

# **Future of Work**

## **AI observatory at the workplace**

November 2021 - GPAI Paris Summit



**GPAI** | THE GLOBAL PARTNERSHIP  
ON ARTIFICIAL INTELLIGENCE

*Please note that this report was developed by experts of the Global Partnership on Artificial Intelligence's Working Group on the Future of Work.  
The report reflects the personal opinions of GPAI experts and does not necessarily reflect the views of the experts' organizations, GPAI, the OECD or their respective members.*

|                                                                             |           |
|-----------------------------------------------------------------------------|-----------|
| <b>Introduction .....</b>                                                   | <b>4</b>  |
| <b>Methodological considerations .....</b>                                  | <b>6</b>  |
| Interview approach .....                                                    | 6         |
| Students' community .....                                                   | 6         |
| Improving the questionnaire .....                                           | 8         |
| Choice of use cases .....                                                   | 10        |
| Linking our survey with the work of the OECD .....                          | 10        |
| <b>General data on the survey .....</b>                                     | <b>11</b> |
| <b>Main findings .....</b>                                                  | <b>13</b> |
| Establishing methodological principles of a POC .....                       | 13        |
| Encourage and improve the integration of academic research .....            | 17        |
| Define the right trade-offs between usability and user involvement .....    | 19        |
| Build a situated explainability of an AI system.....                        | 20        |
| Develop a general AI training independent of a particular application ..... | 22        |
| Accompany use-cases with an independent ethics committee .....              | 24        |
| Diversify design teams to reduce bias in data .....                         | 25        |
| <b>Recommendations from use-cases.....</b>                                  | <b>26</b> |

## Introduction

To build a better future for workers collaborating with AI, to be more inclusive on various criteria such as disability, gender, ethnicity... a mandatory initial step is observation. The aim is to capture what is happening in the real context of workplaces: observe AI at the workplace, gather as diverse as possible use cases, conduct qualitative analyses of its impact in different situations, geographies, sectors, users.

The collection of use cases by GPAI ensures it will be neutral and trustworthy. These two last criteria are central as the Observation platform will allow the experts to conduct further research and, for example, analyze the reality of AI in companies through: (1) the impact of cultural specificities in the way AI is implemented at the workplace based on a large number of use cases across geographies and cultural contexts, and (2) the possible changes in the way in which AI systems are implemented from ongoing observations. This will provide insight to GPAI members on an improved human-centered approach of AI at the workplace and enlighten decision-makers, whether they are politicians or in the private sector.

In 2020, several actions had already been completed:

- The design of a questionnaire to describe use cases. This questionnaire considers the process of social integration of AI systems around five dimensions:
  - a. Motivations for the AI system implementation;
  - b. Participation of workers and their representatives in the process of defining, designing and developing the AI system;
  - c. Role of the Human-Machine Interaction (HMI) in the implementation of the AI system;
  - d. Consideration of ethics in the design process;
  - e. Impact of the AI system on employment, work and organizations
- 54 interviews in 8 countries had been conducted with AI system designers, executives, managers and users of AI at work.
- These use cases had revealed 5 main functions of AI at work: creation of new knowledge, generalization of knowledge, matching, prediction and augmentation.
- Four major observations were obtained:
  - a. A diversity of augmentation processes in organizations;
  - b. The centrality of workers in the design of a professional application of AI;
  - c. A large heterogeneity in the consideration of ethical issues;
  - d. A modification in the value of human work that, depending on the context, weaken or strengthen the role of experts.
- At that time, we had set two short- and medium-term objectives. In the short term, we were aiming to focus our efforts on the representativeness of respondents in order to create a full picture of the social effects of AI systems in workplaces and organizations, and eventually to observe similarities and differences between these actors. In the medium term, we were aiming to develop a collaborative working method between the committees of the Working Group (WG) Future of Work (FoW). The catalog would thus be able to empirically feed in-depth analysis.

In 2021, the project has undertaken three main actions:

- Improving the questionnaire by integrating the objectives defined by committees of the WG focusing on training, biases and human machine interactions. This resulted to make the survey more usable by experts for further research analysis. It was decided to introduce changes in the survey on a step-by-step basis to ensure consistency of the approach and to measure the informative value of each question.
- The creation of a students' community together with the status of GPAI Junior investigator given to each student. Students collect use cases at workplace in their respective countries, in the form of interviews. The experts of FoW oversee the students, helping them to find use cases and



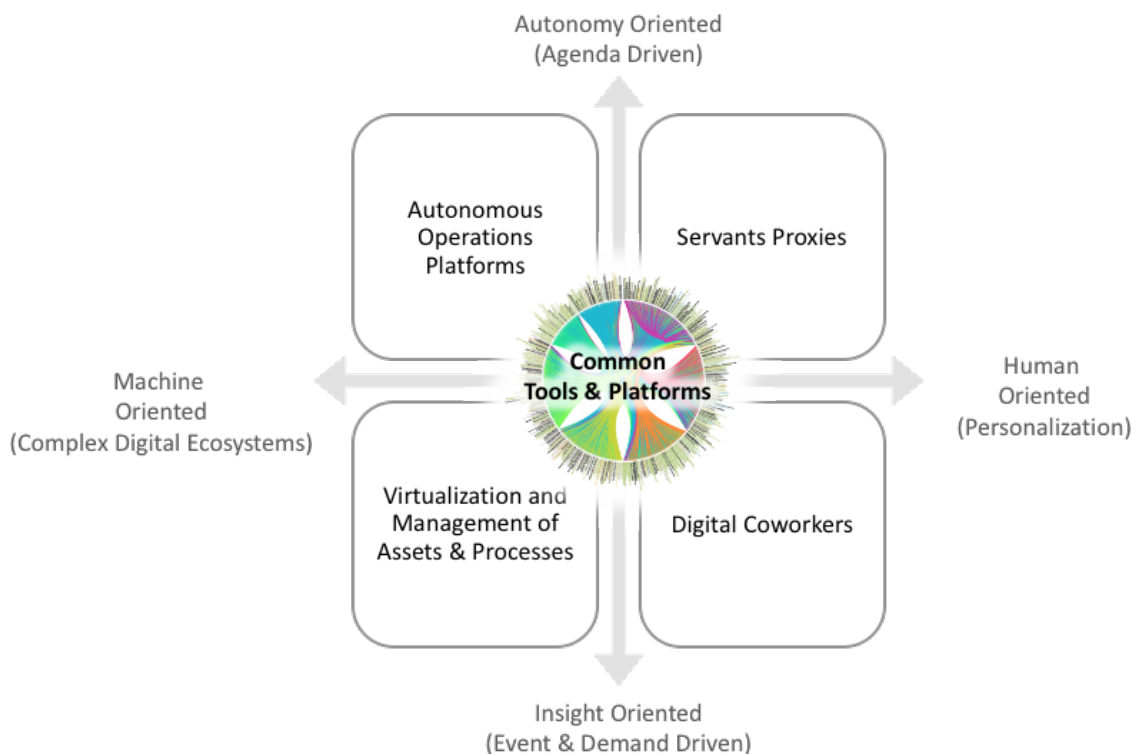
analyze the interviews. An additional activity of the students is to gather information for the next generation of students so that the community will further expand in time and include more countries. The first generation of the students' community included 5 students, each responsible for conducting 8 interviews in their own country. This allowed increasing the catalogue of use cases from 80 in 2020 to 110 in 2021.

- Organizing the survey material around a constructive AI taxonomy intended to facilitate the establishment of a better interviewing strategy and a better structuring of the knowledge produced. This taxonomy is based on two dimensions that allow classifying AI-systems. The first dimension describes what type of behavior is expected from the AI-system. The two extremes are:
  - **Autonomy** – the system is expected to act autonomously, without the need for human collaboration or supervision.
  - **Insight** – the system is expected to provide information, supporting an autonomous agent (human or non-human).

The second dimension designates the immediate beneficiary/user of the AI-system, that can be a machine, another system and a human. The two extremes are:

- **Enhanced Human** – the system delivers information, service or product to a human being.
- **Smart Digital Ecosystems** – the system provides value to a technological system.

The taxonomy is displayed on figure 1.



