
PoE-World: Compositional World Modeling with Products of Programmatic Experts

Wasu Top Piriyaikulij¹ Yichao Liang² Hao Tang¹
Adrian Weller^{2,3} Marta Kryven⁴ Kevin Ellis¹
Cornell University¹ University of Cambridge²
The Alan Turing Institute³ Dalhousie University⁴

Abstract

Learning how the world works is central to building AI agents that can adapt to complex environments. Traditional world models based on deep learning demand vast amounts of training data, and do not flexibly update their knowledge from sparse observations. Recent advances in program synthesis using Large Language Models (LLMs) give an alternate approach which learns world models represented as source code, supporting strong generalization from little data. To date, application of program-structured world models remains limited to natural language and grid-world domains. We introduce a novel program synthesis method for effectively modeling complex, non-gridworld domains by representing a world model as an exponentially-weighted product of programmatic experts (PoE-World) synthesized by LLMs. We show that this approach can learn complex, stochastic world models from just a few observations. We evaluate the learned world models by embedding them in a model-based planning agent, demonstrating efficient performance and generalization to unseen levels on Atari’s Pong and Montezuma’s Revenge. We release our code and display the learned world models and videos of the agent’s gameplay at <https://topwasu.github.io/poe-world>.

1 Introduction

How should an intelligent agent represent the dynamics of the natural world? We want a representation that is efficiently learnable, yet flexible enough to handle stochasticity and partial observability, and which supports planning and decision-making. Neural network world models such as Dreamer [1] are radically flexible, but demand enormous training data (compared to humans [2]). Symbolic world models such as WorldCoder [3] instead generate a Python program to represent how the world works. These programmatic world models are data-efficient, because program synthesis requires less data than neural network training— but struggle to scale beyond simple gridworlds, as they do a discrete combinatorial search to find a single large program describing everything about how the world works.

We take inspiration from a longstanding view in philosophy and cognitive science of the mind as a community of interacting experts [4, 5, 6, 7]. This modular organization is evident across multiple scales in natural intelligence, from functional specialization of brain systems [8] to the distinct learning trajectories of specific skills [9, 10]. We integrate this modular perspective with the computational paradigm of learning as Program Synthesis, which models learned concepts as symbolic programs [11, 12, 13, 14, 15, 16]. We extend this line of work by proposing a new computational account of learning world models: as the acquisition of context-specific expert programs, which are refined through practice and reused compositionally to support flexible, goal-directed behavior.

Algorithmically, our key idea is to *decompose the problem of learning a world program into learning hundreds of small programs*. Each of these learned programs encodes a different causal law, which we probabilistically aggregate to predict future observations (Figure 1a). This makes our world

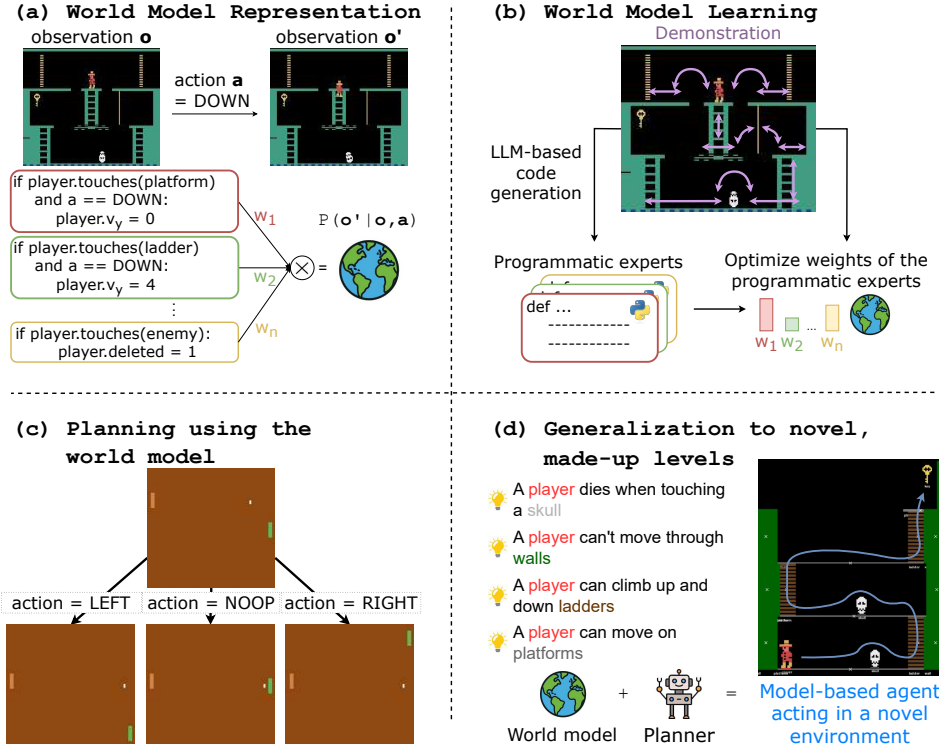


Figure 1: (a) World models predict the next state given a state-action history. We do this with a product of experts of many small programs. (b) The learner is given a short (<1 min) demonstration of gameplay as input, and uses it to synthesize initial world-model programs. These programs are refined online in later environment interactions. (c) World models support planning by imagining future states. (d) Symbolic programs encode abstract knowledge that generalizes to new game levels.

knowledge more modular, and also more learnable, because we no longer search for a single monolithic program handling everything at once. The resulting system, which we call PoE-World (**P**roduct of **P**rogrammatic **E**xperts), can build elaborate world models that accurately support planning and reinforcement learning (RL) from even a brief demonstration in complex Atari games, such as Montezuma’s Revenge. PoE-World handles stochasticity because the product of programs is probabilistic, and, as we show, further handles partial observability. To the best of our knowledge, this is the first time a symbolic world model has been learned for environments of this complexity.

Importantly, although PoE-World models fine-grained pixel-level movement, it does *not* model pixel-level visual appearance, instead assuming symbolic observations from an object detector. Unlike model-based reinforcement learning, PoE-World does not attempt efficient exploration, but focuses on faithfully learning from a demonstrated trajectory (Figure 1b). Lastly, while the world models are ultimately used for planning (Figure 1c), PoE-World fundamentally focuses on world modeling, and not on solving the challenging computational problem of planning itself.

Despite these limitations, we view PoE-World as addressing a central learning problem: Given limited demonstrations of a new environment, quickly assemble a working world model that compositionally generalizes to new situations (Figure 1d). We highlight the following contributions:

1. The PoE-World representation and learning algorithm for symbolic world models.
2. An empirical study of PoE-World on two representative Atari games, Pong and Montezuma’s Revenge, demonstrating its superior learning efficiency compared to deep RL, and improved scalability compared to state-of-the-art symbolic model-based RL: PoE-World can synthesize 4000+ line programs that generalize zero-shot to novel game levels and game variations.
3. Demonstration of how to use PoE-World’s world models for planning-based decision making, and as a simulated pretraining environment for deep RL.

2 Background: World Model Learning

A sequential decision-making problem can be described as $(\mathcal{O}, \mathcal{A}, P, R)$ where $\mathcal{O}$ is an observation space, $\mathcal{A}$ is an action space, P is an environment dynamics $P = p_{env}(o_{t+1}|o_{1:t}, a_{1:t})$ where $o \in \mathcal{O}$ and $a \in \mathcal{A}$, and R is a reward function. This full-history environment formulation is mathematically equivalent to a Partially Observable Markov Decision Process (POMDP) formulation [17]. In the setting where P is unknown to an agent, the agent learns by interacting with the environment and observing transitions (o_t, a_t, o_{t+1}, r_t) at timestep t .

We focus on learning the world model $\hat{P} = p_{model}(o_{t+1}|o_{1:t}, a_{1:t})$ which approximates the true, unknown environment dynamics $P = p_{env}(o_{t+1}|o_{1:t}, a_{1:t})$ given observed trajectories $D = \{\tau_i\}_{i=1}^n$ where each trajectory is a sequence of observations and actions $\tau = (o_{1:T+1}, a_{1:T})$ for some T ($r_{1:T}$ is also a part of a trajectory, but we drop it for simplicity). The learned world model will later be used by a model-based agent, either through lookahead planning or RL training, to act in the environment.

We treat the world modeling problem as an optimization problem, specifically empirical risk minimization:

$$p_{model}^* = \arg \min_{p_{model}} \sum_{(o_{1:T+1}, a_{1:T}) \in D} \sum_{t=1}^T \ell(p_{model}; o_{1:t+1}, a_{1:t}) \quad (1)$$

where ℓ is a loss function, such as the negative log likelihood function $\ell(p_{model}; o_{1:t+1}, a_{1:t}) = -\log p_{model}(o_{t+1}|o_{1:t}, a_{1:t})$.

Previous works in the model-based reinforcement learning literature [18, 19, 1, 20, 21, 22, 23] have used various deep neural network architectures, including convolutional and recurrent neural networks, transformers, and diffusion models, to define p_{model} as a parametric model $p_{model} = p_{\theta}$. Then, θ is optimized via gradient descent. A weakness of such approaches is poor sample efficiency and generalization. For example, Diamond [22] has failure modes such as imagining that the player can walk through a wall or teleport, even after training on almost 100 hours of observations.

Other works leverage Large Language Models (LLMs) to synthesize a code world model [3, 24, 25], which can be seen as searching for programmatic p_{model} . These LLM-based code generation algorithms are more sample efficient and extrapolate more systematically than deep learning approaches. However, they have yet to succeed beyond simple text-based and gridworld games.

3 Modeling the World as a Product of Programmatic Experts (PoE-World)

To address the limitations of existing world model learning methods, we propose representing world models as exponentially-weighted **Products of programmatic Experts (PoE-World)**, enabling sample-efficient and scalable learning of probabilistic world models that leverages LLM code generation. Figure 1 visualizes both the representation and learning algorithm discussed in this section.

3.1 World Model Representation: Product of Programmatic Experts (PoE-World)

World models represented as exponentially weighted **Products of programmatic Experts (PoE-World)** can be described mathematically as follows:

$$p_{\theta}(o_{t+1}|o_{1:t}, a_{1:t}) \propto \prod_i p_i^{expert}(o_{t+1}|o_{1:t}, a_{1:t})^{\theta_i} \quad (2)$$

where p_i^{expert} are programs and θ_i are their associated scalar weights.

