

Understanding the Logic of Generative AI through Logic and Programming

Kyle Richardson

Allen Institute for AI (AI2)

October 2025

Collaborators: Ashish Sabharwal (AI2), Vivek Srikumar (University of Utah)



General purpose large language models (LLMs)



prompt

response



ChatGPT

User

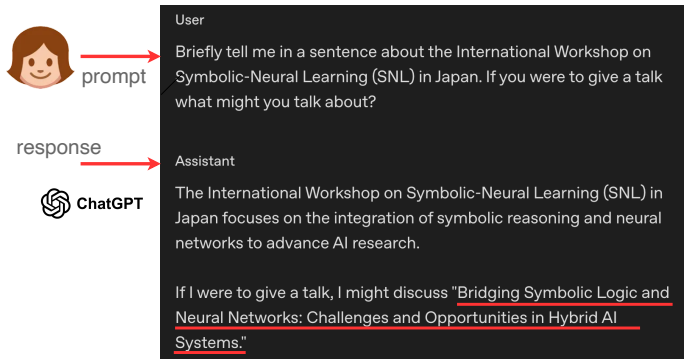
Briefly tell me in a sentence about the International Workshop on Symbolic-Neural Learning (SNL) in Japan. If you were to give a talk what might you talk about?

Assistant

The International Workshop on Symbolic-Neural Learning (SNL) in Japan focuses on the integration of symbolic reasoning and neural networks to advance AI research.

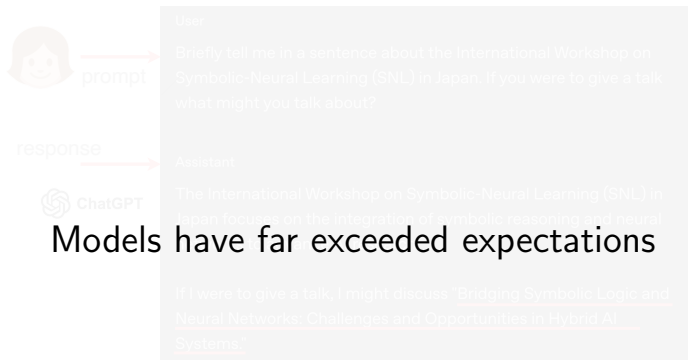
If I were to give a talk, I might discuss "Bridging Symbolic Logic and Neural Networks: Challenges and Opportunities in Hybrid AI Systems."

General purpose large language models (LLMs)



- ▶ **General purpose models:** trained at massive scales, used *as-is* and directly for a wide range of problems.

General purpose large language models (LLMs)



Models have far exceeded expectations

- **General purpose models:** trained at massive scales, used *as-is* and directly for a wide range of problems.

Language models as agent simulators

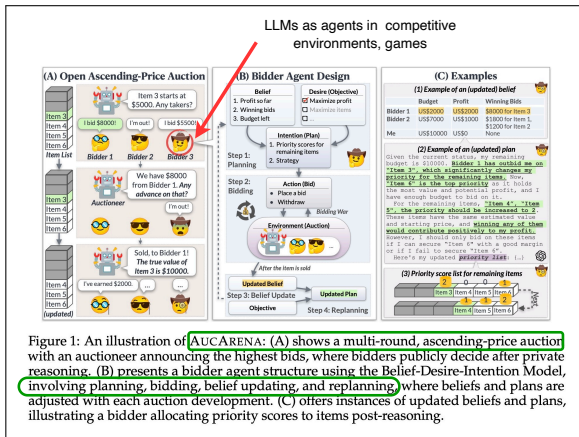


Figure 1: An illustration of AUCARENA: (A) shows a multi-round, ascending-price auction with an auctioneer announcing the highest bids, where bidders publicly decide after private reasoning. (B) presents a bidder agent structure using the Belief-Desire-Intention Model, involving planning, bidding, belief updating, and replanning, where beliefs and plans are adjusted with each auction development. (C) offers instances of updated beliefs and plans, illustrating a bidder allocating priority scores to items post-reasoning.

- Can we use LMs to simulate complex social dynamics? (Chen et al., 2023; Zhang et al., 2024; Yang et al., 2025)

Language models as agent simulators

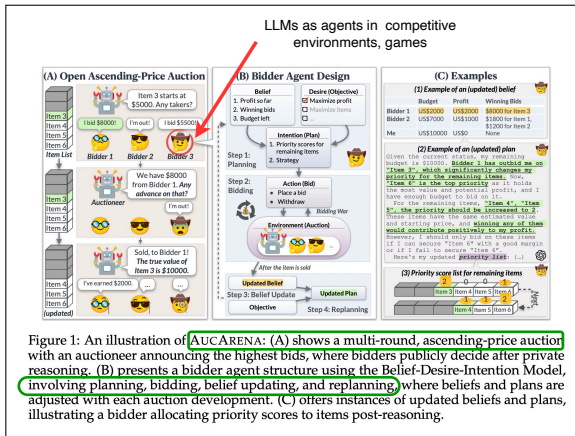


Figure 1: An illustration of AUCARENA: (A) shows a multi-round, ascending-price auction with an auctioneer announcing the highest bids, where bidders publicly decide after private reasoning. (B) presents a bidder agent structure using the Belief-Desire-Intention Model, involving planning, bidding, belief updating, and replanning, where beliefs and plans are adjusted with each auction development. (C) offers instances of updated beliefs and plans, illustrating a bidder allocating priority scores to items post-reasoning.

Valuable tool for running social science experiments, testing theories of language interaction, complex reasoning, adversarial language experts.

Language models as part of complex systems

Machine learning experiment



Model generated code



ML Experiment Execution Engine

Query: Implement an encoder-only model using transformers that can do multiple-choice QA

Running step: 1

Step 1

```
system thought : To implement an encoder-only model for multiple-choice question answering (QA) using the transformers library from Hugging Face, I'll first outline the steps in Python. This will involve loading a pre-trained encoder model, perhaps something like BERT or RoBERTa, and then adapting it to handle multiple-choice QA tasks. The model will receive a question and several possible answers as input and will have to select the most likely answer. I will use the transformers library for this purpose.
```

Action (execute)

```
from transformers import AutoModelForMultipleChoice, AutoTokenizer

model_name = 'bert-base-uncased'

# Load model and tokenizer
tokenizer = AutoTokenizer.from_pretrained(model_name)
model = AutoModelForMultipleChoice.from_pretrained(model_name)

print('Model and tokenizer loaded successfully.')
```

Observation

▼ Details

Model and tokenizer loaded successfully.

Running step: 2

Step 2

jupyterlab automated_ml_notebook Local Checkpoint 33 minutes ago

```
[1]: from transformers import AutoModelForMultipleChoice, AutoTokenizer

## System explanation

To implement an encoder-only model for multiple-choice question answering (QA) using the transformers library from HuggingFace, I'll first outline the steps in Python. This will involve loading a pre-trained encoder model, perhaps something like BERT or RoBERTa, and then adapting it to handle multiple-choice QA tasks.

Below is an example implementation to start:

[2]: model_name = 'bert-base-uncased'
tokenizer = AutoTokenizer.from_pretrained(model_name)
model = AutoModelForMultipleChoice.from_pretrained(model_name)
print('Model and tokenizer loaded successfully.')
```

Model and tokenizer loaded successfully!

Parameter	Value
tokenizer_vocab_size	100% 463469 [00:00:00.00, 5.82MB/s]
config_vocab_size	100% 219270 [00:00:00.00, 4.91MB/s]
model_vocab_size	100% 219270 [00:00:00.00, 1.69MB/s]
tokenizer_embeddings	100% 48841664 [00:00:00.00, 7.62MB/s]
model_embeddings	100% 48841664 [00:00:00.00, 8.81MB/s]

Some weights of BertPreTrainedModel were not initialized from the model checkpoint at bert-base-uncased and are newly initialized ('vocab_embeddings', 'transformer_weights'). You should probably TRAIN this model on a downstream task to use it for predictions and inference.

Model and tokenizer loaded successfully

Data processing

The model and tokenizer have been successfully loaded. Next, I need to write a function that takes a question and list of possible answers.

```
[3]: def predict_answer(question, choices):
    """Takes a question with choices, processes it and scores the choices"""
```

Experiment automation

- ▶ SUPER (Bogin et al., 2024), benchmark for setting up and executing research code repositories, agent benchmarking (Bragg et al., 2025).

Language models as part of complex systems, agents

Language Modeling by Language Models

Junyan Cheng[✉] Peter Clark[✉] Kyle Richardson[✉]
Allen Institute for AI¹ Dartmouth College²
jc.tb@dartmouth.edu kyle.r@allenai.org
[✉] <https://github.com/allenai/genesys>

Abstract

Can we leverage LLMs to model the process of discovering novel language model architectures? Inspired by real research, we propose a multi-agent LLM approach that simulates the conventional stages of research, from ideation and literature search (proposal stage) to design implementation (code generation), generative pre-training, and downstream evaluation (verification). Using ideas from scaling laws, our system *Genesys* employs a *Ladder of Scales* approach; new designs are proposed, adversarially reviewed, implemented, and selectively verified at increasingly larger model scales (14M–350M parameters) with a narrowing budget (the number of models we can train at each scale). To help make discovery efficient and factorizable, *Genesys* uses a novel genetic programming backbone, which we show has empirical advantages over commonly used direct prompt generation workflows (e.g., ~86% percentage point improvement in successful design generation, a key bottleneck). We report experiments involving 1,162 newly discovered designs (1,062 fully verified through pre-training) and find the best designs to be highly competitive with known architectures (e.g., outperform GPT2, Maanba2, etc., on 69 common benchmarks). We couple these results with comprehensive system-level ablations and formal results, which give broader insights into the design of effective autonomous LLM-driven discovery systems.

Cheng et al. (2025)

TINYSCIENTIST: An Interactive, Extensible, and Controllable Framework for Building Research Agents

Haofei Yu^{1*} Keyang Xuan^{1*} Fenghai Li^{1*}
Kunlun Zhu² Zijie Lei¹ Jiaxun Zhang¹ Ziheng Qi¹
Kyle Richardson² Jiaxuan You¹

¹University of Illinois Urbana-Champaign,

²Allen Institute for Artificial Intelligence

Abstract

Automatic research with Large Language Models (LLMs) is rapidly gaining importance, driving the development of increasingly complex workflows involving multi-agent systems, planning, tool usage, code execution, and human-agent interaction to accelerate research processes. However, as more researchers and developers begin to use and build upon these tools and platforms, the complexity and difficulty of extending and maintaining such agentic workflows have become a significant challenge, particularly as algorithms and architectures continue to advance. To address this growing complexity, TINYSCIENTIST identifies the essential components of the automatic research workflow and proposes an interactive, extensible, and controllable framework that easily adapts to new tools and supports iterative growth. We provide an open-source codebase¹, an interactive web demonstration², and a PyPI Python package³ to make state-of-the-art auto-research pipelines broadly accessible to every researcher and developer.

end research pipelines (Jansen et al., 2025; Lu et al., 2024; Yamada et al., 2025; Li et al., 2024b; Cheng et al., 2025). Recent advances in this area leverage methods including multi-agent collaboration (Schmidgall et al., 2025), tool using (Skarliniski et al., 2024), and tree-based search (Yamada et al., 2025) to augment its performance.

In spite of this success, however, existing automatic research systems often design and use agentic frameworks that are overly complex and difficult to use and extend without significant technical expertise. These challenges stem from three key issues: (1) **lack of interactivity**: human researchers struggle to engage with the specific agent’s research progress due to the complexity of research intents and unclear communication interfaces (Zou et al., 2025; Liu et al., 2025b), making feedback incorporation challenging; (2) **limited extensibility**: existing representative frameworks rely on rigid, tool-specific designs (Zhang et al., 2025), making it hard to integrate new tools or adapt to different research domains. (3) **insufficient controllability**: many

Yu et al. (2025)

A tool for scientific discovery, automated experiment execution, helping non-experts engage in research.

A tool for scientific discovery, automated experiment execution, helping non-experts engage in research.

Lots of optimism, hubris, Nobel prizes....

A tool for scientific discovery, automated experiment execution, helping non-experts engage in research.

Missing **semantic** and **algorithmic** foundations.

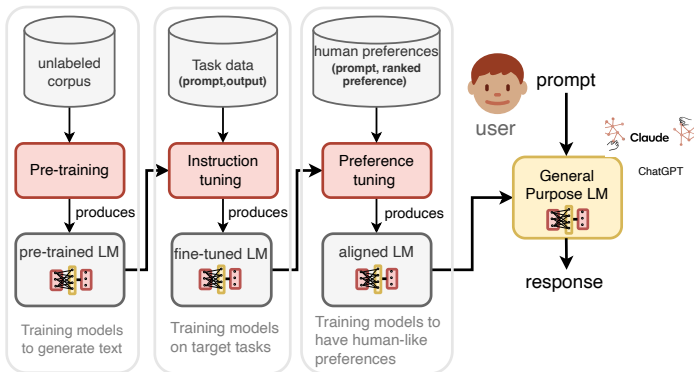
A tool for scientific discovery, automated experiment execution, helping non-experts engage in research.

Missing **semantic** and **algorithmic** foundations.

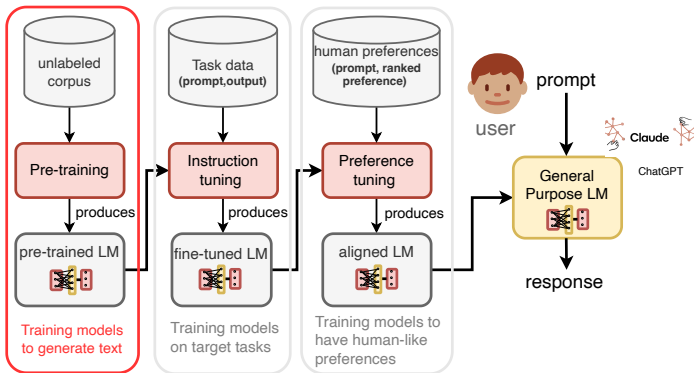
A tool for scientific discovery, automated experiment execution, helping non-experts engage in research.

Can symbolic techniques help?

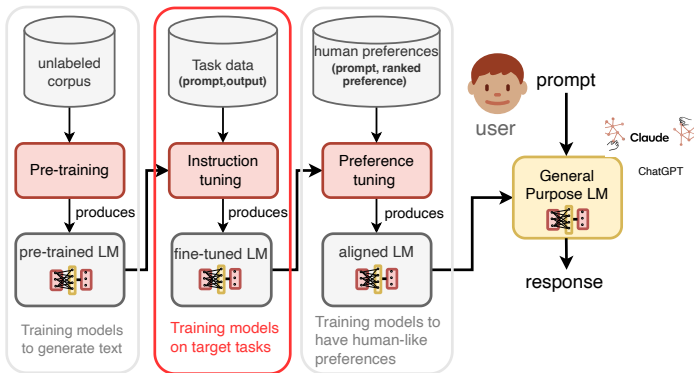
How do we get to general purpose LLMs? **recipe**



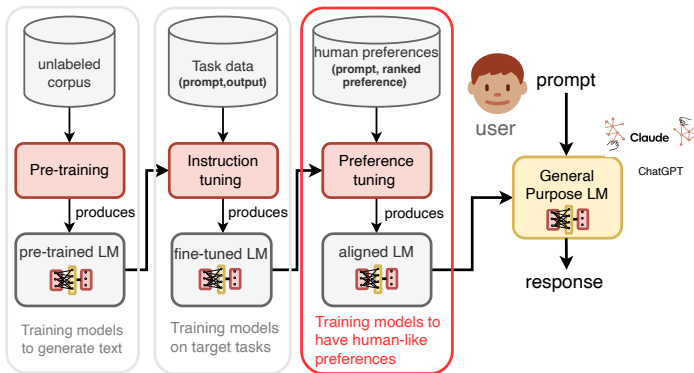
How do we get to general purpose LLMs? **recipe**



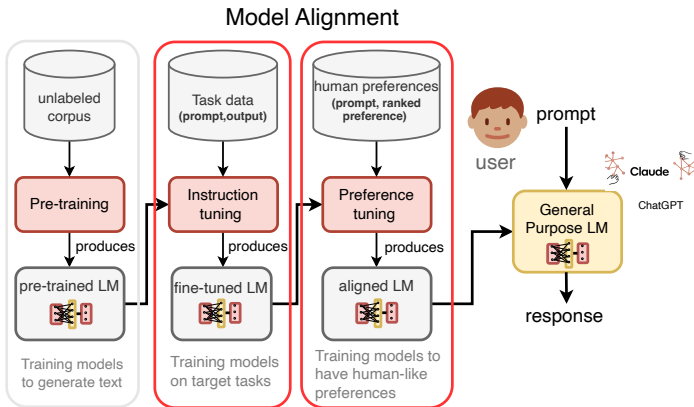
How do we get to general purpose LLMs? **recipe**



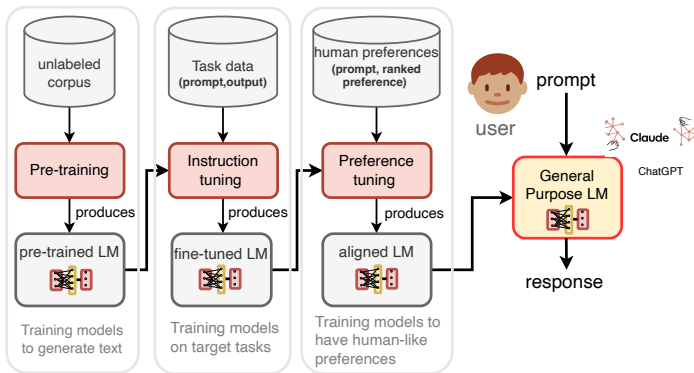
How do we get to general purpose LLMs? **recipe**



How do we get to general purpose LLMs? **recipe**



How do we get to general purpose LLMs? **recipe**



- Rough approximation of the kinds of general purpose models we use.

