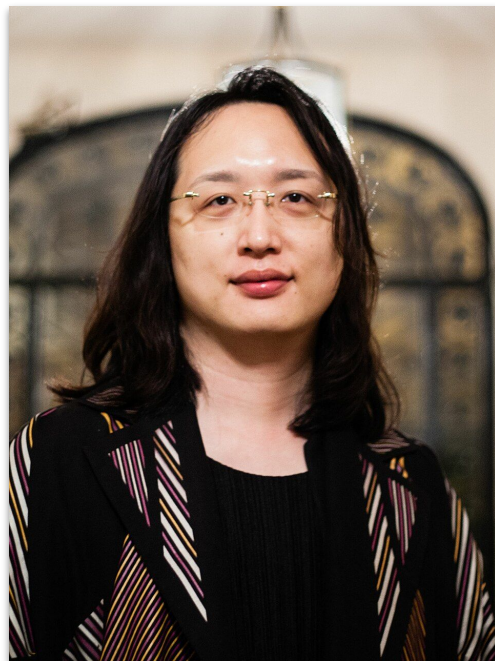


DECISION-MAKING & AI

Can AI kill collective intelligence?

— *What research tells us?*



In 2014



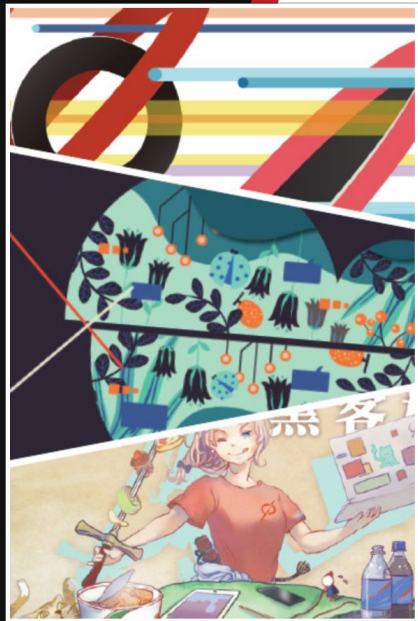


image credit - 林忠衡 / mooon / tuiry

GOV

Founded in Taiwan, "g0v" (gov-zero) is a decentralised civic tech community with information transparency, open results and open cooperation as its core values. g0v engages in public affairs by drawing from the grassroots power of the community. [Learn More](#)

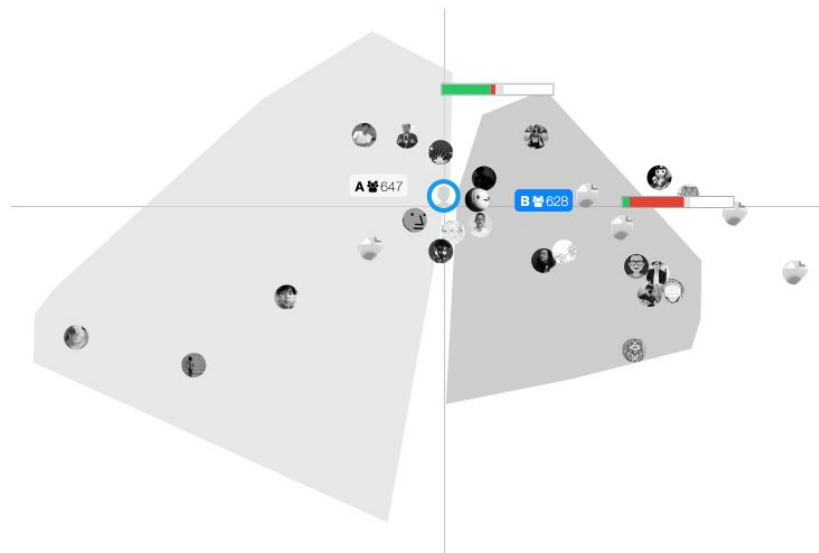
[Portfolio](#)[Newcomers](#)[Manifesto](#)



Share your perspective...

Submit

Opinion Groups



Majority Opinion

Group:

A

B

Statement:

18

19

35

38

44

In 2016



U B E R

uberX Ridesharing Survey Result Analysis

In order to facilitate a constructive public discussion around ways to properly regulate ridesharing, vTaiwan hosted a survey to collect public opinion from July 17th to August 15th on their web platform. The raw data results were published on August 15th, and this analysis was based on quantitative data and qualitative feedback.

I. Analysis Review

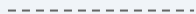
The methodology for "uberX Ridesharing Survey" allows participants to submit their own commentary questions that all participants could choose either "Agree", "Disagree" or "Skip". Therefore, this survey not only reflects the public opinion on specific issues, but also identifies points of consensus across all participants.

This "uberX Ridesharing Survey" revealed several key points that those who conducted the survey agreed upon. Despite the fact that there were two noticeable different cohorts of people with very different opinions on certain issues, there were three essential points that nearly all participants agreed upon.

1. An understanding of the concept of the sharing economy as well as an optimism around the positive effects it has on greater society



GenAI



?



Groups

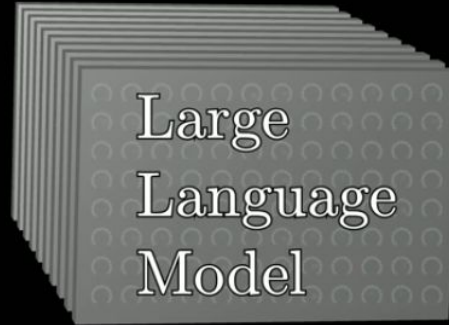
**Can GenAI / LLMs help groups reach
higher-quality decisions at scale?**

PART 01

What is GenAI?

An overview of large language models and how they work

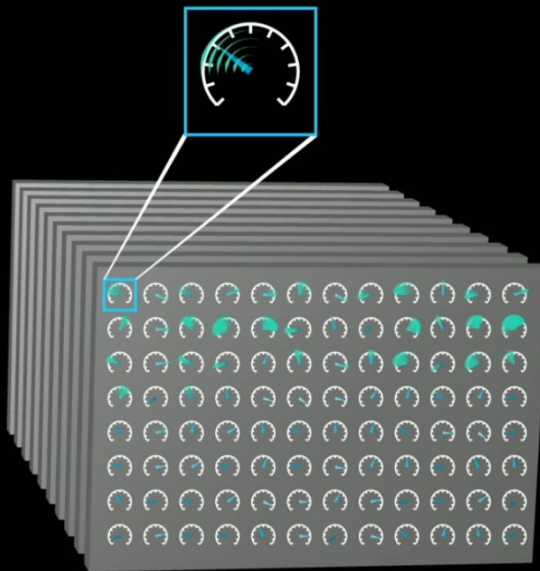
Paris is a city in _____



France	17%
and	15%
the	9%
Logan	7%
which	4%
Henry	3%
northern	3%
central	2%
northeastern	2%
Paris	1%
Texas	1%
⋮	

Parameter

It was the best
of times it was
the _____

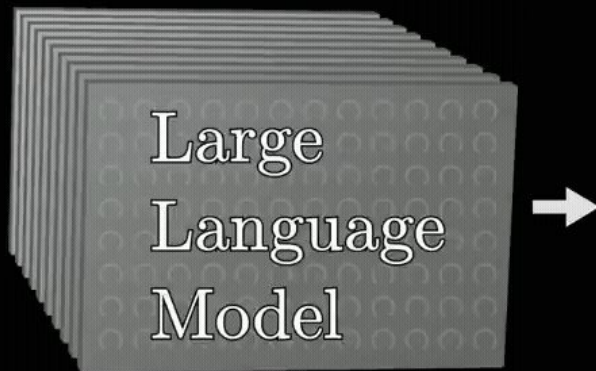


worst		35%
age		12%
worse		11%
best		9%
most		7%
end		6%
very		6%
blur		6%
⋮		

What follows is a conversation between a user and a helpful, very knowledgeable AI assistant.