This representation enables modularity and compositionality. It allows composing many small programs into a full world model to capture complicated environmental dynamics (see Figure 2). We can think of each program as an expert which expresses opinions about particular aspects of the world. For example, in the context of modeling a video game environment, one expert might encode “if a player touches a skull, then the player dies,” while another encodes “if an action is LEFT when the player is on a platform, the player’s x-axis velocity is -2.” The former expert does not express opinions on the player’s movement, while the latter does not mention conditions for player’s death.

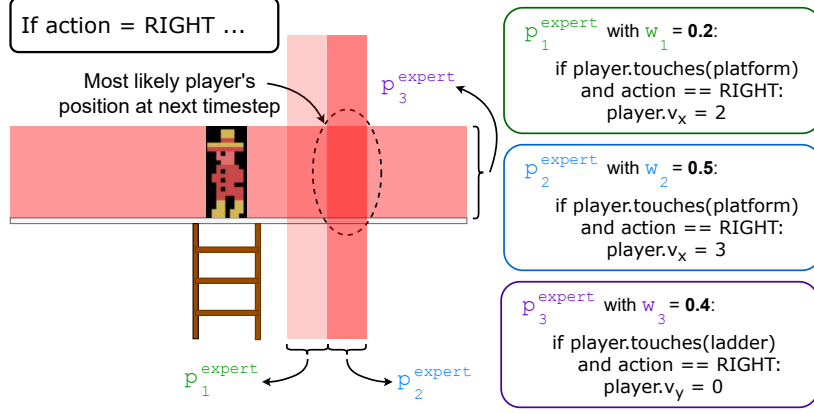


Figure 2: A “heat map” illustration of how simple Python programs are interpreted as distributions, and how they are combined into a single distribution over next-timestep object locations.

A factored state representation yields tractable inference. We assume an object-centric state where each object has a bounding box and velocities, each stored as a separate slot, attribute, or feature. We treat each such feature as conditionally independent, given the history so far. This independence assumption makes Equation (2) tractable, because we can compute a separate normalizing constant for each feature. Concretely, let f index the different object features of the next observation o_{t+1} , and let o^f be the indexing of o with feature f . Then,

$$p_{\theta}(o_{t+1}|o_{1:t}, a_{1:t}) = \prod_f \frac{1}{Z_f} \prod_i p_i^{expert}(o_{t+1}^f|o_{1:t}, a_{1:t})^{\theta_i} \quad Z_f = \sum_{o^f} \prod_i p_i^{expert}(o^f|o_{1:t}, a_{1:t})^{\theta_i}$$

Benefits of the full-history formulation over POMDP formulation. The full-history environment formulation is formally equivalent to a POMDP which instead compresses the history into a Markov latent state. We use the history formulation because it makes our world models more modular. Learning global latent variables would entangle all the experts, because every expert would condition on the latent state. Therefore, learning a new latent variable (e.g. “how long the player has been falling”) changes the input/output space of every expert, necessitating global joint updates to the structure of every program (e.g. an expert for “is the player dead” would need to be updated). The history formulation allows independent learning of independent mechanisms.

Hard constraints. Atari (and the real world) is too complex to perfectly simulate with any effectively learnable program. Therefore, our probabilistic model tends to over-approximate the set of possible futures, giving fuzzy approximate predictions. For example, in Montezuma’s Revenge, instead of perfectly modeling the physics of falling downward and landing on the ground, we predict a generic downward trajectory. Ideally that trajectory would perfectly enforce the constraint that the player lands flat on the ground, and never sinks into the ground, but a fuzzy stochastic expert for falling downward could violate that constraint. Therefore, to sharpen the model’s outputs, we further learn a collection of hard constraints, $\{c_j\}$, where $c_j : \mathcal{O} \rightarrow \{0, 1\}$:

$$p_{\theta}(o_{t+1}|o_{1:t}, a_{1:t}) \propto \prod_i p_i^{expert}(o_{t+1}|o_{1:t}, a_{1:t})^{\theta_i} \cdot \mathbb{I} \left[\bigvee_j c_j(o_{t+1}) \right] \quad (3)$$

We further discuss hard constraints and provide concrete examples in Appendix A.1.

Multi-timestep predictions. We represent the programmatic experts and the world as distributions over the next-timestep observations in eq. (2). It is likewise possible to reformulate a multi-timestep expert $p^{expert}(o_{t+1:t+H}|o_{1:t}, a_{1:t})$ as a product of next-timestep experts by assuming that the predictions of the multi-step expert at different timestep are independent:

$$p^{expert}(o_{t+1:t+H}|o_{1:t}, a_{1:t}) = \prod_{k=1}^H p^{expert}(o_{t+k}|o_{1:t}, a_{1:t}) \quad (4)$$

3.2 World Model Learning: Program Synthesis and Weight Optimization

We begin with a demonstration trajectory, learn a world model, and then begin to act in the world according to that model. As the agent acts, it collects more trajectory data, which it uses to update or “debug” its model. Concretely, learning proceeds as follows:

- Step 1: Synthesize programmatic experts $\{p_i^{expert}\}_{i=1}^m$ given observed trajectories $D = \{\tau_i\}_{i=1}^n$
- Step 2: Fit the weights θ of the experts according to eq. (5) with a gradient-based optimizer
- Step 3: Remove the experts with weights below threshold δ
- Step 4: Repeat Step 1-3 every time the observed trajectories get updated

Generating program experts. Following previous works [3, 24, 25], we adopt Large Language Models (LLMs) as our Python program generator. We input a small batch of transitions $(o_{t:t+H+1}, a_{t:t+H})$ to the LLM prompt to produce the programmatic experts $\{p_i^{expert}\}$.

Figure 2 shows how we interpret small Python programs as distributions over the observations. While we could ask LLMs to synthesize probabilistic programs to specify the distributions, we find it much more effective to have LLMs synthesize simple, deterministic Python programs, presumably because generic Python code is far more prevalent in LLM training data. We assume that an observation is represented as a list of objects, where each object has the following attributes: x/y velocity and visibility. Then, a distribution over observations is a distribution over each object’s attributes. To interpret a Python program as a distribution, we assume that all object attributes are conditionally independent given full history, as mentioned in Section 3.1, and convert all object attributes set by the program to distributions with single peaks at the given values. We then add noise to the distributions to ensure non-zero probabilities over alternative values. Any object attributes whose values are not set by the program follow uniform distributions over all possible values: consequently, experts with a single if-condition (fig. 2) yield a uniform distribution when the if-condition is not satisfied.

Gradient-based Weights Optimization. Once we have $\{p_i^{expert}\}_{i=1}^m$, we can perform maximum likelihood estimation to obtain θ :

$$\theta^* = \arg \max_{\theta} \sum_{(o_{1:T+1}, a_{1:T}) \in D} \sum_{t=1}^T \log p_{\theta}(o_{t+1} | o_{1:t}, a_{1:t}) \quad (5)$$

This equation instantiates eq. (1) by letting p_{model} have a parametric form $p_{model} = p_{\theta}$ and choosing negative log likelihood as the loss function. We can use any gradient-based optimizer to optimize the weights. We use L-BFGS [26], which worked better than Adam [27] and SGD, because we have small data and few parameters.

Finally, the experts whose weights are below a threshold δ are removed from the world model. We repeat this loop every time there are new observations. Appendix A.1 contains full algorithm details.

3.3 World Model Usage: RL in Simulation and Planning with World Model

An important goal of world modeling is to aid decision-making. We consider two such ways of using world models. First, a world model can serve a simulator for reinforcement learning. This makes RL policy learning more sample efficient, because we can quickly learn a world model from real environment interactions, which then substitutes the actual environment. Subsequent policy learning can take place in the world model, obviating the need for further interaction with the real environment. In practice, we also continue reinforcement learning after pretraining in the world model. Formally, we learn a policy $\pi : \mathcal{O}^* \rightarrow \mathcal{A}$ that inputs an observation history and outputs an action.

Alternatively, world models can be used for lookahead planning. Given a reward function R , we plan for a horizon of H timesteps by searching for an optimal action sequence, given our previous observations $o_{1:t}$ and actions $a_{1:t-1}$:

$$a_{t:t+H}^* = \arg \max_{a_{t:t+H}} \mathbb{E}_{p_{\theta}(o_{t+1:t+H} | o_{1:t}, a_{1:t+H})} \left[\sum_{k=0}^{H-1} R(o_{t+k+1}; o_{1:t+k}, a_{1:t+k}) \right]$$

where $p_{\theta}(o_{t+1:t+H} | o_{1:t}, a_{1:t+H}) = \prod_{k=0}^{H-1} p_{\theta}(o_{t+k+1} | o_{1:t+k}, a_{1:t+k})$.

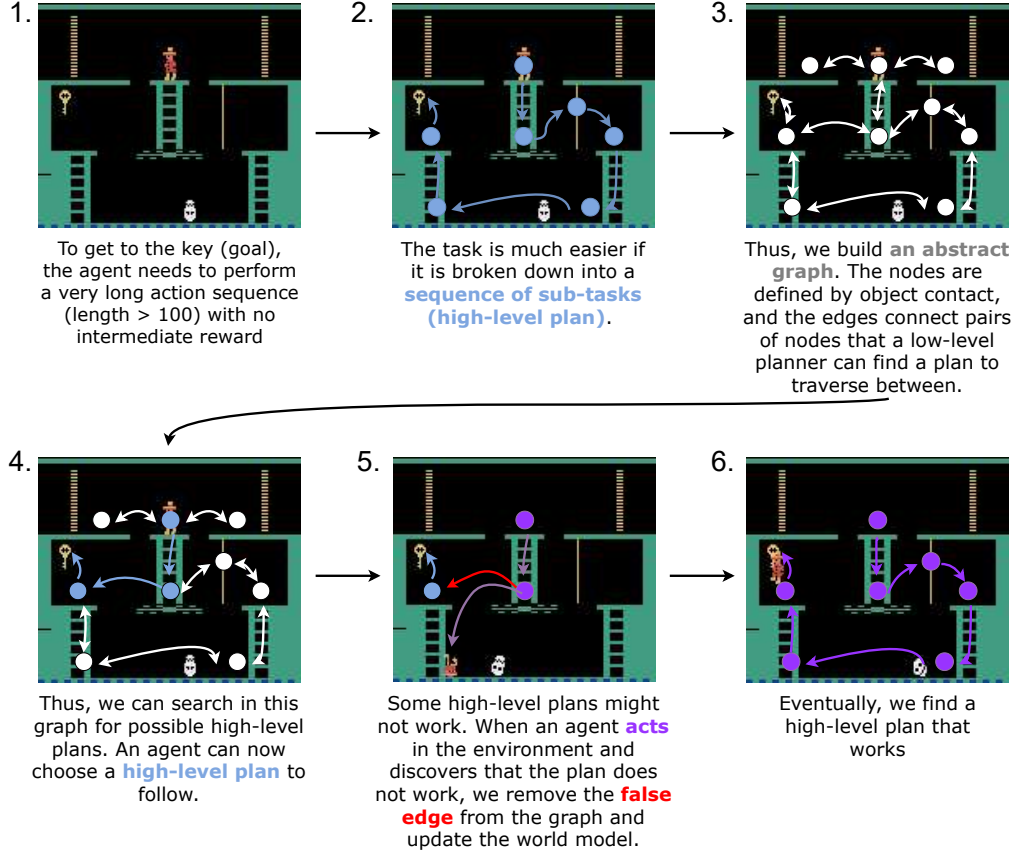


Figure 3: A sequence of illustrations that demonstrates how our hierarchical planner works.

