

The Decision Maker's Dilemma

or how I stopped struggling with possible choices

Changkun Ou

@changkun

IDC 2022 Spring

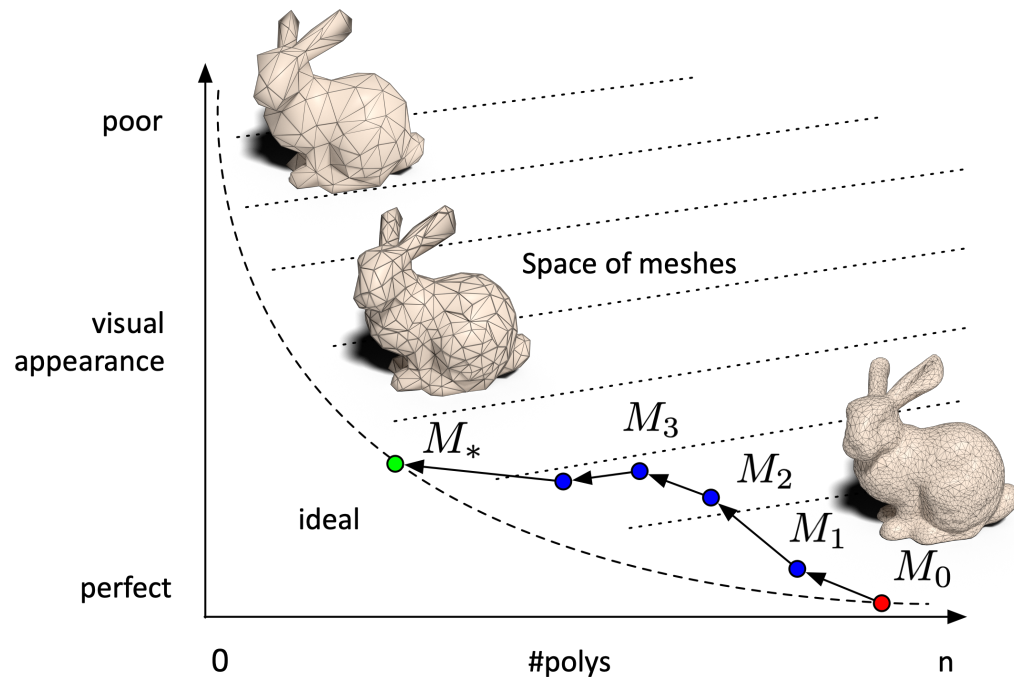
LMU Munich Media Informatics

April 2022

Germany

The Human-Machine Interaction Loop

Objective. (Polygon reduction) aims to reduce the number of faces while preserving the overall **visual appearance**.



The Human-Machine Interaction Loop

Objective. (Polygon reduction) aims to reduce the number of faces while preserving the overall **visual appearance**.

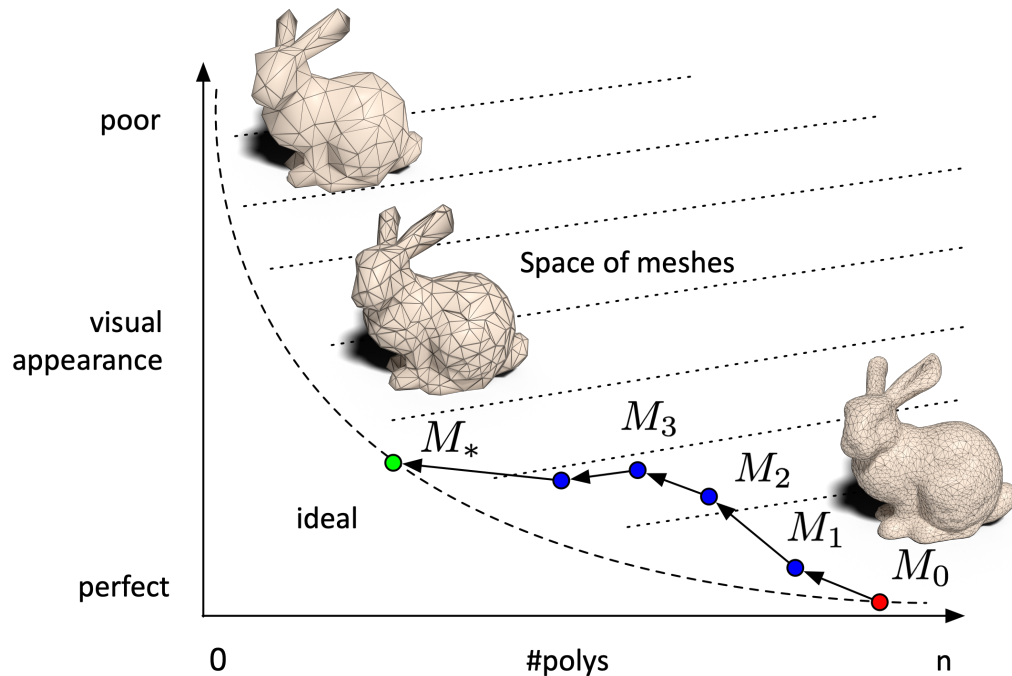
The human-machine Interaction loop:

HA $\rightarrow$ MC $\rightarrow$ HI $\rightarrow$ HA $\rightarrow$ MC $\rightarrow$ HI $\rightarrow$...

HA: Human action

MC: Machine computation

HI: Human inspection



The Human-Machine Interaction Loop

Objective. (Polygon reduction) aims to reduce the number of faces while preserving the overall **visual appearance**.

The human-machine Interaction loop:

HA $\rightarrow$ MC $\rightarrow$ HI $\rightarrow$ HA $\rightarrow$ MC $\rightarrow$ HI $\rightarrow$...

