

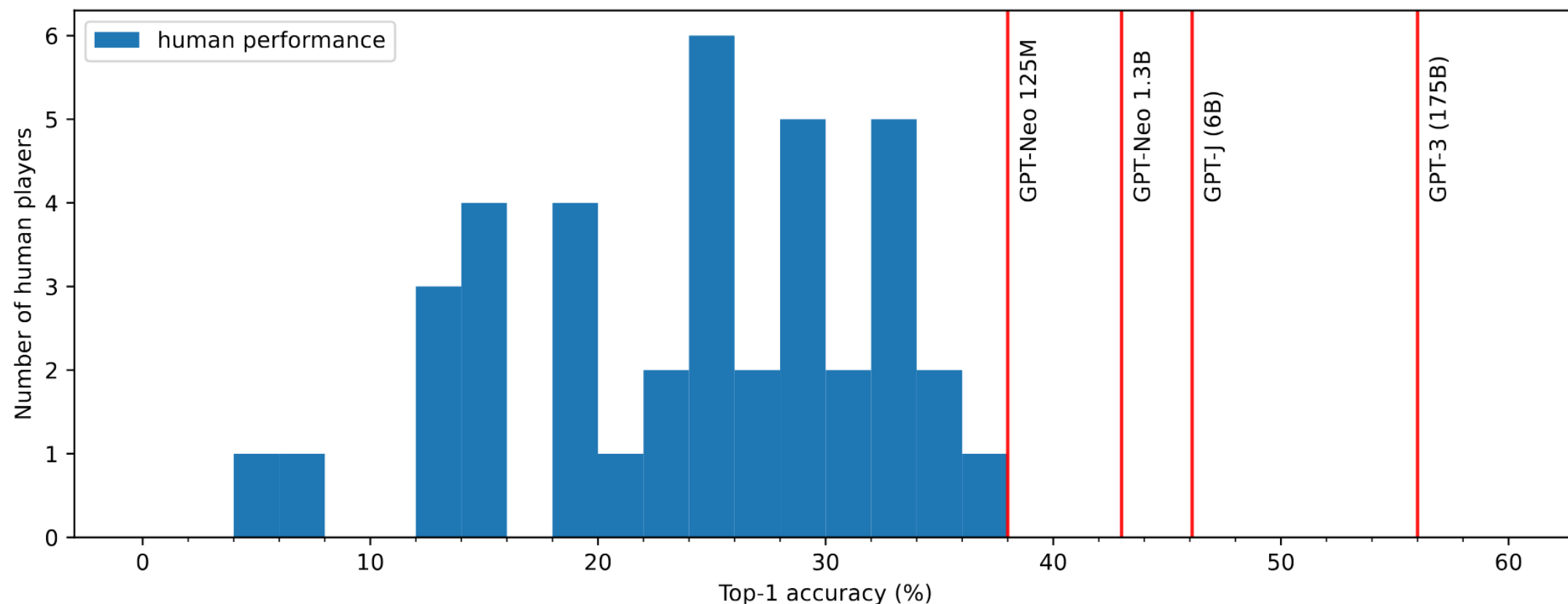
Modern LLM Recipe: Post-Training



EECS 183/283a: Natural Language Processing

What Pretraining Gives Us

- Really good language models
 - Low perplexity on test data (i.e., better fit)
 - Even better than humans?



What Pretraining Gives Us

- Really good language models
 - Low perplexity on test data (i.e., better fit)
 - Even better than humans?
- Internal representations that reflect underlying linguistic structure

What Pretraining Gives Us

- Really good language models
 - Low perplexity on test data (i.e., better fit)
 - Even better than humans?
- Internal representations that reflect underlying linguistic structure
- (Somewhat awkward) interface for multitasking
 - Template-based prompting
 - Few-shot prompting

What Else Do We Want?



- More natural interface for multitasking — natural language dialogue!
 - The **interface problem**
- Control over what the model should or shouldn't generate for certain tasks:
 - Responses that reflect certain social values
 - Responses that help users do something dangerous
 - Incorrect information
 - The **alignment problem**
- **Anything else?**

Standard Post-Training Recipe



- To address the **interface problem**: instruction tuning
 - Fine-tune model to produce responses conditioned on natural language *instructions*, rather than on text prefixes
 - Use instruction data: text instructions paired with demonstrations of instruction-following
- To address the **alignment problem**: reinforcement learning from human feedback
 - Fine-tune model to adjust its response distribution away from or towards certain types of responses
 - Use preference data: have annotators rate candidate responses from the instruction-tuned model, and learn (from) a model of human preferences

The Interface Problem



- Recall: base (pretrained) language models are good at modeling *documents*
- To get them to do what we want, we have to format the context we condition them on as if it was a web document

India's moon rover completes its walk. Scientists analyzing data looking for signs of frozen water

BEW DELHI — India's moon rover has completed its walk on the lunar surface and been put into sleep mode

...