## Methodological considerations

### Interview approach

In 2020, the questionnaire was used in two different ways:

- Either as a standalone questionnaire: the respondent answers to the questions on a template or through an online version. An additional exchange between the respondent and the expert may also complete or clarify some of the answers.
- Or as a guide for an exchange between the respondent and the expert, who fills in the questionnaire himself.

The first approach was very effective because it allowed for a large number of responses and reduced the workload for the interviewers. However, it often produced data that were poor in quality and therefore difficult to use. We then decided to systematize the interview survey. These interviews last an hour and a quarter on average, focus on a real use case, an AI system that is being integrated into an organization at the proof-of-concept or production stage, and organized around the 5 dimensions of the questionnaire:

- Motivations for AI system implementation.
- Participation of workers and their representatives in the process of defining, designing and developing the AI system.
- Role of the Human Machine Interaction (HMI) in the implementation of the AI system.
- Consideration of ethics in the design process.
- Impact of the AI system on employment, work and organizations.

The interview process also relies on the agility of the interviewer. Rather than stringing questions one after the other, interviewers are encouraged to prioritize certain dimensions according to three criteria:

- Function of the AI system. Regarding to that function, the interview should focus on certain dimensions of the questionnaire. For example, in case of chatbots, the interviewer should concentrate on human-machine interactions while for AI for recruitment assistance, the interviewer may concentrate on bias and discrimination.
- Recurring questions. Some of them can be removed based on the first answer. For example, the questionnaire regularly points to the involvement of social partners at different stages of AI system integration. When the first response is negative, there is no need to ask such question again.
- Enlightening on a topic. The interviewer can then leave time to further develop this subject. For that purpose, the questionnaire includes a series of potential follow-up questions. The interviewer can also formulate additional questions.

### Students' community

In forming this community of students, the working group had several goals. The first is to increase the number and quality of use cases in the catalog. The second is to offer a high-level international experience that enriches the students' skills. The third is to prepare the future generation of GPAI Experts. The first generation of the student community has gathered:

- Two French students: Anne-Charlotte Mariel and Louison Carroué.
- An Italian: Sarah de Martino.
- A French woman living in Canada: Justine Dirma.
- A Mexican living in Spain: Alejandra Rojas.



All of them were engaged in a doctoral thesis. Our recruitment process respected two criteria:

- Interviewing skills.
- Knowledge about AI, especially about AI issues in the workplace.

The community of students were accompanied by the WG in two ways: plenary meetings allowed to share experiences and methodological discussions; experts from their countries helped them identify use cases.

The mandate of the first generation included several tasks:

- Conduct 8 interviews on real use cases. These interviews were to be conducted in the respondent's native language to facilitate the exchange.
- Transcribe and translate (in English) these interviews to make them available to the Working Group.
- Contribute to the analysis of the interviews. This analysis work was done during a workshop during which the students had to identify the best practices that, according to them, emerged from their use cases. The table below summarizes the workshop's focus.
- Contribute to the improvement of our methodology so that we can extend it to the second generation of the student community, which will include more students.

All students received personal feedback to allow them to situate their contribution to this report. Empowering students is one of the goals of this community.

In 2022, we hope to integrate about 20 students from new countries. The first-generation of students will be able to continue their activity, if they wish, with the status of coach.

| GOOD PRACTICES: INSPIRING PRACTICES FOR RESPONSIBLE AI<br>(based on OECD principles <sup>1</sup> )                                                                                                                              |                                                                                                                                              |                                                                                                                                   |                                                                                                                                                                                                                                  |                                                                                                                                                                              |                                                                                                                                                              |                                                                                                                                                                                                                                                                                                                    |
|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Design process: the process described seems particularly relevant to achieving responsible AI (co-design, open innovation, integration of independent external actors, participation of social partners, etc.). Questions 2,3,4 | <b>Employees' personal data:</b> The use of employees' personal data reconciles respect for dignity and system efficiency. <b>Question 5</b> | <b>HMI:</b> The HMI develop human capacities, facilitate work, decision making, integrate the human in the loop <b>Question 6</b> | <b>Factors:</b> The integration of ethics in the design and use of the system contributes to responsible AI (the respondent demonstrates that ethical issues have been integrated in this use case) <b>Questions 7, 8, 9, 10</b> | <b>Impact assessment:</b> A real upstream impact analysis has been undertaken, the method used could be duplicated to other cases... <b>Questions 11, 12, 13, 14, 15, 16</b> | <b>Implementation:</b> The implementation processes facilitate the acceptability of AI while respecting the principles of responsible AI. <b>Question 17</b> | <b>Reviews and adjustments:</b> The system is a success for the company and for the users, it is engaged in an improvement process that involves the workers, it reconciles the achievement of the goals and empowerment of the workers, it has unexpected benefits... <b>Questions 18, 19, 20, 21, 22, 23, 24</b> |

<sup>1</sup> "Inclusive and sustainable growth well-being", "Human values & fairness", "Transparency", "Robustness & Safety", "Accountability".

## Improving the questionnaire

As explained, the questionnaire has evolved to make it more relevant to the work of the various committees of the Working Group FoW. Two precautions were taken. First, we kept the original structure around the five dimensions. Second, we decided to limit changes to one theme per year. The common objective of these two methodological precautions is to ensure continuity and consistency in the collection of data in order to allow longitudinal analysis. In this sense, in 2021, the changes were primarily focused on the questions related to the HMI, which lacked precision in the previous version.

- Questions about HMI in 2020:  
*Is HMI intended – in what respect? (empowerment of employees, traceability, explainability, etc.)*
- Questions about HMI in 2021:  
*Is HMI currently involved in your work?*

With the potential follow-up questions:

- a. If the HMI technology is not yet implemented, is it intended to be applied in the company? In what respect: empowerment of employees, traceability, explainability, etc.
- b. What kind of HMI technologies do you use? (bot, chatbot, social robot, cobot or other kind?) (one to one or in group?)
- c. What kind of interactions do you have with these technologies? (In face-to-face, by phone, by internet?) (language interaction [spoken, written], physical interaction [facial, gestural, touch, multimodal] or both language and physical interaction?)
- d. Are HMI technologies useful for your work? How much of your time is spent interacting? (100% 75% 50% 25%)
- e. What is your assessment about the following issues of the work with human-like cobots and chatbots? (autonomy v. obedience, replacement v. augmentation, creativity v. dependency)
- f. If the HMI technologies do not fully meet the expected work or present some errors, do you have procedures for reporting the anomaly to management?
- g. Does the system help in making decisions? Which opportunities resulted from it? (work done easier, quicker or better)
- h. Do you like to interact with HMI technologies?
- i. Which risks are you expecting from HMI technologies? (high, medium, low or no risk)
- j. What are the most important social values (positive and/or negative) of working with human-like cobots and chatbots? (trust, transparency, explainability, tolerance, fun, traceability, scalability, empowerment, integration, security, or others)

We have also detailed the questions based on the results of the previous survey in order to improve the quality of the data and to facilitate and homogenize the work of the students' community. The following table displays the topics, the committees that deal with them in priority and the main questions related to these topics in the questionnaire. The questions in bold were included in 2021.



| Topics                         | Committee        | Main questions                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                            |
|--------------------------------|------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| AI system definition           | All              | 1. What sort of AI system is used?                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                        |
| Process of planning            | All              | 2. What are the purpose and goals of an AI application in the company? (Process or product optimization, new business model, automation, substitution of jobs?<br>3. Are workers/representative bodies involved in setting goals of the AI application?<br>4. Is cooperation with researchers / developers and external experts given?                                                                                                                                                                                                                                                    |
| Employee's personal data       | Committee 4      | 5. Are employees' personal data required for operational use or affected by operational use? (if yes, what kind of data...)                                                                                                                                                                                                                                                                                                                                                                                                                                                               |
| Human-Machine Interaction      | Committee 3      | 6. Is HMI currently involved in your work?                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                |
| The ethical factors considered | Committee 4, all | 7. Is the transparency of the AI system for the company (and for the user in the company) required and given?<br>8. How is Data quality addressed?<br>9. How is the issue of accountability addressed?<br>10. <b>Is the system auditable?</b>                                                                                                                                                                                                                                                                                                                                             |
| Impact assessment              | Committee 5      | 11. What working areas / working groups were affected in respect of the number and quality of jobs (reorganizations etc.)?<br>12. Which impact (bias)?<br>13. Were there Impacts on qualification demands and skill management?<br>14. Were there impacts on the workload, working conditions and health management?<br>15. Were there impacts regarding the use of personal data of workers (privacy, data protection and trade-offs; realize benefits to employees)?<br>16. Were there regulations on using personal data and if so, in what regard?                                    |
| Implementation                 | Committee 2, 5.  | 17. What are the required skills? What are the measures put in place for training?                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                        |
| Review and adjustments         | All              | 18. <b>Do you find that the system makes mistakes? (many, moderately, not at all)? Can it be trusted? (totally, moderately, not at all)?</b><br>19. Are there experiences, reviews and adjustments (Ex Post Evaluation)?<br>20. How is success for this use case measured?<br>21. <b>What worked less well in the use case?</b><br>22. Effects on number of jobs, quality of jobs, job satisfaction, workload, skills?<br>23. Are there unintended outcomes for workers situation and prospects?<br>24. Are there opportunities and ways to redesign the AI system and work organization? |
| Other comments                 | All              | Message to be sent to the GPAI, Question from the respondent...                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                           |

These main questions are accompanied by potential follow-up questions, that are dependent on the respondent type (user, developer, manager, social partner, all). The whole questionnaire is presented in the appendix.