HA: Human action

MC: Machine computation

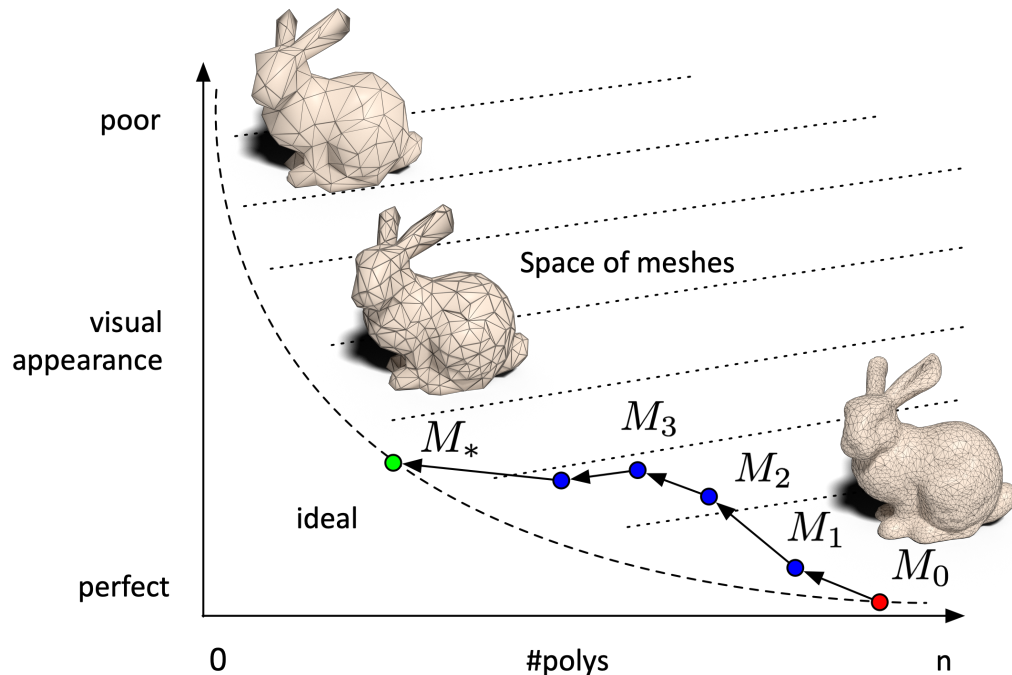
HI: Human inspection

Why not learn human actions?

Illumination galleries [Marks et al., 1997]

Material galleries [Brochu et al., 2007]

Animation galleries [Brochu et al., 2010]



The Human-Machine AI Interaction Loop

User task. Indicate a rank of given models based on their overall *visual appearance*.

Local Evaluation

(Almost-)Well informed

The human-AI interaction loop:

HC → AO → HI → HC → AO → HI → ...

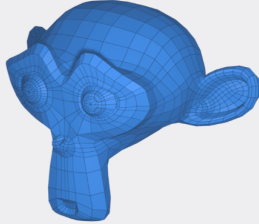
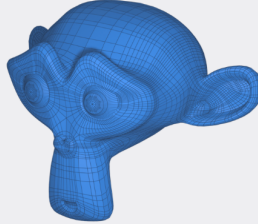
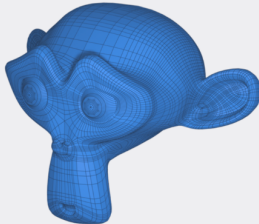
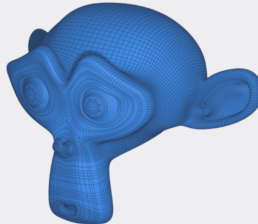
HA: Human ranking

AO: AI optimization

HI: Human inspection

Geometry galleries

Task: Rank the current four models.

A  current number of faces: 1,498 (-90.79%)	B  current number of faces: 9,214 (-62.36%)
C  current number of faces: 7,270 (-69.52%)	D  current number of faces: 20,358 (-21.31%)

D

Excellent

C

Good

A

B

Fair

Poor

Terrible

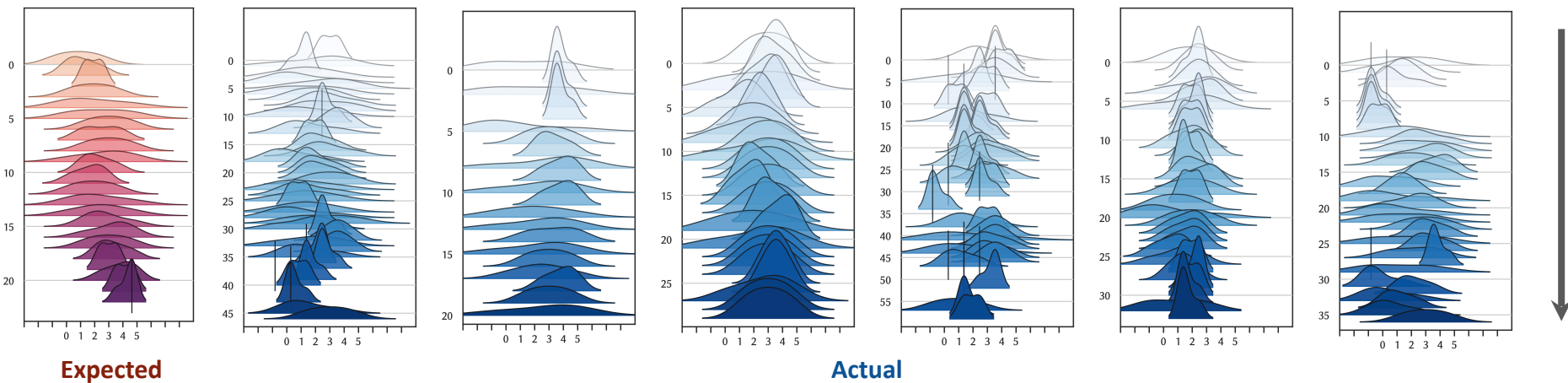
Invalid

SKIP

SUBMIT

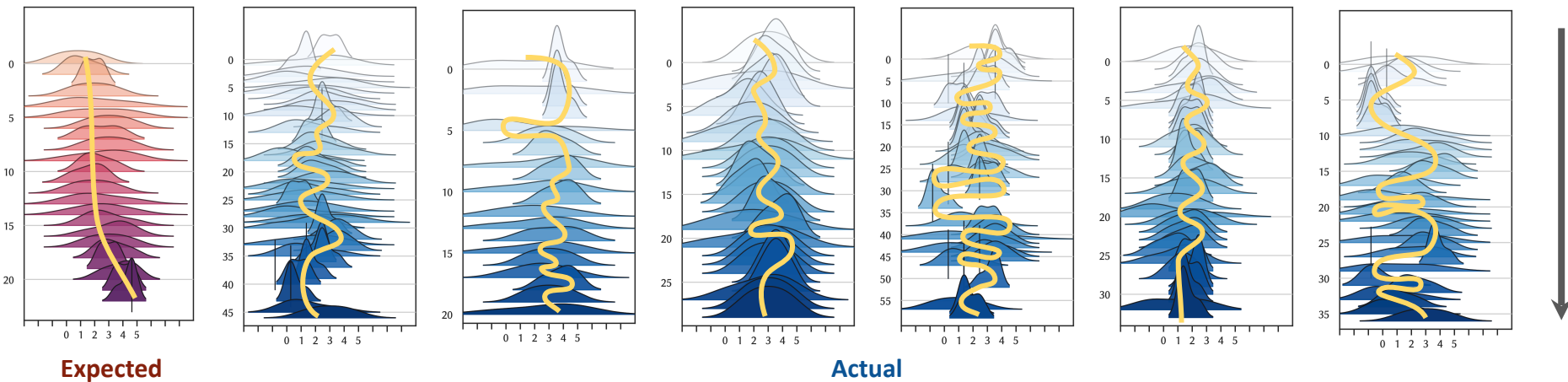
Preference Stability [Ou et al., 2021-]