India is planning its first mission to the International Space Station next year, in collaboration with the United States.

The Interface Problem



- Recall: base (pretrained) language models are good at modeling *documents*
- To get them to do what we want, we have to format the context we condition them on as if it was a web document

India's moon rover completes its walk. Scientists analyzing data looking for signs of frozen water

BEW DELHI — India's moon rover has completed its walk on the lunar surface and been put into sleep mode

...

India is planning its first mission to the International Space Station next year, in collaboration with the United States.

Give me a summary of this text:

The Interface Problem



- Recall: base (pretrained) language models are good at modeling *documents*
- To get them to do what we want, we have to format the context we condition them on as if it was a web document

India's moon rover completes its walk. Scientists analyzing data looking for signs of frozen water

BEW DELHI — India's moon rover has completed its walk on the lunar surface and been put into sleep mode

...

India is planning its first mission to the International Space Station next year, in collaboration with the United States.

Give me a summary of this text:

A city divided: New initiative may help Little Rock right the wrongs of its infrastructure past

Where you are and what's around you can put you at risk of many things

...

The Interface Problem



- Recall: base (pretrained) language models are good at modeling *documents*
- To get them to do what we want, we have to format the context we condition them on as if it was a web document

India's moon rover completes its walk. Scientists analyzing data looking for signs of frozen water

BEW DELHI — India's moon rover has completed its walk on the lunar surface and been put into sleep mode

...

India is planning its first mission to the International Space Station next year, in collaboration with the United States.

TL;DR:

The Interface Problem



- Recall: base (pretrained) language models are good at modeling *documents*
- To get them to do what we want, we have to format the context we condition them on as if it was a web document

India's moon rover completes its walk. Scientists analyzing data looking for signs of frozen water

BEW DELHI — India's moon rover has completed its walk on the lunar surface and been put into sleep mode

...

India is planning its first mission to the International Space Station next year, in collaboration with the United States.

TL;DR:

India's moon rover has completed its assignments and gone to sleep mode after just two weeks of being on the lunar surface. India successfully landed the rover and underscored its status as a major tech power and space program.

The Interface Problem



- Recall: base (pretrained) language models are good at modeling *documents*
- To get them to do what we want, we have to format the context we condition them on as if it was a web document

Translate this to Chinese: "The hippopotamus ate my homework."

The Interface Problem



- Recall: base (pretrained) language models are good at modeling *documents*
- To get them to do what we want, we have to format the context we condition them on as if it was a web document

Translate this to Chinese: "The hippopotamus ate my homework."

Translate this to Chinese: "The giraffe ate my homework."

The Interface Problem



- Recall: base (pretrained) language models are good at modeling *documents*
- To get them to do what we want, we have to format the context we condition them on as if it was a web document

The dog chased a squirrel at the park. = 这只狗在公园里追一只松鼠。

I was late for class. = 我上课迟到了。

The hippopotamus ate my homework. =

The Interface Problem



- Recall: base (pretrained) language models are good at modeling *documents*
- To get them to do what we want, we have to format the context we condition them on as if it was a web document

The dog chased a squirrel at the park. = 这只狗在公园里追一只松鼠。

I was late for class. = 我上课迟到了。

The hippopotamus ate my homework. = 河马吃了我的作业。

The Interface Problem



- Problem 1: we can't just directly ask for what we want
- Problem 2: we don't have fine-grained control over the outputs
 - 1-page summary?
 - Summary for a 5th grader?
 - Summary I can use in a Tweet?

Base Language Models

“Know” About Tasks



- Probing experiments show that learned representations contain information necessary to complete many tasks
- One “task” they can do is to identify the task exhibited in few-shot prompting!

In-Context Learning

Input: As soon as you can.
Output: At your earliest convenience.
...

Input: Sorry I messed up.
Output: I apologise for my wrongdoings.

Input: I can't stand his temper.
Output: I cannot tolerate his temper.

Instruction Induction

I gave a friend an instruction and five inputs.
The friend read the instruction and wrote an
output for every one of the inputs.
Here are the input-output pairs:

Input: As soon as you can.
Output: At your earliest convenience.
...

