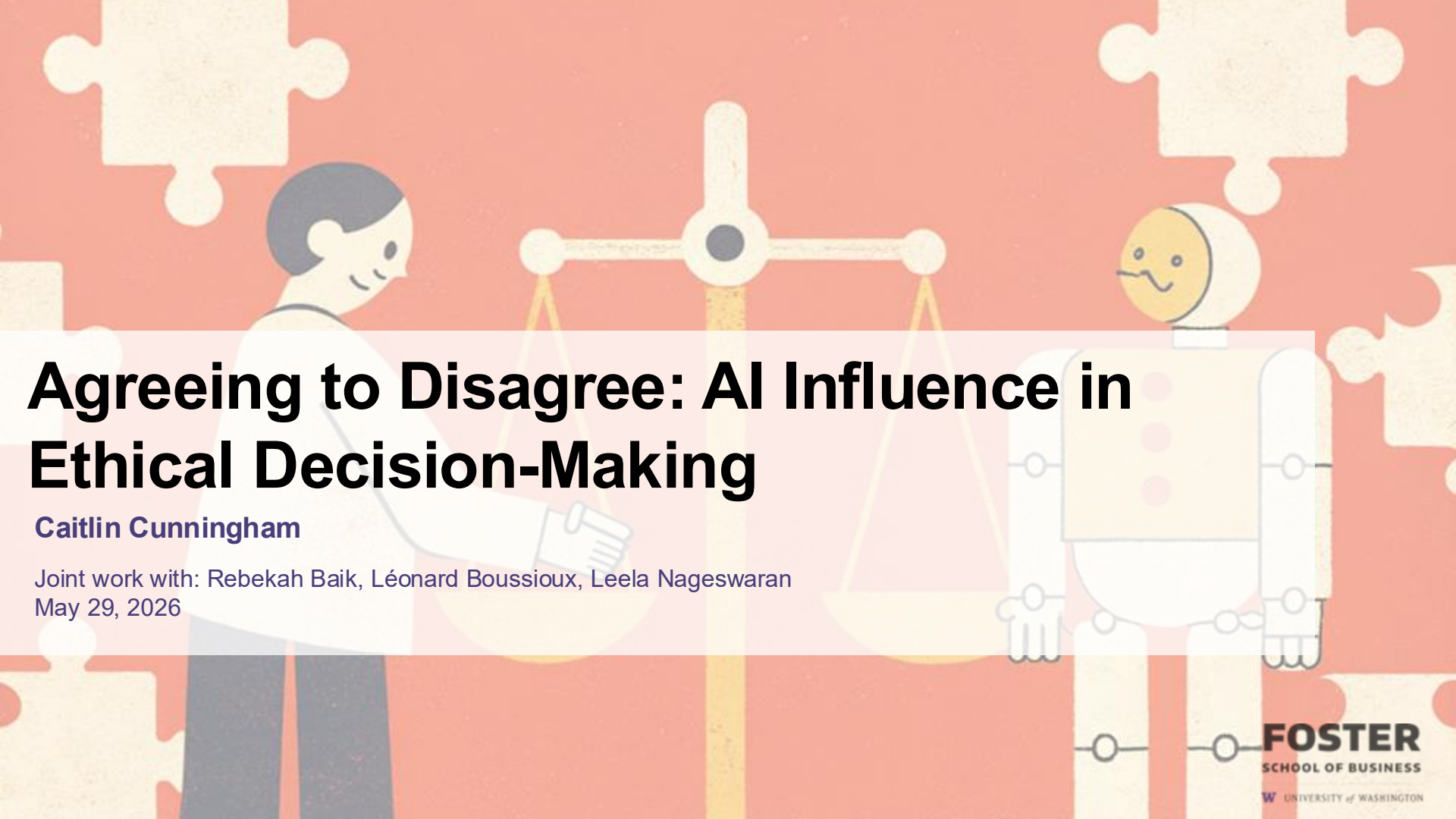


An illustration on a red background with large, light-colored puzzle pieces. In the center, a balance scale is shown. On the left pan, a person with dark hair, wearing a white lab coat and a white glove, is leaning over. On the right pan, a white robot with a yellow head and a smiling face is standing. The scale is balanced.

Agreeing to Disagree: AI Influence in Ethical Decision-Making

Caitlin Cunningham

Joint work with: Rebekah Baik, Léonard Boussieux, Leela Nageswaran
May 29, 2026

FOSTER
SCHOOL OF BUSINESS

W UNIVERSITY of WASHINGTON

Ethical Contexts: Highly Subjective

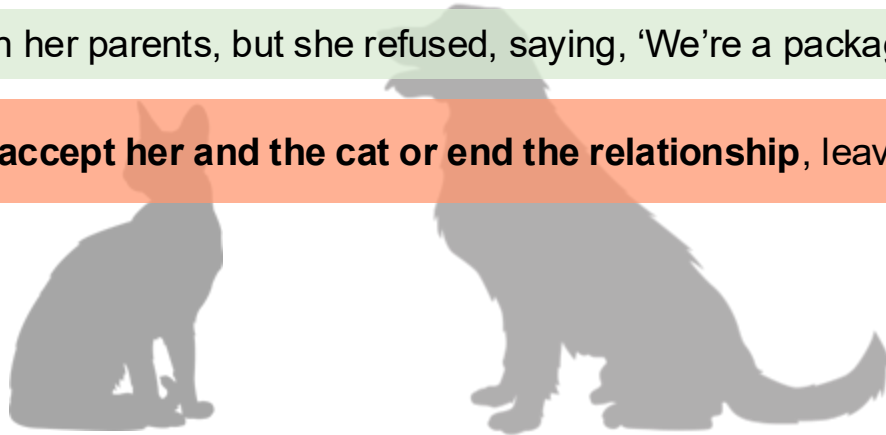
Am I in the wrong...

if I don't want my girlfriend to bring her cat when she moves in with me?

I love my girlfriend and want to marry her, but I hate her cat. It jumps everywhere, steps on me, and cries loudly all night.

I suggested the cat live with her parents, but she refused, saying, 'We're a package deal.'

She demanded I choose—**accept her and the cat or end the relationship**, leaving me stuck.



Getting Personal with AI



Watch Live



Subscribe

Sign In

ChatGPT Is Blowing Up Marriages as Spouses Use AI to Attack Their Partners

"My family is being ripped apart, and I firmly believe this phenomenon is central to why."



By **Maggie Harrison Dupré** / Published **Sep 18, 2025 11:05 AM EDT**

FORTUNE

HEALTH • CHATGPT

Gen Z is increasingly turning to affordable on-demand therapy, but mental health therapists say there are dangers many aren't considering

By **Marco Quiroz-Gutierrez** ✓
Reporter

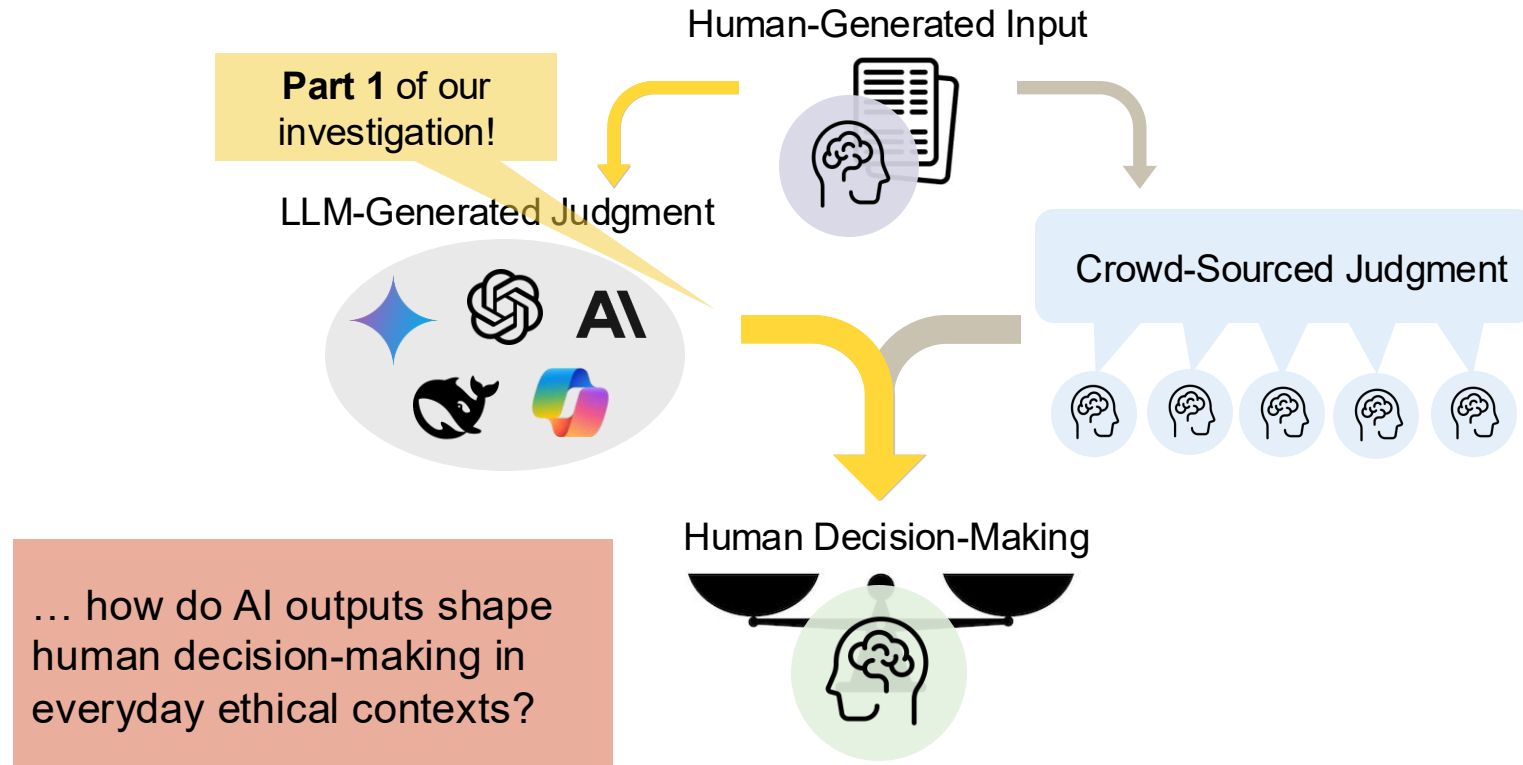
June 1, 2025 at 12:10 PM UTC



or
vice

Share Save

Human-AI Interaction in Ethical Contexts



Dataset: Reddit Anecdotes from Scruples

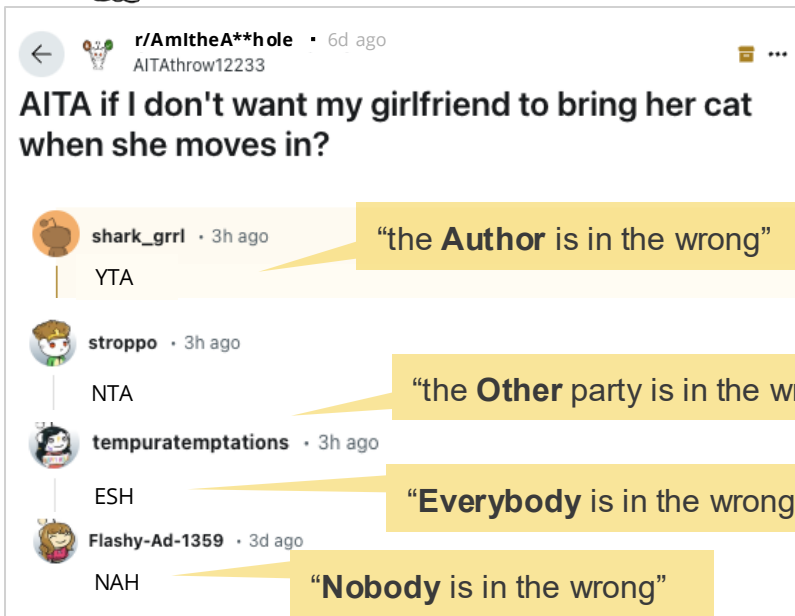


Over 32,000 real-life, ethical anecdotes shared on Reddit (Lourie et al. 2021).