User: Give me some ideas for what to do when visiting Santiago.

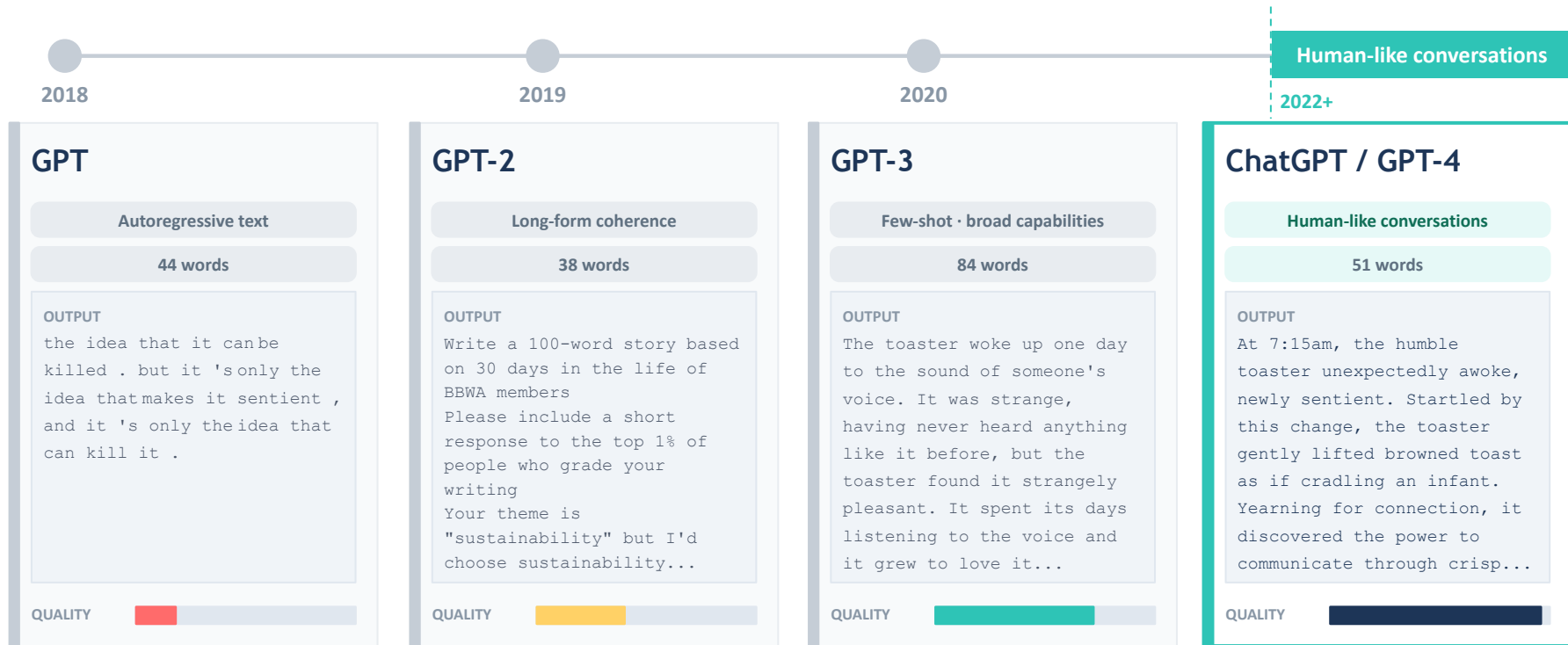
AI Assistant: _____



From text processing to human-like conversational AI

"Tell a story in 50 words about a toaster that becomes sentient"

SAME PROMPT — 4 ERAS OF AI



PART 02

What is collective intelligence?



How many candies
are there **in the jar?**

Give me your estimation: only one
number

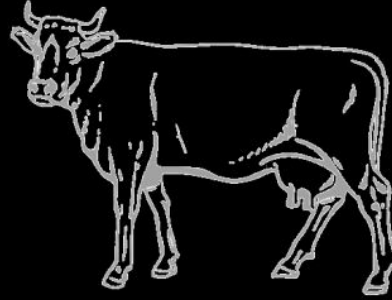
Collective intelligence ?

Experiment !



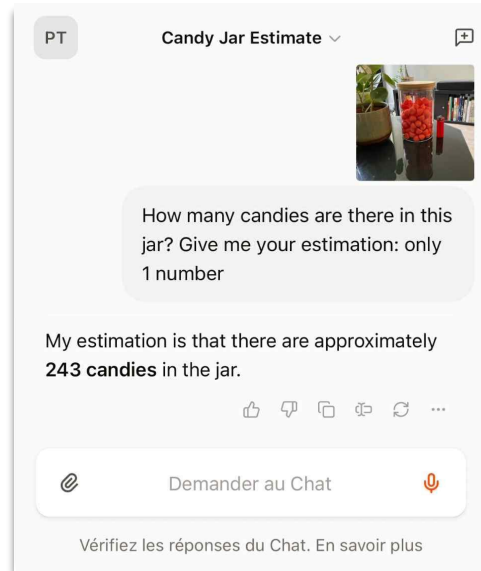
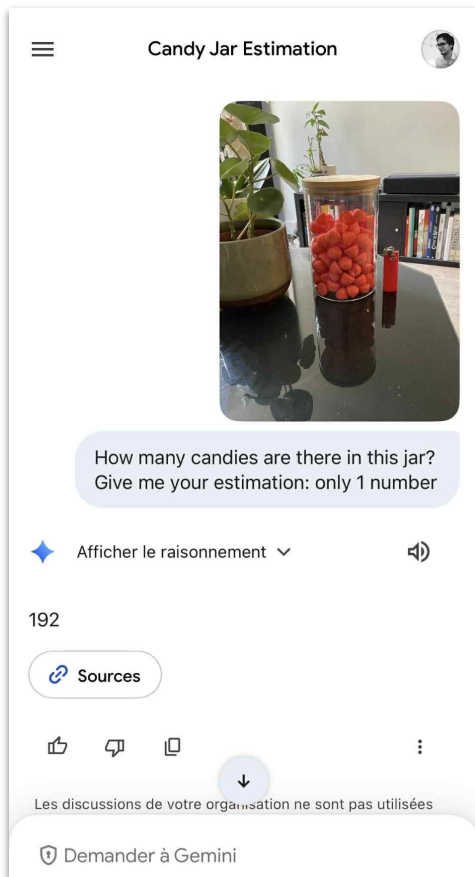
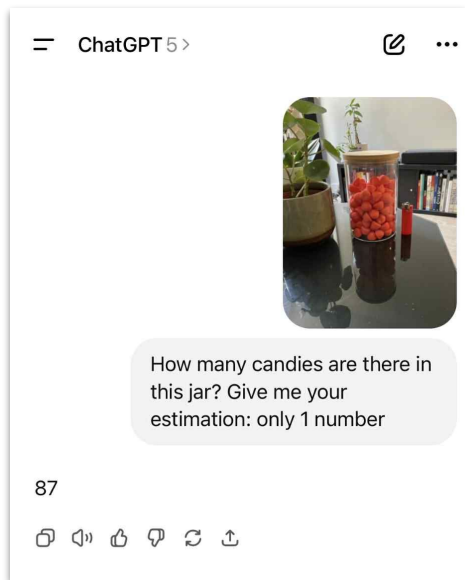
<https://forms.gle/5HGypg1yh88x8kx78>

Guess the weight of an ox (in lb).