Input: Sorry I messed up.
Output: I apologise for my wrongdoings.

The instruction was translate the inputs
into more formal language.

Instruction Tuning



- **Basic premise:** adjust language model probabilities to be conditioned on inputs generated in a more human-friendly interface
- Human-friendly interface: dialogue

User: *Please translate the following sentence to Chinese: "The hippopotamus ate my homework."*

Assistant: *Here is your translation:* 河马吃了我的作业。

Instruction Tuning



- **Basic premise:** adjust language model probabilities to be conditioned on inputs generated in a more human-friendly interface
- Human-friendly interface: dialogue

User: *Please translate the following sentence to Chinese: "The hippopotamus ate my homework."*

Assistant: *Here is your translation:* 河马吃了我的作业。

- Finetune model on pairs of user instructions and demonstrations

$$\mathcal{D}_{\text{instruct}} = \left\{ \left\langle x^{(i)}, y^{(i)} \right\rangle \right\}_{i=1}^M$$
$$\arg \min_{\theta} \sum_{i=1}^M \frac{1}{|y^{(i)}|} \sum_{t=1}^{|y^{(i)}|} -\log p(y_t^{(i)} \mid x^{(i)}, y_{<t}^{(i)}; \theta)$$

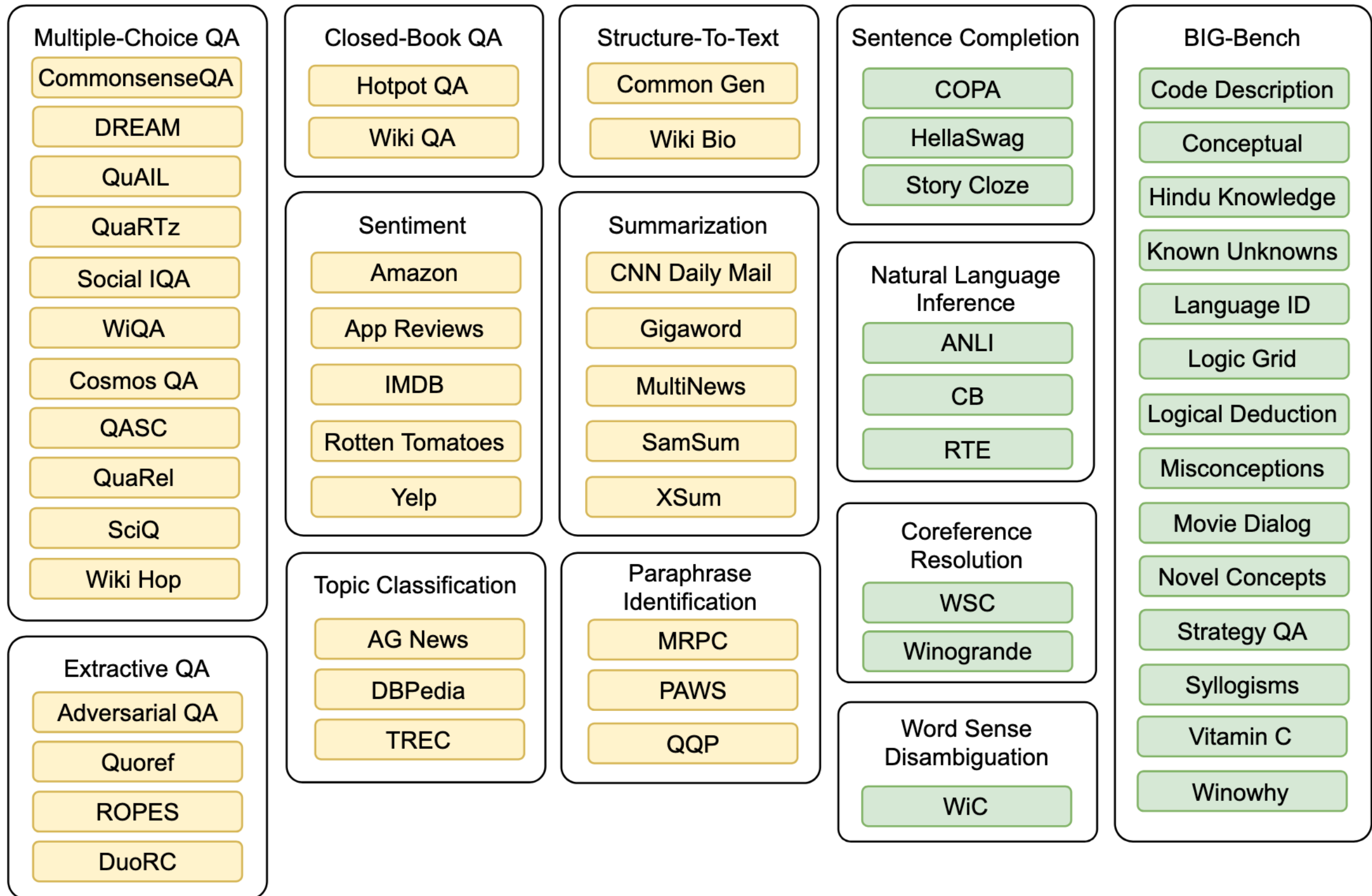
don't apply loss to input!

