

Sample-Efficient Personalized Reward Modeling for Pluralistic Alignment

Ramya Korlakai Vinayak

Assistant Professor

Department of Electrical & Computer Engineering
Affiliated with Dept. of CS and Dept. of Statistics

ramya@ece.wisc.edu



Joint work w/ Daiwei Chen, Yi Chen, Aniket Rege, Zhi Wang, Geelon So, Gokcan Tatli, Greg Canal, Blake Mason, Rob Nowak

Outline

- Motivation: Why should we care about heterogeneity in human preferences?
- Reward modeling for preference learning
- Personalizable reward modeling
- Learning metric and multiple user preferences simultaneously
- Sample efficiency by leveraging commonalities:
Personalizable reward modeling for pluralistic alignment (PAL)
- Experimental results
- Future directions



Generated by ChatGPT

References

- D. Chen, Y. Chen, A. Rege, Z. Wang, R. K. Vinayak, “PAL: Sample-Efficient Personalized Reward Models for Pluralistic Alignment”, ICLR 2025
- Z. Wang, G. So, R. K. Vinayak, “Metric Learning from Limited Pairwise Preference Comparisons”, UAI 2024.
- G. Canal, B. Mason, R. K. Vinayak, R. Nowak, “One for All: Simultaneous Metric and Preference Learning over Multiple Users”, Neurips 2022
- G. Tatli, Y. Chen, R. K. Vinayak “Learning Population of Preferences via Pairwise Comparison Queries”, AISTATS 2024
- G. Tatli, R. Nowak, R. K. Vinayak “Metric Learning in RKHS”, UAI 2025
- Z. Wang, C. Chen, R. K. Vinayak “Bridging Lifelong Learning and Multitask Representation Learning via Algorithmic and Complexity Measure”, ALT 2026

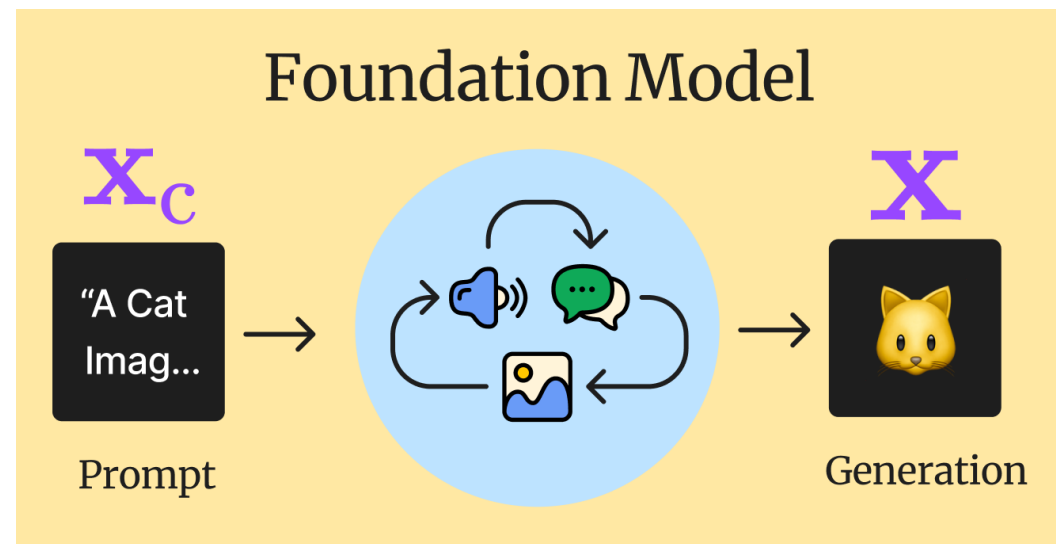
Outline

- Motivation: Why should we care about heterogeneity in human preferences?



Generated by ChatGPT

Aligning AI/ML to human preferences



“Alignment”:
further tuning using a lot
of human feedback

Prompt: A cat Image

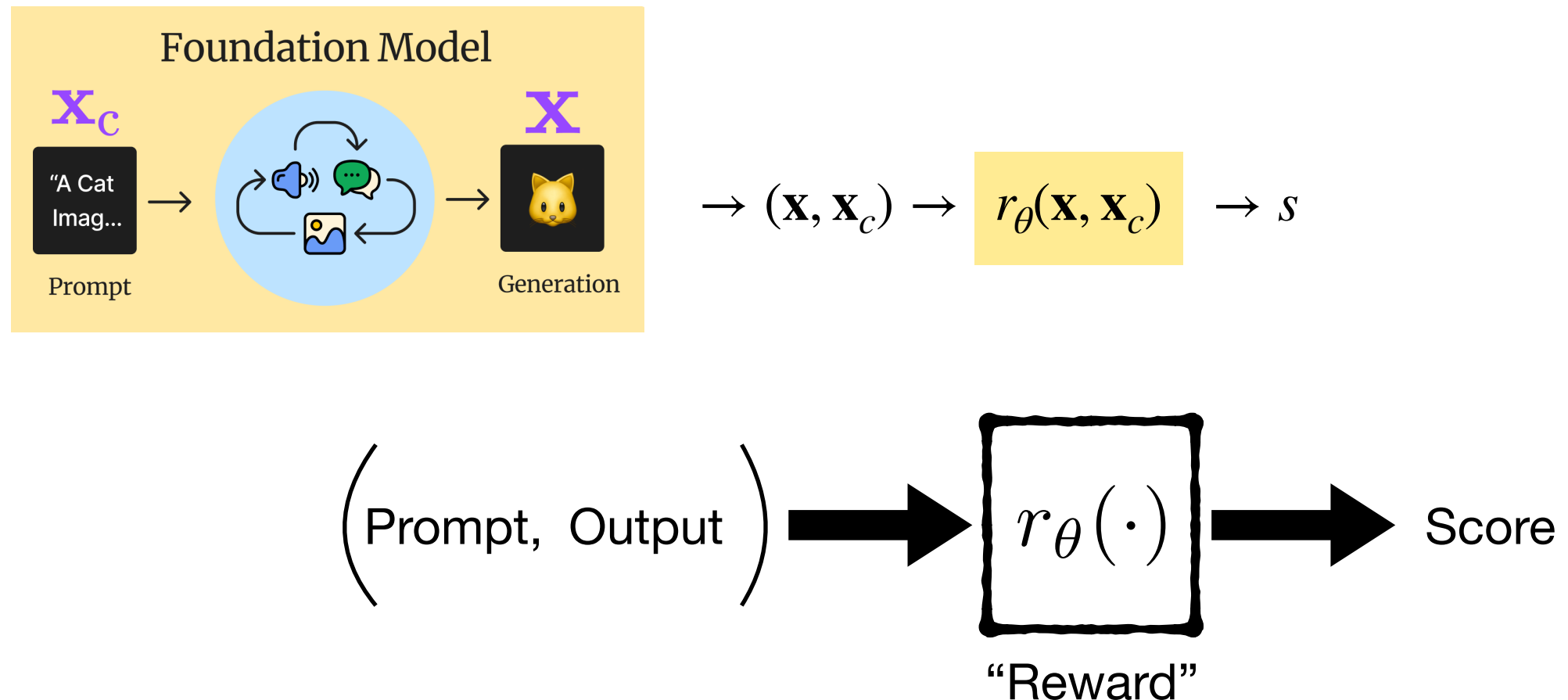
Two outputs: 🐱 , 🐈

👤 : I prefer 🐈 over 🐱

Pairwise
Comparisons

- Large models trained on internet-scale data — are not readily deployable safely in the real world
- They are often very heavily fine-tuned with human feedback

Aligning AI/ML to human preferences



- Learn a reward function that will score the “outputs” from the pre-trained model that are preferred more by humans with higher score
- Use reward function to either (i) fine-tune the model to provide more preferred outputs or (ii) output the top ranked among multiple outputs during inference

“Plurality”: Human preferences are heterogeneous

- Alignment of AI: Whose preferences are we aligning AI to?
- We are not a monolith: different people have different preferences
- How should we model heterogeneous human preferences?
- Can we learn *personalizable* rewards with few samples per user?

2402.05070v1 [cs.AI] 7 Feb 2024

A Roadmap to Pluralistic Alignment

Taylor Sorensen¹ Jared Moore² Jillian Fisher^{1,3} Mitchell Gordon^{1,4} Niloofar Mireshghallah¹
Christopher Michael Rytting¹ Andre Ye¹ Liwei Jiang^{1,5} Ximing Lu¹ Nouha Dziri⁵ Tim Althoff¹
Yejin Choi^{1,5}

Abstract

With increased power and prevalence of AI systems, it is ever more critical that AI systems are designed to serve *all*, i.e., people with diverse values and perspectives. However, aligning models to serve *pluralistic* human values remains an open research question. In this piece, we propose a roadmap to pluralistic alignment, specifically using language models as a test bed. We identify and formalize three possible ways to define and operationalize pluralism in AI systems: 1) *Overton* pluralistic models that present a spectrum of reasonable responses; 2) *Steerably pluralistic* models that can steer to reflect certain perspectives; and 3) *Distributionally pluralistic* models that are well-calibrated to a given population in distribution. We also propose and formalize three possible classes of *pluralistic benchmarks*: 1) *Multi-objective* benchmarks, 2) *Trade-off steerable* benchmarks, which incentivize models to steer to arbitrary trade-offs, and 3) *Jury-pluralistic* benchmarks which explicitly model diverse human ratings. We use this framework to argue that current alignment techniques may be flawed.

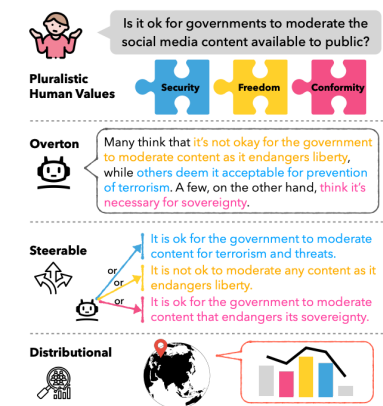
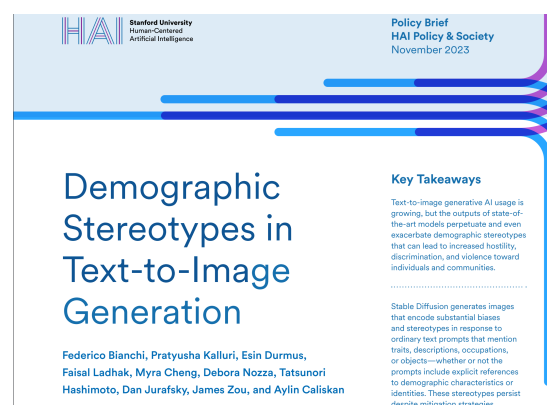


Figure 1. Three kinds of pluralism in models.



OPINIONS ABOUT ARTIFICIAL INTELLIGENCE – BY COUNTRY

Q. Let's now talk about products and services using artificial intelligence (AI). Artificial intelligence refers to computers and robots doing things that traditionally require using human intelligence. How much do you agree or disagree with the following? % "Agree"

	Argentina	Australia	Brazil	Canada	China	France	Germany	India	Indonesia	Italy	Japan	South Korea	Malaysia	Mexico	Netherlands	Poland	Russia	Saudi Arabia	Spain	Sweden	Switzerland	Taiwan	Thailand	UK	USA
I have a good understanding of what artificial intelligence is	64%	64%	59%	60%	69%	59%	70%	67%	71%	56%	62%	59%	57%	62%	73%	62%	61%	70%	74%	61%	69%	70%	60%	63%	58%
Products and services using artificial intelligence will profoundly change my daily life in the next 5-10 years	60%	60%	52%	52%	61%	44%	67%	60%	63%	44%	50%	40%	40%	51%	74%	53%	53%	76%	60%	71%	53%	71%	56%	60%	52%
Products and services using artificial intelligence make my life easier	60%	59%	46%	49%	62%	44%	70%	67%	71%	41%	58%	39%	40%	50%	72%	54%	52%	76%	73%	71%	47%	74%	58%	64%	57%
Products and services using artificial intelligence have more benefits than drawbacks	52%	51%	37%	38%	57%	32%	63%	70%	64%	37%	52%	31%	38%	49%	73%	50%	42%	62%	61%	63%	53%	70%	46%	53%	57%
I know which types of products and services use artificial intelligence	58%	47%	38%	37%	58%	36%	59%	76%	62%	37%	46%	34%	37%	38%	69%	43%	52%	60%	61%	61%	61%	52%	57%	69%	57%
I trust companies that use artificial intelligence as much as I trust other companies	50%	51%	36%	40%	50%	34%	54%	76%	71%	37%	42%	30%	34%	35%	48%	48%	39%	46%	61%	61%	51%	52%	71%	69%	56%
Products and services using artificial intelligence have profoundly changed my daily life in the past 5-10 years	49%	53%	37%	37%	52%	32%	54%	73%	58%	33%	49%	32%	33%	38%	67%	43%	50%	62%	62%	60%	60%	61%	56%	72%	56%
Products and services using artificial intelligence make me nervous	38%	33%	51%	42%	33%	47%	38%	30%	39%	37%	48%	32%	30%	32%	53%	33%	35%	37%	38%	38%	38%	31%	48%	52%	52%

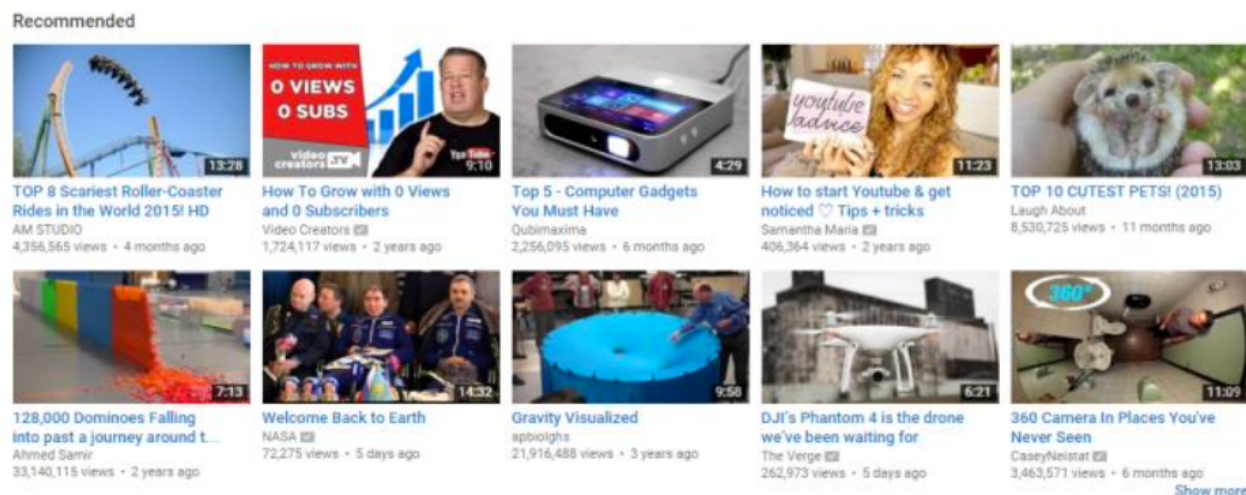
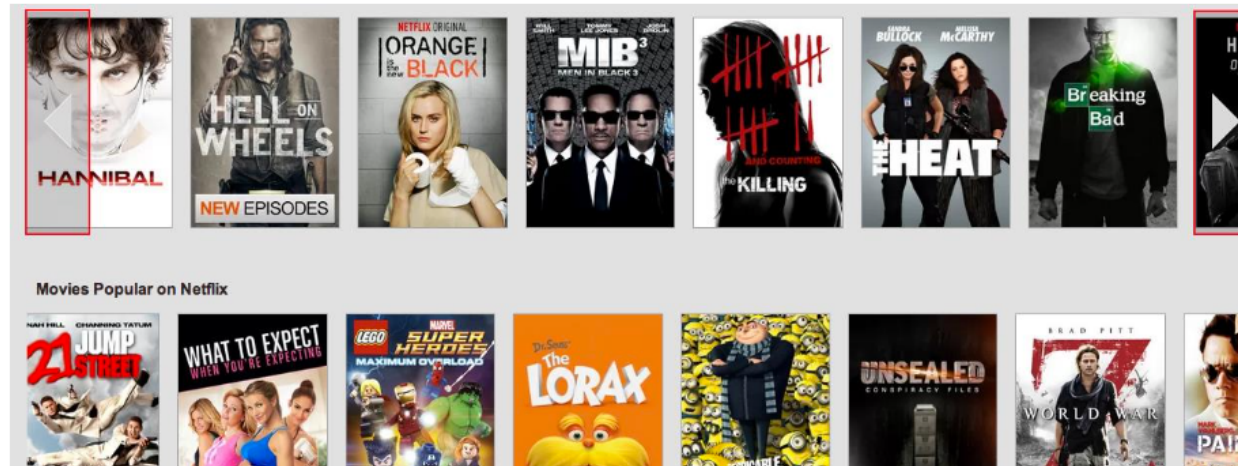
Tao et al., Cultural bias and cultural alignment of large language models." *PNAS nexus* 3.9, 2024

Bianchi et. al., Demographic Stereotypes in Text-to-image models, Stanford University HAI Policy Brief, 2023

World Economic Forum. "Artificial intelligence: Do we trust it?," *World Economic Forum*, Jan. 12, 2022.