## Choice of use cases

The choice of use cases did not follow any particular logic, other than following opportunities according to the networks of the experts. We did not give a sectoral priority to students but we insisted on the need to obtain feedback from end-users, as this type of respondent was very little represented in our 2020 sample. This request was however very difficult to meet:

- The majority of use cases remain in the Proof-of-Concept stage. These PoC are aimed at early adopters, selected for their anticipated adherence to the innovation. Their opinions tend to align with those of the project management.
- The project management tends to control the communication on these PoC with a strong focus on the positive aspects.

As a result, few end-users were interviewed in 2021. This is a difficulty as end-users, enthusiastic or doubtful, bring testimonies that allow to better understand how AI systems fit into the workplace.

In 2022, we will try to get more end users, knowing this will still be difficult. We can make the assumption that this will be gradually facilitated by the commoditization of AI systems in the workplace. For the moment, organizations are cautious and try to control the communication on the feedback of their first experiments. Doing this, they are avoiding conclusions on Proofs of Concept, from which we are expecting to learn lessons in order to progress towards more efficient systems. We can also rely on the on-going Fair Work with AI project conducted by the WG in order to make progress on end-user feedback.

## Linking our survey with the work of the OECD

Throughout the year, we exchanged with the OECD so that our project will be complementary to their work on AI at work. Indeed, the OECD is undertaking a sectoral survey (finance and manufacturing), focusing on large companies in Canada, Germany, Japan, USA and Austria. It was agreed between the two survey teams not to interview the same companies in the same countries. However, the schedules of the two studies were relatively different, which limited the risk of overlap. We also exchange information on our methodologies and field experiences.

## General data on the survey

As of the date of publication of this report, the 2020 catalog has been analyzed and initial findings are visualized in Figures 2 through 4 below.

Figure 2 shows the distribution of use cases per region.

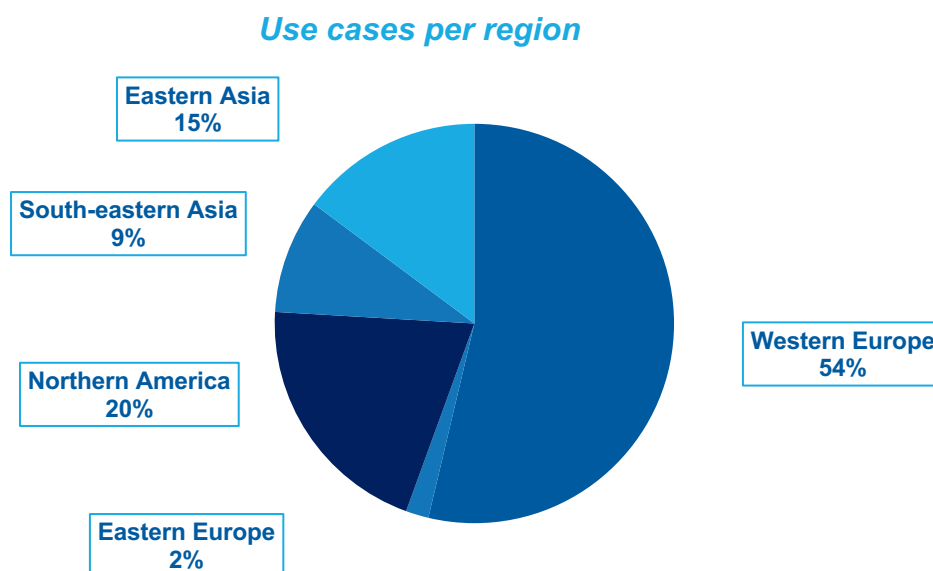
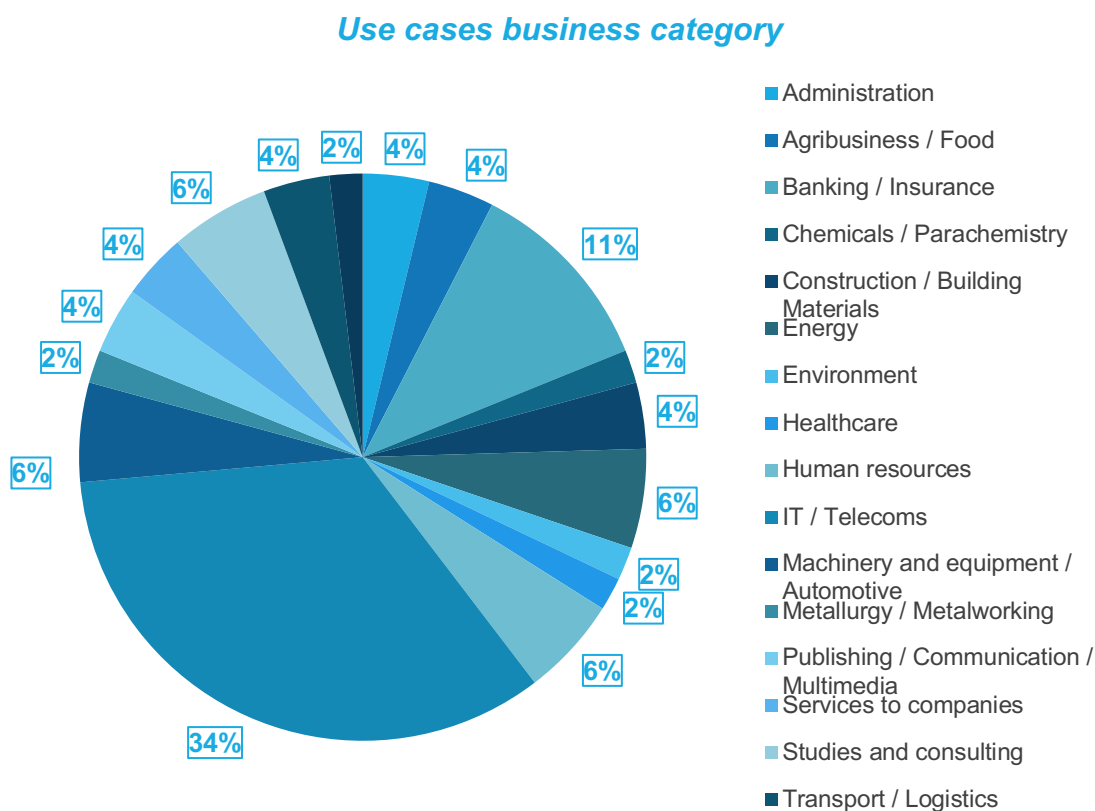


