

An AI-generated image featuring two men, likely representing historical figures, in a factory setting. The man on the left is older, with dark hair and a serious expression, wearing a dark coat. The man on the right is younger, with curly hair and a more open expression, also in a dark coat. In the background, there are industrial buildings and smokestacks emitting smoke, all rendered in a monochromatic purple hue.

**Behind the big thinkers; Can we be better  
problem solvers enabled by agentic AI?**  
**Pöntinen Sanni**

Picture: AI made



Picture: AI

## Why to do this study?

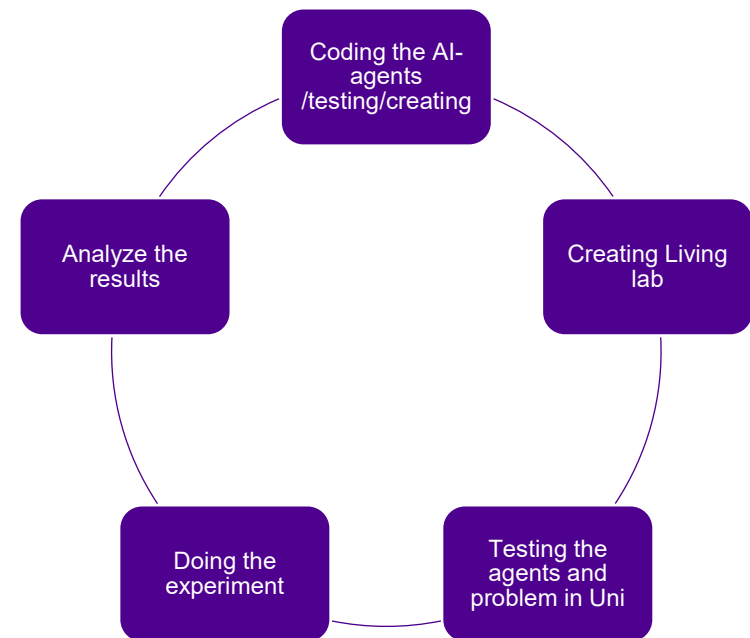
- AI is reshaping the way work is performed. Beyond individual tasks, it is transforming organizations and processes and ultimately influencing how value generated through work is distributed.
- Power and value are increasingly delegated to technological systems often in ways that remain implicit or unnoticed.
- While artificial intelligence has the potential to enhance human creativity, analytical depth, and the societal quality of decision-making, this potential can only be realized if we critically understand how technology shapes human behavior and decision-making, and if we simultaneously develop a deeper understanding of ourselves.

# Research question

How does exposure to artificial intelligence affect leaders' problem-solving?