- High probability mismatch between *expected* and *actual* ranking behavior



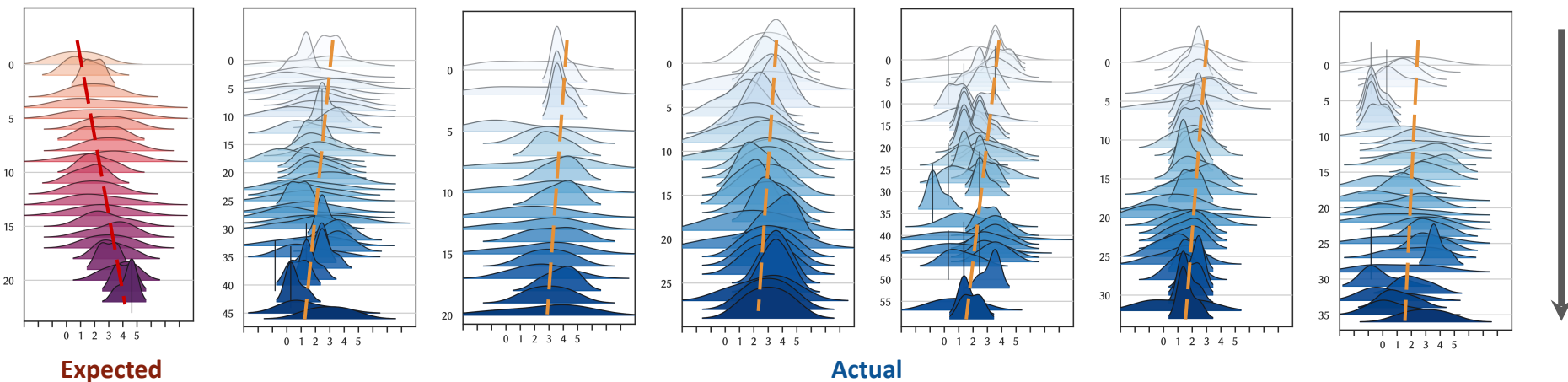
Preference Stability [Ou et al., 2021-]

- High probability mismatch between *expected* and *actual* ranking behavior
- Very unstable and significant decreasing



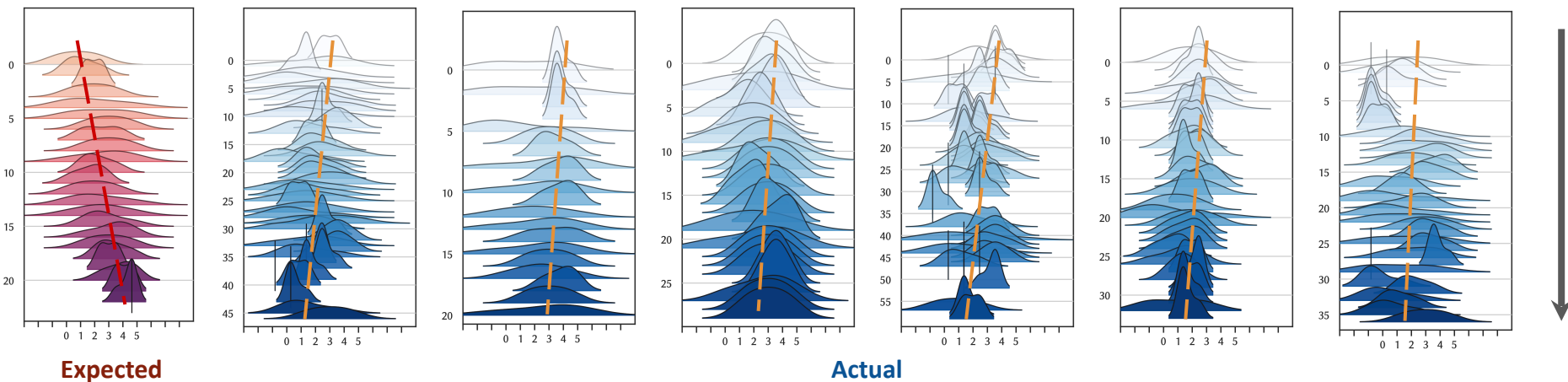
Preference Stability [Ou et al., 2021-]

- High probability mismatch between *expected* and *actual* ranking behavior
- Very unstable and significant decreasing
- Globally or even locally inconsistent/conflicting choice behavior



Preference Stability [Ou et al., 2021-]

- High probability mismatch between *expected* and *actual* ranking behavior
- Very unstable and significant decreasing
- Globally or even locally inconsistent/conflicting choice behavior
- Explanations: human errors and AI errors
 - Heuristics (anchoring, availability, representatives), decision noises (level, stable pattern, transient)
 - Algorithm assumption violation



Involving Human Decisions in Polygon Reduction

A decision concerning the following objectives:

- Reduction ratio: informed on the user interface

Involving Human Decisions in Polygon Reduction

A decision concerning the following objectives:

- Reduction ratio: informed on the user interface
- Surface quality: poor visual correlation, measure using *Chamfer* distance



Surface/Distance
Quality

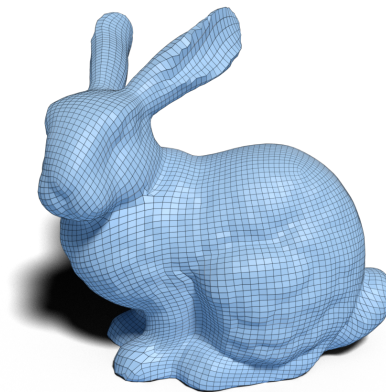
Involving Human Decisions in Polygon Reduction

A decision concerning the following objectives:

- Reduction ratio: informed on the user interface
- Surface quality: poor visual correlation, measure using *Chamfer* distance
- Wireframe quality: high visual correlation, measure using average *cell* quality



Surface/Distance
Quality



Wireframe/Cell
Quality

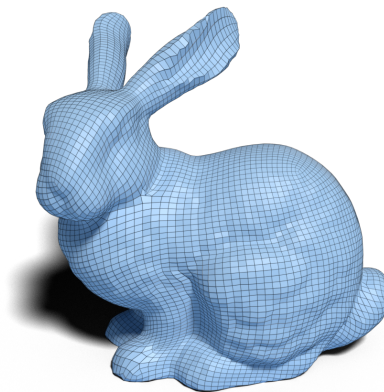
Involving Human Decisions in Polygon Reduction