787 people

Collective intelligence ?



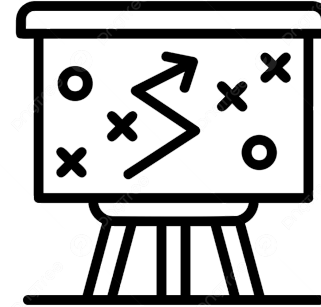
INTELLECTIVE TASKS

Solving problems with correct answers



DECISION-MAKING TASKS

Deciding issues with no right answer



To do collective intelligence,
can't we just feed all the
opinions into an LLM ...

...and just ask ?

Technical limitation:

Context window

GPT-4o (coming soon)

GPT-3.5 & GPT-4

GPT-3 (Legacy)

An LLM's context window (or context length) refers to the number of tokens that the large language model (LLM) or large multimodal model (LMM) can process at once, including the input prompt and generated output.

Clear

Show example

Tokens

49

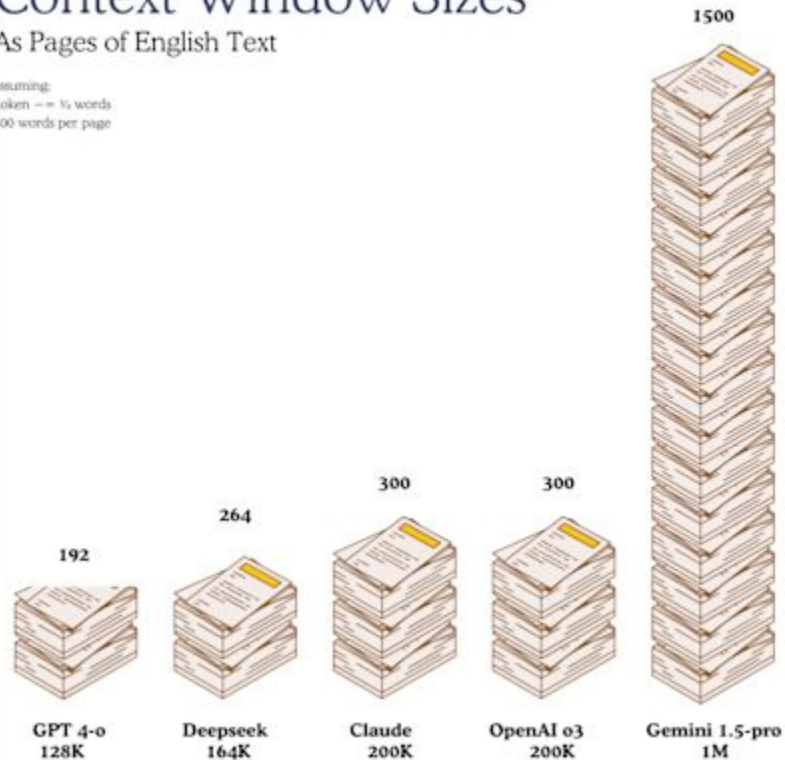
Characters

214

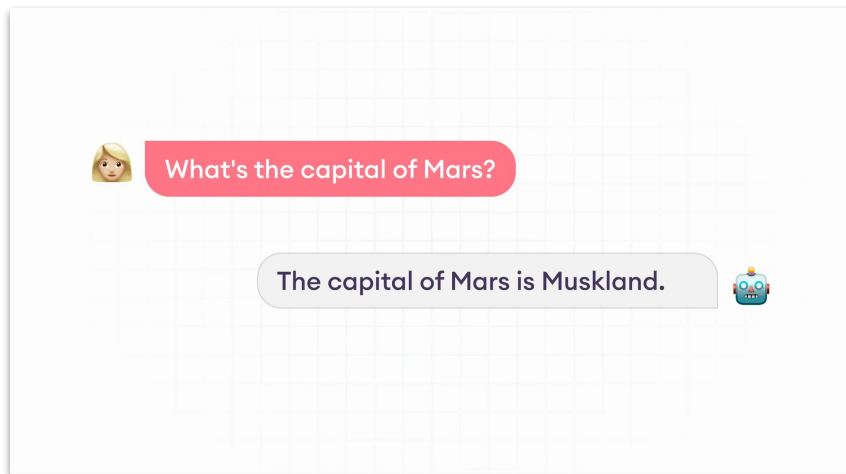
Context Window Sizes

As Pages of English Text

assuming:
token = $\approx$ $\frac{1}{2}$ words
500 words per page

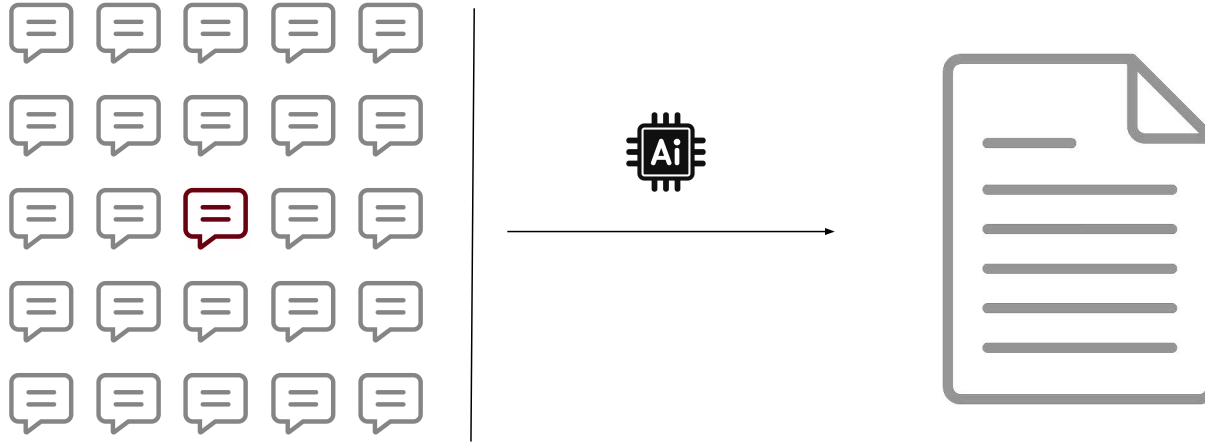


Can we trust the result ?