To help our world models guide long-horizon decision-making, we implement a hierarchical planner inspired by task and motion planning (TAMP). It first plans in a high-level abstract state space defined by object contact, then lowers those plans into actual Atari button presses using the learned world model (Figure 3, Appendix A.2). The resulting agent we denote **PoE-World + Planner**.

4 Experimental Results

Domains and Evaluation. We evaluate our agent, **PoE-World + Planner**, against other methods on Atari’s Pong and Montezuma’s Revenge (MR), using the Arcade Learning Environment [28]. We use OCArari [29] to parse each image frame as a list of objects, each with an object category, a bounding box, and velocities.¹ Both games are partially observed: the current state and action cannot uniquely determine the next state. A demonstration of fewer than 1000 frames is created for each game, but these demonstrations are not successful gameplays: they serve only to illustrate the core causal mechanics. In Montezuma’s Revenge, our demonstration never achieves positive reward.

To test compositional extrapolation, we created alternative versions of both games, called Pong-Alt and Montezuma’s Revenge-Alt (Figure 4), both of which recombine and rearrange the types of objects seen in the training demonstration. Pong-Alt increases the number of objects (3 balls and 3 enemies). Montezuma’s Revenge-Alt adds more enemies (which the player has to jump over) and ladders, while changing the map to resemble the game Kangaroo (see Appendix A.3 for details). We do not provide demonstrations for these alternative versions of the games.

¹OCArari does not reliably extract objects so we have to make custom changes to OCArari for every single Atari game we apply our method to, which makes it infeasible to run on the full Atari suite. We choose Pong and Montezuma’s Revenge (arguably one of the hardest games in Atari—[30] calls it a “grand challenge”) as representative games for which we manually patch OCArari to reliably detect objects.

Method	Score			
	Pong	Pong-Alt	MR	MR-Alt
Random Agent	-20.67 ± 0.33	-20.00 ± 0.58	0.00 ± 0.00	0.00 ± 0.00
PPO @ 100k env steps [31]	-21.00 ± 0.00	-18.66 ± 0.88	0.00 ± 0.00	0.00 ± 0.00
LLM as Agent (ReAct) [32]	-20.00 ± 0.00	-20.67 ± 0.33	0.00 ± 0.00	0.00 ± 0.00
WorldCoder + Planner [3]	-17.00 ± 3.00	-19.00 ± 1.00	0.00 ± 0.00	0.00 ± 0.00
PoE-World + Planner (Ours)	-12.33 ± 0.88	-13.67 ± 0.67	100 ± 0.00	100 ± 0.00
PPO @ 20m env steps [31]	17.00 ± 0.58	1.33 ± 2.03	0.00 ± 0.00	0.00 ± 0.00

Table 1: Scores on Pong and Montezuma’s Revenge (MR) and their alternate versions. For PoE-World and WorldCoder, brief demonstrations on Pong and MR are given to initialize the world models. Their agents then train for at most 3k steps before evaluation.

Baselines. **PPO** [31] is a go-to, widely-used model-free RL algorithm. It optimizes a lower bound of the policy’s performance using gradient descent. **LLM as Agent (ReAct)** [32] directly uses an LLM as a policy. ReAct prompts LLMs to use extra chain-of-thought [33] “thinking” actions before selecting an action. **WorldCoder** [3] is an LLM agent that models the world as a single Python program. It uses an LLM code generation and repair algorithm called REx [34] to refine its world model to achieve high predictive accuracy on observed trajectories. More details in Appendix A.4.

Agent Results. Figure 5 and Table 1 show the scores of different agents on Pong, Pong-Alt, Montezuma’s Revenge, and Montezuma’s Revenge-Alt. The scores on Pong and Pong-Alt indicates the difference in points achieved by the player and the enemies when the game ends at 21 points. The scores on Montezuma’s Revenge and Montezuma’s Revenge-Alt become positive if and only if the agent succeeds in collecting the key. As shown in Table 1, our agent, **PoE-World + Planner**, performs best across all environments in the low-data regime, particularly when PPO baseline is allowed to have 100,000 training environment interactions, the standard budget for sample-efficient agents on Atari 2600 [20]. In Figure 5, we keep training PPO for more steps, finding that it takes over a million steps for PPO to surpass PoE-World + Planner. Moreover, PPO with 20M training steps never achieves positive score on Montezuma’s Revenge. PoE-World + Planner is the only method that manages to obtain positive reward on Montezuma’s Revenge in both base and alternative versions.

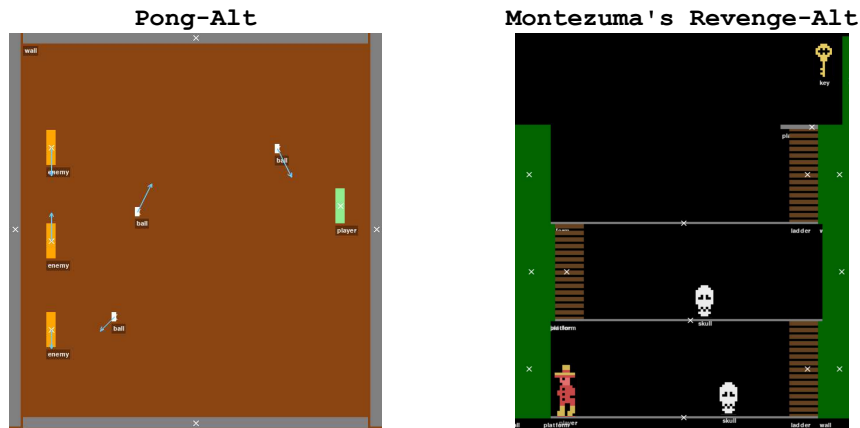


Figure 4: Screenshots of the alternative environments. Pong-Alt’s objective is to hit three balls past three enemies to score points, while Montezuma’s Revenge-Alt requires players to avoid moving skulls by jumping over them and climbing up ladders to reach and collect the key.

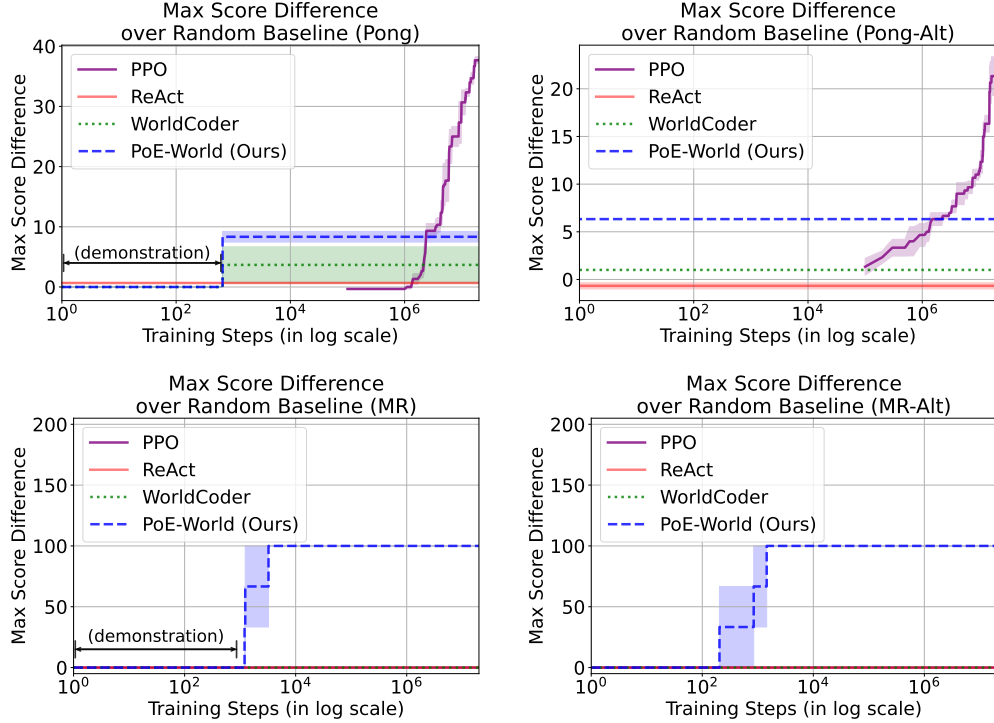


Figure 5: Maximum score differences over random baseline achieved at different number of training steps (in log scale). The text "(demonstration)" annotates the length of demonstration used to initialize WorldCoder and PoE-World, which is not displayed for the Alt games since we use the world models initialized by the demonstrations of the corresponding base games.

Instead of planning, can we use PoE-World to learn a policy?

Training a policy avoids the test-time compute cost of planning. In Figure 6, we show that pre-training a policy inside PoE-World’s world model accelerates policy learning: we first run PPO “in simulation” (in our world model), and then fine-tune in the actual Atari environment. The fine-tuned PPO achieves significantly higher score than vanilla PPO at most training steps, and while the vanilla PPO takes 1M training steps to do better than a random agent, the fine-tuned PPO takes only 200k training steps. Asymptotically the pretrained and randomly-initialized policies converge to the same value: World-model pretraining is effectively a way of warm-starting policy training.

