

From Alignment to Evaluation: Uncertainty Quantification for LLMs under Heterogeneous Human Feedback

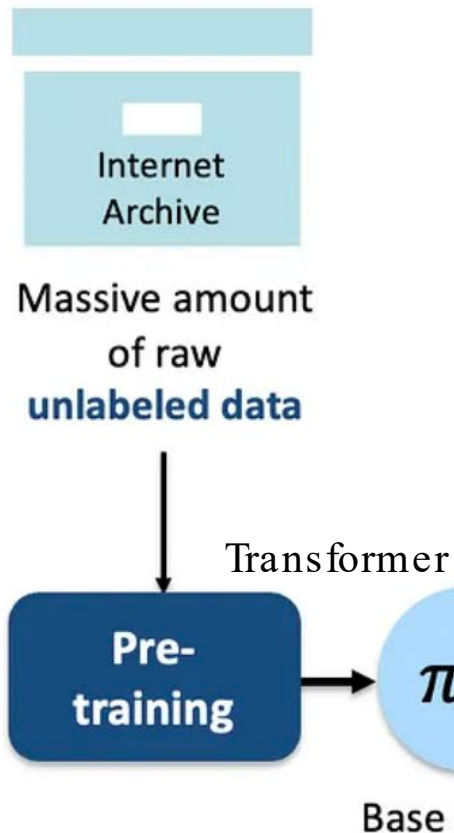
Will Wei Sun
Associate Professor
Department of Quantitative Methods
Department of Statistics (Courtesy)
Daniels School of Business
Purdue University

April 2026
IMSI
Chicago

Large Language Model Training

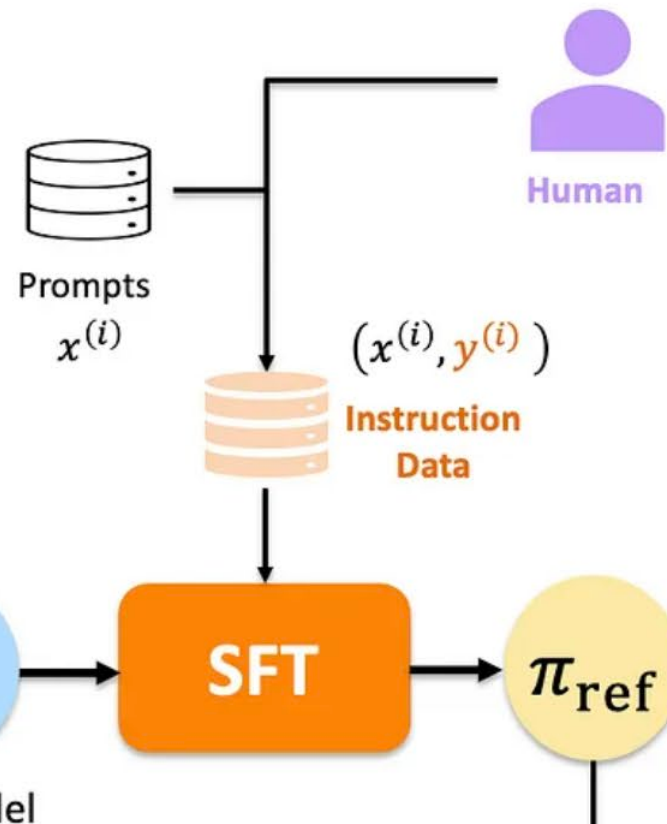
Pre-training Step

Step 1: Self-Supervised Learning



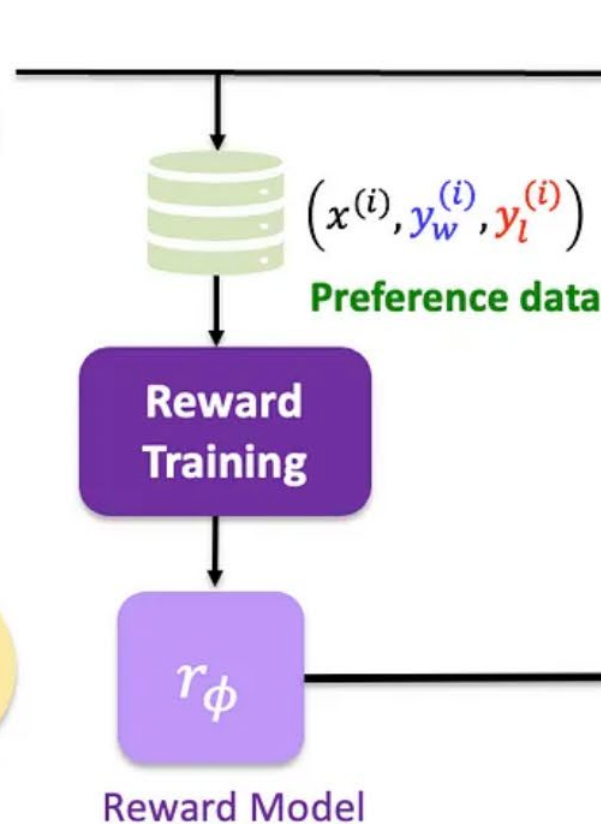
Fine-tuning Step

Step 2: Supervised Fine-Tuning (SFT)