A decision concerning the following objectives:

- Reduction ratio: informed on the user interface
- Surface quality: poor visual correlation, measure using *Chamfer* distance
- Wireframe quality: high visual correlation, measure using average *cell* quality
- Rendering quality: high visual correlation, measure using *SSIM* and *PSNR*



Surface/Distance
Quality

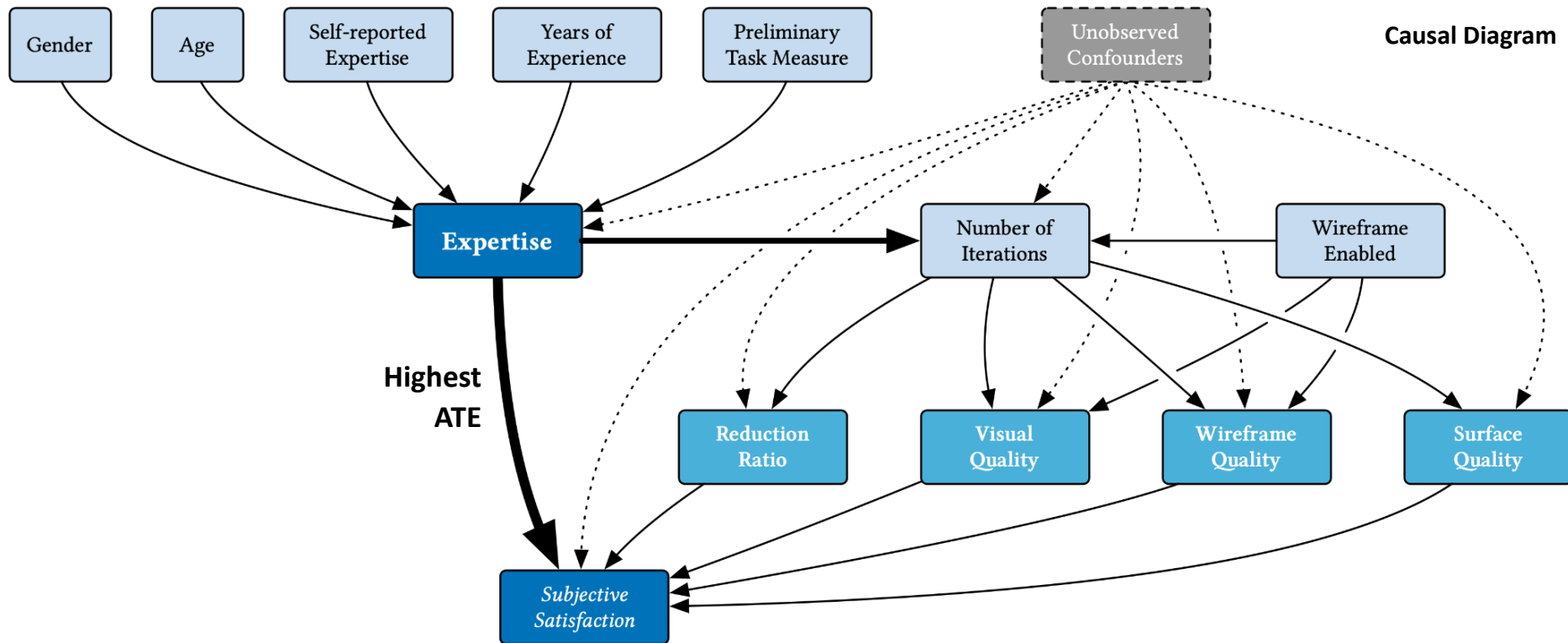


Wireframe/Cell
Quality

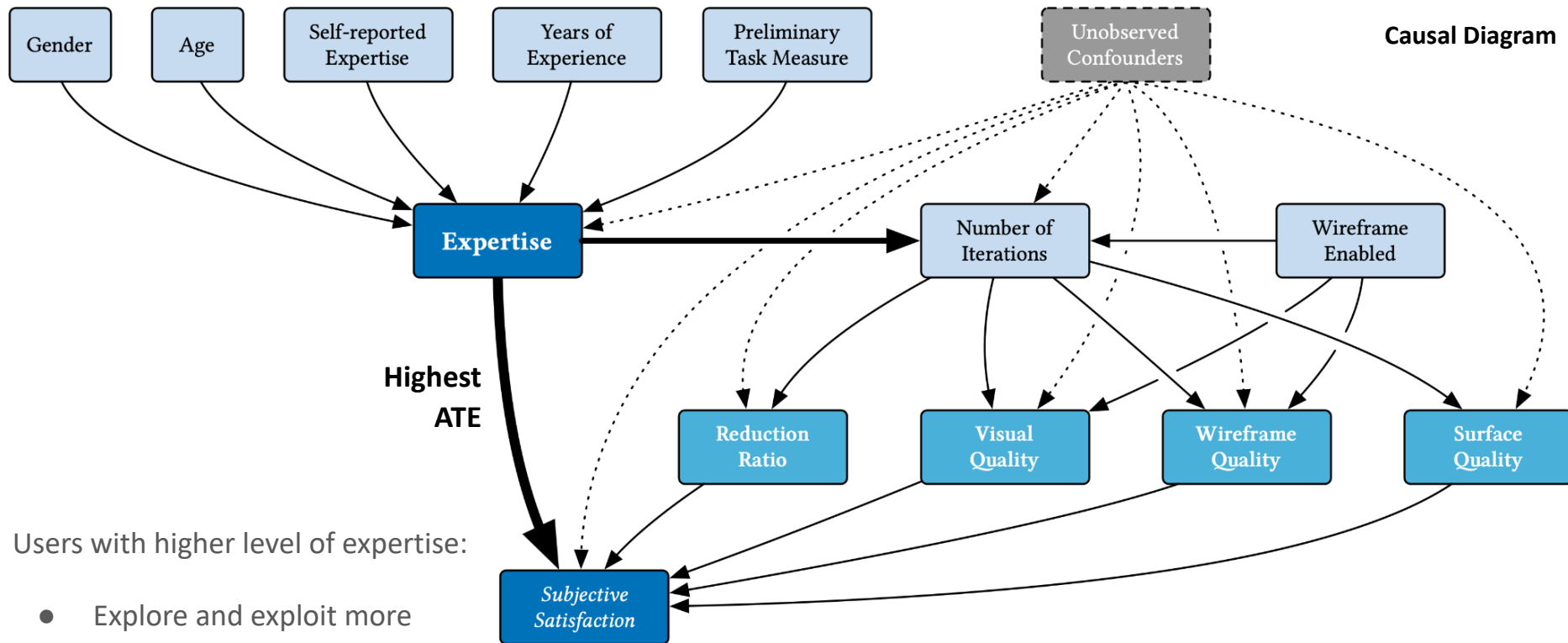


Visual/Rendering
Quality

Expertise Considered Harmful [Ou and Butz, 2022-]



