

A Markov Chain Approach to Preference Alignment

Takuya Koriyama Tengyuan Liang

July 4, 2026

The University of Chicago

Alignment is central to modern LLM fine-tuning

- Modern LLMs are first trained to predict the next token on large-scale text corpora.
- To make them useful assistants, we further fine-tune them to align with human preferences.
- A common source of supervision is *pairwise human preference*:

response A $\prec$ response B “B is preferred to A.”



Preference alignment

- Let μ_{ref} be a reference distribution (e.g., pre-trained LLM)
- Given a dataset of pairwise human preference,

$$A \prec B$$

$$B \prec C$$

$$D \prec B$$

$$E \prec A$$

$$\vdots$$

we aim to fine-tune μ_{ref} .

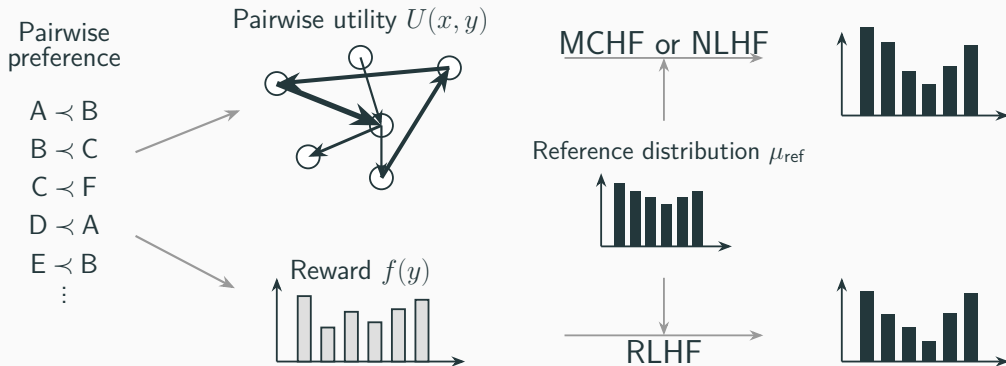
Existing approaches

- Reinforcement Learning from Human Feedback (RLHF) (e.g. Christiano et al. (2017))
- Nash Learning from Human Feedback (NLHF) (Munos et al. (2024))

Our proposal

- Markov Chain from Human Feedback (MCHF)

Overview



Reinforcement Learning from Human Feedback (RLHF)

Step 1: Reward modeling

- Let $\mathcal{P}(x \prec y)$ be the probability that a human prefers y to x .
- Assume BTL¹ model: $\exists$ reward function f such that

$$\mathcal{P}(x \prec y) = \text{sigmoid}(f(y) - f(x)).$$

- Given pairwise preference data $\{x_i \prec y_i\}_{i=1}^n$, we estimate f by the MLE $\hat{f}$:

¹Bradley–Terry (1952) or Luce (1959) Model

Reinforcement Learning from Human Feedback (RLHF)

Step 2: KL regularized distributional optimization

- Given a reward estimate $\hat{f}$, solve

$$\mu_{\text{RL}} = \arg \max_{\mu} \int \mu(dy) \hat{f}(y) - \text{KL}(\mu \mid \mu_{\text{ref}})$$

- The solution has an explicit formula:

$$\mu_{\text{RL}}(dy) = \frac{\exp(\hat{f}(y)) \mu_{\text{ref}}(dy)}{\int \exp(\hat{f}(z)) \mu_{\text{ref}}(dz)} \propto \exp(\hat{f}(z)) \mu_{\text{ref}}(dz)$$

but sampling from it directly is not feasible.

- In practice, μ is parameterized by a generative model μ_{θ} (e.g., diffusion model, transformer), and we optimize θ by gradient descent using RL algorithms.

Reward modeling enforces transitive preference

- Given a human preference $\mathcal{P}(\cdot \prec \cdot)$, define the average order $\prec_*$ as

$$x \prec_* y \quad \text{if and only if} \quad \mathcal{P}(x \prec y) > \frac{1}{2}$$

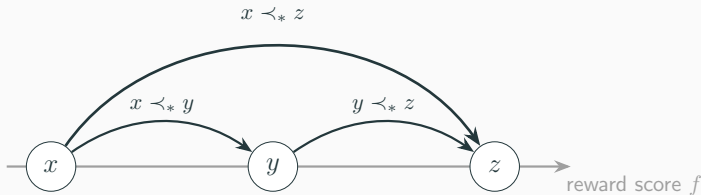
- Recall the BTL model assumes

$$\mathcal{P}(x \prec y) = \text{sigmoid}(f(y) - f(x)) = \frac{e^{f(y)}}{e^{f(y)} + e^{f(x)}},$$

so that

$$x \prec_* y \quad \Leftrightarrow \quad f(x) < f(y)$$

- Thus, the resulting average order $\prec_*$ must be transitive:



Non-transitive preference: Condorcet Paradox (1785)

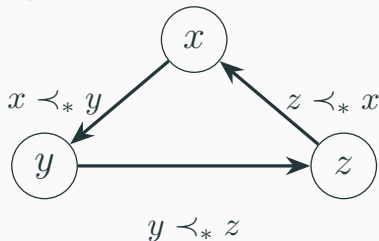
- Consider three outputs x, y, z and three annotators:

Annotator	Preference ranking
1	$x \prec y \prec z$
2	$y \prec z \prec x$
3	$z \prec x \prec y$

- Then,

$$\mathcal{P}(x \prec y) = \mathcal{P}(y \prec z) = \mathcal{P}(z \prec x) = \frac{2}{3} > \frac{1}{2}$$

- The resulting average order $\prec_*$ is non-transitive:



Step 1: Pairwise utility modeling

- Instead of the BTL model

$$\mathcal{P}(x \prec y) = \text{sigmoid}(-f(x) + f(y)),$$

we assume $\exists$ antisymmetric pairwise utility U such that

$$\mathcal{P}(x \prec y) = \text{sigmoid}(U(x, y))$$

and estimate U by the MLE

Nash Learning from Human Feedback (NLHF)

Step 2: Given an antisymmetric pairwise utility U , solve

$$\mu_{\text{NL}} \in \arg \max_{\mu} \min_{\nu} F(\mu, \nu),$$

$$\text{where } F(\mu, \nu) \equiv \iint U(x, y) \nu(dx) \mu(dy) - \text{KL}(\mu \mid \mu_{\text{ref}}) + \text{KL}(\nu \mid \mu_{\text{ref}})$$

- Since U is antisymmetric, F is also antisymmetric: $F(\mu, \nu) = -F(\nu, \mu)$
- By standard minimax theorem, μ_{NL} is a unique solution satisfying

$$\mu_{\text{NL}} = \arg \max_{\mu} F(\mu, \nu = \mu_{\text{NL}})$$

- In practice, we solve μ_{NL} by iterations: $\mu^{t+1} = \arg \max_{\mu} F(\mu, \nu = \mu^t)$ with $\mu^0 = \mu_{\text{ref}}$

Nash Learning from Human Feedback (NLHF)

$$\mu_{\text{NL}} \in \arg \max_{\mu} \min_{\nu} \left\{ \iint U(x, y) \nu(dx) \mu(dy) - \text{KL}(\mu \mid \mu_{\text{ref}}) + \text{KL}(\nu \mid \mu_{\text{ref}}) \right\}.$$

- NLHF formulates alignment as a simultaneous game; μ and ν compete at the same time
- However, alignment can be naturally defined as a non-simultaneous game:

Given x , choose $y = \arg \max_y U(x, y)$

- However, we have to pick y such that the marginal distribution of y is close to μ_{ref}
- Motivated by these points, we propose a Markov chain approach

Our contribution 1: a Markov chain approach

- Given μ_{ref} and a pairwise utility $U(x, y)$, define the Markov kernel $P(x, dy)$:

$$P(x, dy) = \frac{1}{Z(x)} \exp(U(x, y)) \mu_{\text{ref}}(dy)$$

where $Z(x)$ is a normalizing constant s.t. $\int P(x, dy) = 1$ for all x .

- Given a current distribution μ , we align it via conditional sampling

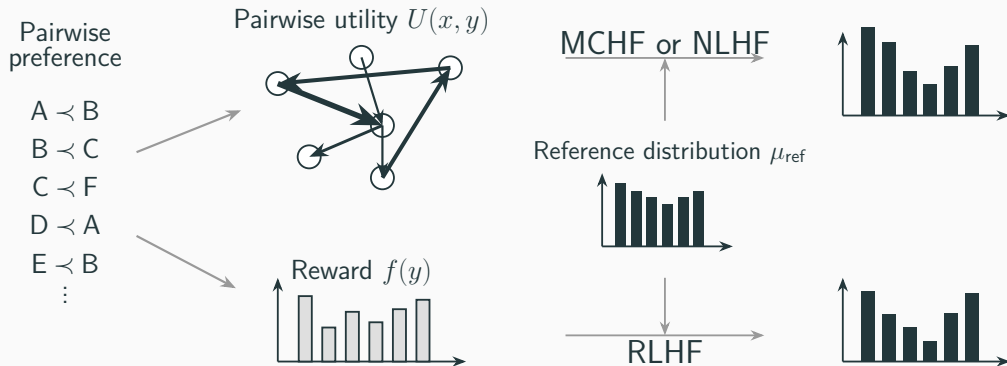
$$\mu_{\text{new}}(dy) = \mu P(dy) = \int \mu(dx) P(x, dy)$$