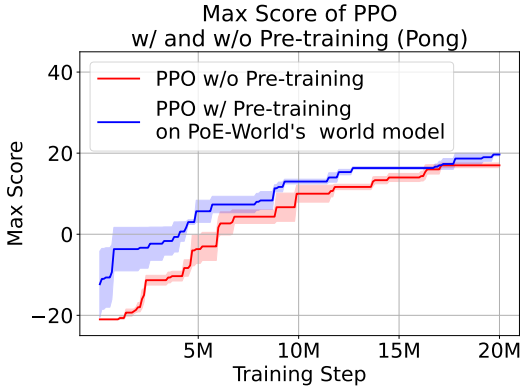


Figure 6: PPO with and without pre-training within our world model.

Next Observation Prediction Results. Table 2 and Table 3 show next observation and next observation’s object attributes prediction accuracies under the symbolic world modeling approaches, respectively. PoE-World outperforms the baselines on most settings, except for the test observation (random frames) for Pong where all methods perform similarly because a random agent rarely succeeds in hitting the ball, and so knowledge of game mechanics is not comprehensively tested.

The role of hard constraints. We investigate the role of hard constraints in our world model representation in Table 4. In practice, hard constraints act to rule out “physically impossible” scenes. PoE-World ended up enforcing constraints just for MR, since most possible scenes in Pong are already physically possible. As shown in Table 4, removing hard constraints causes the agent’s performance

Method	Next Observation Prediction Accuracy					
	Train		Test (1000 Random Frames)			
	Pong	MR	Pong	Pong-Alt	MR	MR-Alt
WorldCoder	0.33 ± 0.01	0.36 ± 0.02	0.06 ± 0.01	0.01 ± 0.01	0.10 ± 0.01	0.12 ± 0.03
PoE-World (ours)	0.51 ± 0.01	0.75 ± 0.00	0.06 ± 0.01	0.03 ± 0.00	0.31 ± 0.00	0.43 ± 0.00

Table 2: Next observation prediction accuracy on the training demonstration and 1000 random frames.

Method	Next Observation’s Object Attributes Prediction Accuracy					
	Train		Test (1000 Random Frames)			
	Pong	MR	Pong	Pong-Alt	MR	MR-Alt
WorldCoder	0.82 ± 0.00	0.75 ± 0.01	0.66 ± 0.01	0.67 ± 0.02	0.58 ± 0.00	0.65 ± 0.02
PoE-World (ours)	0.88 ± 0.00	0.93 ± 0.00	0.66 ± 0.00	0.76 ± 0.00	0.76 ± 0.00	0.83 ± 0.00

Table 3: Next observation’s object attributes prediction accuracies on the training demonstration and 1000 random frames. Non-moving objects, such as walls, platform, etc. are excluded.

Method	Planning successes		Next Observation Prediction Accuracy (Test)	
	MR	MR-Alt	MR	MR-Alt
PoE-World + Planner w/o Hard Constraints	2/9	2/9	0.30 ± 0.01	0.43 ± 0.00
PoE-World + Planner	5/9	4/9	0.31 ± 0.00	0.43 ± 0.00

Table 4: Planning successes (out of 9 tries), where the goal is for the player to grab the key, and next observation prediction accuracies on 1000 random test frames of our agent with and without hard constraints on Montezuma’s Revenge (MR) and its alternate version (MR-Alt).

to drop on both the base and alternative versions of Montezuma’s Revenge. Interestingly, we found no significant changes in next observation prediction accuracies. We hypothesize that hard constraints do not necessarily turn bad predictions into good ones. Instead, they perform damage control—they refine poor predictions just enough to make them usable for long-horizon planning.

Qualitative difference between WorldCoder’s and PoE-World’s world models. World models produced by PoE-World consists of 4000+ lines of code for Montezuma’s Revenge compared to WorldCoder’s less than 100 lines. This difference is reflected in their ability to capture the underlying causal laws: PoE-World models accurately represent important causal laws like character movement constraints with respect to platforms and ladders, while WorldCoder fails to model these mechanics, predicting that the player can “fly” around the map without any constraints. Moreover, WorldCoder models tend to hallucinate, e.g., imagining nonexistent bullet firing abilities in Montezuma’s Revenge, potentially because there is no granular way to downweigh buggy parts of a world model. PoE-World, in contrast, prunes irrelevant experts with low weights, yielding more precise world models.

5 Related Work

World models as programs has been explored in several recent works that motivate PoE-World. Similar to our work, WorldCoder [3] and CodeWorldModels [24] use LLMs to write a Python transition function. Unlike our work, they learn a single monolithic program, which severely limits scalability: our world models have an order of magnitude more code needed to model complex environments, as well as handle partial observability and non-determinism. Recent works have also explored synthesizing high-level abstract world models as programs [35, 36] to support robotic planning by integrating visual perception with symbolic reasoning. They could synergize with our work, as we have focused on learning low-level world models describing the motion of objects.

Earlier works [37, 38] learn world models in restricted (non-Turing complete) domain specific programming languages, performing well on benchmarks co-designed with their Domain Specific Language. Even earlier, Schema Networks learn a conceptually related factor-graph world model [39]. AIXI [40] is a theoretical model of reinforcement learning which considers all possible Turing

machines, mathematically related to all these works. Interestingly, AIXI also works with a history space as a conceptually elegant way of modeling partial observability. Other studies generate programs for world models primarily from natural language instructions, rather than from example interactions with the environment [41, 42, 43, 44].

Hierarchical Planning. Our approach to segmenting complex continuous tasks into symbolic states is inspired by hierarchical [45] and task-and motion [46, 47] planning. Hierarchical representations in Reinforcement Learning (RL) are often expressed as options—temporal abstractions that allow agents to reason at multiple time scales by grouping sequences of primitive actions [48]. Options can be manually specified by design or learned from experience [49]. Task and Motion Planning (TAMP) automatically discover symbolic states and action abstractions that enable generalizable symbolic plans grounded in continuous motion, however the applications of TAMP to Atari-style domains are still lacking due to complex dynamics. Similar to recent work [29, 50], we define a high-level state space in terms of object contact relations—which is key to supporting symbolic abstraction.

Alignment with human cognition. Our system, composed of simpler, specialized programmatic experts that give rise to complex behavior, is inspired by a view of mind as a community of interacting agents—a recurring theme in philosophy and cognitive science [5, 6, 52, 7]. Our modeling approach to modeling objects and actions aligns with empirical studies of event segmentation [53, 54] and hierarchical planning [55, 56]. Hierarchical state-spaces based on motion cues, such as ours, predict how people draw [13] and interpret social interactions [57], attesting to the cognitive alignment of our approach. Likewise, our programmatic representations of actions align with studies that demonstrate human concept learning to be akin to mental programs [58, 59, 11, 12], which recent work models by LLM-based program synthesis.

6 Discussion and Limitations

Compositionality. Recombining pieces of knowledge to generalize and extrapolate is a core feature of symbolic systems and, arguably, also of human cognition, ranging from natural language to abstract reasoning to everyday thought [60, 52, 61, 62]. Our approach is in this spirit. It factors its knowledge into small experts whose predictions can combine to extrapolate to scenes with more objects recomposed into new arrangements. More formally, our approach generalizes to novel "entity compositions" and "relational compositions," terms used by [63]. It should be noted, however, that our compositional factoring is orthogonal to our use of symbolic code: Programs can be monolithic [3], and neural nets can be factored [64]. Nonetheless, compositionality proved critical to scalably learning the underlying symbolic program. Synthesizing a single monolithic program is, in our view, intractable not just for the real world, but even for Atari.

Limitations. We make important assumptions, and only address part of the full model-based reinforcement learning problem. Symbolic programs expect symbolic inputs: We do not learn straight from pixels. RL involves exploration, decision-making, and reward function learning, but we do not address those problems here. However, we speculate our approach could unlock better methods for exploration: A program-structured world model exposes an interpretable interface for describing beliefs about how the world works, and efficiently exploring the world is analogous to testing the program that encodes the world model. Therefore ideas from software testing and program analysis could, in theory, be brought to bear, enabling new approaches to model-based exploration.

Acknowledgments

We thank Edward Gu and Taha Jafry for their contributions to earlier versions of this project. We are also grateful to Atharv Sonwane, Yilun Du, and Tom Silver for their valuable discussions and feedback on the paper, and we appreciate the OCArari team’s assistance and support with the OCArari package. This work was supported by an NSF CAREER grant.