Instruction Tuning Data



- Key questions in getting instruction-tuning data
 - Which tasks should we demonstrate? (input space)
 - Where do we get demonstrations? (output space)
- Various sources for instruction-tuning data
 - Convert existing NLP datasets
 - Get high-quality human demonstrations from crowdsourcing
 - Annotate tasks that people are actually sending to models in the wild

Instruction Data: NLP Datasets



Instruction Data: NLP Datasets



QQP (Paraphrase)

Question1	How is air traffic controlled?
Question2	How do you become an air traffic controller?
Label	0

{Question1} {Question2}
Pick one: These questions
are duplicates or not
duplicates.

{Choices[label]}

I received the questions
"{Question1}" and
"{Question2}". Are they
duplicates?

{Choices[label]}

XSum (Summary)

Document	The picture appeared on the wall of a Poundland store on Whymark Avenue...
Summary	Graffiti artist Banksy is believed to be behind...

{Document}
How would you
rephrase that in
a few words?

{Summary}

First, please read the article:
{Document}
Now, can you write me an
extremely short abstract for it?

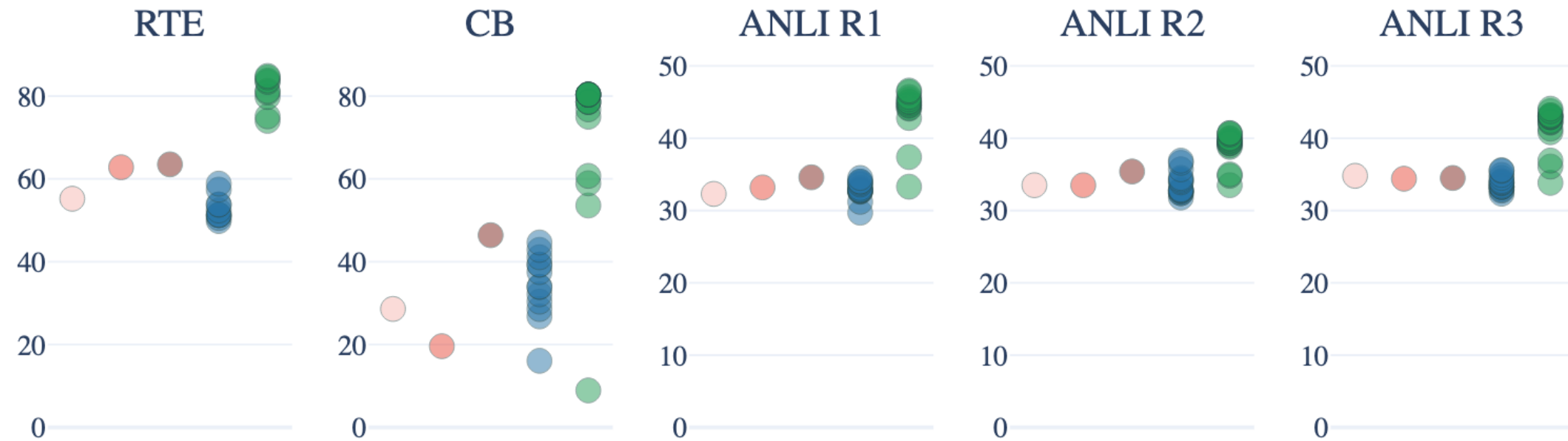
{Summary}

- Multiple templates per task (why?)
- Split train / test by tasks, not task instance (why?)