Expertise Considered Harmful [Ou and Butz, 2022-]



Users with higher level of expertise:

- Explore and exploit more
- Significantly less satisfactory

Pareto Optimality [Pareto, 1912]

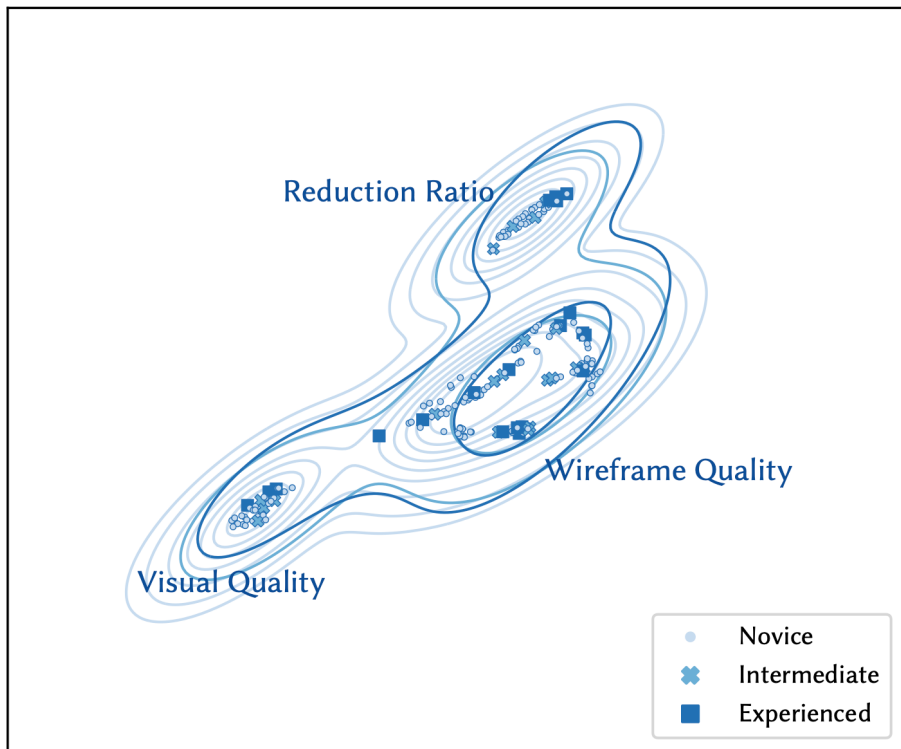
Definition. A situation where **no objective can be better** without making **at least one objective worse**.

Approx. Pareto optimality. A situation where **no objective can be better** without making **at least one objective worse not more than δ** .

The distributions of final satisfactory models.

Each of the cluster satisfies a 0.05 approximate Pareto optimality.

The surface quality objective is not perceivable significantly by participants.



[Ou and Butz, 2022]

Pareto Optimality [Pareto, 1912]

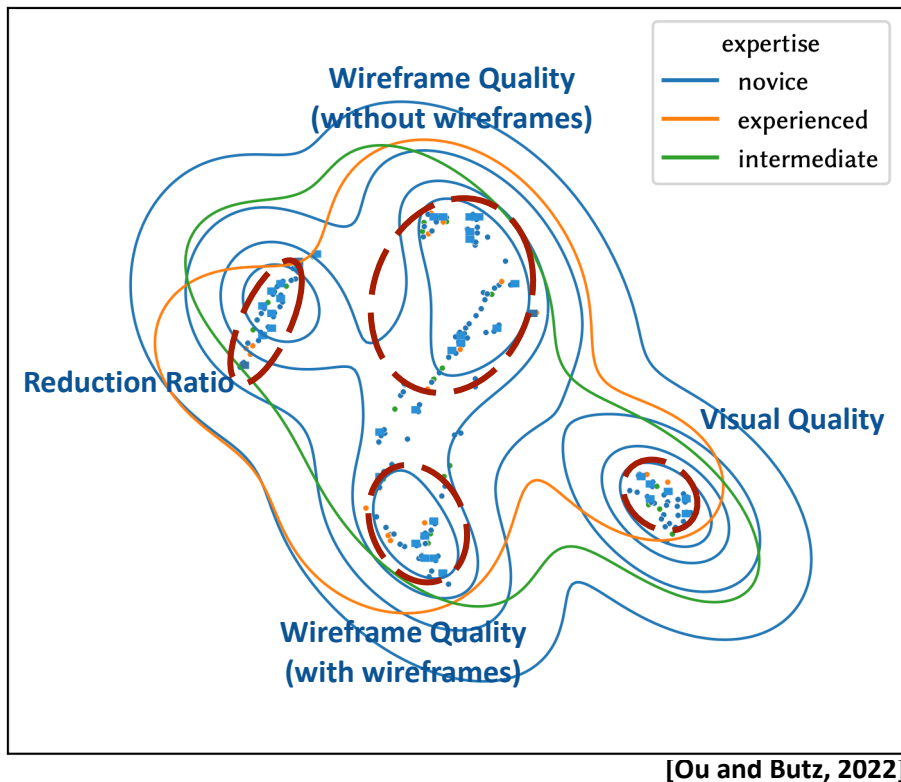
Definition. A situation where **no objective can be better** without making **at least one objective worse**.

Approx. Pareto optimality. A situation where **no objective can be better** without making **at least one objective worse not more than δ** .