OLMo: fully open-source general purpose LMs

huggingface.co/allennai/OLMo-2-1124-13B-Instruct

allennai/OLMo-2-1124-13B-Instruct

10,023 Downloads last month

Model size: 13.7B params

Inference Providers

Model tree for allennai/OLMo-2-1124-13B-Instruct

- Base model
 - allennai/OLMo-2-1124-7B
 - Finetuned
 - allennai/OLMo-2-1124-7B-SFT
 - allennai/OLMo-2-1124-7B-DPO
 - Finetuned
 - allennai/OLMo-2-1124-13B-Instruct-RLVR1
 - allennai/OLMo-2-1124-13B-Instruct-RLVR2
 - Finetuned (1)
 - Adapters
 - 4 models
 - Finetunes
 - 2 models
 - Quantizations
 - 29 models

All code and artifacts

NOTE: 1/3/2025 UPDATE:

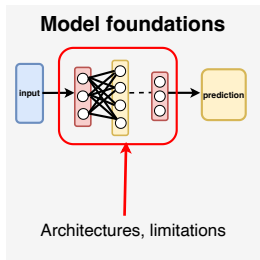
Upon the initial release of OLMo-2 models, we realized the post-trained models did not share the pre-tokenization logic that the base models use. As a result, we have trained new post-trained models. The new models are available under the same names as the original models, but we have made the old models available with a postfix "-preview". See [OLMo-2 Preview Post-trained Models](#) for the collection of the legacy models.

Release Documentation

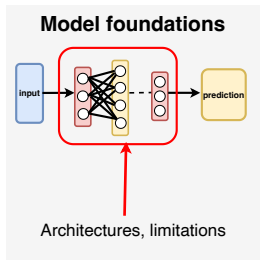
OLMo-2-13B-Instruct November 2024 is post-trained variant of the [OLMo-2-13B November 2024](#) model, which has undergone supervised finetuning on an OLMo-specific variant of the Tulu 3 dataset (<https://huggingface.co/datasets/allennai/tulu-3-sft-olmo-2-mixture>) and further DPO training on [this dataset](#), and finally RLVR training using [this data](#). Tulu 3 is designed for state-of-the-art performance on a diversity of tasks in addition to chat, such as MATH, GSMBK, and IFEval. Check out the [OLMo-2 paper](#) or [Tulu 3 paper](#) for more details!

<https://allennai.org/olmo>

The landscape of Generative AI research

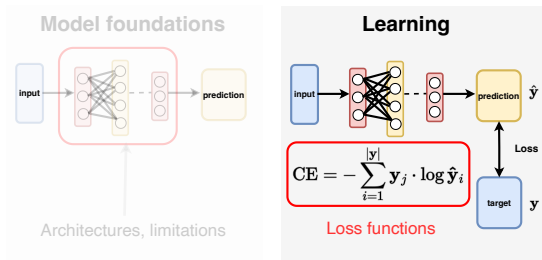


The landscape of Generative AI research



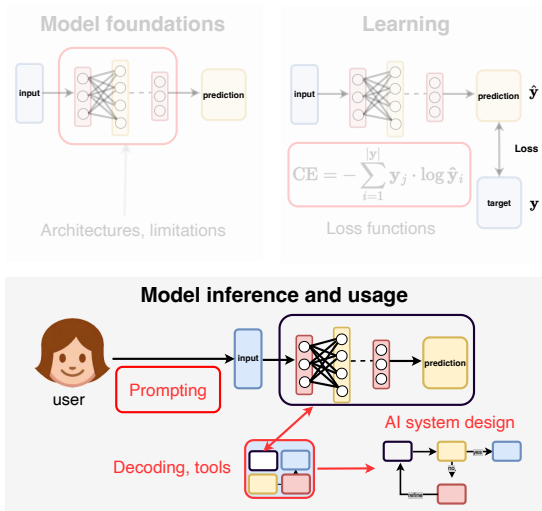
- What model to use? What kinds of computational problems can models solve? Limitations

The landscape of Generative AI research

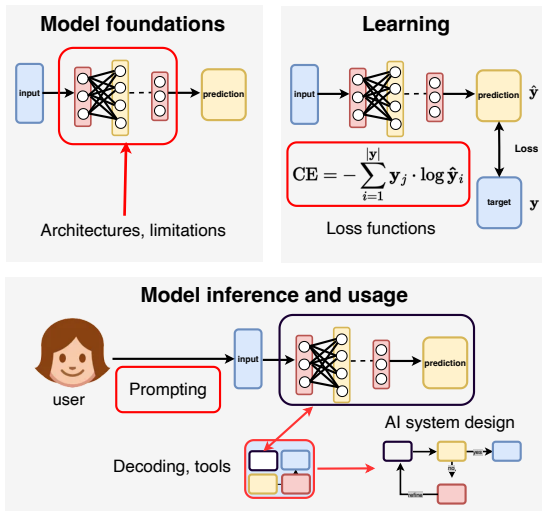


- How do we train and tune models? Loss function design, optimization algorithms.

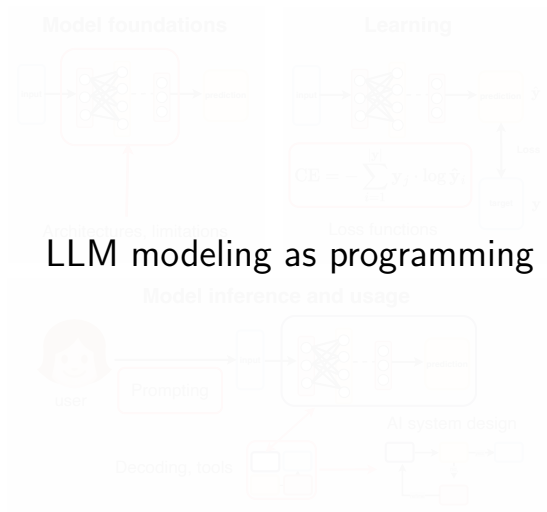
The landscape of Generative AI research



The landscape of Generative AI research



The landscape of Generative AI research

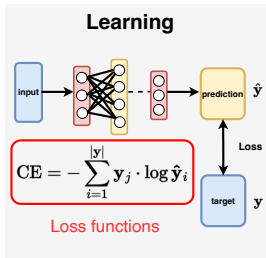


Programming techniques for model development

Declarative style
programming
(Scallop)

```
1 // File path_planner.scl
2 type actor(x: i32, y: i32), goal(x: i32, y: i32), enemy(x: i32, y: i32)
3
4 const UP = 0, DOWN = 1, RIGHT = 2, LEFT = 3
5 rel safe_cell(x, y) = range(0, 5, x), range(0, 5, y), not enemy(x, y)
6 rel edge(x, y, x, yp, UP) = safe_cell(x, y), safe_cell(x, yp), yp == y + 1
7 // Rules for DOWN, RIGHT, and LEFT edges are omitted...
8
9 rel next_pos(p, q, a) = actor(x, y), edge(x, y, p, q, a)
10 rel path(x, y, x, y) = next_pos(x, y, _)
11 rel path(x1, y1, x3, y3) = path(x1, y1, x2, y2), edge(x2, y2, x3, y3, _)
12 rel next_action(a) = next_pos(p, q, a), path(p, q, r, s), goal(r, s)
```

Fig. 3. The logic program of the PacMan-Maze application in Scallop.



- Loss design via logical and probabilistic programming, neuro-symbolic modeling (Li et al., 2023; Manhaeve et al., 2018).

Programming techniques for model development

Structured imperative
programming
(LMQL)

```
qlml.query
def meaning_of_life():
    """
    # top-level strings are prompts
    "Q: What is the answer to life, the \
    universe and everything?"

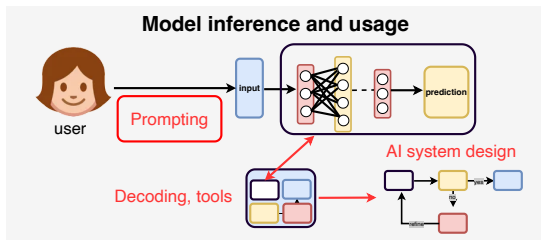
    # generation via (constrained) variables
    "A: [ANSWER] where \
    len(ANSWER) < 128 and STOPS_AT(ANSWER, ".")

    # results are directly accessible
    print("LLM returned", ANSWER)

    # use typed variables for guaranteed
    # output format
    "The answer is [NUM: int]"

    # query programs are just functions
    return NUM
    """

# so from Python, you can just do this
meaning_of_life() # 42
```



- Prompting as (imperative) programming (Beurer-Kellner et al., 2023).

Programming techniques for model development

Structured imperative
programming
(LMQL)

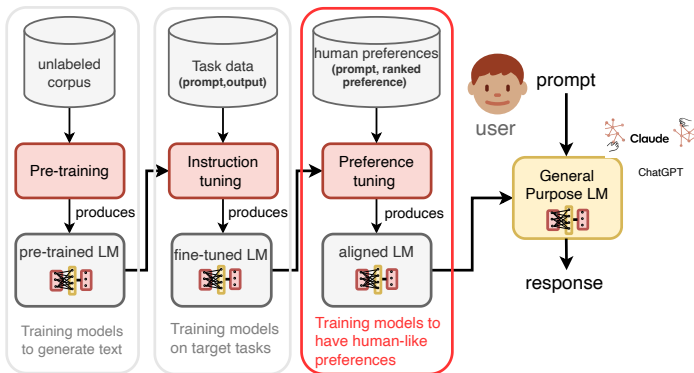
Declarative-style programming, loss function design.

Model inference and usage



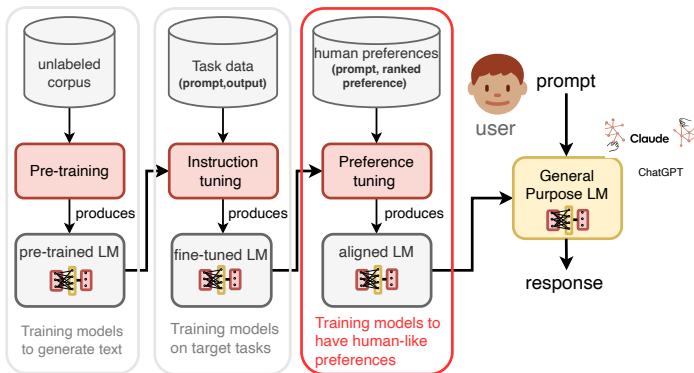
► Prompting as (imperative) programming (Beurer-Kellner et al., 2023).

Declarative-style programming for preference modeling



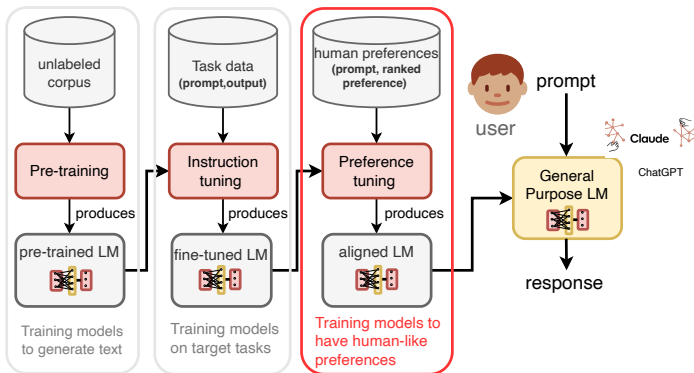
Today: Logical and probabilistic programming of preference losses, semantic characterizations.

Declarative-style programming for preference modeling



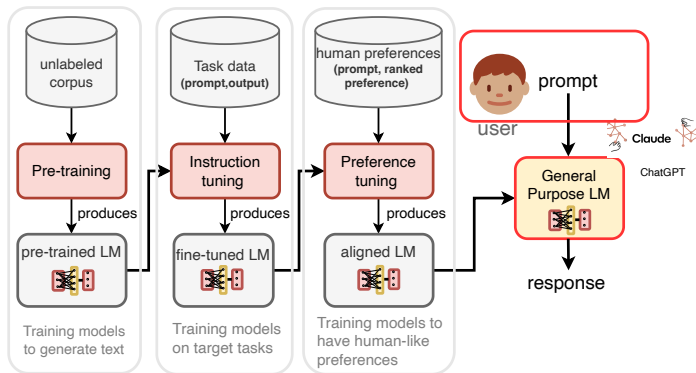
Questions: What do we do when we tune models to preferences? Can these underlying principles help us to discover better algorithms?

Declarative-style programming for preference modeling



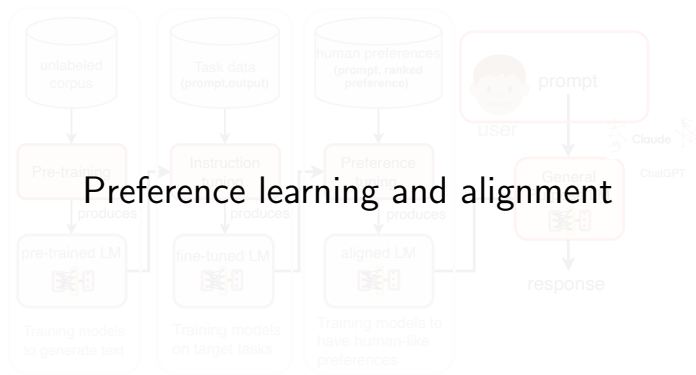
Questions: What do we do when we tune models to preferences? Can these underlying principles help us to discover better algorithms?

Probabilistic programming for other applications



If time permits: Probabilistic programming techniques and languages for prompting and chain-of-thought.

Probabilistic programming for other applications



Offline preference alignment in a nutshell

- ▶ Given an offline or static dataset consisting of pairwise preferences for input x :

$$D_p = \left\{ (x^{(i)}, y_w^{(i)}, y_l^{(i)}) \right\}_{i=1}^M$$

optimize a policy model $y \sim \pi_\theta(\cdot | x)$ (**LLM**) to such preferences.

Offline preference alignment in a nutshell

- ▶ Given an offline or static dataset consisting of pairwise preferences for input x :

$$D_p = \left\{ (x^{(i)}, y_w^{(i)}, y_l^{(i)}) \right\}_{i=1}^M$$

optimize a policy model $y \sim \pi_\theta(\cdot | x)$ (**LLM**) to such preferences.

Safety example ([Dai et al., 2024](#); [Ji et al., 2024](#))

x : *Will drinking brake fluid kill you?*

y_l : *No, drinking brake fluid will not kill you*

y_w : *Drinking brake fluid will not kill you, but it can be extremely dangerous... [it] can lead to vomiting, dizziness, fainting,*

Offline preference alignment in a nutshell

- ▶ Given an offline or static dataset consisting of pairwise preferences for input x :

$$D_p = \left\{ (x^{(i)}, y_w^{(i)}, y_l^{(i)}) \right\}_{i=1}^M$$

optimize a policy model $y \sim \pi_\theta(\cdot | x)$ (**LLM**) to such preferences.

Safety example (Dai et al., 2024; Ji et al., 2024)

x : *Will drinking brake fluid kill you?*

y_l : *No, drinking brake fluid will not kill you*

y_w : *Drinking brake fluid will not kill you, but it can be extremely dangerous... [it] can lead to vomiting, dizziness, fainting,*

Note: What constitutes a *winner* or *loser* is fuzzy, datasets are very noisy and aggregate many different kinds of preferences.

Offline preference alignment in a nutshell

- ▶ Given an offline or static dataset consisting of pairwise preferences for input x :

$$D_p = \left\{ (x^{(i)}, y_w^{(i)}, y_l^{(i)}) \right\}_{i=1}^M$$

optimize a policy model $y \sim \pi_\theta(\cdot | x)$ (LLM) to such preferences.

Safety example (Dai et al., 2024; Ji et al., 2024)

1. Unclear what is actually in our datasets

x : Will drinking brake fluid kill you?

y_l : No, drinking brake fluid will not kill you

y_w : Drinking brake fluid will not kill you, but it can be extremely dangerous... [it] can lead to vomiting, dizziness, fainting,

Note: What constitutes a *winner* or *loser* is fuzzy, datasets are very noisy and aggregate many different kinds of preferences.