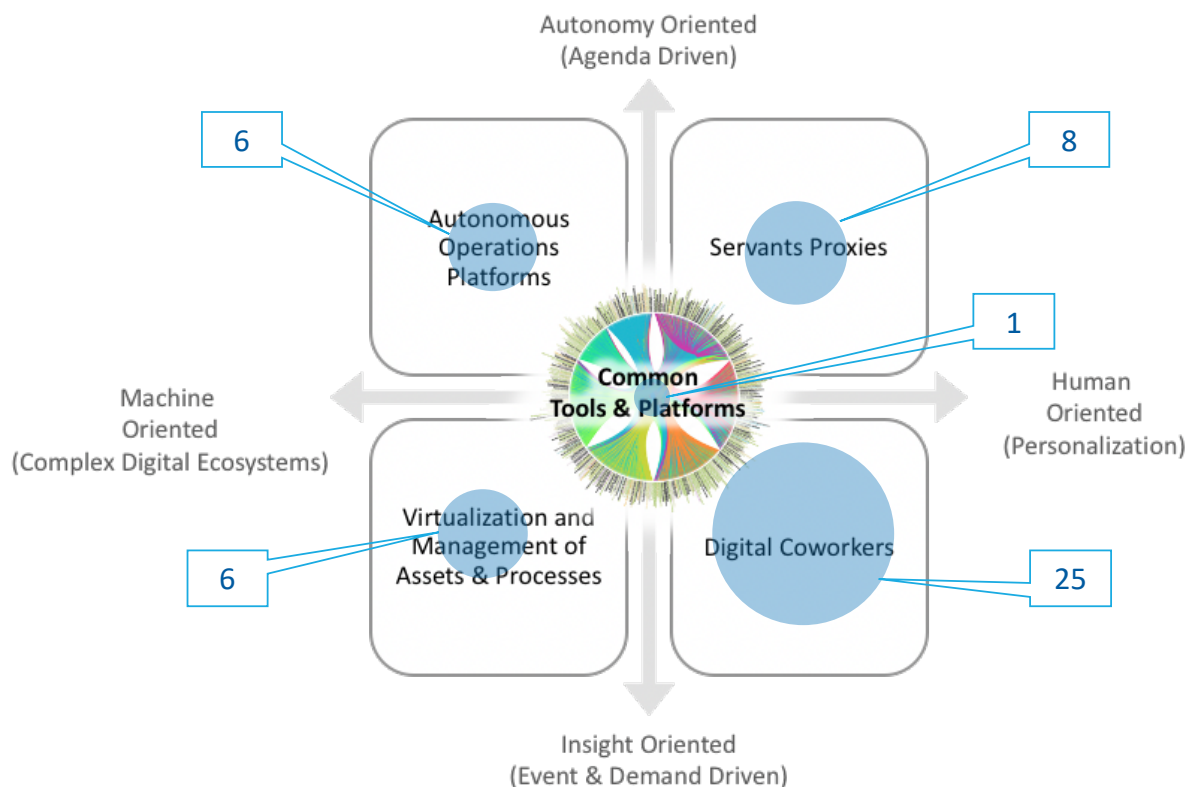
Figure 3 shows the distribution of use cases according to business categories. Seventeen categories were considered.



Lastly, Figure 4 shows the results of the classification of use cases according to the taxonomy described above. This map offers an easy understanding of artificial intelligence techniques and tools described in the catalogue. From the 54 use cases of the 2020 catalogue, 48 have been classified and 6 have not, due to a lack of reliable information. The size of the blue circles is proportional to the number of use cases of classes.

Five classes are considered along the two axes:

- **Servant Proxies** - solutions that replace the cognitive work of people in service relationships with other people, machines or infrastructure (e.g. Smart Home hubs, autonomous vehicles, digital assistants in the areas of sales and customer service, care robots, concierge robots).
- **Digital Coworkers** - solutions that expand / support people's cognitive work by providing knowledge and information supporting decision-making, solving non-trivial problems, ...
- **Autonomous Operations Platforms** - autonomous cyber-physical platforms offering technological and business services (automated factories and warehouses, autonomous transport systems).
- **Virtualization and Management of Assets & Processes** - solutions enabling the creation of digital images/simulations (digital twins) of various assets (tangible - buildings, machines, cities - and intangible - processes, systems, ...) in order to perform various types of operations on them (event prediction, configuration optimization, ...).
- **Common Tools & Platforms** - horizontal tools and platforms facilitating the creation of solutions from other application classes (ML components, low-code environments focused AI solutions, ...).



## Main findings

The project conducts a collection of real-world use cases of AI systems at work, with the objective to diversify the sample and aggregate a large number of cases. However, the following recommendations are based on interview reports that are obviously subjective. Some of these interviews have allowed to highlight inspiring practices or recurring problems shared in this report. These recommendations are far from being definitive conclusions and should be understood as empirical photographs that can help to better understand the present of AI and to better think about the future of work.

This leads to set out seven recommendations around three main themes:

- Facilitating the social implementation of AI at work.
- Empowering workers.
- Fair AI.

## Establishing methodological principles of a POC

The majority of the use cases in the catalog are Proof of Concept (PoC). A PoC is a demonstration of feasibility, i.e. a concrete and preliminary experimental realization, short or incomplete, illustrating a certain method or idea in order to demonstrate or not the feasibility, with a budget accessible to a project manager. Located very early in the development process of a new product or process, the PoC is usually considered as an important step on the way to a fully functional prototype. Thus, the organizations concerned by the use cases are currently in the discovery phase of AI through experiments whose outcomes could be decisive for a wide diffusion of AI at work.

### Most POCs achieve their objectives...

A common strategy of the project leaders of these PoCs is to target the use cases on consensual tasks: security, improvement of the available information, production of new knowledge, reduction of the drudgery of tasks, improvement of quality... The objective is that this first use case does not trigger hostility from the workers by putting them in difficulty. For instance, these PoCs concern:

- Automate a task to which workers are not attached  
e.g.: Automate the prediction of gas consumption
- Automate a tedious task that is a factor of musculoskeletal disorder or stress  
e.g.: Automate threat detection on surveillance camera videos to reduce the time spent behind screens
- Automate a task that has been poorly performed until now but that wastes time  
e.g.: Detect the presence and size of graffiti on a train before anticipating the need for cleaning at the maintenance center
- Automate a task that workers do not know how to do  
e.g.: Detect the level of humidity in a tile factory to reduce the non-quality of the product
- Automate a task with results immediately superior to the previous methods  
e.g.: Identify the presence of protected bird species in wind farms to reduce fatal accidents
- Automate a task identified by future users as relevant to reduce the repetitive nature of their work  
e.g.: Reduce the load of answering repetitive emails which represented up to 80% of the work time
- Automate surveillance tasks to refocus workers on interventions  
e.g.: Automating the monitoring of elephants in a zoo to free up time for care

Beyond the technical success, the realization of a PoC has two recognized virtues. The most obvious one is the discovery of AI and its potentialities, a demystification that only a PoC fully realizes.



**End-User interview:** *“And on the weak signals, we saw things. From a purely technical point of view, we were rather surprised by the relative simplicity of implementing a sophisticated concept. The interface used by the provider convinced me of the progress of the mastery of these subjects, the performance was not there but the handling of the tools was relatively accessible. I didn't expect AI to be so easy to do. I thought it was a universe of untold complexity. It is complex from a theoretical point of view but the manipulation of the software tools seemed relatively accessible to me. There is a certain maturity of the tools, not necessarily from the point of view of their performance but of their implementation.”*

The realization of a PoC is also a positive experience for the participants. It produces an organizational learning effect because it engages a process of formalization of the knowledge and know-how of an organization.

**End-user interview:** *“What can be considered as a beneficial effect of this PoC is that it forced us to ask ourselves questions, to think about our job, about the user relationship, about what an incident is, about what a reliable signal is. It was a stimulus, it didn't lead to a measure or a different way of doing our work, but it improved our understanding of the flows we had to manage.”*

The PoCs studied therefore generally meet the established KPIs and are good experiences for all participants. However, the vast majority of them are not continued as a project.

#### ...but do not lead to an implementation project

There are several reasons why AI systems are not deployed after an experimental phase. A significant factor was the pandemic that changed the innovation strategy of many organizations.

In a few cases, the PoC was also considered a failure.