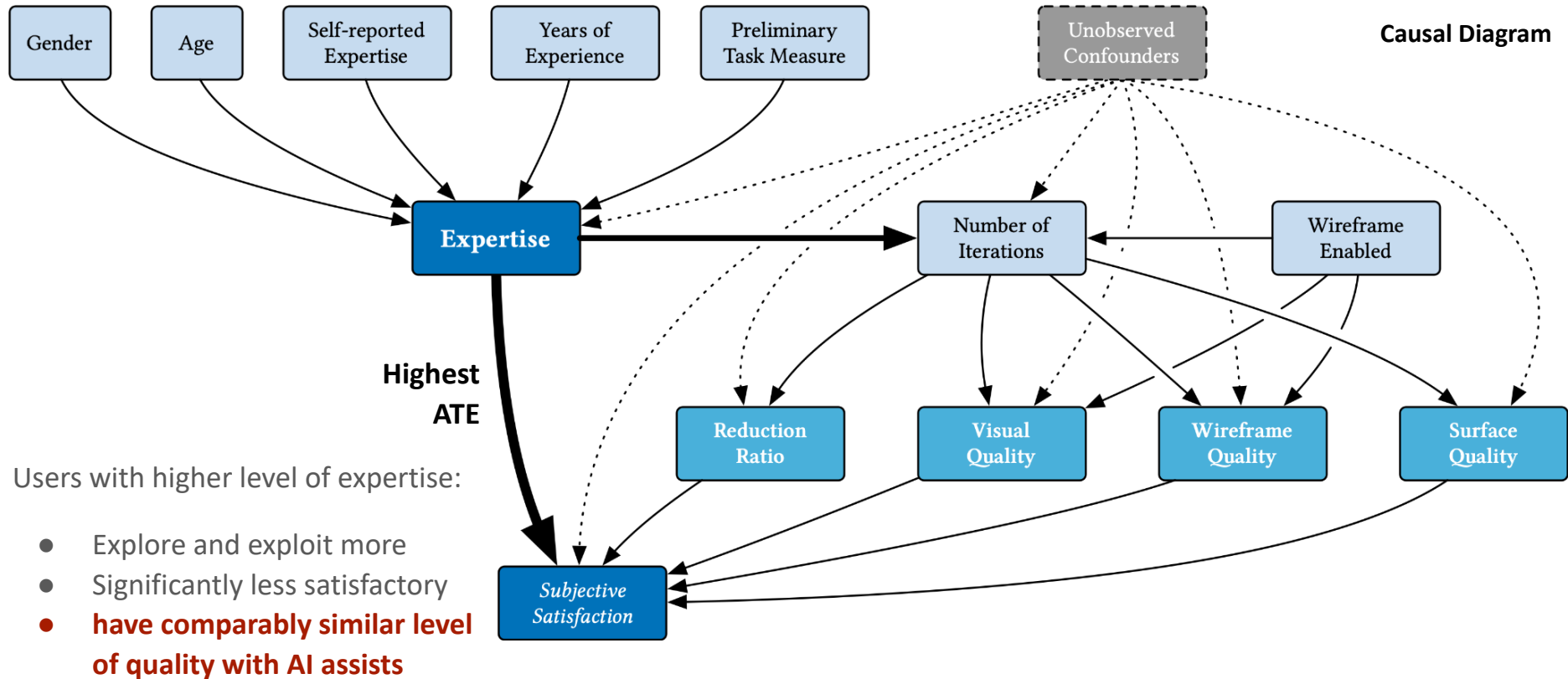
The distributions of final satisfactory models.

Each of the cluster satisfies a 0.05 approximate Pareto optimality.

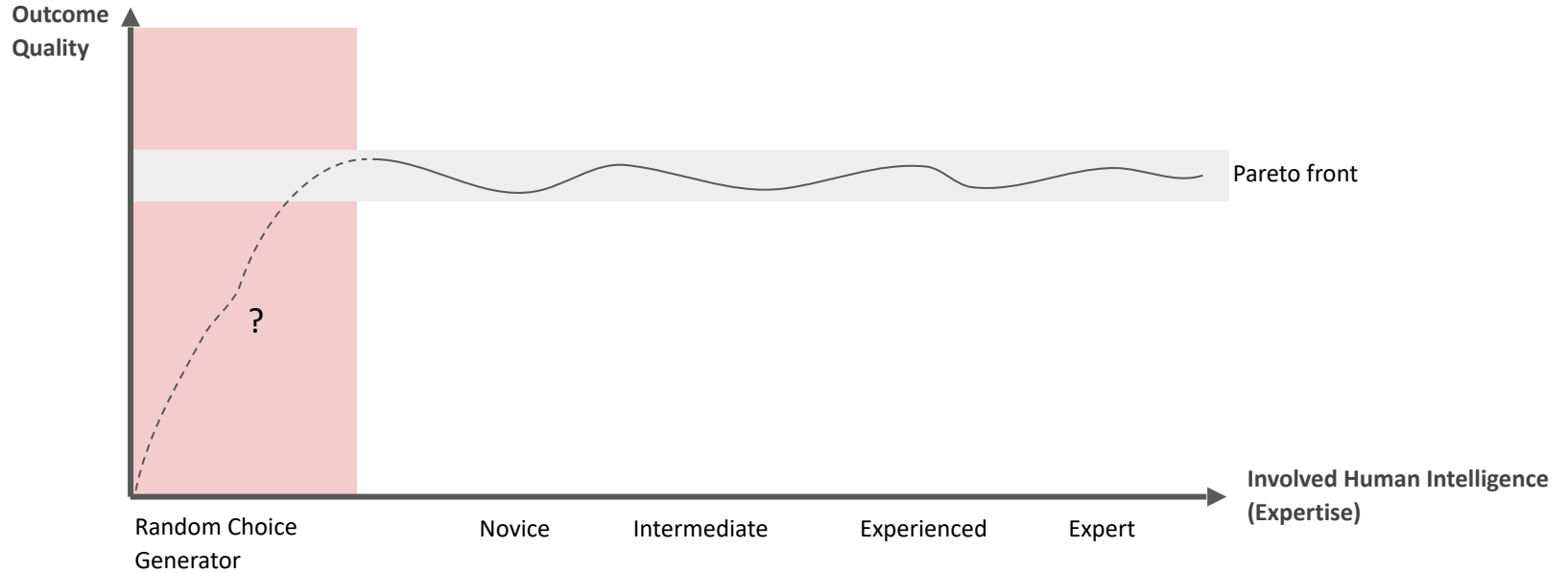
The surface quality objective is not perceivable significantly by participants.



