

LLM Unlearning Methods

Bi-objective Goal

- 1) **Unlearn** model knowledge on **unlearn set**;
- 2) **Retain** model performance on **retain set**.

question	What is the full name of the geology author born in Karachi, Pakistan on 06/30/1975?	What is the capital of Japan?
original	The author's name is Hina Ameen .	The capital of Japan is Tokyo .
unlearned	As of now, the full name of the authors is not mentioned .	The capital of Japan is Tokyo .

LLM Unlearning Methods

$$\min_{\theta} \mathcal{L}_u(\mathcal{D}_u; \theta) + \mathcal{L}_r(\mathcal{D}_r; \theta)$$

Unlearn Objective Retain Objective

Each objective possesses **unique properties**, yet there is no **unified toolkit** available to comprehend them.

Gradient Effect (G-effect)

From a **gradient perspective**, analyzing the impacts of **unlearning objectives** on **model performance**.

unlearn: $\mathcal{R}(\mathcal{D}_u; \theta_u) \gg \mathcal{R}(\mathcal{D}_u; \theta_o)$ **retain**: $\mathcal{R}(\mathcal{D}_r; \theta_u) \leq \mathcal{R}(\mathcal{D}_r; \theta_o)$

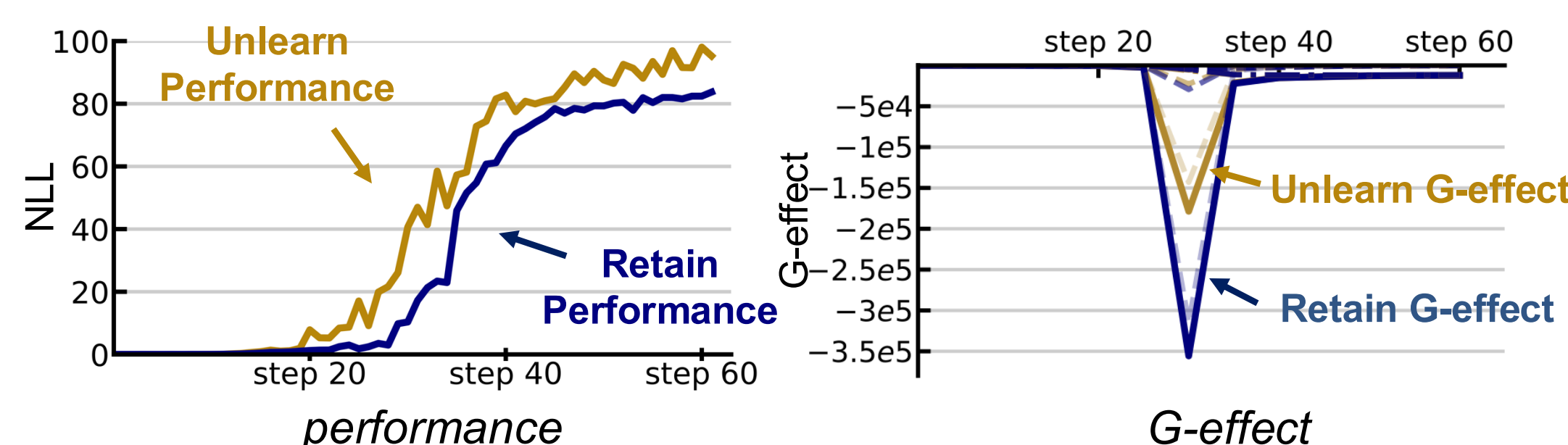
objective gradient

$$e = \nabla_{\theta} \mathcal{L}(\mathcal{D}_u; \theta)^T \nabla_{\theta} \mathcal{R}(\mathcal{D}; \theta)$$

performance gradient $\mathcal{L}$ benefits $\mathcal{R}$ $\mathcal{L}$ damages $\mathcal{R}$

Unlearn G-effect, negative e_u is preferred for removal;
Retain G-effect, positive e_r is preferred for retention.