**Manager interview:** *“It was an experiment, it was limited, but when we made the last feedback we agreed that it was not interesting to invest further in this approach because the added value was low compared to what was expected. It was an arbitrage of Return on Investment that was not sufficient. There was an interest in the thing, an interest in the approach, but there was a burden of development, work, analysis and training that was disproportionate to the results. The project did not clearly bring any perceptible added value for us. The whole economic balance of the project was not positive.  
It's difficult to analyse... Perhaps there were difficulties linked to the material, to the quality of the data, to the quality of the system. We made a global qualification by saying that it didn't bring us what we hoped for, i.e. a new vision, new points that we didn't know. What worked/not worked needs to be analysed very methodically. When we said that we wanted to detect citizens' dissatisfactions, well maybe it wasn't the right question, maybe we don't have the right material to do it, maybe there was nothing to find, maybe when there are serious events it's not on the requests that we see them coming, so no matter how hard we look we'll never find them. The aim was to look for something that we couldn't see, but maybe it just doesn't exist.”*

Except for the organizations affected by the pandemic and the rare failures, the PoCs obtain conclusive results and sometimes even perform well without leading to deployment projects. It is therefore important to understand the blocking factors that hinder the implementation of AI at work. Our survey shows that these PoCs reveal a certain number of challenges, i.e. transformation issues, which organizations that have carried out and succeeded in POCs are not ready to face.

#### Reorganize: AI systems imply rethinking the organization of the activity

The realization of a PoC is a temporary project that mobilizes different participants who are encouraged to cooperate for the success of the experimentation. To go into production, this temporary cooperation must be made permanent. This implies a reorganization of the organization's components, which is a major transformation.



The AI system provider explains how the need for data in sufficient quantity, quality and diversity challenges organizations

**Designer Interview:** *“A lot of people use AI to do R&D without worrying about putting it into production. When we put into production, we ask ourselves new questions: what data? Why do we do it? What variables? What are the optimal configurations? This is the MLOps branch (Machine Learning Operations). It introduces a new discipline in companies. We are moving from centric process to data centric. In the first case, we create the application and then we look for the necessary data. Before, you only talked to people who had the relevant data. In this case, the data lack is transverse, and we have to ask ourselves about the meaning and origin of each category of data. We are changing modes of organization; people are obliged to work together. This challenges the organization's silos.”*

In other cases, the AI system implies a complete rethinking of the workstations to make its integration possible, as demonstrated by **Case study 1**.

### Case- study 1: Rethinking the ergonomics of a workstation

This AI system is intended for incident analysis in industrial facilities where risk management issues are crucial. Two types of reports coexist: on the one hand, quantitative, numerical reports; on the other hand, qualitative comments written by the experts. AI automates the evaluation of the consistency between these two sources. This allows users to understand the gap between the measurement of risks and their perception in order to better understand them. After a co-development process between the digital transformation teams and the business experts, the application achieves an excellent level of performance. However, it is rarely used. Indeed, the focus on the application's efficiency has obscured its integration into the workstation. The agent already manages several interfaces, the AI adds one more, which makes its task more cumbersome. Despite its ability to succeed in its mission, the AI system does not fit into the activity, which it disrupts excessively. *“We have learned a lot from this experience,”* explains the project manager. *“We must now better address the issues of ergonomics, think beyond the tool to understand the entire workstation. For example, we are thinking about the use of voice to interact more naturally with the various tools, to make the human-machine relationship more fluid and less laborious”*. Parameterization is also a problem. In order to reduce the risk of error, designers tend to adjust their system to ensure that it does not miss a problem. This results in a large number of “false positives”, the notification of errors that are not errors. Put another way, AIs, especially when applied to critical issues, tend to cry wolf a little too often for fear of missing it when it actually shows up. It's not just about designing a suitable application, but also a suitable relationship to the application. The acceptance of AI is linked to the ergonomic qualities of the device, as well as to its ability to integrate with the context and experience of the user, and to arouse positive emotions. The system must be useful, usable and accessible, i.e. compatible with the needs of individuals and the specificities of their activity.

### Socialize: AI systems destabilize the value system associated with the activity

An AI system can destabilize the value system of an organization and work, i.e. a set of personal and/or collective norms and references, which influence the attitude and behavior of those who adhere to them. The integration of AI systems into the value system of an organization is essential to develop the trust of workers. This is particularly the case with the use of chatbots. All the designers fear the system could go wrong and be made public: *“The answers generated by the bot are not the result of machine learning, only of the symbolic. In other words, we wrote all the answers in advance, the bot doesn't formulate them by itself. Our client wanted to avoid the black box and a “Tay” experience, i.e. a chatbot that becomes racist and sexist by learning from users”*. This solution reduces the risks for the client but also for the designer: *“we already had, once, an inappropriate answer. The user photographed it, and it started circulating on social networks. This is a nightmare for us, because it destroys credibility and trust in our product”*.

In other use cases, the AI system poses problems of attribution of responsibility in case of problems. In **Case study 2**, the PoC did go into production, but after significant difficulties that led to changes in the initial project.



## Case study 2: An irresponsible cobot

This AI system is associated with a cobot that assists laboratory workers in a yogurt production plant. Its articulated arm is equipped with a vision system to read barcodes and a learning module to improve its ability to grasp objects. Its mission? To take care of all the so-called "low added value" tasks of the process: to recover the products in the boxes, to scan, to check in the database that the product held is the one to be tested. He opens the product and reproduces the operators' protocol. The cobot automates repetitive tasks to reduce musculoskeletal disorders. The operators keep the most valorizing dimension, the analysis, they use their head, not their arms and hands: "There were forbidden words, explains the project manager: "We were not allowed to talk about intelligence. It is the operators who are intelligent, not the machines. The analysis remains 100% human, we could perhaps have automated it, but that was not desired. However, the social acceptability of this solution is fragile and takes time to build. Indeed, the manufacture of dairy products is subject to strict sanitary constraints that generate strong issues of traceability and responsibility. However, the cobot generates a zone of lawlessness that slows down its acceptability. Who is responsible if it makes a mistake? Which department is ready to take his guardianship? Negotiations are tough because the consequences of health scandals in the sector are serious. Finally, the quality manager agreed to act as guarantor to unblock the project. But the cobot poses another difficulty. He is technically designed to evolve among his human "colleagues". However, this property is not part of the criteria for the security certification required for such integration. The legal system does not know how to handle an accident at work with this robot, because it cannot define responsibilities to repair the damage suffered. "This has resulted in a very safe implementation. The cobot is confined in a delimited space, which can be viewed but is not fenced in, as is usually the case. It remains visible thanks to the bay windows so that it can also be a showcase for everyone to see. The professional community is not ready to welcome the cobot completely. The cobot is not enclosed or hidden, but it is isolated. Liability issues are a major barrier to the social entry of artificial intelligence into the workplace.

### Practice: AI systems transform, generate or destroy professional practices

The systems transform the activities of a job. It can also produce new practices or destroy them. In this way, we look very concretely at what technology "allows/enables" or "forces" us to do, but also at what it "prevents" or "does not allow" us to do, and this, on different dimensions of the activity.

- **The AI system generates an uninteresting task**  
e.g: This chatbot helps employees to identify their soft skills. Based on this assessment, employees can identify opportunities for career development within the company. When the chatbot doesn't have the answer in its knowledge base, an HR employee receives an alert that a conversation has failed. The employee must then enrich the knowledge base by writing an answer to the question. This mission has been named "AI trainer". The first AI trainer was associated with the entire PoC. He was interested in this project, which made it easier to accept this additional workload. Although not very interesting, it made sense with the whole project. When this employee left, his replacement found no interest in this task and gradually abandoned it. The chatbot was closed because it was no longer performing well.
- **The AI system destroys professional practices**  
e.g: This AI system is a voice assistant at the workstation for technicians in industry. They interact with it, usually via a headset, to remedy situations where the lack of information or utilities forces them to an unwanted interruption of their activity. The voice assistant interfaces with a business knowledge base, corporate information, email or the user's phone, and also includes tools such as a calculator, countdown timer or measurement converter. It can also be used to generate feedback or to verbalize observations that are captured in files, all in natural language, without leaving one's workstation and stopping one's task to grab a tablet, computer or smartphone. The company that designed and markets it guarantees gains in productivity, quality and safety, i.e. consistency in industrial performance. But for this customer, a textile manufacturer, this promise is not entirely acceptable. The voice assistant could allow technicians to report data in spreadsheets directly from their workstations, without going to the dedicated data entry room. But this is not the only function of this room, which is also an "airlock" for breathing and exchanges. The production environment is indeed noisy, which causes auditory stress and reduces the possibilities of discussion. This room offers a punctual improvement of their working conditions, allows a break without being a break. The technicians are attached to it; removing it would certainly make them more productive, but at the cost of affecting their well-being at work.