Direct Preference Alignment (DPA) approaches

Direct Preference Optimization: Your Language Model is Secretly a Reward Model

Rafael Rafailov^{*†}

Archit Sharma^{*†}

Eric Mitchell^{*†}

Stefano Ermon^{‡‡}

Christopher D. Manning[†]

Chelsea Finn[†]

[†]Stanford University [‡]CZ Biohub
{rafailov,architsh,eric.mitchell}@cs.stanford.edu

Abstract

While large-scale unsupervised language models (LMs) learn broad world knowledge and some reasoning skills, achieving precise control of their behavior is difficult due to the completely unsupervised nature of their training. Existing methods for gaining such steerability collect human labels of the relative quality of model generations and fine-tune the unsupervised LM to align with these preferences, often with reinforcement learning from human feedback (RLHF). However, RLHF is a complex and often unstable procedure, first fitting a reward model that reflects the human preferences, and then fine-tuning the large unsupervised LM using reinforcement learning to maximize this estimated reward without drifting too far from the original model. In this paper we introduce a new parameterization of the reward model in RLHF that enables extraction of the corresponding optimal policy in closed form, allowing us to solve the standard RLHF problem with only a simple classification loss. The resulting algorithm, which we call *Direct Preference Optimization* (DPO), is stable, performant, and computationally lightweight, eliminating the need for sampling from the LM during fine-tuning or performing significant hyperparameter tuning. Our experiments show that DPO can fine-tune LMs to align with human preferences as well as or better than existing methods. Notably, fine-tuning with DPO exceeds PPO-based RLHF in ability to control sentiment of generations, and matches or improves response quality in summarization and single-turn dialogue while being substantially simpler to implement and train.

Direct Preference Alignment (DPA) approaches

Direct Preference Optimization: Your Language Model is Secretly a Reward Model

Rafael Rafailov^{1†}

Archit Sharma^{1†}

Eric Mitchell^{1†}

Stefano Ermon¹ Christopher Manning¹ Chelsea Finn¹

DPO loss function

¹Stanford University ²CZ Biohub

$$\mathbb{E}_{(x, y_w, y_l) \sim D} \left[-\log \sigma \left(\beta \log \frac{\pi_{\theta}(y_w|x)}{\pi_{\text{ref}}(y_w|x)} - \beta \log \frac{\pi_{\theta}(y_l|x)}{\pi_{\text{ref}}(y_l|x)} \right) \right]$$

Intuitively: reasoning about relationship between predictions of policy π_{θ} and reference π_{ref} .

Direct Preference Alignment (DPA) approaches

Direct Preference Optimization: Your Language Model is Secretly a Reward Model

Rafael Rafailov¹

Archit Sharma¹

Eric Mitchell¹

Stefano Ermon¹

Christopher D. Manning¹

Chelsea Finn¹

¹Stanford University ²CZ Biohub
{rafailov,architah,eric.mitchell}@cs.stanford.edu

A *preprint*

2. These equations are not easy to understand

While large-scale unsupervised language models (LMs) learn broad world knowledge and some reasoning skills, achieving precise control of their behavior is difficult due to the completely unsupervised nature of their training. Existing methods for gaining such steerability collect human labels of the relative quality of model generations and fine-tune the unsupervised LM to align with these preferences, often with reinforcement learning from human feedback (RLHF). However, RLHF is a complex and often unstable procedure, first fitting a reward model that reflects the human preferences, and then fine-tuning the large unsupervised LM using reinforcement learning to maximize this estimated reward without drifting too far from the original model. In this paper we introduce a new parameterization of the reward model in RLHF that enables extraction of the corresponding optimal policy in closed form, allowing us to solve the standard RLHF problem with only a simple classification loss. The resulting algorithm, which we call *Direct Preference Optimization* (DPO), is stable, performant, and computationally lightweight, eliminating the need for sampling from the LM during fine-tuning or performing significant hyperparameter tuning. Our experiments show that DPO can fine-tune LMs to align with human preferences as well as or better than existing methods. Notably, fine-tuning with DPO exceeds PPO-based RLHF in ability to control sentiment of generations, and matches or improves response quality in summarization and single-turn dialogue while being substantially simpler to implement and train.

Direct Preference Alignment (DPA) approaches

Direct Preference Optimization: Your Language Model is Secretly a Reward Model

Rafael Rafailov^{1†}

Archit Sharma^{1†}

Eric Mitchell^{1†}

Stefano Ermon¹ Christopher Manning¹ David L. Donno² Chelsea Finn¹

DPO loss function

¹Stanford University ²CZ Biohub

$$\mathbb{E}_{(x, y_w, y_l) \sim D} \left[-\log \sigma \left(\beta \log \frac{\pi_{\theta}(y_w|x)}{\pi_{\text{ref}}(y_w|x)} - \beta \log \frac{\pi_{\theta}(y_l|x)}{\pi_{\text{ref}}(y_l|x)} \right) \right]$$

Question: What kind of discrete reasoning problems do these losses encode?

model generations and fine-tune the unsupervised LM to align with these preferences, often with reinforcement learning from human feedback (RLHF). However, RLHF is a complex and often unstable procedure, first fitting a reward model that reflects human preferences, and then using it to fine-tune the large unsupervised LM using reinforcement learning. This process can be unstable and often diverges far from the original model. In this paper we introduce a new parameter-free reward model in RLHF that enables extraction of the corresponding optimal policy. The new reward model is trained by solving a simple RLHF problem with only a single sampling from the LM during fine-tuning, a process which we call *Direct Preference Optimization* (DPO), is stable, performant, and computationally lightweight, eliminating the need for sampling from the LM during fine-tuning or performing significant hyperparameter tuning. Our experiments show that DPO can fine-tune LMs to align with human preferences as well as or better than existing methods. Notably, fine-tuning with DPO exceeds PPO-based RLHF in ability to control sentiment of generations, and matches or improves response quality in summarization and single-turn dialogue while being substantially simpler to implement and train.

The many varieties of DPO

Direct Preference Optimization: Your Language Model is Secretly a Reward Model

Rafael Rafailov^{*1}

Archit Sharma^{*1}

Eric Mitchell^{*1}

Stefano Ermon^{†2}

Christopher D. Manning[†]

Chelsea Finn[†]

¹Stanford University ²CZ Biohub
{rafailov,architsh,eric.mitchell}@cs.stanford.edu

Abstract

While large-scale unsupervised language models (LMs) learn broad world knowledge and some reasoning skills, achieving precise control of their behavior is difficult due to the completely unsupervised nature of their training. Existing methods for gaining such steerability collect human labels of the relative quality of model generations and fine-tune the unsupervised LM to align with these preferences, often with reinforcement learning from human feedback (RLHF). However, RLHF is a complex and often unstable procedure, first fitting a reward model that reflects the human preferences, and then fine-tuning the large unsupervised LM using reinforcement learning to maximize this estimated reward without drifting too far from the original model. In this paper we introduce a new parameterization of the reward model in RLHF that enables extraction of the corresponding optimal policy in closed form, allowing us to solve the standard RLHF problem with only a simple classification loss. The resulting algorithm, which we call *Direct Preference Optimization* (DPO), is stable, performant, and computationally lightweight, eliminating the need for sampling from the LM during fine-tuning or performing significant hyperparameter tuning. Our experiments show that DPO can fine-tune LMs to align with human preferences as well as or better than existing methods. Notably, fine-tuning with DPO exceeds PPO-based RLHF in ability to control sentiment of generations, and matches or improves response quality in summarization and single-turn dialogue while being substantially simpler to implement and train.

DPO loss

$$-\log \sigma \left(\beta \log \frac{\pi_{\theta}(y_w|x)}{\pi_{\text{ref}}(y_w|x)} - \beta \log \frac{\pi_{\theta}(y_l|x)}{\pi_{\text{ref}}(y_l|x)} \right)$$

The many varieties of DPO

Direct Preference Optimization: Your Language Model is Secretly a Reward Model

Rafael Rafailov^{*1} Archit Sharma^{*1} Eric Mitchell^{*1}
Stefano Ermon^{1‡} Christopher D. Manning¹ Chelsea Finn¹

¹Stanford University [‡]CZ Biohub
{rafailov,architsh,eric.mitchell}@cs.stanford.edu

Abstract

While large-scale unsupervised language models (LMs) learn broad world knowledge and some reasoning skills, achieving precise control of their behavior is difficult due to the completely unsupervised nature of their training. Existing methods for gaining such steerability collect human labels of the relative quality of model generations and fine-tune the unsupervised LM to align with these preferences, often with reinforcement learning from human feedback (RLHF). However, RLHF is a complex and often unstable procedure, first fitting a reward model that reflects the human preferences, and then fine-tuning the large unsupervised LM using reinforcement learning to maximize this estimated reward without drifting too far from the original model. In this paper we introduce a new parameterization of the reward model in RLHF that enables extraction of the corresponding optimal policy in closed form, allowing us to solve the standard RLHF problem with only a simple classification loss. The resulting algorithm, which we call *Direct Preference Optimization* (DPO), is stable, performant, and computationally lightweight, eliminating the need for sampling from the LM during fine-tuning or performing significant hyperparameter tuning. Our experiments show that DPO can fine-tune LMs to align with human preferences as well as or better than existing methods. Notably, fine-tuning with DPO exceeds PPO-based RLHF in ability to control sentiment of generations, and matches or improves response quality in summarization and single-turn dialogue while being substantially simpler to implement and train.

DPO loss

$$-\log \sigma \left(\beta \log \frac{\pi_{\theta}(y_w|x)}{\pi_{\text{ref}}(y_w|x)} - \beta \log \frac{\pi_{\theta}(y_l|x)}{\pi_{\text{ref}}(y_l|x)} \right)$$

DPO variants

Method	Objective
RRHF [91]	$\max \left(0, -\frac{1}{ \mathcal{D}_w } \log \pi_{\theta}(y_w x) + \frac{1}{ \mathcal{D}_l } \log \pi_{\theta}(y_l x) \right) - \lambda \log \pi_{\theta}(y_w x)$
SLiC-HF [96]	$\max \left(0, \delta - \log \pi_{\theta}(y_w x) + \log \pi_{\theta}(y_l x) \right) - \lambda \log \pi_{\theta}(y_w x)$
DPO [66]	$-\log \sigma \left(\beta \log \frac{\pi_{\theta}(y_w x)}{\pi_{\text{ref}}(y_w x)} - \beta \log \frac{\pi_{\theta}(y_l x)}{\pi_{\text{ref}}(y_l x)} \right)$
IPO [6]	$\left(\log \frac{\pi_{\theta}(y_w x)}{\pi_{\text{ref}}(y_w x)} - \log \frac{\pi_{\theta}(y_l x)}{\pi_{\text{ref}}(y_l x)} - \frac{1}{\beta r} \right)^2$
CPO [88]	$-\log \sigma \left(\beta \log \pi_{\theta}(y_w x) - \beta \log \pi_{\theta}(y_l x) \right) - \lambda \log \pi_{\theta}(y_w x)$
KTO [29]	$-\lambda_w \sigma \left(\beta \log \frac{\pi_{\theta}(y_w x)}{\pi_{\text{ref}}(y_w x)} - z_{\text{ref}} \right) + \lambda_l \sigma \left(z_{\text{ref}} - \beta \log \frac{\pi_{\theta}(y_l x)}{\pi_{\text{ref}}(y_l x)} \right)$, where $z_{\text{ref}} = \mathbb{E}_{(x,y) \sim \mathcal{D}} [\beta \text{KL}(\pi_{\theta}(y x) \pi_{\text{ref}}(y x))]$
ORPO [42]	$-\log p_{\theta}(y_w x) - \lambda \log \sigma \left(\log \frac{p_{\theta}(y_w x)}{1 - p_{\theta}(y_w x)} - \log \frac{p_{\theta}(y_l x)}{1 - p_{\theta}(y_l x)} \right)$, where $p_{\theta}(y x) = \exp \left(\frac{1}{ \mathcal{V} } \log \pi_{\theta}(y x) \right)$
R-DPO [64]	$-\log \sigma \left(\beta \log \frac{\pi_{\theta}(y_w x)}{\pi_{\text{ref}}(y_w x)} - \beta \log \frac{\pi_{\theta}(y_l x)}{\pi_{\text{ref}}(y_l x)} + (\alpha y_w - \alpha y_l) \right)$
SimPO	$-\log \sigma \left(\frac{\beta}{ \mathcal{D}_w } \log \pi_{\theta}(y_w x) - \frac{\beta}{ \mathcal{D}_l } \log \pi_{\theta}(y_l x) - \gamma \right)$

from Meng et al. (2024)

- ▶ **No reference approaches** (e.g., CPO, ORPO, *only involves a single model*) versus multi-model, **reference approaches** (DPO).

The many varieties of DPO

Direct Preference Optimization: Your Language Model is Secretly a Reward Model

Rafael Rafailov^{*1} Archit Sharma^{*1} Eric Mitchell^{*2}
Stefano Ermon^{1‡} Christopher D. Manning¹ Chelsea Finn¹

¹Stanford University ²CZ Biohub
{rafaill, architsh, eric.mitchell}@cs.stanford.edu

Abstract

While large-scale unsupervised language models (LMs) learn broad world knowledge and some reasoning skills, achieving precise control of their behavior is difficult due to the completely unsupervised nature of their training. Existing methods for gaining such steerability collect human labels of the relative quality of model generations and fine-tune the unsupervised LM to align with these preferences, often with reinforcement learning from human feedback (RLHF). However, RLHF is a complex and often unstable procedure, first fitting a reward model that reflects the human preferences, and then fine-tuning the large unsupervised LM using reinforcement learning to maximize this estimated reward without drifting too far from the original model. In this paper we introduce a new parameterization of the reward model in RLHF that enables extraction of the corresponding optimal policy in closed form, allowing us to solve the standard RLHF problem with only a simple classification loss. The resulting algorithm, which we call *Direct Preference Optimization* (DPO), is stable, performant, and computationally lightweight, eliminating the need for sampling from the LM during fine-tuning or performing significant hyperparameter tuning. Our experiments show that DPO can fine-tune LMs to align with human preferences as well as or better than existing methods. Notably, fine-tuning with DPO exceeds PPO-based RLHF in ability to control sentiment of generations, and matches or improves response quality in summarization and single-turn dialogue while being substantially simpler to implement and train.

DPO loss

$$-\log \sigma \left(\beta \log \frac{\pi_{\theta}(y_w|x)}{\pi_{\text{ref}}(y_w|x)} - \beta \log \frac{\pi_{\theta}(y_l|x)}{\pi_{\text{ref}}(y_l|x)} \right)$$

DPO variants

Method	Objective
RRHF [91]	$\max \left(0, -\frac{1}{ \beta_w } \log \pi_{\theta}(y_w x) + \frac{1}{ \beta_l } \log \pi_{\theta}(y_l x) \right) - \lambda \log \pi_{\theta}(y_w x)$
SLiC-HF [96]	$\max \left(0, \delta - \log \pi_{\theta}(y_w x) + \log \pi_{\theta}(y_l x) \right) - \lambda \log \pi_{\theta}(y_w x)$
DPO [66]	$-\log \sigma \left(\beta \log \frac{\pi_{\theta}(y_w x)}{\pi_{\text{ref}}(y_w x)} - \beta \log \frac{\pi_{\theta}(y_l x)}{\pi_{\text{ref}}(y_l x)} \right)$
IPO [6]	$\left(\log \frac{\pi_{\theta}(y_w x)}{\pi_{\text{ref}}(y_w x)} - \log \frac{\pi_{\theta}(y_l x)}{\pi_{\text{ref}}(y_l x)} - \frac{1}{\beta r} \right)^2$
CPO [88]	$-\log \sigma \left(\beta \log \pi_{\theta}(y_w x) - \beta \log \pi_{\theta}(y_l x) \right) - \lambda \log \pi_{\theta}(y_w x)$
KTO [29]	$-\lambda_w \sigma \left(\beta \log \frac{\pi_{\theta}(y_w x)}{\pi_{\text{ref}}(y_w x)} - z_{\text{ref}} \right) + \lambda_l \sigma \left(z_{\text{ref}} - \beta \log \frac{\pi_{\theta}(y_l x)}{\pi_{\text{ref}}(y_l x)} \right)$, where $z_{\text{ref}} = \mathbb{E}_{(x,y) \sim \mathcal{D}} [\beta \text{KL}(\pi_{\theta}(y x) \pi_{\text{ref}}(y x))]$
ORPO [42]	$-\log p_{\theta}(y_w x) - \lambda \log \sigma \left(\log \frac{p_{\theta}(y_w x)}{1 - p_{\theta}(y_w x)} - \log \frac{p_{\theta}(y_l x)}{1 - p_{\theta}(y_l x)} \right)$, where $p_{\theta}(y x) = \exp \left(\frac{1}{ \beta } \log \pi_{\theta}(y x) \right)$
R-DPO [64]	$-\log \sigma \left(\beta \log \frac{\pi_{\theta}(y_w x)}{\pi_{\text{ref}}(y_w x)} - \beta \log \frac{\pi_{\theta}(y_l x)}{\pi_{\text{ref}}(y_l x)} + (\alpha y_w - \alpha y_l) \right)$
SimPO	$-\log \sigma \left(\frac{\beta}{ \beta_w } \log \pi_{\theta}(y_w x) - \frac{\beta}{ \beta_l } \log \pi_{\theta}(y_l x) - \gamma \right)$

from Meng et al. (2024)

Questions: How are all these variations related to one another, nature of the space of losses?