Community members vote on which party is in the wrong:

- The Author (the one telling the story)
- The Other (the other person in the story)
- Everybody
- or Nobody

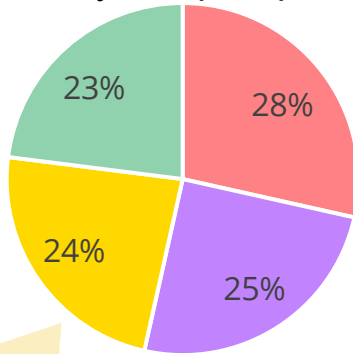
Votes are transformed into “*portions of blame*,” giving us a probability distribution that sums to 1.



Normalized Entropy as a Measure of Controversy

The entropy between “*portions of blame*” assigned by Reddit community members captures the **controversy** level of the anecdote.

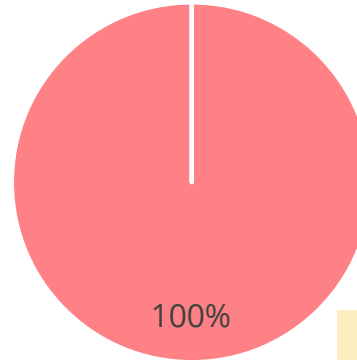
Am I wrong for expecting my mom to want to support me through my first delivery and post partum?



Example of **high entropy (0.99)**
→ **highly controversial**
(no clear consensus)

■ Author ■ Other ■ Everybody ■ Nobody

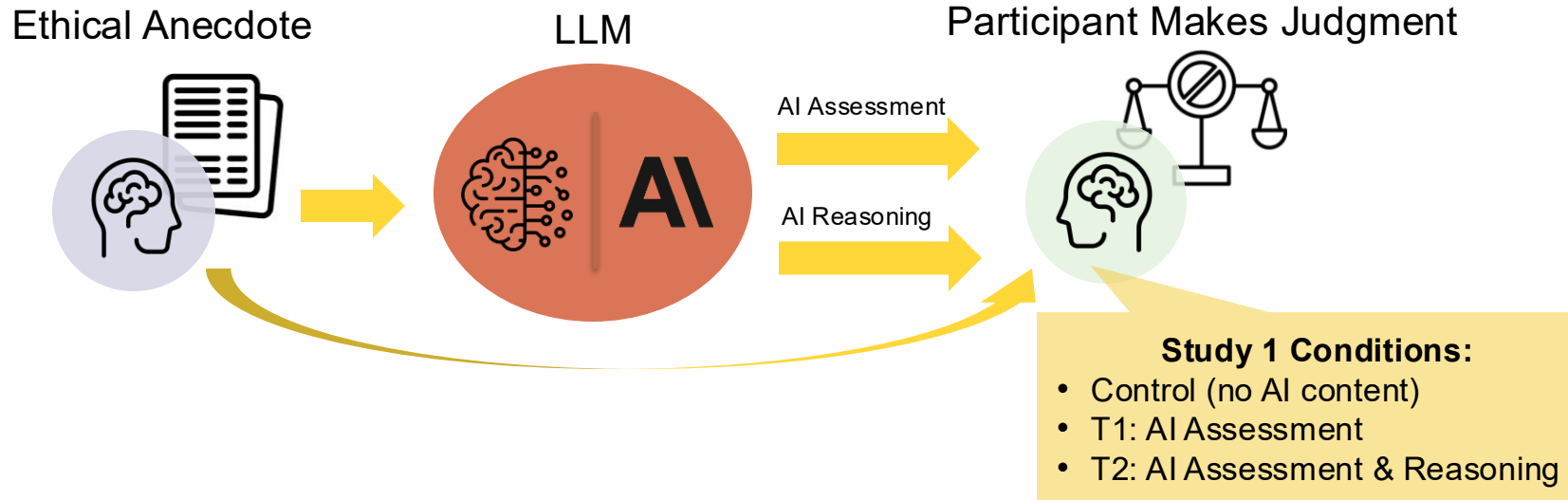
Am I wrong for punching someone who made me mad?



Example of **low entropy (0.0)**
→ **not controversial**
(clear consensus amongst community members)

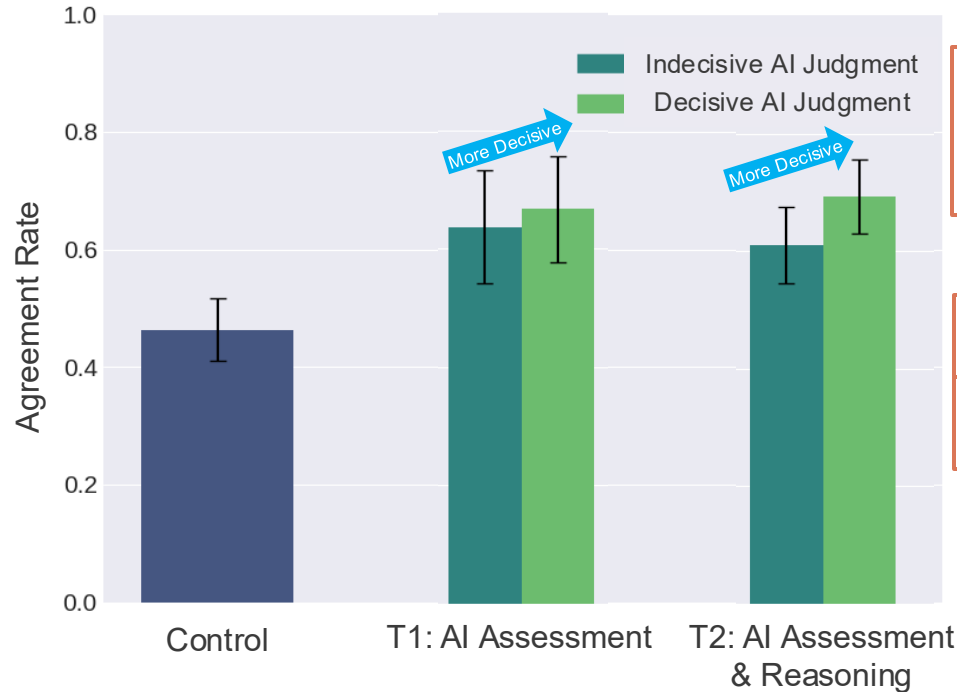
Lab Experiment: Study 1

How do AI judgments (and written reasoning behind those judgments) shape human decision-making in everyday ethical contexts?



95 student participants judge 10 anecdotes: 950 total judgments in Study 1

Exposure to AI Influences Agreement



AI Decisiveness is calculated as 1 minus the entropy between the LLM assigned portions of blame.

DV: Agrees with AI Decision

Baseline

Control

T1: AI Assessment

0.827***

T2: AI Assessment & Reasoning

0.817

Anecdote Controversy

0.283

AI Decisiveness

2.306***

(0.471)

AI Decisiveness × Controversy

-4.688**

(1.630)

Constant

-0.251

(0.149)

Anecdote Controls

✓

Observations

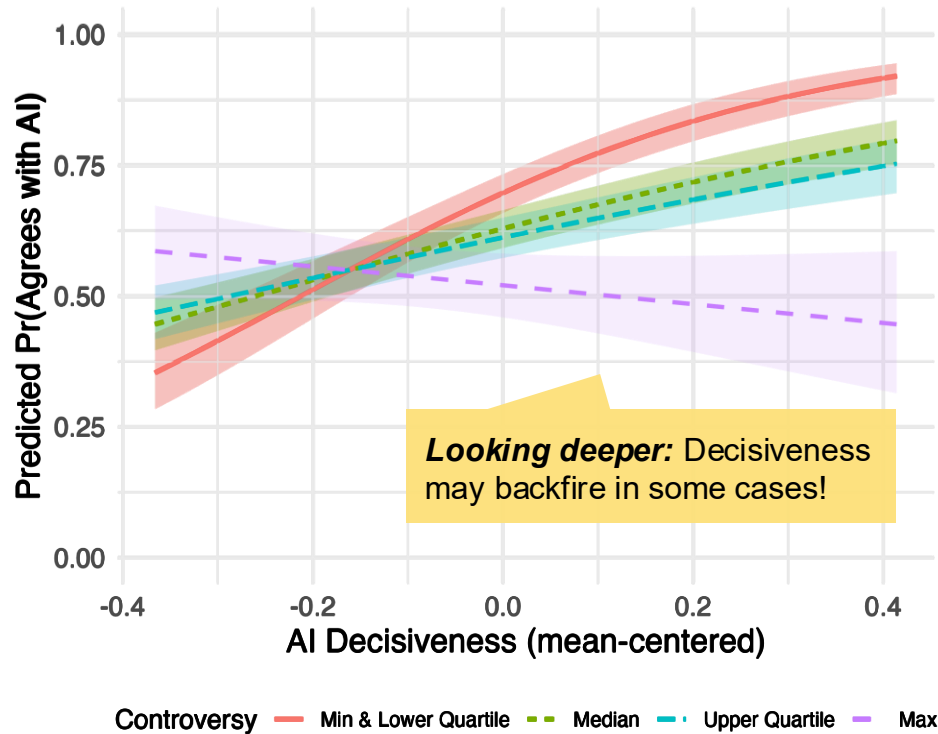
950

Note: Standard errors in parentheses.

*** $p < 0.001$, ** $p < 0.01$, * $p < 0.05$, · $p < 0.1$.

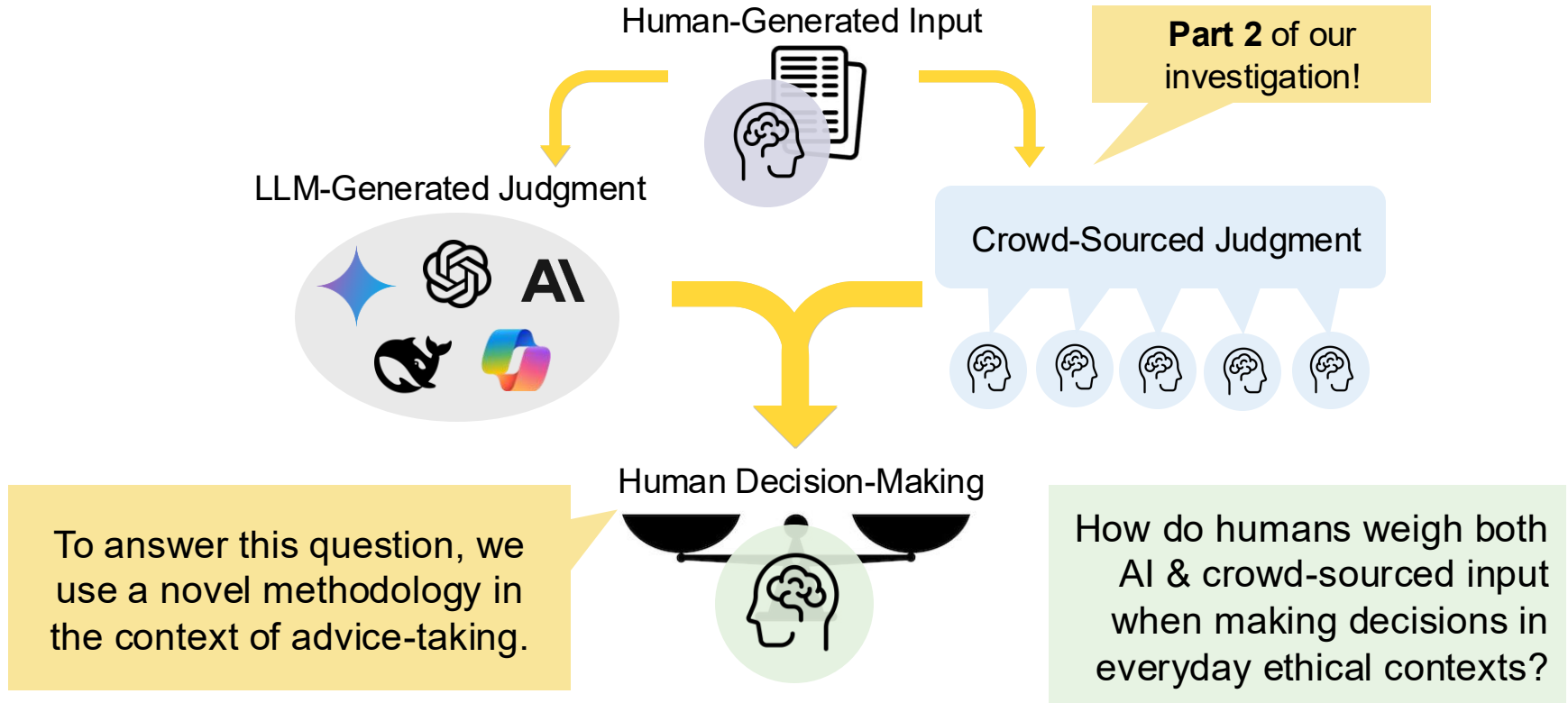
At first glance: If the AI is **decisive** (gives polarized judgments), humans agree more.