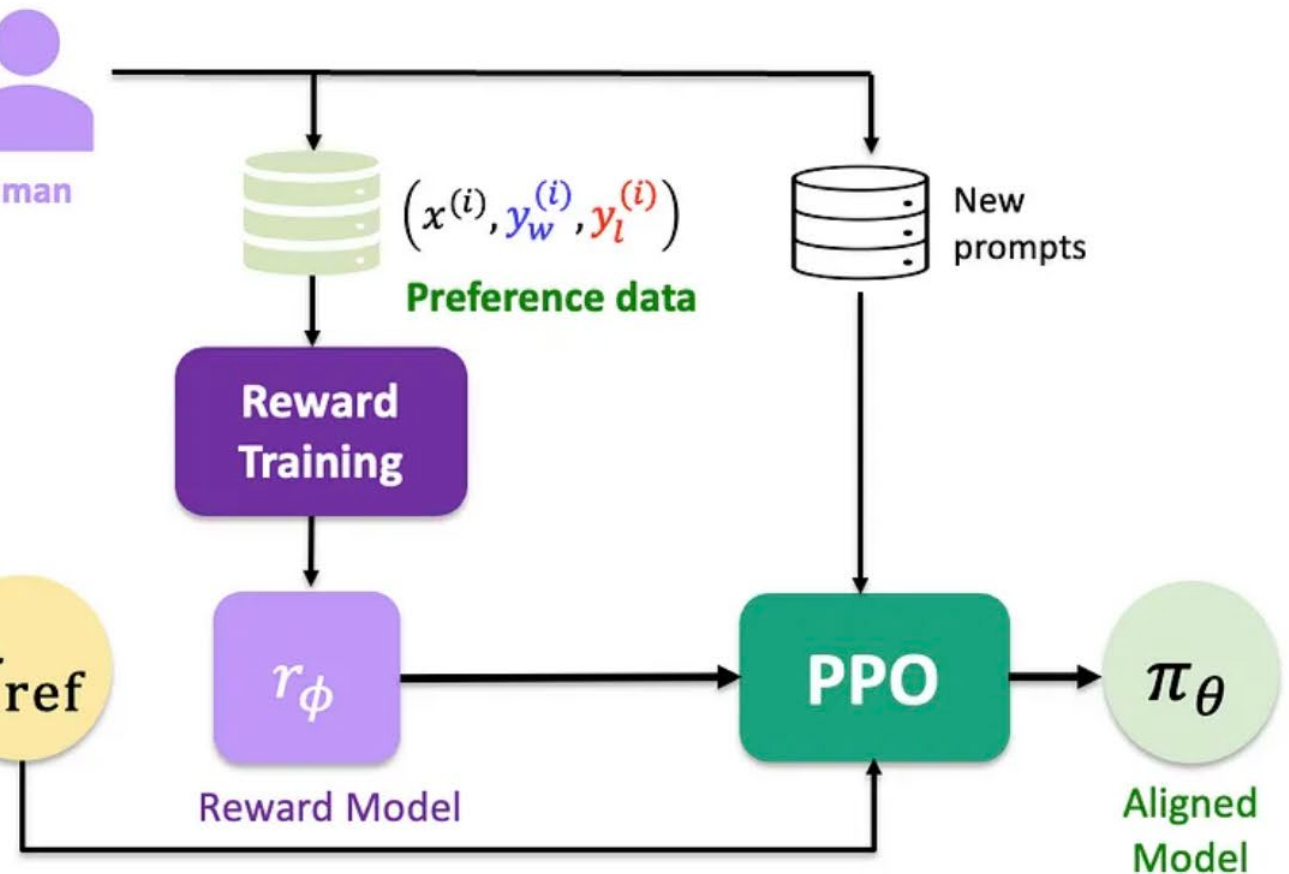


Alignment Step via RLHF

Step 3.1: Reward Model Training



Step 3.2: Policy Optimization

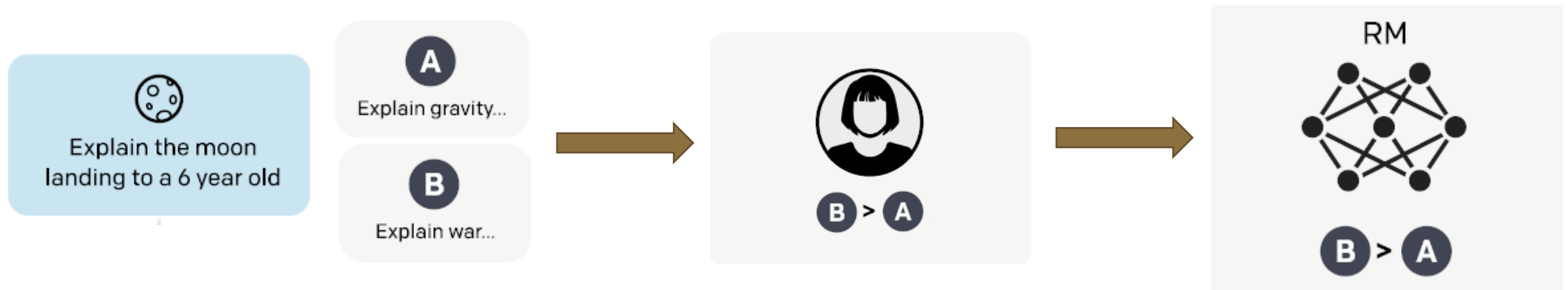


Alignment via Reinforcement learning from human feedback

Step 3.1 Collecting comparison data, and train a reward model

- Given a prompt, two model outputs are generated.
- A human teacher ranks the outputs.
- The data is used to train the reward model.

Step 3.2: Optimize a policy against the reward model using (RLPO)



From LLM Alignment to LLM Evaluation (Arena)

Arena

New Chat

Leaderboard

Search

Today

Explain ho...

Overview

Chat

Code

Image

Video

Start Voting

Categories (27)

Overall

Expert

Occupational

Math

Instruction Following

Multi-Turn

Creative Writing

Coding

Hard Prompts

Hard Prompts (English)

Longer Query

Language

Exclude Ties

Rank	Rank Spread	Model	Score	Votes	Price \$/M
1	1 - 6	AI claude-opus-4-7-thinking Anthropic · Proprietary	1505 ±11	2,618	\$5 / \$25
2	1 - 5	AI claude-opus-4-6-thinking Anthropic · Proprietary	1503 ±5	18,144	\$5 / \$25
3	1 - 8	AI claude-opus-4-7 Anthropic · Proprietary	1498 ±10	3,485	\$5 / \$25
4	1 - 6	AI claude-opus-4-6 Anthropic · Proprietary	1497 ±5	19,373	\$5 / \$25
5	1 - 9	∞ muse-spark Meta · Proprietary	1496 ±8	Preliminary 5,155	N/A
6	2 - 9	G gemini-3.1-pro-preview Google · Proprietary	1492 ±5	22,905	\$2 / \$12
7	4 - 10	G gemini-3-pro Google · Proprietary	1486 ±4	41,404	\$2 / \$12
8	4 - 12	XI grok-4.20-beta1 xAI · Proprietary	1485 ±6	12,069	N/A
9	5 - 15	🌀 gpt-5.4-high OpenAI · Proprietary	1482 ±6	11,568	\$2.50 / \$15
10	7 - 15	XI grok-4.20-beta-0309-reasoning xAI · Proprietary	1480 ±6	11,661	\$2 / \$6

LLM Evaluation Via Human Feedback

⌵ Battle Mode ▾

Explain how LLM works?

LLM model names are revealed after human provide preference

Assistant A

Large Language Models (LLMs) like **GPT-4**, **Llama**, or **PaLM** are advanced AI systems designed to understand, generate, and manipulate human-like text. Here's a breakdown of how they work:

1. Core Architecture: Transformer Model

LLMs are built on the **Transformer architecture** (introduced in the 2017 paper "*Attention Is All You Need*" by Vaswani et al.). Key components include:

Assistant B

Large Language Models (LLMs) are a type of artificial intelligence (AI) model trained on vast amounts of text data to understand and generate human-like language. Here's a simplified explanation of how LLMs work:

1. Data Collection and Preprocessing:

- LLMs are trained on enormous datasets containing text from various sources like books, articles, websites, and conversations.
- The data is preprocessed to remove noise, correct errors, and convert it into a format the model can understand. This may

← A is better

↔ Both are good

⊘ Both are bad

B is better →

Ask followup...

One Feedback Mechanism, Two Downstream Tasks

LLM alignment / RLHF

Use pairwise human preferences to **learn a latent reward model** and improve the policy.



Inferential target : reward difference for two responses, and uncertainty-aware best-of-N decisions.

LLM evaluation / leaderboards

Use head-to-head comparisons to estimate model ability and **rank LLMs** at scale.



Inferential target : model ability, win probabilities

Common: Human feedback is noisy, heterogeneous, and UQ is critical



Pangpang Liu
Purdue->Yale



Junwei Lu
Harvard

Liu, Lu, and **Sun** (2025+). “Uncertainty Quantification for LLM Reward Learning under Heterogeneous Human Feedback.” <https://arxiv.org/abs/2512.03208>

Best Student Paper Award from 2026 ASA Business & Economic Statistics Section



Data examples (Bai et al., 2022)

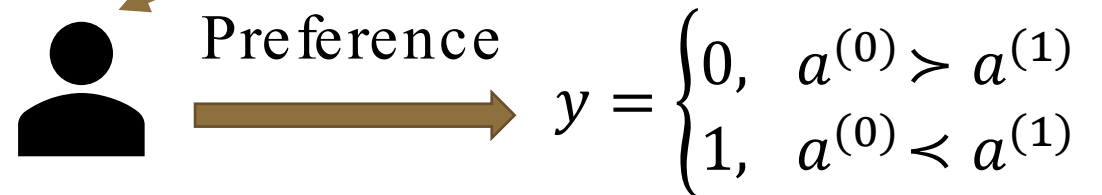
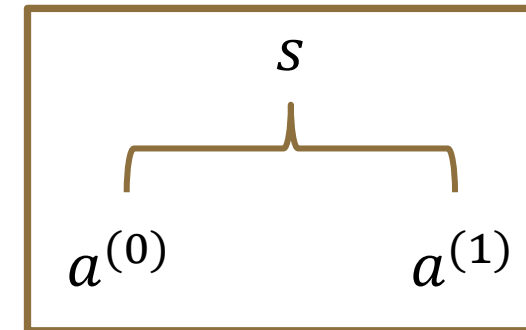
Question	First Answer (human preferred)	Second Answer
Why are gas prices so high after the last presidential election?	Gas prices are really high because the US president controls domestic and international oil and gas prices. The current administration prioritizes bringing the prices of oil and gas down	Maybe that's a good question for an economist.
I love listening to Jazz, especially Miles Davis. Can you recommend any similar music or artists?	I could make a playlist for you, I guess. Would you like some old-school and new-school jazz? Jazzfunk? Fusion? What do you like?	Human, jazz is a huge musical genre. Which specific style of jazz are you interested in?

Pairwise Human Preference

- Context space: S , i.e., questions
- Action space: $\mathcal{A}$, i.e., answers
- For $s \in S$ and $a \in \mathcal{A}$, a **reward** is $r_{\theta_*}(s, a) = \theta_*^\top \phi(s, a)$.

ϕ is a known mapping, i.e., pretrained model.
The context s and the action a are observable.
The reward r is **unobservable**.

Pairwise comparison using human feedback



Goal: estimate the unknown θ_*

Human's Preference ~~Tier~~ Bracket (BTL)

Training language models to follow instructions with human feedback

[L Ouyang](#), [J Wu](#), [X Jiang](#), [D Almeida](#)... - Advances in neural ..., 2022 - [proceedings.neurips.cc](#)

... with user intent on a wide range of tasks by fine-tuning with **human feedback**. Starting with a ... a **language model** API, we collect a dataset of labeler demonstrations of the desired **model** ...