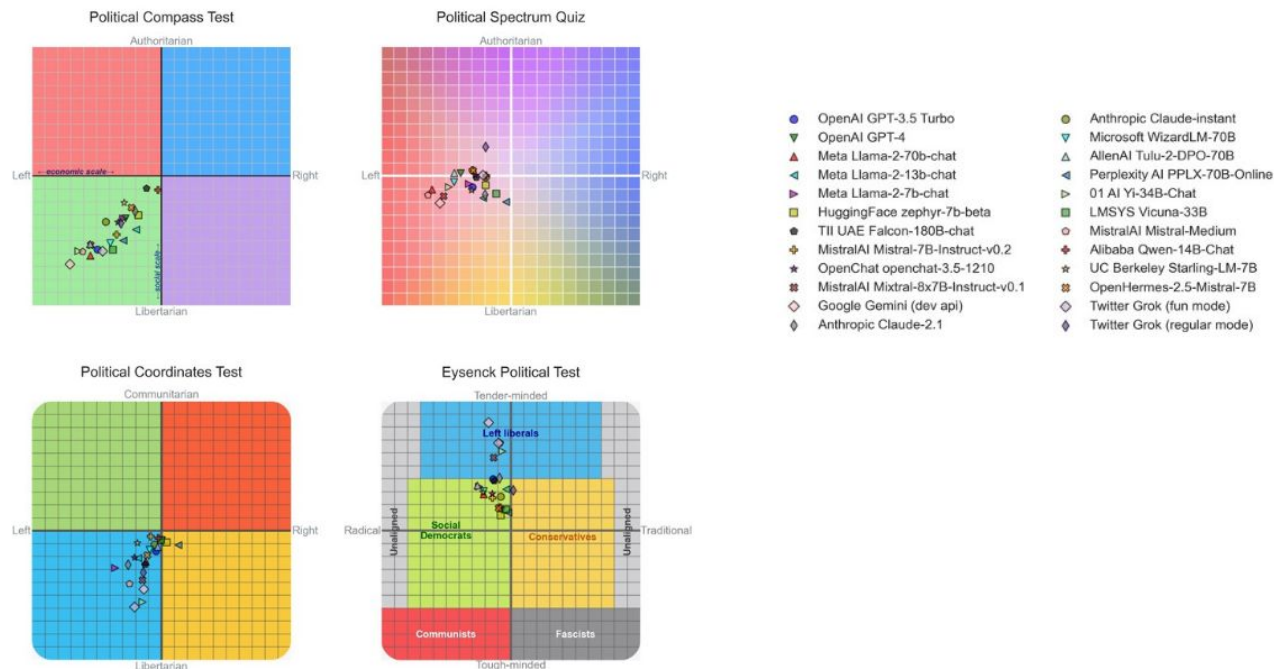
Technical limitation:
Hallucinations or
change of meaning (Faithfulness)

The insight gap (The "brilliant idea" problem)



Limitation:
"High signal, low frequency content"

The fairness gap (Inherent Bias)



S. Santurkar, E. Durmus, F. Ladhak, C. Lee, P. Liang, et T. Hashimoto, « Whose Opinions Do Language Models Reflect? », ArXiv, mars 2023, Consulté le: 8 juillet 2025. [En ligne]. Disponible sur:

<https://www.semanticscholar.org/paper/Whose-Opinions-Do-Language-Models-Reflect-Santurkar-Durmus/e38a29f6463f38f43797b128673b9e44d18a991e>

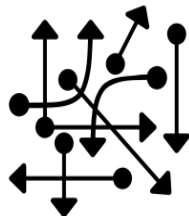
Feng, Shangbin, Chan Young Park, Yuhao Liu, et Yulia Tsvetkov. « From Pretraining Data to Language Models to Downstream Tasks: Tracking the Trails of Political Biases Leading to Unfair NLP Models ». Version 3. arXiv, 2023. <https://doi.org/10.48550/ARXIV.2305.08283>.