<https://www.weforum.org/stories/2022/01/artificial-intelligence-ai-technology-trust-survey/>

Recommendations are everywhere



FOR YOU

Recommended stories



Facebook thought pandemic online shopping would last forever. It didn't.

By Naomi Nix and Jaclyn Peiser

Outline

- Motivation: Why should we care about heterogeneity in human preferences?
- Reward modeling for preference learning



Generated by ChatGPT

How is alignment done from human preferences?

2.2.3 Fitting the Reward Function

We can interpret a reward function estimate $\hat{r}$ as a preference-predictor if we view $\hat{r}$ as a latent factor explaining the human's judgments and assume that the human's probability of preferring a segment σ^i depends exponentially on the value of the latent reward summed over the length of the clip:³

$$\hat{P}[\sigma^1 \succ \sigma^2] = \frac{\exp \sum \hat{r}(o_t^1, a_t^1)}{\exp \sum \hat{r}(o_t^1, a_t^1) + \exp \sum \hat{r}(o_t^2, a_t^2)}. \quad (1)$$

We choose $\hat{r}$ to minimize the cross-entropy loss between these predictions and the actual human labels:

$$\text{loss}(\hat{r}) = - \sum_{(\sigma^1, \sigma^2, \mu) \in \mathcal{D}} \mu(1) \log \hat{P}[\sigma^1 \succ \sigma^2] + \mu(2) \log \hat{P}[\sigma^2 \succ \sigma^1].$$

Deep Reinforcement Learning from Human Preferences

Paul F Christiano
OpenAI
paul@openai.com

Jan Leike
DeepMind
leike@google.com

Tom B Brown
nottombrown@gmail.com

Miljan Martic
DeepMind
miljanm@google.com

Shane Legg
DeepMind
legg@google.com

Dario Amodei
OpenAI
damodei@openai.com

Training language models to follow instructions with human feedback

Long Ouyang* Jeff Wu* Xu Jiang* Diogo Almeida* Carroll L. Wainwright*

Pamela Mishkin* Chong Zhang Sandhini Agarwal Katarina Slama Alex Ray

John Schulman Jacob Hilton Fraser Kelton Luke Miller Maddie Simens

Amanda Askell† Peter Welinder Paul Christiano*†

Jan Leike*

Ryan Lowe*

OpenAI

Direct Preference Optimization: Your Language Model is Secretly a Reward Model

Rafael Rafailov*†

Archit Sharma*†

Eric Mitchell*†

Stefano Ermon†‡

Christopher D. Manning†

Chelsea Finn†

*Stanford University †CZ Biohub
{rafailov, architsh, eric.mitchell}@cs.stanford.edu

• Bradley-Terry Model, 1952

$$\Pr(\mathbf{x}_l \succ \mathbf{x}_r | \mathbf{x}_c) = \frac{e^{r_\theta(\mathbf{x}_l; \mathbf{x}_c)}}{e^{r_\theta(\mathbf{x}_l; \mathbf{x}_c)} + e^{r_\theta(\mathbf{x}_r; \mathbf{x}_c)}}$$

$\mathbf{x}_c$: prompt or context

$\mathbf{x}_l$, $\mathbf{x}_r$: generated outputs

RANK ANALYSIS OF INCOMPLETE BLOCK DESIGNS

I. THE METHOD OF PAIRED COMPARISONS

BY RALPH ALLAN BRADLEY AND MILTON E. TERRY*

Virginia Agricultural Experiment Station of the Virginia Polytechnic Institute

Limitations

- Bradley-Terry Model, 1952

RANK ANALYSIS OF INCOMPLETE BLOCK DESIGNS

I. THE METHOD OF PAIRED COMPARISONS

BY RALPH ALLAN BRADLEY AND MILTON E. TERRY*

Virginia Agricultural Experiment Station of the Virginia Polytechnic Institute

Probability that i beats j :
(i is preferred over j)

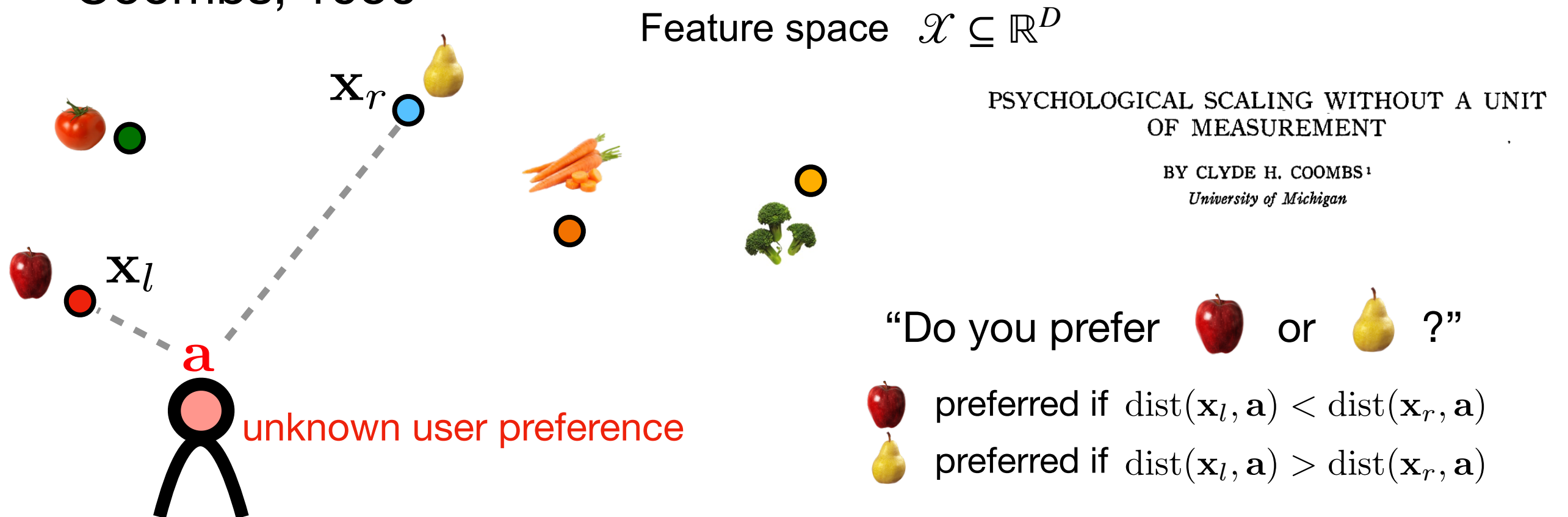
$$P(i \succ j) = \frac{e^{r_i}}{e^{r_i} + e^{r_j}}$$

- **Assumes a single ordering** of the items/alternatives being compared
- By using BTL model, **the alignment methods are inherently assuming a single ordering of the outputs by humans**
- Differences in elicited preferences are “noisy” answers where the level of noise depends on how close the weights/scores of the two alternates are with a specific **sigmoid noise model**

Thurstone 1927, Bradley & Terry 1952, Luce 1959, Joe 1988, Hunter 2004, Ammar & Shah 2012, Neghabhan et. al. 2012, Heckel et. al. 2016

Preference models: Ideal Point Model

- Coombs, 1950



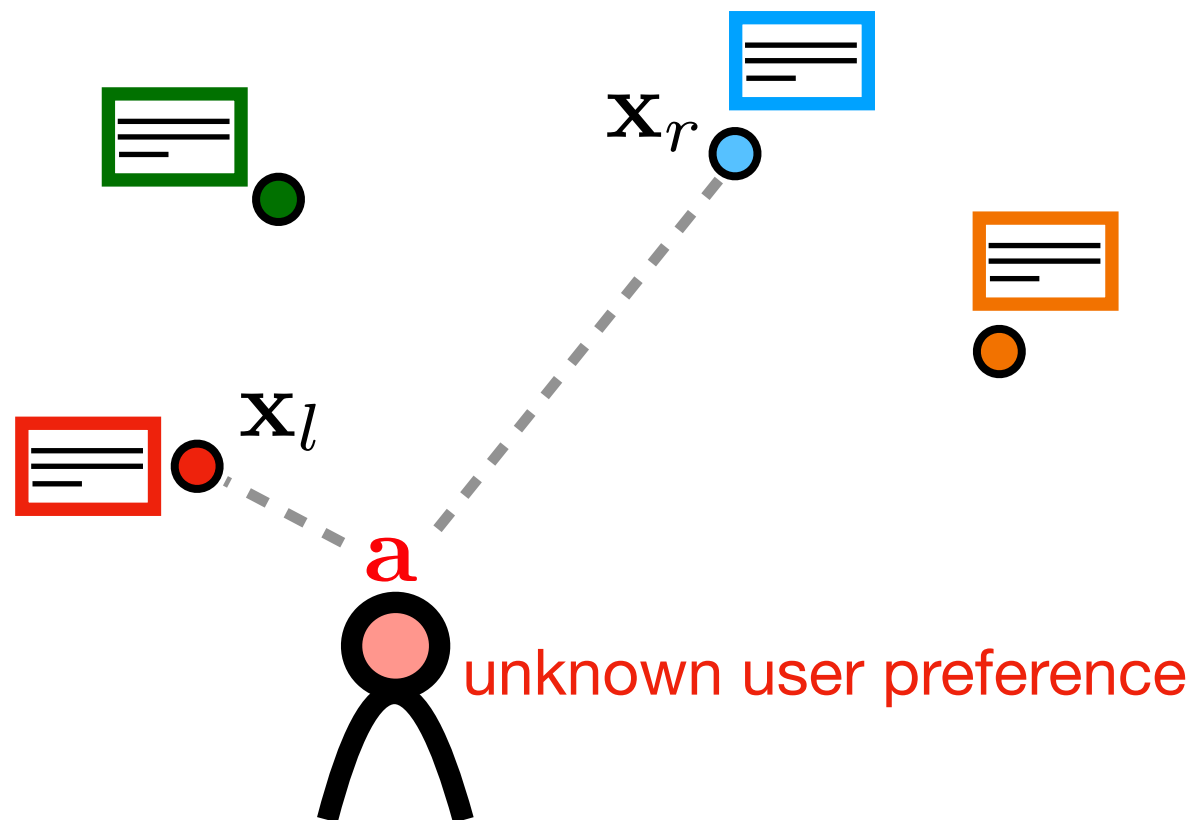
$$P(\text{🍏 preferred over 🍐}) = h(\text{dist}(\mathbf{x}_l, \mathbf{a}) - \text{dist}(\mathbf{x}_r, \mathbf{a}))$$

h : strict monotonic function; when diff of dist is zero, the prob is 1/2, otherwise noise $< 1/2$

Preference models: Ideal Point Model

$$\mathbf{x}_l := \begin{bmatrix} \mathbf{x}_l \\ \mathbf{x}_c \end{bmatrix} \quad \mathbf{x}_r := \begin{bmatrix} \mathbf{x}_r \\ \mathbf{x}_c \end{bmatrix}$$

$\mathbf{x}_c$: context



“Do you prefer  or  ?”

 preferred if $\text{dist}(\mathbf{x}_l, \mathbf{a}) < \text{dist}(\mathbf{x}_r, \mathbf{a})$

 preferred if $\text{dist}(\mathbf{x}_l, \mathbf{a}) > \text{dist}(\mathbf{x}_r, \mathbf{a})$

$$P(\text{red-bordered icon preferred over blue-bordered icon}) = h(\text{dist}(\mathbf{x}_l, \mathbf{a}) - \text{dist}(\mathbf{x}_r, \mathbf{a}))$$

h : strict monotonic function; when diff of dist is zero, the prob is 1/2, otherwise noise < 1/2

$$P(\text{red-bordered icon preferred over blue-bordered icon}) = \frac{1}{1 + e^{\text{dist}(\mathbf{x}_r, \mathbf{a}) - \text{dist}(\mathbf{x}_l, \mathbf{a})}}$$

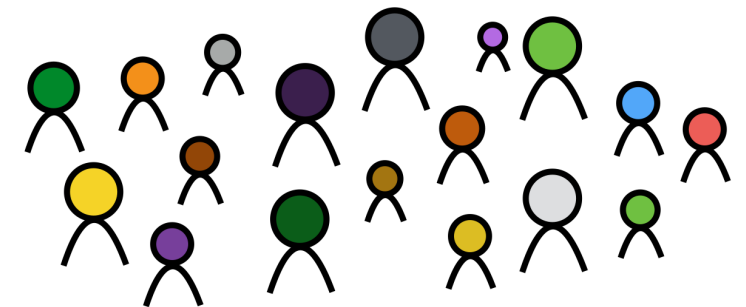
↑
Difference of scores

Past works

- Universal models: a model (preference point) that fits on average for all individuals
- Distance function is assumed to be known

However....