- The Markov kernel then induces an iterative alignment

$$\mu_{\text{ref}} \mapsto \mu_{\text{ref}} P \mapsto \cdots \mu_{\text{ref}} P^\infty \equiv \mu_{\text{MC}},$$

where μ_{MC} is a stationary distribution of P (inspired by PageRank)

Meta comparison of MCHF, NLHF, RLHF



Meta comparison of NLHF and MCHF via implementation perspective

Consider the two algorithms:

$$\text{NLHF : } \mu_{\text{NL}}^{t+1} = \arg \max_{\mu} F(\mu, \nu = \mu_{\text{NL}}^t)$$

$$\text{MCHF : } \mu_{\text{MC}}^{t+1} = \mu_{\text{MC}}^t \text{P}$$

- NLHF repeatedly solves a RLHF solution with iterate-dependent reward

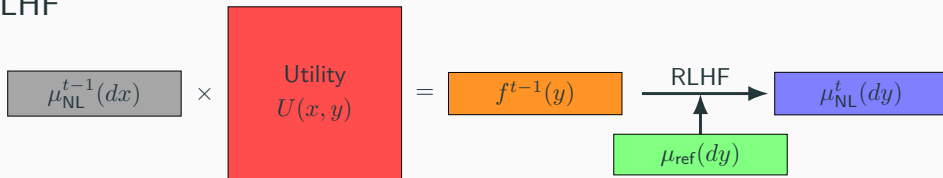
$$\mu_{\text{NL}}^{t+1} = \arg \max_{\mu} \int f_{\text{NL}}^t(y) \mu(dy) - \text{KL}(\mu \mid \mu_{\text{ref}}), \quad f_{\text{NL}}^t(y) = \int U(x, y) \mu_{\text{NL}}^t(dx)$$

- MCHF repeatedly applies one fixed conditional sampler:

$$\text{P}(x, dy) \propto \exp(U(x, y)) \mu_{\text{ref}}(dy)$$

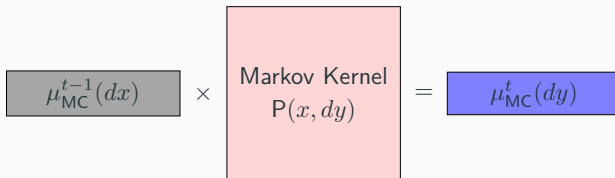
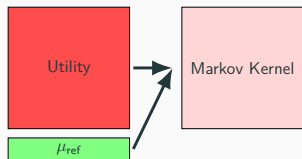
Meta comparison of NLHF and MCHF via implementation perspective

NLHF



MCHF

(Construct the Markov kernel at the beginning)



Our contributions 2: Unify MCHF, NLHF, and RLHF

Three methods are based on different premises:

- RLHF μ_{RL} : Reward-based approach
- NLHF μ_{NL} : Game-theoretic approach
- MCHF μ_{MC} : Markovian approach

Q: how can we compare these methods in a mathematically rigorous framework?

MCHF and RLHF reduce to RLHF when U is separable

Proposition

Let $\mu_{\text{MC}}(U)$ and $\mu_{\text{NL}}(U)$ be the MCHF/NLHF equilibria based on U . Then,

$$\left(\exists f, U(x, y) = -f(x) + f(y) \right) \Rightarrow \left(\mu_{\text{MC}}(U) = \mu_{\text{NL}}(U) = \mu_{\text{RL}}(f) \right)$$

where μ_{RL} is the RLHF solution with reward f :

$$\mu_{\text{RL}}(f)(dy) \propto \exp(f(y)) \mu_{\text{ref}}(dy)$$

- Recall

$$\mathcal{P}(x \prec y) = \text{sigmoid}(U(x, y))$$

so $U(x, y) = -f(x) + f(y)$ means we impose BTL assumption.

Our contributions 2: Unify MCHF, NLHF, and RLHF

- Q. What if $U(x, y) \approx -f(x) + f(y)$ (but not necessarily in equality)?
- We address this question by introducing a seminorm

$$\|U\|_{\oplus} = \inf_{g, f \in L^{\infty}(\mu_{\text{ref}})} \|U - g \oplus f\|_{\infty}, \quad \text{where} \quad (g \oplus f)(x, y) = g(x) + f(y)$$

- When $\|U\|_{\oplus}$ is small, there exists a common correction term such that

$$\text{for } * \in \{\text{MC}, \text{NL}\}, \quad \mu_*(U) \approx \mu_{\text{RL}}(\hat{f}) + \text{Correction},$$

where

$$\hat{f}(y) = \int \mu_{\text{ref}}(dx) U(x, y), \quad \text{Correction} = \text{a Linear functional of } U - (-\hat{f}) \oplus \hat{f}$$

1. Definition of seminorm $\|U\|_{\oplus}$ and some basic properties
2. Refined convergence analysis for MCHF and NLHF
The smaller $\|U\|_{\oplus}$, the faster MCHF and NLHF iterations converge to their own equilibria
3. Asymptotic analysis of MCHF and NLHF when $\|U\|_{\oplus} \rightarrow 0$.