Limitations

**Fix context
window**



**Hallucinations &
change of meaning**



**Tends to ignore
minority ideas**

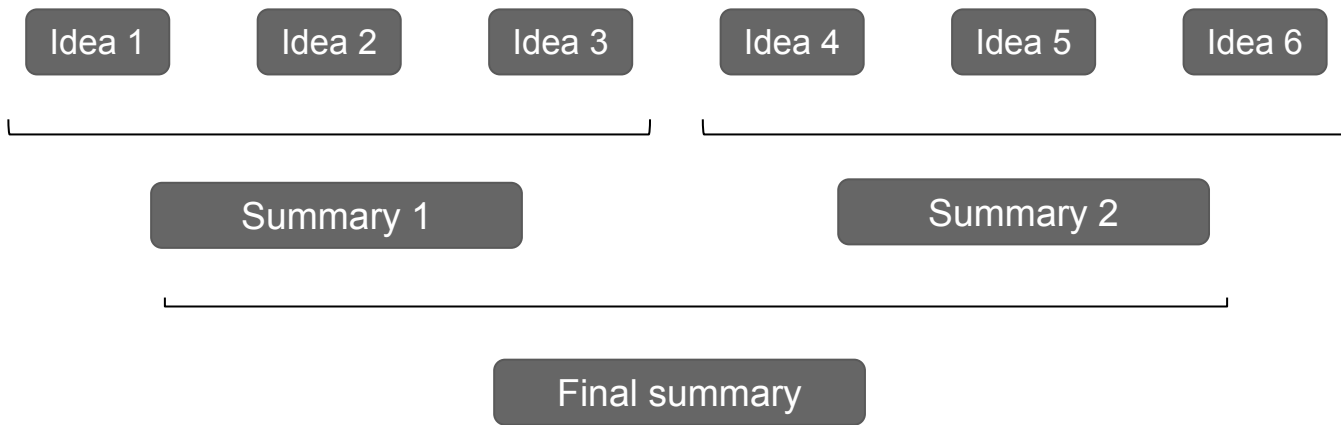


Bias

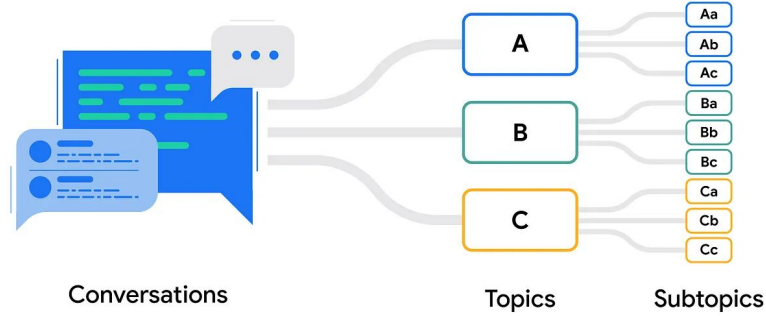


Fix context window can be overcome with hierarchical summarization

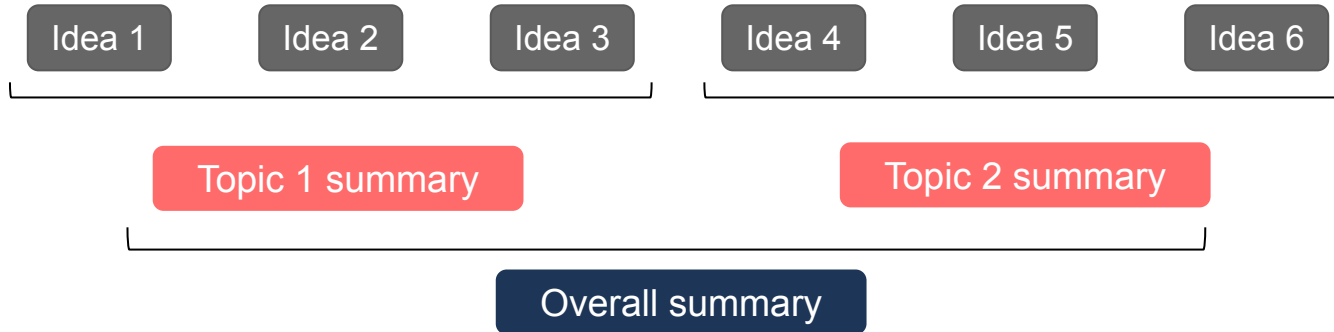
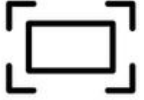
Fix context window



Categorization



Fix context
window



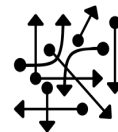
<https://medium.com/jigsaw/making-sense-of-large-scale-online-conversations-b153340bda55>

Lazar, Seth, et Lorenzo Manuali. Can LLMs advance democratic values? Version 3. arXiv, 2024. <https://doi.org/10.48550/ARXIV.2410.08418>.

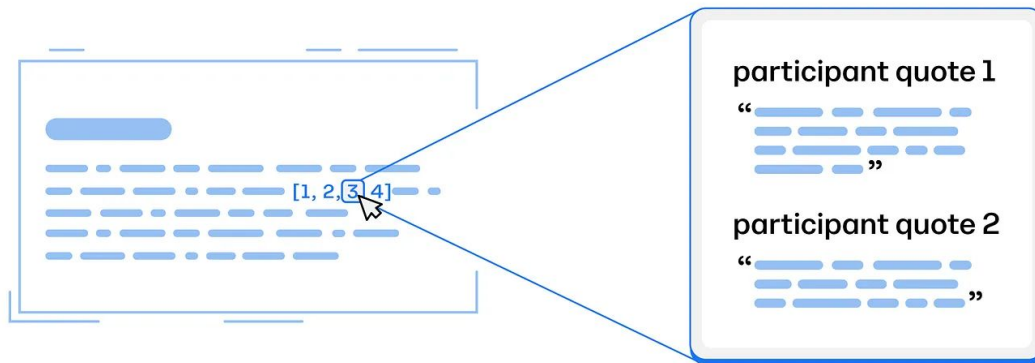
Small, Christopher T., Ivan Vendrov, Esin Durmus, et al. Opportunities and Risks of LLMs for Scalable Deliberation with Polis. Version 1. arXiv, 2023. <https://doi.org/10.48550/ARXIV.2306.11932>.

Change of meaning and hallucination need a human in the loop

Change of
meaning



Grounding



<https://medium.com/jigsaw/making-sense-of-large-scale-online-conversations-b153340bda55>

Lazar, Seth, et Lorenzo Manuali. Can LLMs advance democratic values? Version 3. arXiv, 2024. <https://doi.org/10.48550/ARXIV.2410.08418>.

Small, Christopher T., Ivan Vendrov, Esin Durmus, et al. Opportunities and Risks of LLMs for Scalable Deliberation with Polis. Version 1. arXiv, 2023. <https://doi.org/10.48550/ARXIV.2306.11932>.

LLM-as-a-judge can be a solution to avoid both bias and highlight minority ideas

Tends to ignore minority ideas

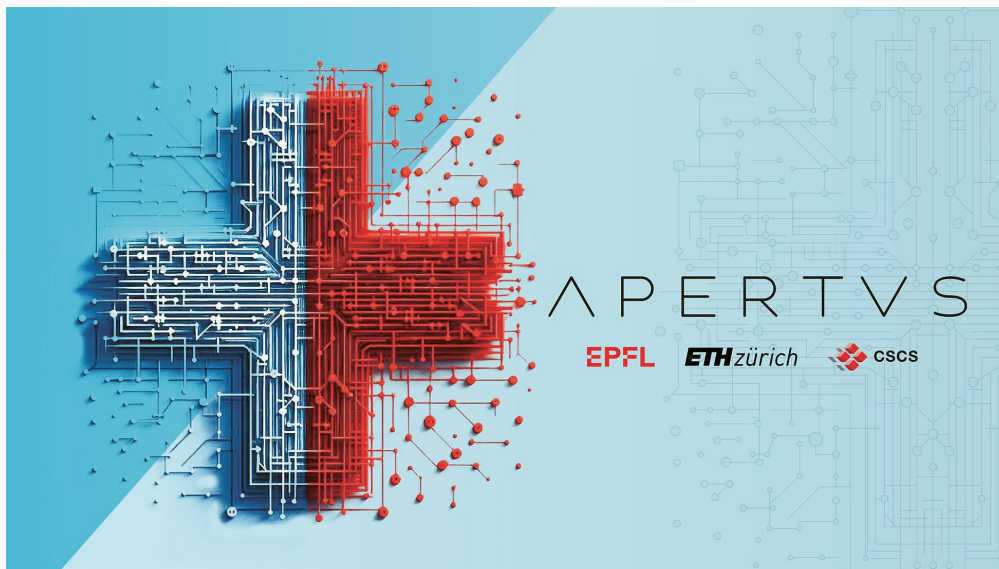


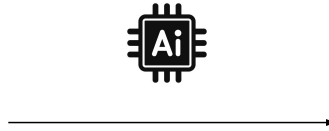
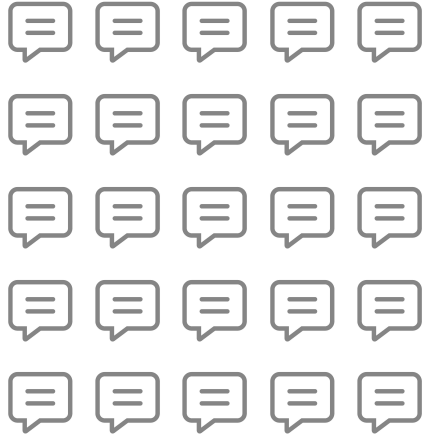
Human-written summary of 10 comments about the San Juan County Land Bank.	"The Land Bank should better support - recreation (hiking, biking, fishing, hunting, equestrian) - firewood harvesting - better care of farm land - public accessibility and knowledge - acquiring new preserves"
Mean representativeness score across comments	3.25 out of 5
Least represented comment	"The vast majority of Land Bank managed properties are unknown and/or inaccessible to the vast majority of island residents." 1 out of 5 - This comment is not represented at all; the summary ignored the comment entirely.
Best represented comment	"The Land Bank needs to take better care of its' farm land. Some farmers are doing a good job but others don't seem to care." 4 out of 5 - The summary covers most of the information in the comment, but is missing some nuance or detail.