- People have heterogeneous tastes (preferences)
- Distance function is often unknown



Outline

- Motivation: Why should we care about heterogeneity in human preferences?
- Reward modeling for preference learning
- Personalizable reward modeling

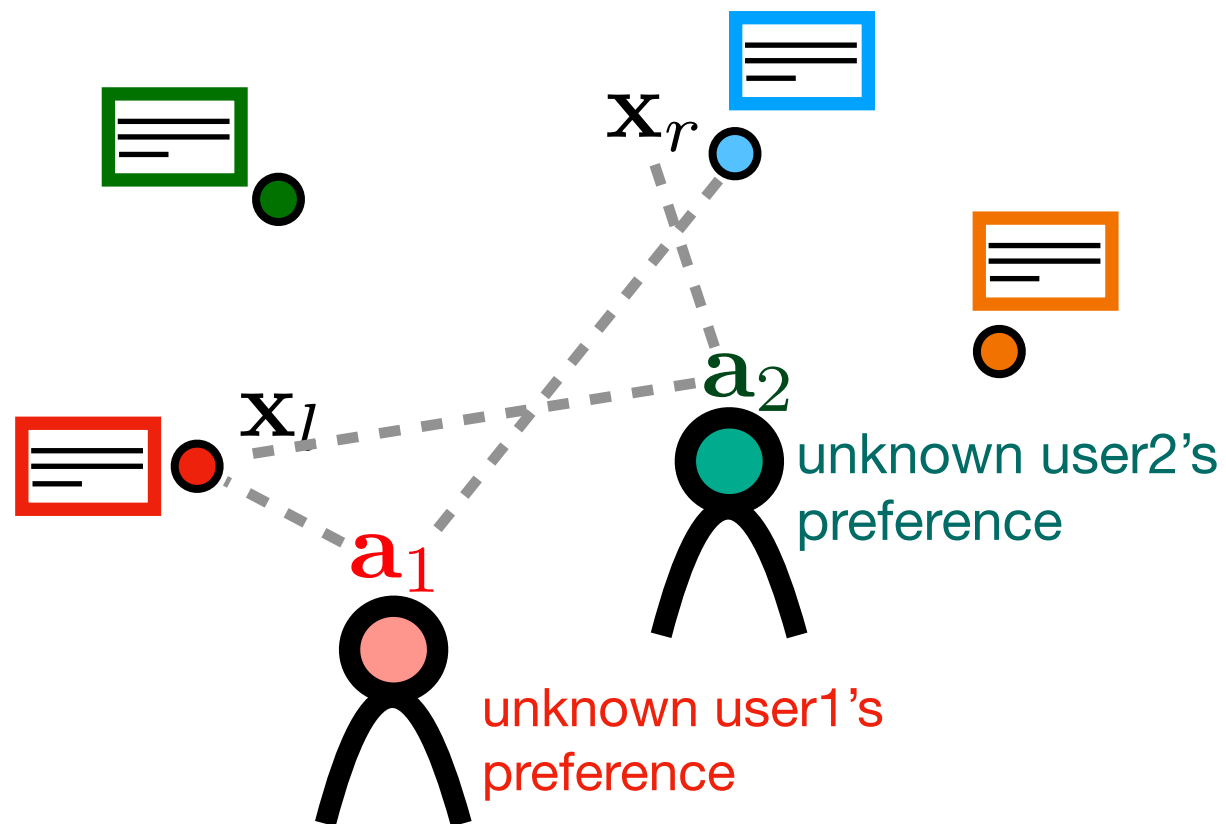


Generated by ChatGPT

Personalized Preferences

$$\mathbf{x}_l := \begin{bmatrix} \mathbf{x}_l \\ \mathbf{x}_c \end{bmatrix} \quad \mathbf{x}_r := \begin{bmatrix} \mathbf{x}_r \\ \mathbf{x}_c \end{bmatrix}$$

$\mathbf{x}_c$: context



“Do you prefer  or  ?”

 preferred if $\text{dist}(\mathbf{x}_l, \mathbf{a}_i) < \text{dist}(\mathbf{x}_r, \mathbf{a}_i)$

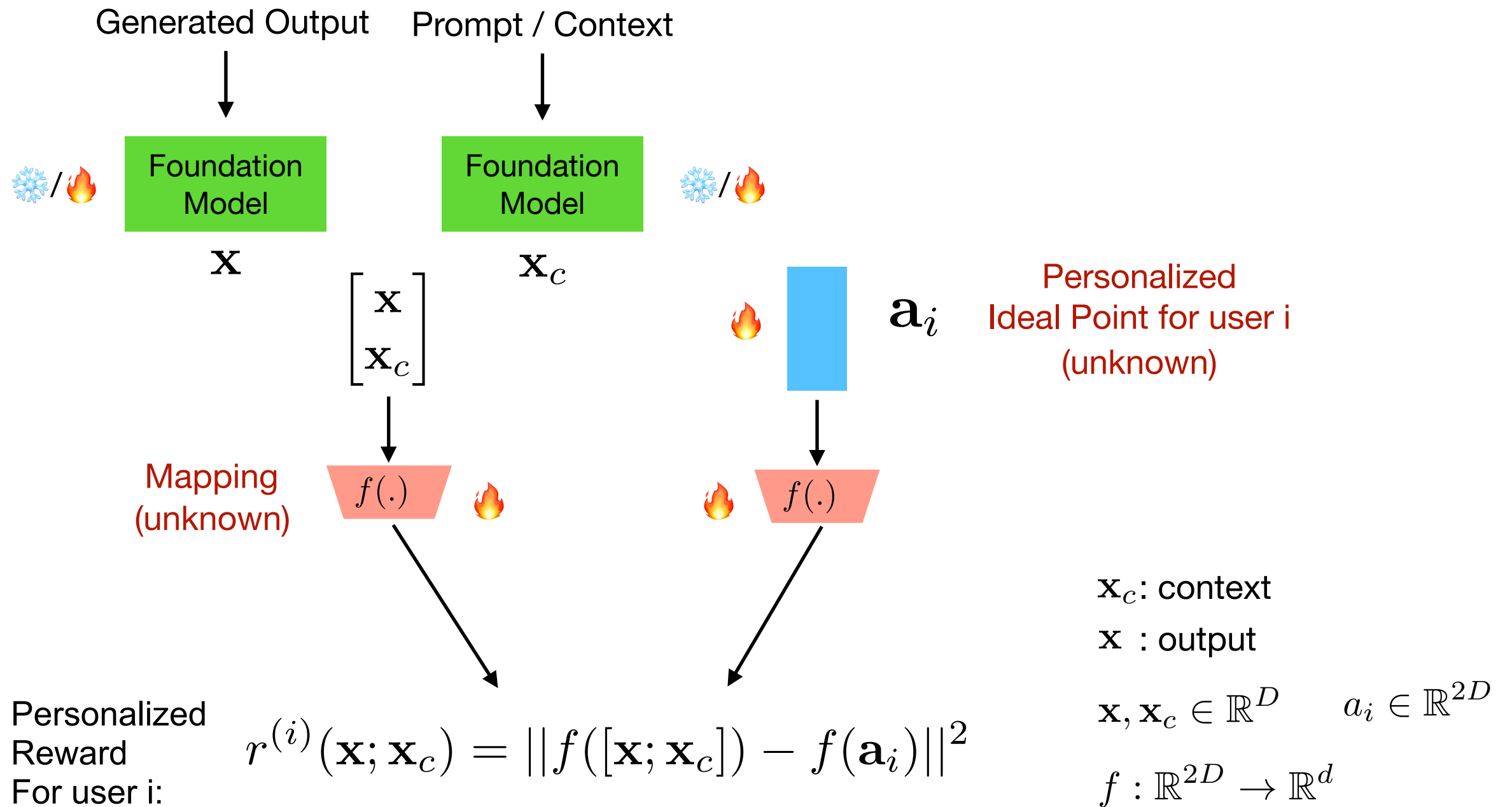
 preferred if $\text{dist}(\mathbf{x}_l, \mathbf{a}_i) > \text{dist}(\mathbf{x}_r, \mathbf{a}_i)$

$i = 1, 2$

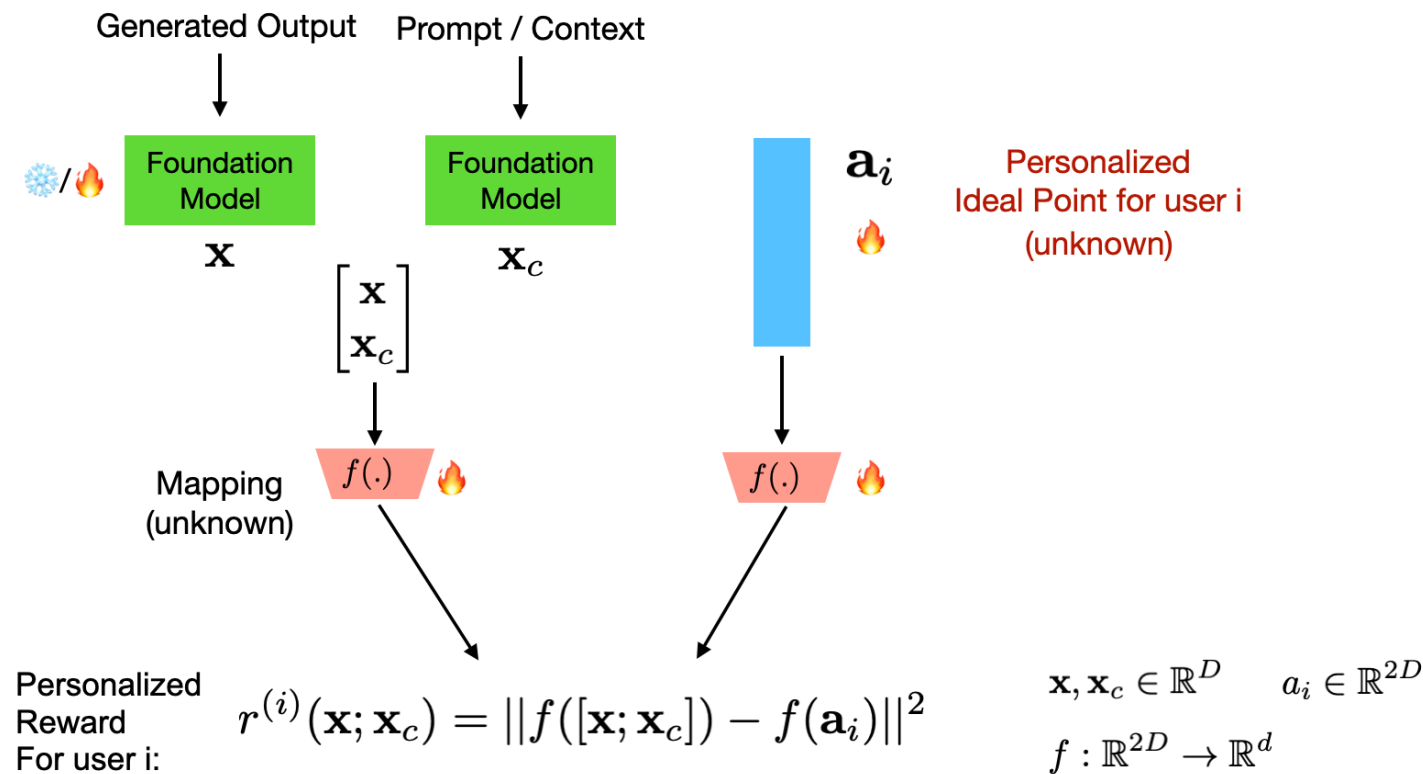
$$P(\text{red box icon preferred over blue box icon} \mid \text{user } i) = h(\text{dist}(\mathbf{x}_l, \mathbf{a}_i) - \text{dist}(\mathbf{x}_r, \mathbf{a}_i))$$

h : strict monotonic function; when diff of dist is zero, the prob is 1/2, otherwise noise < 1/2

Personalizable Reward Modeling



Learning Personalized Preferences



- Learn the unknown mapping f and unknown preference points from pairwise comparison data

$$\arg \min_{f, \mathbf{a}_i} \sum_p \text{loss} \left(y_{p,i} \left(||f(\begin{bmatrix} \mathbf{x}_{l_{p,i}} \\ \mathbf{x}_{c_{p,i}} \end{bmatrix}) - f(\mathbf{a}_i)||^2 - ||f(\begin{bmatrix} \mathbf{x}_{r_{p,i}} \\ \mathbf{x}_{c_{p,i}} \end{bmatrix}) - f(\mathbf{a}_i)||^2 \right) \right)$$

- If we learn individually, each user needs to provide too much data to learn all the components

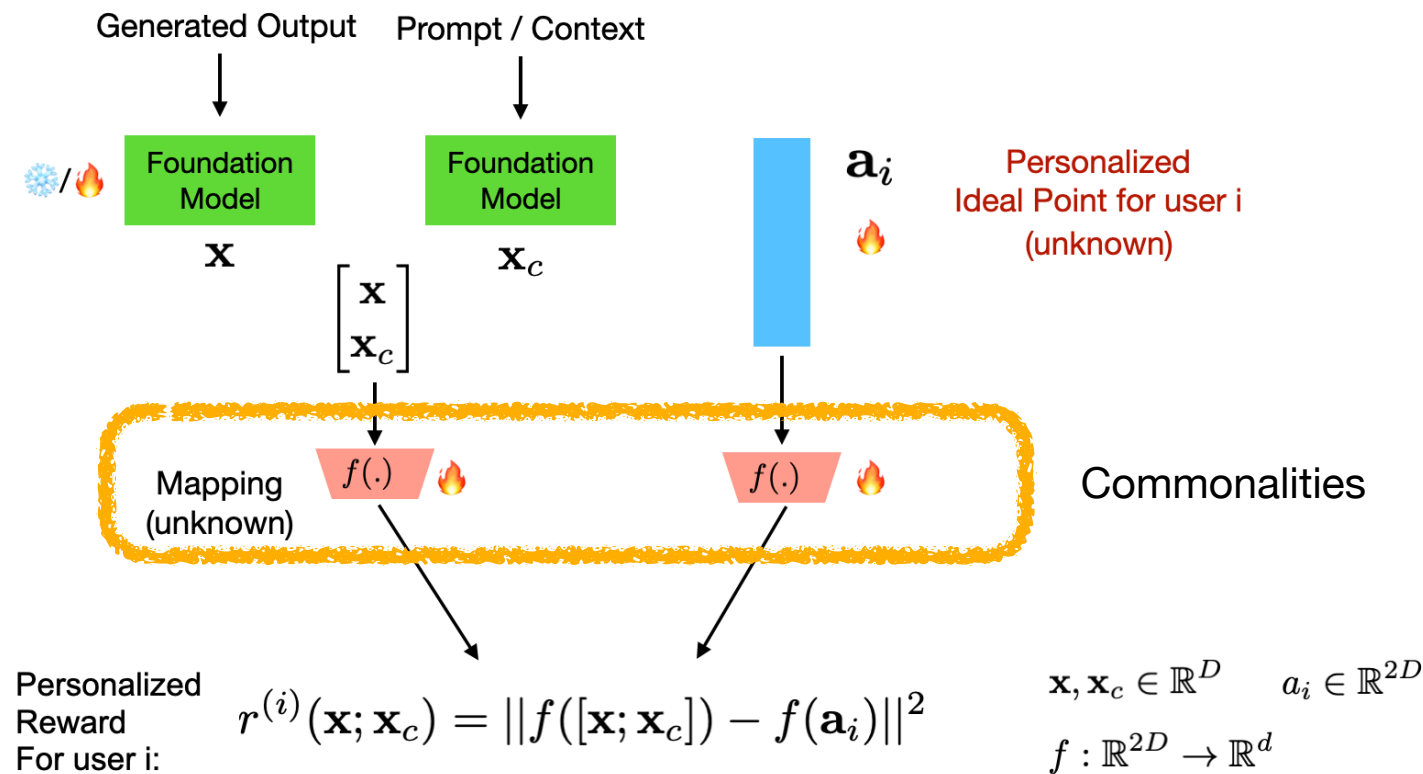
Outline

- Motivation: Why should we care about heterogeneity in human preferences?
- Reward modeling for preference learning
- Personalizable reward modeling
- Learning metric and multiple user preferences simultaneously