**Definition of seminorm:
measuring non-transitive
structure of pairwise preference**

- Fix a reference measure μ_{ref}
- Assume that the pairwise utility U encodes human preference as

$$U(x, y) > 0 \Leftrightarrow \text{Humans prefer } y \text{ to } x$$

- Throughout this talk, U is given, and we assume $U \in L^\infty(\mu_{\text{ref}} \otimes \mu_{\text{ref}})$.

Definition of the seminorm $\|\cdot\|_{\oplus}$

Definition

Define the seminorm $\|\cdot\|_{\oplus}$ on $L^{\infty}(\mu_{\text{ref}} \otimes \mu_{\text{ref}})$ as

$$\|U\|_{\oplus} = \inf_{g, f \in L^{\infty}(\mu_{\text{ref}})} \|U - g \oplus f\|_{\infty} \quad \text{where} \quad (g \oplus f)(x, y) = g(x) + f(y)$$

- $\|U\|_{\oplus} \leq \|U\|_{\infty}$ by the definition

Proposition

Suppose U is antisymmetric. Then,

$$\|U\|_{\oplus} \asymp \|\Delta U\|_{\infty},$$

where

$$\Delta U = U - (-\hat{f}) \oplus \hat{f}, \quad \hat{f}(y) = \int \mu_{\text{ref}}(dx) U(x, y).$$

Furthermore,

$$\|U\|_{\oplus} \asymp \operatorname{ess\,sup}_{x,y,z} |U(x, y) + U(y, z) + U(z, x)|$$

- $A \asymp B$ means $A \leq c_1 \cdot B$ and $c_2 \cdot A \leq B$ for absolute positive constants c_1, c_2 .
- $\operatorname{ess\,sup}_{x,y,z}$ is the essential supremum with respect to $\mu_{\text{ref}} \otimes \mu_{\text{ref}} \otimes \mu_{\text{ref}}$.

Refined convergence analysis for MCHF and NLHF

Theorem

Let P be the Markov kernel based on U :

$$P(x, dy) = \frac{1}{Z(x)} \exp(U(x, y)) \mu_{\text{ref}}(dy)$$

Then, there exists a unique stationary distribution

$$\mu_{\text{MC}} = \mu_{\text{MC}} P$$

and

$$\forall t \geq 1, \quad d_{\text{TV}}(\mu_{\text{ref}} P^t, \mu_{\text{MC}}) \leq c(\|U\|_{\oplus})^t \cdot d_{\text{TV}}(\mu_{\text{ref}}, \mu_{\text{MC}})$$

where

$$c(r) = \min(r, 1 - e^{-2r}) \in [0, 1).$$

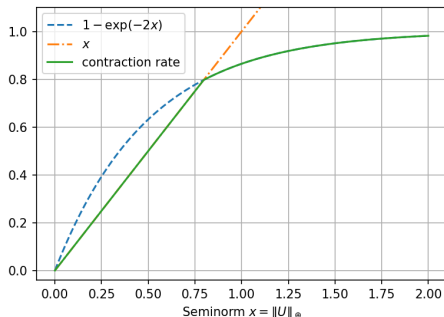
Refined convergence analysis for MCHF

$$\forall t \geq 1, \quad d_{\text{TV}}(\mu_{\text{ref}} \mathbf{P}^t, \mu_{\text{MC}}) \leq c(\|U\|_{\oplus})^t \cdot d_{\text{TV}}(\mu_{\text{ref}}, \mu_{\text{MC}})$$

where

$$c(r) = \min(r, 1 - e^{-2r}) \in [0, 1).$$

- $c(r)$ is increasing in $r \Rightarrow$ the smaller $\|U\|_{\oplus}$ is, the faster the Markov chain converges



Theorem

Consider the mirror-descent iteration

$$\begin{aligned}\nu_t &= \arg \min_{\nu} \left\{ \iint U(x, y) \nu(dx) \mu_t(dy) + \text{KL}(\nu \mid \mu_{\text{ref}}) \right\}, \\ \mu_{t+1} &= \arg \max_{\mu} \left\{ \iint U(x, y) \nu_t(dx) \mu(dy) - \text{KL}(\mu \mid \mu_{\text{ref}}) - \eta^{-1} \text{KL}(\mu \mid \mu_t) \right\}.\end{aligned}$$

If $0 < \eta \leq \|U\|_{\oplus}^{-2}$, then

$$\text{KL}(\mu_{\text{NL}} \mid \mu_t) \leq \frac{1}{(1 + \eta)^t} \text{KL}(\mu_{\text{NL}} \mid \mu_{\text{ref}}).$$

- Previous works (e.g. Munos et al. (2024)) allow η at most $\eta = O(\|U\|_{\infty}^{-2})$,
- Since $\|U\|_{\oplus} \leq \|U\|_{\infty}$, our work provides refined convergence analysis.