The many varieties of DPO

Direct Preference Optimization: Your Language Model is Secretly a Reward Model

Rafael Rafailov^{1*} Archit Sharma^{2*} Eric Mitchell^{2*}
Stefano Ermon^{2†} Christopher D. Manning² Chelsea Finn¹

¹Stanford University ²CZ Biohub
(rafailov,architsh,eric.mitchell)@cs.stanford.edu

Abstract

While large-scale unsupervised language models have reached the state-of-the-art on many tasks, edge and some reasoning tasks remain difficult due to the completely unsupervised nature of their training. Existing methods for gaining such capability either harness human labels of the relative quality of model generations and fine-tune the unsupervised LM to align with these preferences, often with reinforcement learning from human feedback (RLHF). However, RLHF is a complex and often unstable procedure, first fitting a reward model that reflects the human preferences, and then fine-tuning the large unsupervised LM using reinforcement learning to maximize this estimated reward without drifting too far from the original model. In this paper we introduce a new parameterization of the reward model in RLHF that enables extraction of the corresponding optimal policy in closed form, allowing us to solve the standard RLHF problem with only a simple classification loss. The resulting algorithm, which we call *Direct Preference Optimization* (DPO), is stable, performant, and computationally lightweight, eliminating the need for sampling from the LM during fine-tuning or performing significant hyperparameter tuning. Our experiments show that DPO can fine-tune LMs to align with human preferences as well as or better than existing methods. Notably, fine-tuning with DPO exceeds PPO-based RLHF in ability to control sentiment of generations, and matches or improves response quality in summarization and single-turn dialogue while being substantially simpler to implement and train.

DPO loss

$$-\log \sigma \left(\beta \log \frac{\pi_{\theta}(y_w|x)}{\pi_{\theta}(y_l|x)} - \beta \log \frac{\pi_{\theta}(y_l|x)}{\pi_{\theta}(y_r|x)} \right)$$

DPO variants

Method	Objective
GRPO [42]	$-\log \pi_{\theta}(y_w x) + \frac{1}{\beta} \log \pi_{\theta}(y_l x) + \frac{1}{\beta} \log \pi_{\theta}(y_r x) - \lambda \log \pi_{\theta}(p_w x)$
DPO [66]	$-\log \sigma \left(\beta \log \frac{\pi_{\theta}(y_w x)}{\pi_{\theta}(y_l x)} - \beta \log \frac{\pi_{\theta}(y_l x)}{\pi_{\theta}(y_r x)} \right)$
DPO [6]	$\left(\log \frac{\pi_{\theta}(y_w x)}{\pi_{\theta}(y_l x)} - \log \frac{\pi_{\theta}(y_l x)}{\pi_{\theta}(y_r x)} - \beta \right)^2$
CPO [88]	$-\log \sigma \left(\beta \log \pi_{\theta}(p_w x) - \beta \log \pi_{\theta}(p_l x) \right) - \lambda \log \pi_{\theta}(p_w x)$
KTO [29]	$-\lambda_{\text{opt}} \sigma \left(\beta \log \frac{\pi_{\theta}(p_w x)}{\pi_{\theta}(p_l x)} - \beta_{\text{opt}} \right) + \lambda_{\text{opt}} \left(\lambda_{\text{opt}} - \beta \log \frac{\pi_{\theta}(p_w x)}{\pi_{\theta}(p_l x)} \right)$, where $\lambda_{\text{opt}} = E_{p_w, p_l \sim p} [\text{SRL}(\pi_{\theta}(p_w x), \pi_{\theta}(p_l x))]$
GRPO [42]	$-\log \pi_{\theta}(p_w x) - \lambda \log \sigma \left(\log \sqrt{\frac{\pi_{\theta}(p_w x)}{\pi_{\theta}(p_l x)}} - \log \sqrt{\frac{\pi_{\theta}(p_l x)}{\pi_{\theta}(p_r x)}} \right)$, where $p_{\theta}(p x) = \exp \left(\frac{1}{\beta} \log \pi_{\theta}(p x) \right)$
R-DPO [64]	$-\log \sigma \left(\beta \log \frac{\pi_{\theta}(p_w x)}{\pi_{\theta}(p_l x)} - \beta \log \frac{\pi_{\theta}(p_l x)}{\pi_{\theta}(p_r x)} + (\alpha(p_w) - \alpha(p_l)) \right)$
SimPO	$-\log \sigma \left(\frac{\beta}{\alpha(p_l)} \log \pi_{\theta}(p_w x) - \frac{\beta}{\alpha(p_l)} \log \pi_{\theta}(p_l x) - \gamma \right)$

from Meng et al. (2024)

Formalization of preference losses

The many varieties of DPO

Direct Preference Optimization: Your Language Model is Secretly a Reward Model

Rafael Rafailov^{1*} Archit Sharma^{2*} Eric Mitchell^{2*}
Stefano Ermon^{2†} Christopher D. Manning² Chelsea Finn²

¹Stanford University ²CZ Biohub
(rafailov,architsh,eric.mitchell)@cs.stanford.edu

Abstract

While edge models are increasingly deployed in production, training methods for gaining such extremely scarce human feedback of the relative quality of model generations and fine-tune the unsupervised LM to align with these preferences, often with reinforcement learning from human feedback (RLHF). However, RLHF is a complex and often unstable procedure, first fitting a reward model that reflects the human preferences, and then fine-tuning the large unsupervised LM using reinforcement learning to maximize this estimated reward without drifting too far from the original model. In this paper we introduce a new parameterization of the reward model in RLHF that enables extraction of the corresponding optimal policy in closed form, allowing us to solve the standard RLHF problem with only a simple classification loss. The resulting algorithm, which we call *Direct Preference Optimization* (DPO), is stable, performant, and computationally lightweight, eliminating the need for sampling from the LM during fine-tuning or performing significant hyperparameter tuning. Our experiments show that DPO can fine-tune LMs to align with human preferences as well as or better than existing methods. Notably, fine-tuning with DPO exceeds PPO-based RLHF in ability to control sentiment of generations, and matches or improves response quality in summarization and single-turn dialogue while being substantially simpler to implement and train.

DPO loss

$$-\log \sigma \left(\beta \log \frac{\pi_{\theta}(y_w|x)}{\pi_{\theta}(y_l|x)} - \beta \log \frac{\pi_{\theta}(y_l|x)}{\pi_{\theta}(y_r|x)} \right)$$

DPO variants

Method	Objective
GRPO [42]	$-\log \pi_{\theta}(y_w x) - \lambda \log \sigma \left(\log \frac{\pi_{\theta}(y_w x)}{\pi_{\theta}(y_l x)} - \log \frac{\pi_{\theta}(y_l x)}{\pi_{\theta}(y_r x)} \right)$
DPO [66]	$-\log \sigma \left(\beta \log \frac{\pi_{\theta}(y_w x)}{\pi_{\theta}(y_l x)} - \beta \log \frac{\pi_{\theta}(y_l x)}{\pi_{\theta}(y_r x)} \right)$
DPO [61]	$\left(\log \frac{\pi_{\theta}(y_w x)}{\pi_{\theta}(y_l x)} - \log \frac{\pi_{\theta}(y_l x)}{\pi_{\theta}(y_r x)} - \beta \right)^2$
CPO [88]	$-\log \sigma \left(\beta \log \pi_{\theta}(y_w x) - \beta \log \pi_{\theta}(y_l x) \right) - \lambda \log \pi_{\theta}(y_r x)$
KTO [29]	$-\lambda_{\text{ref}} \sigma \left(\beta \log \frac{\pi_{\theta}(y_w x)}{\pi_{\theta}(y_l x)} - \beta_{\text{ref}} \right) + \lambda_{\text{ref}} \left(\beta_{\text{ref}} - \beta \log \frac{\pi_{\theta}(y_w x)}{\pi_{\theta}(y_l x)} \right)$ where $\beta_{\text{ref}} = E_{p(y x)} [\text{DRL}(\pi_{\theta}(y x), \pi_{\text{ref}}(y x))]$
GRPO [42]	$-\log \pi_{\theta}(y_w x) - \lambda \log \sigma \left(\log \frac{\pi_{\theta}(y_w x)}{\pi_{\theta}(y_l x)} - \log \frac{\pi_{\theta}(y_l x)}{\pi_{\theta}(y_r x)} \right)$ where $\pi_{\theta}(y x) = \exp \left(\frac{1}{\beta} \log \pi_{\theta}(y x) \right)$
R-DPO [64]	$-\log \sigma \left(\beta \log \frac{\pi_{\theta}(y_w x)}{\pi_{\theta}(y_l x)} - \beta \log \frac{\pi_{\theta}(y_l x)}{\pi_{\theta}(y_r x)} \right) + (\alpha(y_w) - \alpha(y_l))$
SimPO	$-\log \sigma \left(\frac{\beta}{\alpha(y_l)} \log \pi_{\theta}(y_w x) - \frac{\beta}{\alpha(y_l)} \log \pi_{\theta}(y_l x) - \gamma \right)$

from Meng et al. (2024)

Going away from these opaque equations

Preference learning as a discrete reasoning problem

Loss Function ℓ

$$-\log \sigma \left(\log \frac{\pi_{\theta}(y_w|x)}{\pi_{\text{ref}}(y_w|x)} - \log \frac{\pi_{\theta}(y_l|x)}{\pi_{\text{ref}}(y_l|x)} \right)$$

Preference learning as a discrete reasoning problem

Loss Function ℓ

$$-\log \sigma \left(\log \frac{\pi_{\theta}(y_w|x)}{\pi_{\text{ref}}(y_w|x)} - \log \frac{\pi_{\theta}(y_l|x)}{\pi_{\text{ref}}(y_l|x)} \right)$$

Two models, four predictions

Preference learning as a discrete reasoning problem

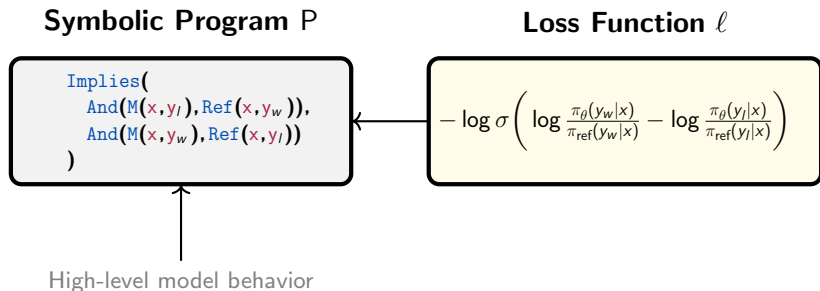
Loss Function ℓ

$$-\log \sigma \left(\log \frac{\pi_{\theta}(y_w|x)}{\pi_{\text{ref}}(y_w|x)} - \log \frac{\pi_{\theta}(y_l|x)}{\pi_{\text{ref}}(y_l|x)} \right)$$

Two models, four predictions

- **Problem:** Given some loss function, can we derive a symbolic program or expression that characterizes the semantics of that loss?

Preference learning as a discrete reasoning problem



- **Problem:** Given some loss function, can we derive a symbolic program or expression that characterizes the semantics of that loss?

Preference learning as a discrete reasoning problem

Symbolic Program P

```
Implies(  
  And(M(x, yI), Ref(x, yw)),  
  And(M(x, yw), Ref(x, yI))  
)
```

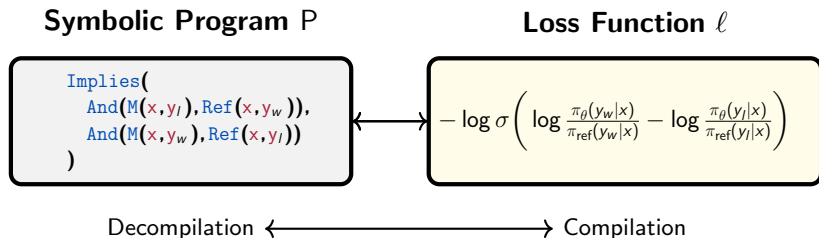
Loss Function ℓ

$$-\log \sigma \left(\log \frac{\pi_{\theta}(y_w|x)}{\pi_{\text{ref}}(y_w|x)} - \log \frac{\pi_{\theta}(y_I|x)}{\pi_{\text{ref}}(y_I|x)} \right)$$

Decompilation $\longleftrightarrow$ Compilation

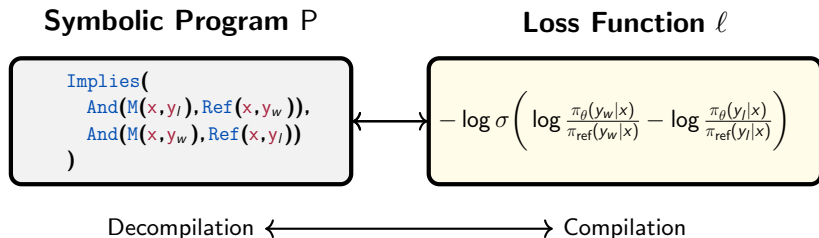
- **Problem:** Given some loss function, can we derive a symbolic program or expression that characterizes the semantics of that loss?

Preference learning as a discrete reasoning problem



- **Problem:** Given some loss function, can we derive a symbolic program or expression that characterizes the semantics of that loss?

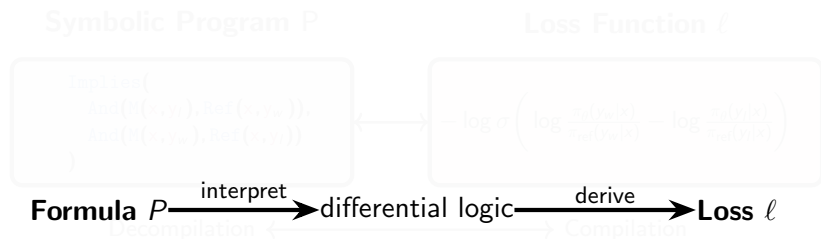
Preference learning as a discrete reasoning problem



- **Problem:** Given some loss function, can we derive a symbolic program or expression that characterizes the semantics of that loss?

1. **Compilation:** Translating specifications into loss, well studied.

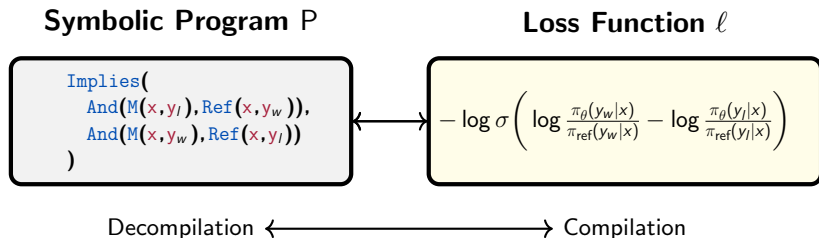
Preference learning as a discrete reasoning problem



- **Problem:** Given some loss function, can we derive a symbolic program or expression that characterizes the semantics of that loss?

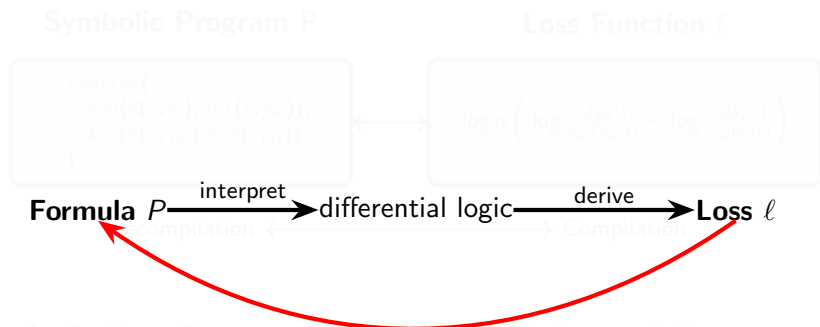
1. **Compilation:** Translating specifications into loss, well studied.

Preference learning as a discrete reasoning problem



- **Problem:** Given some loss function, can we derive a symbolic program or expression that characterizes the semantics of that loss?
 1. **Compilation:** Translating specifications into loss, well studied.
 2. **Decompilation:** Losses to specifications (inverse), less explored.

Preference learning as a discrete reasoning problem



- **Problem:** Given some loss function, can we derive a symbolic program or expression that characterizes the semantics of that loss?

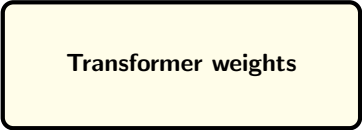
Loss as logic (Richardson et al., 2025)

1. **Compilation:** Translating specifications into loss, well studied.
2. **Decompilation:** Losses to specifications (inverse), less explored.

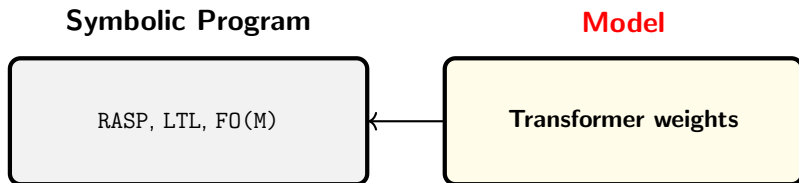
Formal analysis via decompilation in general

Model

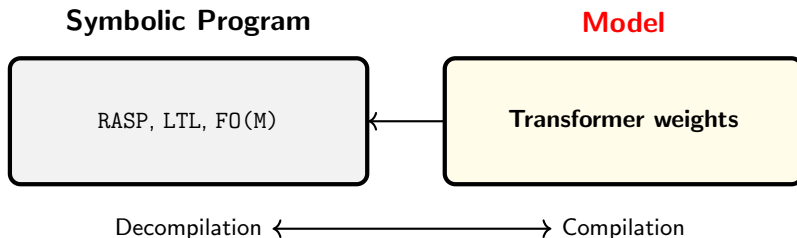
Transformer weights



Formal analysis via decompilation in general



Formal analysis via decompilation in general



- We know what the *target languages* are (Weiss et al., 2021; Merrill and Sabharwal, 2023; Yang and Chiang, 2024), how to compile, decompile (Friedman et al., 2023).