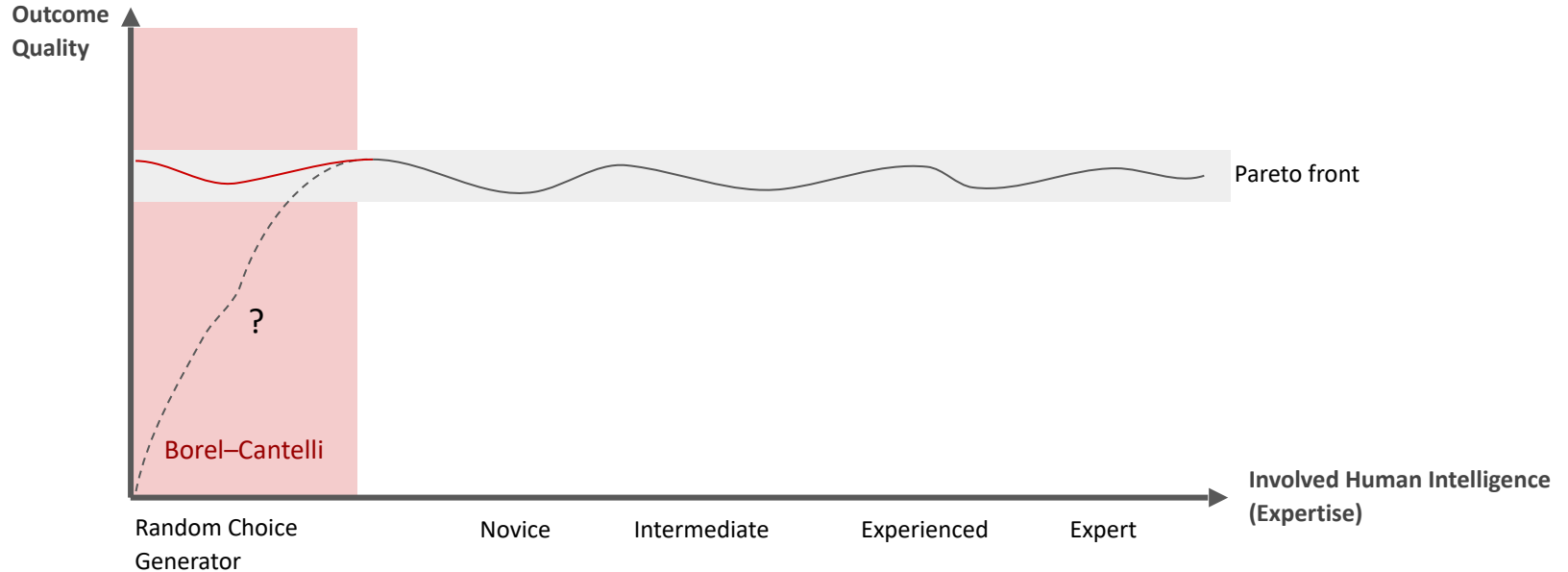
Expertise Considered Harmful [Ou and Butz, 2022-]



Harmful doesn't mean Unhelpful



Harmful doesn't mean Unhelpful

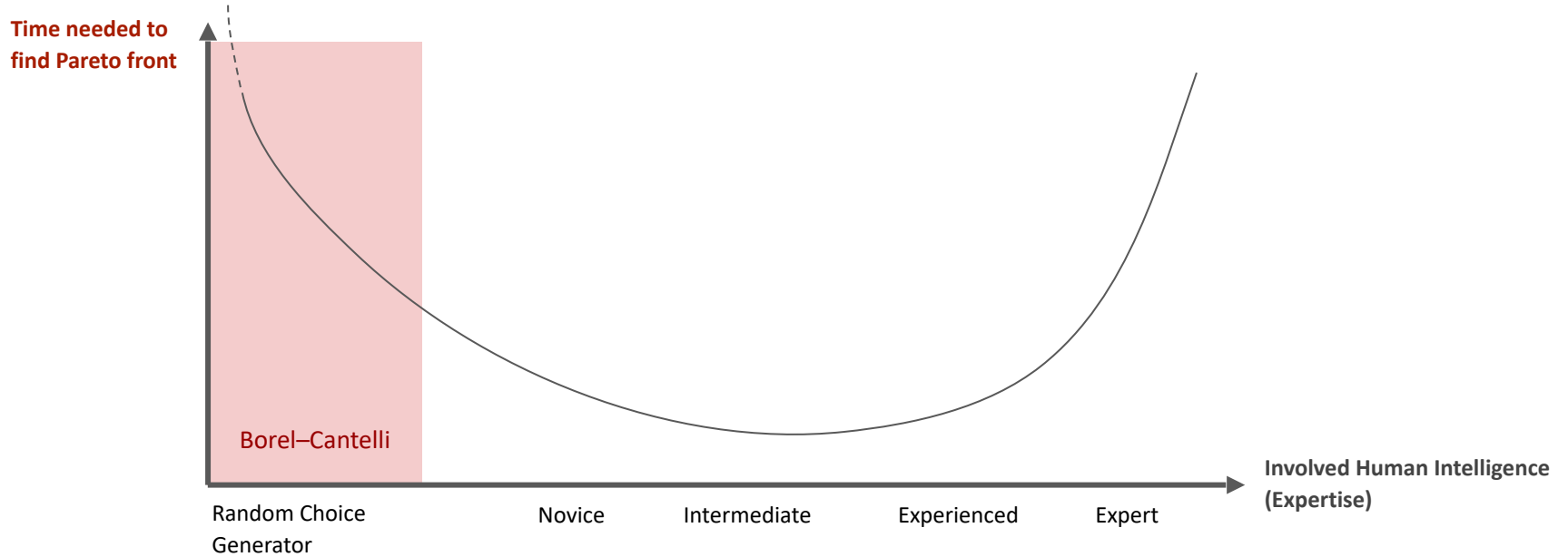


The Borel–Cantelli lemma [Borel, 1909] [Cantelli, 1917].

With infinite amount of events, the probability of observing any meaningful result is 1.0

Strictly speaking, the event happens almost surely if the Lebesgue measure is 1.