Exposure to AI Influences Agreement



DV: Agrees with AI Decision	
<i>Baseline</i>	<i>Control</i>
T1: AI Assessment	0.827*** (0.177)
T2: AI Assessment & Reasoning	0.993*** (0.180)
Anecdote Controversy	-0.757** (0.283)
AI Decisiveness	2.306*** (0.471)
AI Decisiveness × Controversy	-4.688** (1.630)
Constant	-0.251 (0.149)
Anecdote Controls	✓
Observations	950
<i>Note: Standard errors in parentheses.</i>	
*** $p < 0.001$, ** $p < 0.01$, * $p < 0.05$, · $p < 0.1$.	

Human-AI Interaction in Ethical Contexts



Revealed Preference Model of Advice Weighting

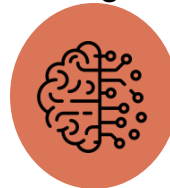
We model participants' weights of each advice source as a preference that is revealed through their judgment regarding the party (*Author* or *Other*) that is most at fault.

Participants select the alternative with higher utility.

Crowd-Sourced Judgment



AI Judgment



$$\text{Latent utility: } U_{ij}^{(k)} = W_{\text{Crowd},i} \cdot p_{\text{Crowd},j}^{(k)} + W_{\text{AI},i} \cdot p_{\text{AI},j}^{(k)} + \varepsilon_{ij}^{(k)}$$

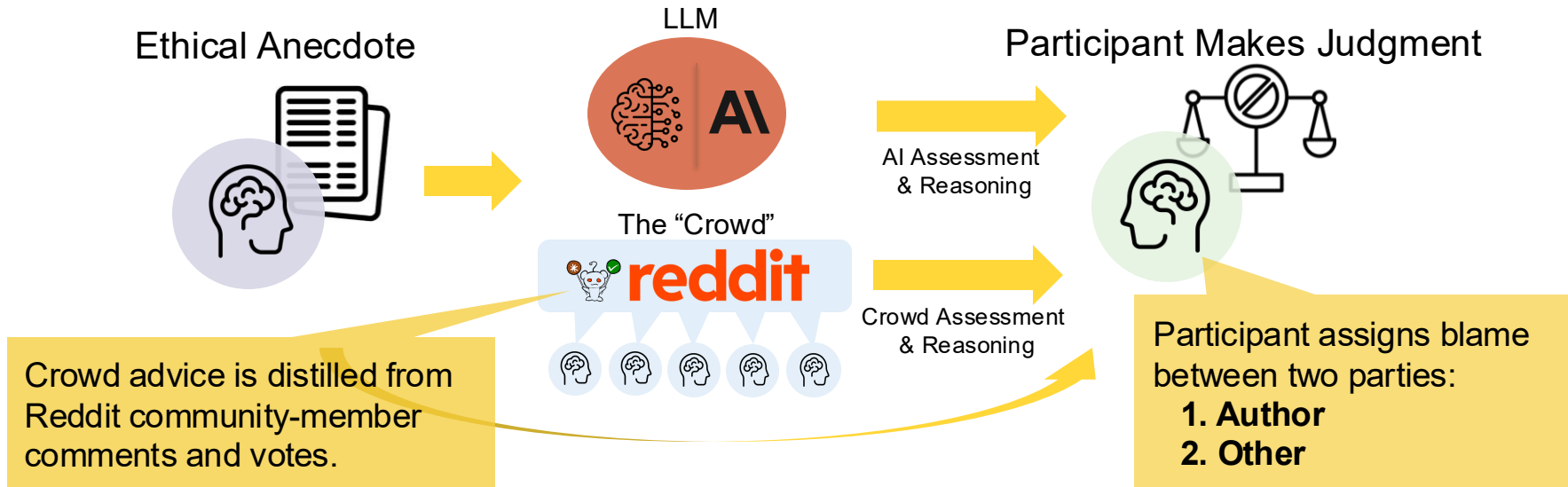
$k \in \{ \text{Author}, \text{Other} \}$

$$Pr_{ij}(\text{choose Author}) = \frac{\exp(U_{ij}^{\text{Author}})}{\exp(U_{ij}^{\text{Author}}) + \exp(U_{ij}^{\text{Other}})}$$

Weights vary randomly across participants.

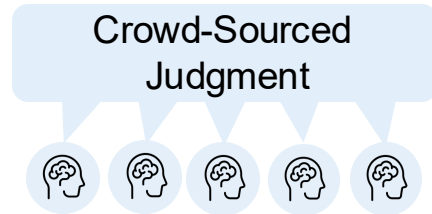
Lab Experiment: Study 2

How do humans weigh both AI & crowd-sourced input when making decisions in ethical contexts? To answer this question, we use the following lab experiment.

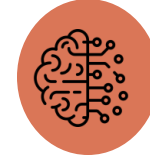


147 Prolific participants judge 10 anecdotes: 1,470 total judgments in Study 2

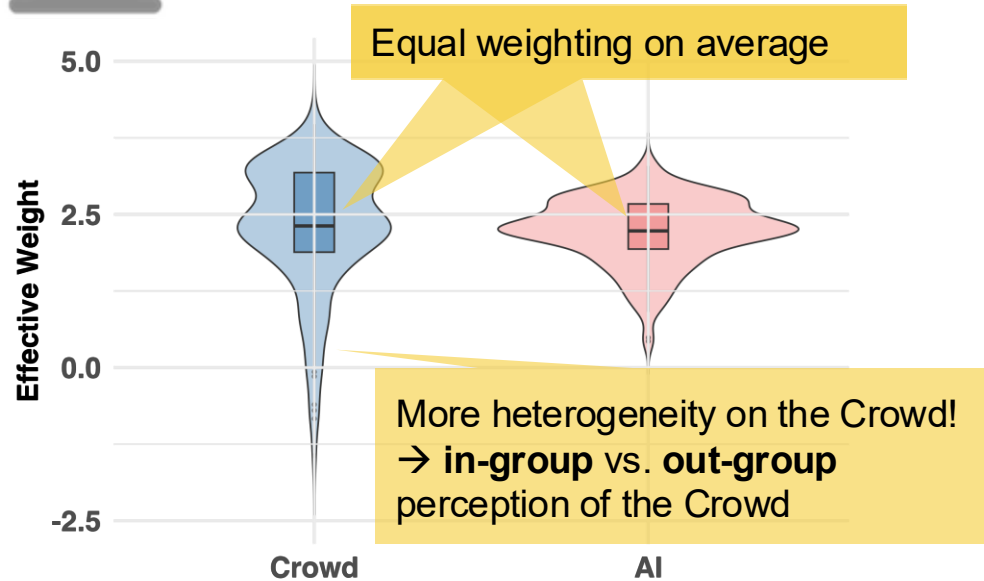
Revealed Advisor Preferences: Advice Weights



AI Judgment



All Treatment	
Weight on Crowd ($\hat{\beta}_{\text{Crowd}}$)	2.310*** (0.267)
Weight on AI ($\hat{\beta}_{\text{AI}}$)	2.225*** (0.240)
SD of Random Slope (Crowd)	1.684*** (0.293)
SD of Random Slope (AI)	1.148*** (0.313)
Log Likelihood	-471.34
Observations	1,044
Note: Standard errors in parentheses. *** $p < 0.001$, ** $p < 0.01$, * $p < 0.05$, · $p < 0.1$.	



Mechanism Analysis: Distinguishing Similarity in Reasonings

But what leads to differences in advisor preference across contexts?

$$U_{ij}^{(k)} = W_{Crowd,i} \cdot p_{Crowd,j}^{(k)} + W_{AI,i} \cdot p_{AI,j}^{(k)} + \dots + \varepsilon_{ij}^{(k)}$$

$$W_{Crowd \times CS,i} \left(p_{Crowd,j}^{(k)} \cdot CS_j \right) + W_{AI \times CS,i} \left(p_{AI,j}^{(k)} \cdot CS_j \right)$$

AI Assessment of Fault Probabilities:

Author: 67%

Other: 36%

AI Reasoning:

The author bears 67% of the responsibility in this scenario...

... [content cropped for brevity] ...

High divergence:
AI offers unique advice
→ higher AI weight

...]

Cosine Similarity

[0.0249015465,
0.02637814916,
...]

**text embedding of
Crowd Reasoning**

Crowd's Assessment of Fault Probabilities:

Author: 78%

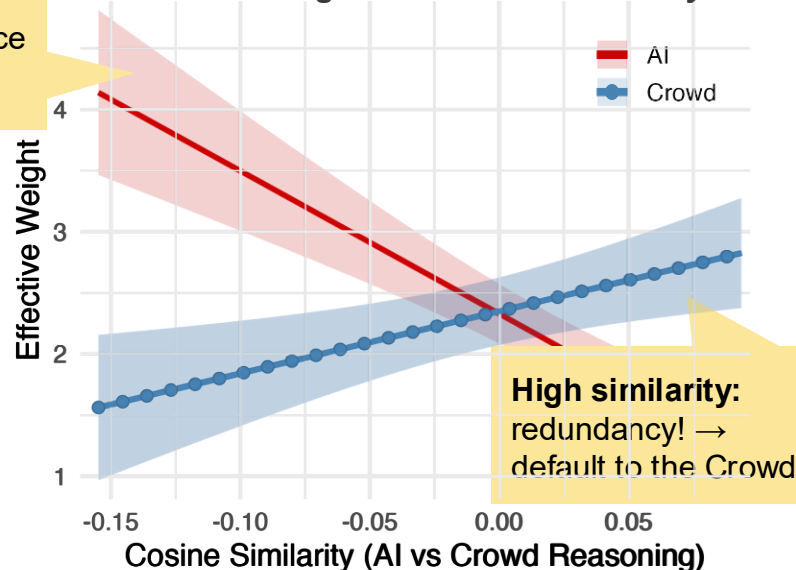
Other: 22%

Crowd's Reasoning:

The majority view holds that the author was wrong for revealing ...

... [content cropped for brevity] ...

Advisor Weights vs. Cosine Similarity



Note: Cosine Similarity is mean-centered for interpretability.

Managerial Implications

AI outputs are persuasive in moral scenarios → need guardrails to avoid one-sided nudging.

LLM-Generated Judgment



Human-Generated Input



Crowd-Sourced Judgment



Human Decision-Making



Similar advice → Crowd preferred!

Use scalable AI to provide meaningfully unique advice.

Heterogeneity in reliance on the advisors → offer personalization instead of a one-size-fits-all.

You be the judge! Who is in the wrong?

**Is the boyfriend in the wrong
if he doesn't want his girlfriend to bring her cat when she moves in?**

"The boyfriend is
in the wrong."



"The girlfriend is
in the wrong."