Generated by ChatGPT

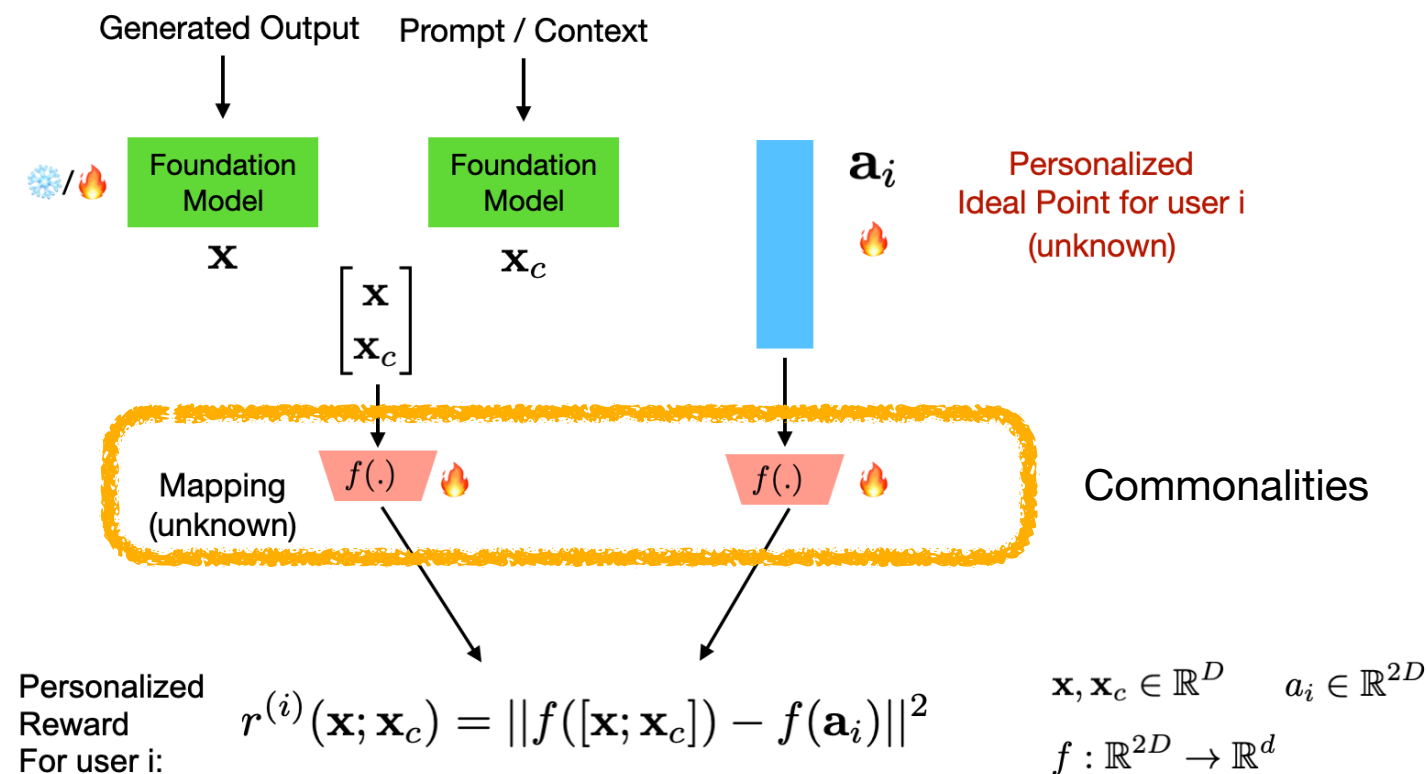
Simultaneous Metric and Preference Learning



- Leverage everyone's data to learn the common mapping f
- Use individual data to learn the personalized preference

$$\arg \min_{f, \mathbf{a}_1, \dots, \mathbf{a}_N} \sum_i \sum_p \text{loss} \left(y_{p,i} \left(||f([\mathbf{x}_{l_{p,i}}; \mathbf{x}_{c_{p,i}}]) - f(\mathbf{a}_i)||^2 - ||f([\mathbf{x}_{r_{p,i}}; \mathbf{x}_{c_{p,i}}]) - f(\mathbf{a}_i)||^2 \right) \right)$$

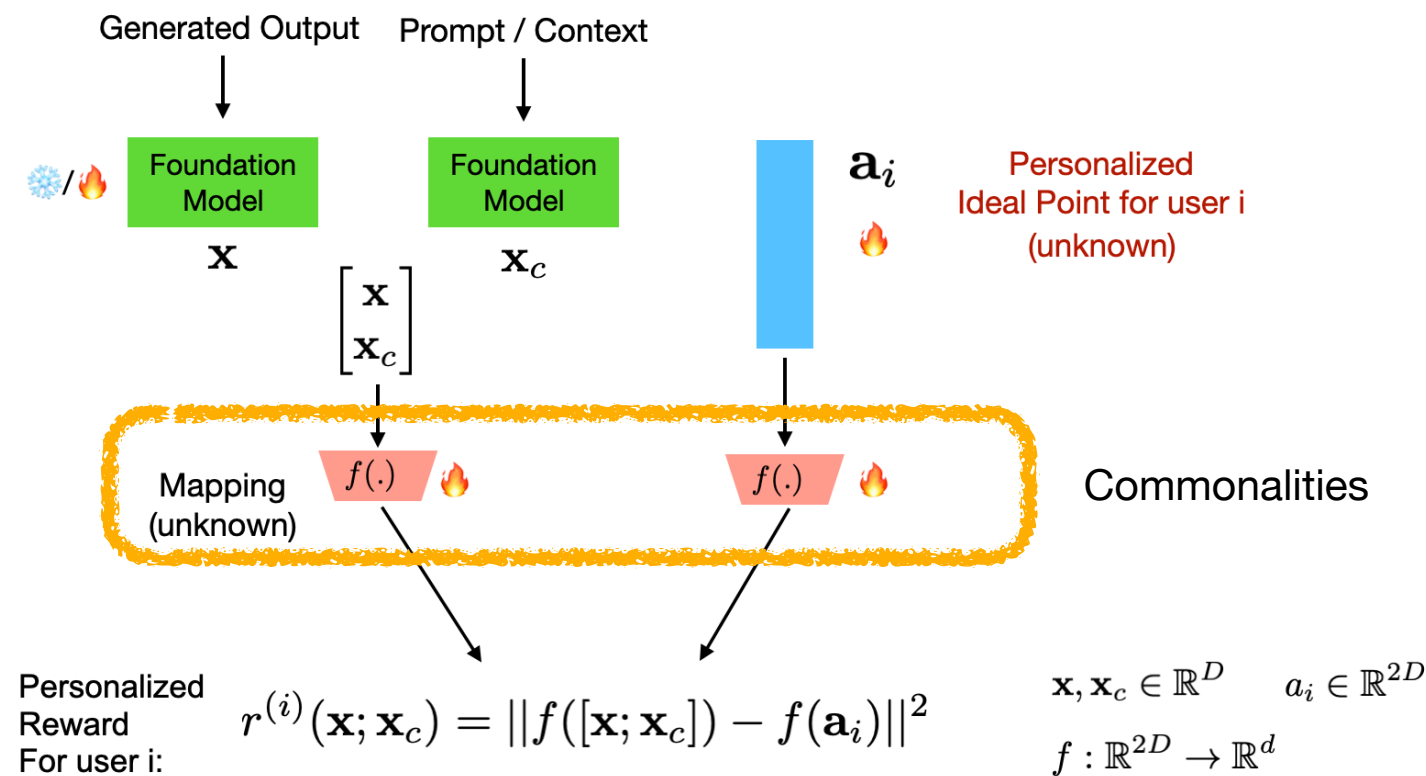
Simultaneous Metric and Preference Learning



- Cost of learning f is amortized across all users
- With sufficiently large number of users in the dataset, each user only needs to provide enough data to learn their latent preference, i.e., ideal point

$O(\text{\#dimensions})$

Metric Learning from Pairwise Comparisons

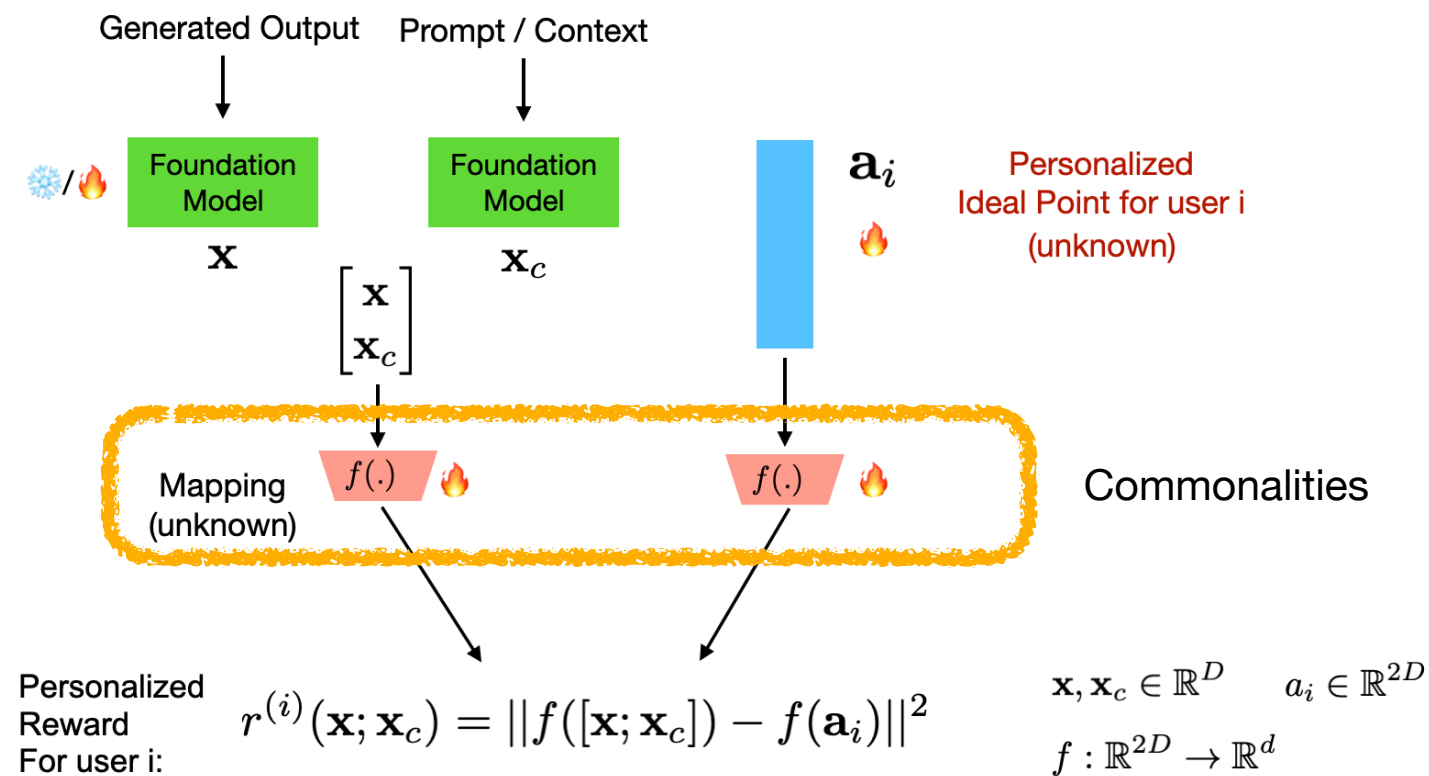


- Each user needs to provide enough data to learn their latent ideal point $O(\text{\#dimensions})$
- High-dimensional — costly

Can we learn the unknown mapping with fewer samples than $O(\text{\#dimensions})$?

Wang, So, Vinayak, “Metric Learning from Limited Pairwise Preference Comparisons”, UAI 2024, arXiv:2403.19629, 2024.

Metric Learning from Pairwise Comparisons



Can we learn the unknown mapping without localizing the preferences?

NO when there is no low-dimensional structure to leverage

Wang, So, Vinayak, "Metric Learning from Limited Pairwise Preference Comparisons", UAI 2024, arXiv:2403.19629, 2024.