G-effect vs. Performance



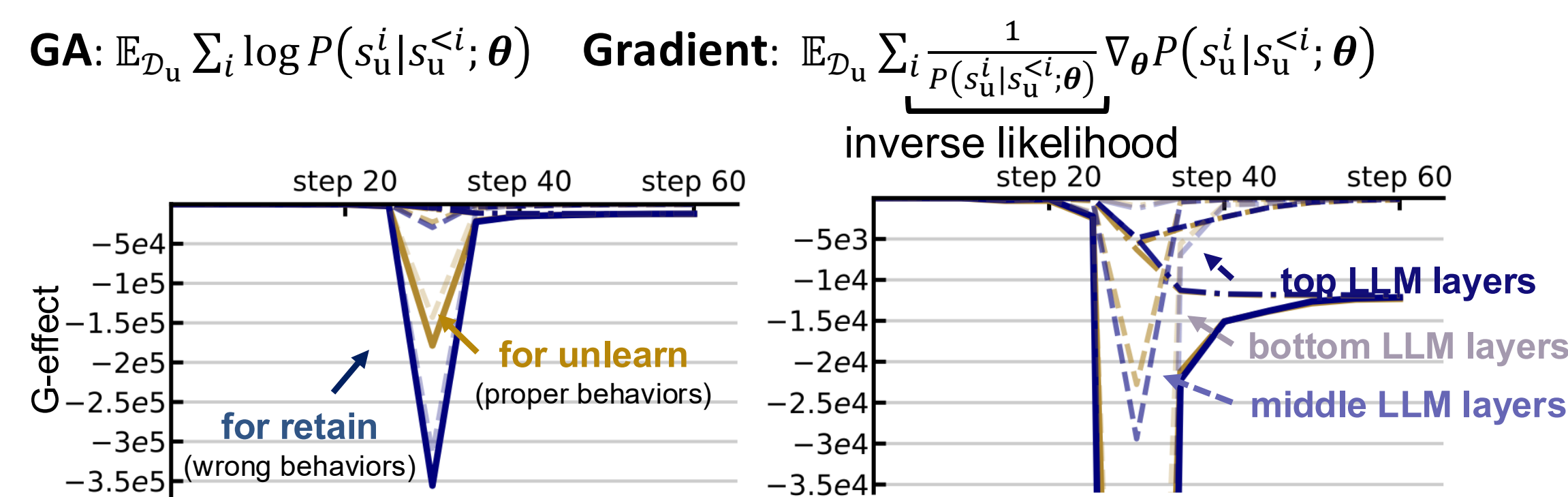
- We can **disentangle** the impacts of unlearn and retain objectives through G-effect, yet hard for metrics.

- G-effect allows us to quantify **positive/negative** impacts even for the bi-objectives with opposing impacts.
- Gradients offer **deeper insights** into the effects of data, layers, or sub-components within objectives.

Case Studies

With the G-effect, we test a set of **unlearn objectives (GA, NPO, PO, RMU)** and **retain objectives (NLL, KL, RR)** separately.

Unlearn Objective 1. Gradient Ascent (GA)



Observation 1. Unlearning **negatively** impacts retention.