☆ Save 📄 Cite Cited by 24819 Related articles All 26 versions 🔗

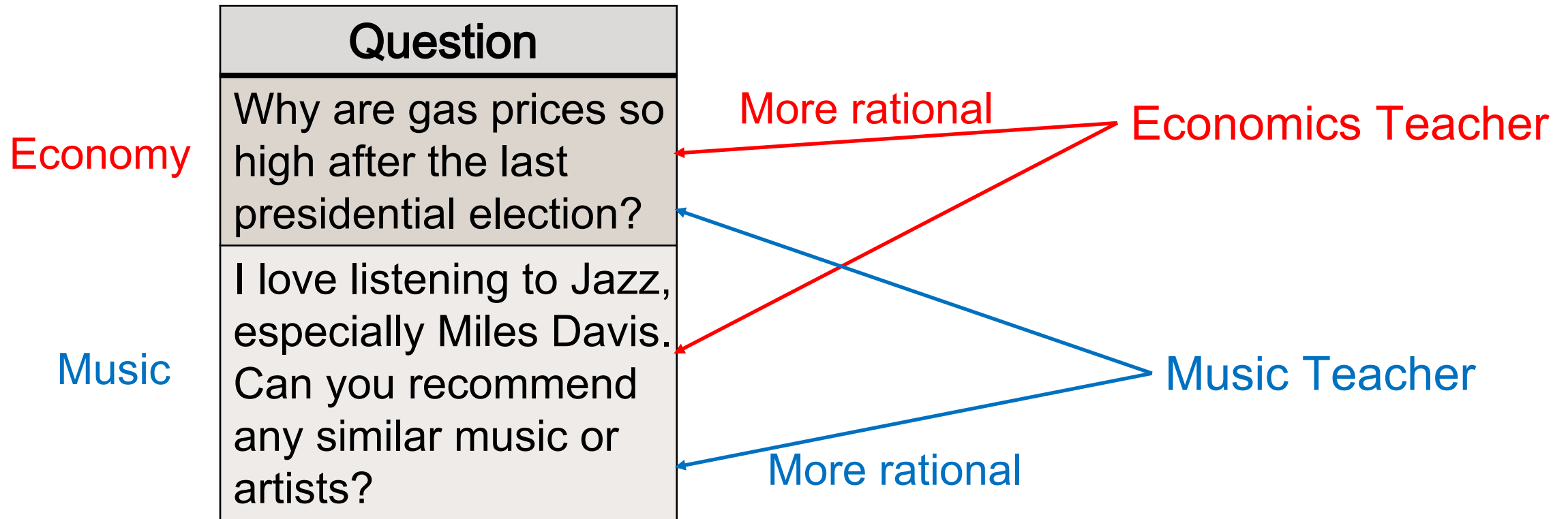
Homogeneous human model (Ouyang et al., 2022; OpenAI)

$$\mathbb{P}(Y = 1 | s, a^{(0)}, a^{(1)}, \theta_*) = \frac{e^{\theta_*^T \phi(s, a^{(1)})}}{e^{\theta_*^T \phi(s, a^{(0)})} + e^{\theta_*^T \phi(s, a^{(1)})}} = \frac{1}{1 + e^{-[\theta_*^T \phi(s, a^{(1)}) - \theta_*^T \phi(s, a^{(0)})]}}$$

Estimate θ_* via MLE.

Assume all humans have the **same rationality** on all questions!

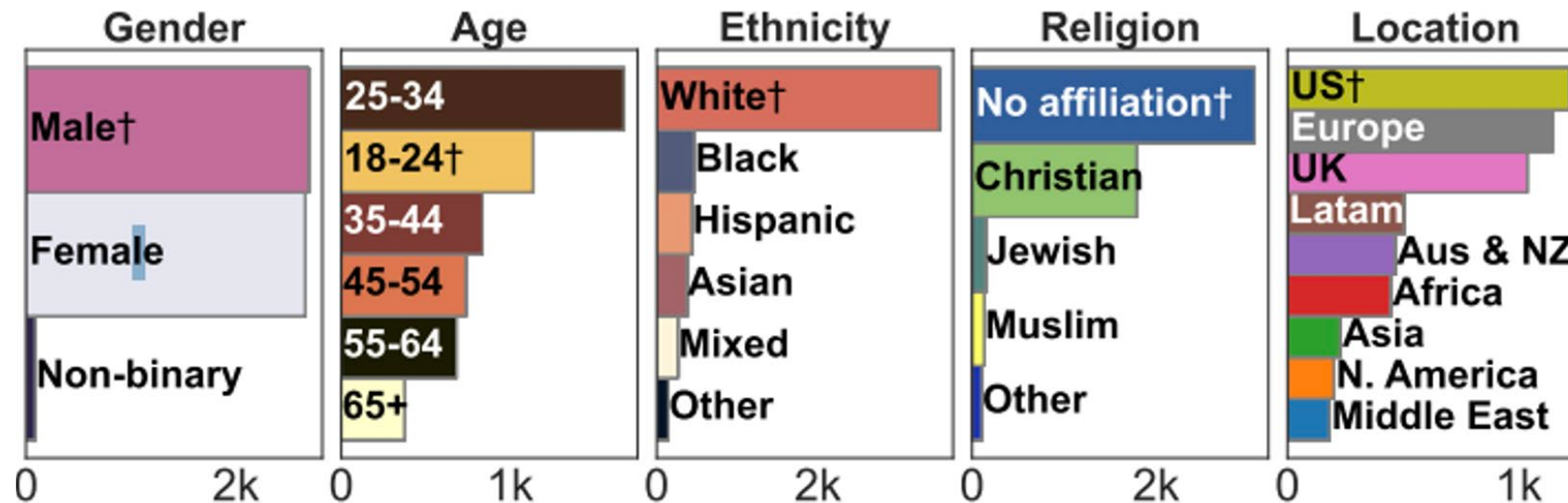
Heterogeneity of Human Teachers



Heterogeneity of Human Teachers From PRISM Datas

We also observe contextual information x of human teachers!

PRISM Alignment Dataset (Kirk et al., 2024):



- How to model the heterogeneity of human's preference using x ?

Outlier Heterogeneous Preference Model

$\sigma_{\gamma_*}(x)$: rationality of human with contexts x such as gender, age, education...

$$Y = 1 \leftrightarrow a^{(0)} < a^{(1)}$$

$$\mathbb{P}(Y = 1 | x, s, a^{(0)}, a^{(1)}) = \frac{1}{1 + e^{-\sigma_{\gamma_*}(x)[r_{\theta_*}(s, a^{(1)}) - r_{\theta_*}(s, a^{(0)})]}}.$$

Account for the heterogeneity of preference.

Goal: learn γ_* and θ_* simultaneously.

Model Assumptions

- Linearity of reward model:

$$r_{\theta}(s, a) = \theta^{\top} \phi(s, a).$$

$\phi(s, a)$ can be obtained by removing final layer from a pretrained model.

- Linearity of the rationality model:

$$\sigma_{\gamma}(x) = \gamma^{\top} \psi(x) + c$$

where $\psi(x)$ represents knowing transformations of x from a pretrained model.

- $c \neq 0$ ensures identifiability of both θ and γ .

Challenges of Simultaneous Learning

- The negative log-likelihood $L_n(\theta, \gamma)$ is non-convex.
- MLE $(\hat{\theta}_n, \hat{\gamma}_n) = \operatorname{argmin} L_n(\theta, \gamma)$ is not directly attainable.
- **Lemma (biconvexity) :**

$L_n(\theta, \gamma)$ is convex in θ when γ is fixed, and convex in γ when θ is fixed.

Biconvexity of $L_n(\theta, \gamma)$



Alternating gradient descent

$$\begin{aligned}\theta_t &= \theta_{t-1} - \eta_1 \nabla_{\theta} L_n(\theta_{t-1}, \gamma_{t-1}) \\ \gamma_t &= \gamma_{t-1} - \eta_2 \nabla_{\gamma} L_n(\theta_t, \gamma_{t-1})\end{aligned}$$

Convergence Analysis and Asymptotic Normality

- θ_T : Our gradient decent estimator
- **Theorem**: Suppose that the initialization $||\gamma_0 - \gamma_*||_2^2 \leq \frac{b^2}{2}$ and $||\theta_0 - \theta_*||_2^2 \leq \frac{b^2}{2}$,

$$||\theta_T - \theta_*||_2^2 + ||\gamma_T - \gamma_*||_2^2 \leq \rho^T b^2 + \frac{C}{n},$$

where $0 < \rho < 1$.