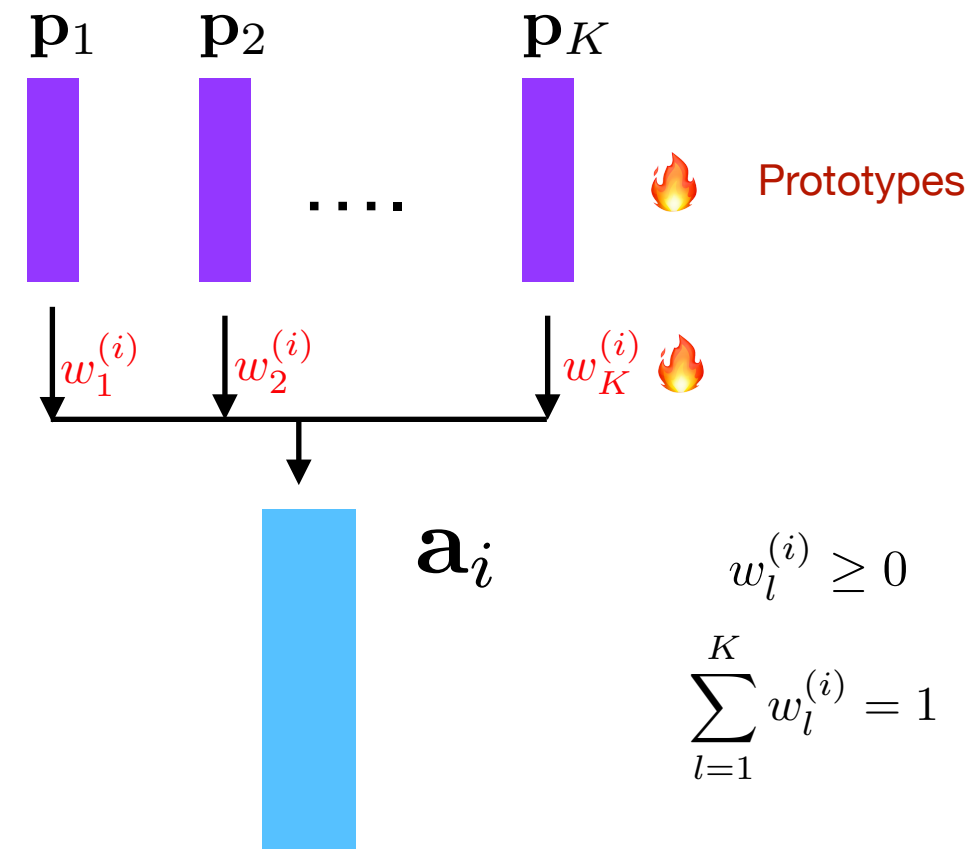
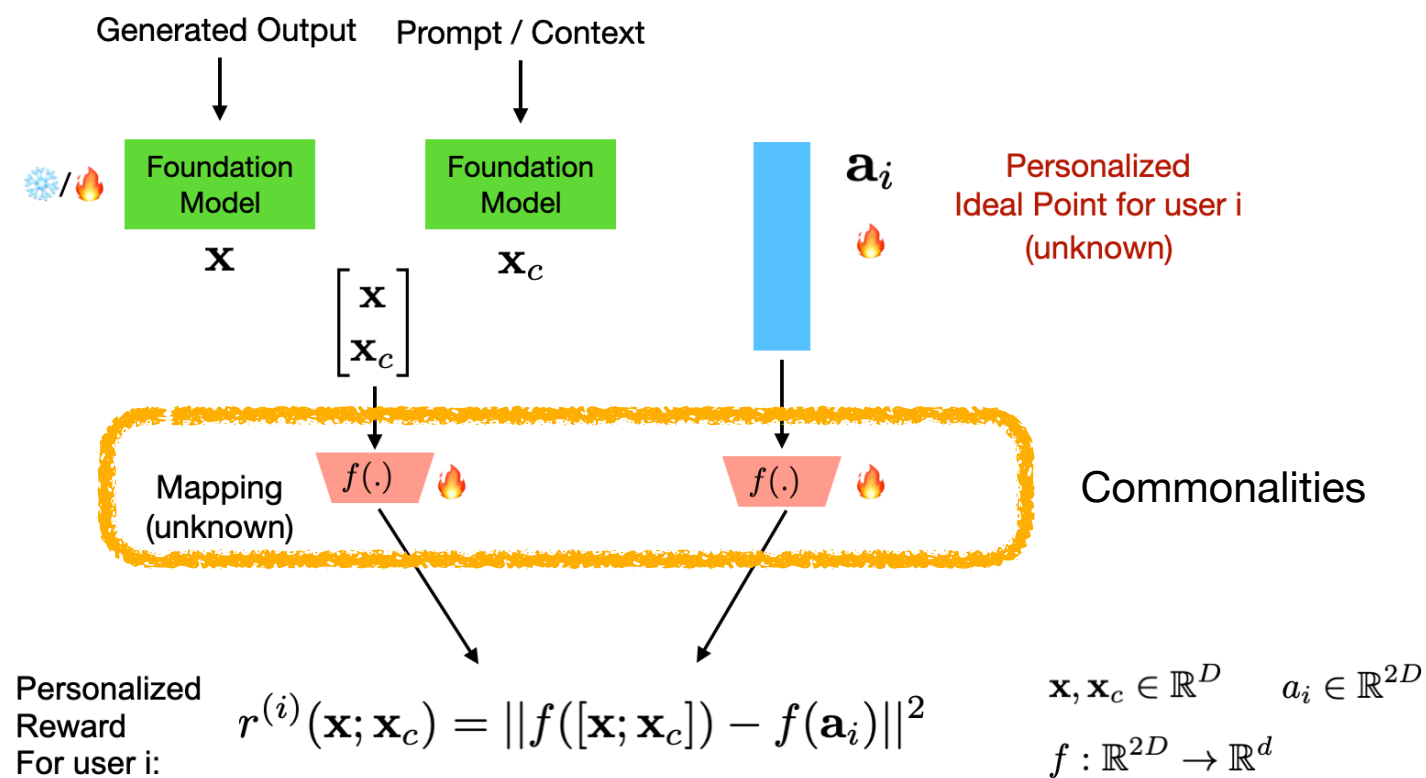
Outline

- Motivation: Why should we care about heterogeneity in human preferences?
- Reward modeling for preference learning
- Personalizable reward modeling
- Learning metric and multiple user preferences simultaneously
- Sample efficiency by leveraging commonalities:
Personalizable reward modeling for pluralistic alignment (PAL)

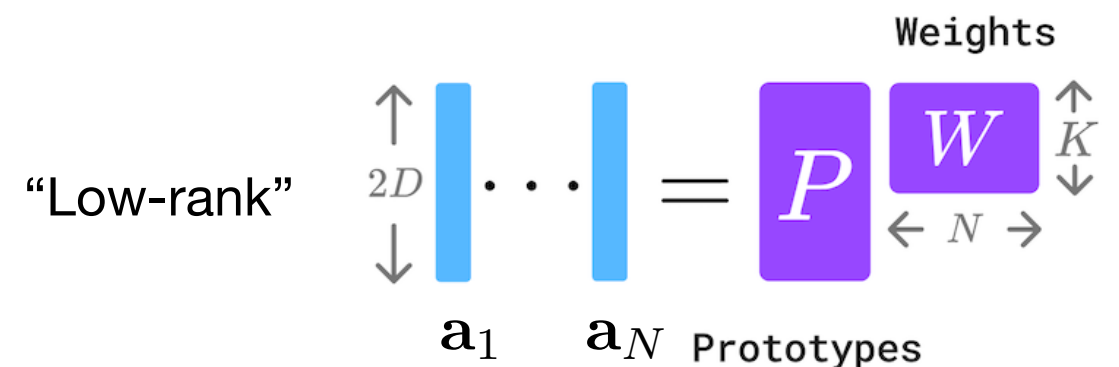


Generated by ChatGPT

Personalizable Reward Modeling: PAL A



- Each user's preference is modeled as a mixture of K prototypical preferences

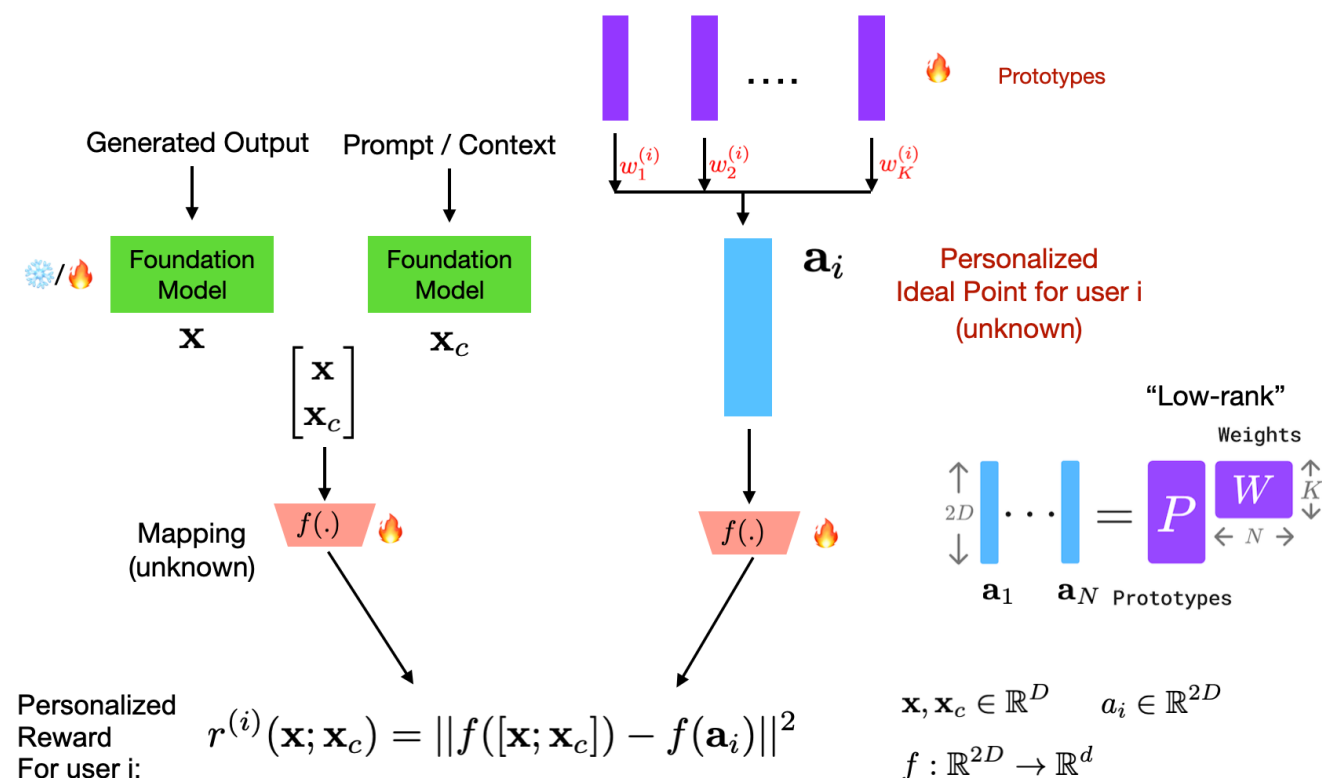


Learning Personalizable Rewards: PAL A

- With sufficient number of users in the dataset, each user needs to only provide enough data to learn their **weights**

$O(\# \text{ of Prototypes})$

$$K \ll D$$



- Leveraging commonalities and low-rank structure in user preferences further reduces the samples needed per user

$$\arg \min_{f, \mathbf{p}_1, \dots, \mathbf{p}_K, w^{(1)}, \dots, w^{(N)}} \sum_i \sum_p \text{loss} \left(y_{p,i} \left(\|f(\begin{bmatrix} \mathbf{x}_{l_p,i} \\ \mathbf{x}_{c_p,i} \end{bmatrix}) - f(\mathbf{a}_i)\|^2 - \|f(\begin{bmatrix} \mathbf{x}_{r_p,i} \\ \mathbf{x}_{c_p,i} \end{bmatrix}) - f(\mathbf{a}_i)\|^2 \right) \right)$$

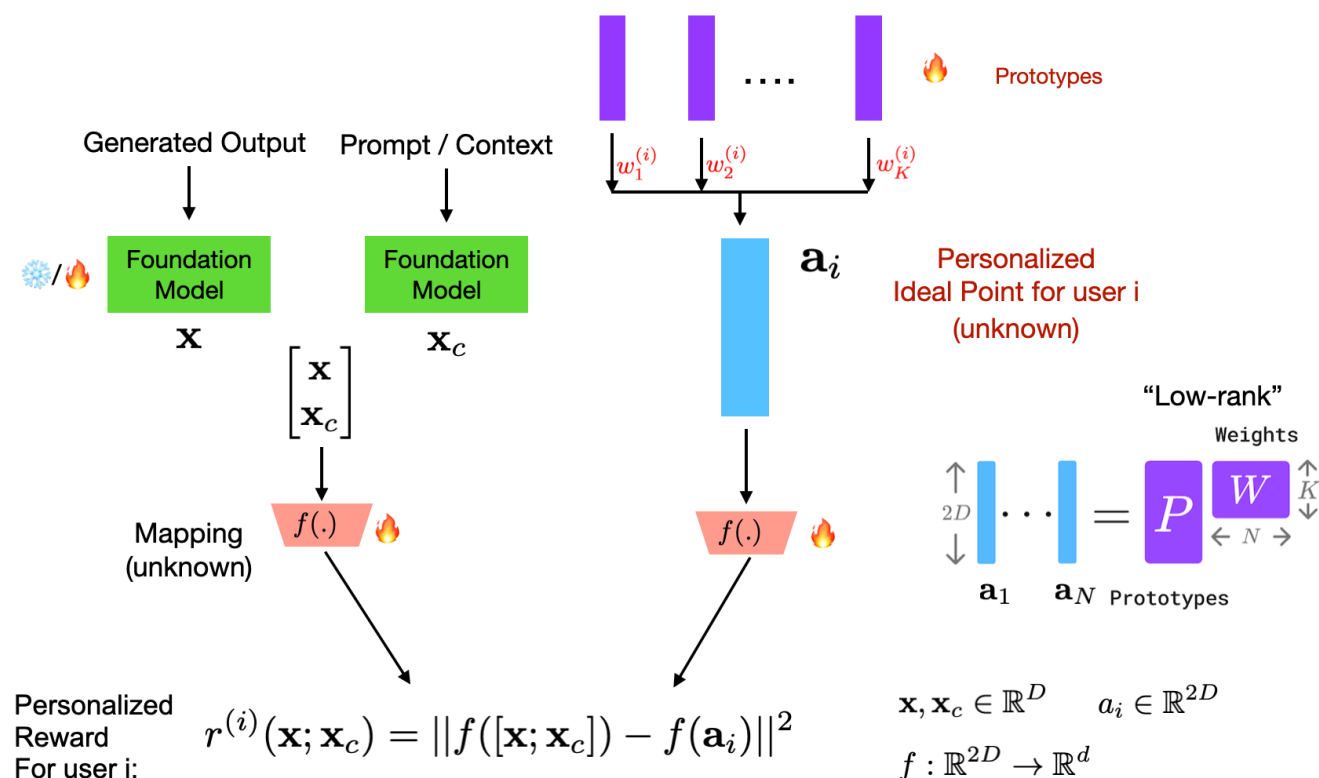
where $\mathbf{a}_i = \sum_{l=1}^K w_l^{(i)} \mathbf{p}_l$ $\sum_{l=1}^K w_l^{(i)} = 1$ $w_l^{(i)} \geq 0 \quad \forall l, i$

Few-shot Generalization to New Users

- New user needs to only provide enough data to learn their **weights**

$O(\# \text{ of Prototypes})$

$$K \ll D$$

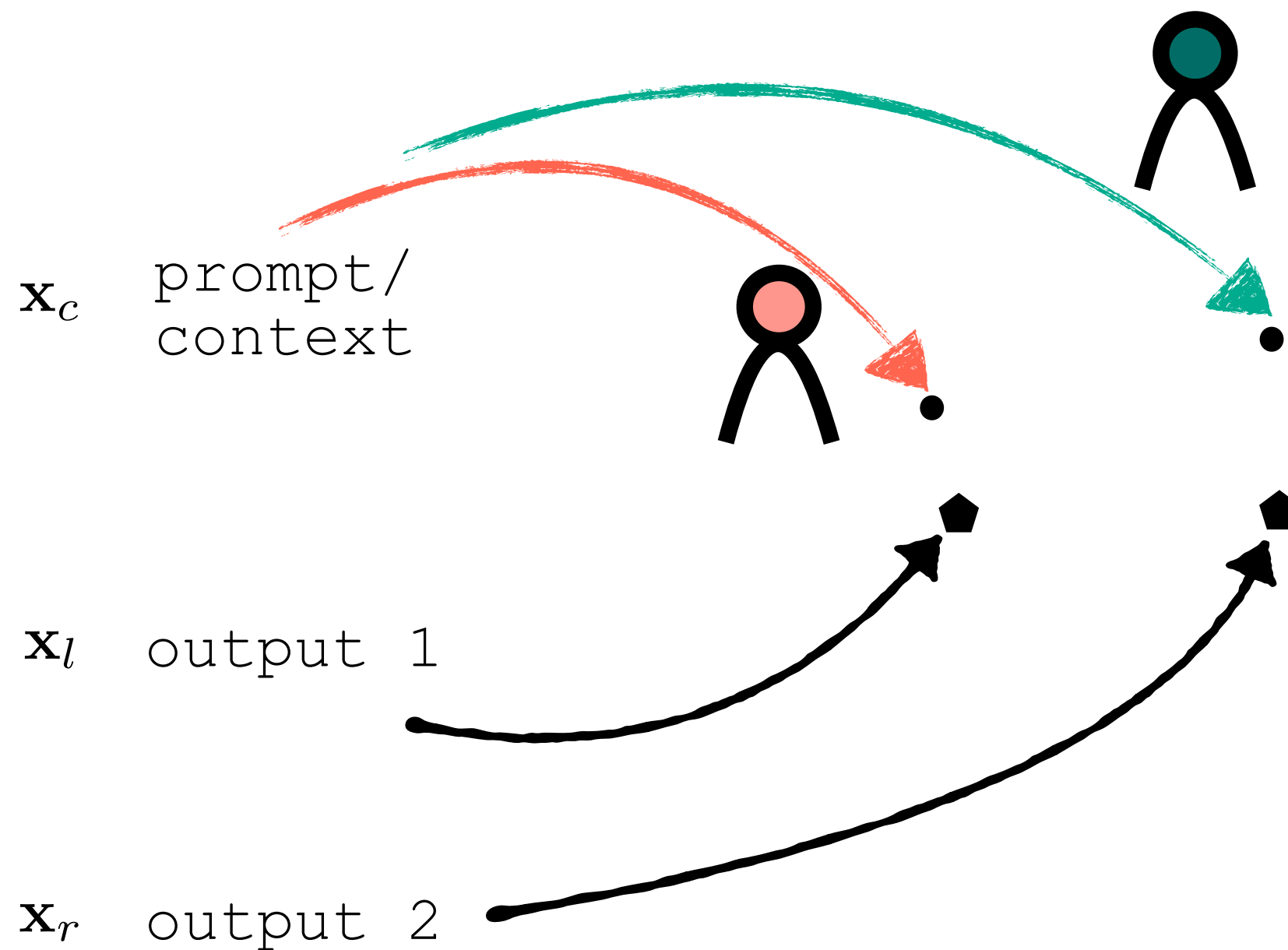


- Mapping f and prototypes are kept fixed
- Few-shot learning to only learn the weights for the new user