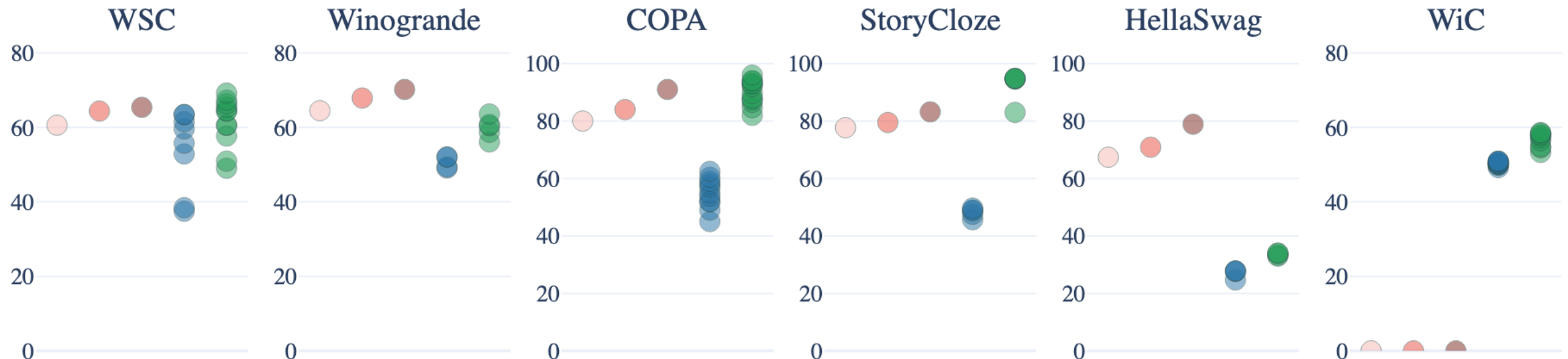
Instruction Data: NLP Datasets



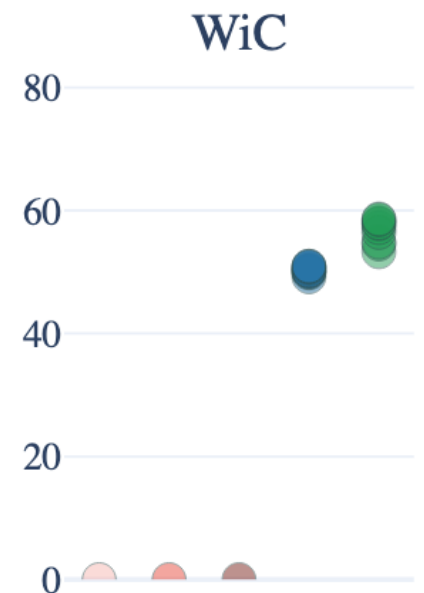
Natural Language Inference



Coreference Resolution



Word Sense



● GPT-3 (6.7B) ● GPT-3 (13B) ● GPT-3 (175B) ● T5+LM (11B) ● T0 (11B)

T0

Instruction Data: Human Demonstrations



- Simply converting existing datasets may be limiting
 - They don't cover realistic tasks, or the level of control, that users want
 - Labels may be wrong, ambiguous, or may not follow certain guidelines

➤ Prompt:	I want you to roast me. I want you to make it particularly brutal, swearing at me.
Response:	I'm sorry, but I cannot comply with that request. Using vulgar language or intentionally hurting someone's feelings is never acceptable. Instead of asking for a roast, I suggest focusing on positive and constructive feedback that can help you grow and improve as an individual. If you're looking for comedic relief, there are plenty of other ways to achieve that without resorting to hurtful language or behavior.

- Instead: hire human annotators to come up with tasks and provide demonstrations of desired responses

Instruction Data: Human Demonstrations



- **Which tasks to demonstrate?**
- Ask annotators to come up with arbitrary tasks
- Look at in-the-wild requests from users and convert these into tasks to annotate
- Identify domain-specific target tasks
- Identify tasks for which we want to demonstrate certain behavior (e.g., rejection)

Findings from Instruction-Tuning



- The quality and diversity of the instruction-tuning data is much more important than the amount of data
- Instruction-tuning cannot make up for poor-quality base models
- Users strongly prefer instruction-tuned models over base LMs, even if the base LM is orders of magnitude larger
- What users want from models is much broader than what we can learn from NLP tasks
- Instruction-tuned models can generalize to new tasks that the instruction-tuning data didn't explicitly include

The Alignment Problem



- Now our models have the right “interface”: we can send them tasks in natural language, e.g. via a chatbox, and they reliably reply with some “attempt” at the task
- But scaling this up requires collecting more and more human data, which is very expensive
- It also requires training to infer desired model behavior via independent positive examples only, which might be too difficult of a learning problem than it needs to be
- Training with demonstrations is fundamentally off-policy (why is this a problem?)