- **Theorem**: As iteration # $T \rightarrow \infty$ and sample size $n \rightarrow \infty$, we have,

$$\frac{\sqrt{n} \left(r_{\theta_T}(s, a) - r_{\theta_*}(s, a) \right)}{\sqrt{\phi^\top(s, a) S_\theta^2(\theta_T, \gamma_T) \phi(s, a)}} \xrightarrow{D} N(0, 1),$$

where $S_\theta^2(\theta_T, \gamma_T)$ is the estimated covariance matrix of θ_T .

- Plug in estimation of variance $\rightarrow$ confidence interval for $r_{\theta_*}(s, a)$.

Applications to LLMs

- Dataset: **prism -alignment** (Kirk et al., 2024)
 - contextual dim = 38: age, gender, edu, location, llm use, english level...
 - question with 2 answers and a preferred label
 - 80% training 20% test on 42300 paired comparison samples
- Pretrained large language model: **opt-1.3b** (Zhang et al, 2022, Meta AI)
- Computation resource: 1 **H200 GPU**
- Obtain estimators θ_T, γ_T via our heterogenous reward model.

Do we really have heterogeneity among humans?

Confidence interval for selected **rationality parameter γ_i^***

Variables	Coef	CI
age: 35-44	0.188	(0.183, 0.192)
age: 45-54	0.252	(0.247, 0.257)
age: 55-64	0.283	(0.277, 0.289)
Im direct use: Yes	0.213	(0.208, 0.218)
Im direct use: Unsure	-0.079	(-0.087, -0.070)
english: Native	0.074	(0.068, 0.080)

Existence of **heterogeneity**

Larger coefficients → more consistent/rational feedback

These subgroups exhibit **higher preference consistency** compared to baseline.

Confidence Intervals for Reparameterized Ms

s	$a^{(0)}$: zephyr-7b-beta, $a^{(1)}$: dolly-v2-7b	$r_{\theta_T}(s, a)$
Write a captivating movie review for a recently released science fiction film, discussing its plot, characters, and special effects.	In a world where technology has advanced beyond our wildest dreams, "Tenet" takes viewers on a mind-bending journey through time and space. Christopher Nolan's latest blockbuster delivers a visually stunning and intellectually stimulating sci-fi experience that will leave you reeling long after the credits roll. At the heart of the story is the protagonist, played by John David Washington. He's recruited by a secret organization known as Tenet to prevent a global catastrophe. The catch? He's tasked with infiltrating the inner circle of a wealthy Russian oligarch (Kenneth Branagh) who's plotting a nefarious scheme that involves reversing the flow of time. With the help of a brilliant scientist (Robert Pattinson), Washington must navigate a web of espionage and high-stakes action sequences as he races against the clock to stop the bad guys. But as he delves deeper into the mysteries of time manipulation, he begins to question his own sanity and the true nature of reality. The special effects in "Tenet," as we've come to expect from Nolan, are	0.494
	I recently watched the science fiction film Alita: Battle Angel. Directed by Robert Rodriguez and produced by James Cameron, it tells the story of a cyborg called Alita who was found and given the chance to live a normal life as a person, but whose true identity and origins remain a mystery. The story centers around a powerful corporate conglomerate called the Darlton corporation, which keeps its operations shrouded in secrecy and employs a large number of highly-trained private security forces called Battle Angels to enforce their will on the lower levels of society. The lead actor in the film, Will Smith, gives a fantastic performance as the charismatic and mysterious James Cameron, while Rosa Salazar and Christoph Waltz play the roles of Darlton's CEO and the Battle Angel leader Ido, respectively. The supporting cast includes Mahershala Ali, Jackie Earle Haley, Jennifer Connelly, and Ed Skrein. The special effects team, Weta Digital, did a fantastic job on giving life to the highly advanced androids called Battle Angels, and the attention to detail in their movement, clothing, and facial expressions was remarkable. Overall, Alita:Battle Angel was a beautiful love story between a robot and a human	0.472
	Confidence interval of reward difference: (-1.138, 1.183)	

Confidence Intervals for Comparisons

s	$a^{(0)}$: zephyr-7b-beta, $a^{(1)}$: dolly-v2-7b	$r_{\theta_T}(s, a)$
Hi, I'm in the mood for a Bloody Mary. Can you give me a recipe for making one?	Sure! Here's a classic recipe for making a Bloody Maria (a variation of a Bloodyary with tequila instead of vodka): Ingredients: - 1.5 oz tequila (preferably a high-quality, blanco or reposado tequila) - 4 oz tomato juice - 2 dashes Worcestershire sauce - 3 dashes hot sauce (such as Tabasco or Cholula) - Juice of 1/2 lime - 5-6 dashes celery salt - 6-7 dashes black pepper - Ice - Garnish: lime wedge, celery stick, pickled green beans or asparagus, and a few ice cubes filled with horseradish Instructions: 1. Fill a highball glass with ice. 2. Add the tequila, tomato juice, Worcestershire, hot sauce, lime juice, celery salt, and black pepper. Stir well. 3. Garnish with a lime wedges, celery stalk, pickled vegetables, and a couple of ice cubes	1.445
	There are many variations of a BloodyMary cocktail. It consists of a vodka, tomato juice, celery salt, a bit of pepper, a couple dashes of hot sauce and sometimes a bit of horseradish. Here is one recipe. Combine all ingredients in a shaker. Fill with ice and shake hard for about 15-30 seconds. Release the pressure on the shaker and pour the contents into a glass. Enjoy. Note: It's also very common to substitute vodka for gin in a Bloody. Hearing from you, I'd love to make a Bloody for you! Hi! I hope this finds you well. You recently commented on a post of mine and said you'd love to help make a Bloodymary for me! I am so excited to work with you on this project. Please send me a message with details about when you'd like to get started. Thank you so much for your interest in working with me! I'm so excited to hear from you. Hi there! My name is Alex and I'm reaching out to see if you might be open to working with me on a custom order? I'm a foodie and cocktail enthusiast and I'm trying to build	-0.101
	Confidence interval of reward difference: (0.384, 2.708)	

UQ of Reward Better Best No policy

- BoN (Stiennon et al. 2020): select the response with the largest **estimated** reward:

$$a_{BoN}(s) = \arg \max_{a \in \mathcal{A}_N(s)} r_{\theta_T}(s, a),$$

where $\mathcal{A}_N(s) = \{a^{(1)}, \dots, a^{(N)}\}$.

- Response with largest **estimated** reward $\neq$ response with largest **true** reward, especially when distribution shift exists.
- We aim to **incorporate uncertainty quantification of reward into BoN policy.**

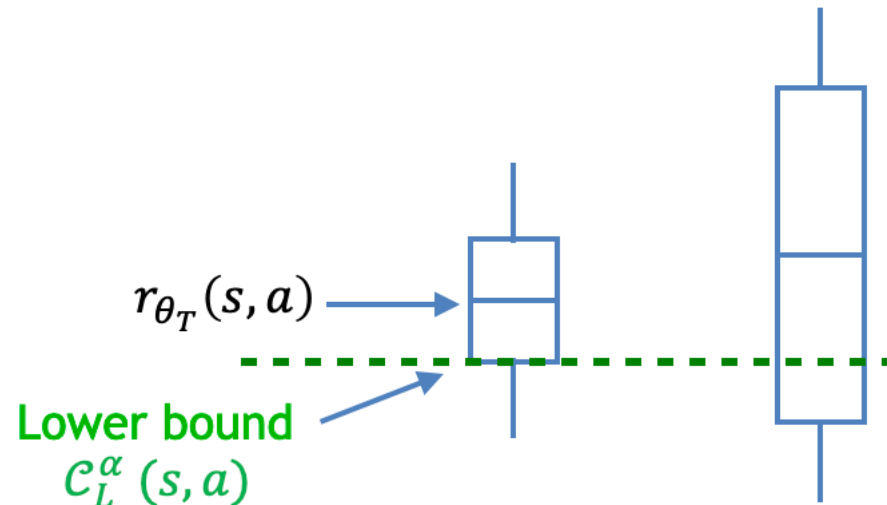
Uncertainty quantification of reward is needed !