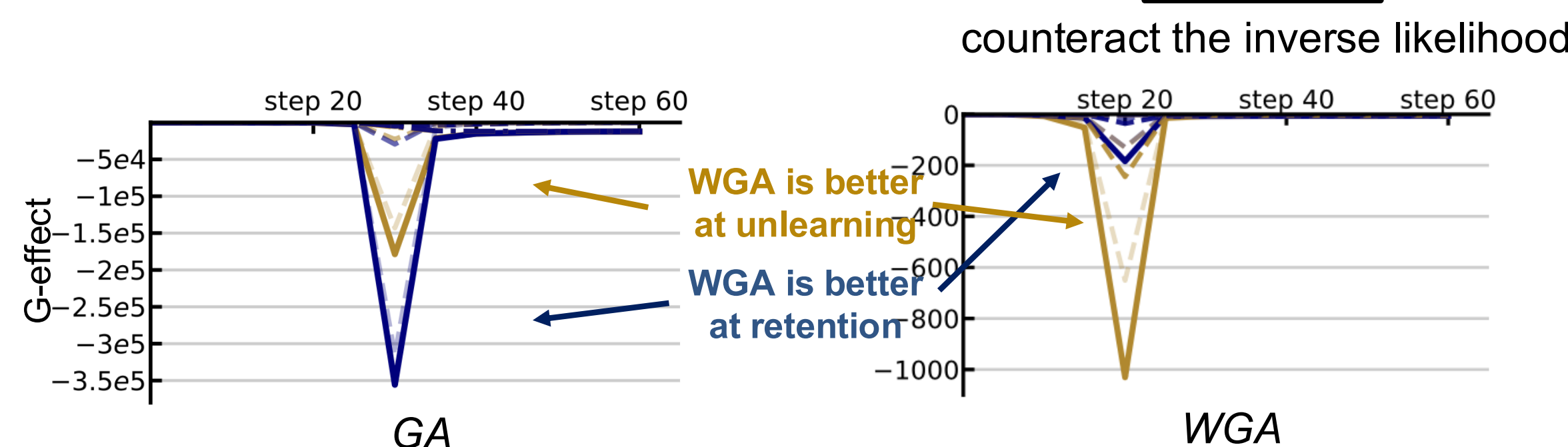
Reason. The inverse likelihood wrongly focuses more on sufficiently unlearned tokens, leading to **over-unlearning**.

Observation 2. Unlearning **affects bottom layers** of LLMs more.

Reason. Large gradients **accumulate** due to the chain rule, a general scenario holds for many other unlearning objectives.

Improvement 1. Weighted GA (WGA)

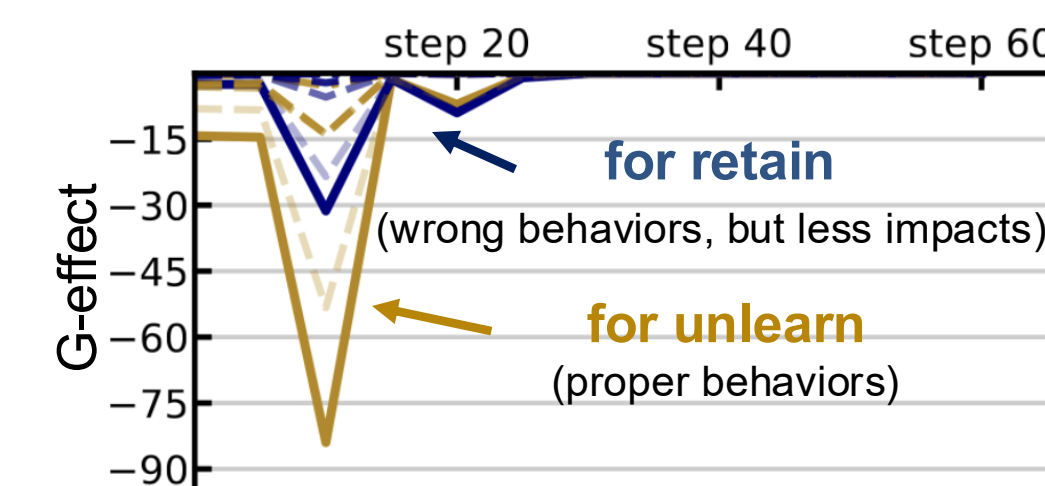
WGA: $\mathbb{E}_{\mathcal{D}_u} \sum_i P(s_u^i | s_u^{<i}; \theta)^{\alpha} \log P(s_u^i | s_u^{<i}; \theta)$ Gradient: $\mathbb{E}_{\mathcal{D}_u} \sum_i P(s_u^i | s_u^{<i}; \theta)^{\alpha-1} \nabla_{\theta} P(s_u^i | s_u^{<i}; \theta)$



Unlearn Objective 2. Negative Preference Optimization (NPO)

NPO: $\mathbb{E}_{\mathcal{D}_u} \frac{1}{\beta} \log(1 + \left(\frac{p(s_u; \theta)}{p(s_u; \theta_o)}\right)^{\beta})$ Gradient: $\mathbb{E}_{\mathcal{D}_u} \sum_i \frac{2P(s_u; \theta)^{\beta}}{P(s_u; \theta)^{\beta} + P(s_u; \theta_o)^{\beta}} \nabla_{\theta} \log P(s_u; \theta)$