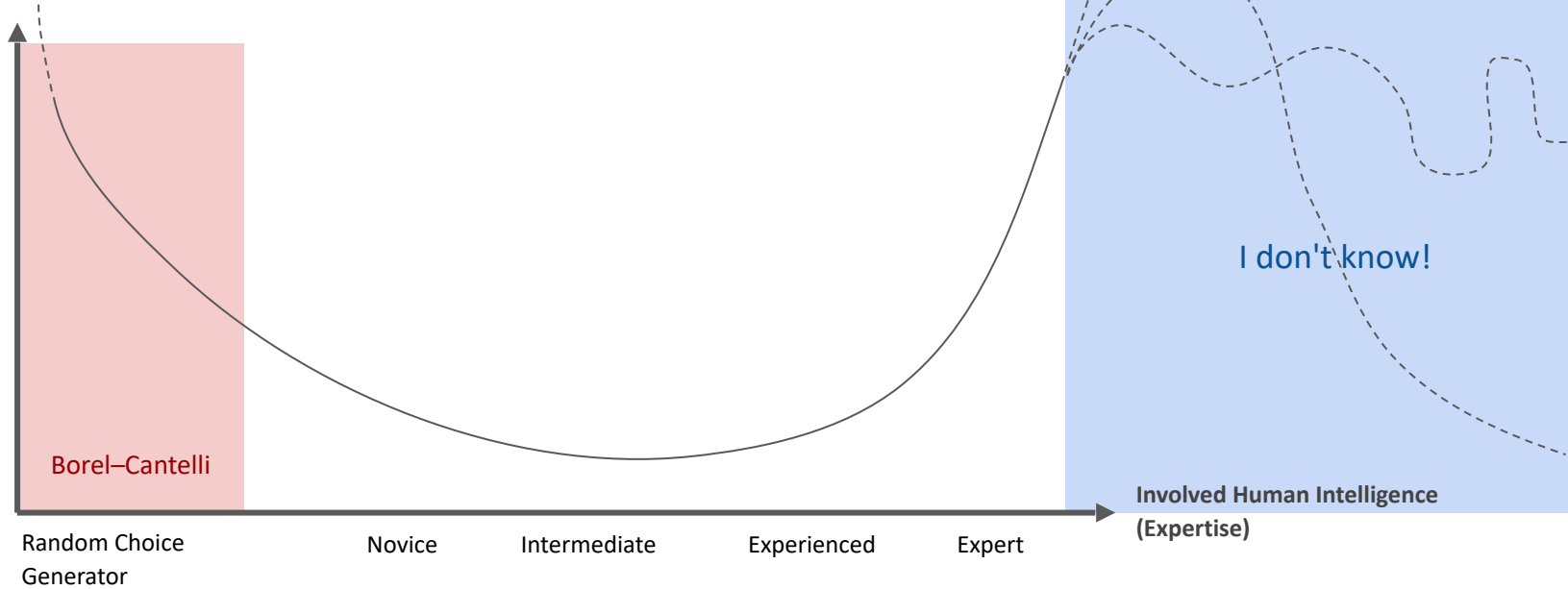
Harmful doesn't mean Unhelpful (cont.)



How could we compare expert and random generator in this case?

Harmful doesn't mean Unhelpful (cont.)

Time needed to
find Pareto front



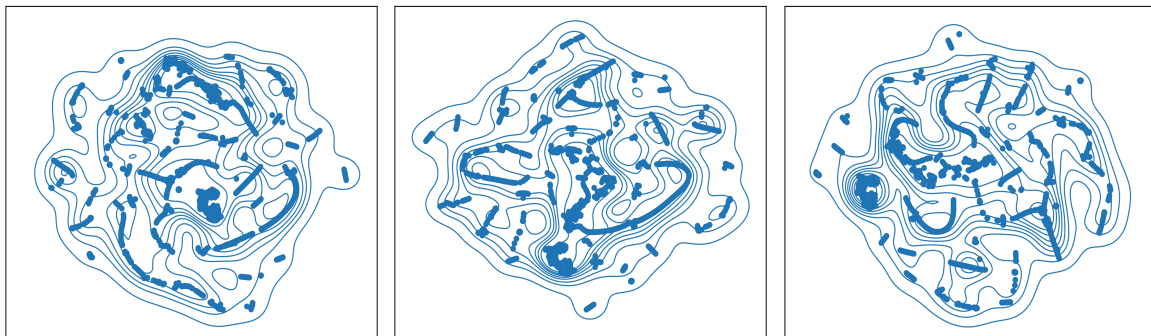
Preference Elicitation, Aggregation, and Manipulation

Individual choices regarding N objectives:

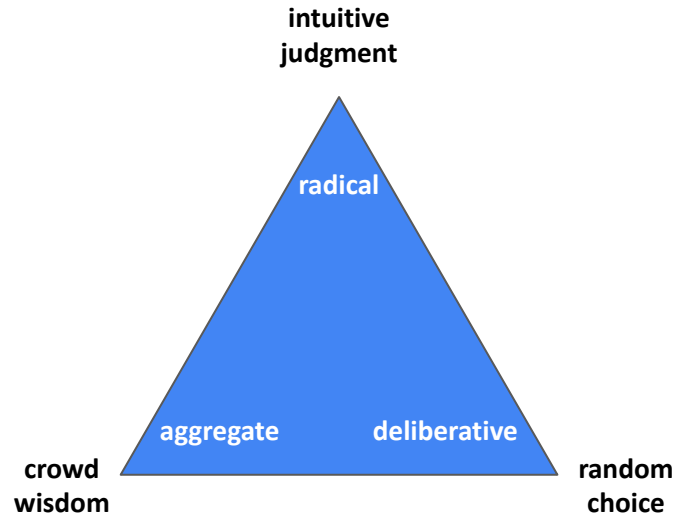
- $N = 0$: random, or choose based on prior
- $N = 1$: maximizing the objective, or satisficing
- $N = 2$: every optimized option is a Pareto frontier if objectives are orthogonal
- $N > 2$: bounded number of choices are Pareto frontiers

Aggregated crowd choices

⇒ The Arrow's Impossibility Theorem (no perfect voting)



The Decision Maker's Dilemma



Do you want *follow the intuition*, or *use queried majority vote*,
or *just make a random choice*?

Connecting Theories

Psychophysics [Engen, 1988]

Preference, decision, and choice [Aristotle, 40 BC], [Hausman, 2011]

Bounded rationality [Simon, 1955] Heuristics [Tversky and Kahneman 1974]

Satisficing, maximizing, happiness [Schwartz, 2002] Social choice [Lewis et al., 2014] [Gershman et al., 2015]

Bounded optimality [Russell and Subramanian, 1995] Provably beneficial AI [Russel, 2019]

Computational rationality [Lewis et al., 2014] [Gershman et al., 2015]

Paxos consensus [Lamport, 2001]

Axiomatic set theory [Jech, 2003]

Summary & Discussion

I argue

- Any claimed (rational) decisions are subjective (either aggregative, deliberative, or radical)
- "Bias" is largely misused under AI context (both human bias or AI bias) but better be replaced by "belief" or "prior"
- Making a decision among Pareto frontiers is nothing different than predicting the future
- "defer to human, ask permission" might not be the optimal solutions

Summary & Discussion

I argue

- Any claimed (rational) decisions are subjective (either aggregative, deliberative, or radical)
- "Bias" is largely misused under AI context (both human bias or AI bias) but better be replaced by "belief" or "prior"
- Making a decision among Pareto frontiers is nothing different than predicting the future
- "defer to human, ask permission" might not be the optimal solutions

Interesting philosophical difficulties

- Why did people make a certain choice?
- What will people do when they cannot tell a difference?
- What will people do when they do not know enough?
- Do we, as human beings, really have objectives/purposes?
- Where is the boundary between subjective preference and objective rationality?
- Is it really commensurable when inferring preferences?