References

- [1] Danijar Hafner, Jurgis Pasukonis, Jimmy Ba, and Timothy Lillicrap. Mastering diverse domains through world models. *arXiv preprint arXiv:2301.04104*, 2023.
- [2] Pedro Tsividis, Thomas Pouncy, Jaqueline L Xu, Joshua B Tenenbaum, and Samuel J Gershman. Human learning in atari. In *AAAI Spring Symposia*, volume 1, 2017.
- [3] Hao Tang, Darren Yan Key, and Kevin Ellis. Worldcoder, a model-based LLM agent: Building world models by writing code and interacting with the environment. In *The Thirty-eighth Annual Conference on Neural Information Processing Systems*, 2024.
- [4] Jerry A Fodor. *The modularity of mind*. MIT press, 1983.
- [5] Marvin Minsky. *Society of mind*. Simon and Schuster, 1986.
- [6] Daniel Clement Dennett and Daniel C Dennett. *Consciousness explained*. Penguin uk, 1993.
- [7] Gary F Marcus. *The algebraic mind: Integrating connectionism and cognitive science*. MIT press, 2003.
- [8] Michael S Gazzaniga, Richard B Ivry, and GR Mangun. *Cognitive neuroscience. the biology of the mind*. sixth edition, 2025.
- [9] John R Anderson. Acquisition of cognitive skill. *Psychological review*, 89(4):369, 1982.
- [10] Sannyuya Liu, Qing Li, Xiaoxuan Shen, Jianwen Sun, and Zongkai Yang. Automated discovery of symbolic laws governing skill acquisition from naturally occurring data. *Nature Computational Science*, 4(5):334–345, 2024.
- [11] Brenden M Lake, Ruslan Salakhutdinov, and Joshua B Tenenbaum. Human-level concept learning through probabilistic program induction. *Science*, 350(6266):1332–1338, 2015.
- [12] Brenden M Lake and Steven T Piantadosi. People infer recursive visual concepts from just a few examples. *Computational Brain & Behavior*, 3(1):54–65, 2020.
- [13] Lucas Tian, Kevin Ellis, Marta Kryven, and Josh Tenenbaum. Learning abstract structure for drawing by efficient motor program induction. *Advances in Neural Information Processing Systems*, 33:2686–2697, 2020.
- [14] Matthias Hofer, Tuan Anh Le, Roger Levy, and Josh Tenenbaum. Learning evolved combinatorial symbols with a neuro-symbolic generative model. *arXiv preprint arXiv:2104.08274*, 2021.
- [15] Joshua S Rule, Steven T Piantadosi, Andrew Cropper, Kevin Ellis, Maxwell Nye, and Joshua B Tenenbaum. Symbolic metaprogram search improves learning efficiency and explains rule learning in humans. *Nature Communications*, 15(1):6847, 2024.
- [16] Yichao Liang, Josh Tenenbaum, Tuan Anh Le, et al. Drawing out of distribution with neuro-symbolic generative models. *Advances in Neural Information Processing Systems*, 35:15244–15254, 2022.
- [17] Karl Johan Åström. Optimal control of markov processes with incomplete state information i. *Journal of mathematical analysis and applications*, 10:174–205, 1965.
- [18] David Ha and Jürgen Schmidhuber. World models. *arXiv preprint arXiv:1803.10122*, 2018.
- [19] Julian Schrittwieser, Ioannis Antonoglou, Thomas Hubert, Karen Simonyan, Laurent Sifre, Simon Schmitt, Arthur Guez, Edward Lockhart, Demis Hassabis, Thore Graepel, et al. Mastering atari, go, chess and shogi by planning with a learned model. *Nature*, 588(7839):604–609, 2020.
- [20] Łukasz Kaiser, Mohammad Babaeizadeh, Piotr Miłoś, Błażej Osipiński, Roy H Campbell, Konrad Czechowski, Dumitru Erhan, Chelsea Finn, Piotr Kozakowski, Sergey Levine, Afroz Mohiuddin, Ryan Sepassi, George Tucker, and Henryk Michalewski. Model based reinforcement learning for atari. In *International Conference on Learning Representations*, 2020.
- [21] Vincent Micheli, Eloi Alonso, and François Fleuret. Transformers are sample-efficient world models. In *The Eleventh International Conference on Learning Representations*, 2023.
- [22] Eloi Alonso, Adam Jelley, Vincent Micheli, Anssi Kanervisto, Amos Storkey, Tim Pearce, and François Fleuret. Diffusion for world modeling: Visual details matter in atari. In *Thirty-eighth Conference on Neural Information Processing Systems*, 2024.
- [23] Guangyao Zhou, Sivaramakrishnan Swaminathan, Rajkumar Vasudeva Raju, J Swaroop Guntupalli, Wolfgang Lehmach, Joseph Ortiz, Antoine Dedieu, Miguel Lázaro-Gredilla, and Kevin Murphy. Diffusion model predictive control. *arXiv preprint arXiv:2410.05364*, 2024.
- [24] Nicola Dainese, Matteo Merler, Minttu Alakuijala, and Pekka Marttinen. Generating code world models with large language models guided by monte carlo tree search. In *The Thirty-eighth Annual Conference on Neural Information Processing Systems*, 2024.
- [25] Zergham Ahmed, Joshua B Tenenbaum, Christopher J Bates, and Samuel J Gershman. Synthesizing world models for bilevel planning. *arXiv preprint arXiv:2503.20124*, 2025.

- [26] Dong C Liu and Jorge Nocedal. On the limited memory bfgs method for large scale optimization. *Mathematical programming*, 45(1):503–528, 1989.
- [27] Diederik P Kingma and Jimmy Ba. Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980*, 2014.
- [28] M. G. Bellemare, Y. Naddaf, J. Veness, and M. Bowling. The arcade learning environment: An evaluation platform for general agents. *Journal of Artificial Intelligence Research*, 47:253–279, jun 2013.
- [29] Quentin Delfosse, Jannis Blüml, Bjarne Gregori, Sebastian Sztiwernia, and Kristian Kersting. Ocatari: Object-centric atari 2600 reinforcement learning environments. *arXiv preprint arXiv:2306.08649*, 2023.
- [30] Adrien Ecoffet, Joost Huizinga, Joel Lehman, Kenneth O Stanley, and Jeff Clune. First return, then explore. *Nature*, 590(7847):580–586, 2021.
- [31] John Schulman, Filip Wolski, Prafulla Dhariwal, Alec Radford, and Oleg Klimov. Proximal policy optimization algorithms. *arXiv preprint arXiv:1707.06347*, 2017.
- [32] Shunyu Yao, Jeffrey Zhao, Dian Yu, Nan Du, Izhak Shafran, Karthik Narasimhan, and Yuan Cao. React: Synergizing reasoning and acting in language models. *arXiv preprint arXiv:2210.03629*, 2022.
- [33] Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Fei Xia, Ed Chi, Quoc V Le, Denny Zhou, et al. Chain-of-thought prompting elicits reasoning in large language models. *Advances in Neural Information Processing Systems*, 35:24824–24837, 2022.
- [34] Hao Tang, Keya Hu, Jin Zhou, Si Cheng Zhong, Wei-Long Zheng, Xujie Si, and Kevin Ellis. Code repair with llms gives an exploration-exploitation tradeoff. *Advances in Neural Information Processing Systems*, 37:117954–117996, 2024.
- [35] Yichao Liang, Nishanth Kumar, Hao Tang, Adrian Weller, Joshua B Tenenbaum, Tom Silver, João F Henriques, and Kevin Ellis. Visualpredicator: Learning abstract world models with neuro-symbolic predicates for robot planning. *arXiv preprint arXiv:2410.23156*, 2024.
- [36] Tom Silver, Rohan Chitnis, Joshua Tenenbaum, Leslie Pack Kaelbling, and Tomás Lozano-Pérez. Learning symbolic operators for task and motion planning. In *2021 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, pages 3182–3189. IEEE, 2021.
- [37] Pedro A. Tsividis, Joao Loula, Jake Burga, Nathan Foss, Andres Campero, Thomas Pouncy, Samuel J. Gershman, and Joshua B. Tenenbaum. Human-level reinforcement learning through theory-based modeling, exploration, and planning, 2021.
- [38] Ria Das, Joshua B. Tenenbaum, Armando Solar-Lezama, and Zenna Tavares. Combining functional and automata synthesis to discover causal reactive programs. *Proc. ACM Program. Lang.*, 7(POPL), January 2023.
- [39] Ken Kanksy, Tom Silver, David A Mély, Mohamed Eldawy, Miguel Lázaro-Gredilla, Xinghua Lou, Nimrod Dorfman, Szymon Sidor, Scott Phoenix, and Dileep George. Schema networks: Zero-shot transfer with a generative causal model of intuitive physics. In *International conference on machine learning*, pages 1809–1818. PMLR, 2017.
- [40] Marcus Hutter. *Universal artificial intelligence: Sequential decisions based on algorithmic probability*. Springer Science & Business Media, 2005.
- [41] Fan-Yun Sun, SI Harini, Angela Yi, Yihan Zhou, Alex Zook, Jonathan Tremblay, Logan Cross, Jiajun Wu, and Nick Haber. Factorsim: Generative simulation via factorized representation. In *The Thirty-eighth Annual Conference on Neural Information Processing Systems*, 2024.
- [42] Feiyu Zhu and Reid Simmons. Bootstrapping cognitive agents with a large language model. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 38, pages 655–663, 2024.
- [43] Lionel Wong, Jiayuan Mao, Pratyusha Sharma, Zachary S Siegel, Jiahai Feng, Noa Korneev, Joshua B Tenenbaum, and Jacob Andreas. Learning adaptive planning representations with natural language guidance. *arXiv preprint arXiv:2312.08566*, 2023.
- [44] Lin Guan, Karthik Valmeekam, Sarath Sreedharan, and Subbarao Kambhampati. Leveraging pre-trained large language models to construct and utilize world models for model-based task planning. In *Thirty-seventh Conference on Neural Information Processing Systems*, 2023.
- [45] Leslie Pack Kaelbling and Tomás Lozano-Pérez. Hierarchical task and motion planning in the now. In *2011 IEEE International Conference on Robotics and Automation*, pages 1470–1477. IEEE, 2011.
- [46] Aidan Curtis, George Matheos, Nishad Gothoskar, Vikash Mansinghka, Joshua Tenenbaum, Tomás Lozano-Pérez, and Leslie Pack Kaelbling. Partially observable task and motion planning with uncertainty and risk awareness. *arXiv preprint arXiv:2403.10454*, 2024.
- [47] Aidan Curtis, Tom Silver, Joshua B Tenenbaum, Tomás Lozano-Pérez, and Leslie Kaelbling. Discovering state and action abstractions for generalized task and motion planning. In *Proceedings of the AAAI conference on artificial intelligence*, volume 36, pages 5377–5384, 2022.