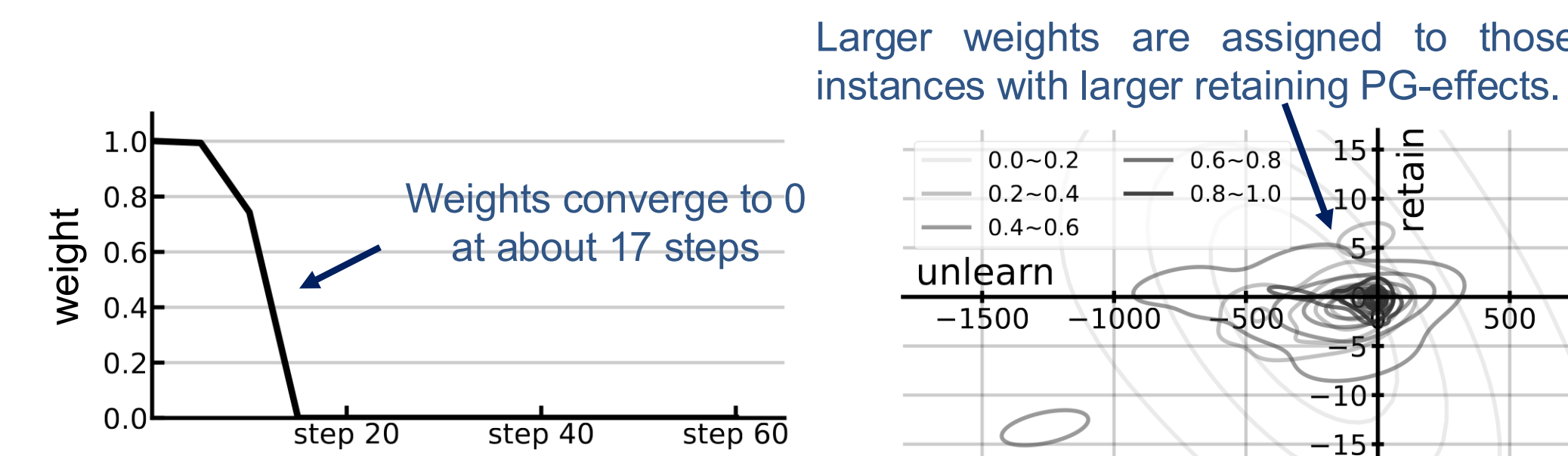
w_{npo} reweighting



Observation 3. NPO has **fewer negative impacts** on retention.

Reason. NPO improves GA via w_{npo} **reweighting term**.

How to understand the reweighting mechanism w_{npo} within NPO?



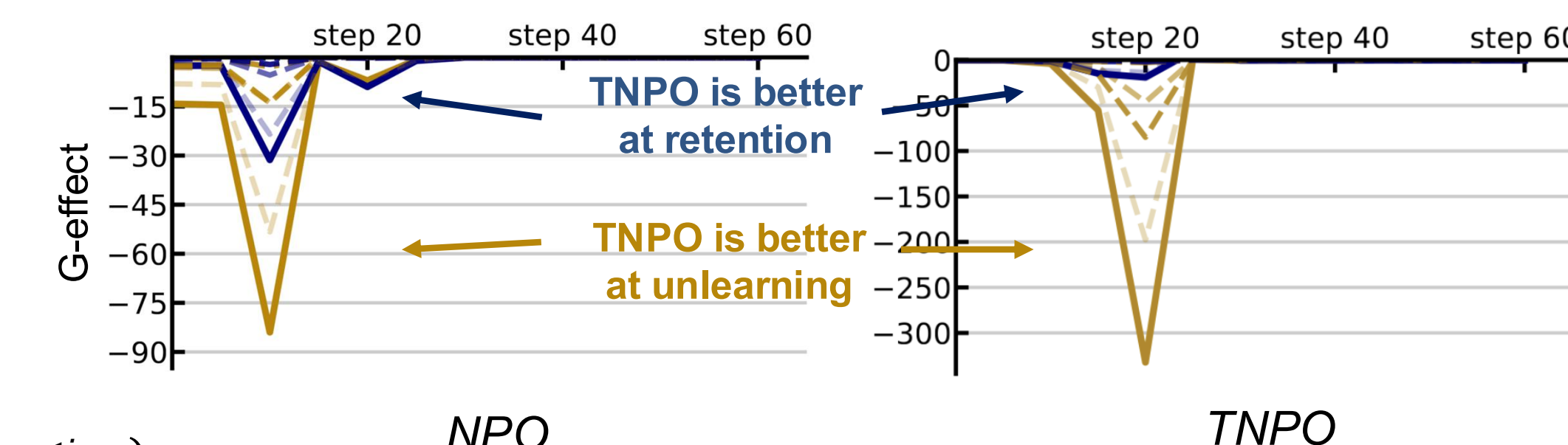
Observation 4. The NPO weight serves as **early stopping**.