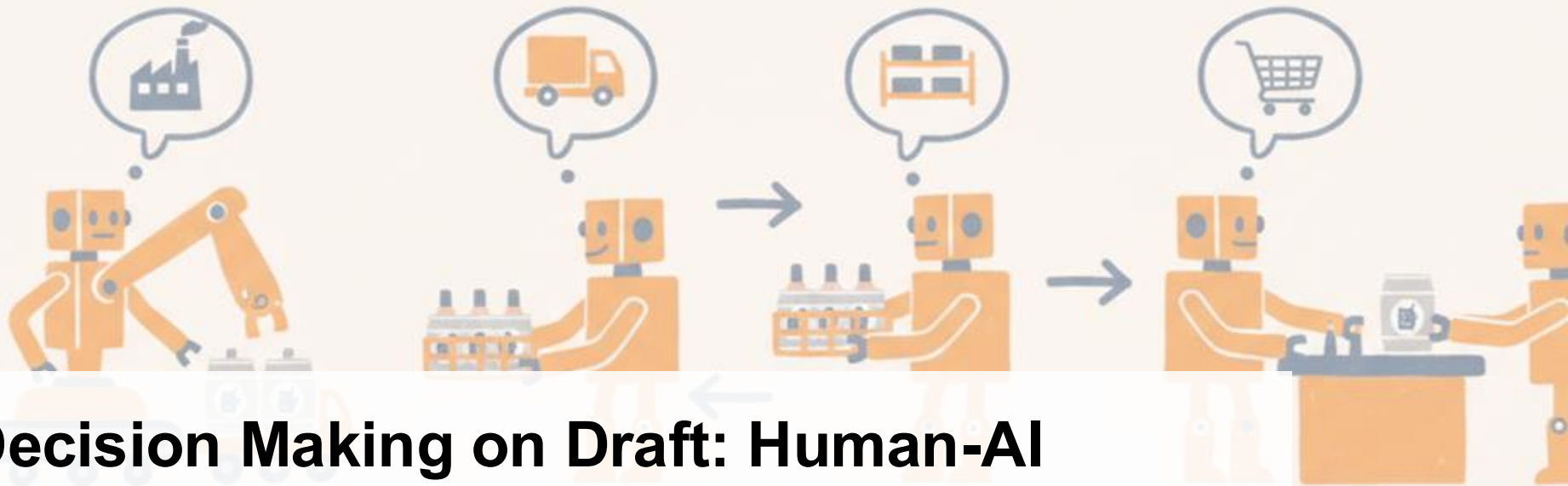
The background features a stylized illustration. In the center, a human figure with dark hair and a purple shirt is shaking hands with a white robot figure. They are positioned on a large, dark grey scale of justice. The human is on the left pan, and the robot is on the right pan. The scale is balanced. The background is a warm orange color with large, light-colored gear shapes. The text 'Thank you!' is centered over the handshake, with the email address 'csc27@uw.edu' below it.

Thank you!

csc27@uw.edu

FOSTER
SCHOOL OF BUSINESS

W UNIVERSITY of WASHINGTON



Decision Making on Draft: Human-AI Collaboration in the Beer Game

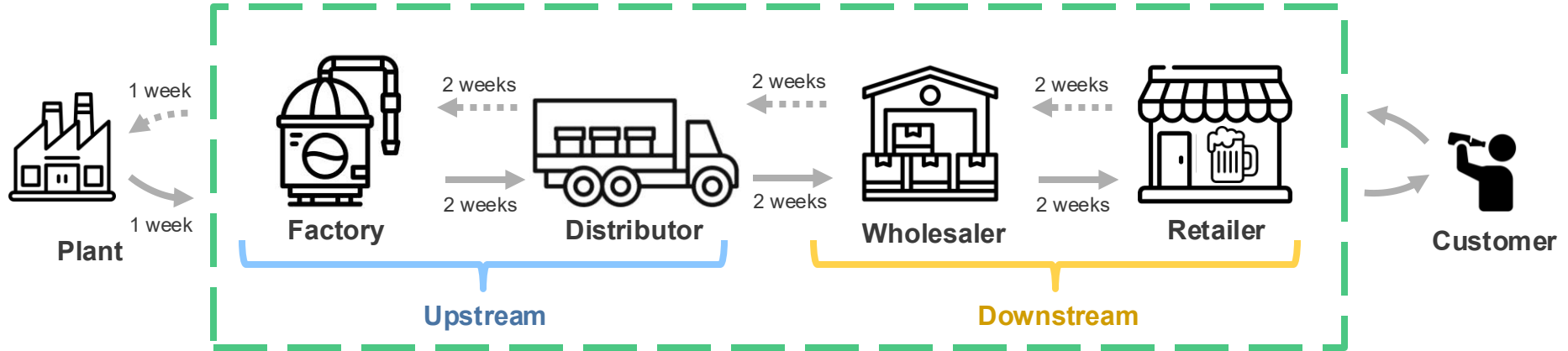
Caitlin Cunningham

Joint work with: Rebekah Baik, Leela Nageswaran, Léonard Boussieux

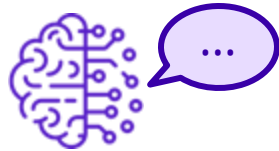
FOSTER
SCHOOL OF BUSINESS

W UNIVERSITY of WASHINGTON

Human-AI Collaboration in Supply Chains



Qualitative AI Advisor



Quantitative AI Advisor



AI & Ethics in Business

**Newsrooms Are
Considering**

By Meg Little Reilly,
Community News at
Published Apr 22, 2024,



**WHO, IT
on AI us**

11 July 2025 | Ne

≡ **W I R**

CAROLINE HASKINS

People Who Say They're Experiencing AI Psychosis Beg the FTC for Help

The Federal Trade Commission received 200 complaints mentioning ChatGPT between November 2022 and August 2025. Several attributed delusions, paranoia, and spiritual crises to the chatbot.

**The
Guardian**

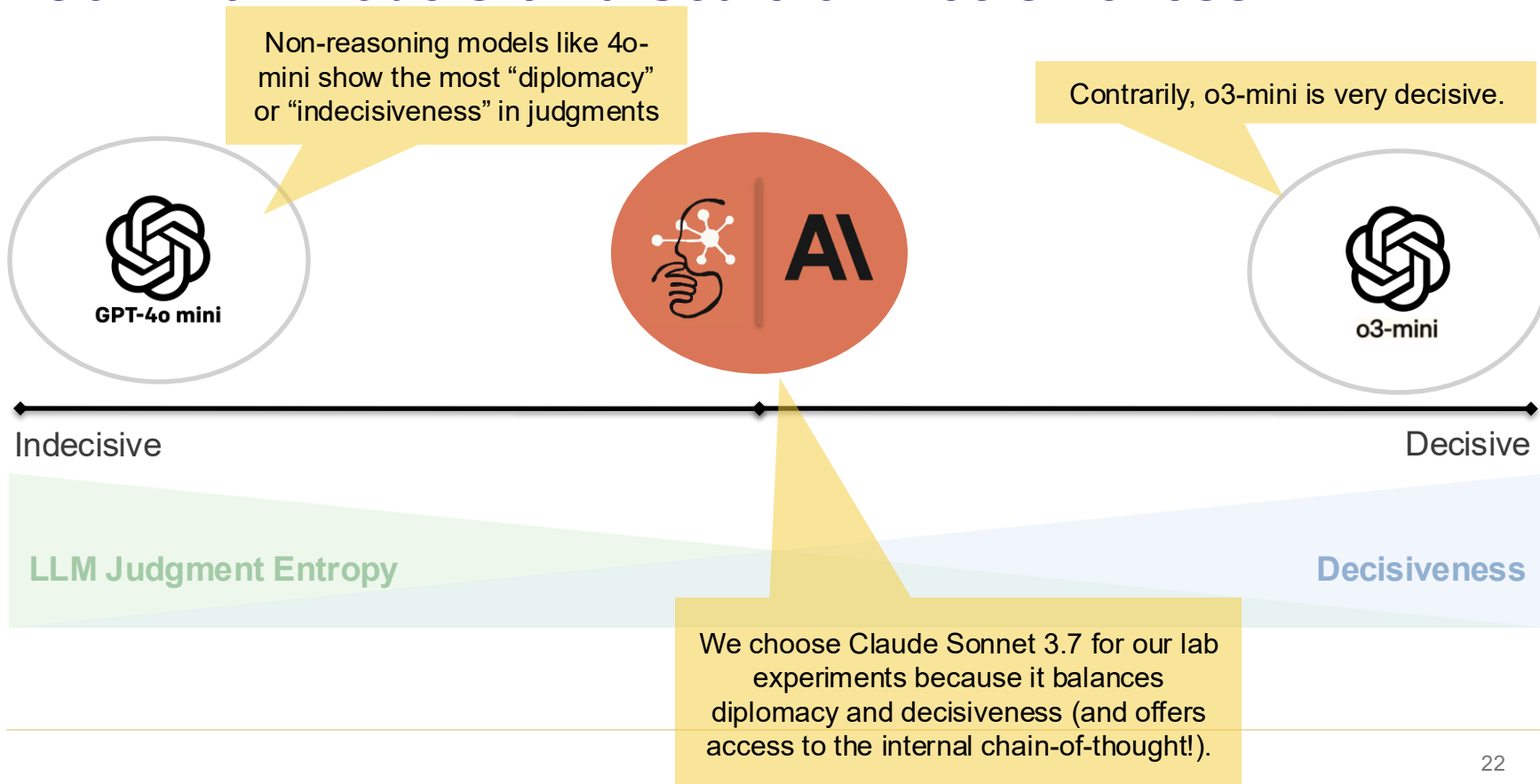
'Sliding into an abyss': experts warn over rising use of AI for mental health support

Rachel Hall

Sat 30 Aug 2025 06:00 EDT

Therapists say they are seeing negative impacts of people increasingly turning to AI chatbots for help

Common Models on a Scale of ‘Decisiveness’



Study 1: Appendix Slides

Study 1 Conditions

Treatment 1: participant weighs the AI's assessment when judging

AI Decision: Author

AI Assessment of Fault Probabilities:

Author: 70%

Other: 30%

Everybody: 0%

Nobody: 0%

AI Reasoning:

The author bears the majority of blame in this situation. While having productivity concerns is valid, denying a reasonable lunch break around noon (a standard time for meals) and asking someone to wait an additional 2 hours (until approximately 2PM) is inconsiderate of basic human needs. The author was prioritizing job performance over the shadow's wellbeing. The previous offer at 9AM isn't equivalent, as that's breakfast time, not lunch.

Treatment 2: participant weighs both the AI's assessment and reasoning

Am I wrong for delaying my coworker's lunch?

I work as a TV tech and am a relatively new hire (approximately three months) who has been performing well. As a result, I was assigned a very newly hired tech to shadow me yesterday.

... [content cropped for brevity] ...

Am I in the wrong? I simply wanted to complete a significantly delayed job and avoid looking bad at work.

Who is at fault in this situation?

Author

Other

Everybody

Nobody

1

2

3

4

5

How confident are you in your response?



How helpful was the AI assessment

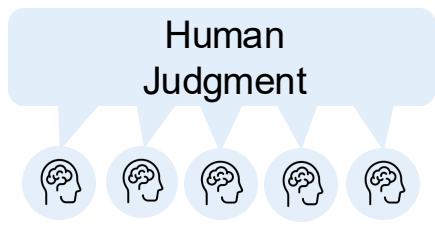


Control: participant judges without any advice

Revealed Preference Model of Advice Weighting

We model participants' weights of each advice source as a preference that is revealed through their judgment regarding the party that is most at fault.

Participants select the alternative with higher utility.



AI Judgment



$$\text{Latent utility: } U_{ij}^{(k)} = W_{Human,i} \cdot p_{Human,j}^{(k)} + W_{AI,i} \cdot p_{AI,j}^{(k)} + \varepsilon_{ij}^{(k)}$$

Study 1: $k \in \{Author, Other, Everybody, Nobody\}$

Study 2: $k \in \{Author, Other\}$

$$Pr_{ij}(\text{choose Author}) = \frac{\exp(U_{ij}^{Author})}{\sum_k \exp(U_{ij}^{Author})}$$

Weights vary randomly across participants.