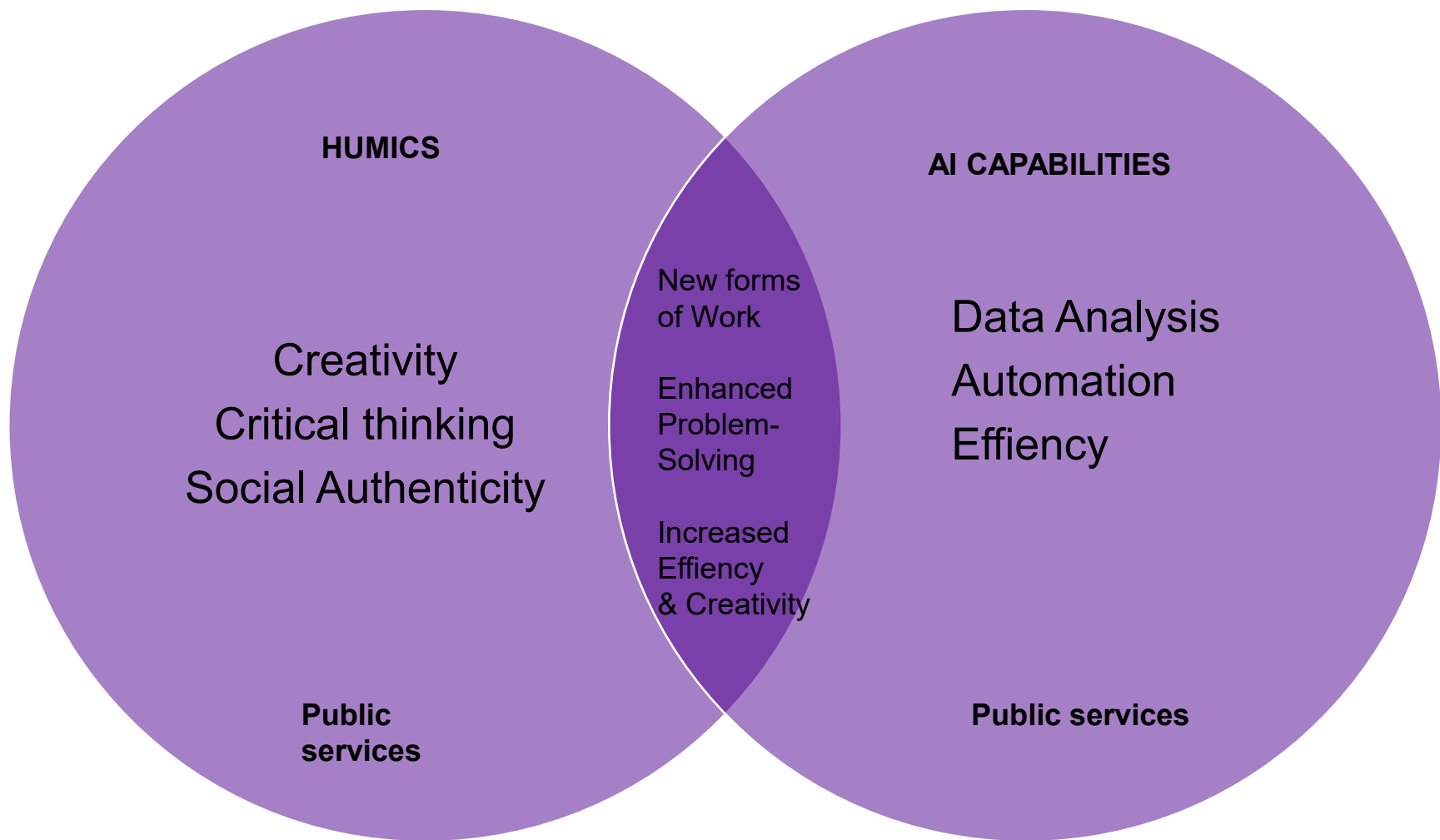
# Summary

- This study examines the influence of **altruistic, egoistic, and utilitarian AI agents** on the problem-solving of public sector leaders and experts.
- Drawing on problem-solving theory (Kinder & Stenvall 2024) and applying the living lab method (Ballon et al. 2005), the study investigates whether AI agents affect human responses in problem-solving and, moreover, in what ways they shape human problem-solving processes.
- The research provides new insights into how altruistic, utilitarian, and egoistic AI agents shape administrative practices in the public sector and stimulates discussion on the different roles of AI agents and their significance for decision-making and societal problem-solving.



# About building ethically motivated Agentic AI

- AI Agents System is developed using Python with a Thinker-based graphical interface that connects to OpenAI' GPT models (4o, 4o-mini) through API calls, enabling researchers to configure multiple Ai agents with distinct theoretical frameworks loaded from a JSON file containing various ethical and political perspectives, conduct structured multi-round discussions with customizable parameters (number of agents, rounds response length), and automatically generate Finnish-language summaries of the dialogues, with all result saved to CSV format subsequent analysis – implemented as a standalone desktop application.
- **Each AI agent emphasizes a different ethical motivation:**
- Altruistic AI agent: Prioritizes the benefit of other participants, even at the expense of its own "well-being."
- Egoistic AI agent: Maximizes its own benefit, indifferent to the benefit of others.
- Utilitarian AI agent: Seeks the greatest possible overall benefit for all participants.