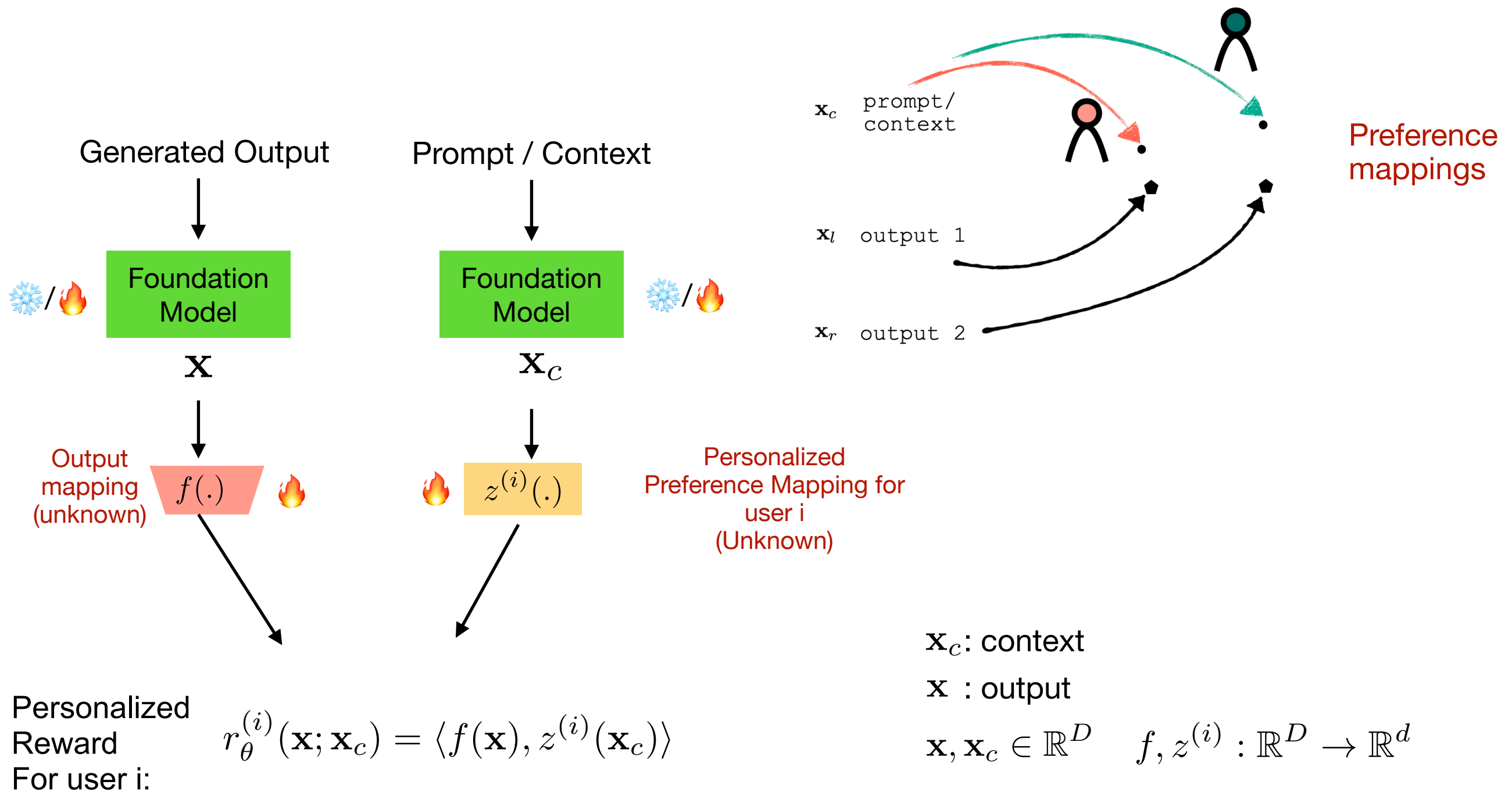
$$\arg \min_{w^{(i)}} \sum_p \text{loss} \left(y_{p,i} \left(||f([\mathbf{x}_{l_{p,i}}; \mathbf{x}_{c_{p,i}}]) - f(\mathbf{a}_i)||^2 - ||f([\mathbf{x}_{r_{p,i}}; \mathbf{x}_{c_{p,i}}]) - f(\mathbf{a}_i)||^2 \right) \right)$$

$$\text{where } \mathbf{a}_i = \sum_{l=1}^K w_l^{(i)} p_l \quad \sum_{l=1}^K w_l^{(i)} = 1 \quad w_l^{(i)} \geq 0 \quad \forall l, i$$

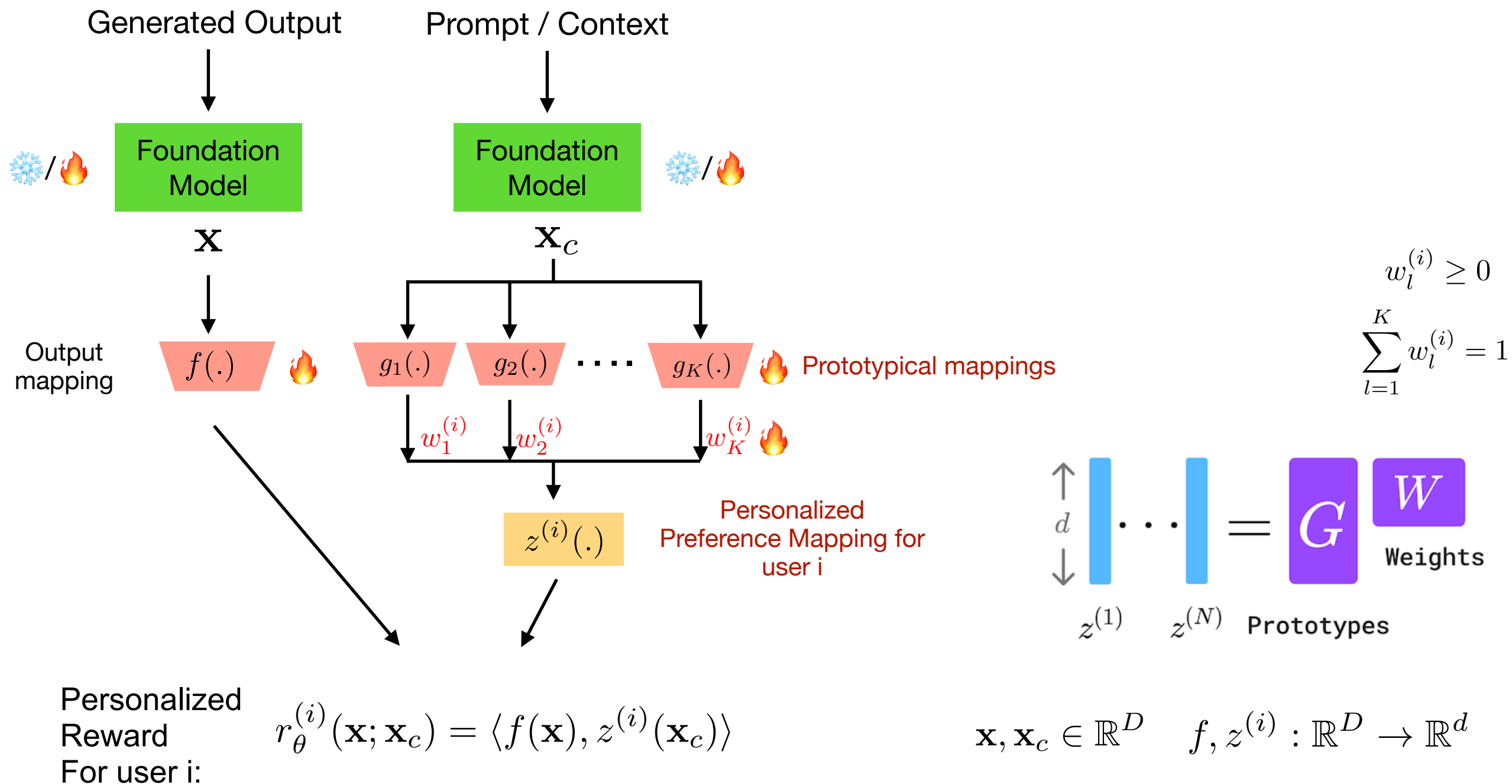
Capturing diverse human preferences via preference mappings