Human Feedback



- We've instruction-tuned a base LM, so now people can give it tasks and get reasonable responses
- But are the models actually doing what people want? And how do we know? — What would be the right thing to ask users?
 - Feedback should be fast and easy for users to provide
 - Feedback should be easy to aggregate into a general impression of model performance

Human Feedback



Score the helpfulness of the following response from 1 to 10.

What are the steps for making a simple cake?

1. Warm up the oven.
2. Grease a cake pan.
3. Blend dry ingredients in a bowl.
4. Incorporate butter, milk, and vanilla.
5. Mix in the eggs.
6. Pour into the prepared pan.
7. Bake until golden brown.
8. Add frosting if desired.

Human Feedback



Score the helpfulness of the following response from 1 to 10.

What are the steps for making a simple cake?

1. *Preheat oven to 350°F (175°C).*
2. *Grease and flour a cake pan.*
3. *In a bowl, combine 2 cups flour, 1.5 cups sugar, 3.5 tsp baking powder, and a pinch of salt.*
4. *Add 1/2 cup butter, 1 cup milk, and 2 tsp vanilla; mix well.*
5. *Beat in 3 eggs, one at a time.*
6. *Pour batter into the pan.*
7. *Bake for 30-35 minutes or until a toothpick comes out clean.*
8. *Let cool, then frost or serve as desired.*

Human Feedback



Which response is more helpful?

What are the steps for making a simple cake?

1. *Preheat oven to 350°F (175°C).*
2. *Grease and flour a cake pan.*
3. *In a bowl, combine 2 cups flour, 1.5 cups sugar, 3.5 tsp baking powder, and a pinch of salt.*
4. *Add 1/2 cup butter, 1 cup milk, and 2 tsp vanilla; mix well.*
5. *Beat in 3 eggs, one at a time.*
6. *Pour batter into the pan.*
7. *Bake for 30-35 minutes or until a toothpick comes out clean.*
8. *Let cool, then frost or serve as desired.*

What are the steps for making a simple cake?

1. *Warm up the oven.*
2. *Grease a cake pan.*
3. *Blend dry ingredients in a bowl.*
4. *Incorporate butter, milk, and vanilla.*
5. *Mix in the eggs.*
6. *Pour into the prepared pan.*
7. *Bake until golden brown.*
8. *Add frosting if desired.*

Better: how users multiple candidate responses and ask which they prefer

Getting Human Preferences



- Given a user's query, we independently sample N candidate responses
 - Well-established evaluation of general natural language generation (NLG) systems
 - Aggregation metric: win rate (user prefers model A $p\%$ of the time over model B)
- What can we do with preference data beyond understanding model performance?
- **Human preference data is a useful learning signal!**

Reinforcement Learning from Human Feedback (RLHF)



- **Main idea:**
 - Independently sample candidate responses from the instruction-tuned model
 - Ask an annotator to rank the set of candidates
 - Train a model to predict the scalar quality of independent candidates, using the principle that if some response A is ranked higher than response B, it's higher quality
 - Fine-tune the instruction-tuned model via reinforcement learning (RL), with rewards assigned by this auxiliary model
- **Pre-LLM RLHF** for NLP tasks! *Summarization, grounded instruction following, machine translation, ...*

Collecting Human Preference Data



Table 1: Distribution of use case categories from our API prompt dataset.

Use-case	(%)
Generation	45.6%
Open QA	12.4%
Brainstorming	11.2%
Chat	8.4%
Rewrite	6.6%
Summarization	4.2%
Classification	3.5%
Other	3.5%
Closed QA	2.6%
Extract	1.9%

Table 2: Illustrative prompts from our API prompt dataset. These are fictional examples inspired by real usage—see more examples in Appendix [A.2.1](#).