Study 1: Aggregate Results

	(1)	(2)
Weight on Human ($\hat{\beta}_{\text{control}}$)	1.563*** (0.236)	1.599*** (0.238)
Weight on AI ($\hat{\beta}_{\text{AI}}$)	2.352*** (0.204)	—
Weight on AI with No Reasoning ($\hat{\beta}_{\text{AI,T1}}$)	—	2.202*** (0.253)
Weight on AI with Reasoning ($\hat{\beta}_{\text{AI,T2}}$)	—	2.503*** (0.298)
SD of Random Slope (Human)	0.005 (0.580)	0.002 (0.594)
SD of Random Slope (AI)	1.068*** (0.259)	—
SD (AI with No Reasoning)	—	0.831* (0.403)
SD (AI with Reasoning)	—	1.354*** (0.366)
Log Likelihood	−567.50	−567.14
Observations	605	605
Note: Standard errors in parentheses. *** $p < 0.001$, ** $p < 0.01$, * $p < 0.05$, · $p < 0.1$		

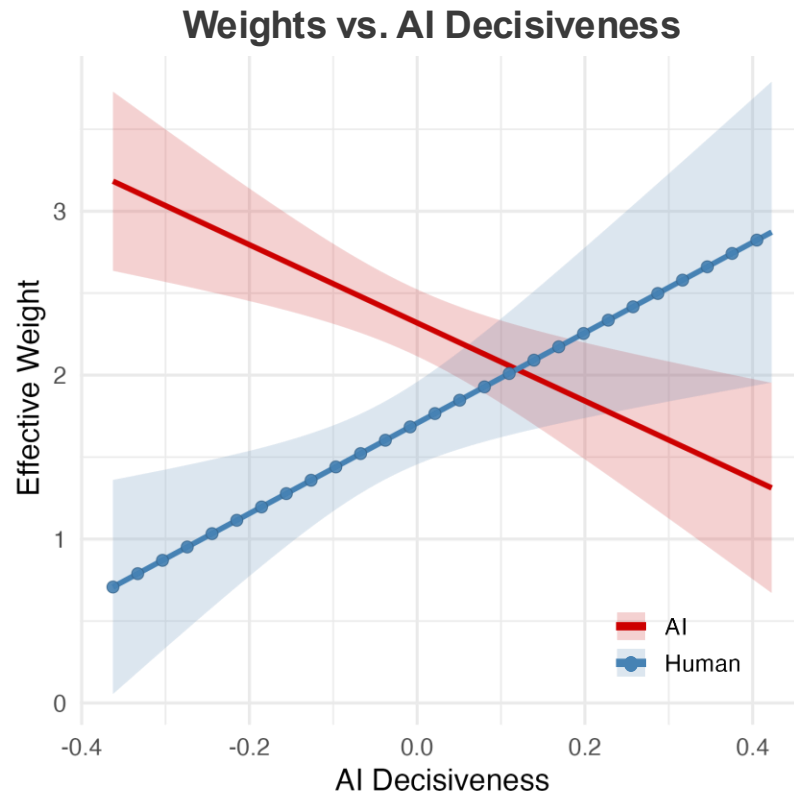
AI receives higher weight on average!

Explanations increase weighting on average.

Study 1: Decisiveness Subsets

	Indecisive AI	Decisive AI
Weight on Human ($\hat{\beta}_{\text{human}}$)	1.370*** (0.272)	2.498*** (0.517)
Weight on AI ($\hat{\beta}_{\text{AI}}$)	2.511*** (0.266)	1.902*** (0.296)
SD of Random Slope (Human)	0.006 (0.849)	0.005 (1.334)
SD of Random Slope (AI)	1.214** (0.402)	0.930* (0.432)
Observations	350	255

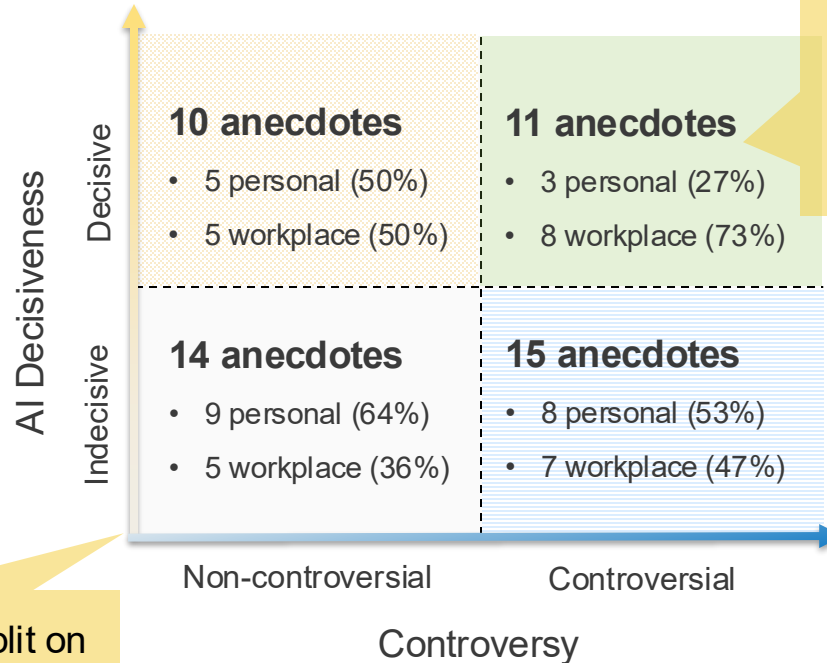
Note: Standard errors in parentheses.
*** $p < 0.001$, ** $p < 0.01$, * $p < 0.05$, $p < 0.1$.



Study 1: Controversy Subsets

	Non-Controversial	Controversial
Weight on Human ($\hat{\beta}_{\text{human}}$)	1.479*** (0.345)	1.596*** (0.322)
Weight on AI ($\hat{\beta}_{\text{AI}}$)	2.398*** (0.257)	2.216*** (0.260)
SD of Random Slope (Human)	0.021 (2.311)	0.005 (0.809)
SD of Random Slope (AI)	0.856. (0.483)	0.819. (0.451)
Log Likelihood	-273.29	-297.29
Observations	296	309
Note: Standard errors in parentheses. *** $p < 0.001$, ** $p < 0.01$, * $p < 0.05$, · $p < 0.1$.		

Study 1: Anecdote Categories by Decisiveness & Controversy



The AI is weighted **least** at the cross-section of Decisiveness and Controversy. This cross shows the highest proportion of **workplace** anecdotes!

Both cross-sections split on median Decisiveness and median Controversy.

Study 1 Additional Outcomes: Perceived Confidence & Helpfulness

AI output does increase confidence...

... but the inclusion of reasoning does not add significant additional benefit beyond the AI assessment.

However, reasoning does significantly contribute to perceived helpfulness of advice.

	(1) DV: Confidence	(2) DV: Confidence	(3) DV: Helpfulness
<i>Baseline</i>	<i>Control</i>	<i>T1: AI Assessment</i>	<i>T1: AI Assessment</i>
T: All Treatment (T1 + T2)	0.096* (0.048)		
T2: AI Assessment & Reasoning		-0.098 (0.067)	0.734*** (0.105)
Anecdote Controversy	0.017 (0.080)	0.028 (0.097)	0.001 (0.152)
AI Decisiveness	0.444*** (0.132)	0.593*** (0.163)	0.384 (0.255)
Constant	3.779*** (0.111)	3.863*** (0.126)	2.473*** (0.195)
Controls:			
Anecdote Text Length	✓	✓	✓
Anecdote Category	✓	✓	✓
Anecdote Sentiment	✓	✓	✓
Model	Mixed-effects OLS	Mixed-effects OLS	Mixed-effects OLS
Observations	950	605	605
Log Likelihood	-1,036.386	-650.152	-916.385
AIC	2,090.77	1,318.31	1,850.77
BIC	2,134.48	1,357.95	1,890.42
Note: Standard errors in parentheses. *** $p < 0.001$, ** $p < 0.01$, * $p < 0.05$, $p < 0.1$.			

Study 2: Appendix Slides

Study 2 Conditions: Dual Advisor Presentation

Am I wrong for telling my manager that a colleague is going to hand in their notice?

Okay so me and my colleague are pretty close. One late shift, she tells me in confidence that she plans to leave. Cool.

... [content cropped for brevity] ...

It got me thinking though. Should I have lied to my manager?
Am I in the wrong?

Who is at fault in this situation?

AI Assessment of Fault Probabilities:

Author: 67%

Other: 36%

AI Reasoning:

The author bears 67% of the responsibility in this scenario...

... [content cropped for brevity] ...

Crowd's Assessment of Fault Probabilities:

Author: 78%

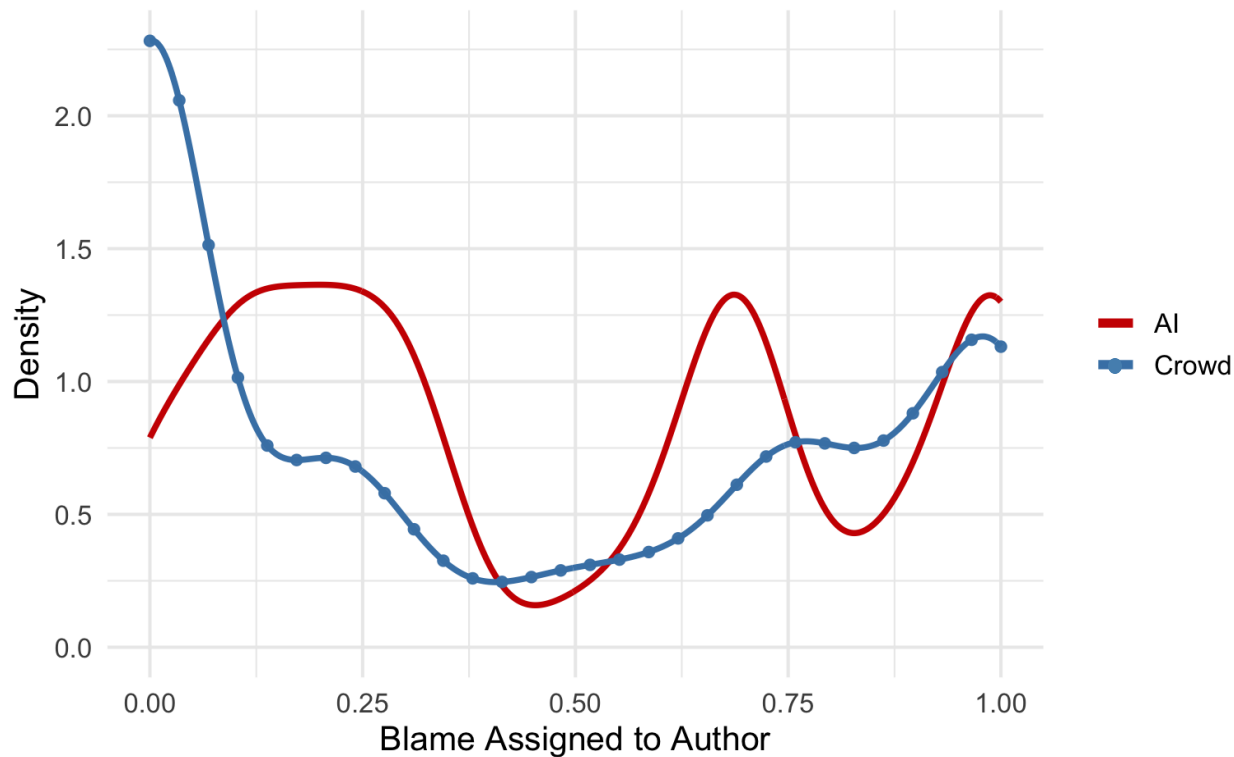
Other: 22%

Crowd's Reasoning:

The majority view holds that the author was wrong for revealing a colleague's confidential plans...

... [content cropped for brevity] ...

Study 2: Author Blame Density Plot



Study 2: Influential AI? The Double-edged Sword of Decisiveness

$$U_{ij}^{(k)} = W_{Human,i} \cdot p_{Human,j}^{(k)} + W_{AI,i} \cdot p_{AI,j}^{(k)} + \dots + \varepsilon_{ij}^{(k)}$$

$$\underbrace{W_{Human \times AI_D,i} \left(p_{Human,j}^{(k)} \cdot AI_Decisiveness_j \right) + W_{AI \times AI_D,i} \left(p_{AI,j}^{(k)} \cdot AI_Decisiveness_j \right)}_{\text{gavel icon}}$$

All Treatment	
Weight on Human ($\hat{\beta}_{human}$)	1.563*** (0.236)
Weight on AI ($\hat{\beta}_{AI}$)	2.352*** (0.204)
SD of Random Slope (Human)	0.005 (0.580)
SD of Random Slope (AI)	1.068*** (0.259)
Observations	605
Note: Standard errors in parentheses. *** $p < 0.001$, ** $p < 0.01$, * $p < 0.05$, $p < 0.1$.	



