

$$J(\theta) = E_{y \sim p_{\theta}(\cdot|x)} [r(x, y)]$$

$$= \sum_y r(x, y) p_{\theta}(y|x)$$

$$\frac{dJ}{d\theta} = \sum_y r(x, y) \frac{d}{d\theta} p_{\theta}(y|x)$$

↳ Q: can we just sample a few y and compute the average of $r(x, y) \frac{d}{d\theta} p_{\theta}(y|x)$?

↳ A: no, because

$$E_{y \sim p_{\theta}(\cdot|x)} \left[r(x, y) \frac{d}{d\theta} p_{\theta}(y|x) \right]$$

from sampling!

$$= \sum_y \underbrace{p_{\theta}(y|x)} r(x, y) \frac{d}{d\theta} p_{\theta}(y|x)$$

dividing by this now yields log trick!

So from last class:

$$A(x, y) = r(x, y) - \underset{\downarrow}{b(x)}$$

baseline:

1. mean batch stats
2. group stats
3. trained value fn or critic
 - ↳ RLHF

Value fn:

↳ takes prompt x and partial response $y_{<i}$, predicts $r(x, y)$

↳ just like reward model, scalar head trained on top of LM, L2 loss

↳ unlike RM, Value fn trained alongside policy model!

$$J(\theta) = \mathbb{E}_{x, y \sim p_\theta} \left[A(x, y) - \beta \underbrace{D_{KL}(p_\theta \| p_{SFT})}_{\log p_\theta(y|x) - \log p_{SFT}(y|x)} \right]$$

↪ reverse KL penalty

$$= \sum_{i=1}^n \log p_\theta(y_i|x) - \log p_{SFT}(y_i|x)$$

We're not done yet!

↪ want to make small, stable changes when updating p_θ

↪ **policy ratio** $p_i = \frac{p_\theta(y_i|x)}{p_{\theta_{old}}(y_i|x)}$

$p_i = 1$
 ↪ $\theta = \theta_{old}$ for first step
 on batch, different
 after subsequent steps,

↳ how much more/less likely
is this token under the new
policy vs. the policy that generated
it?

↳ we can weight our advantage
by P_i , and control update
by clipping P_i

Let's denote

$$A_t^{KL}(x, y) = A(x, y) - \beta D_{KL}(P_\theta || P_{SFT})$$

then

$$J(\theta) = E \left[\sum_{i=1}^n \min \left(P_i A_t^{KL}(x, y), \right. \right. \\ \left. \left. \text{clip}(P_i, 1-\epsilon, 1+\epsilon) A_t^{KL}(x, y) \right) \right]$$

↑ clipping
hyperparam
= 0.1, 0.2

intuition:

↳ if $A_t^{kl} > 0$, we like this token and want to increase its prob.

↳ max increase of $(1+\epsilon)A_t^{kl}$

↳ if $A_t^{kl} < 0$, we don't like the token

↳ max decrease $(1-\epsilon)A_t^{kl}$

↳ no gradient thru clipped term!

example:

$\epsilon = 0.2$, so we clip $(p_i, 0.8, 1.2)$

$A_i = 2.0$, $p_i = 1.3$

↳ unclipped = 2.6

clipped = 2.4 $\Rightarrow$ choose min,

no gradient,
token already more
likely under new θ

$A_i = 2.0$, $p_i = 1.1$

↳ unclipped = 2.2, clipped = 2.4

↳ normal gradient, increase prob of token

$$A_i = -2.0, p_i = 0.6$$

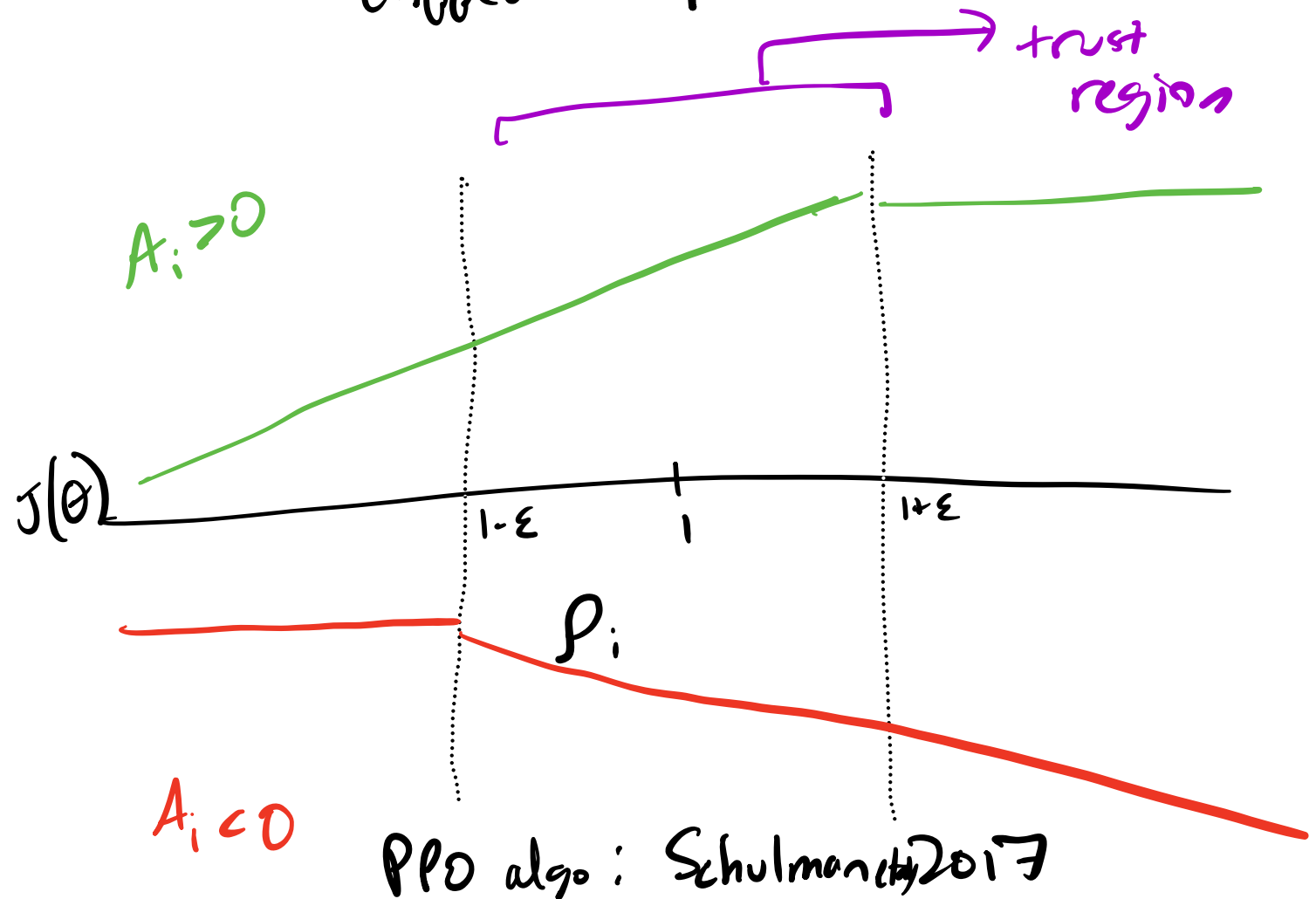
↳ unclipped = -1.2

clipped = -1.6 $\Rightarrow$ choose min, no gradient, token already less likely under new θ

$$A_i = -2.0, p_i = 0.9$$

↳ unclipped = -1.8 $\rightarrow$ normal grad, decrease prob of token

clipped = -1.6



GRPO : group-relative PPO

↳ baseline comes from group Statistics

↳ for each prompt x , sample a group of responses $y_1, y_2, \dots, y_G$

↳ compute mean reward μ_x and std dev of reward σ_x over group

↳ for any specific y_i in group,

$$A(x, y_i) = \frac{r(x, y_i) - \mu_x}{\sigma_x}$$

↳ z-score

Deepseek R1:

- ↳ GRPO

- ↳ remove reward model,
instead use simple verifiable tasks

- ↳ e.g. Math problems

 - ↳ prompt x : $37 + 53 = ?$

 - ↳ response y_1 : 32 $\Rightarrow$ reward 0

 - response y_2 : 47 $\Rightarrow$ reward 0

 - response y_3 : 90 $\Rightarrow$ reward 1

- ↳ incentivize thinking / reasoning
before answering