Small- $\|U\|_{\oplus}$ asymptotics

Small- $\|U\|_{\oplus}$ perturbation: equilibrium level

Theorem

For each U , let p_{MC} and p_{NL} be the densities of MCHF and NLHF based on U ,

$$p_*(\cdot) = \frac{d\mu_*(U)}{d\mu_{\text{ref}}}(\cdot) \quad \text{for } * \in \{\text{MC}, \text{NL}\},$$

and let p_{RL} be the density of RLHF based on the reward $\hat{f}$,

$$p_{\text{RL}}(\cdot) = \frac{\exp(\hat{f}(\cdot))}{\int \mu_{\text{ref}}(dz) \exp(\hat{f}(z))}, \quad \text{where } \hat{f}(y) = \int U(x, y) \mu_{\text{ref}}(dx)$$

and let $\Delta U = U - (-\hat{f}) \oplus \hat{f}$. Then,

$$\text{for } * \in \{\text{MC}, \text{NL}\}, \quad \|p_* - p_{\text{RL}} - p_{\text{RL}} \odot \Delta U^* p_{\text{RL}}\|_1 = o(\|U\|_{\oplus})$$

where $\odot$ is the entrywise product and $\Delta U^* p(\cdot) = \int \mu_{\text{ref}}(dx) \Delta U(x, \cdot) p(x)$.

Small- $\|U\|_{\oplus}$ perturbation: equilibrium level

$$\text{for } * \in \{\text{MC}, \text{NL}\}, \quad \|p_* - p_{\text{RL}} - p_{\text{RL}} \odot \Delta U^* p_{\text{RL}}\|_1 = o(\|U\|_{\oplus})$$

- MCHF and NLHF have the same correction term. Thus, by the triangle inequality,

$$\|p_{\text{MC}} - p_{\text{NL}}\|_1 = o(\|U\|_{\oplus})$$

- MCHF and NLHF both incorporate the non-transitive component of U :

$$\Delta U = U - (-\hat{f}) \oplus \hat{f}$$

on top of the RLHF density p_{RL} based on the reward $\hat{f}(y) = \int U(x, y) \mu_{\text{ref}}(dx)$

Theorem

Let p_{MC}^t and p_{NL}^t be the densities generated by the algorithms for MCHF and NLHF defined by:

$$\mu_{\text{MC}}^t = \mu_{\text{MC}}^{t-1} \mathbf{P}, \quad \mu_{\text{NL}}^t = \arg \max_{\mu} \left\{ \iint U(x, y) \mu_{\text{NL}}^{t-1}(dx) \mu(dy) - \text{KL}(\mu \mid \mu_{\text{ref}}) \right\}$$

initialized at $\mu_{\text{MC}}^0 = \mu_{\text{NL}}^0 = \mu_{\text{ref}}$. Then, for any antisymmetric U such that $\|U\|_{\oplus} \rightarrow 0$,

$$\begin{aligned} \|p_*^1 - p_{\text{RL}}\|_1 &= o(\|U\|_{\oplus}), \\ \sup_{t \geq 2} \|p_*^t - p_{\text{RL}} - p_{\text{RL}} \odot \Delta U^* p_{\text{RL}}\|_1 &= o(\|U\|_{\oplus}). \end{aligned}$$

Small- $\|U\|_{\oplus}$ perturbation: iteration level

$$\text{for } * \in \{\text{MC}, \text{NL}\}, \quad \begin{aligned} \|p_*^1 - p_{\text{RL}}\|_1 &= o(\|U\|_{\oplus}), \\ \sup_{t \geq 2} \|p_*^t - p_{\text{RL}} - p_{\text{RL}} \odot \Delta U^* p_{\text{RL}}\|_1 &= o(\|U\|_{\oplus}). \end{aligned}$$