Indecisive AI

1.370***
(0.272)
2.511***
(0.266)
0.006
(0.849)
1.214**
(0.402)

350

Decisive AI

2.498***
(0.517)
1.902***
(0.296)
0.005
(1.334)
0.930*
(0.432)

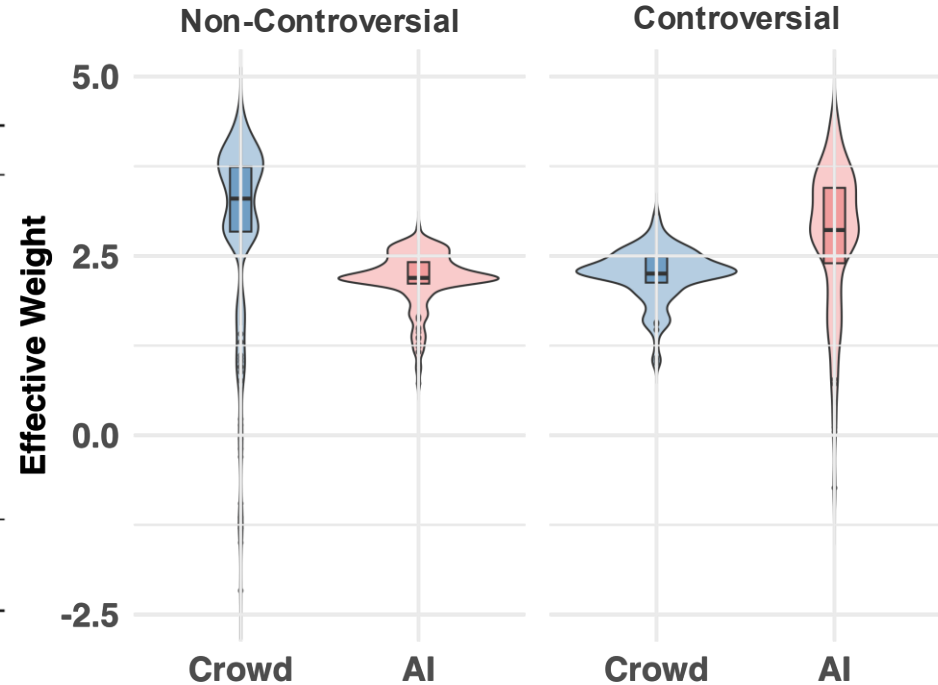
255

+ weight

Study 2: Advisor Preferences Across Controversy

	Non-Controversial	Controversial
Weight on Crowd ($\hat{\beta}_{\text{crowd}}$)	2.845*** (0.510)	2.245*** (0.367)
Weight on AI ($\hat{\beta}_{\text{AI}}$)	2.173*** (0.407)	2.758*** (0.432)
SD of Random Slope (Crowd)	2.390*** (0.551)	1.089 (0.556)
SD of Random Slope (AI)	1.176 (0.668)	1.981*** (0.494)
Observations	502	542

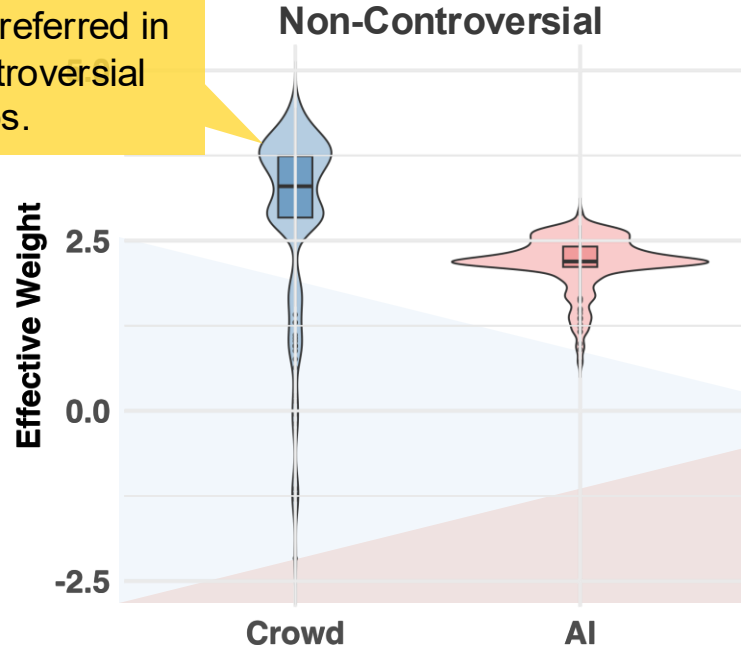
Note: Standard errors in parentheses. *** $p < 0.001$, ** $p < 0.01$, * $p < 0.05$, $\cdot p < 0.1$.



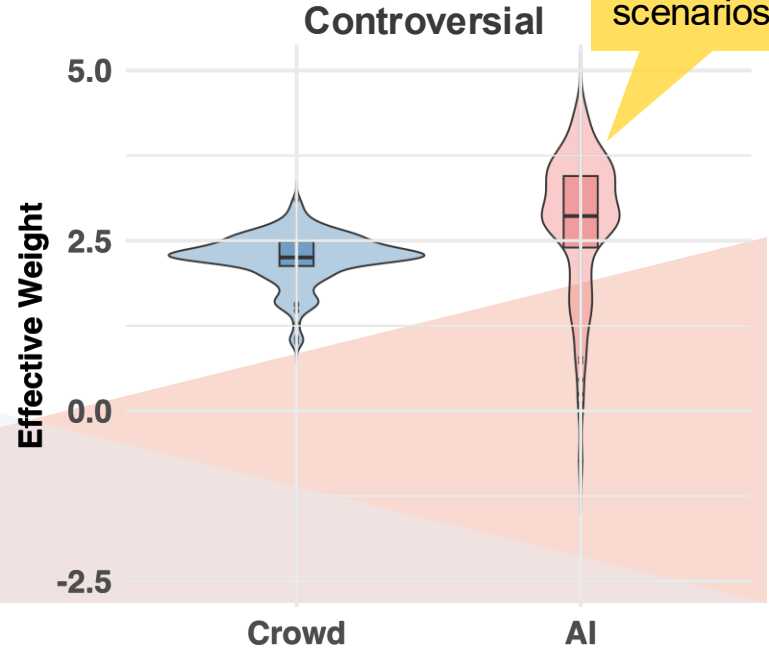
Advisor Preferences in the Face of Controversy

But how does advisor preference change across anecdote dimensions, like **controversy**?

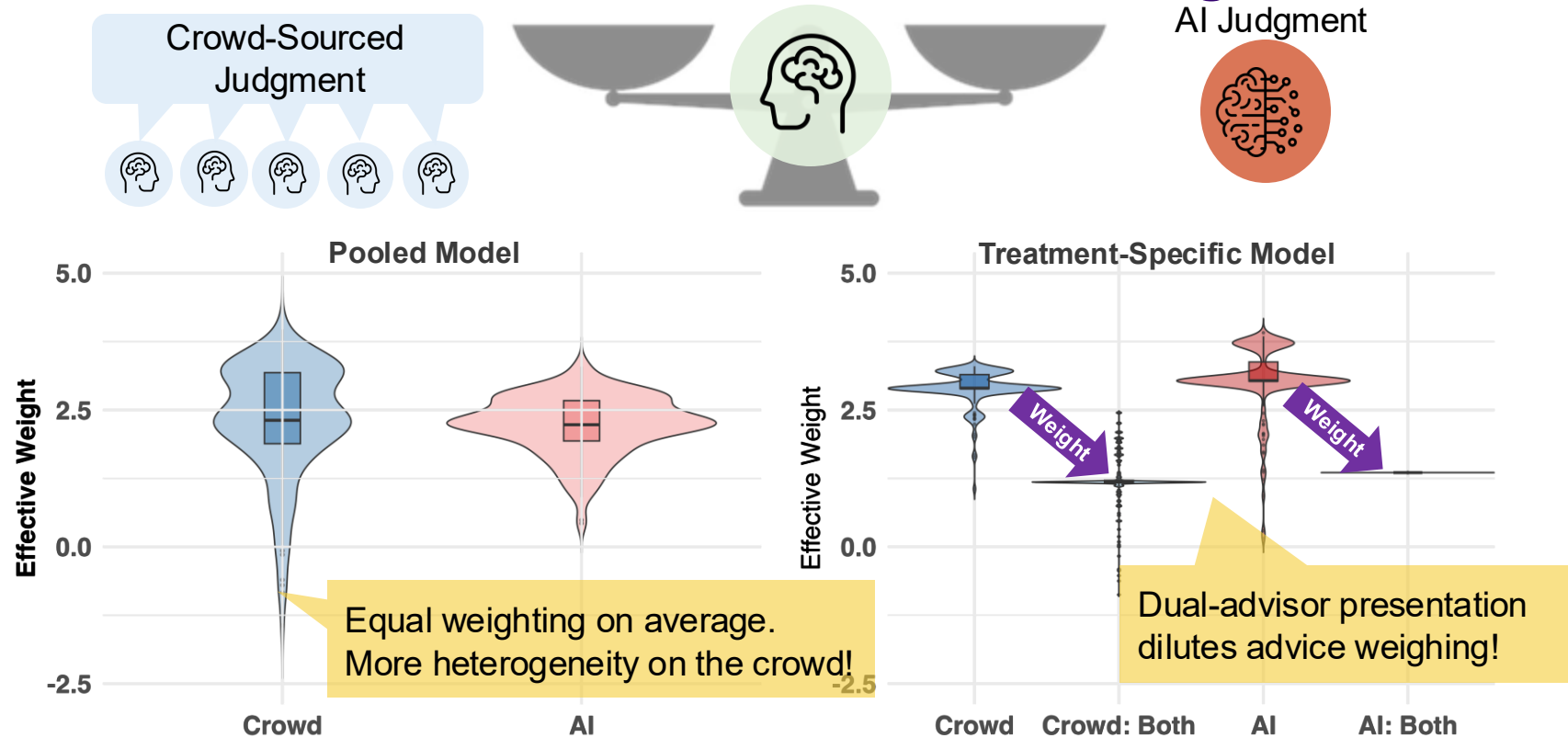
Crowd preferred in non-controversial scenarios.



AI preferred in controversial scenarios!