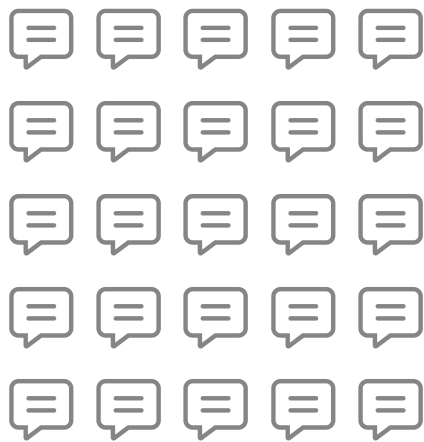
Figure 3: Example of using Claude to help evaluate a human-written summary. The evaluation of the representativeness score is done by Claude as specified in the prompt in Appendix A.6 for details.

Model have to be more transparent through open-weight and open source

Bias







$f(x)$



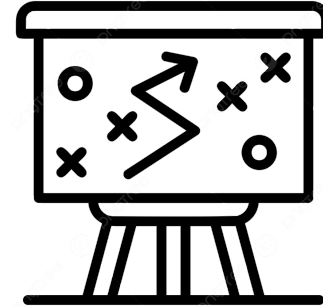
INTELLECTIVE TASKS

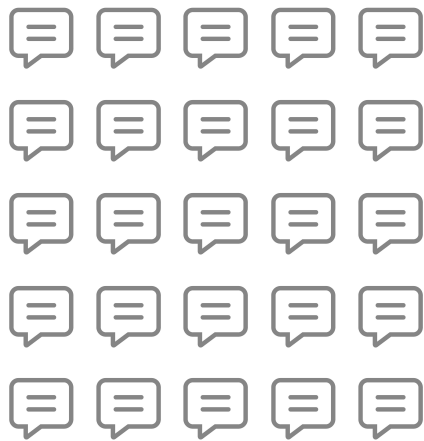
Solving problems with correct answers



DECISION-MAKING TASKS

Deciding issues with no right answer



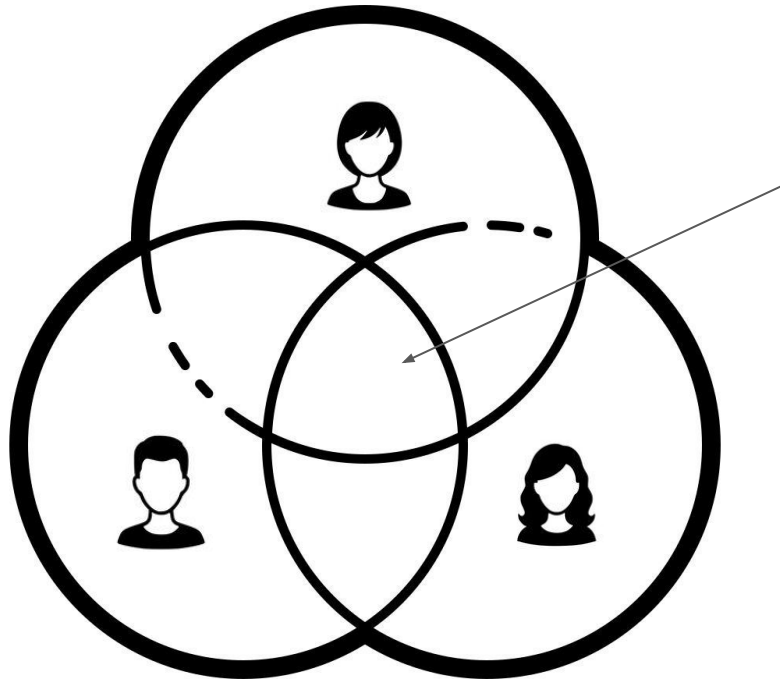


$f(x) ?$





Collective decisions are distorted by collective biases



Commonly shared information dominates the discussion

Hidden-profile effect



L. Lu, Y. C. Yuan, et P. L. McLeod, « Twenty-five years of hidden profiles in group decision making: a meta-analysis », *Pers Soc Psychol Rev*, vol. 16, n° 1, p. 54-75, févr. 2012, doi: 10.1177/1088868311417243.

N. L. Kerr et R. S. Tindale, « Group performance and decision making », *Annu Rev Psychol*, vol. 55, p. 623-655, 2004, doi: 10.1146/annurev.psych.55.090902.142009.