- **First step:** both MCHF and NLHF learn the RLHF solution with reward

$$\hat{f}(y) = \int U(x, y) \mu_{\text{ref}}(dx).$$

- **Second step onward:** both incorporate the non-transitive residual

$$\Delta U = U - (-\hat{f}) \oplus \hat{f}.$$

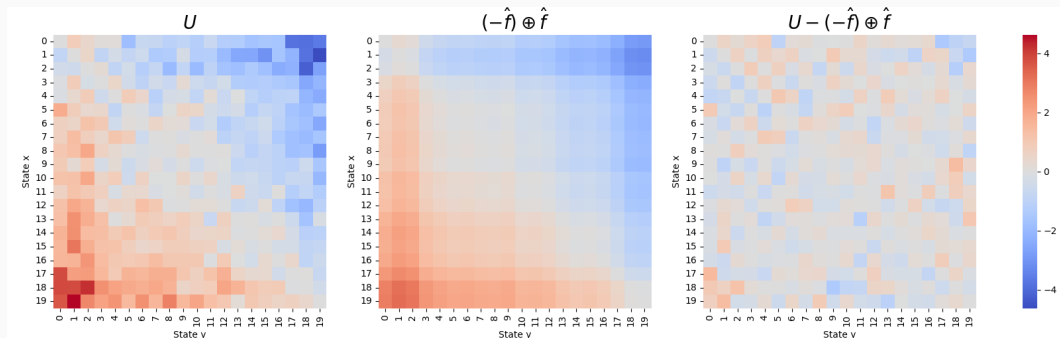
- **Practical implication:** when $\|U\|_{\oplus}$ is small, the main first-order benefit of iterative alignment appears within the first two steps.

Numerical illustration

Set $\mu_{\text{ref}} = \text{Uniform}[n]$ and take U as $U(x, y) = \text{sigmoid}^{-1}(\mathcal{P}(x \prec y))$, where

$$\mathcal{P}(x \prec y) = \Phi\left(R(y) - R(x) + E(x, y) - E(y, x)\right),$$

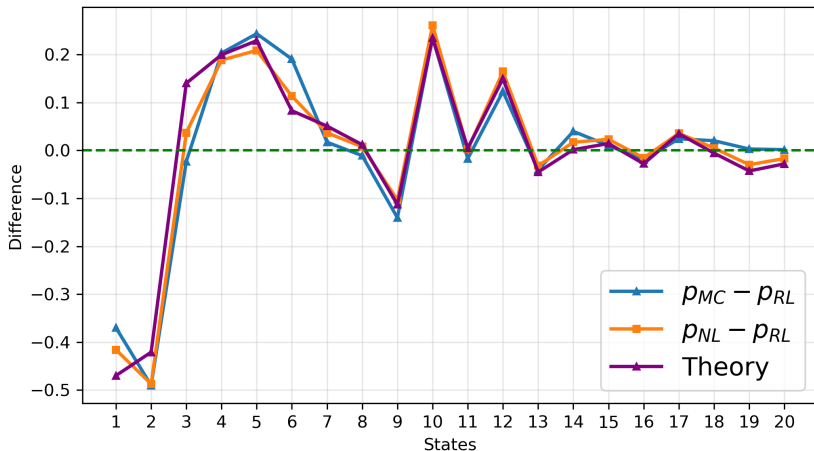
where Φ is the CDF of $N(0, 1)$, $R \sim N(0, I_n)$, and $E \sim N(0, I_{n \times n})$.



Numerical simulation

Based on the previous U , verify

$$\forall * \in \{\text{MC}, \text{NL}\}, \quad \|p_* - p_{\text{RL}} - p_{\text{RL}} \odot \Delta U^* p_{\text{RL}}\|_1 = o(\|U\|_{\oplus})$$



Conclusion

Takeaways and future work

1. MCHF turns a pairwise utility into a preference-guided Markov kernel.
2. MCHF and NLHF iterations both converge geometrically to their own equilibria fast when $\|U\|_{\oplus}$ is small.
3. In the small- $\|U\|_{\oplus}$ asymptotics, MCHF and NLHF share the same first-order expansion around RLHF.

Future work

1. Practical implementation of the conditional sampler $P(x, dy)$
2. Calculate $\|U\|_{\oplus}$ in real data

Paper: *A Markov Chain Approach to Preference Alignment*

arXiv:2606.22652

- Munos et al. (2024). Nash Learning from Human Feedback.
- Christiano et al. (2017). Deep reinforcement learning from human preferences.

Backup slides

Benefit of MCHF over NLHF: inference time refinement

Consider the Condorcet Paradox example

Annotator	Preference ranking
1	$x \prec y \prec z$
2	$y \prec z \prec x$
3	$z \prec x \prec y$

and construct U as the logit transform of average preference. Then, U can be written as

$$U = c \begin{bmatrix} 0 & 1 & -1 \\ -1 & 0 & 1 \\ 1 & -1 & 0 \end{bmatrix}$$

for a constant $c > 0$.

Q. Fix a distribution μ and suppose each annotator samples outputs from μ until they get their best output. Then, how should we choose μ from U so that the expected time to get the best output is small?