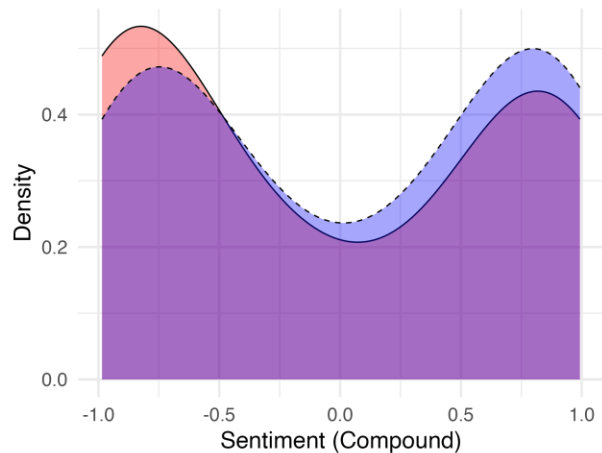


Revealed Advisor Preferences: Advice Weights



Study 2: Density of NLP Metrics by Advisor

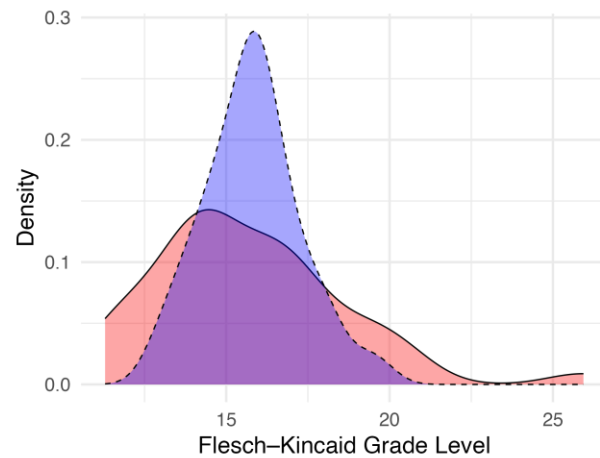
VADER Sentiment



Reading Ease



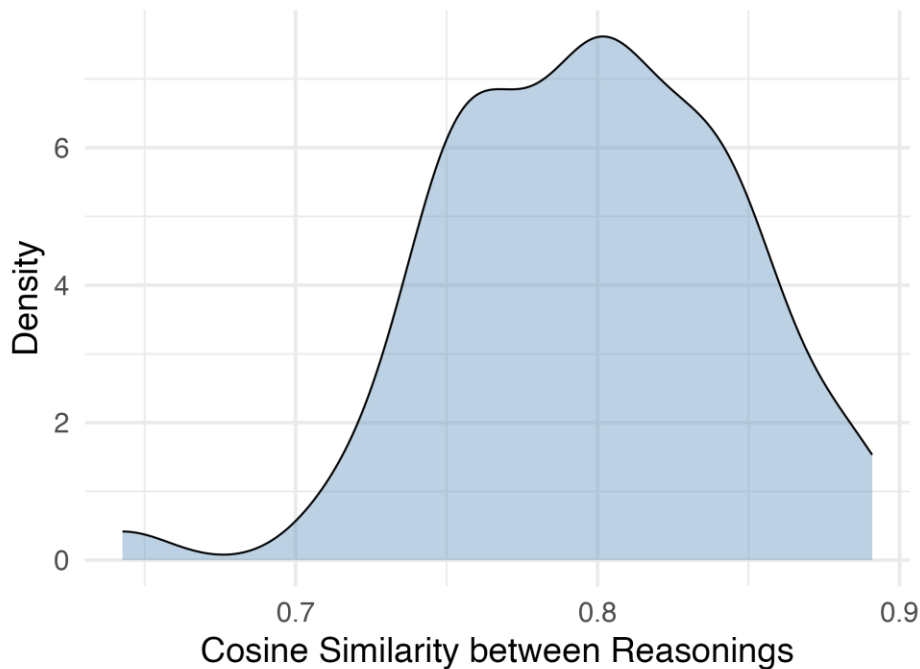
Writing Complexity



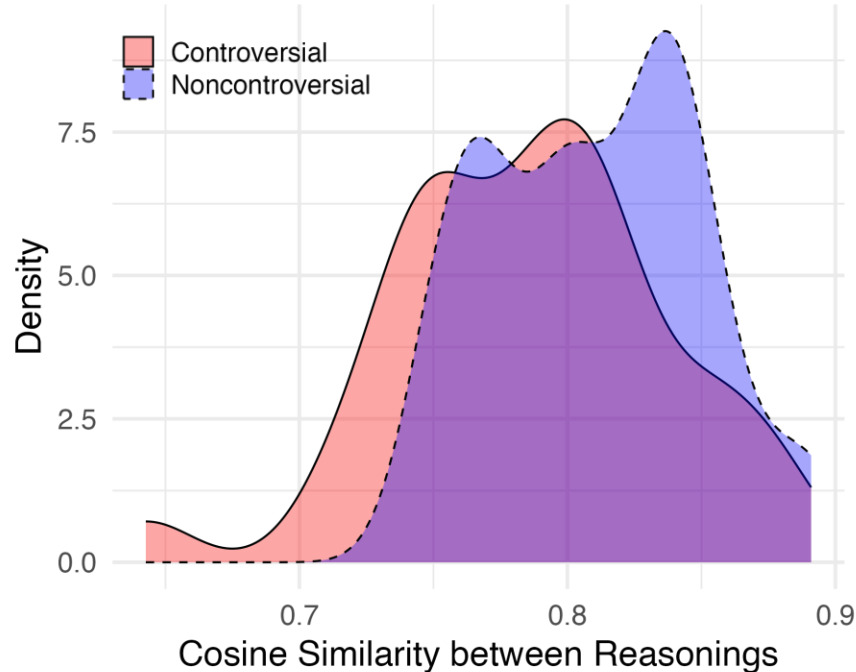
AI Crowd

Study 2: Density of Cosine Similarity of Reasonings

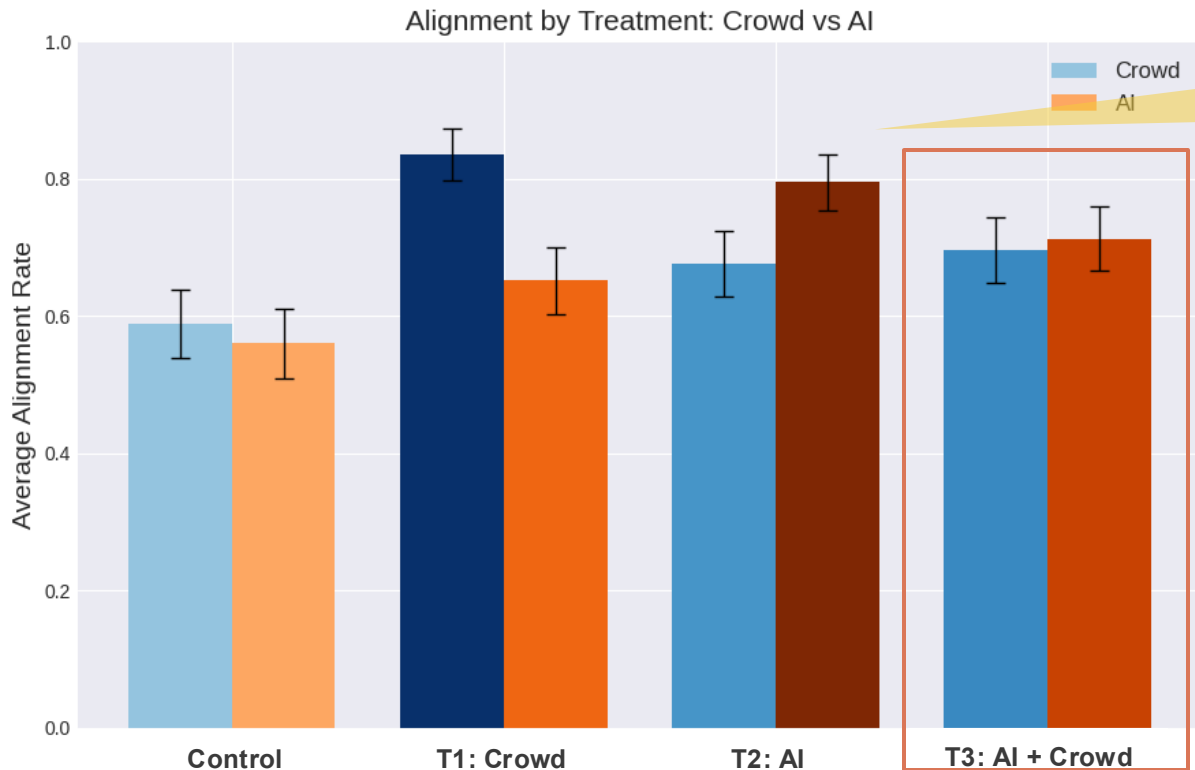
Aggregate



By Controversy



Study 2 Logit Results: Alignment with Crowd & AI

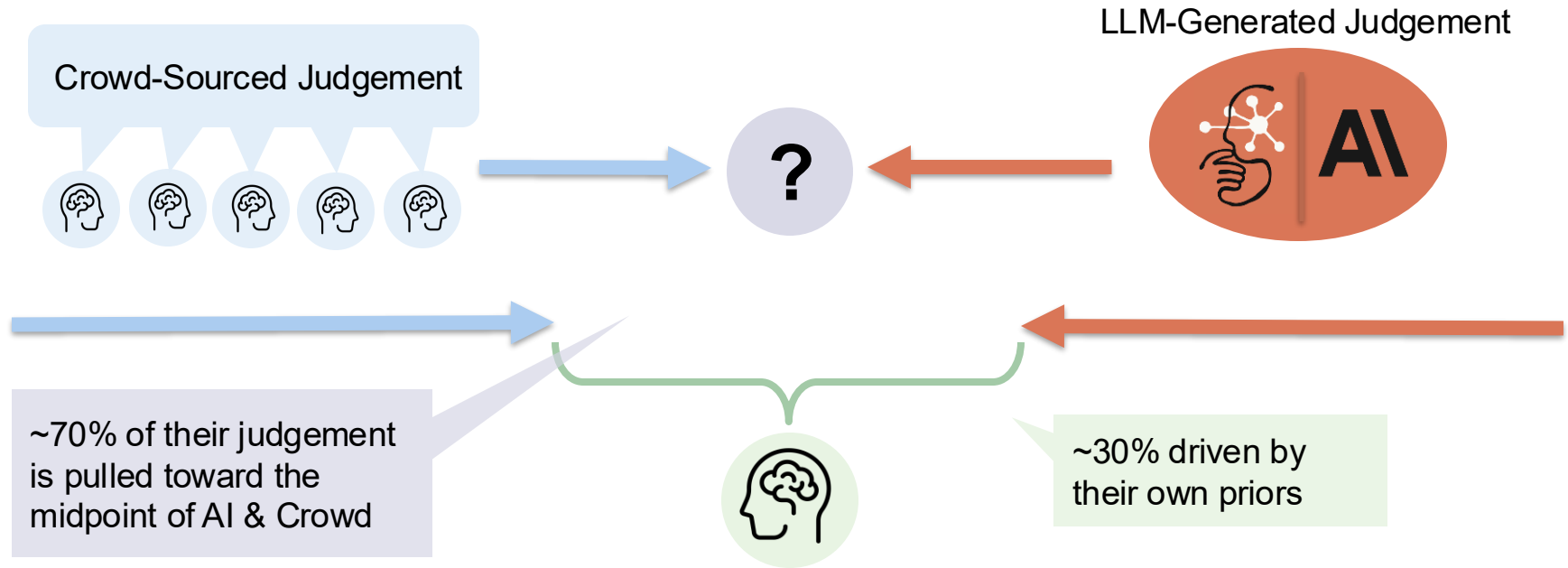


Treatments are strong predictors of alignment with the Crowd (T1 or T3) or the AI (T2 or T3).

Alignment is close when exposed to both sources!