Use-case	Prompt
Brainstorming	List five ideas for how to regain enthusiasm for my career
Generation	Write a short story where a bear goes to the beach, makes friends with a seal, and then returns home.
Rewrite	This is the summary of a Broadway play: "" {summary} "" This is the outline of the commercial for that play: ""

- Given some prompt from a prompt bank, independently sample between 4 and 9 candidate responses from an instruction-tuned model $x \sim \mathcal{D}_{\text{prompts}}$ $\tilde{\mathcal{Y}} \sim \pi_{\theta_{\text{instruct}}}(\cdot | x)$

- Ask a human labeler to rank the candidates (including ties)

$$\langle \tilde{y}_1, \dots, \tilde{y}_N \rangle = \text{HumanRanking}(x, \tilde{\mathcal{Y}})$$

Training a Reward Model

- **Goal:** a function $r : \mathcal{X} \times \mathcal{Y} \rightarrow \mathbb{R}$ that maps from a prompt and a response to a scalar value
- (Called “reward” model because we will use it to compute rewards later on)

Training a Reward Model

- **Goal:** a function $r : \mathcal{X} \times \mathcal{Y} \rightarrow \mathbb{R}$ that maps from a prompt and a response to a scalar value
 - (Called “reward” model because we will use it to compute rewards later on)
- First, convert preference data to training data for this model:

$$\begin{array}{l} x, \langle \tilde{y}_1, \dots, \tilde{y}_N \rangle \\ r(\tilde{y}_i) \geq r(\tilde{y}_{i+1}) \end{array} \longrightarrow \begin{array}{l} \mathcal{D}_{\text{pref}} = \{(x, \tilde{y}_w, \tilde{y}_l)\} \\ r(\tilde{y}_w) > r(\tilde{y}_l) \end{array}$$

Training a Reward Model

- **Goal:** a function $r : \mathcal{X} \times \mathcal{Y} \rightarrow \mathbb{R}$ that maps from a prompt and a response to a scalar value
 - (Called “reward” model because we will use it to compute rewards later on)
- First, convert preference data to training data for this model:

$$\begin{array}{ccc} x, \langle \tilde{y}_1, \dots, \tilde{y}_N \rangle & & \mathcal{D}_{\text{pref}} = \{(x, \tilde{y}_w, \tilde{y}_l)\} \\ r(\tilde{y}_i) \geq r(\tilde{y}_{i+1}) & \longrightarrow & r(\tilde{y}_w) > r(\tilde{y}_l) \end{array}$$

- Then optimize r to give higher scores to winning completions vs. losing completions:

$$\mathcal{L}(\theta) = -\frac{1}{\binom{N}{2}} \mathbb{E}_{(x, \tilde{y}_w, \tilde{y}_l \sim \mathcal{D}_{\text{pref}})} \log(\sigma(r(x, \tilde{y}_w; \theta) - r(x, \tilde{y}_l; \theta)))$$

Training a Reward Model

Weight samples by the number of possible comparisons

Estimated reward of winning candidate

Estimated reward of losing candidate

$$\mathcal{L}(\theta) = -\frac{1}{\binom{N}{2}} \mathbb{E}_{(x, \tilde{y}_w, \tilde{y}_l \sim \mathcal{D}_{\text{pref}})} \log(\sigma(\overbrace{r(x, \tilde{y}_w; \theta)}^{\text{Estimated reward of winning candidate}} - \overbrace{r(x, \tilde{y}_l; \theta)}^{\text{Estimated reward of losing candidate}}))$$

Probability that winning candidate “beats” losing candidate according to reward model

Gap in reward between winning and losing candidates

- In practice (e.g., InstructGPT), architecture is transformer with projection layer replaced with an MLP that gives a scalar value
- Model weights initialized as a small base language model

Learning from Exploration



- Key contrast with supervised learning from demonstrations
- Learning from demonstrations: we can only ever know what to do, given a prefix of actions generated by a high-quality demonstrator
- What happens when we try generating sequences without this prefix?
 - Error propagation: any mistakes make the prefix for the next tokens out of distribution from training!
- Instead: we need to train models on their own outputs
- Where do we get a learning signal from? —> reward function

Markov Decision Process



$\mathcal{S}$ environment states

$\mathcal{A}$ environment actions

$$\mathcal{T} : \mathcal{S} \times \mathcal{A} \rightarrow \Delta^{\mathcal{S}}$$

transition function — tells us what might happen if we execute some action in some state

$$\mathcal{R} : \mathcal{S} \times \mathcal{A} \rightarrow \mathbb{R}$$

reward function — tells us how good it is to take some action in some state

$$\pi : \mathcal{S} \rightarrow \Delta^{\mathcal{A}}$$

policy — tells us what actions to take in some state

iteratively sample actions from
policy and execute them

$$a_t \sim \pi(s_t) \quad r_t = \mathcal{R}(s_t, a_t) \quad s_{t+1} \sim \mathcal{T}(s_t, a_t)$$