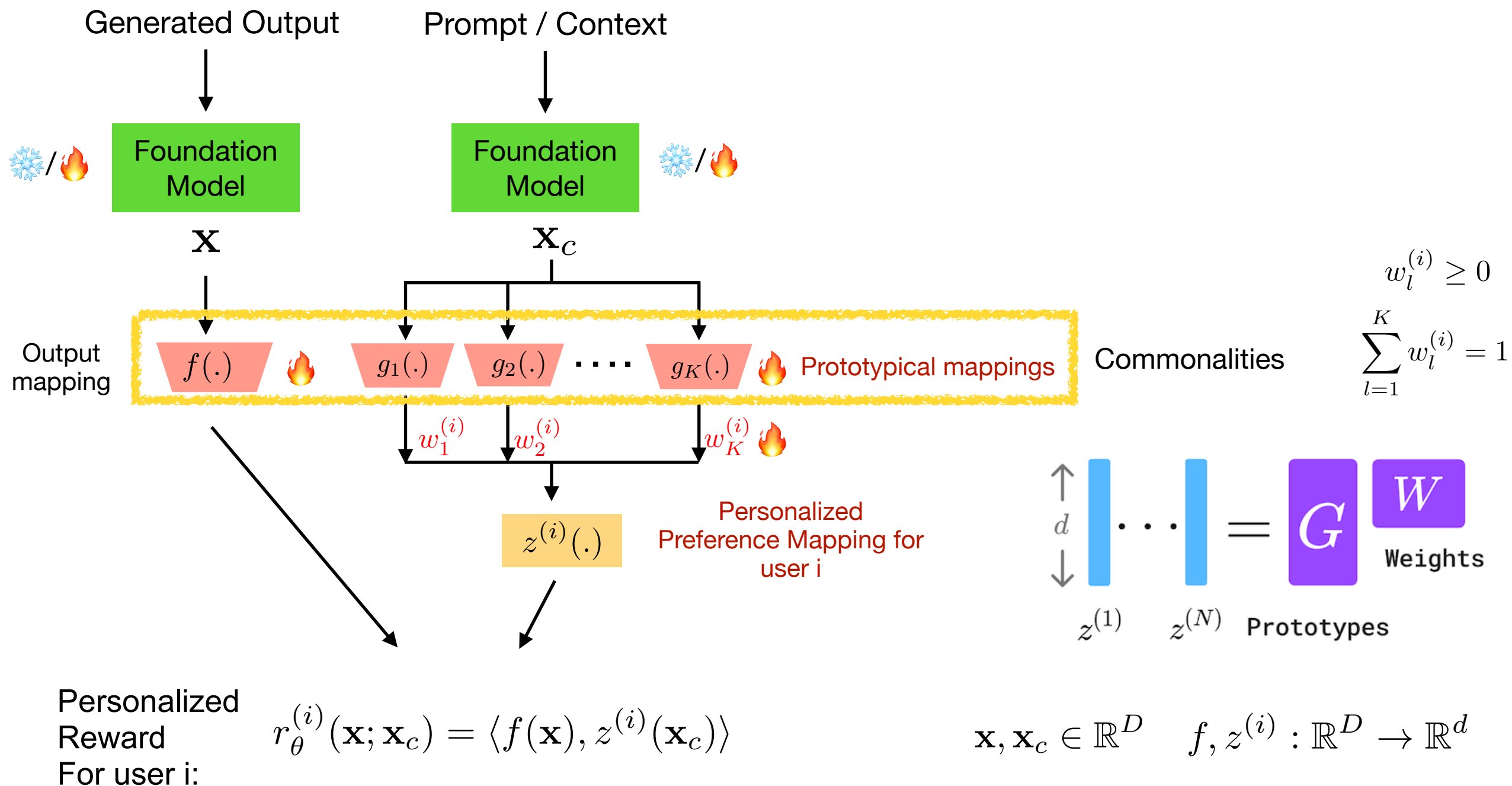
Personalizable Reward Modeling: PAL B



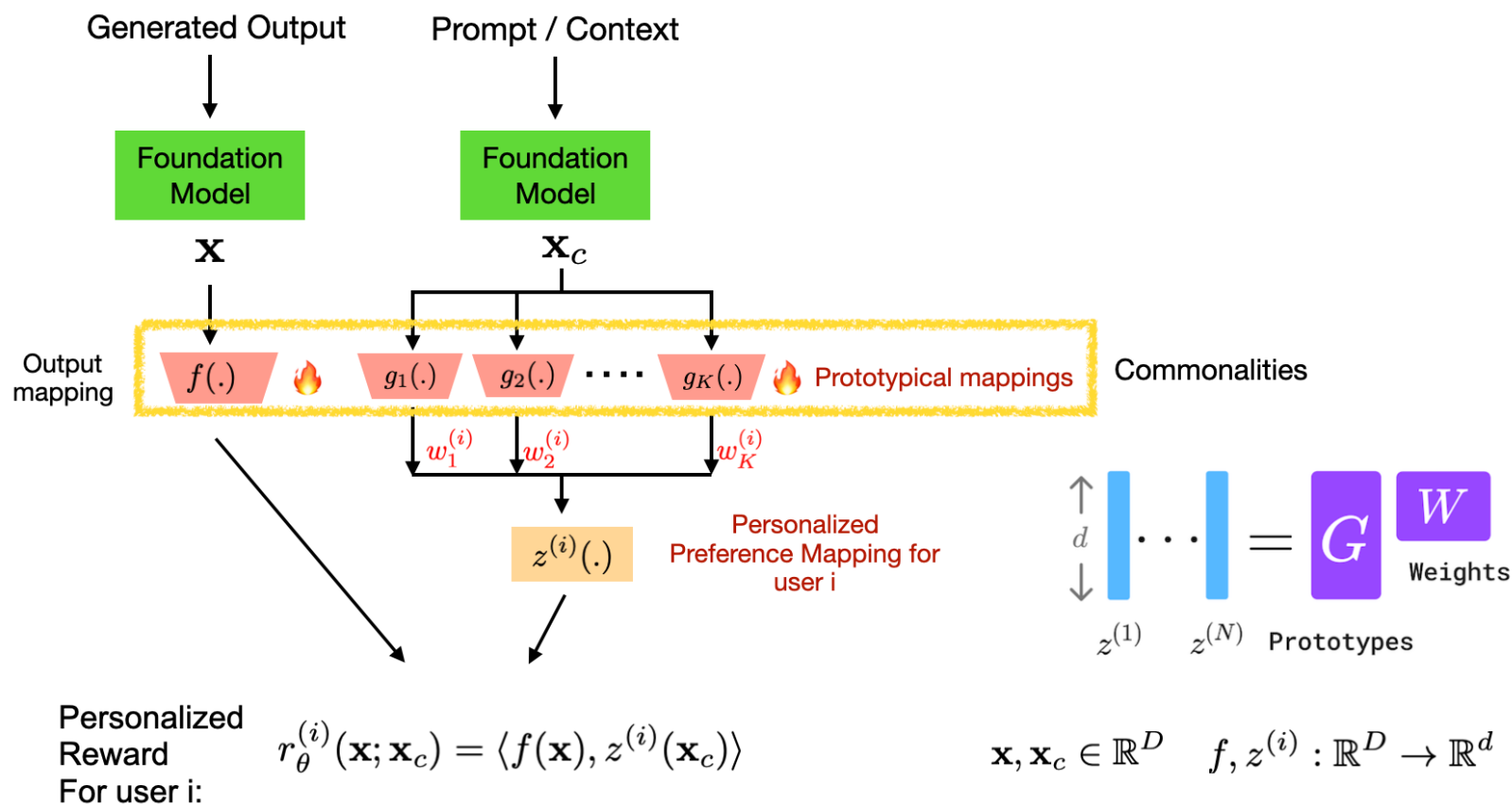
Personalizable Reward Modeling: PAL B



Personalizable Reward Modeling: PAL B



Learning Personalizable Rewards: PAL B



- Leveraging commonalities and low-rank structure in user preferences reduces the samples needed per user

$$\arg \min_{f, g_1, \dots, g_K, w^{(1)}, \dots, w^{(N)}} \sum_i \sum_p \text{loss} \left(y_{p,i} \left(\langle f(\mathbf{x}_{l_{p,i}}), z^{(i)}(\mathbf{x}_{c_{p,i}}) \rangle - \langle f(\mathbf{x}_{r_{p,i}}), z^{(i)}(\mathbf{x}_{c_{p,i}}) \rangle \right) \right)$$

where $\mathbf{a}_i = \sum_{l=1}^K w_l^{(i)} p_l \quad \sum_{l=1}^K w_l^{(i)} = 1 \quad w_l^{(i)} \geq 0 \quad \forall l, i$

Learning Personalizable Rewards: PAL B

Algorithm 2 Learning algorithm for PAL-B

Require: Dataset $\mathcal{D} = \bigcup_{i=1}^N \{(\mathbf{x}_{j,l}^{(i)}, \mathbf{x}_{j,r}^{(i)}; \mathbf{x}_{j,c}^{(i)}, y_j^{(i)})\}_{j=1}^{m_i}$, loss function ℓ , mapping function f_θ , prototype mapping functions $\{g_{\theta_k}\}_{k=1}^K$, user weights $\{\mathbf{w}^{(i)} := [w_1^{(i)}, \dots, w_K^{(i)}]\}_{i=1}^N$.

- 1: **for** each iteration **do**
- 2: **Sample** a mini-batch $\{(\mathbf{x}_{j,l}^{(i)}, \mathbf{x}_{j,r}^{(i)}; \mathbf{x}_{j,c}^{(i)}, y_j^{(i)})\}$ ▷ random pairs, not ordered by users
- 3: **User Ideal Point (conditioned on prompts):**
- 4: $\mathbf{a}^{(i)} = \left[g_{\theta_1}(\mathbf{x}_{c,j}^{(i)}), \dots, g_{\theta_K}(\mathbf{x}_{c,j}^{(i)}) \right]^\top \cdot \mathbf{w}^{(i)}$
- 5: **Distance:**
- 6: $d_{l,j}^{(i)} = \langle f_\theta(\mathbf{x}_{l,j}^{(i)}), \mathbf{a}^{(i)} \rangle, \quad d_{r,j}^{(i)} = \langle f_\theta(\mathbf{x}_{r,j}^{(i)}), \mathbf{a}^{(i)} \rangle$
- 7: **Loss:** $\psi_j^{(i)}(\mathbf{x}_{l,j}^{(i)}, \mathbf{x}_{r,j}^{(i)}; \mathbf{x}_{c,j}^{(i)}, y_j^{(i)}) = \ell(y_j^{(i)} \cdot (d_{l,j}^{(i)} - d_{r,j}^{(i)}))$
- 8: **Update Step:** $\operatorname{argmin}_{\theta, \mathbf{P}, \{\mathbf{w}^{(i)}\}_{i=1}^N} \sum \psi_j^{(i)}(\mathbf{x}_{l,j}^{(i)}, \mathbf{x}_{r,j}^{(i)}; \mathbf{x}_{c,j}^{(i)}, y_j^{(i)})$
- 9: **end for**

Outline

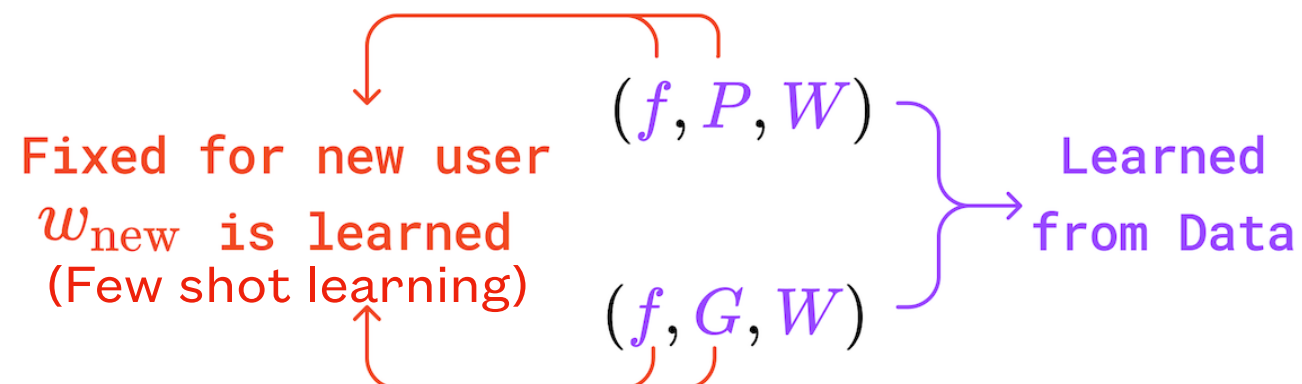
- Motivation: Why should we care about heterogeneity in human preferences?
- Reward modeling for preference learning
- Personalizable reward modeling
- Learning metric and multiple user preferences simultaneously
- Sample efficiency by leveraging commonalities:
Personalizable reward modeling for pluralistic alignment (PAL)
- Experimental results



Generated by ChatGPT

Training and Generalization

- Learn the unknown mapping f and unknown prototypes from pairwise comparisons



Two types of generalization for preference test accuracy:

1. **Seen** accuracy: on users present in train set for unseen pairs
2. **Unseen** accuracy: on users not present in train set

Adapt to Unseen Users

- Part of data to localize users
- Only learn weights W
- Fix mappings and prototypes

Experiments: Pick-a-pic dataset

Text-to-Image generation

Model	Params	v1 Acc (%)
CLIP-H/14	986M	59.23
ImageReward	447M	61.10
HPS	986M	66.70
PickScore	986M	71.85
PAL-A-Tiny	8.4M	69.29 $\pm$ 0.6
PAL-B-Tiny	6.3M	71.13 $\pm$ 0.3

PickScore fine tunes last 20 layers of CLIP-H

Our models: train 2-layer MLP on the CLIP-H embeddings

K = 1

Model	Train Dataset	Test Accuracy on Pick-a-Pic v2 (%)	
		No-leakage	Leakage
CLIP-H/14	-	62.57	58.59
PickScore	pickapic v1	68.04	74.16*
PAL-B-Tiny on CLIP-H	pickapic v1	70.02 $\pm$ 0.39	79.32 $\pm$ 1.68*
PAL-B-Tiny on CLIP-H	pickapic v2	70.51 $\pm$ 0.22	68.67 $\pm$ 0.51
PAL-B-Tiny on PickScore	pickapic v2	70.16 $\pm$ 0.19	74.79 $\pm$ 0.13*

Experiments: Summary dataset

Language generation

Model	Seen Acc (%)
Vanilla DPO	58.91
P-DPO Individual	91.04
P-DPO Cluster ($K = 5$)	91.12
PAL-B-Tiny ($K = 2$)	79.54 ± 0.54
PAL-B-Large ($K = 2$)	92.82 ± 0.95

$$r_{\theta}^{(i)}(\mathbf{x}; \mathbf{x}_c) = \langle f(\mathbf{x}), z^{(i)}(\mathbf{x}_c) \rangle \quad \mathbf{x}, \mathbf{x}_c \in \mathbb{R}^D \quad f, z^{(i)} : \mathbb{R}^D \rightarrow \mathbb{R}^d$$

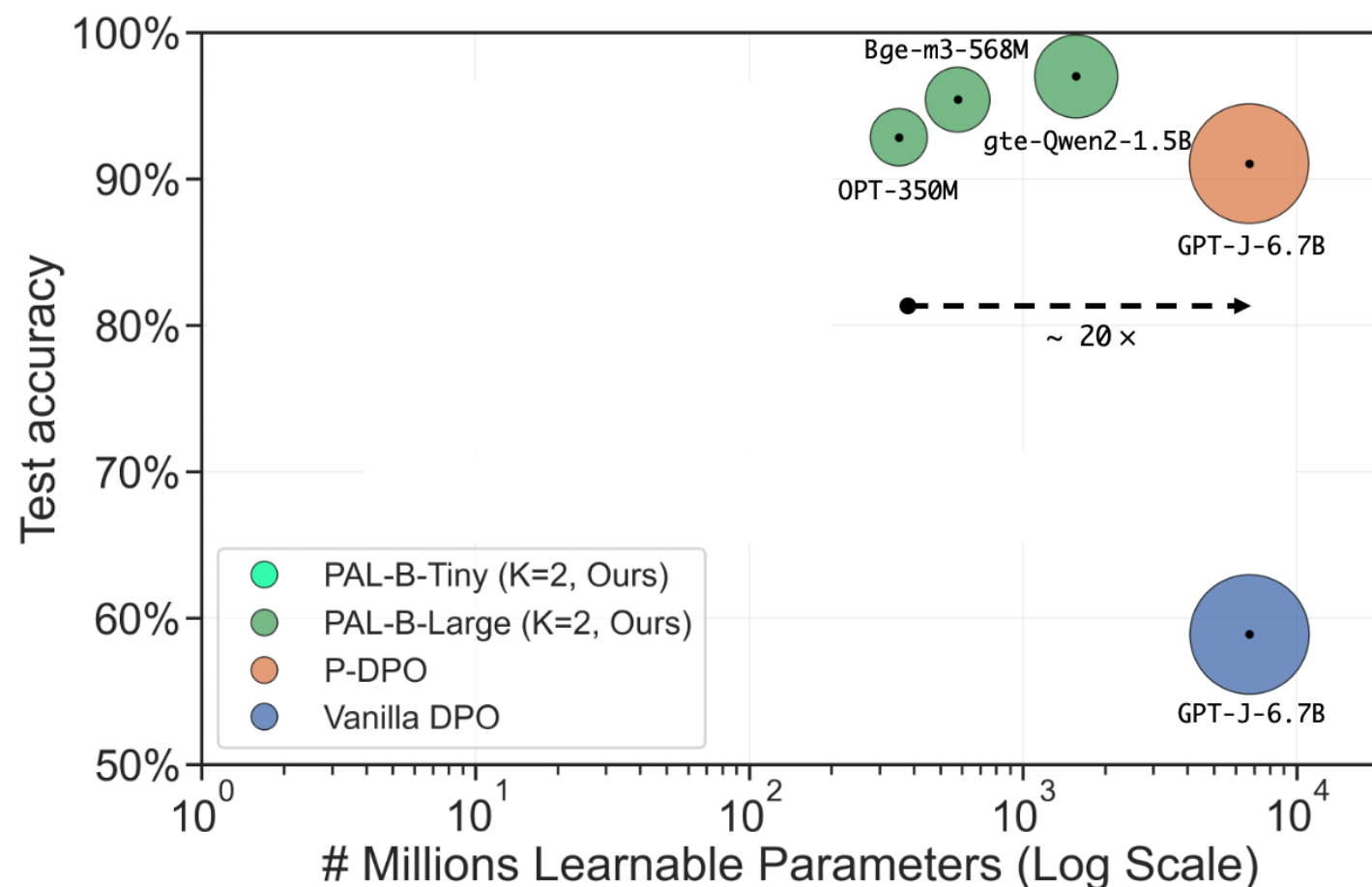
$$z^{(i)}(\mathbf{x}_c) = \sum_{k=1}^K w_k^{(i)} g_k(\mathbf{x}_c) \quad w_k^{(i)} \geq 0 \quad \sum_{k=1}^K w_k^{(i)} = 1$$

Li et. al. uses supervised fine tuning on a *6B-parameter language model*

We train 2-layer MLP (~600K parameters) + a simpler language model for representations (with 350M parameters, OPT-350)

D. Chen, Y. Chen, A. Rege, Z. Wang, R. K. Vinayak, "PAL: Sample-Efficient Personalized Reward Models for Pluralistic Alignment", ICLR 2025

<https://openreview.net/pdf?id=1kFDrYCuSu>



Experiments: Summary dataset

Language generation

Model	Seen Acc (%)
Vanilla DPO	58.91
P-DPO Individual	91.04
P-DPO Cluster ($K = 5$)	91.12
PAL-B-Tiny ($K = 2$)	79.54 ± 0.54
PAL-B-Large ($K = 2$)	92.82 ± 0.95

$$r_{\theta}^{(i)}(\mathbf{x}; \mathbf{x}_c) = \langle f(\mathbf{x}), z^{(i)}(\mathbf{x}_c) \rangle \quad \mathbf{x}, \mathbf{x}_c \in \mathbb{R}^D \quad f, z^{(i)} : \mathbb{R}^D \rightarrow \mathbb{R}^d$$