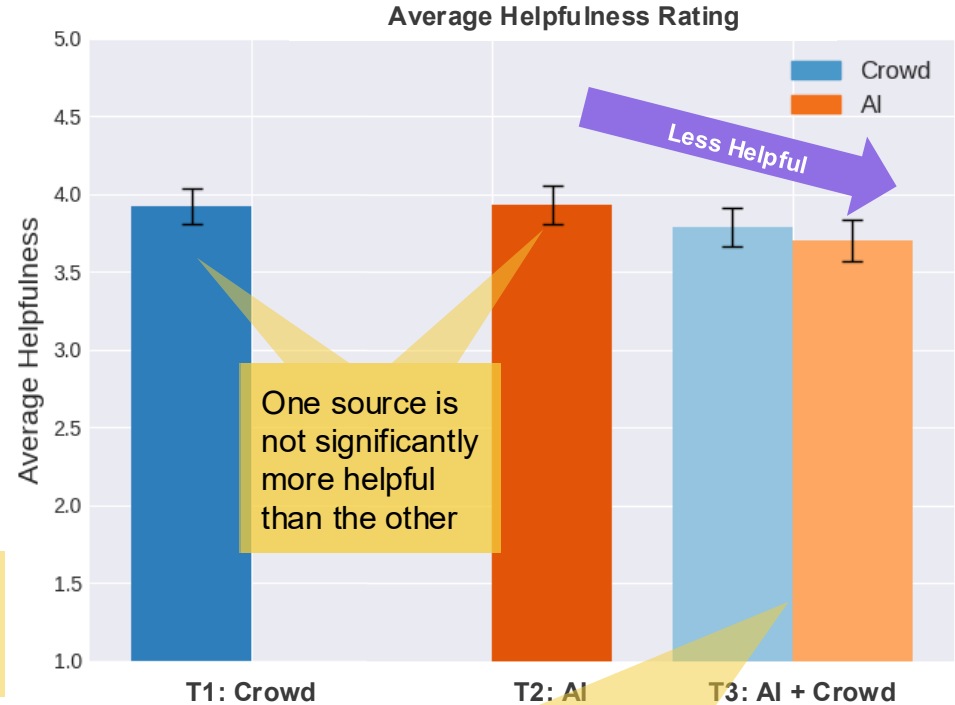
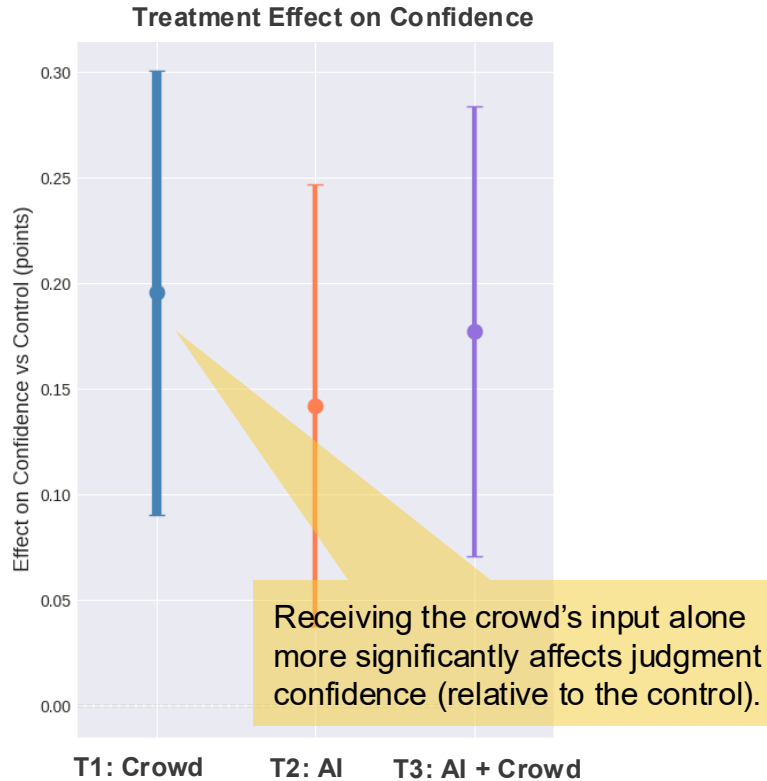
But how do participants incorporate the different sources?

Study 2: “Splitting the Difference”



Participants blend both sources, but they also incorporate their own idiosyncratic judgement. This suggests advice integration is structured but not mechanical averaging.

Study 2: Perceived Confidence and Helpfulness



When seen with the crowd advice, AI advice is perceived as significantly **less** helpful ($p < 0.01$). Crowd advice is not reliably lower ($p > 0.05$).

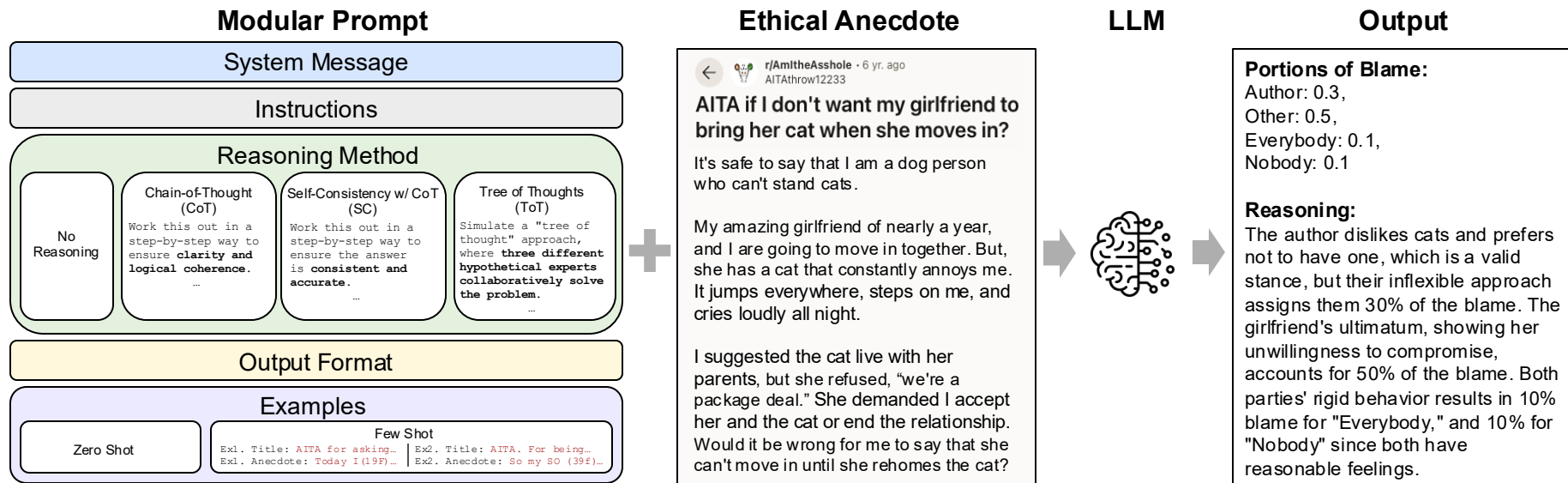
Study 2 Additional Outcomes: Helpfulness

	(1) T1 vs. T2: DV Helpfulness	(2) T3: DV Helpfulness
<i>Baseline</i>	<i>Advisor: AI</i>	<i>Advisor: AI</i>
Advisor: Crowd	-0.018 (0.119)	0.085 (0.063)
AI Decisiveness (shown)	0.260* (0.110)	0.246* (0.096)
Crowd Decisiveness (shown)	0.256* (0.104)	0.124 (0.092)
Constant	3.774*** (0.113)	3.498*** (0.148)
Controls:		
Anecdote Text Length	✓	✓
Anecdote Category	✓	✓
Anecdote Sentiment	✓	✓
Model	Mixed-effect OLS	Mixed-effect OLS
Observations	$N_{T1} + N_{T2} = 740$	$N_{T3} * 2 = 710$

Lack of significance →
one advisor is not more
helpful than the other

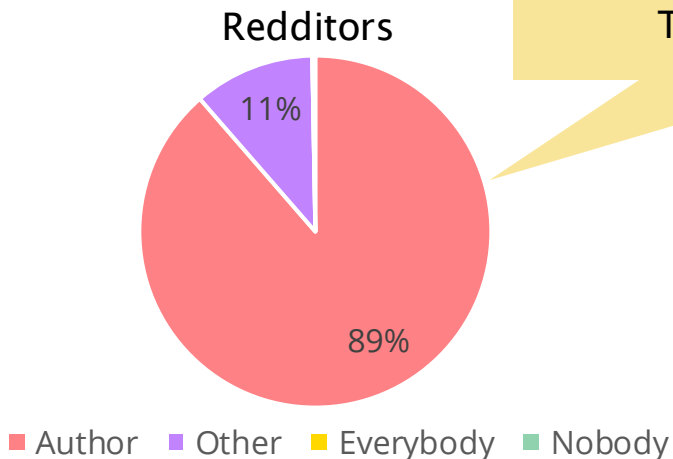
Notes: Standard errors in parentheses. *** $p < 0.001$, ** $p < 0.01$, * $p < 0.05$, $p < 0.1$.
Decisiveness terms are interacted with visibility masks; effects apply when that source is shown.
In T3 both sources are visible (“shown”), so the participant makes two helpfulness ratings per anecdote judgment.

Prompting Framework



What of the Cat-Hating Boyfriend?

Am I in the wrong if I don't want my girlfriend
to bring her cat when she moves in?

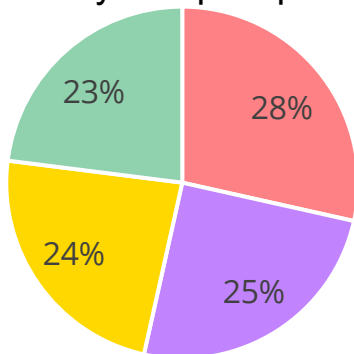


Controversy (normalized entropy) = 0.27
Total community votes: 1059

Normalized Entropy Calculation

Example of **high entropy (0.99)** → **highly controversial**
(no clear consensus)

Am I wrong for expecting my mom
to want to support me through my
first delivery and post partum?



■ Author ■ Other ■ Everybody ■ Nobody

Entropy formula:

$$H = - \sum_{\{i\}}^n p_i \ln(p_i)$$

Steps:

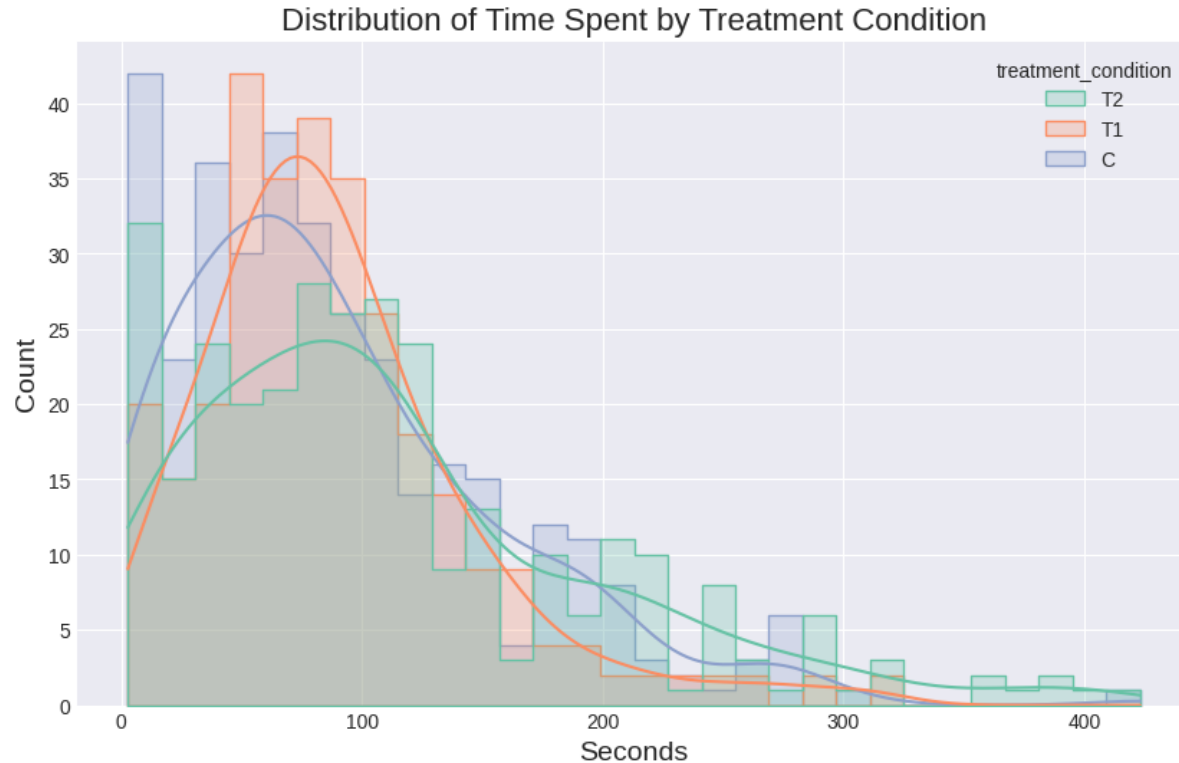
1. Calculate blame portions using votes: 0.28, 0.25, 0.24, 0.23
2. Substitute blame portions into formula:

$$H = -[0.28 \ln(0.28) + 0.25 \ln(0.25) + 0.24 \ln(0.24) + 0.23 \ln(0.23)]$$
$$H = \sim 1.385$$

3. Normalize using min and max entropy across all anecdotes

$$H = \sim 1.385 \rightarrow \text{Normalized Entropy} = 0.99$$

Study 1: Distributions of Time Spent per Anecdote



Study 2: Time Spent per Anecdote by Treatment