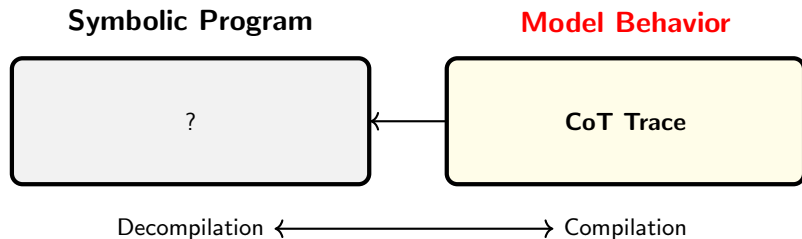
Formal analysis via decompilation in general

Model Behavior

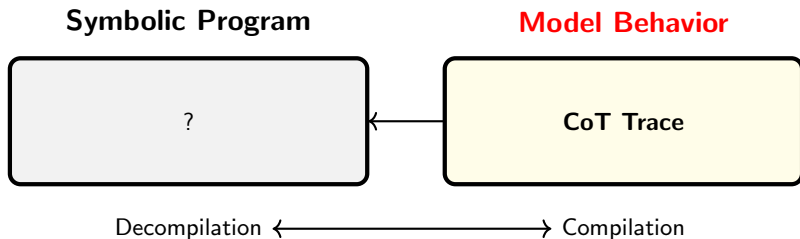


CoT Trace

Formal analysis via decompilation in general



Formal analysis via decompilation in general



- Not always clear what the target programming language is or should be.

Language model programming: the languages and formal interfaces used for for doing such analysis ([Richardson and Wijnholds, 2025](#)).

Language model programming: ESSLLI 2025

Lecturers

[Kyle Richardson](#) (Allen Institute for AI)

[Gijs Wijnholds](#) (Leiden Institute of Advanced Computer Science)

Slides

[lecture 1](#): course overview, language modeling basics, [transformers](#), [RASP](#).

[lecture 2](#): declarative approaches to model training and fine-tuning, the [semantic loss](#) and [weighted model counting](#), [other](#) approaches.

[lecture 3](#): high-level programming techniques for [direct preference alignment](#) and [LLM alignment](#), [formal characterizations](#) of known loss functions.

[lecture 4](#): [declarative and probabilistic approaches](#) to test-time inference, [LLM self-correction](#), [consistency](#), distilling LLMs to tractable models, [logic programming](#).

[lecture 5](#): Advanced prompting, [chain-of-thought](#), [imperative model programming](#), [\(discrete\) probabilistic programming](#).

background [logic notes](#), [extended notes on transformers](#)

extra lectures [Prompting as programming](#), [Grammar-constrained decoding](#)

Helpful Resources

Below are some pointers to code resources:

- languages [\[scallop\]](#), [\[problog\]](#), [\[pyDatalog\]](#), [\[lmql\]](#), [\[rasp\]](#), [\[NumPy RASP\]](#), [\[deepproblog\]](#)
- automated reasoning tools/circuits [\[Z3 solver\]](#), [\[python-sat\]](#), [\[pysdd\]](#), [\[circuit\]](#)
- NLP and general ML [\[transformers\]](#), [\[PyTorch\]](#), [\[pylon-lib\]](#), [\[hf datasets\]](#), [\[hf hub\]](#)
- other useful utilities [\[sympy\]](#)

Useful tutorials: [Transformers from scratch](#) (some examples/ideas used in lecture 1), [Lectures on Probabilistic Programming](#), [Tractable Probabilistic Models](#)

<https://github.com/yakazimir/LMProgramming>

Lecturers

[Ilya Richardson](#) (Allen Institute for AI)

[Rita Vihitahy](#) (Leiden Institute of Advanced Computer Science)

Slides

[lecture 1](#): course overview, language modeling basics, [transformers](#), [RASP](#)

[lecture 2](#): declarative approaches to model training and fine-tuning, the [cross-entropy loss](#) and [weighted model counting](#), other approaches.

[lecture 3](#): high-level programming techniques for [direct preference alignment](#) and [LLM alignment](#), formal characterizations of known loss functions.

[lecture 4](#): [declarative and probabilistic approaches](#) to test-time inference, [LLM self-correction](#), [consistency](#), distilling LLMs to tractable models, [logic programming](#).

[background](#) [logic](#), [syntax](#), [semantics](#), [losses](#), [on transformers](#)

[extra lectures](#) [Translating to programming](#), [Quantifier elimination](#), [decoding](#)

Helpful Resources

Below are some pointers to code resources:

- languages: [pytorch](#), [pytorch-lightning](#), [pytorch-jit](#), [pytorch-train](#), [PyTorch Deep Learning](#)
- automated reasoning toolchains: [Z3 solver](#), [python-smt](#), [smtlib](#), [smtlib2](#)
- NLP and general ML: [transformers](#), [PyTorch](#), [datautils](#), [NLTK](#), [NLTK](#), [NLTK](#)
- other useful utilities: [pymc3](#)

Useful tutorials: [Transformers from scratch](#) (some examples/ideas used in lecture 1), [Lectures on Probabilistic Programming](#), [Tractable Probabilistic Models](#)

What is the right programming language for preference?

<https://github.com/yakazimir/LMProgramming>

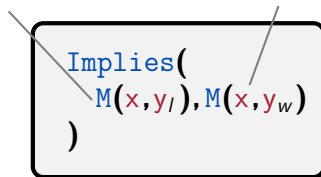
What do these programs tell us?

```
Implies(  
  M(x, yl), M(x, yw)  
)
```

What do these programs tell us?

Model predicts loser

Model predicts winner



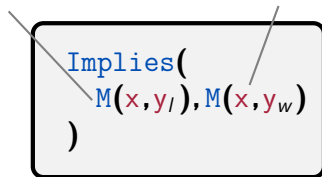
```
Implies(  
  M(x, yl), M(x, yw)  
)
```

Conceptually: Model predications are logical propositions, Boolean variables inside of formulas, weighted by prediction probability.

What do these programs tell us?

Model predicts loser

Model predicts winner



$$w(\mathbf{M}(\mathbf{x}, \mathbf{y})) = \pi_{\mathbf{M}}(\mathbf{y} \mid \mathbf{x})$$

Conceptually: Model predications are logical propositions, Boolean variables inside of formulas, weighted by prediction probability.

What do these programs tell us?

Model predicts loser Model predicts winner

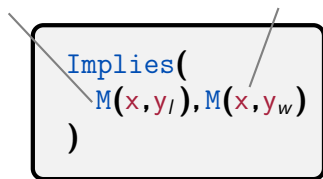
```
Implies(  
  M(x, yl), M(x, yw)  
)
```

Whenever the model deems the loser to be a valid generation, it should deem the winner to be valid too.

Conceptually: Predictions are connected through Boolean operators, express constraints on predictions; ρ_θ as formulas.

Uncovering the natural logic of these algorithms

Model predicts loser Model predicts winner

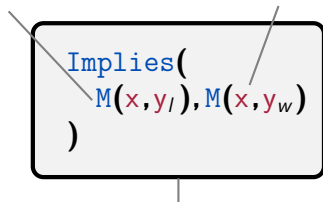


Whenever the model deems the loser to be a valid generation, it should deem the winner to be valid too.

Assumption: Every loss function has an internal logic that can be expressed in this way, we want to uncover that logic.

Uncovering the natural logic of these algorithms

Model predicts loser Model predicts winner



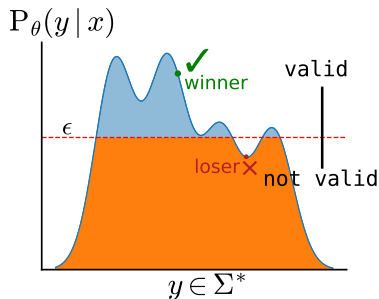
Whenever the model deems the loser to be a valid generation, it should deem the winner to be valid too.

Running example: This program and semantics is foundational to many DPO-style losses.

Uncovering the natural logic of these algorithms

```
Implies(  
  M(x, y_l), M(x, y_w)  
)
```

Whenever the model deems the loser to be a valid generation, it should deem the winner to be valid too.

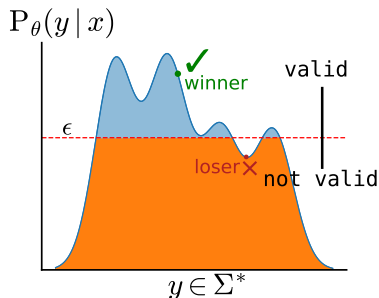


Model behavior: Programs tell us about the structure of the model's output distribution (right).

Uncovering the natural logic of these algorithms

```
Implies(  
  M(x, yl), M(x, yw)  
)
```

```
And(  
  M(x, yw),  
  Not(M(x, yl)))
```

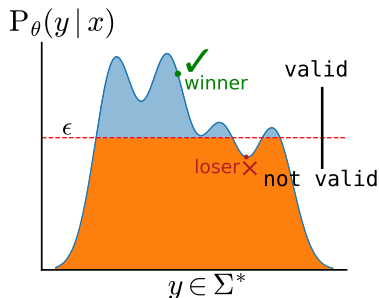


Model behavior: Programs tell us about the structure of the model's output distribution (right).

Uncovering the natural logic of these algorithms

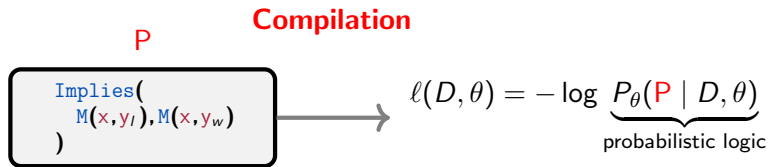
```
Implies(  
  M(x, yl), M(x, yw)  
)
```

```
And(  
  M(x, yw),  
  Not(M(x, yl)))
```



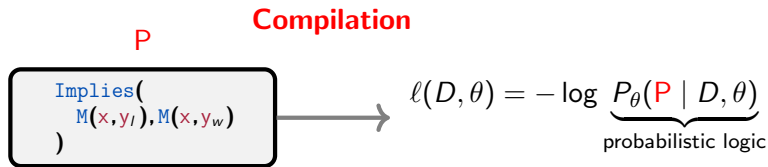
Observation: The second program is more strict than the first, involves semantic entailment.

Compilation and decompilation again



Whenever the model deems the loser to be a valid generation, it should deem the winner to be valid too.

Compilation and decompilation again

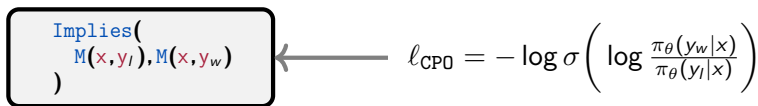


Whenever the model deems the loser to be a valid generation, it should deem the winner to be valid too.

What we did: defined a novel probabilistic logic for preference modeling, interpret formulas in that logic to derive differentiable losses.

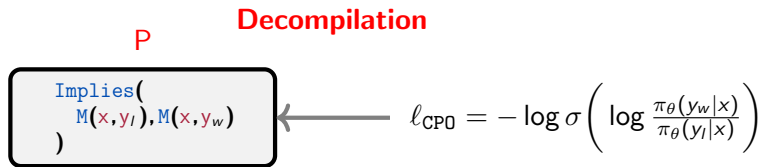
Compilation and decompilation again

P **Decompilation**



Whenever the model deems the loser to be a valid generation, it should deem the winner to be valid too.

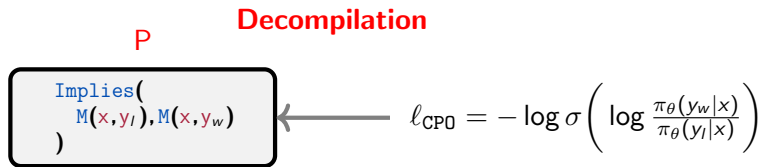
Compilation and decompilation again



Whenever the model deems the loser to be a valid generation, it should deem the winner to be valid too.

$$\underbrace{\ell_{\text{CP0}}(D, \theta) = -\log P_{\theta}(\textcolor{red}{P} \mid D, \theta)}_{\text{correctness property}}$$

Compilation and decompilation again



Whenever the model deems the loser to be a valid generation, it should deem the winner to be valid too.

$$\underbrace{\ell_{\text{CPD}}(D, \theta) = -\log P_{\theta}(\mathbf{P} \mid D, \theta)}_{\text{correctness property}}$$

The second thing we did: Defined a mechanical procedure for decompilation, proved its correctness, invariance to choice of f .

Illustration of approach and results (Richardson et al., 2025)

Input Loss ℓ_{ORPO}

$$-\log \sigma \left(\log \frac{\text{Odds}_{\theta}(y_w|x)}{\text{Odds}_{\theta}(y_l|x)} \right)$$

Illustration of approach and results (Richardson et al., 2025)

Input Loss ℓ_{ORPO}

$$-\log \sigma \left(\log \frac{\text{Odds}_{\theta}(y_w|x)}{\text{Odds}_{\theta}(y_l|x)} \right)$$



$$\frac{\rho_{\theta}^t}{\rho_{\theta}^b} = \frac{P_{\theta}(y_w|x)(1-P_{\theta}(y_l|x))}{P_{\theta}(y_l|x)(1-P_{\theta}(y_w|x))}$$

Core loss equation

Illustration of approach and results (Richardson et al., 2025)

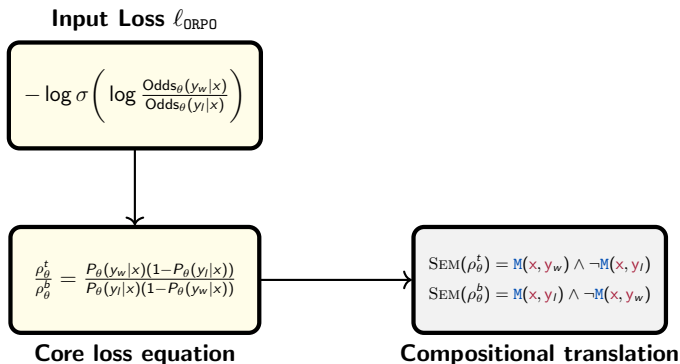


Illustration of approach and results (Richardson et al., 2025)

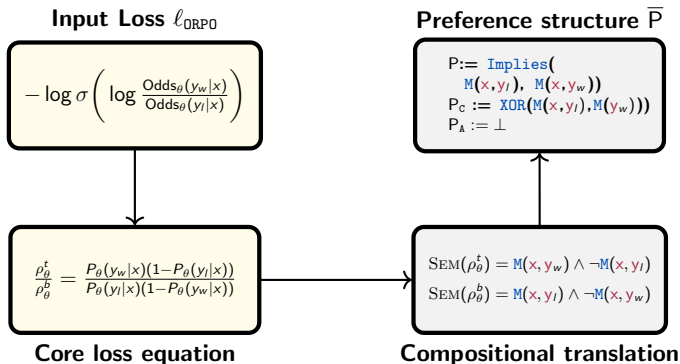


Illustration of approach and results (Richardson et al., 2025)

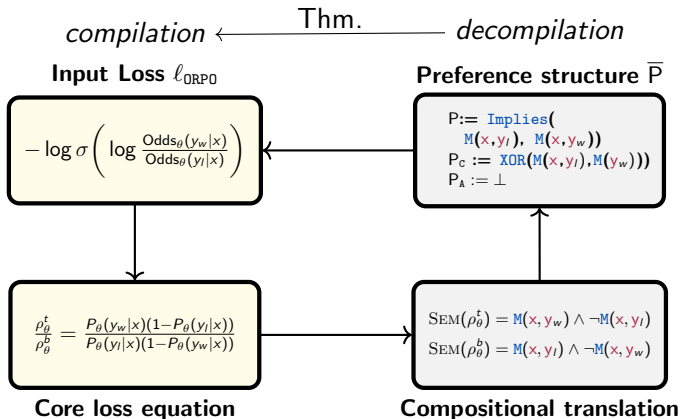
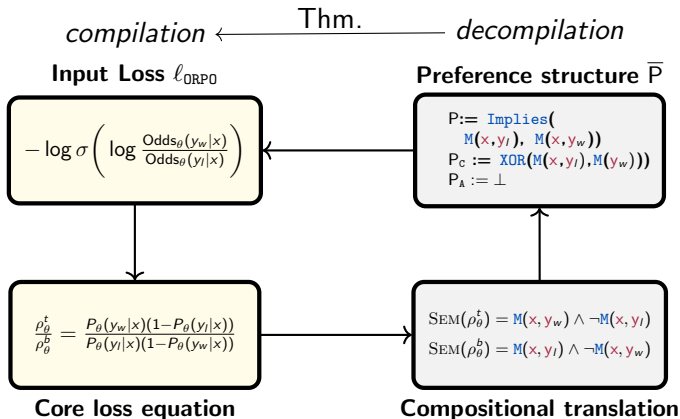
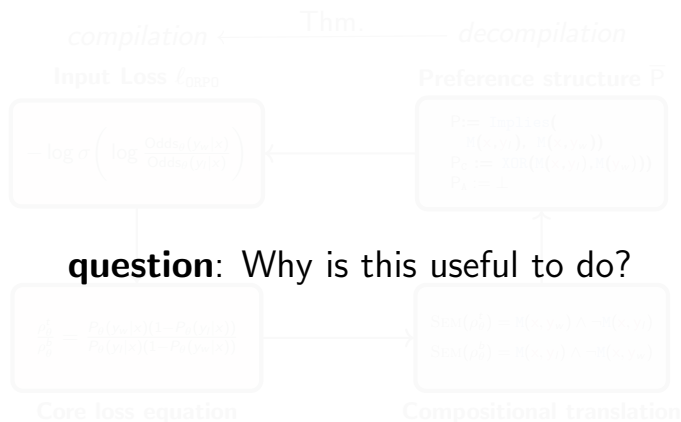


Illustration of approach and results (Richardson et al., 2025)



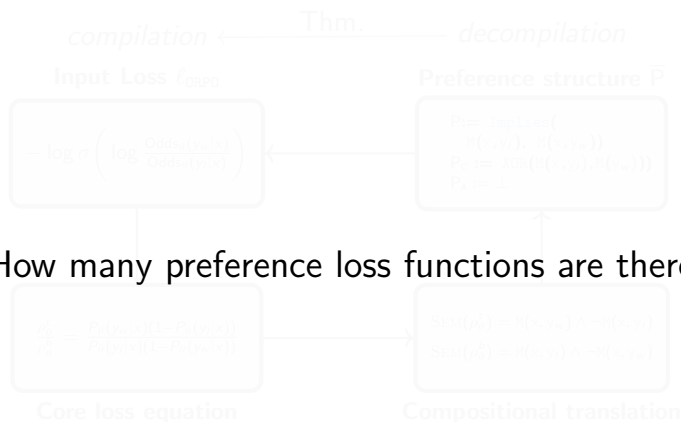
- **Preference structure**, a core construct in our logic, encoding for preference losses, has a natural Boolean interpretation.