- [48] Richard S Sutton, Doina Precup, and Satinder Singh. Between mdps and semi-mdps: A framework for temporal abstraction in reinforcement learning. *Artificial intelligence*, 112(1-2):181–211, 1999.
- [49] Pierre-Luc Bacon, Jean Harb, and Doina Precup. The option-critic architecture. In *Proceedings of the AAAI conference on artificial intelligence*, volume 31, 2017.
- [50] Jannis Blüml, Cedric Derstroff, Bjarne Gregori, Elisabeth Dillies, Quentin Delfosse, and Kristian Kersting. Deep reinforcement learning via object-centric attention. *arXiv preprint arXiv:2504.03024*, 2025.
- [51] Satinder Singh, Michael James, and Matthew Rudary. Predictive state representations: A new theory for modeling dynamical systems. *arXiv preprint arXiv:1207.4167*, 2012.
- [52] Brenden M Lake, Tomer D Ullman, Joshua B Tenenbaum, and Samuel J Gershman. Building machines that learn and think like people. *Behavioral and brain sciences*, 40:e253, 2017.
- [53] Tan T Nguyen, Matthew A Bezdek, Samuel J Gershman, Aaron F Bobick, Todd S Braver, and Jeffrey M Zacks. Modeling human activity comprehension at human scale: prediction, segmentation, and categorization. *PNAS nexus*, 3(10):pgae459, 2024.
- [54] Jeffrey M Zacks. Event perception and memory. *Annual review of psychology*, 71(1):165–191, 2020.
- [55] Marta Kryven, Suhyou Yu, Max Kleiman-Weiner, Tomer Ullman, and Joshua Tenenbaum. Approximate planning in spatial search. *PLOS Computational Biology*, 20(11):e1012582, 2024.
- [56] Jan Balaguer, Hugo Spiers, Demis Hassabis, and Christopher Summerfield. Neural mechanisms of hierarchical planning in a virtual subway network. *Neuron*, 90(4):893–903, 2016.
- [57] Tianmin Shu, Marta Kryven, Tomer D Ullman, and Joshua B Tenenbaum. Adventures in flatland: Perceiving social interactions under physical dynamics. In *42d Proceedings of the Annual Meeting of the Cognitive Science Society*, 2020.
- [58] Marta Kryven, Cole Wyeth, Aidan Curtis, and Kevin Ellis. Cognitive maps are generative programs. *Cognitive Sciences Society*, 2025.
- [59] Sugandha Sharma, Aidan Curtis, Marta Kryven, Josh Tenenbaum, and Ila Fiete. Map induction: Compositional spatial submap learning for efficient exploration in novel environments. *International Conference of Learning Representations*, 2022.
- [60] Nouha Dziri, Ximing Lu, Melanie Sclar, Xiang Lorraine Li, Liwei Jiang, Bill Yuchen Lin, Sean Welleck, Peter West, Chandra Bhagavatula, Ronan Le Bras, Jena D. Hwang, Soumya Sanyal, Xiang Ren, Allyson Ettinger, Zaid Harchaoui, and Yejin Choi. Faith and fate: Limits of transformers on compositionality. In *Thirty-seventh Conference on Neural Information Processing Systems*, 2023.
- [61] Brenden M Lake and Marco Baroni. Human-like systematic generalization through a meta-learning neural network. *Nature*, 623(7985):115–121, 2023.
- [62] Elizabeth S Spelke and Katherine D Kinzler. Core knowledge. *Developmental science*, 10(1):89–96, 2007.
- [63] Atharva Sehgal, Arya Grayeli, Jennifer J. Sun, and Swarat Chaudhuri. Neurosymbolic grounding for compositional world models, 2023.
- [64] Anirudh Goyal, Alex Lamb, Jordan Hoffmann, Shagun Sodhani, Sergey Levine, Yoshua Bengio, and Bernhard Schölkopf. Recurrent independent mechanisms. In *International Conference on Learning Representations*, 2021.
- [65] Marlos C. Machado, Marc G. Bellemare, Erik Talvitie, Joel Veness, Matthew J. Hausknecht, and Michael Bowling. Revisiting the arcade learning environment: Evaluation protocols and open problems for general agents. *Journal of Artificial Intelligence Research*, 61:523–562, 2018.
- [66] Ruocheng Wang, Eric Zelikman, Gabriel Poesia, Yewen Pu, Nick Haber, and Noah D Goodman. Hypothesis search: Inductive reasoning with language models. In *The Twelfth International Conference on Learning Representations*, 2024.
- [67] Linlu Qiu, Liwei Jiang, Ximing Lu, Melanie Sclar, Valentina Pyatkin, Chandra Bhagavatula, Bailin Wang, Yoon Kim, Yejin Choi, Nouha Dziri, and Xiang Ren. Phenomenal yet puzzling: Testing inductive reasoning capabilities of language models with hypothesis refinement. *arXiv preprint arXiv:2310.08559*, 2023.
- [68] Top Piriyakulkij, Cassidy Langenfeld, Tuan Anh Le, and Kevin Ellis. Doing experiments and revising rules with natural language and probabilistic reasoning. *Advances in Neural Information Processing Systems*, 37:53102–53137, 2024.
- [69] Samuel Acquaviva. Overcoming the expressivity-efficiency tradeoff in program induction. Master’s thesis, Massachusetts Institute of Technology, 2024.
- [70] Jake Bruce, Michael D Dennis, Ashley Edwards, Jack Parker-Holder, Yuge Shi, Edward Hughes, Matthew Lai, Aditi Mavalankar, Richie Steigerwald, Chris Apps, et al. Genie: Generative interactive environments. In *Forty-first International Conference on Machine Learning*, 2024.

- [71] Geoffrey E Hinton. Training products of experts by minimizing contrastive divergence. *Neural computation*, 14(8):1771–1800, 2002.
- [72] Caelan Reed Garrett, Rohan Chitnis, Rachel Holladay, Beomjoon Kim, Tom Silver, Leslie Pack Kaelbling, and Tomás Lozano-Pérez. Integrated task and motion planning. *Annual review of control, robotics, and autonomous systems*, 4(1):265–293, 2021.
- [73] Siyuan Zhou, Yilun Du, Jiaben Chen, Yandong Li, Dit-Yan Yeung, and Chuang Gan. RoboDreamer: Learning compositional world models for robot imagination. In Ruslan Salakhutdinov, Zico Kolter, Katherine Heller, Adrian Weller, Nuria Oliver, Jonathan Scarlett, and Felix Berkenkamp, editors, *Proceedings of the 41st International Conference on Machine Learning*, volume 235 of *Proceedings of Machine Learning Research*, pages 61885–61896. PMLR, 21–27 Jul 2024.
- [74] Yilun Du, Shuang Li, and Igor Mordatch. Compositional visual generation with energy based models. In *Advances in Neural Information Processing Systems*, 2020.
- [75] Zhutian Yang, Jiayuan Mao, Yilun Du, Jiajun Wu, Joshua B. Tenenbaum, Tomás Lozano-Pérez, and Leslie Pack Kaelbling. Compositional Diffusion-Based Continuous Constraint Solvers. In *Conference on Robot Learning*, 2023.
- [76] A Paszke. Pytorch: An imperative style, high-performance deep learning library. *arXiv preprint arXiv:1912.01703*, 2019.
- [77] Rémi Coulom. Efficient selectivity and backup operators in monte-carlo tree search. In *International conference on computers and games*, pages 72–83. Springer, 2006.
- [78] Cameron B Browne, Edward Powley, Daniel Whitehouse, Simon M Lucas, Peter I Cowling, Philipp Rohlfshagen, Stephen Tavener, Diego Perez, Spyridon Samothrakis, and Simon Colton. A survey of monte carlo tree search methods. *IEEE Transactions on Computational Intelligence and AI in games*, 4(1):1–43, 2012.
- [79] Antonin Raffin, Ashley Hill, Adam Gleave, Anssi Kanervisto, Maximilian Ernestus, and Noah Dormann. Stable-baselines3: Reliable reinforcement learning implementations. *Journal of machine learning research*, 22(268):1–8, 2021.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: The claims made in the abstract and introduction accurately reflect the paper's contributions and scope.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: The paper discusses the limitations of the work performed by the authors.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [\[NA\]](#)

Justification: The paper does not include theoretical results.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [\[Yes\]](#)

Justification: We have fully described the algorithms used in the paper, with details included in the appendix.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: The code is available at <https://anonymous.4open.science/r/poe-world-neurips-628F/>

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyper-parameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: Algorithm details are described in the appendix.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [Yes]

Justification: The paper reports error bars computed over 5 seeds.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).

- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: We have provided OpenAI API cost for our experiments and compute resources used for our experiments.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: The research conducted in the paper conforms, in every aspect, with the NeurIPS Code of Ethics.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [NA]

Justification: Our work is foundational research and has no direct societal impact.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.

- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: The paper poses no such risks.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: The creators of original data are all properly credited, and the license and terms of use of the data are explicitly mentioned and properly respected.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.

- If this information is not available online, the authors are encouraged to reach out to the asset’s creators.

13. **New assets**

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [NA]

Justification: The paper currently does not release new assets.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. **Crowdsourcing and research with human subjects**

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: The paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. **Institutional review board (IRB) approvals or equivalent for research with human subjects**

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: The paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. **Declaration of LLM usage**

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

Answer: [Yes]

Justification: This paper describes the usage of LLMs in our method.

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.

A Appendix

A.1 PoE-World algorithm details

As mentioned in Section 3.2, our learning algorithm PoE-World alternates between two main steps: programmatic expert synthesis and gradient-based weight optimization. Given a trajectory, we do batch processing of size 10—we first learn $p_\theta(o_{11}|o_{1:10}, a_{1:10})$, then $p_\theta(o_{21}|o_{1:20}, a_{1:20})$, and so on.

Programmatic expert synthesis. To synthesize programmatic experts, we first need to turn the observations into natural language observations that can be input into LLMs. A sequence of consecutive observation transitions $(o_{t:t+H+1}, a_{t:t+H})$ can be transformed into a text showing the first input object list, the list of actions, and the changes to the input object list for the following timesteps, all in natural language. Table 5 shows an example text representation.

We note that in the text representation, only changes to a single object type (which is ‘player’ in Table 5) is displayed at a time. This choice is to enforce modularity between the object types: each expert put non-uniform distribution only on a single attribute of objects with a specific object type. This means Equation (2) can be further rewritten as:

$$p_\theta(o_{t+1}|o_{1:t}, a_{1:t}) \propto \prod_{obj-type} \prod_i p_i^{obj-type_expert}(o_{t+1}|o_{1:t}, a_{1:t})^{\theta_i} \quad (6)$$

where $\{p_i^{obj-type_expert}\}$ are the experts associated with each object type.

We then implement multiple (10 in total) LLM-based synthesis modules. For each object type, these modules take as input a sequence of consecutive observation transitions and output a set of programmatic experts. The output programmatic experts are pooled into a single set once all modules finish running. Having multiple synthesis modules allows each module to focus on different aspects of the environment, e.g., how objects move passively, how object moves when interacts with objects, how objects get created and deleted, etc. We describe one of our modules, ActionSynthesizer, below.

ActionSynthesizer focuses on synthesizing experts that explain how each action affects objects when the objects are interacting (touching) other objects (see example experts of this kind in Figure 2). It takes in just a single transition $(o_{t:t+1}, a_t)$. It turns this transition into a text representation as discussed above (Table 5). Then, it prompts a LLM to output causal explanations for the object changes (see the prompt Table 6). An example causal explanation is “the player objects that touch an ladder object set their y-axis velocity to -4”. With a set of causal explanations, we prompt a LLM to turn each of them into a program (see the prompt Table 7 and Table 8).