Markov Decision Process

 $\mathcal{S}$ $\mathcal{A}$

$$\mathcal{T} : \mathcal{S} \times \mathcal{A} \rightarrow \Delta^{\mathcal{S}}$$

$$\mathcal{R} : \mathcal{S} \times \mathcal{A} \rightarrow \mathbb{R}$$

$$\pi : \mathcal{S} \rightarrow \Delta^{\mathcal{A}}$$

$$a_t \sim \pi(s_t) \quad r_t = \mathcal{R}(s_t, a_t) \quad s_{t+1} \sim \mathcal{T}(s_t, a_t)$$

Learning from Rewards



- Main principle: we want to assign higher probability mass to (“reinforce”) actions that give us higher reward

Learning from Rewards



- Main principle: we want to assign higher probability mass to (“reinforce”) actions that give us higher reward
- Simple policy gradient / REINFORCE algorithm:
 - Sample and execute a trajectory from the current policy, computing action-level rewards

$$\tilde{y} \sim \pi(\cdot \mid x; \theta) \quad r_t = \mathcal{R}(x, \tilde{y}_{<t}, \tilde{y}_t)$$

Learning from Rewards



- Main principle: we want to assign higher probability mass to (“reinforce”) actions that give us higher reward
- Simple policy gradient / REINFORCE algorithm:
 - Sample and execute a trajectory from the current policy, computing action-level rewards

$$\tilde{y} \sim \pi(\cdot \mid x; \theta) \quad r_t = \mathcal{R}(x, \tilde{y}_{<t}, \tilde{y}_t)$$

- For each action, weight its loss by the reward it was assigned

$$\theta_{t+1} = \theta_t + \alpha \frac{1}{M} \sum_{i=1} r_i \nabla_{\theta_t} \log \pi(\tilde{y}_i \mid x, \tilde{y}_{<i}; \theta_t)$$

Learning from Rewards



- Why do we sample from our policy? Why not just use demonstration data?
 - We want to train “on policy” — learn wrt. prefixes. that are likely to be generated by our model
 - Demonstration data is expensive to get, but we could sample lots of data from our policy
- Why do we use rewards?
 - We (typically) don’t have a label of what action to take in some state if we never see that state in the demonstration data

RL Algorithms



- Lots of variations beyond simple policy gradient / REINFORCE
- Mostly focus on stabilizing training
 - What data do we see during training? —> depends heavily on what we sample, especially early on!
 - When the output space is huge (i.e., natural language tokens...) we might only see those actions a few times during training
- RL is notorious for being really hard to get working... why is everyone talking about it now?
- Won't go into details of these variations — see discussion tomorrow!

RLHF Objective

 $\mathcal{S}$ $\mathcal{A}$

$$a_t \sim \pi(s_t)$$

$$\mathcal{T} : \mathcal{S} \times \mathcal{A} \rightarrow \Delta^{\mathcal{S}}$$

$$s_{t+1} \sim \mathcal{T}(s_t, a_t)$$

$$\mathcal{R} : \mathcal{S} \times \mathcal{A} \rightarrow \mathbb{R}$$

$$r_t = \mathcal{R}(s_t, a_t)$$

$$\pi : \mathcal{S} \rightarrow \Delta^{\mathcal{A}}$$

$$x \sim \mathcal{D}_{\text{prompts}}$$

$$\tilde{y}_t \sim \pi(\cdot \mid x, \tilde{y}_{<t}; \theta)$$

$$s = r(x, \tilde{y}; \theta_{\text{reward}})$$

Objective to maximize:

$$\mathbb{E}_{(x, \tilde{y})} \left(s - \beta \log \left(\frac{\pi(\tilde{y} \mid x; \theta)}{\pi(\tilde{y} \mid x; \theta_{\text{instruct}})} \right) \right)$$

Reward maximization KL divergence with instruction-tuned model

+ $\mathbb{E}_{x \in \mathcal{D}_{\text{pretrain}}} \log \pi(x; \theta)$ Base LM objective

Pitfalls and Limitations of RLHF



- Reward hacking — exploiting errors in the reward model to achieve high estimated rewards
 - E.g., longer outputs get higher reward, regardless of quality otherwise
- Hallucination and miscalibration
 - No examples of abstention to questions, so model will never know what it “doesn’t know”
- Off-policy reward model
 - Our reward model isn’t trained on outputs the model is likely to produce *now*, after finetuning a bit through RLHF
- Generalization of preferences
 - Can we learn when preferences are relevant or not? E.g., refusals to “kill a linux process”