Illustration of approach and results (Richardson et al., 2025)



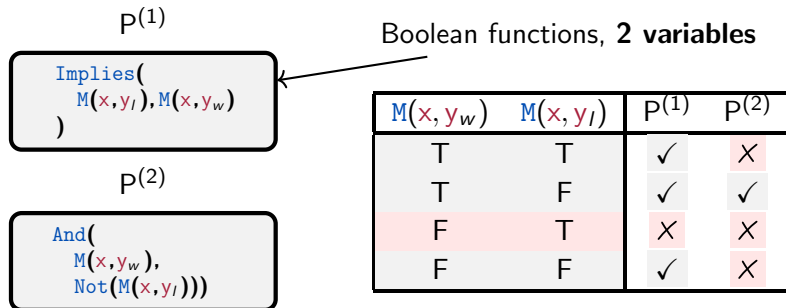
- Preference structure, a core construct in our logic, encoding for preference losses, has a natural Boolean interpretation.

Illustration of approach and results (Richardson et al., 2025)

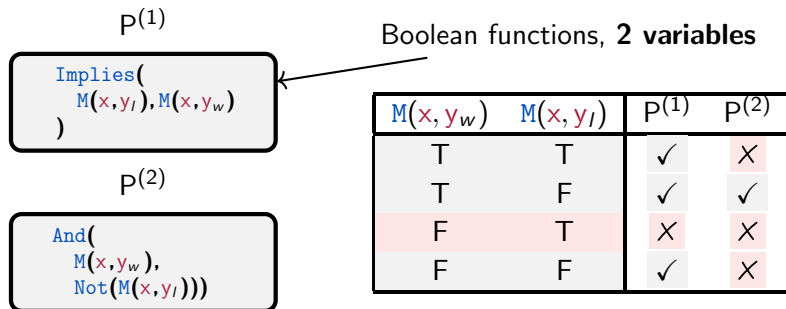


- Preference structure, a core construct in our logic, encoding for preference losses, has a natural Boolean interpretation.

Why is this useful? understanding the space

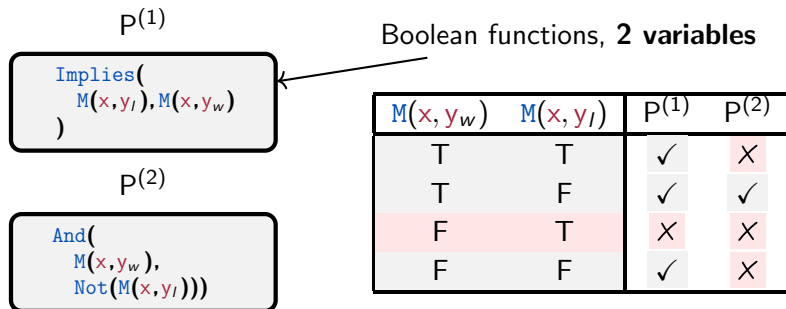


Why is this useful? understanding the space



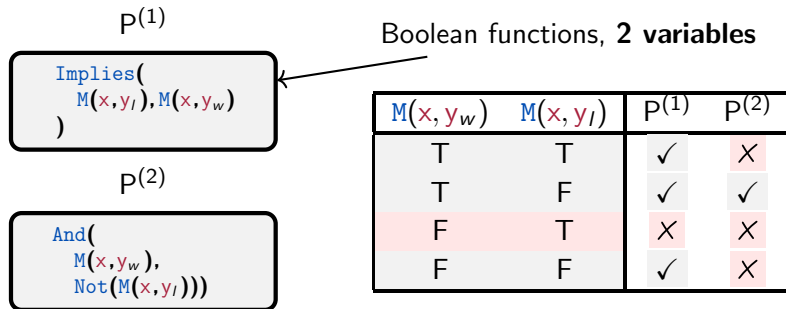
- ▶ Every program (in our logic) is pair of Boolean functions (in n variables),
corr. to ✓ and X, leads to 4^{2^n} possible loss functions.

Why is this useful? understanding the space



Loss creation will end up being equivalent to drawing different sets of ✓ s and X (or blank marks) in a truth table.

Why is this useful? understanding the space



no reference: 256 losses

Loss creation will end up being equivalent to drawing different sets of ✓ s and X (or blank marks) in a truth table.

Loss functions as truth tables

```
Implies(  
  And(M(x,yI),Ref(x,yw)),  
  And(M(x,yw),Ref(x,yI))  
)
```

4 variables

Ref(x,y _w)	M(x,y _I)	Ref(x,y _I)	M(x,y _w)
F	F	F	F
F	F	F	T
F	F	T	F
F	F	T	T
F	T	F	F
F	T	F	T
F	T	T	F
F	T	T	T
T	F	F	F
T	F	F	T
T	F	T	F
T	F	T	T
T	T	F	F
T	T	F	T
T	T	T	F
T	T	T	T

w/ reference: 4,294,967,296 losses

Loss functions as truth tables

```
Implies(  
  And(M(x,yI),Ref(x,yw)),  
  And(M(x,yw),Ref(x,yI))  
)
```

4 variables

answer: loads.

Ref(x,y _w)	M(x,y _I)	Ref(x,y _I)	M(x,y _w)
F	F	F	F
F	F	F	T
F	F	T	F
F	F	T	T
F	T	F	F
F	T	F	T
F	T	T	F
F	T	T	T
T	F	F	F
T	F	F	T
T	F	T	F
T	F	T	T
T	T	F	F
T	T	F	T
T	T	T	F
T	T	T	T

w/ reference: 4,294,967,296 losses

Loss functions as truth tables

```
Implies(  
  And(M(x,y_l),Ref(x,y_w)),  
  And(M(x,y_w),Ref(x,y_l))  
)
```

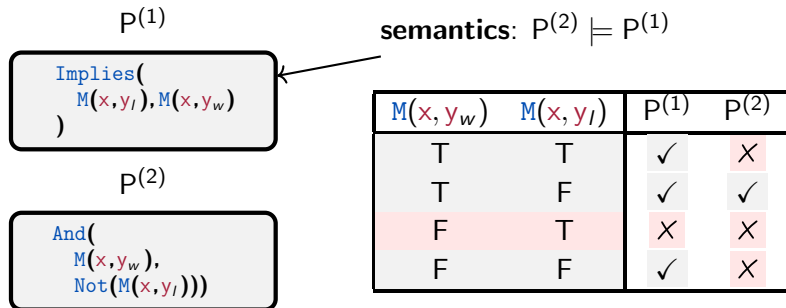
question: How are losses related to one another?

4 variables

Ref(x,y_w)	M(x,y_l)	Ref(x,y_l)	M(x,y_w)
F	F	F	F
F	F	F	T
F	F	T	F
F	F	T	T
F	T	F	F
F	T	F	T
F	T	T	F
F	T	T	T
T	F	F	F
T	F	T	F
T	F	T	T
T	T	F	F
T	T	F	T
T	T	T	F
T	T	T	T

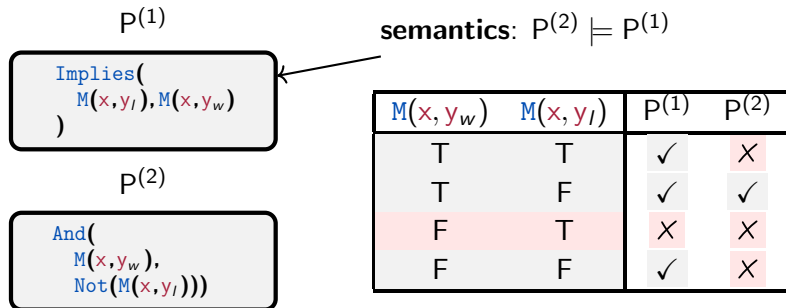
w/ reference: 4,294,967,296 losses

Why is this useful? understanding the structure



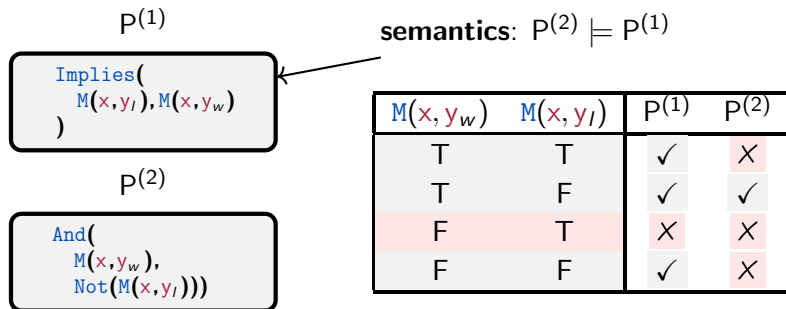
Proposition (Xu et al., 2018): Loss behavior is monotonic w.r.t semantic entailment: if $P^{(2)} \models P^{(1)}$ then $\ell(D, \theta, P^{(2)}) \geq \ell(D, \theta, P^{(1)})$.

Why is this useful? understanding the structure



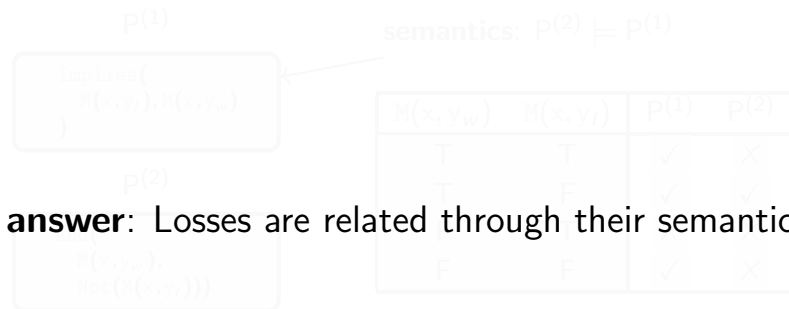
Proposition (Xu et al., 2018): Loss is equivalent under semantic equivalence: If $P^{(2)} \equiv P^{(1)}$ then $\ell(D, \theta, P^{(2)}) = \ell(D, \theta, P^{(1)})$.

Why is this useful? understanding the structure



Theorem: $\ell(D, \theta, P^{(2)}) > \ell(D, \theta, P^{(1)})$ (the loss of $P^{(1)}$ is contained in the loss of $P^{(2)}$).

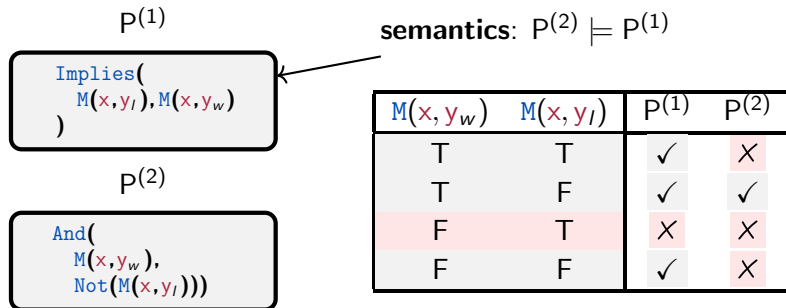
Why is this useful? understanding the structure



answer: Losses are related through their semantics

Theorem: $\ell(D, \theta, p^{(2)}) \geq \ell(D, \theta, p^{(1)})$ (the loss of $p^{(1)}$ is contained in the loss of $p^{(2)}$).

Why is this useful? understanding the structure



Practical strategy: Start with empirically successful losses, modify semantics (make more or less constrained), then experiment accordingly.

Deriving new losses symbolically, from first principles

Symbolic Program

```
Implies(  
  And(M(x, yI), Ref(x, yw)),  
  And(M(x, yw), Ref(x, yI))  
)
```

DPO Loss

$$-\log \sigma \left(\log \frac{\pi_{\theta}(y_w|x)}{\pi_{\text{ref}}(y_w|x)} - \log \frac{\pi_{\theta}(y_I|x)}{\pi_{\text{ref}}(y_I|x)} \right)$$



Deriving new losses symbolically, from first principles

Symbolic Program

```
Implies(  
  And(M(x, yI), Ref(x, yw)),  
  And(M(x, yw), Ref(x, yI))  
)
```

 modify

```
Implies(  
  And(M(x, yI), Ref(x, yw)),  
  M(x, yw)  
)
```

DPO Loss

$$-\log \sigma \left(\log \frac{\pi_{\theta}(y_w|x)}{\pi_{\text{ref}}(y_w|x)} - \log \frac{\pi_{\theta}(y_I|x)}{\pi_{\text{ref}}(y_I|x)} \right)$$

Deriving new losses symbolically, from first principles

Symbolic Program

```
Implies(  
  And(M(x, yI), Ref(x, yw)),  
  And(M(x, yw), Ref(x, yI))  
)
```

modify

```
Implies(  
  And(M(x, yI), Ref(x, yw)),  
  M(x, yw)  
)
```

DPO Loss

$$-\log \sigma \left(\log \frac{\pi_{\theta}(y_w|x)}{\pi_{\text{ref}}(y_w|x)} - \log \frac{\pi_{\theta}(y_l|x)}{\pi_{\text{ref}}(y_l|x)} \right)$$

Novel loss

$$-\log \sigma \left(\log \frac{\pi_{\theta}(y_w|x)\pi_{\text{ref}}(y_l|x)(1-\pi_{\theta}(y_l|x))}{\pi_{\theta}(y_l|x)\pi_{\text{ref}}(y_w|x)(1-\pi_{\theta}(y_w|x))} \right)$$

Deriving new losses symbolically, from first principles

Symbolic Program

```
Implies(  
  And(M(x, yI), Ref(x, yw)),  
  And(M(x, yw), Ref(x, yI))  
)
```

↓ modify

```
Implies(  
  And(M(x, yI), Ref(x, yw)),  
  M(x, yw)  
)
```

DPO Loss

$$-\log \sigma \left(\log \frac{\pi_{\theta}(y_w|x)}{\pi_{\text{ref}}(y_w|x)} - \log \frac{\pi_{\theta}(y_I|x)}{\pi_{\text{ref}}(y_I|x)} \right)$$

Novel loss

$$-\log \sigma \left(\log \frac{\pi_{\theta}(y_w|x)\pi_{\text{ref}}(y_I|x)(1-\pi_{\theta}(y_I|x))}{\pi_{\theta}(y_I|x)\pi_{\text{ref}}(y_w|x)(1-\pi_{\theta}(y_w|x))} \right)$$

- High-level programming language for defining new losses.

Deriving new losses symbolically, from first principles

Symbolic Program

```
Implies(  
  And(M(x,y),Ref(x,y_w)),  
  And(M(x,y_w),Ref(x,y))  
)
```

DPO Loss

$$-\log \sigma \left(\log \frac{\pi_{\theta}(y_w|x)}{\pi_{\text{ref}}(y_w|x)} - \log \frac{\pi_{\theta}(y_l|x)}{\pi_{\text{ref}}(y_l|x)} \right)$$

questions: How does our logic work? What do we see?

```
Implies(  
  And(M(x,y_l),Ref(x,y_w)),  
  M(x,y_w)  
)
```

$$-\log \sigma \left(\log \frac{\pi_{\theta}(y_w|x)\pi_{\text{ref}}(y_l|x)(1-\pi_{\theta}(y_l|x))}{\pi_{\theta}(y_l|x)\pi_{\text{ref}}(y_w|x)(1-\pi_{\theta}(y_w|x))} \right)$$

How does the logic work?

$M(x, y_w)$	$M(x, y_l)$	CPO	ORPO	unCPO
T	T	✓ X		✓
T	F	✓	✓	✓
F	T	X	X	X
F	F			✓

P

```
Implies(  
   $M(x, y_l)$ ,  $M(x, y_w)$   
)
```

Whenever the model deems the loser to be a valid generation, it should deem the winner to be valid too.

How does the logic work?

