

From last class (PPO):

$b = \text{batch size}$
 $L = \text{seq length}$

1. given a batch of prompts $x_1, x_2, \dots, x_b$
we generate corresponding responses $y_1, y_2, \dots, y_b$ from policy model
2. for each example (x_i, y_i) , we compute reward $r(x_i, y_i)$
3. compute advantages by subtracting baseline / KL penalty

$$A(x, y) = r(x, y) - b$$

$\downarrow$
 $b \times 1$

$\downarrow$
varies
 $b \times 1$ if using batch mean

$b \times 2$ using value f_v

$$A^{kl}(x, y) = A(x, y) - \beta D_{kl}(P_\theta || P_{\text{ref}})$$

$\nearrow$
 $b \times L$

$b \times L$

4. compute policy ratios, take gradient steps

$$\rho_i = \frac{p_{\theta}(y_i | y_{<i}, x)}{p_{\theta_{old}}(y_i | y_{<i}, x)} \Rightarrow 1 \text{ before the first gradient steps}$$

θ_{old} frozen
 ρ of dimensionality $b \times L$

$$J(\theta) = \sum_{i=1}^L \min \left(\rho_i A^{kl}(x, y), \text{clip}(\rho_i, 1-\epsilon, 1+\epsilon) A^{kl}(x, y) \right)$$

$\hookrightarrow +A \Rightarrow \text{increase } \rho_i$

$-A \Rightarrow \text{decrease } \rho_i$

policy ratio.

$\hookrightarrow$ before any grad steps, $\rho_i = 1$

bc $p_{\theta_{old}} = p_{\theta}$

$\hookrightarrow 1-4$

$\hookrightarrow$ we normally take multiple steps on one batch

$\hookrightarrow$ why? can't we just take one step w/ a big LR

- ↳ Stability, if we take one big step, we'll move too far and many tokens will fall into clipped region
- ↳ rollouts (i.e. generating y from P_θ) are very expensive, compared to just doing multiple rounds of backprop

PPO $\rightarrow$ GRPO

- ↳ remove reward model, replace w/ verifiable reward
 - ↳ RL gets cheaper
 - ↳ reduce reward hacking
 - ↳ can't reliably compute reward for many tasks
- ↳ remove value fn, replace w/ group mean
- ↳ KL not included in advantage computation, but included as separate penalty
 - ↳ often completely dropped

1. given a prompt x ,
generate a group of G different
responses: $y_1, y_2, \dots, y_G$ → 4-32

2. compute reward for each response
in the group $r(x, y_i)$

3. compute group mean and std. dev

$$\mu_x = \frac{1}{G} \sum_{i=1}^G r(x, y_i) \quad \sigma_x$$

4. compute group relative advantages → $b \times 1$

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

→ $b \times 1$

5. compute loss, do gradient updates

$$J(\theta) = \frac{1}{6} \sum_{i=1}^6 \left[\sum_{t=1}^6 \min \left(P_t A_t^{GR}(x, y_i), \text{clip}(P_t, 1-\epsilon, 1+\epsilon) A_t^{GR}(x, y_i) \right) \right]$$

$$- \beta D_K(P_\theta \| P_{SFT})$$

can be dropped

okay but I still have $r(x, y_i)$

↳ verifiable reward

↳ binary (0 or 1)

↳ correctness, proper formatting

↳ graded by some simple rules

x'. compute 219 + 408

$r(x, y)$: if $y = 627 \Rightarrow r(x, y) = 1$
 else $r(x, y) = 0$ } correctness reward

↳ thinking tokens $\Rightarrow$ reasoning

↳ special subsequence of y
used for reasoning

x' : let $2x+5=17$. solve for x . } problem

thinking request format { put your thinking steps between $\langle \text{think} \rangle \langle / \text{think} \rangle$ tokens. box your final answer $\langle \text{answer} \rangle \langle / \text{answer} \rangle$

y : $\langle \text{think} \rangle$ well okay so $2x+5=17$, that means $2x=12$, so $x=6$. $\langle / \text{think} \rangle$
 $\langle \text{answer} \rangle 6 \langle / \text{answer} \rangle$

↳ reward computed only over the final answer

↳ thinking tokens still part of backbone

↳ separate format reward