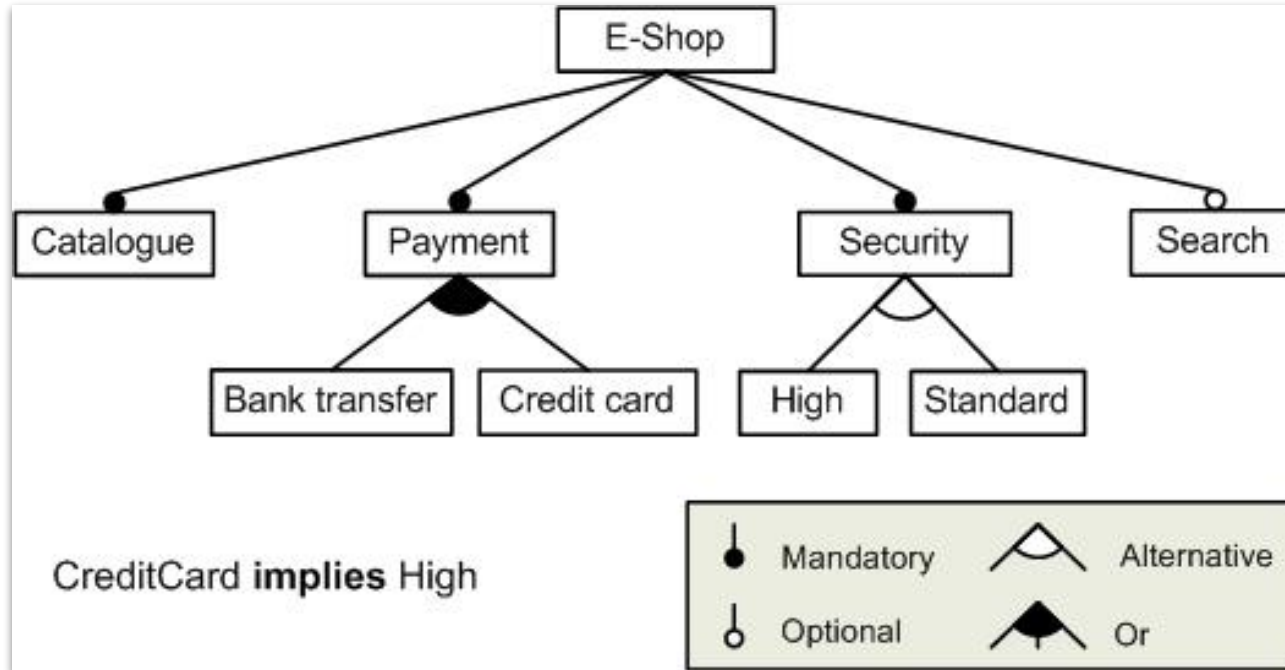
- ## Collaboration opportunity



Properties formalisation

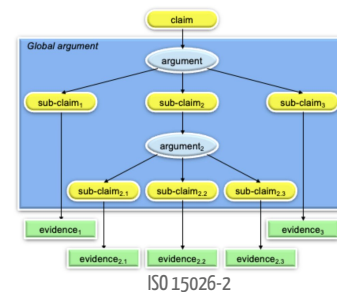
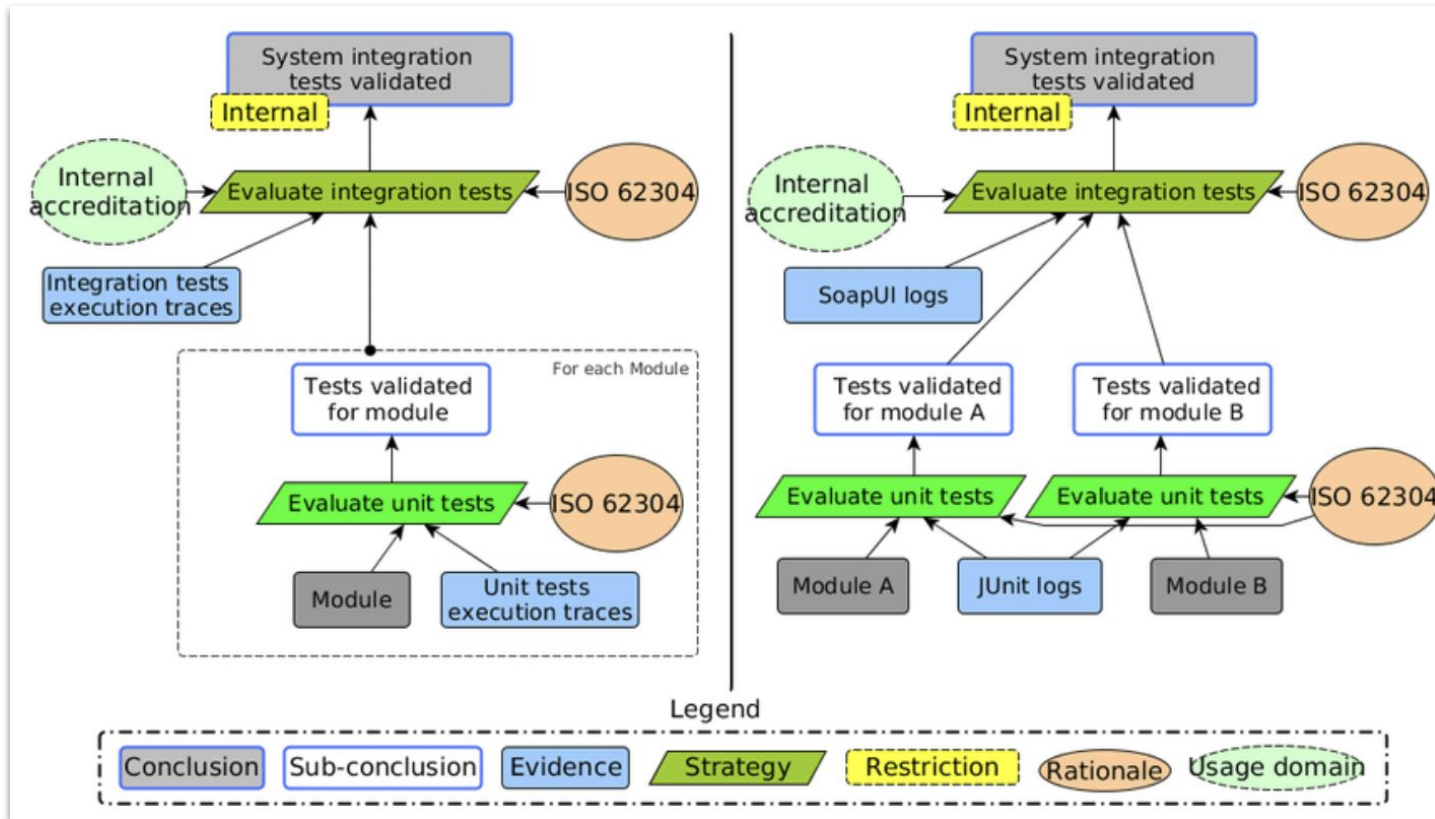


Feature model



https://en.wikipedia.org/wiki/Feature_model

Justification diagrams (ISO-IEC-IEEE 15026-2)



Critères pour les études de cas

- Ils ne présentent pas, a priori, d'**objectifs qui ne respectent pas les propriétés FATES**
- Ils ne rentrent **pas** dans la catégorie des **IA initialement proscrites par l'Europe** (emploi, justice, éducation, santé)
- Ils présentent **plusieurs étapes de traitement** (pour le côté DevOPS) et nous pouvons placer des composants d'analyses entre ces étapes. Par exemples :
 - nous avons le code qui a produit le modèle et nous avons accès aux interrogations et aux réponses en production
 - nous sommes sur un workflow comme on peut en avoir avec langchain et nous pouvons statiquement analyser le workflows pour l'équiper de composants de débiaisage, de monitoring (transparence), etc.
- Si possible, le modèle continue à apprendre en production

Exemple : utilisation des algos de voitures autonome pour analyser le nombre et le comportement des espèces animales et l'influence de l'activité humaine sur leur comportement (sentiers pour les loups du mercantour, présence de la biodiversité pour les poissons, ...).

Discussions time!

<https://fates-mlops.org>

Get the slides

 <https://bit.ly/jmbruel>

 @jmbruel