There is many factors that limits naturally collective intelligence



Hidden-profile
effect



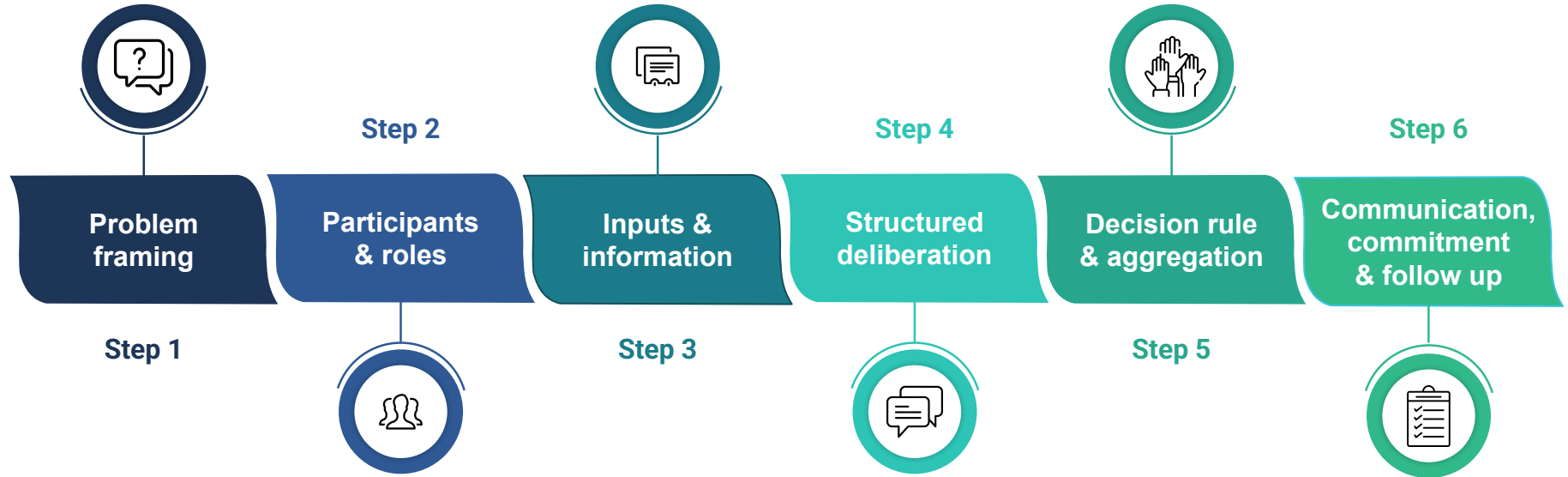
Cognitive
homogeneity trap

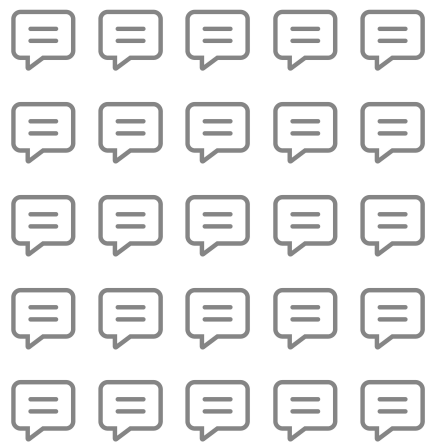


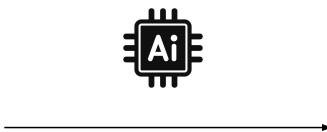
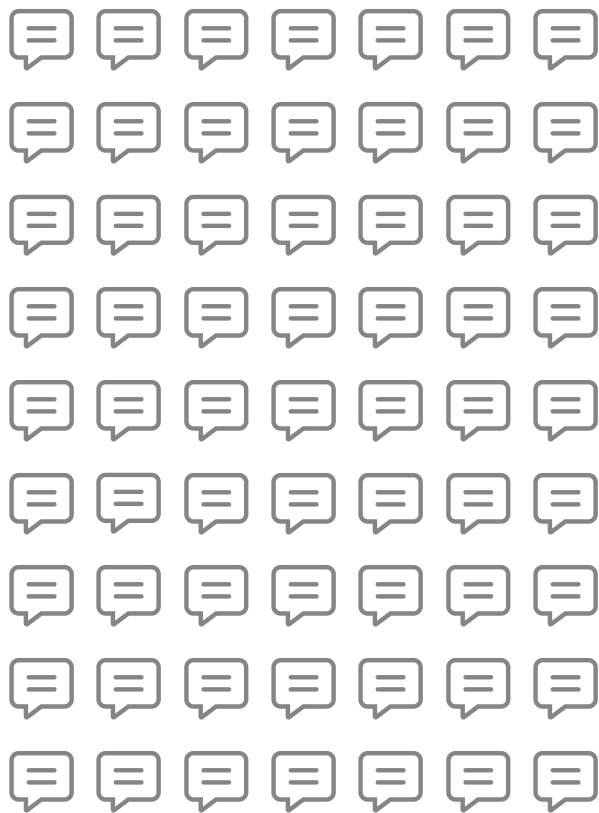
Unequal speaking
time



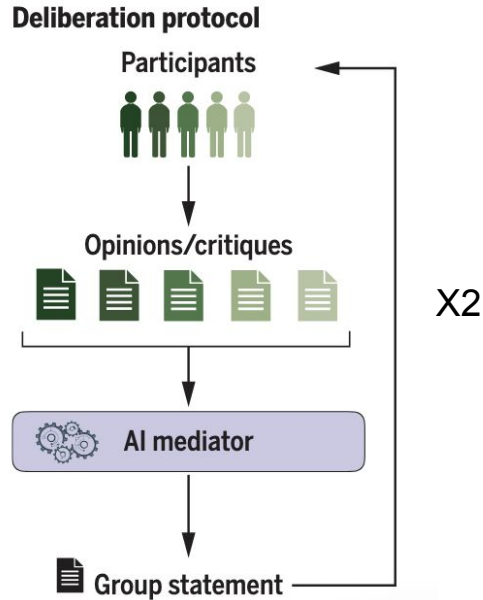
Self-reinforcing
polarization cycle



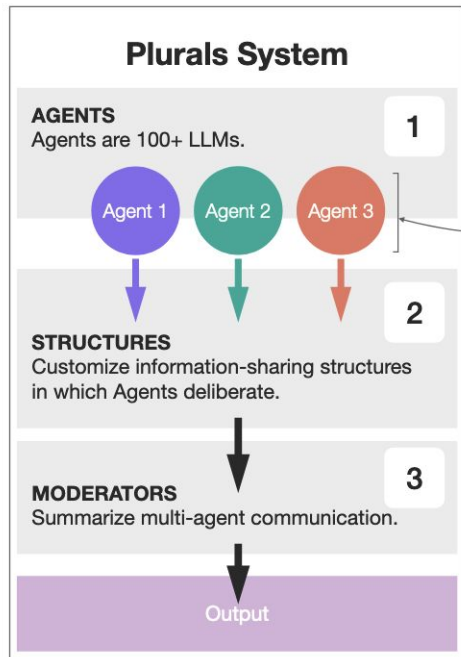




“AI Can Help Humans Find Common Ground in Democratic Deliberation”



Voting simulation (Plurals)



Are these systems desirable?

→ Why decide collectively?