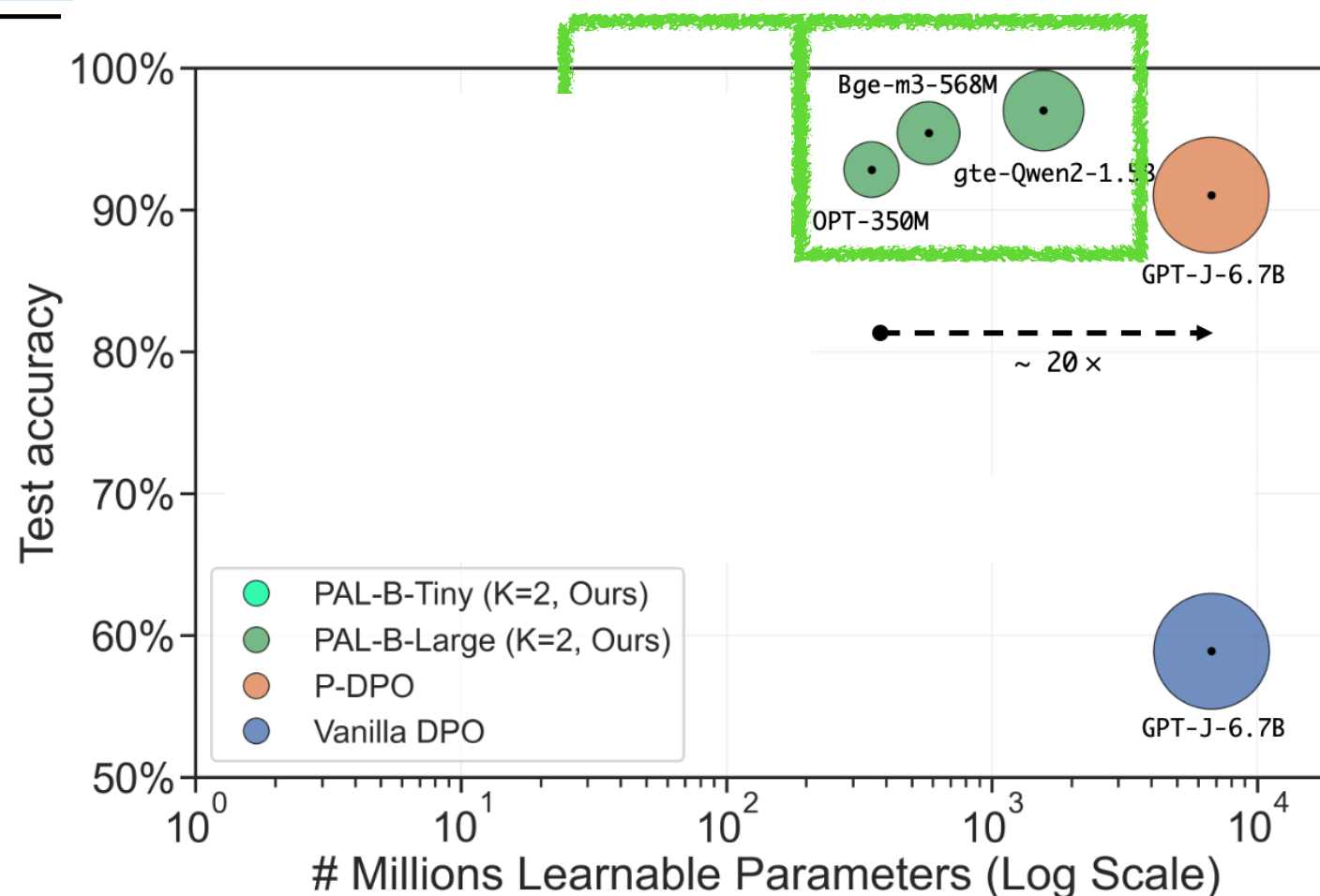
$$z^{(i)}(\mathbf{x}_c) = \sum_{k=1}^K w_k^{(i)} g_k(\mathbf{x}_c) \quad w_k^{(i)} \geq 0 \quad \sum_{k=1}^K w_k^{(i)} = 1$$

Li et. al. uses supervised fine tuning on a *6B-parameter language model*

We train 2-layer MLP (~600K parameters) + a simpler language model for representations (with 350M parameters, OPT-350)

D. Chen, Y. Chen, A. Rege, Z. Wang, R. K. Vinayak, "PAL: Sample-Efficient Personalized Reward Models for Pluralistic Alignment", ICLR 2025

<https://openreview.net/pdf?id=1kFDrYCuSu>



Experiments: Summary dataset

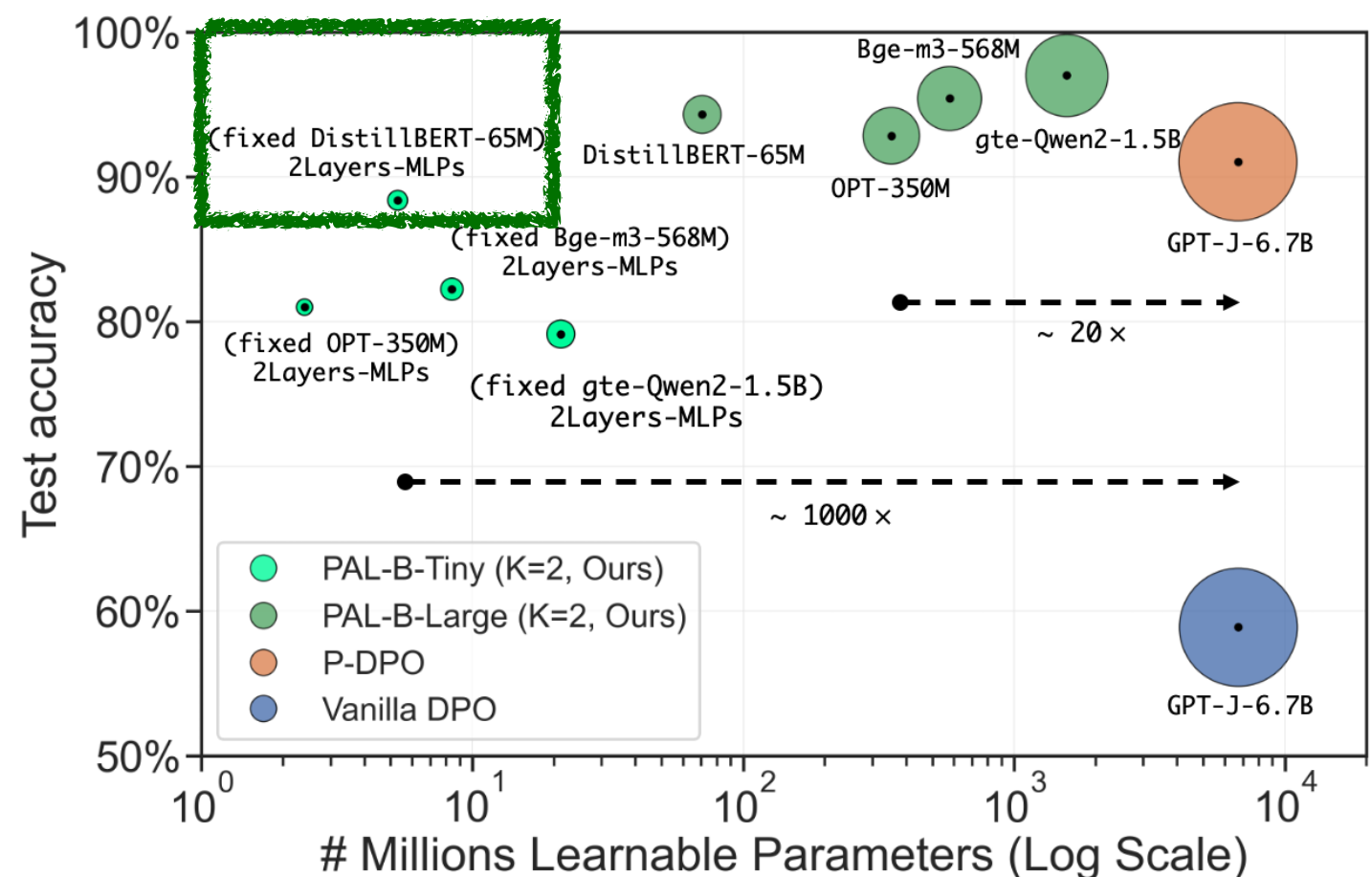
Language generation

Model	Seen Acc (%)
Vanilla DPO	58.91
P-DPO Individual	91.04
P-DPO Cluster ($K = 5$)	91.12
PAL-B-Tiny ($K = 2$)	79.54 ± 0.54
PAL-B-Large ($K = 2$)	92.82 ± 0.95

With Fixed DistillBERT-65M, just a 2-layer MLP gets to nearly 90% accuracy!

D. Chen, Y. Chen, A. Rege, Z. Wang, R. K. Vinayak, "PAL: Sample-Efficient Personalized Reward Models for Pluralistic Alignment", ICLR 2025

<https://openreview.net/pdf?id=1kFDrYCuSu>



Experiments: Summary dataset

Language generation

Model	Seen Acc (%)
Vanilla DPO	58.91
P-DPO Individual	91.04
P-DPO Cluster ($K = 5$)	91.12
PAL-B-Tiny ($K = 2$)	79.54 ± 0.54
PAL-B-Large ($K = 2$)	92.82 ± 0.95

Generalize to new users via few shot learning

Unseen user generalization

Methods	Test accuracy (Unseen user)
P-DPO individual	55.34%
P-DPO K=5	54.55%
Vanilla DPO	55.37%
PAL-B K=2 N=10	$91.63\% \pm 1.43\%$
PAL-B K=2 N=20	$91.72\% \pm 1.40\%$
PAL-B K=2 N=50	$92.02\% \pm 1.10\%$
PAL-B K=2 N=100	$91.97\% \pm 1.91\%$

Advantages of our modeling approach

- **Sample-Efficient:** can be learned with few samples per user
- **Modular design:** leverages commonalities while adapting to individual preferences,
- **Versatile:** can be applied to different modalities, e.g., text-to-image, text-to-text, classical recommendation systems
- **Few-shot generalization:** can be adapted to new users with few samples
- **High-performance with fewer parameters:** we obtain SoTA results with 20x to 100x fewer parameters

Outline

- Motivation: Why should we care about heterogeneity in human preferences?
- Reward modeling for preference learning
- Personalizable reward modeling
- Learning metric and multiple user preferences simultaneously
- Sample efficiency by leveraging commonalities:
Personalizable reward modeling for pluralistic alignment (PAL)
- Experimental results
- Future directions



Generated by ChatGPT

Limitations of current datasets

Currently available open datasets for alignment or reward learning suffer from severe limitations for learning under heterogeneity

- Many datasets do not contain anonymized ids
- Collected using a clear rubric on how to evaluate the outputs: so these are crowdsourcing the researchers preference rather than collecting people's preferences
- Highly imbalanced in terms of number of examples per user:
A small subset of users provide a lot of the comparison data while a large set of people provide very few comparisons

Need better datasets and benchmarks!

Future directions

PAL-B

$$\begin{aligned} r_{\theta}^{(i)}(\mathbf{x}, \mathbf{x}_c) \\ = \sum_{l=1}^K w_l^{(i)} \langle f(\mathbf{x}), g_l(\mathbf{x}_c) \rangle \end{aligned}$$

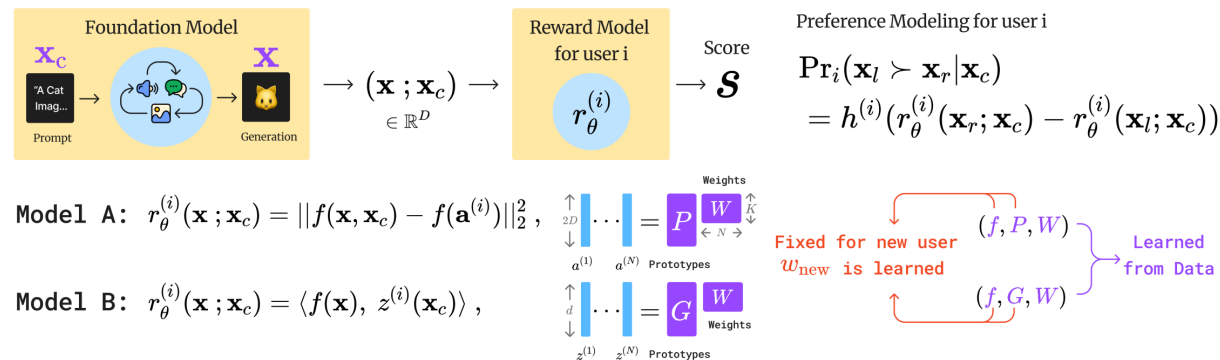
- More general approaches for reward modeling

$$r_{\theta}^{(i)}(\mathbf{x}, \mathbf{x}_c) = \sum_{l=1}^K w_l^{(i)} r_{\theta_l}(\mathbf{x}, \mathbf{x}_c)$$

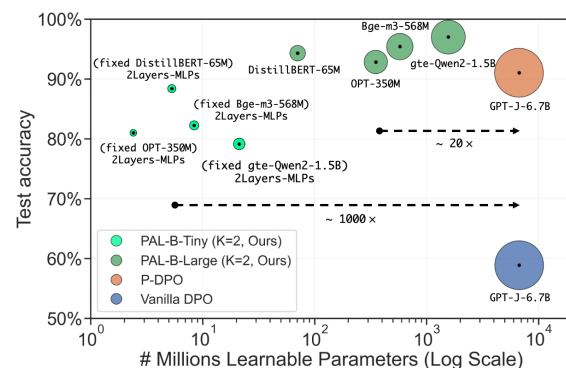
Chen et al. 2024, Ponder et al. 2024, Wang et al. 2024, Bose et al. 2025,

- Need better datasets and benchmarks!
- Adapting on the fly
 - Z. Wang, C. Zhang, R. K. Vinayak “Bridging Lifelong and Multi-Task Representation Learning via Algorithm and Complexity Measure”, Algorithmic Learning Theory (ALT) 2026
- Need better understanding of generalization
 - OOD generalization for users
 - OOD generalization for prompts
- Global safety constraints + Personal preferences
- Conversations, dialogues and interactive collaborations with AI

Sample-Efficient Personalized Reward Modeling for Pluralistic Alignment



- Modular
- Versatile
- Sample-Efficient



Thank you!
Questions?
 ramya@ece.wisc.edu
 @ramyavinayak

<https://github.com/RamyaLab/pluralistic-alignment>



- D. Chen, Y. Chen, A. Rege, Z. Wang, R. K. Vinayak, "PAL: Sample-Efficient Personalized Reward Modeling for Pluralistic Alignment", ICLR 2025
- Canal, Mason, Vinayak, Nowak, "One for all: Simultaneous metric and preference learning over multiple users", Neurips 2022
- Wang, So, Vinayak, "Metric Learning from Limited Pairwise Preference Comparisons", UAI 2024, arXiv:2403.19629, 2024.
- G. Tatli, Y. Chen, R. K. Vinayak "Learning Population of Preferences via Pairwise Comparison Queries", AISTATS 2024
- Z. Wang, C. Chen, R. K. Vinayak "Bridging Lifelong Learning and Multitask Representation Learning via Algorithmic and Complexity Measure", ALT 2026