- **The AI system transforms professional practices**

e.g: In this aeronautical company, an AI system was designed to facilitate a sensitive operation, the adjustment of a door. Normally, the operator has to make about ten adjustments to reach a "tolerance position", which specifies the desired location, as well as the authorized deviation from this standard. The application must reduce these adjustment occurrences to two by taking into account more data than it is accurate. The companion is equipped with a tablet with a main screen that shows him the schematic of the door. He fills in certain areas with numerical values. A calculation starts, and then the application tells him what to do, i.e. which actuator to touch. The application then asks him to re-measure the same values. The AI system tells him if he is within the tolerance. The interface is textual. It takes the form of an instruction. At the beginning, the users were associated with its conception, until they sensed what was at stake. Now, this sensitive operation is available to all.

**End-user interview:** *"We know that the people who do this have a very high added value, but we don't want the setting to rest on them. The intelligence is in the machine, managers want to be able to put anyone on the task. The person is « the hand » of the application. Despecialization and versatility degrade the value of expertise. We dissociate the skill from the job. For me, the system should flatter the expertise instead of reducing it. [...] We saw that the users continued on paper during the test phase, they did not use the tablet. There was 10 and 20% use. This is derisory compared to the time spent on development."*

Our study shows that the performance demonstrated by an AI system is a necessary but not sufficient condition. AI systems challenge organizations, and the difficulty of these challenges is a barrier to the sustainable entry of AI systems in work environments. It is therefore necessary to broaden the measure of success of a PoC to include extra-technological issues on at least three levels:

- **Organizational issues:** today, PoCs are tested in simulated organizational conditions that do not correspond to the ordinary organization, nor the necessary organization correlative to the implementation of the system. Considering these issues from the PoC phase will facilitate the scaling up.
- **Socialization issues** of the AI system: upstream and during the PoC, a mapping of the values of the organization and the professions is necessary to identify and anticipate the social penetration issues of the AI system.
- **Professional practice issues:** the realization of a PoC can be an exciting project that obscures the medium-term impacts on work. The PoC can also crystallize power issues that will contribute to discredit the AI system.

## Encourage and improve the integration of academic research

Many use cases are the subject of cooperation with academic research. These partnerships enrich the projects in two main ways:

- **Partnerships focus on technical issues.** In this case, the purpose of the partnership is to improve the performance of the AI system itself.

**Designer interview:** *"AI cannot be developed without the Universities. They provide knowledge of areas that I would not be able to exploit by myself, because they give you a spectrum of knowledge, as a programmer you come out with knowledge of cameras or other technical things. They explain what the real problem is, what data is needed, how they have extracted the data, then you have to sort and filter to consume it within the model. They tell you everything, how they do things and you help them".*



**Designer interview:** *"Cooperation is given in our company. We cooperate with developers and data scientists that are experts such as university professors. The same algorithm that we use has been created by a professor and a group of researchers. We are currently working in synergy with several groups of researchers, e.g. university employees, that are not anymore in the academic world but who are interested in continuing the study of AI systems".*

**Manager interview:** *"Yes, despite our knowledge on dam technology and energy production, we have no knowledge in AI, so the external researchers were essential for the project development".*

- **Partnerships focus on societal issues.** In this case, the purpose of the partnership is to address social issues and to improve its acceptability.

**Designer interview:** *"We worked on the user-centred approach with a professional ergonomist. We went through the literature on two subjects: chatbot and industry 4.0, from the angle of putting humans at the heart of the industry. Within the framework of European projects, we also work with a sociologist and an anthropologist specialized in the integration of AI systems in work environments. They are in charge of experimenting our solution in factories and interviewing users. The goal is for us to better understand how workers react to our solution in order to evolve our positioning and parameterize certain aspects of human-machine cooperation".*

Despite important contributions of academic research to the development of AI systems dedicated to work, three types of reservations are expressed about this collaboration:

- The interests of economic, practical actors diverge from those of researchers, academic.

**Designer interview:** *"We don't need to work with researchers because our issues are very practical. Our interests are not aligned".*

- The operating methods of economic actors and academic researchers are too different. This penalizes projects.

**Manager interview:** *"A development project must follow a schedule, with milestones that set deliverables. Researchers don't know how to work in industrial project mode".*

**Designer interview:** *"We don't do it because our projects must be quickly achievable. Research deadlines are too long, they are not profitable in the short term".*

- The researchers' contributions do not meet the standards of transparency, explainability and auditability imposed by the organization.

**Manager interview:** *"I included researchers so that our conversational agent would improve the understanding of word meanings, beyond syntactic analysis. I wanted lexical analysis. It was about understanding meaning better, such as irony. Sometimes a "thank you" has different meanings. Their solution seemed to work, but their system was not at all transparent. They were taking hundreds of conversations to train their algorithm, but they couldn't really explain to us how they got to that level of performance. Our company made certain requirements that led us to reject the "black box". We couldn't integrate their system".*

- Researchers are not sufficiently interested in the application dimension of AI systems and do not take into account the business.

**Designer interview:** *"We were competing with a group of math researchers who had no knowledge of the business and weren't looking for any. I asked a lot of questions to learn, and we had a person on site who connected the trade to us. That's why we won".*

## Define the right trade-offs between usability and user involvement

Like digital applications, AI systems have quality of user experience requirements. It is generally accepted that the acceptability of the system is increased by the quality of the human-machine interaction. This quality is often synonymous with "ease of use". The easier the user experience of an AI system is, i.e. fluid, user-friendly, intuitive, ergonomic, the faster the AI system will integrate professional practices.

**Manager interview:** *"It is an application, an interface that we created and adapted to our needs. We often have discussions about the display and the wording, what is the simplest for the user".*

**Manager interview:** *"There is HMI, and for me it's fundamental for the operation. Specially for the image cases, the AI systems results are easy to be interpreted. When working with outputs as graphs or tables, it's not that obvious. It's a challenge for the development team to make the interface as easy as possible for the user, especially when presenting numerical data. It must be as more interactive as possible with a simple and straightforward for the user".*

It is clear that user experience is a parameter to consider in the design of an AI system. However, the user experience of an AI system at work and the user experience of a digital application do not address the same issues of performance, accountability and user empowerment. If user-friendliness, intuitiveness and ergonomics can enchant the user, they can also generate a certain passivity and lead to a disengagement synonymous with disempowerment. In this sense, it is important that HMIs consider good levels of compromise between ease of use and the cognitive engagement of the user.

### Case study 3: Empowering the professional

#### *Do not feed the competition between professional and machine*

This application of AI in dermatology assists general practitioners in the diagnosis of skin cancer, achieving the best results in the world: "We want to augment the human with AI. We want to accelerate processes, deliver targeted knowledge at a specific time." The system, designed by a team of doctors in mathematics and artificial intelligence, uses neural networks to recognize and analyze images. This is a sensitive activity for the profession, as the marketing of AI applied to the medical field communicates on detection performances superior to humans. But the designer of this system refuses to feed this competition: "We try not to tend towards the confrontation on the performances, it impacts negatively the profession. We never compare our performance to that of the professionals".

#### *Joining forces with professionals*

This choice also justifies their positioning in the therapeutic relationship: "If we had chosen to make an application for the general public, we would certainly have had an interesting early detection, but at the expense of the professionals who would have been weakened. They have studied for many years to reach this level. With us, AI becomes the ally of the professional, it is he who has the strength of AI on his side. The consumer should not have the power of diagnosis". The augmented patient would also have been alone at the time of the cancer diagnosis delivered by the AI: "He would have received a very anxiety-provoking information, being isolated, with initial reactions dictated by fear. We would have had two weakened people, which is not in line with our vision of AI. With our application, the doctor has more time to discuss with the patient, the time saved is invested in the relationship. We want to strengthen the therapeutic alliance, with a more qualitative consultation."

#### *Organizing diagnostic interaction*

The interaction between the doctor and the tool is guided by similar concerns. The professional takes a photo and receives a risk analysis in the form of a color: green, everything is fine; orange, additional analysis; red, mandatory transfer to a dermatologist who prioritizes the appointment. But the interface also organizes an exchange, leaving space for the doctor's subjectivity. It first communicates on the margin of error of the AI, on the quality of the image and especially explains its choice of color. It explains on which elements it is based, with classes and subclasses of reasoning. The professional can also comment on the result, qualify it, express his doubts, formulate a counter-diagnosis. All of this information is sent to the expert dermatologist. For him, the application facilitates the sorting, identifies emergencies and patients who really need a consultation.