Benefit of MCHF over NLHF: inference time refinement

Annotator	Preference ranking
1	$x \prec y \prec z$
2	$y \prec z \prec x$
3	$z \prec x \prec y$

$$\implies U = c \begin{bmatrix} 0 & 1 & -1 \\ -1 & 0 & 1 \\ 1 & -1 & 0 \end{bmatrix}$$

- Suppose $\mu_{\text{ref}} = \text{Uniform}\{x, y, z\}$. Then $\mu_{\text{NL}} = \mu_{\text{ref}}$ and hence the hitting time to the best output for each annotator is 3.
- For a fixed $\gamma > 0$, consider the Markov chain

$$\mu_{\text{ref}} \rightarrow \mu_{\text{ref}} P_{\gamma} \rightarrow \mu_{\text{ref}} P_{\gamma}^2 \cdots \quad \text{where} \quad P_{\gamma}(x, dy) \propto \exp(\gamma U(x, y)) \mu_{\text{ref}}(dy).$$

Then, one can show that the hitting time is a strictly decreasing function of $\gamma > 0$.

- We extend this phenomenon to general μ_{ref} and consider the hitting time to a general set A , showing that we can reduce the hitting time if and only if there is enough *energy* $U(x, y)^2$ from A^c to A .

MCHF reduces to RLHF when U takes an additive form

Proof

- Recall that μ_{MC} is the stationary distribution of the Markov kernel:

$$P(x, dy) = \frac{\exp(U(x, y))\mu_{\text{ref}}(dy)}{\int \exp(U(x, z))\mu_{\text{ref}}(dz)}.$$

- If $U(x, y) = -f(x) + f(y)$,

$$P(x, dy) = \frac{\exp(f(y))\mu_{\text{ref}}(dy)}{\int \exp(f(z))\mu_{\text{ref}}(dz)} = \mu_{\text{RL}}(f)(dy)$$

- Thus, the stationary distribution μ_{MC} is $\mu_{\text{RL}}(f)$.

NLHF reduces to RLHF when U takes an additive form

Proof

- Recall

$$\mu_{\text{NL}} \in \arg \max_{\mu} \min_{\nu} F(\mu, \nu), \quad F(\mu, \nu) \equiv \iint U(x, y) \nu(dx) \mu(dy) - \text{KL}(\mu \mid \mu_{\text{ref}}) + \text{KL}(\nu \mid \mu_{\text{ref}})$$

- If $U(x, y) = -f(x) + f(y)$,

$$\begin{aligned} F(\mu, \nu) &= - \int f(x) \nu(dx) + \int f(y) \mu(dy) - \text{KL}(\mu \mid \mu_{\text{ref}}) + \text{KL}(\nu \mid \mu_{\text{ref}}) \\ &= \left(\int f(y) \mu(dy) - \text{KL}(\mu \mid \mu_{\text{ref}}) \right) - \left(\int f(x) \nu(dx) - \text{KL}(\nu \mid \mu_{\text{ref}}) \right) \end{aligned}$$

- Thus,

$$\mu_{\text{NL}} = \arg \max_{\mu} \int f(y) \mu(dy) - \text{KL}(\mu \mid \mu_{\text{ref}}) = \mu_{\text{RL}}(f)$$

When local regularization is removed

Theorem

With $\eta = +\infty$, the NLHF updates satisfy

$$d_{\text{TV}}(\nu_t, \nu_{\text{NL}}) \leq \|U\|_{\oplus} d_{\text{TV}}(\mu_t, \mu_{\text{NL}}), \quad d_{\text{TV}}(\mu_{t+1}, \mu_{\text{NL}}) \leq \|U\|_{\oplus} d_{\text{TV}}(\nu_t, \nu_{\text{NL}}).$$

Therefore, if $\|U\|_{\oplus} < 1$,

$$d_{\text{TV}}(\mu_t, \mu_{\text{NL}}) \leq \|U\|_{\oplus}^{2t} d_{\text{TV}}(\mu_{\text{ref}}, \mu_{\text{NL}}).$$

Proof idea for MCHF contraction

1. Minorization / coupling:

$$\begin{aligned} \exists \text{ prob. measure } \mu_*, \exists \delta > 0 \quad & P(x, dy) \geq \delta \cdot \mu_*(dy) \\ \Rightarrow \quad & d_{\text{TV}}(\mu P, \nu P) \leq (1 - \delta) d_{\text{TV}}(\mu, \nu). \end{aligned}$$

2. Dobrushin coefficient:

$$\sup_{\mu \neq \nu} \frac{d_{\text{TV}}(\mu P, \nu P)}{d_{\text{TV}}(\mu, \nu)} = \text{ess sup}_{x, x'} d_{\text{TV}}(P(x, \cdot), P(x', \cdot)).$$

Backup: stability under utility perturbations

Let $B_{\oplus}^{\infty}(R) = \{U : \|U\|_{\oplus} \leq R\}$. For $* \in \{\text{MC}, \text{NL}\}$, there exists C_R such that

$$\sup_{U, \hat{U} \in B_{\oplus}^{\infty}(R)} \frac{d_{\text{TV}}(\mu_*(U), \mu_*(\hat{U}))}{\|U - \hat{U}\|_{\infty}} \leq C_R.$$

- Robustness matters because the utility is estimated from finite preference data.
- The map $U \mapsto \mu_*(U)$ is stable on bounded $\|\cdot\|_{\oplus}$ balls.