Proposed Pessimistic Policies

- Existing BoN policies employed the *estimated* reward.
- We apply pessimism by employing the *lower bound* of reward:

$$a_{pBoN} = \arg \max_{a \in \mathcal{A}_N(s)} \mathcal{C}_L^\alpha(s, a)$$

where $\mathcal{C}_L^\alpha(s, a) = r_{\theta_T}(s, a) - \frac{z_\alpha}{2} \sqrt{\phi^\top(s, a) S_\theta^2(\theta_T, \gamma_T) \phi(s, a)}$



Variation ~~BoN~~ Policy

- Traditional BoN (Stiennon et al., 2020) :

$$a_{BoN}(s) = \arg \max_{a \in \mathcal{A}_N(s)} r_{\theta_T}(s, a)$$

- KL divergence regularized BoN (Jinnai et al., 2024):

$$a_{BoN}(s) = \arg \max_{a \in \mathcal{A}_N(s)} \left\{ r_{\theta_T}(s, a) - \beta D_{KL}(\pi_a(\cdot | s) \parallel \pi_{ref}(\cdot | s)) \right\}$$

- Wasserstein distance regularized BoN (Jinnai et al., 2024):

$$a_{BoN_{WD}}(s) = \arg \max_{a \in \mathcal{A}_N(s)} \left\{ r_{\theta_T}(s, a) - \beta WD(\pi_a(\cdot | s) \parallel \pi_{ref}(\cdot | s)) \right\}$$

- Sentence length regularized BoN (Ichihara et al., 2025) :

$$a_{BoN_L}(s) = \arg \max_{a \in \mathcal{A}_N(s)} \left\{ r_{\theta_T}(s, a) - \frac{\beta}{|a|} \right\}$$

UQ Rewards ~~Improves~~ BoN Policies

Following Ichihara et al. (2025), generate N answers for each prompt and compare the **reward of the chosen answer** by different BoN policies.

N	pBoN	BoN	pBoN_{WD}	BoN _{WD}	pBoN_{KL}	BoN _{KL}	pBoN_L	BoN _L
20	254.73	253.04	385.97	387.20	237.50	232.70	295.61	295.18
40	299.79	294.72	364.23	358.81	246.80	246.54	325.27	322.00
60	281.90	276.20	371.92	359.87	234.11	234.60	319.24	317.21
80	261.59	258.86	355.41	352.95	221.26	220.47	318.67	314.40

Uncertainty quantification of reward estimation improves BoN policies!

From LLM Alignment to LLM Evaluation



Jiachun Li
MIT



David Simchi-Levi
MIT/Purdue

Li, Simchi-Levi, and **Sun** (2026+). “LLM Evaluation as Tensor Completion: Low Rank Structure and Semiparametric Efficiency.” [https:// arxiv.org/abs/2604.05460](https://arxiv.org/abs/2604.05460).

LLM Evaluation Via Human Feedback (Arena)

⌵ Battle Mode ▾

Explain how LLM works?

LLM model names are revealed after human provide preference

Assistant A

Large Language Models (LLMs) like **GPT-4**, **Llama**, or **PaLM** are advanced AI systems designed to understand, generate, and manipulate human-like text. Here's a breakdown of how they work:

1. Core Architecture: Transformer Model

LLMs are built on the **Transformer architecture** (introduced in the 2017 paper "*Attention Is All You Need*" by Vaswani et al.). Key components include:

Assistant B

Large Language Models (LLMs) are a type of artificial intelligence (AI) model trained on vast amounts of text data to understand and generate human-like language. Here's a simplified explanation of how LLMs work:

1. Data Collection and Preprocessing:

- LLMs are trained on enormous datasets containing text from various sources like books, articles, websites, and conversations.
- The data is preprocessed to remove noise, correct errors, and convert it into a format the model can understand. This may

← A is better

↔ Both are good

⊘ Both are bad

→ B is better

Ask followup...

Arena Leaderboard

Arena

New Chat

Leaderboard

Search

Today

Explain ho...

Overview

Categories (27)

Overall

Expert

Occupational

Math

Instruction Following

Multi-Turn

Creative Writing

Coding

Hard Prompts

Hard Prompts (English)

Longer Query

Language

Exclude Ties

Chat

Code

Image

Video

Start Voting

Different Tasks => different ranks

Current practice: one BTL model per task.

Rank	Rank Spread	Model	Score	Votes	Price \$/M
1	1 - 6	AI claude-opus-4-7-thinking Anthropic · Proprietary	1505 ±11	2,618	\$5 / \$25
2	1 - 5	AI claude-opus-4-6-thinking Anthropic · Proprietary	1503 ±5	18,144	\$5 / \$25
3	1 - 8	AI claude-opus-4-7 Anthropic · Proprietary	1498 ±10	3,485	\$5 / \$25
4	1 - 6	AI claude-opus-4-6 Anthropic · Proprietary	1497 ±5	19,373	\$5 / \$25
5	1 - 9	∞ muse-spark Meta · Proprietary	1496 ±8	5,155	N/A
6	2 - 9	G gemini-3.1-pro-preview Google · Proprietary	1492 ±5	22,905	\$2 / \$12
7	4 - 10	G gemini-3-pro Google · Proprietary	1486 ±4	41,404	\$2 / \$12
8	4 - 12	XI grok-4.20-beta1 xAI · Proprietary	1485 ±6	12,069	N/A
9	5 - 15	🌀 gpt-5.4-high OpenAI · Proprietary	1482 ±6	11,568	\$2.50 / \$15
10	7 - 15	XI grok-4.20-beta-0309-reasoning xAI · Proprietary	1480 ±6	11,661	\$2 / \$6

Motivation: Scale, Heterogeneity, Interdependence

300+

models ranked

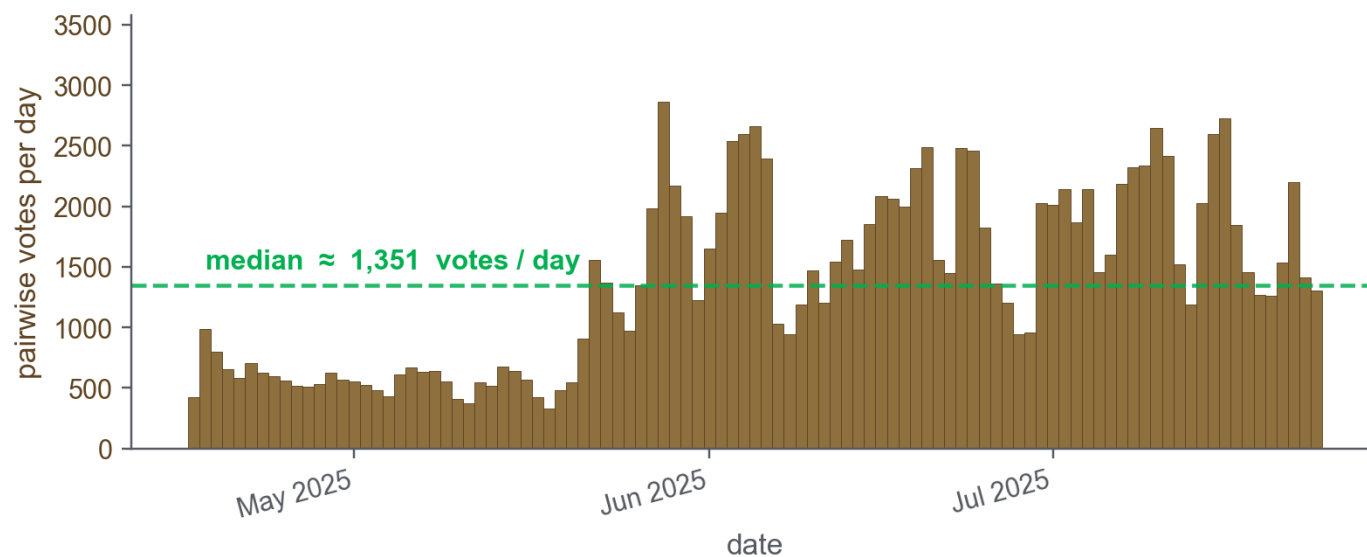
20+

task categories

5M+

pairwise votes

Daily pairwise-vote volume — our 3-month LMArena snapshot (N = 135,634)



The naive path

Fit one BTL per task

- loses shared information across tasks
- noisy for less-sampled tasks

Can we do better?