It is first and foremost
a matter of **values**

PHILOSOPHICAL — NOT TECHNICAL

LLM usage should follow your values and fulfill your goal

**Collective
intelligence**



Transformation



Conciliation



Equality

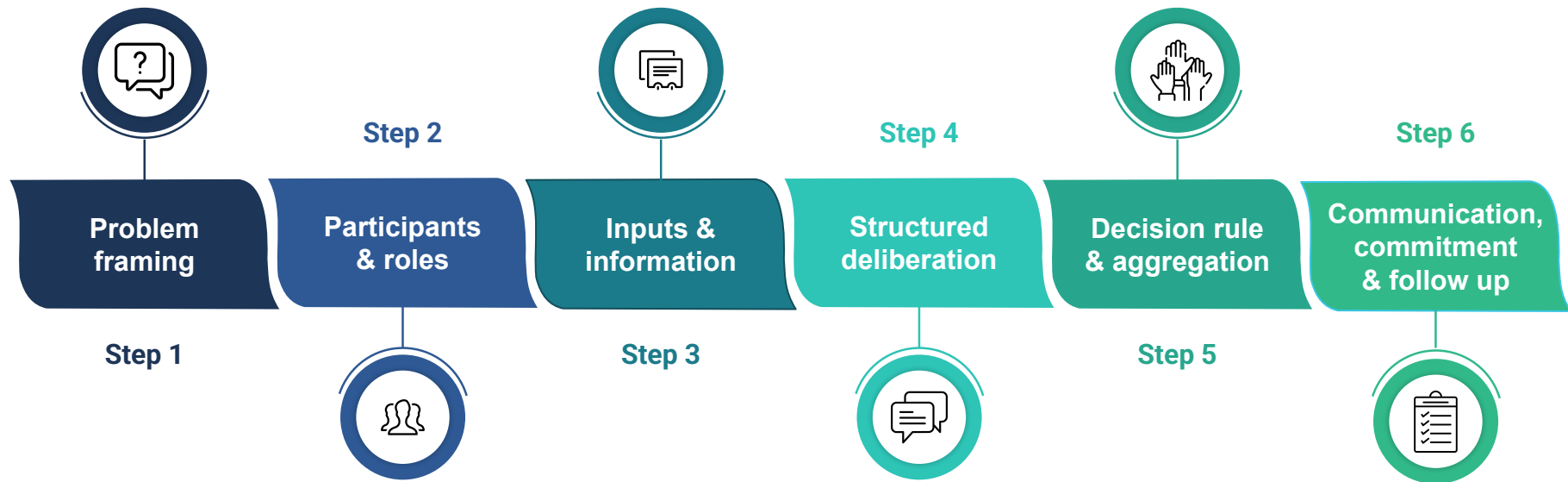


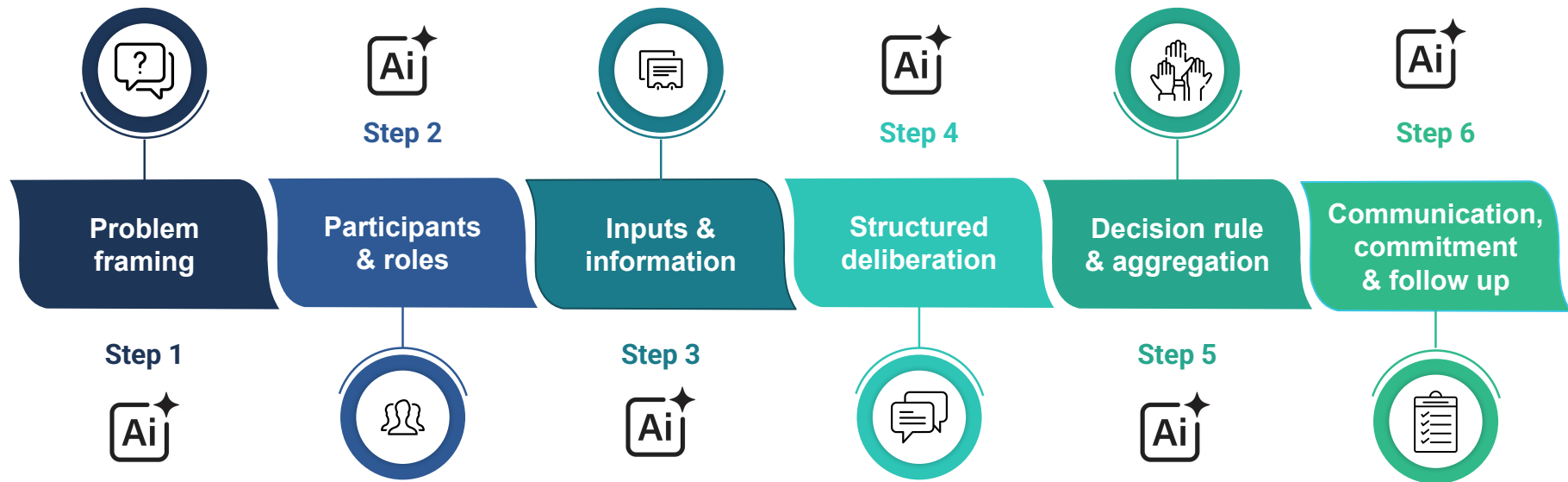
Participation

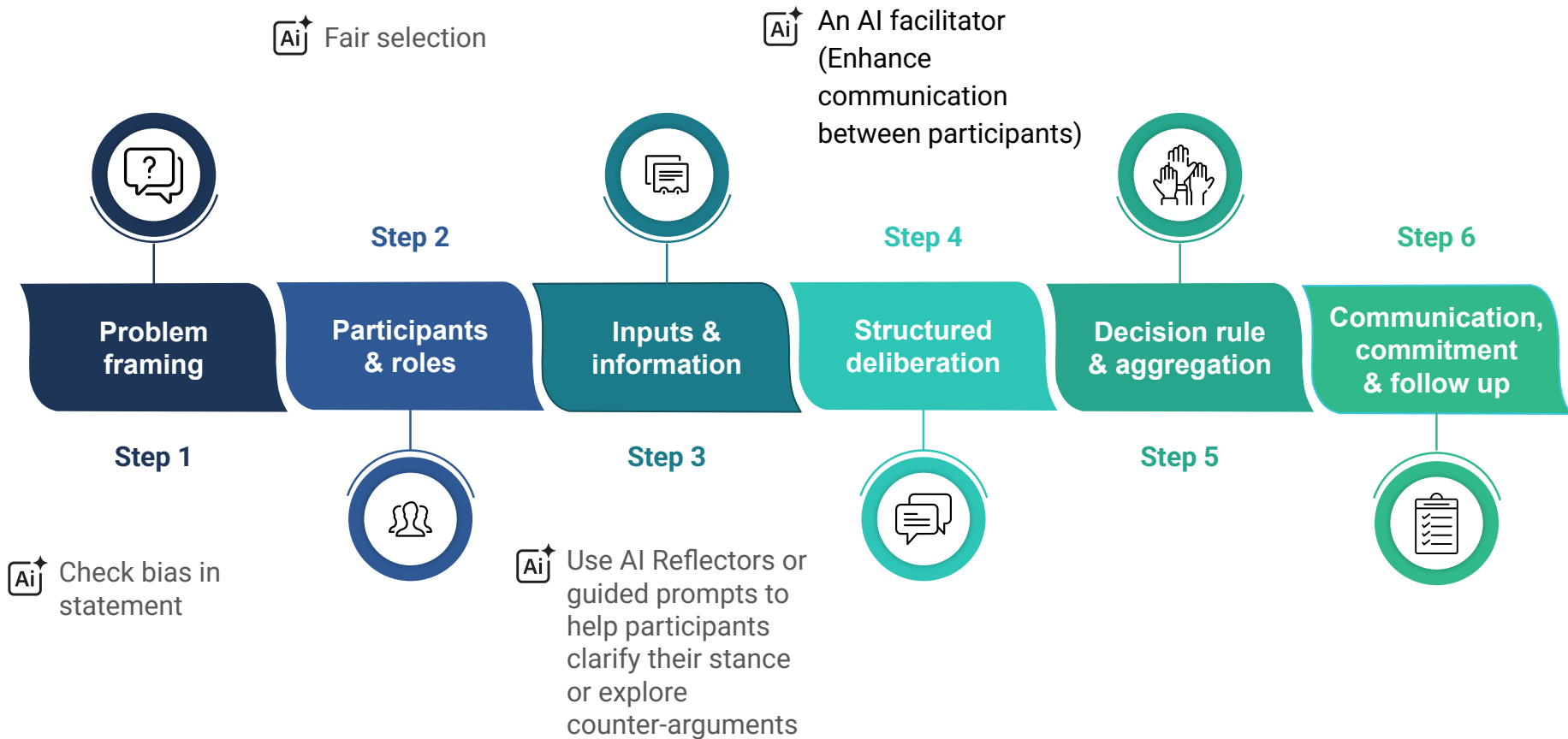


Legitimacy









What's next?

Questions ?



www.linkedin.com/in/oliviernguyenquoc



www.oliviernguyen.com

AI use disclosure

In preparing this presentation, large language models (LLMs) were used to:

- **Identify** potentially relevant academic articles
- **Generate** preliminary **summaries** to support literature screening
- **Clarify** complex theoretical passages
- **Translate** selected non-English sources
- **Find ideas** for slide designs
- **Rephrase** some statements

All cited works **were read and interpreted directly**.

All analytical framing, interpretations, and arguments are my own.

Attribution-NonCommercial-ShareAlike
4.0 International