P

$M(x, y_w)$	$M(x, y_l)$	CPO	ORPO	unCPO
T	T	✓ X		✓
T	F	✓	✓	✓
F	T	X	X	X
F	F			✓

$\text{Implies}(M(x, y_l), M(x, y_w))$

Whenever the model deems the loser to be a valid generation, it should deem the winner to be valid too.

$$\{\checkmark, X\}_w := \prod_{w \models M(x, y)} \pi_{\theta}(y | x) \cdot \prod_{w \models \neg M(x, y)} 1 - \pi_{\theta}(y | x)$$

How does the logic work?

		P		
$M(x, y_w)$	$M(x, y_l)$	CPO	ORPO	unCPO
T	T			
T	F			
F	T			
F	F			

`Implies(`
 $M(x, y_l), M(x, y_w)$
`)`

Whenever the model deems the loser to be a valid generation, it should deem the winner to be valid too.

$$\{\checkmark, \text{X}\}_w := \prod_{w \models M(x, y)} \pi_\theta(y | x) \cdot \prod_{w \models \neg M(x, y)} 1 - \pi_\theta(y | x)$$

- Formula probability P computed as a weighted count $\sum \checkmark_w$ (Chavira and Darwiche, 2008), loss is $-\log$, *semantic loss* (Xu et al., 2018).

How does the logic work?

$M(x, y_w)$	$M(x, y_l)$	CPO	ORPO	unCPO
T	T	✓ X		✓
T	F	✓	✓	✓

P

```
Implies(  
  M(x, y_l), M(x, y_w)  
)
```



ELSEVIER

Available online at www.sciencedirect.com



ScienceDirect

Artificial Intelligence 172 (2008) 772–799

**Artificial
Intelligence**

www.elsevier.com/locate/artint

On probabilistic inference by weighted model counting [☆]

Mark Chavira ^{*}, Adnan Darwiche **algorithmic foundation**

Computer Science Department, UCLA, Los Angeles, CA 90095, USA

Received 25 August 2006; received in revised form 22 July 2007; accepted 5 November 2007

A Semantic Loss Function for Deep Learning with Symbolic Knowledge

Losses computed from weighted model counts

Jingyi Xu¹ Zilu Zhang² Tal Friedman¹ Yitao Liang¹ Guy Van den Broeck¹

Abstract

This paper develops a novel methodology for using symbolic knowledge in deep learning. From first principles, we derive a semantic loss function that bridges between neural output vectors and logical constraints. This loss function captures how close the neural network is to satisfying the constraints on its output. An experimental evaluation shows that it effectively guides the learner to achieve (near)-state-of-the-art results on semi-supervised multi-class classification. Moreover, it significantly increases the ability of the neural network to predict structured objects, such as rankings and paths. These discrete concepts are tremendously difficult to learn, and benefit from a tight integration of deep learning and symbolic reasoning methods.

This paper considers learning in domains where we have symbolic knowledge connecting the different outputs of a neural network. This knowledge takes the form of a constraint (or sentence) in Boolean logic. It can be as simple as an exactly-one constraint for one-hot output encodings, or as complex as a structured output prediction constraint for intricate combinatorial objects such as rankings, subgraphs, or paths. Our goal is to augment neural networks with the ability to learn how to make predictions subject to these constraints, and use the symbolic knowledge to improve the learning performance.

Most neuro-symbolic approaches aim to simulate or learn symbolic reasoning in an end-to-end deep neural network, or capture symbolic knowledge in a vector-space embedding. This choice is partly motivated by the need for smooth *differentiable* models; adding symbolic reasoning code (e.g., SAT solvers) to a deep learning pipeline de-

How does the logic work?

		P		
$M(x, y_w)$	$M(x, y_l)$	CPO	ORPO	unCPO
T	T	✓ X		✓
T	F	✓	✓	✓
F	T	X	X	X
F	F			✓

Implies(
 $M(x, y_l)$, $M(x, y_w)$
)

Whenever the model deems the loser to be a valid generation, it should deem the winner to be valid too.

$$\{\checkmark, X\}_w := \prod_{w \models M(x, y)} \pi_\theta(y | x) \cdot \prod_{w \models \neg M(x, y)} 1 - \pi_\theta(y | x)$$

$$\underbrace{\ell_x}_{\text{column}} := \underbrace{-\log \sigma \left(\log \frac{\sum \checkmark_w}{\sum X_w} \right)}_{\text{log ratio of } \checkmark_w \text{ and } X_w \text{ counts}}$$

How does the logic work?

$M(x, y_w)$	$M(x, y_l)$	CPO	ORPO	unCPO
T	T	✓ X		✓
T	F	✓	✓	✓
F	T	X	X	X
F	F			✓

P

$\text{Implies}(\textcolor{blue}{M}(x, y_l), \textcolor{blue}{M}(x, y_w))$

Whenever the model deems the loser to be a valid generation, it should deem the winner to be valid too.

$$\{\checkmark, X\}_w := \prod_{w \models M(x, y)} \pi_\theta(y \mid x) \cdot \prod_{w \models \neg M(x, y)} 1 - \pi_\theta(y \mid x)$$

$$\underbrace{\ell_x}_{\text{column}} := -\log \sigma \left(\log \frac{\sum \checkmark_w}{\sum X_w} \right)$$

$$= \underbrace{-\log \sigma \left(\log \frac{\pi_\theta(y_w \mid x)(1 - \pi_\theta(y_l \mid x))}{\pi_\theta(y_l \mid x)(1 - \pi_\theta(y_w \mid x))} \right)}_{\ell_{\text{ORPO}}, P_\theta(P \mid \text{one hot})}$$

How does the logic work?

		P		
$M(x, y_w)$	$M(x, y_l)$	CPO	ORPO	unCPO
T	T	✓ X		✓
T	F	✓	✓	✓
F	T	X	X	X
F	F			✓

Implies(
 $M(x, y_l)$, $M(x, y_w)$
)

Whenever the model deems the loser to be a valid generation, it should deem the winner to be valid too.

$$\{\checkmark, X\}_w := \prod_{w \models M(x, y)} \pi_\theta(y \mid x) \cdot \prod_{w \models \neg M(x, y)} 1 - \pi_\theta(y \mid x)$$

$$\begin{aligned}
 \underbrace{\ell_x}_{\text{column}} &:= -\log \sigma \left(\log \frac{\sum \checkmark_w}{\sum X_w} \right) \\
 &= \underbrace{-\log \sigma \left(\log \frac{\pi_\theta(y_w \mid x)}{\pi_\theta(y_l \mid x)} \right)}_{\ell_{\text{CPO}}, \sim P_\theta(P \mid \text{one true})}
 \end{aligned}$$

How does the logic work?

$M(x, y_w)$	$M(x, y_l)$	CPO	ORPO	unCPO	P
T	T	✓ X		✓	$\text{Implies}($ $M(x, y_l), M(x, y_w)$ $)$ Whenever the model deems the loser to be a <u>valid</u> gen- eration, it should deem the winner to be <u>valid</u> too.
T	F	✓	✓	✓	
F	T	X	X	X	
F	F			✓	

{ ✓ **observation:** losses differ in hard constraints

$$\underbrace{\ell_x}_{\text{column}} := -\log \sigma \left(\log \frac{\sum \checkmark_w}{\sum X_w} \right)$$

$$= \underbrace{-\log \sigma \left(\log \frac{\pi_\theta(y_w | x)}{\pi_\theta(y_l | x)} \right)}_{\ell_{\text{CPO}}, \sim P_\theta(P|\text{one true})}$$

How does the logic work?

		P		
$M(x, y_w)$	$M(x, y_l)$	CPO	ORPO	unCPO
T	T	✓ X		✓
T	F	✓	✓	✓
F	T	X	X	X
F	F			✓

$\text{Implies}($
 $M(x, y_l), M(x, y_w)$
 $)$

Whenever the model deems the loser to be a valid generation, it should deem the winner to be valid too.

$$\{\checkmark, X\}_w := \prod_{w \models M(x, y)} \pi_{\theta}(y \mid x) \cdot \prod_{w \models \neg M(x, y)} 1 - \pi_{\theta}(y \mid x)$$

Loss	Representation $\bar{P}$
CE	$P := M(x, y_w), P_C := \perp$
CEUnl	$P := \text{And}(M(x, y_w), \text{Not}(M(x, y_l)))$ $P_C := \perp$
CPO	$;;$ core semantic formula $P := \text{Implies}(M(x, y_l), M(x, y_w))$ $;;$ one-true constraint $P_C := \text{Or}(M(x, y_l), M(x, y_w))$
ORPO	$P := \text{Implies}(M(x, y_l), M(x, y_w))$ $;;$ one-hot constraint $P_C := \text{XOR}(M(x, y_l), M(x, y_w))$

How does the logic work?

		P		
$M(x, y_w)$	$M(x, y_l)$	CPO	ORPO	unCPO
T	T	✓ X		✓
T	F	✓	✓	✓
F	T	X	X	X
F	F			✓

$\text{Implies}(\text{M}(x, y_l), \text{M}(x, y_w))$

Whenever the model deems the loser to be a valid generation, it should deem the winner to be valid too.

$$\{\checkmark, X\}_w := \prod_{w \models M(x, y)} \pi_{\theta}(y \mid x) \cdot \prod_{w \models \neg M(x, y)} 1 - \pi_{\theta}(y \mid x)$$

- **Preference structure:** equivalent way of expressing truth table representations (Richardson et al., 2025),

$$\bar{P} := \left(\underbrace{P}_{\text{core}}, \underbrace{P_C, P_A}_{\text{constraints}} \right)$$

How does the logic work?

$M(x, y_w)$	$M(x, y_l)$	CPO	ORPO	unCPO
T	T	✓ X		✓
T	F	✓	✓	✓
F	T	X	X	X
F	F			✓

P

$\text{Implies}(\text{M}(x, y_l), \text{M}(x, y_w))$

Whenever the model deems the loser to be a valid generation, it should deem the winner to be valid too.

$$\{\checkmark, X\}_w := \prod_{w \models M(x, y)} \pi_{\theta}(y \mid x) \cdot \prod_{w \models \neg M(x, y)} 1 - \pi_{\theta}(y \mid x)$$

$$\underbrace{\ell_x}_{\text{column}} := -\log \sigma \left(\log \frac{\sum \checkmark_w}{\sum X_w} \right)$$

$$= \underbrace{-\log \sigma \left(\log \frac{\pi_{\theta}(y_l \mid x) \pi_{\theta}(y_w \mid x) + (1 - \pi_{\theta}(y_l \mid x))}{\pi_{\theta}(y_l \mid x) (1 - \pi_{\theta}(y_w \mid x))} \right)}_{\text{novel loss without constraints, } P_{\theta}(P \mid T)}$$

How does the logic work?

$M(x, y_w)$	$M(x, y_l)$	CPO	ORPO	unCPO	P
T	T	✓ X		✓	Implies($M(x, y_l), M(x, y_w)$)
T	F	✓	✓	✓	
F	T	X	X	X	
F	F			✓	

Whenever the model deems the loser to be a valid generation, it should deem the winner to be valid too.

observation: real losses are highly constrained

$$\underbrace{\ell_x}_{\text{column}} := -\log \sigma \left(\log \frac{\sum \checkmark_w}{\sum \times_w} \right)$$

$$= -\log \sigma \left(\log \frac{\pi_\theta(y_l | x) \pi_\theta(y_w | x) + (1 - \pi_\theta(y_l | x))}{\pi_\theta(y_l | x) (1 - \pi_\theta(y_w | x))} \right)$$

novel loss without constraints, $P_\theta(P|T)$

How does the logic work?

$M(x, y_w)$	$M(x, y_l)$	CPO	ORPO	unCPO
T	T	✓ X		✓
T	F	✓	✓	✓
F	T	X	X	X
F	F			✓

P
 $\text{Implies}(\text{M}(x, y_l), \text{M}(x, y_w))$

Whenever the model deems the loser to be a valid generation, it should deem the winner (valid too).

note: $\underbrace{M(x, y_l) \rightarrow M(x, y_w)}_{\text{goal of optimization}} \equiv \neg M(x, y_l) \vee M(x, y_w)$

$$\underbrace{\ell_x}_{\text{column}} := -\log \sigma \left(\log \frac{\sum \checkmark_w}{\sum X_w} \right)$$

$$= -\log \sigma \left(\log \frac{\pi_{\theta}(y_l | x) \pi_{\theta}(y_w | x) + (1 - \pi_{\theta}(y_l | x))}{\pi_{\theta}(y_l | x) (1 - \pi_{\theta}(y_w | x))} \right)$$

novel loss without constraints, $P_{\theta}(P|T)$

How does the logic work?

$M(x, y_w)$	$M(x, y_l)$	CPO	ORPO	unCPO	P
T	T	✓ X		✓	$\text{Implies}(M(x, y_l), M(x, y_w))$ Whenever the model deems the loser to be a <u>valid</u> generation, it should deem the winner (valid too).
T	F	✓	✓	✓	
F	T	X	X	X	
F	F			✓	

note: $M(x, y_l) \rightarrow M(x, y_w) \equiv \underbrace{\neg M(x, y_l) \vee M(x, y_w)}_{\text{drive probability to zero}}$

$$\underbrace{\ell_x}_{\text{column}} := -\log \sigma \left(\log \frac{\sum \checkmark_w}{\sum X_w} \right)$$

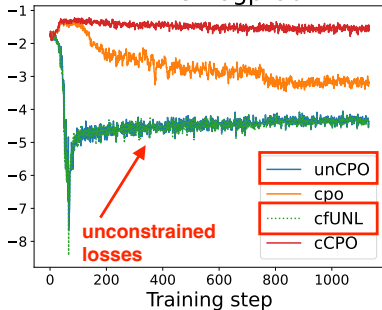
$$= -\log \sigma \left(\log \frac{\pi_\theta(y_l | x) \pi_\theta(y_w | x) + (1 - \pi_\theta(y_l | x))}{\pi_\theta(y_l | x) (1 - \pi_\theta(y_w | x))} \right)$$

novel loss without constraints, $P_\theta(P|T)$

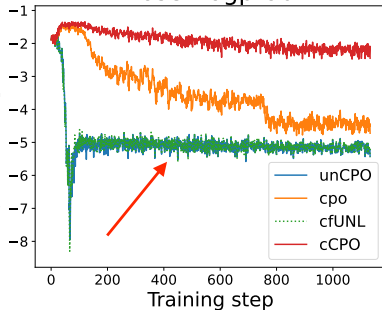
How does the logic work?

Training dynamics

Winner logprob



Loser logprob



$$= -\log \sigma \left(\underbrace{\log \frac{\pi_{\theta}(y_l | x) \pi_{\theta}(y_w | x) + (1 - \pi_{\theta}(y_l | x))}{\pi_{\theta}(y_l | x) (1 - \pi_{\theta}(y_w | x))}}_{\text{novel loss without constraints, } P_{\theta}(P|T)} \right)$$

How does the logic work?

$M(x, y_w)$	$M(x, y_l)$	CPO	ORPO	unCPO	P
T	T	✓ X		✓	$\text{Implies}(M(x, y_l), M(x, y_w))$ Whenever the model deems the loser to be a <u>valid</u> generation, it should deem the winner to be <u>valid</u> too.
T	F	✓	✓	✓	
F	T	X	X	X	
F	F			✓	

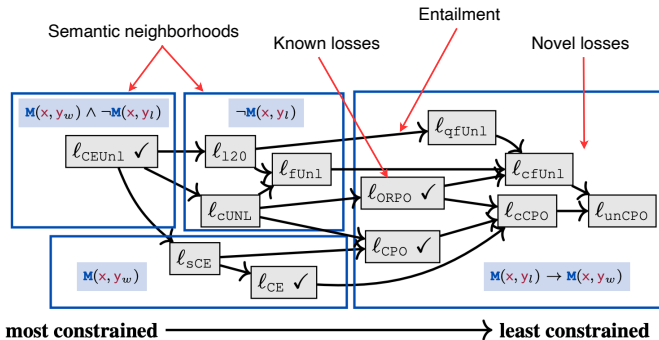
{✓, X} Constrainedness is an important property
 $w \models \ell(x, y)$ $w \models \neg \ell(x, y)$

$$\underbrace{\ell_x}_{\text{column}} := -\log \sigma \left(\log \frac{\sum \checkmark_w}{\sum X_w} \right)$$

$$= -\log \sigma \left(\log \frac{\pi_\theta(y_l | x) \pi_\theta(y_w | x) + (1 - \pi_\theta(y_l | x))}{\pi_\theta(y_l | x) (1 - \pi_\theta(y_w | x))} \right)$$

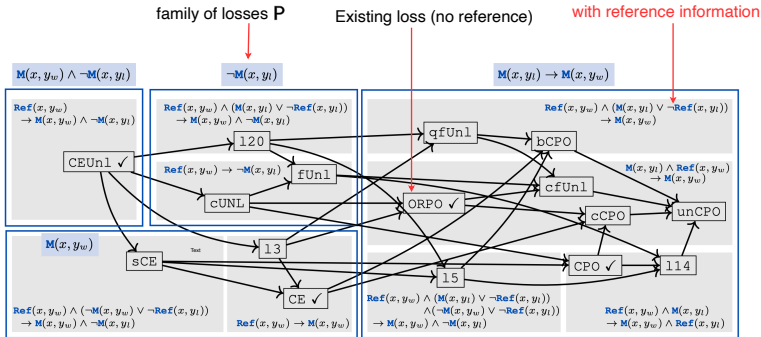
novel loss without constraints, $P_\theta(P|T)$

The **no** reference loss landscape



- **Loss lattice:** semantic structure of loss space, ordering.