#### *Training the professional*

For the general practitioner, it is also a way to train: "They have a short dermatology course, but they are the first to come into contact with possible cancerous lesions. The tool is an aid to their training".

## Build a situated explainability of an AI system

Transparency of AI systems, especially those that use learning methods, is essential for worker trust in AI. This results in two important concepts [Gilpin et al., 2018].

- **Interpretability.** It answers the question "how does an algorithm make a decision?" (what calculations? what internal data?). This consists of providing information representing both the reasoning of the Machine Learning algorithm and the internal representation of the data in a format interpretable by an ML or data expert. The result provided is strongly linked to the data used and requires a knowledge of the data but also of the model.

**Designer interview:** "We helped develop what we call a medication process, go into the data with statistical models and approaches that identifies what areas of the data are causing the group differences. So what causes males and females to score differently and we try to identify those data elements to minimize the group difference that we found and keep running models and approaches to minimize biases. So, let's say we improve the AI technology we use to be less biased and if we're still biased, we medicate back to that assessment".



- **Explainability.** It answers the question "Why?" (What links with the problem posed? What relations with the application elements? ...) This consists in providing information in a complete semantic format that is sufficient in itself and accessible to any user, expert or layman, and whatever his expertise in Machine Learning.

Interpretability is the first step to be carried out in order to achieve explainability. An explainable model is therefore interpretable but the opposite is not automatically true. For non-expert users, the explainability of AI systems is a major axis of empowerment. Explainability allows them to carry out their activity while feeling responsible.

- The functioning of AI systems is neither intuitive nor comparable to other types of systems.

**End-user interview:** *"what was different compared to other tools is that we know more or less how it works, whereas here, there was a real lack of clarity about the results, how the tool obtained a result".*

- The operation of AI systems can therefore be confusing for workers who may try to understand the system on their own, wasting time and risking misinterpretation.

**Manager interview:** *"This was only seen as a time saving on low value-added tasks. But we had not identified that new complex tasks would emerge. The most difficult thing is to apprehend the system's proposal, especially when it does not spontaneously match the one the agent would have made. Sometimes, the way the system apprehends the different elements of an email is confusing and we have to try to retrace its reasoning. It's not easy, it's an increase in competence that we hadn't anticipated. Our solution doesn't just simplify things, it also adds complexity."*

- In the absence of explanations, workers cannot judge the system's proposal. They are then forced to arbitrate between the system's judgment and their own judgment.

**End-user interview:** *"Sometimes the AI system's judgment is so different and so incomprehensible that it just seems like a big joke. We laugh a lot but then we stop using it. If there is a general definition of explainability, each profession or activity imposes a contingent approach to explainability. This "job explainability" must be defined in collaboration with the workers. Several use cases develop explainability by communicating the quality of the data used by the AI system in order to facilitate its understanding by the user".*

**Manager interview:** *"Data quality is a big problem in the industry. A big, big, big problem. We have the problem, because, I don't know in another countries, but mine is not ready to receive artificial intelligence, so the data is horrible. And so, data quality is very important for us, and we use the beginning of the project to do a data quality analysis. We report the data quality analysis to our clients. It is good to do it because our final product is not only the prediction, but also the data quality analysis, both quantitative and qualitative and we report this to clients, so it is very important".*

**Designer interview:** *"We don't have enough quality data to provide a score, then we don't score it. It's outside the threshold of our capabilities, it could be a technology problem, or too short of an answer. It doesn't get scored and this helps our quality control".*

## Develop a general AI training independent of a particular application

The realization of an AI system is often an opportunity for workers to discover artificial intelligence. Thus, the demystification of artificial intelligence and the business experimentation are situated in the same temporality. Demystification and experimentation are done in a similar framework with the same actors. The solution provider is simultaneously the AI trainer. This position is differently apprehended by the providers.

- Some providers reduce the training to the use of their AI system in order to make the worker operational as quickly as possible.

**Designer interview:** *"We did a training session with the other data scientist and the people from the City of Paris who tested the tool to explain how the tool works. We made a small manual to show how to select the pipeline, how to see the results, how to visualize".*

- A short training period is positively associated with the quality of the user experience, i.e. the usability of the system. Sometimes training is considered optional because the system is so intuitive.

**Designer interview:** *"I think looking at an application does not require training. I mean, before doing the activity they would have to look at the application. I think maybe the first workload is to learn and get used to using it".*

- Other providers only talk about tools and deliberately avoid talking about AI to avoid backlash.

**Designer interview:** *"The debates at the political and academic level are very theoretical, far removed from the issues on the ground, which complicates things more than they should. There is a real gap between production and the rest of the system that are speculating, understanding AI through big debates. For production, AI does not exist. They do not use it, they talk about tools, it is a tool that helps. I don't talk about AI much. The closer you get to production, the less we talk about AI".*

- On the other hand, other "supplier-trainers" consider that they have a broader mission of demystification or evangelization of workers. They therefore provide AI discovery training upstream of the development phase of the business-dedicated AI system.

**Designer interview:** *"I think that we can't expect our clients to have knowledge about AI, machine learning, deep learning, this has not happened. So, we've tried to give them this knowledge, because they have to know about the process. And what we always have to do is explain the dashboards and all the metrics that we use in the dashboards, for example, teach them about the predictions and what they mean. They don't need to have previous knowledge about our technology, we can teach them. We always do trainings about the dashboard and if they want to know more about AI, which I really enjoy, we can organize workshops. It's not a problem, it's good for us".*

In terms of training and education, empowering workers implies going beyond the simple use of an AI system dedicated to a profession. AI and its applications are both fascinating and worrying, and a lot of more or less accurate and precise information is circulating. In any case, they raise many questions, fears and uncertainties, especially concerning its impact on work:

- Can organizations neglect these questions when they engage in an AI project?
- Should these questions be treated in the same time frame as a trade application project?
- Should these issues be addressed by providers?



General training on AI and its business impacts that is doubly independent of a particular project and the vendors seems to be a more satisfactory approach to promoting worker empowerment. It would also make workers more effective co-designers of AI systems because they will already have a better understanding of AI and thus anticipate key dimensions of implementing a system in their business.

#### **Case-study 4: Implementing independent AI training**

In 2017, this major industrial group decided to invest heavily in the development of AI systems and to prepare for work transformations: "the evolution of the company's projects, the technical skills of its employees, their relationships at work and in the organization, but also the employment and training frameworks, will have to be anticipated in light of the upcoming changes." Each employee of the company belongs to a trade academy that brings together professionals working in the same activity. The objective is to create a professional community, a space for exchange, information and training dedicated to a particular profession. To prepare for the implementation of AI systems in its businesses, the Internal Services Academy, which brings together support functions, i.e. 12,000 employees, has initiated a process of "cultural" support for the implementation of AI. The aim of the program was not to evangelize but to stimulate debate and create a shared culture around academic knowledge and ethical issues," explains the Academy's manager. We wanted to address the individual as much as the employee, by offering him or her an opportunity to learn, to strengthen his or her knowledge on an issue that is much talked about. These employees will probably work with AI systems later on, but we don't know exactly when. We're not in a specific time frame that puts pressure on us. It's comfortable. We don't have to convince or convert, just acculturate, develop a critical mind. The system is designed by engineers and academics specializing in work and the ethics of technology. Numerous times of "ethical deliberation" are thus organized in order to "re-engage knowledge in a responsible, respectful and vigilant debate on the new opportunities brought by AI in order to generate personal and collective questioning". This system takes the form of workshops in which the participants register freely without going through their hierarchical channel (which must however validate the employee's absence). The workshops last half a day and are divided into four parts: understanding, debating, applying, debating. The understanding is a time of presentation to understand AI from a theoretical and applicative point of view. The first debate, general, raises the major ethical issues of AI and its impacts: the impact of uses, the impact on human dignity, on workers. The applications take the form of exercises during which the participants carry out an activity assisted by an AI system, all integrated into a scenario: a personality analysis system to prepare an appointment, the creation of a chatbot to welcome trainees, a data mining system to analyze customer satisfaction. The second discussion is about the analysis of these experiences. The goal is to express principles of automation of an activity: what is automatable? What is desirable to automate or should remain a human task? How do we delimit trust in an AI system? Under what circumstances is it most reliable? Knowledge, debates and experiments produce an experience prior to any job application project. 1000 employees attended these workshops. The Academy then set up a special manager session. These sessions were based on the same principles, but with the addition of fictitious managerial situations using the design fiction method. They were played by the managers and the managerial postures were decided collectively: "For example, how to manage the situation of a worker who has made a mistake in correcting the recommendation of an AI system that tends to produce too many false positives? What would be the consequences of a sanction on his ability to take initiative? Another example, how to recognize the merit between employees who are assisted by powerful AI systems and those who are not. By making them play out the situations, the managers were able not only to discover the issues, but also to feel them, which is much more powerful as a learning experience. Finally, the Academy created a Small Private Online Course (SPOC) to broaden the audience for this acculturation. This SPOC associated the same researchers and AI system development teams who, at the end of the course, presented their method, their technologies and the first use cases.