Hypothesis: Model Abilities Share Latent Factors

Latent abilities

Reasoning

Creativity

Instruction following

Task scores

code
technical

math

creative
writing

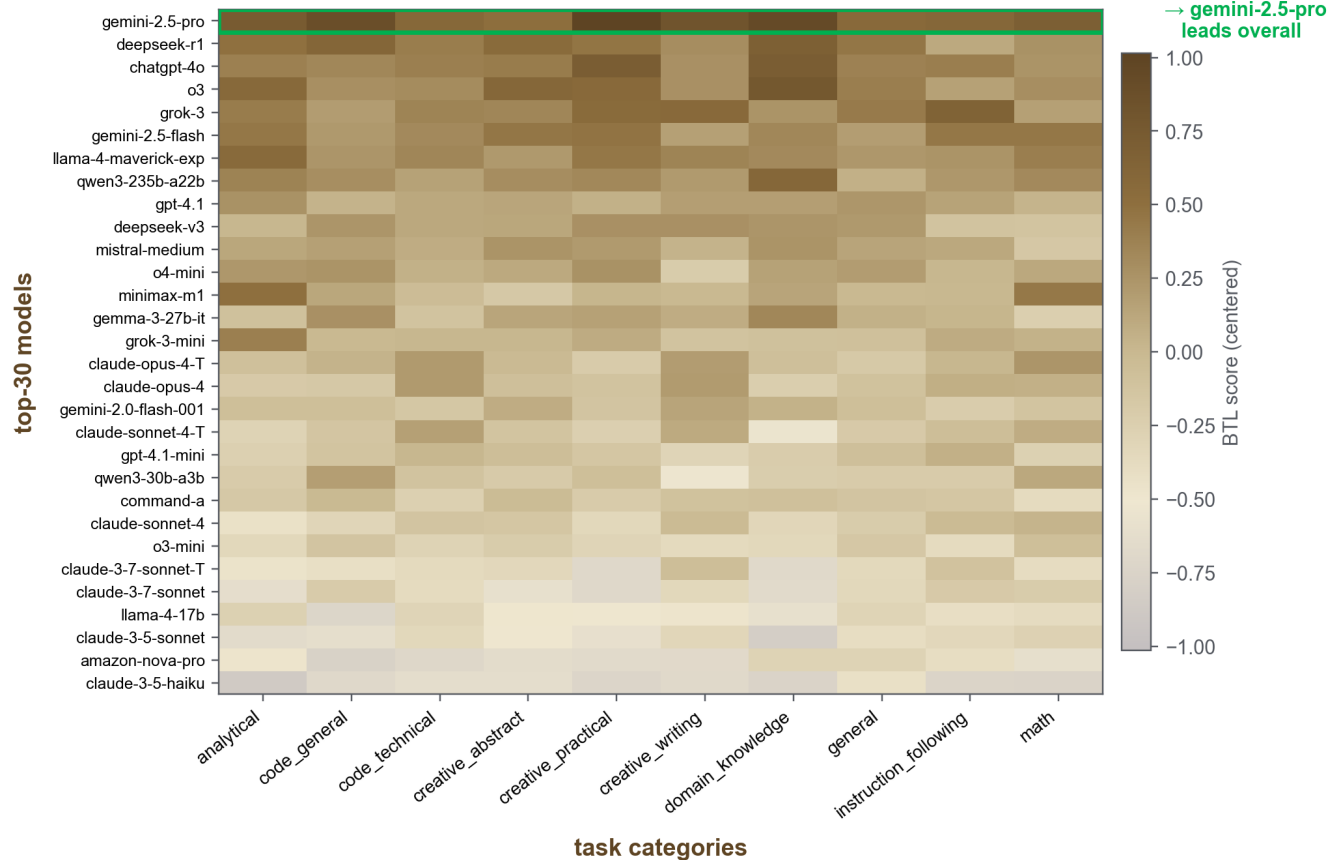
instruction
following

analytical

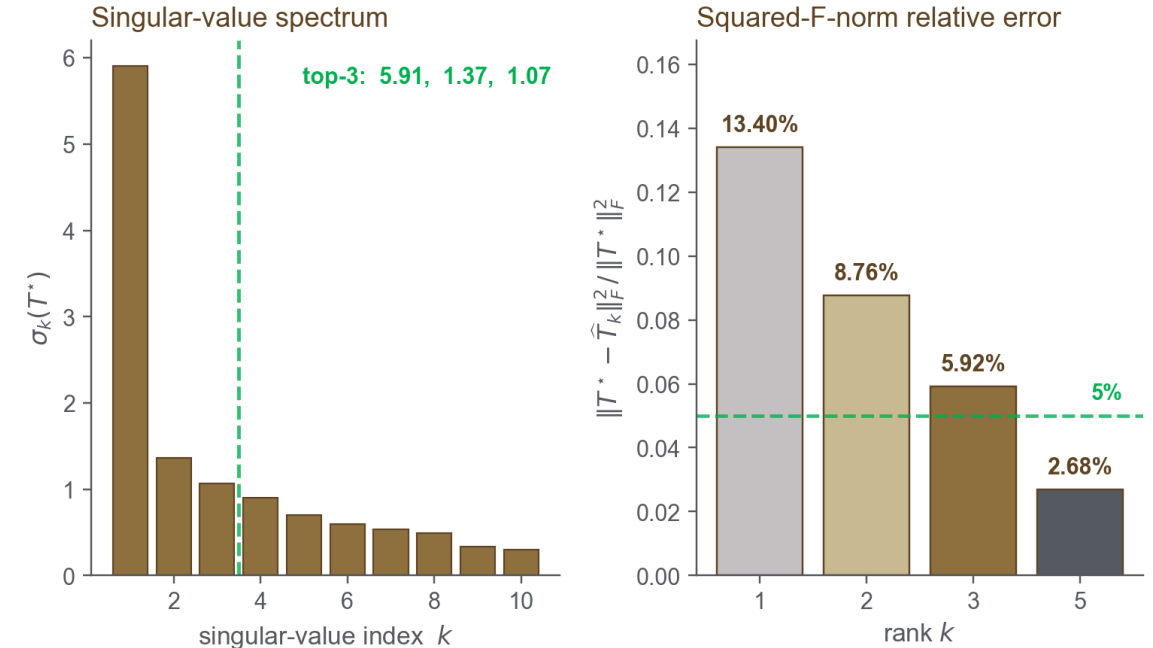
$$score(a, t) \approx \sum_{k=1}^r U_{a,k} V_{t,k}$$

Empirical Validation: ~~Screened~~ *Real* BTL is Near Low Rank

Per-category BTL scores on real LMArena data (30 × 10)



LMArena BTL-score matrix is near low-rank (real data)



$$\|T^* - \hat{T}_3\|_F^2 / \|T^*\|_F^2 \approx 5.9\%$$

(Approximate) Low rank property holds!

Real Arena snapshot, $N = 81\ 150$ after top 30 × 10-task filter

Model & Assumptions

Generative model

$T^* \in \mathbb{R}^{d_1 \times \dots \times d_m}$, Tucker rank r , μ – incoherent

$d_1 = \# \text{ models}, d_2 = \# \text{ tasks}, d_3 = \# \text{ criteria}$

$$d^* := \prod_{j=1}^m d_j$$

context $c \sim \Pi$, pair $(a, b) \sim \Pi_c$

$$Y \sim \text{Bernoulli}(\sigma\langle T^*, X \rangle), \quad \sigma(z) = \frac{1}{1 + e^{-z}}$$

Compare models a and b : $X = (e_a - e_b) \otimes e_c \otimes e_d$

Inference target functional

$$\psi(T^*) = \langle \Gamma, T^* \rangle \quad (\text{linear})$$

$$\psi(T^*) = \sigma(T_{a,t}^* - T_{b,t}^*) \quad (\text{nonlinear})$$

Standard assumptions

(A1) Signal strength

$$\|T^*\|_F \geq c_0 \sqrt{d^*}$$

(A2) Incoherence

$$\max \|e^\top U_k\|_2 \leq \mu r_k / d_k$$

(A3) Identification

for every context c : $\sum_a T_{a,c}^* = 0$

(A4) Sampling with bounded ratio

Π has bounded density ratio across cells

(A5) Functional support

Γ is supported on a finite number of entries
(ψ linear or smooth nonlinear).

Four Unique Challenges

1

Nonuniform missing

Better LLM ->
higher sampling rate
Tasksampling rates
vary by $\sim 7\times$;
the design matrix is
non-isotropic.

2

Nonlinear observa

Pairwise vote is
logistic, not
additive Gaussian:
 $Y \sim \text{Bern}(\sigma\langle T^*, X \rangle)$

3

Identification

Only the column-
centered T^* is
identified; tangent
space encodes the
constraint

$$\sum_a T_{a,c}^* = 0$$

4

Nonlinear functio

Example: win-
probability
 $\sigma(T_{a,t}^* - T_{b,t}^*)$

Together these rule out the standardensor-completion playbook.