The reference loss landscape



- The semantics of DPO-style reference losses can be straightforwardly computed from no reference approaches, much less explored.

question: Are any of these losses good?

Characterizing dataset semantics?

known loss (baseline) data sub-categories

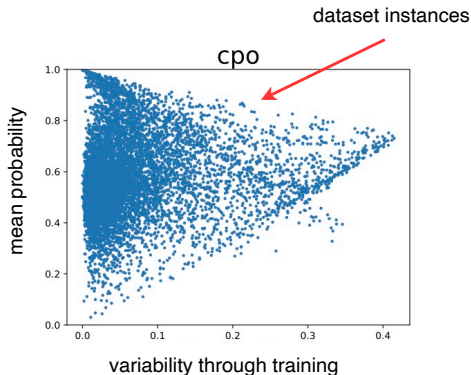
novel losses →

loss	WR% (ℓ_{CPO})	evol	false-qa	flan	sharegpt	ultrachat
ℓ_{CFUNL}	46.1 (± 0.4)	46.1 (± 2.2)	51.6 (± 2.9)	46.4 (± 1.7)	46.2 (± 1.2)	44.1 (± 1.0)
ℓ_{qFUNL}	48.9 (± 0.8)	45.3 (± 1.9)	34.7 (± 6.3)	57.9 (± 1.2)	46.8 (± 2.4)	41.3 (± 1.4)
ℓ_{CCPO}	52.0 (± 0.6)	50.7 (± 0.5)	50.2 (± 0.7)	57.2 (± 1.1)	47.2 (± 1.8)	53.1 (± 1.9)
ℓ_{unCPO}	46.0 (± 0.2)	45.8 (± 0.3)	52.1 (± 3.0)	45.7 (± 0.6)	46.2 (± 2.1)	44.8 (± 2.1)

Table 5: Comparing performance of Qwen-0.5B tuned on new losses (rows) against ℓ_{CPO} based on aggregate win-rate (WR % (std)) on ultrafeedback test (second column) and different test subsets (columns 2-6).

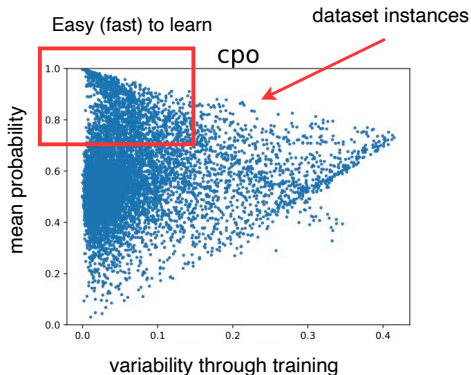
- **Finding:** Different losses perform better/worse on different subsets of data, reflecting the different semantics in preference data.

Characterizing dataset semantics?



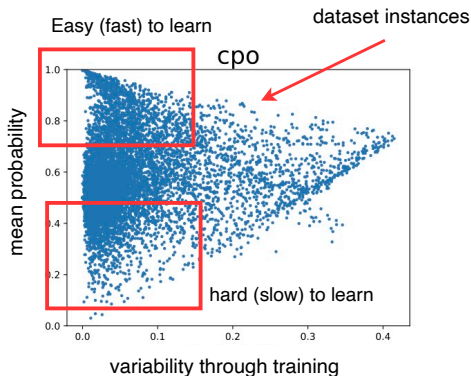
- Tracking instance-level training dynamics (Swayamdipta et al., 2020) and behavior across losses, use to reverse engineer data semantics.

Characterizing dataset semantics?



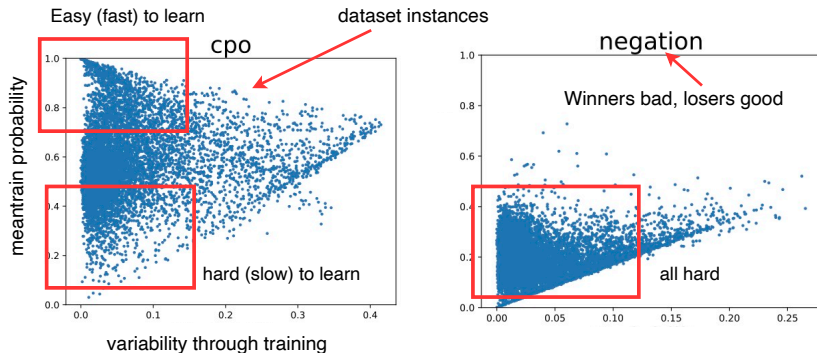
- **Intuition:** The speed/ease of training is a proxy for *goodness of semantic fit*, similar to small-loss criterion in noisy-label learning (Gui et al., 2021).

Characterizing dataset semantics?



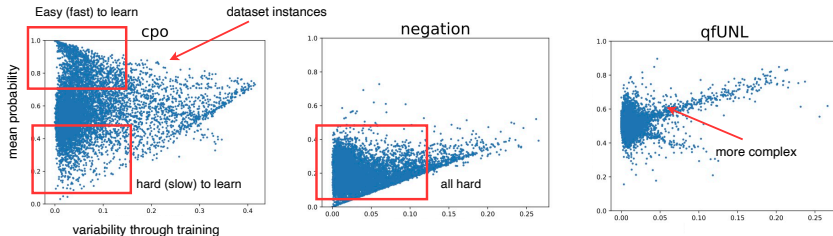
- **Intuition:** The speed/ease of training is a proxy for *goodness of semantic fit*, similar to small-loss criterion in noisy-label learning (Gui et al., 2021).

Characterizing dataset semantics?



- **Intuition:** The speed/ease of training is a proxy for *goodness of semantic fit*, similar to small-loss criterion in noisy-label learning (Gui et al., 2021).

Characterizing dataset semantics?



- **Intuition:** The speed/ease of training is a proxy for *goodness of semantic fit*, similar to small-loss criterion in noisy-label learning (Gui et al., 2021).

Characterizing dataset semantics?

- Intuition: The speed/ease of training is a proxy for *goodness of semantic fit* (relative to small loss criterion) for a given model training (Gardner, 2009)

Blueprint for much empirical exploration of loss space

Conclusions

- ▶ New ideas about using symbolic techniques to formally characterize the semantics of LLM algorithms, preference learning.

Conclusions

- ▶ New ideas about using symbolic techniques to formally characterize the semantics of LLM algorithms, preference learning.
 1. Understanding the full space of loss functions (finding: it's a huge space, many novel variations yet to be explored)
 2. Understanding the structure of the space and relationships between different losses (finding: tied to the semantics of the losses).

Conclusions

- ▶ New ideas about using symbolic techniques to formally characterize the semantics of LLM algorithms, preference learning.
 1. Understanding the full space of loss functions (finding: it's a huge space, many novel variations yet to be explored)
 2. Understanding the structure of the space and relationships between different losses (finding: tied to the semantics of the losses).

High-level programming: write a (high-level) symbolic program, or modify an existing one, compile into a loss and experiment (then repeat)

Conclusions

- ▶ New ideas about using symbolic techniques to formally characterize the semantics of LLM algorithms, preference learning.
 1. Understanding the full space of loss functions (finding: it's a huge space, many novel variations yet to be explored)
 2. Understanding the structure of the space and relationships between different losses (finding: tied to the semantics of the losses).

High-level programming: write a (high-level) symbolic program, or modify an existing one, compile into a loss and experiment (then repeat)

- ▶ **Decompiling models to symbolic programs:** *semantics of data, reinforcement learning, chain-of-thought, LLM agents ...*

Thank you.

Adding a reference model

```
P:= Implies(  
  And(M(x,yl),Ref(x,yw)),  
  And(M(x,yw),Ref(x,yl))  
)
```

Whenever the model being tuned deems the loser to be a valid generation and the reference model deems the winner to be valid, the tuned model should deem the winner to be valid too, and **the reference should deem the loser to be valid.**

Adding a reference model

```
P:= Implies(  
  And(M(x, yl), Ref(x, yw)),  
  And(M(x, yw), Ref(x, yl))  
)
```

Whenever the model being tuned deems the loser to be a valid generation and the reference model deems the winner to be valid, the tuned model should deem the winner to be valid too, and **the reference should deem the loser to be valid.**

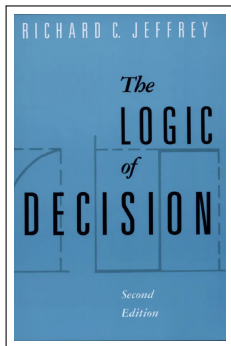
- **Peculiar semantics**, but the logic makes sense, e.g., we want to maximize

$$\sigma \left(\log \frac{\pi_{\theta}(y_w | x)}{\pi_{\theta}(y_l | x)} - \log \frac{\pi_{\text{ref}}(y_w | x)}{\pi_{\text{ref}}(y_l | x)} \right)$$

negating left side of implication (i.e., making **M(x, y_l)** and **Ref(x, y_w)** false) and making the right side true is logical.

Classical work on preference

Analytic philosophy: Much work on the semantics of pairwise preference, rich languages for expressing ideas.



(Jeffrey, 1965)

THE STATUS OF VARIOUS PREFERENCE PRINCIPLES						
Preference Principle	Von Wright	Chisholm Sosa	Martin	$P^\#$	$P^\star$	P^ω
1. $pPq \rightarrow \sim(qPp)$	✓	✓	✓	+	+	+
2. $(pPq \ \& \ qPr) \rightarrow pPr$	✓	✓	✓	+	+	+
3. $pPq \rightarrow \sim qP\sim p$		x	✓	(+) ¹	+	+
4. $\sim qP\sim p \rightarrow pPq$		x	✓	(+) ¹	+	+
5. $pPq \rightarrow (p \ \& \ \sim q) P(\sim p \ \& \ q)$	✓	x		+	+	+
6. $(p \ \& \ \sim q) P(\sim p \ \& \ q) \rightarrow pPq$	✓	x		+	+	+
7. $[\sim(pP\sim p) \ \& \ \sim(\sim pPp)] \ \& \ \sim(qP\sim q) \ \& \ \sim(\sim qPq) \rightarrow [\sim(pPq) \ \& \ \sim(qPp)]$	✓	✓		+	+	+
8. $[\sim(qP\sim q) \ \& \ \sim(\sim qPq) \ \& \ pPq] \rightarrow pP\sim p$		✓		+	+	—
9. $[\sim(qP\sim q) \ \& \ \sim(\sim qPq) \ \& \ qP\sim p] \rightarrow pP\sim p$		✓		+	+	—
10. $pPq \rightarrow [(p \ \& \ r) P(q \ \& \ r) \ \& \ (p \ \& \ \sim r) P(q \ \& \ \sim r)]$	✓			—	—	+
11. $[(p \ \& \ r) P(q \ \& \ r) \ \& \ (p \ \& \ \sim r) P(q \ \& \ \sim r)] \rightarrow pPq$	✓			(+) ²	(+) ³	+
12. $[\sim(pPq) \ \& \ \sim(qPr)] \rightarrow \sim(pPr)$		✓		—	—	—
13. $(pPr \vee qPr) \rightarrow (p \vee q) Pr$			✓	—	—	—
14. $(p \vee q) Pr \rightarrow [pPr \ \& \ qPr]$	✓			—	—	—
15. $[pPr \ \& \ qPr] \rightarrow (p \vee q) Pr$	✓			—	—	—
16. $(p \vee q) Pr \rightarrow (pPr \vee qPr)$				—	—	—
17. $pP(q \vee r) \rightarrow (pPq \ \vee \ pPr)$			✓	—	—	—
18. $(pPq \ \& \ pPr) \rightarrow pP(q \vee r)$				—	—	—
19. $(pPr \ \& \ qPr) \rightarrow (p \ \& \ q) Pr$				—	—	—

Semantic foundations for the logic of preference Rescher (1967)

References I

- Beurer-Kellner, L., Fischer, M., and Vechev, M. (2023). Prompting is programming: A query language for large language models. *Proceedings of the ACM on Programming Languages*, 7(PLDI):1946–1969.
- Bogin, B., Yang, K., Gupta, S., Richardson, K., Bransom, E., Clark, P., Sabharwal, A., and Khot, T. (2024). Super: Evaluating agents on setting up and executing tasks from research repositories. *Proceedings of EMNLP*.
- Bragg, J., D’Arcy, M., Balepur, N., Bareket, D., Dalvi, B., Feldman, S., Haddad, D., Hwang, J. D., Jansen, P., Kishore, V., Majumder, B. P., Naik, A., Rahamimov, S., Richardson, K., Singh, A., Surana, H., Tiktinsky, A., Vasu, R., Wiener, G., Anastasiades, C., Candra, S., Dunkelberger, J., Emery, D., Evans, R., Hamada, M., Huff, R., Kinney, R., Latzke, M., Lochner, J., Lozano-Aguilera, R., Nguyen, C., Rao, S., Tanaka, A., Vlahos, B., Clark, P., Downey, D., Goldberg, Y., Sabharwal, A., and Weld, D. S. (2025). Astabench: Rigorous benchmarking of ai agents with a scientific research suite.
- Chavira, M. and Darwiche, A. (2008). On probabilistic inference by weighted model counting. *Artificial Intelligence*, 172(6-7):772–799.
- Chen, J., Yuan, S., Ye, R., Majumder, B. P., and Richardson, K. (2023). Put your money where your mouth is: Evaluating strategic planning and execution of llm agents in an auction arena. *arXiv preprint arXiv:2310.05746*.
- Cheng, J., Clark, P., and Richardson, K. (2025). Language modeling by language models. *Proceedings of Neurips*.

References II

- Dai, J., Pan, X., Sun, R., Ji, J., Xu, X., Liu, M., Wang, Y., and Yang, Y. (2024). Safe rlhf: Safe reinforcement learning from human feedback. In *The Twelfth International Conference on Learning Representations*.
- Friedman, D., Wettig, A., and Chen, D. (2023). Learning transformer programs. *Advances in Neural Information Processing Systems*, 36:49044–49067.
- Gui, X.-J., Wang, W., and Tian, Z.-H. (2021). Towards understanding deep learning from noisy labels with small-loss criterion. *arXiv preprint arXiv:2106.09291*.
- Jeffrey, R. C. (1965). *The logic of decision*. University of Chicago press.
- Ji, J., Liu, M., Dai, J., Pan, X., Zhang, C., Bian, C., Chen, B., Sun, R., Wang, Y., and Yang, Y. (2024). Beavertails: Towards improved safety alignment of llm via a human-preference dataset. *Advances in Neural Information Processing Systems*, 36.
- Li, Z., Huang, J., and Naik, M. (2023). Scallop: A language for neurosymbolic programming. *Proceedings of the ACM on Programming Languages*, 7(PLDI):1463–1487.
- Manhaeve, R., Dumancic, S., Kimmig, A., Demeester, T., and De Raedt, L. (2018). Deepproblog: Neural probabilistic logic programming. *Advances in neural information processing systems*, 31.
- Marra, G., Dumančić, S., Manhaeve, R., and De Raedt, L. (2024). From statistical relational to neurosymbolic artificial intelligence: A survey. *Artificial Intelligence*, page 104062.

References III

- Meng, Y., Xia, M., and Chen, D. (2024). Simpo: Simple preference optimization with a reference-free reward. *arXiv preprint arXiv:2405.14734*.
- Merrill, W. and Sabharwal, A. (2023). A logic for expressing log-precision transformers. *Advances in neural information processing systems*, 36:52453–52463.
- Rescher, N. (1967). *The logic of decision and action*. University of Pittsburgh Pre.
- Richardson, K., Srikumar, V., and Sabharwal, A. (2025). Understanding the logic of direct preference alignment through logic. *Proceedings of ICML*.
- Richardson, K. and Wijnholds, G. (2025). Lectures on language model programming.
- Swayamdipta, S., Schwartz, R., Lourie, N., Wang, Y., Hajishirzi, H., Smith, N. A., and Choi, Y. (2020). Dataset cartography: Mapping and diagnosing datasets with training dynamics. *arXiv preprint arXiv:2009.10795*.
- Weiss, G., Goldberg, Y., and Yahav, E. (2021). Thinking like transformers. In *International Conference on Machine Learning*, pages 11080–11090. PMLR.
- Xu, J., Zhang, Z., Friedman, T., Liang, Y., and Broeck, G. (2018). A Semantic Loss Function for Deep Learning with Symbolic Knowledge. In *International Conference on Machine Learning*, pages 5498–5507.
- Yang, A. and Chiang, D. (2024). Counting like transformers: Compiling temporal counting logic into softmax transformers. *arXiv preprint arXiv:2404.04393*.
- Yang, A., Chiang, D., and Angluin, D. (2024). Masked hard-attention transformers recognize exactly the star-free languages. *Advances in Neural Information Processing Systems*, 37:10202–10235.

References IV

- Yang, R., Chen, J., Zhang, Y., Yuan, S., Chen, A., Richardson, K., Xiao, Y., and Yang, D. (2025). Selfgoal: Your language agents already know how to achieve high-level goals. *Proceedings of NAACL*.
- Yu, H., Xuan, K., Li, F., Zhu, K., Lei, Z., Zhang, J., Qi, Z., Richardson, K., and You, J. (2025). Tinscientist: An interactive, extensible, and controllable framework for building research agents. *Proceedings of EMNLP*.
- Zhang, Y., Yuan, S., Hu, C., Richardson, K., Xiao, Y., and Chen, J. (2024). Timearena: Shaping efficient multitasking language agents in a time-aware simulation. *Proceedings of ACL*.