## Accompany use-cases with an independent ethics committee

The ethical issues of AI are a major determinant of its acceptability and trust in AI systems. However, in the 2020 FoW report, we showed that behind an apparent consensus on the centrality of ethical issues, the understanding of ethical issues is not homogeneous. Moreover, this insistence on the ethics of AI irritates some AI actors who feel that it is excessive in relation to what AI really is and unfair compared to other more mature industries that are less subject to ethical injunctions. Indeed, while the ethical issues of AI are widely debated by experts around the world and in the media, the ethical issues of AI systems at work are discussed directly between the vendor, designer, and recipients without any real method or expertise. Many countries and organizations have produced regulations, commitments or ethical charters, but there is a missing link in the chain: the implementation of these principles for a particular AI system, in a specific economic and social context. The normative ethics of AI are increasingly strong, but the ethics applied to professional situations are not sufficiently developed.

Yet, the ethical issues of AI systems require a high level of ethical expertise both to analyse a system and to organise a discussion around this system between stakeholders. These skills are often absent from use cases.

- First of all, these are not skills that economic actors are used to soliciting for innovation.
- Secondly, PoC budgets are often limited, which reduces the objectives to the achievement of a performance.
- Finally, AI system providers are often start-ups that have limited their recruitment to technical and commercial functions.

Currently, our survey shows that two sectors have particularly invested in ethical issues: the public sector and the recruitment sector. The former is vigilant about personal data management, the latter is very invested in reducing bias.

Public authorities may have a role to play in encouraging the integration of ethics in the development of an AI system and in its deployment in a profession. They could organise and finance the constitution of independent ethical committees, bringing together a variety of skills, which project leaders could call upon for support. These committees would be competent to analyse the ethical issues and to respect the constraints of a project. In addition to their expertise on a particular AI system, these committees would contribute to increase the culture, understanding and skills of economic actors.

### Case study 5: An independent ethics committee for a risky AI system

This AI system analyses the videos captured by the video surveillance cameras: *“Our system watches 100% in real time the images from the urban security centre. Real-time analysis on about fifteen important points: security (fighting, falling, regrouping), traffic (parking, crossing red lights, forbidden directions, traffic jams, SAMU, firemen), environment (illegal dumping, garbage filling, bulky items), societal functioning (counting bicycles, cars, scooters) to modify the ergonomics of the city. It is a real time dashboard of a city. Our system doesn't recognize people, it doesn't even associate them with human beings. For him, there are only objects. Depending on how these objects combine with each other, the system detects a threat. This threat is notified to the police officer who manages the videos. So, instead of spending time watching videos, he evaluates threats and takes the appropriate decision. Watching videos is not an interesting activity. In addition, it is proven that human attention spans decline rapidly”*. However, this AI application is often compared to the Chinese surveillance system and the start-up faces many suspicions from communities and police. That's why the creator has put in place a comprehensive ethical validation process.

*“1- We have an ethical charter. We don't want to kill people! We are for the citizens, not for the authority.*

*2- We have an ethics committee which meets every two months, with two elected representatives, residents from outside. We always have three guests: a militia member, a philosopher, a ministerial prefect. We choose our subjects together after discussion. They can block a new project.*

*3- We are GDPR by design. A human being is an object, we cannot identify it. It's like a bicycle or a car. We don't do facial recognition and biometric tracking.*

*4- Monthly meeting with the CNIL (in France, commission created by the French Data Protection Act of 6 January 1978. It is responsible for ensuring the protection of personal data contained in computer or paper files and processing, both public and private) to discuss all our new cases.*

*5- We have a lawyer in the company.*

*6- We make an Impact Plan for the CNIL.*

*All our projects are included in it. Because fear hinders technology. You gain by being authentic, especially in relation to the competition.”*

## Diversify design teams to reduce bias in data

In 2021, the catalog of FoW has been enriched with several applications dedicated to recruitment and vision in cities. These uses of AI are generally considered very sensitive because they can generate discrimination due to the biases contained in the system:

- Explicit biases that result from the specific objectives, intentions, human arbitrations, unconscious social representations of a time, the developer or the experts.
- Implicit biases contained in the data mobilized in the learning phases,

In response, statistical methods are used to identify and correct these biases. Another good practice, less mathematical but more social, consists in aiming for a good level of diversity within the product development team.

**Designer interview:** *“Automatic recognition systems’ potential risk is related to the subject of biases. Humans are the ones who train it and if that data is biased then the system can create problems. There have been systems that have had to be removed, for example in control of access to cars. For example, if an AI system hasn't been trained to recognize redheads or black people. If the data is not well prepared, there are biases and with artificial vision issues, biases are created like Tesla, which is automatic. You train Tesla with red signals and you go to a country with blue signals, there are going to be accidents. It can be avoided with different people within teams, the team must be diverse to detect different kind of biases”.*



## Recommendations from use-cases

### The success of a use case:

- **Establishing methodological principles of a PoC beyond the performance of the AI system:** integration issues in organizations are not sufficiently considered, which limits the conversion of these PoC into products.
  - **Reorganize:** AI systems imply rethinking the organization of the activity. *e.g.: a nuclear power plant incident analysis application adds an unmanageable cognitive load for the technician.*
  - **Socialize:** AI systems destabilize the value system associated with the activity. *e.g.: the integration of an AI in the detection of bacteria in a dairy product makes the distribution of responsibilities in case of contamination too complex.*
  - **Practice:** AI systems transform, generate or destroy professional practices. *e.g.: the time saved by the use of a voice assistant reduces the possibilities of exchanges between colleagues and increases the time spent in a noisy room.*
- **Encourage and improve the integration of academic research:** These collaborations are essential but project management can be lacking. Researchers need to integrate business constraints. *e.g.: a research team had to develop a linguistic solution to help a customer-supplier relationship chatbot to capture the ambiguities of the word "thank you" in a conversation. The collaboration was stopped because the researchers could not commit to a deadline and proposed a "black box".*

### Empowering the worker:

- **Define the right trade-offs between usability and user involvement:** HMIs must consider a good level of "cognitive tension" for the user. Ease of use is appreciated but can breed passivity and docility. *e.g.: instead of giving a result, an application organizes a skin disease diagnosis interaction between the system and the doctor.*
- **Build a situated explainability of an AI system:** The explainability of a system must be related to real work situations for AI systems to be accepted and understood. *e.g.: an application for managing citizens' complaints does not communicate to technicians the elements that allow them to make their own judgement. "What was different compared to other tools is that we know more or less how it works, whereas here, there was a real lack of clarity about the results, how the tool obtained a result".*
- **Develop a general AI training independent of a particular application:** Training in AI by designers of the use-case does not promote worker independence and makes them a less effective co-designer of the AI system. *e.g.: a company engaged in upstream AI training and experimentation workshops by academics prior to any deployment.*

### Fair AI:

- **Accompany use-cases with an independent ethics committee:** Public authorities have a role to play as this is not always a market requirement or in the budgets allocated to AI system development. *e.g.: A video surveillance image analysis company has all its new projects assessed by an independent ethics committee.*
- **Diversify design teams to reduce bias in data:** Diverse teams make it easier to incorporate a variety of perspectives and complements pure statistical approaches. *e.g.: a recruitment system developer ensures the social, cultural and gender diversity of its design teams to reduce the biases contained in the algorithms.*