Observation 5. The NPO reweighting mechanism w_{npo} **prioritizes instances** that less damages retention.

Improvement 2. Token-wise NPO (TNPO)

TNPO: $\sum_i w_{\text{tnpo}}^i \log P(s_u^i | s_u^{<i}; \theta)$ with $w_{\text{tnpo}}^i = \frac{2P(s_u^i | s_u^{<i}; \theta)^{\alpha}}{P(s_u^i | s_u^{<i}; \theta)^{\alpha} + P(s_u^i | s_u^{<i}; \theta_o)^{\alpha}}$

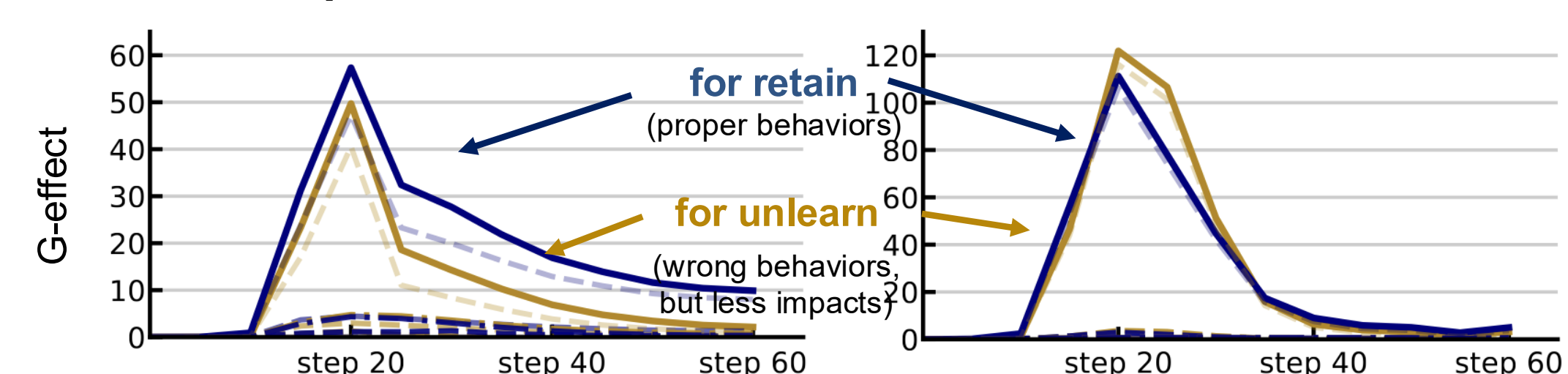
same reweighting scheme yet applied point-wise



Retain Objectives

NLL: $\mathbb{E}_{\mathcal{D}_r} [-\log P(s_r; \theta)]$

KL: $\mathbb{E}_{\mathcal{D}_r} \text{KL}[P(s_r; \theta) || P(s_r; \theta_o)]$



Observation 6. NLL and KL are both effective for retention, while KL can lead to overall larger retain G-effect, thus preferred.

Note. The unlearn G-effect for unlearn objective is much larger than for retain objectives. Thus, we do not need to worry about the side effect on unlearning.