Backup: uniform Fréchet differentiability

Let $p_*(U) = d\mu_*(U)/d\mu_{\text{ref}}$. For $* \in \{\text{MC}, \text{NL}\}$,

$$\lim_{\|E\|_\infty \rightarrow 0} \sup_{\|U\|_\oplus \leq R} \frac{\|p_*(U + E) - p_*(U) - Dp_*(U)[E]\|_1}{\|E\|_\infty} = 0.$$

- This is the technical engine behind the first-order comparison.
- At additive antisymmetric utilities, the MCHF and NLHF derivatives coincide.

Backup: derivative at additive antisymmetric utilities

For $U = -f \oplus f$ and antisymmetric perturbation E ,

$$Dp_*(-f \oplus f)[E] = p_{\text{RL}}(f) \odot E^* p_{\text{RL}}(f), \quad * \in \{\text{MC}, \text{NL}\}.$$

Here

$$p_{\text{RL}}(f)(y) = \frac{\exp(f(y))}{\int \exp(f(y')) \mu_{\text{ref}}(dy')}, \quad E^* h(\cdot) = \int E(x, \cdot) h(x) \mu_{\text{ref}}(dx).$$

Coupling view: MCHF optimizes over a larger class

Suppose μ^0 is a current distribution, and we update it via the following two different rules:

$$\mu_{\text{MC}}^1 = \mu^0 \mathbf{P}, \quad \mu_{\text{NL}}^1 \in \arg \max_{\mu} \int \mu^0(dx) \mu(dy) U(x, y) - \text{KL}(\mu \mid \mu_{\text{ref}})$$

Fix a current distribution μ^0 and consider the couplings induced by MCHF and NLHF

$$\pi_{\text{MC}}(dx, dy) = \mu^0(dx) \mathbf{P}(x, dy), \quad \pi_{\text{NL}}(dx, dy) = \mu^0(dx) \mu_{\text{NL}}^1(dy).$$

Coupling view: MCHF optimizes over a larger class

$$\pi_{\text{MC}}(dx, dy) = \mu^0(dx)P(x, dy), \quad \pi_{\text{NL}}(dx, dy) = \mu^0(dx)\mu_{\text{NL}}^1(dy).$$

The two couplings π_{MC} and π_{NL} maximize the same objective function:

$$F(\pi) = \iint U(x, y)\pi(dx, dy) - \text{KL}(\pi || \mu^0 \otimes \mu_{\text{ref}}),$$

but over different feasible spaces of couplings:

$$\pi_{\text{MC}} = \arg \max_{\int \pi(\cdot, dy) = \mu^0(\cdot)} F(\pi), \quad \pi_{\text{NL}} = \arg \max_{\exists \mu, \pi = \mu^0 \otimes \mu} F(\pi).$$

General hitting-time analysis

For $\gamma \geq 0$, define

$$P_\gamma(x, dy) = \frac{\exp(\gamma U(x, y)) \mu_{\text{ref}}(dy)}{\int \exp(\gamma U(x, z)) \mu_{\text{ref}}(dz)}.$$

For a target set A , let $T_\gamma(A)$ be the expected time to hit A starting from μ_{ref} . At $\gamma = 0$,

$$T_0(A) = \frac{1}{\mu_{\text{ref}}(A)}.$$

Theorem: First derivative

For antisymmetric U ,

$$\dot{T}_0(A) = -\frac{\mu_{\text{ref}}(A^c)}{\mu_{\text{ref}}(A)} \int s(x, A) \mu_{\text{ref}}(dx), \quad s(x, A) = \frac{1}{\mu_{\text{ref}}(A)} \int_A U(x, y) \mu_{\text{ref}}(dy).$$

Second-order refinement when the first-order effect vanishes

If $\int U(x, y) \mu_{\text{ref}}(dy) = 0$ for all x , then $\dot{T}_0(A) = 0$ and

$$\ddot{T}_0(A) = \frac{1}{\mu_{\text{ref}}(A)} \int_{A^c} \{e(x, \mathcal{X}) - e(x, A) - 2\mu_{\text{ref}}(A)s(x, A)^2\} \mu_{\text{ref}}(dx),$$

where

$$e(x, B) = \frac{1}{\mu_{\text{ref}}(B)} \int_B U(x, y)^2 \mu_{\text{ref}}(dy).$$

Consequence

If enough preference energy points from A^c to A , then $\ddot{T}_0(A) < 0$, so small positive γ reduces hitting time.