Other modules follow this template of first prompting a LLM for natural language causal explanations and then prompting a LLM to turn the explanations into programs. We refer the reader to our code at <https://github.com/topwasu/poe-world> for the implementations of other LLM-based module.

The programs synthesized by LLMs use our manually-written helper classes: `Obj`, `ObjList`, `RandomValues`, and `SeqValues`. Their docstrings are passed to the prompt. `Obj` provides a method `Obj.touches` which is the function we manually build in to help determine object contact. `Obj`

```

Example input list of objects:
player object (id = 0) with x-axis velocity = +4,
Interaction -- player object (id = 0) is touching ladder object (id = 2),
Interaction -- player object (id = 0) is touching unknown object (id = 4),

Example list of actions:
NOOP, NOOP, RIGHT

Example output list of object changes:
- The player object (id = 0) sets x-axis velocity to [+0, +0, +2]

```

Table 5: Example text representation of a sequence of consecutive transitions.

contains object attributes: object category, x/y position, x/y velocities. It also contains properties including `center_x`, `center_y`, `left_side`, `right_side`, which are calculated based on the position and velocities, and their setter methods actually modify the velocities under the hood. `ObjList` is a class that represents an object list. It has a method `ObjList.get_objs_by_obj_type` and `ObjList.create_object`. Lastly, `RandomValues` and `SeqValues` are the classes used to mark values set by a LLM. The purpose of these two classes is further discussed below.

After the programmatic experts are synthesized (see examples in Table 9), we interpret them as distributions as discussed in Section 3.2. We tell LLMs in the prompt to set attribute values as instances of `RandomValues`, as opposed to integers. This marks the attributes whose values changed by a LLM so that we can write a program-to-distribution interpreter that puts a single-peak distribution on the attribute whose value is set and uniform distributions on all other attributes. `SeqValues` is similar to `RandomValues`, but is used in multi-timestep predictions scenario, as discussed in Section 3.1.

The LLM used in the steps above is gpt-4o-2024-08-06. We implemented a disk cache for the LLM responses to avoid paying multiple times for the same prompts and seeds.

Hard constraints. The hard constraints (Section 3.1) are also synthesized similar to how we synthesize the programmatic experts. Table 10 shows the constraints learned for Montezuma’s Revenge. In Equation (3), we choose to use a disjunction rather than a conjunction because the physics in video games can be peculiar and unrealistic—a player might have their body overlap with a platform when climbing down a ladder attached to that platform. In the real world, we believe a conjunction would be a better choice.

Gradient-based weight optimization. Once we have the expert distributions $\{p_i^{expert}\}$, we can optimize their weights $\{w_i\}$ according to Equation (5). We use the L-BFGS optimizer [26] implemented in PyTorch [76] with strong Wolfe line search, learning rate = 1, number of epochs = 4, and without mini-batching. We also include a L1 regularization loss with weight = 1 so that the weights do not get too big.

We note that the weight optimization is done without taking into account the hard constraints since we would like to use gradient-based approaches.

We prune programs with weights lower than $\delta = 0.01$ after the weight optimization is done, and we prune constraints that contradict with the observations or explain less than 1% of the observations.

A.2 Planner

In a game like Montezuma’s Revenge, planning in the actual, low-level action space is hard because the number of actions required to get to the first positive reward is very high—at least 100. This means the size of search space is 8^{100} as there are 8 possible actions: NOOP, UP, DOWN, LEFT, RIGHT, FIRE, LEFTFIRE, RIGHTFIRE.

Thus, inspired by task and motion planning (TAMP), our hierarchical planner interleaves planning in the low-level motion action space with planning in a high-level abstract action space. It learns an abstract graph where the nodes are abstract states defined by object contact, and the edges represent whether the world model believes the player can traverse between the two nodes. Then, we search for a path in the abstract graph that takes the player to the goal. In Atari, a goal is usually for a player to touch a goal object (key, ball, platforms, etc.). The discovered path in the abstract graph is a high-level plan—a sequence of subgoals—for a low-level planning agent. A low-level planning agent performs online planning to choose an action and then execute it in the environment.

The planning algorithm can be described step-by-step as follows:

- Step 1: Learn an abstract graph by running a low-level motion planner in simulation on all pairs of nodes. A (ordered) pair of nodes has an edge between them if we can find at least one low-level plan to traverse between them.
- Step 2: Search for a path in the abstract graph. We use breadth-first search (BFS) here to find a path to the goal with the shortest length. The path is our high-level plan. If there is no path to the goal, go back to step 1.
- Step 3: Attempt to follow the high-level plan, completing each subgoal in order, with a low-level planning agent.

Step 4: If there is a "false" edge, update the world model, remove that edge from the graph, then go back to step 2. Otherwise, the algorithm stops, and the agent has achieved the goal.

Figure 3 shows simplified illustrations of how our hierarchical planner works.

We now discuss the implementation of our low-level motion planner:

Low-level motion planner. We implemented two low-level motion planners. The first is a variant of Monte Carlo Tree Search (MCTS) [77, 78]. It follows the same set of procedures as vanilla MCTS with two differences: first, similar to the MCTS algorithm used in [3], it approximates the value of a node using a heuristic function instead of doing a random rollout in the simulate step. The heuristic function is the Manhattan distance between the current position of the player object and the position of a goal object. We find that this function is a good estimate of how good an observation is when trying to achieve a goal. Second, the value of a node is updated as the maximum value of its children nodes, instead of the expected value, in the backpropagation step. Intuitively, using the maximum value encourages the planner to be more optimistic. The exploration parameter for MCTS is initially equal to 1 for Montezuma’s Revenge and Montezuma’s Revenge-Alt and 10 for Pong and Pong-Alt, and it increases by 10 times every 1000 iterations of MCTS.

We also implement “sticky actions” by extending the action space that MCTS searches on so that it includes repeated sequences of primitive actions with lengths 1, 4, and 8. Thus, the extended action space has $3n$ actions where n is the number of primitive actions. We include these repeated action chunks in our action space to make planning easier. Playing Atari games rarely requires players to change actions at every timestep, so repeated action chunks can be helpful.

The second low-level motion planner is a greedy search with the same heuristic function used in MCTS. At each iteration, it greedily finds the best repeated action chunk of length 8 and includes it in the plan. It backtracks if the current state leads to death no matter which action chunk of length 8 the planner chooses. It returns a plan when the player achieves the goal (touches the goal object).

Pong, Pong-Alt, and Montezuma’s Revenge-Alt agents only use the greedy planner, while Montezuma’s Revenge uses both: it first tries to find a plan with MCTS for 4000 iterations and falls back to the greedy planner if MCTS fails.

As discussed earlier, the low-level motion planner is used in two steps: to build the abstract graph by finding a plan to traverse between two abstract nodes in the world model, and to help inform a low-level planning agent that is trying to complete a subgoal.

Low-level planning agent. The agent performs online planning: it uses the low-level planner to find a plan to the goal, and then it takes a sequence of actions and replans. The agent replans when the current plan no longer takes the agent to the goal, so the agent may take several actions before replanning. We further optimize the agent by letting it replan only 40% of the times when the current plan no longer works in Montezuma’s Revenge.

Because of this introduced stochasticity, however, we find that our whole planning pipeline can give different scores in different runs even with the same initial world model, so we treat running the hierarchical planner multiple times with the same initial world model as part of training, and we run the planning algorithm 3 times on the same initial world model to get results on Montezuma’s Revenge and Montezuma’s Revenge-Alt. Our work focuses on world modeling, and we leave it to future work to increase the efficiency and performance of the hierarchical planner.

A.3 Domain details

We evaluate our methods on Atari’s Pong and Montezuma’s Revenge using the Arcade Learning Environment (ALE). The frameskip parameter is set to 3 for both games. We use OCArari to parse each image frame as a list of objects, each with an object category, a bounding box, and velocities. OCArari reverses-engineers the RAM values of Atari games to get the bounding boxes of each object.

We modify OCArari to fix a number of issues in its handling of Pong and Montezuma’s Revenge. This includes fixing bugs in object detection so that it fully detects every object in play, correcting bugs in the bounding boxes so that object interactions correctly correspond to when object bounding boxes touch/overlap, etc. We refer the reader to our code implementation <https://github.com/topwasu/poe-world> for full details.

Details on how we create Pong-Alt and Montezuma’s Revenge-Alt are below:

Pong-Alt is created by layering three Pong environments. We sync the player’s location in all three environments, so that it appears as if we have only one paddle. The enemies and balls, on the other hand, are all at different locations. We end up with one player, three enemies, and three balls.

Montezuma’s Revenge-Alt is created by stacking the lower section of the first room in the original Montezuma’s Revenge to make three platforms of different heights, connected by stairs. We stack up three lower sections of three different Montezuma’s Revenge environments. The first and the third section in the stack are flipped horizontally so that the stairs are on different sides of the room, requiring the player to jump over the skulls to reach the stairs.

ALE code uses GPL-2.0 license, and OCArari code uses MIT license.

A.4 Baseline details

PPO. PPO uses frame stacking = 4. We use the same hyperparameters as the PPO paper [31] and OCArari [29]. We use the stable-baselines3 implementation of PPO with MLP backbone (‘MlpPolicy’) [79]. For the Alt environments, we take the PPO model pretrained 20M steps on the base environments and finetune it on the Alt environments. This pretraining process is done for fair comparison with our method, since our method assumes demonstrations from the base environments when evaluating on the corresponding Alt environments.

ReAct. We write prompts that would alternate between thinking and taking actions, one for Pong and Pong-Alt and another for Montezuma’s Revenge and Montezuma’s Revenge-Alt. The frame observation is transformed into text where we provide each object’s x and y position and a list of all pairwise object interactions. This text observation is input as part of the prompt. In the prompt, we only include the 4 most recent observations (along with the 4 most recent thinking actions and 4 most recent taken actions) as each observation is quite long in text.