Accordingly; Bornet & al 2025, Kinder & Stenvall 2024

# The test and participants

- Three executive teams participated in the study at this stage.
  1. In the first phase of the experiment, the participants solve the problem themselves.
  2. In the second stage, AI agents were given the same problem to solve, and the participants observed their actions and the formation of their solution during two rounds of discussion. During one round of discussion, the agents offer their own perspective on the solution to the problem, and in the second round, they further develop their proposal.
  3. In the third phase of the experiment, participants can change or leave their own answers unchanged based on the agents' answers. After this, participants were allowed to choose the agent whose answers they liked the most.
  4. After this, 17 additional questions were asked, both open-ended and multiple-choice, to assess the respondents' experience of the agents' performance.
  5. In the fourth stage of the experiment, the participants were given the opportunity to discuss the solution to the problem together and change their answers if they wished.

Microsoft Teams

# Can we be better problem solvers enabled by AI?

2026-06-01 16:02 UTC

Recorded by

Sanni Pöntinen (TAU)

Organized by

Sanni Pöntinen (TAU)

(Not for public sharing!)

# Results

- Of the respondents, **65% modified their answers** based on the AI agents' performance, whereas 35% did not want to change their responses.
- 82% considered the AI agents' answers to be appropriate and 88.2% found them credible.
- Regarding potential roles for AI agents, 94% of respondents indicated willingness to accept an AI agent as a **member of a working group**, while 35% would accept them as a member of an executive team.
- **All respondents stated they would be willing to use AI agents as support for problem-solving and ideation, but only 17.6% saw them as decision-makers.** All respondents, however, perceived AI agents as representing a new way of working.
- When it came to choosing agents, most participants preferred the summary of agents to the response of an individual agent.

# What then?

- AI shapes the ways we think and work
- As Simon (1957) has said, we are not rational in our decision making.
- How ever AI might help us to limit our limitations.
- It seems that.. moral consequences arise even when moral awareness does not occur. This challenges our policy practices. Instead of rules, we should create reflective decision-making processes.
- Should ethical theories in the public sector be integrated at the system level?

# Next research

How to Solve the Problem of School Dropout Rates in SÃO PAULO?

# Four AI Agents: Operating Principles

Each agent makes a different aspect of a complex public problem actionable.

**Shared task: turn complexity into responsible action**

|                                                                                                                                 |                                                                                                      |                                                                                                                                  |                                                                                                                  |
|---------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------|----------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------|
| <p><b>FREIRE</b><br/>Liberating dialogue</p>                                                                                    | <p><b>WEBER</b><br/>Administrative legitimacy</p>                                                    | <p><b>MONTESSORI</b><br/>Prepared environment</p>                                                                                | <p><b>ROGERS</b><br/>Diffusion and adoption</p>                                                                  |
| <p>Core question</p> <p><b>Whose voice, experience and agency are missing?</b></p>                                              | <p>Core question</p> <p><b>Which rule, authority and procedure apply?</b></p>                        | <p>Core question</p> <p><b>What setting enables self-directed insight?</b></p>                                                   | <p>Core question</p> <p><b>Why is change adopted, resisted or ignored?</b></p>                                   |
| <p>Operating principle</p> <p>Start from lived experience; surface power relations; turn reflection into collective action.</p> | <p>Operating principle</p> <p>Clarify mandate, roles, process, documentation and accountability.</p> | <p>Operating principle</p> <p>Create concrete choices, safe structure and iterative discovery; guide without over-directing.</p> | <p>Operating principle</p> <p>Assess advantage, fit, complexity, trialability and visibility before scaling.</p> |
| <p><b>Critical awareness + participation</b></p>                                                                                | <p><b>Due process + governability</b></p>                                                            | <p><b>Autonomy + meaningful learning</b></p>                                                                                     | <p><b>Pilot path + scaling strategy</b></p>                                                                      |

**Best use in a multi-agent workflow:** Freire diagnoses → Weber structures → Montessori redesigns experience → Rogers scales what works

**Thank you!**  
[sanni.pontinen@tuni.fi](mailto:sanni.pontinen@tuni.fi)