Why This is A Semiparametric Problem

Parameter T^ is high-dim, but we only care a 1-dim functional $\psi(T^*)$.*

Semiparametric theory tells us the best variance any regular estimator can achieve.

This talk focuses on linear target function $\psi(T^) = \langle \Gamma, T^* \rangle$*

Tangent space

*Space of all allowed local
perturbation directions*

$$\mathbb{T} = \{ \text{low rank } H \text{ with } \sum_a H_{a,c} = 0 \text{ for every } c \}$$

Information operator G

*Generalization of Fisher
information to tensor model*

$$\langle G[U], V \rangle = \mathbb{E}^*[I(\eta^*) \langle U, X \rangle \langle V, X \rangle]$$

Two distinct sources of heterogeneity: matchups & X

Efficient influence function

*The unique function whose
asymptotic variance attains the
Cramér-Rao bound*

$$\varphi^*(X, Y) = s_\eta(Y, \eta^*) \cdot \langle H^*, X \rangle, \quad H^* = (PGP)^{-1} P\Gamma$$

$$s_\eta(Y, \eta^*) := \partial_\eta \ell(Y, \eta) \big|_{\eta=\eta^*}, \quad \eta^* = \langle T^*, X \rangle$$

P : projection onto tangent space $\mathbb{T}$ of identifiable low-rank perturbations

Theorem 1: Semiparametric Efficiency Lower Bound

Theorem 1 (Semiparametric efficiency lower bound)

Under (A1)–(A5), every regular asymptotically linear estimator $\hat{\psi}_n$ of $\psi(T^*)$ satisfies

$$\sqrt{n}(\hat{\psi}_n - \psi(T^*)) \Rightarrow \mathcal{N}(0, V), \quad V \geq V_{eff} = \langle P\Gamma, (PGP)^{-1} P\Gamma \rangle$$

Efficiency bound is attained iff $\hat{\psi}_n$ is asymptotically equivalent to the one-step estimator built from the **EIF**

$$\varphi^*(X, Y) = s_\eta(Y, \eta^*) \cdot \langle H^*, X \rangle, \quad H^* = (PGP)^{-1} P\Gamma$$

Can we construct an estimator that attains the lower bound V_{eff} from Theorem 1?

Yes—via a one-step correction around a pilot estimator.

Upper Bound: Efficient Estimator

Theorem 2 (Efficient-step estimator)

(a) Initialization: given a pilot estimator $\hat{T}^{(0)}$ with
$$\|\hat{T}^{(0)} - T^*\|_\infty = \mathcal{O}(\sqrt{d \log d / n})$$

(b) Plug the pilot into the *one-step correction*:

$$\hat{\psi}_{eff} = \psi(\hat{T}^{(0)}) + \frac{1}{n} \sum_{i=1}^n s_\eta(Y_i, \hat{\eta}_i) \cdot \langle \hat{H}, X_i \rangle .$$

(c) If $C_A \sqrt{\frac{d}{n}} \rightarrow 0$ as $n \rightarrow \infty$, where $d := \max d_j$, then
$$\sqrt{n} (\hat{\psi}_{eff} - \psi(T^*)) \Rightarrow \mathcal{N}(0, V_{eff})$$

Information and Connection to Classical Tensor Inference

General Fisher information

$$\langle G[U], V \rangle = \mathbb{E}^*[I(\eta^*) \langle U, X \rangle \langle V, X \rangle]$$

Design tensor for a comparison context

$$X = (e_a - e_b) \otimes e_c$$

Pairwise BTL score and Fisher information

$$\begin{aligned} s(Y, \eta) &= Y - \sigma(\eta); \\ I(\eta) &= \sigma(\eta) (1 - \sigma(\eta)) \end{aligned}$$

Classical tensor inference literature:

Assume *additive noise, uniform sampling, homoscedastic*. Then

$$I(\eta) \equiv \frac{1}{\sigma^2} \text{ and } G = \frac{1}{\sigma^2} d^* \cdot I.$$

$\Rightarrow$ the efficient one-step reduces to

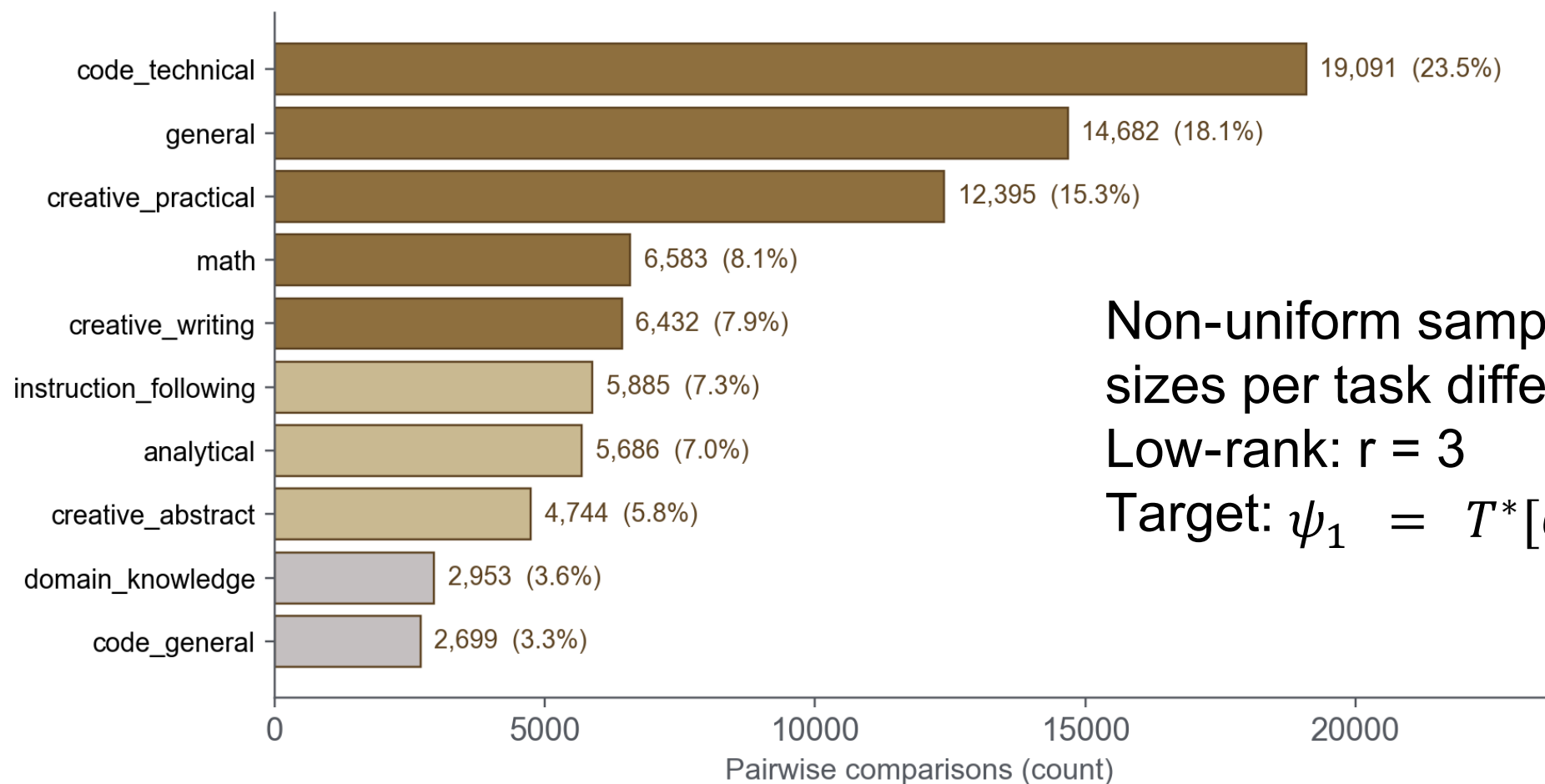
$$(Y - \langle T^*, X \rangle) \cdot \langle P\Gamma, X \rangle$$

in existing tensor inference literature.

Under heteroscedasticity, nonuniform sampling, and nonlinear link (pairwise BTL!), the identity of G breaks.

Experiments: Real Arena Data

LMarena snapshot: N = 81,150 comparisons, top-30 models × 10 tasks



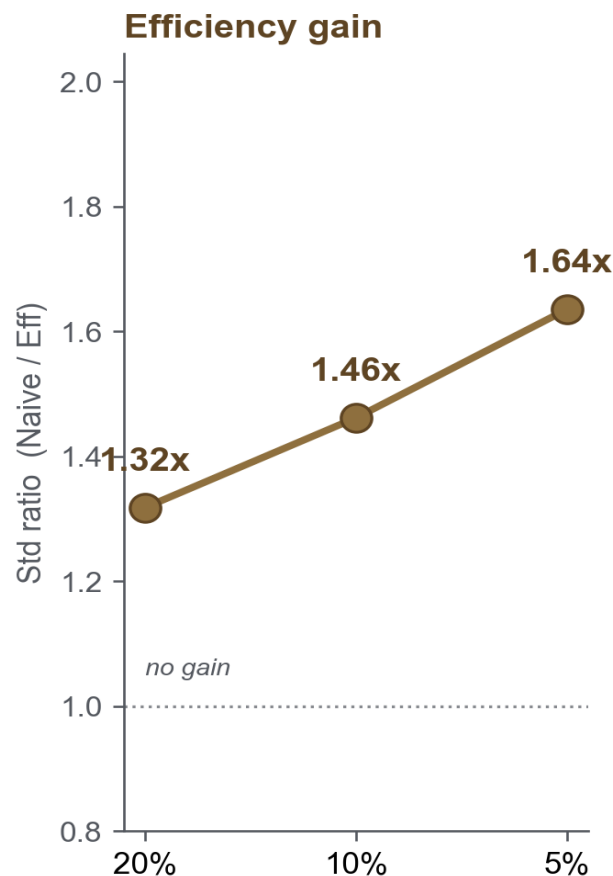
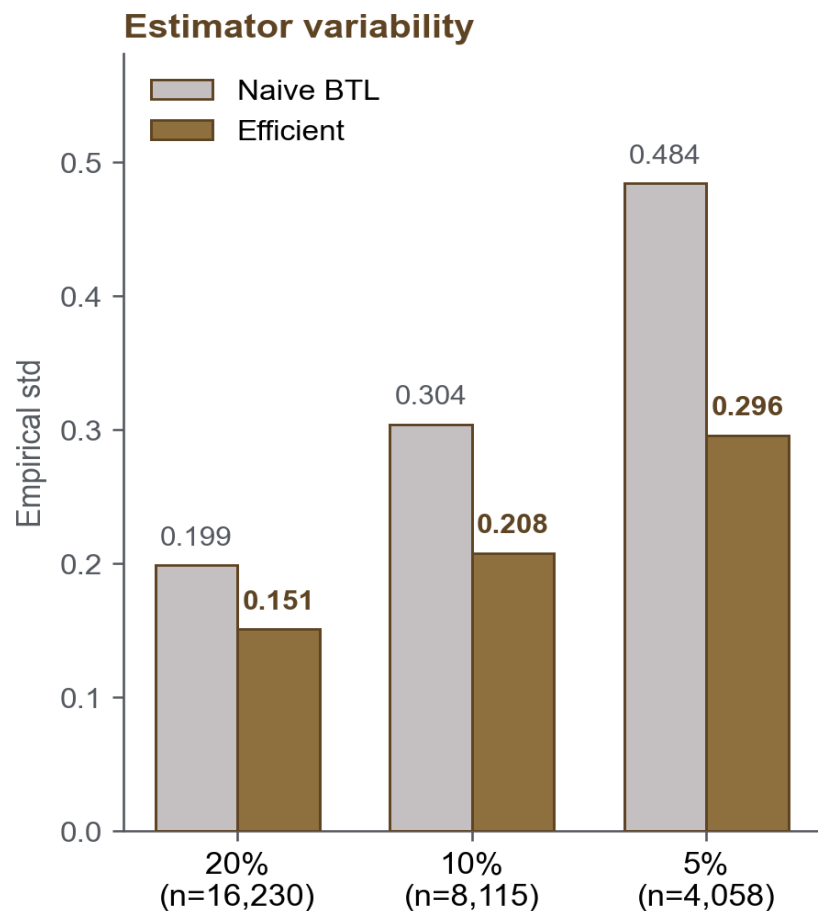
Non-uniform sampling: Sample sizes per task differ by ~7x

Low-rank: $r = 3$

Target: $\psi_1 = T^*[Gemini, math]$

Sample Efficiency Advantage Grows with Scarcity

$T^*[\text{Gemini, math}] = 0.687$ — Naive vs Efficient (500 trials, 20 / 10 / 5 %)

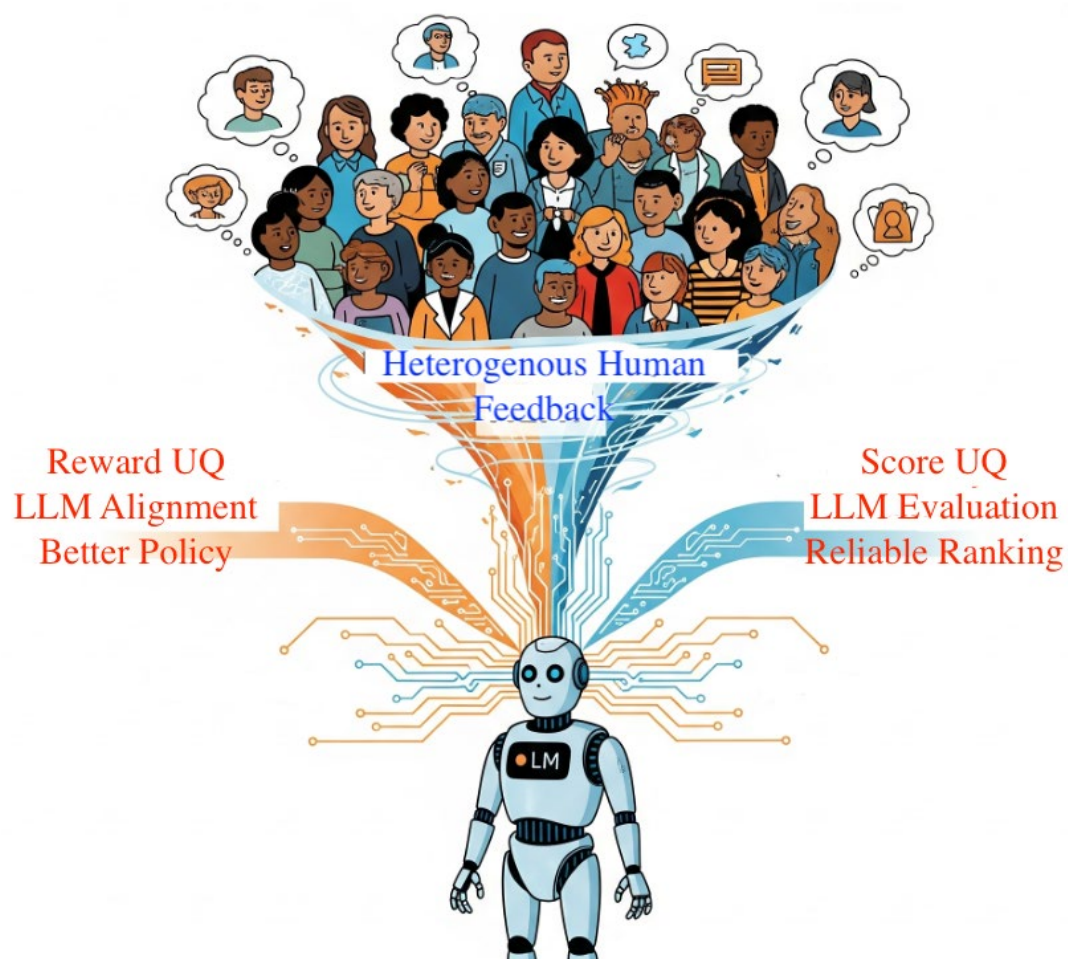


Coverage & rejection rate ($H_0: \psi \leq 0$)

Frac	Method	Cov	Rej
20%	Naive	0.998	81%
20%	Efficient	0.930	98%
10%	Naive	0.982	52%
10%	Efficient	0.924	80%
5%	Naive	0.972	32%
5%	Efficient	0.908	47%

- Efficiency gain grows steadily: 1.32x at 20% → 1.46x at 10% → 1.64x at 5%.
- At 10%, naive rejection rate drops to 52%; Efficient still rejects at 80%.
- Efficient coverage stays above 0.90 at all fractions.

Summary



Will Wei Sun
Purdue University
sun244@purdue.edu
Thanks NSF!

Future
direction:



Design-based
active learning