WorldCoder. We use the official WorldCoder implementation [3] but replace the existing prompts with new ones that are tailored towards our text representation of object-centric Atari frames. The prompts can be found in our codebase <https://github.com/topwasu/poe-world>. The instantiation of WorldCoder in [3] actually has its own planner implemented, but for fair comparison with our method, we use only the world modeling part of that system and combine it with our own planner so that the planner is the same for both WorldCoder and our method. We choose to use our own planner instead of theirs since ours is hierarchical.

WorldCoder code uses MIT license.

A.5 Compute resources and execution time

Compute Resources. For the world modeling part, our experiments are run on 4 CPUs (Cascade-Lake, IceLake, or SapphireRapids) with 64 GB memory. PoE-World and WorldCoder uses a budget of \$20 worth of OpenAI credit per run. For the planner, we also run it mostly on 4 CPUs, but for the part where we need to build an abstract graph by running many low-level planners, we parallelize it on multiple compute jobs on a job scheduling cluster, so we might be using 100 CPUs at a time.

Execution time. PoE-World alone without the planner tends to take around 8 hours to run (this includes the time we need to wait for OpenAI LLM requests). The planner running time varies, but most runs finish under 24 hours.

I'll give you an input list of objects and an output list of object changes, and I want you to list 4 possible reasons for the effects

Here's an example with player objects:

Example input list of objects:

player object (id = 0) with x-axis velocity = +0 and y-axis velocity +2,
Interaction -- player object (id = 0) is touching ladder object (id = 2),
Interaction -- player object (id = 0) is touching unknown object (id = 4),

Example output list of object changes:

- The player object (id = 0) sets x-axis velocity to +0
- The player object (id = 0) sets y-axis velocity to -4

Example reasons:

1. The player objects that touch an unknown object set their x-axis velocity to +0
2. The player objects that touch an unknown object set their y-axis velocity to -4
3. The player objects that touch an ladder object set their x-axis velocity to +0
4. The player objects that touch an ladder object set their y-axis velocity to -4

Please output a list of 4 reasons of the {obj_type} objects for the following input and output list of objects.

Input list of objects:
{input}

Output list of object changes:
{effects}

Please follow these rules for your output:

1. make sure each reason only talks about one object change
2. do not talk about IDs

Table 6: A prompt in ActionSynthesizer used to get LLM to provide causal explanations of the object changes.

We observe that the possible effects of {action} on {obj_type} objects include {obs_lst_txt}

We want to synthesize python functions that implements these effects. The format of the functions should be

```
def alter_{obj_type}_objects(obj_list: ObjList, action: str) -> ObjList:
    {obj_type}_objs = obj_list.get_objs_by_obj_type('{obj_type}') # get all Obj of
obj_type '{obj_type}'
    for {obj_type}_obj in {obj_type}_objs: # {obj_type}_obj is of type Obj
        # You can assume {obj_type}_obj.velocity_x and {obj_type}_obj.velocity_y are
integers.
        pass
    return obj_list
```

And here are the docstrings for relevant classes:

```
class RandomValues:
    Use this class to express the possibility of random values. Example x =
RandomValues([x + 2, x - 2])

    Attributes:
        values (list[ints]): list of possible values

    Methods:
        __init__(values):
            Initialize an instance

class Obj:
    Attributes:
        id (int): id of the object
        obj_type (string): type of the object
        velocity_x (int | RandomValues): x-axis velocity of the object
        velocity_y (int | RandomValues): y-axis velocity of the object
        deleted (int | RandomValues): whether this object gets deleted (1 if it does
and 0 if it does not)

    Methods:
        touches(obj: Obj) -> bool:
            Returns whether this Obj is touching the input obj (True/False)

class ObjList:
    Attributes:
        objs (list of Obj)

    Methods:
        get_objs_by_obj_type(obj_type: str) -> list[Obj]:
            Returns list of objects with the input obj_type

        create_object(obj_type: str, x: int, y: int) -> ObjList:
            Returns a new instance of ObjList with the new object (with obj_type, x,
y) added
```

Table 7: First half of a prompt in ActionSynthesizer used to turn the natural language causal explanations into programs.

Please output {n} different alter_{obj_type}_objects functions that explains each of the {n} possible effects of action '{action}' following these rules:

1. Each function should make changes to one attribute -- this could be the x-axis position, y-axis position, creation of object, or deletion of object.
2. Always use RandomValues to set attribute values. If there are conflicting changes to an attribute, instantiate RandomValues with a list of all possible values for that attribute.
3. Use Obj.touches to check for interactions.
4. Avoid setting each attribute value for each {obj_type} object more than once. For example, use 'break' inside a nested loop.
5. You can assume the velocities of input objects are integers.
6. Please use if-condition to indicate that the effects only happen because of action '{action}'

Format the output as a numbered list.

Table 8: Second half of a prompt in ActionSynthesizer used to turn the natural language causal explanations into programs.

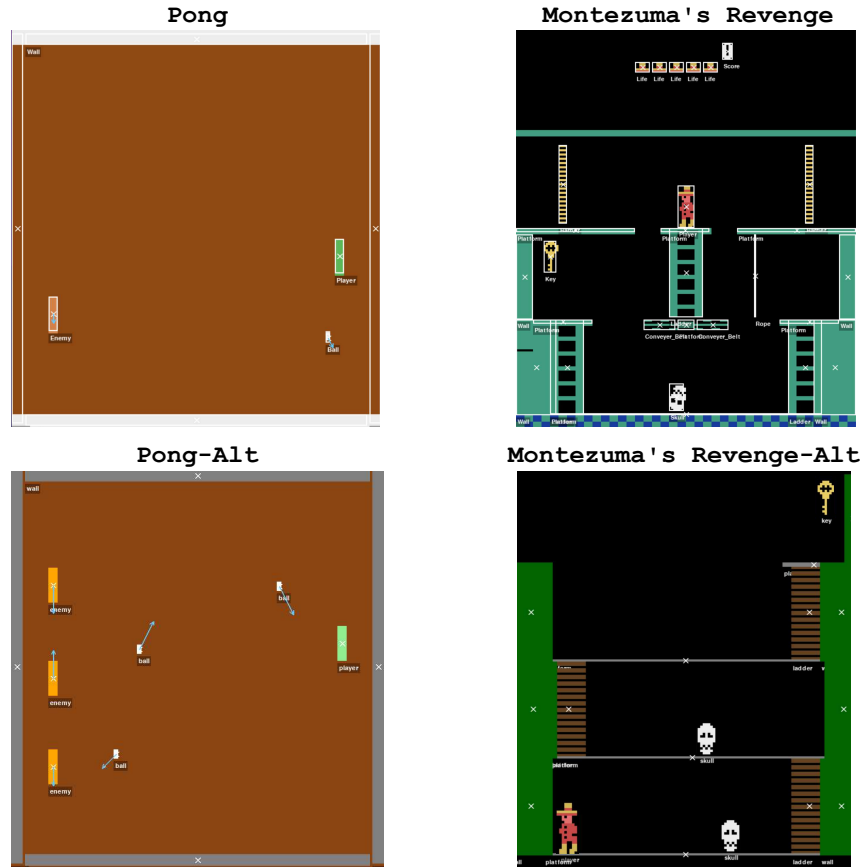


Figure 7: Screenshots of both base and alternative environments for Pong and Montezuma's Revenge.

```

# Example program 1
def alter_player_objects(obj_list: ObjList, action: str, touch_side=3,
touch_percent=0.6) -> ObjList:
    if action == 'NOOP':
        player_objs = obj_list.get_objs_by_obj_type('player')
        conveyer_belts = obj_list.get_objs_by_obj_type('conveyer_belt')
        for player_obj in player_objs:
            for conveyer_belt in conveyer_belts:
                if player_obj.touches(conveyer_belt, touch_side,
touch_percent):
                    player_obj.velocity_x = RandomValues([-1])
                    break
        return obj_list

# Example program 2
def alter_player_objects(obj_list: ObjList, action: str, touch_side=3,
touch_percent=1.0) -> ObjList:
    if action == 'FIRE':
        player_objs = obj_list.get_objs_by_obj_type('player')
        platform_objs = obj_list.get_objs_by_obj_type('platform')
        for player_obj in player_objs:
            for platform_obj in platform_objs:
                if player_obj.touches(platform_obj, touch_side,
touch_percent):
                    player_obj.velocity_y = RandomValues([-6])
                    break
        return obj_list

# Example program 3
def alter_player_objects(obj_list: ObjList, action: str, touch_side=3,
touch_percent=0.3) -> ObjList:
    if action == 'RIGHTFIRE':
        player_objs = obj_list.get_objs_by_obj_type('player')
        platform_objs = obj_list.get_objs_by_obj_type('platform')
        for player_obj in player_objs:
            for platform_obj in platform_objs:
                if player_obj.touches(platform_obj, touch_side,
touch_percent):
                    player_obj.velocity_y = SeqValues([-6, -7, -4, 0, 2, 6,
9])
                    break
        return obj_list

```

Table 9: Example synthesized programmatic experts.

```

# Constraint 1: Player's body must align with the center of the rope
def c1(obj_list: ObjList, _, touch_side=2, touch_percent=0.3) -> tuple:
    touch_ids, satisfied_ids = [], []
    player_objs = obj_list.get_objs_by_obj_type('player') # get all Obj of type
'player'
    rope_objs = obj_list.get_objs_by_obj_type('rope') # get all Obj of type 'rope'
    for player_obj in player_objs: # player_obj is of type Obj
        for rope_obj in rope_objs: # rope_obj is of type Obj
            if player_obj.touches(rope_obj, touch_side, touch_percent):
                touch_ids.append(player_obj.id)
                if player_obj.center_x == rope_obj.center_x:
                    satisfied_ids.append(player_obj.id)
    return touch_ids, satisfied_ids

# Constraint 2: Player's feet must align with the top of the conveyer belt
def c2(obj_list: ObjList, _, touch_side=3, touch_percent=0.1) -> ObjList:
    touch_ids, satisfied_ids = [], []
    player_objs = obj_list.get_objs_by_obj_type('player') # get all Obj of type
'player'
    conveyer_belt_objs = obj_list.get_objs_by_obj_type('conveyer_belt') # get all
Obj of type 'conveyer_belt'
    for player_obj in player_objs: # player_obj is of type Obj
        for conveyer_belt_obj in conveyer_belt_objs: # conveyer_belt_obj is of type
Obj
            if player_obj.touches(conveyer_belt_obj, touch_side, touch_percent):
                touch_ids.append(conveyer_belt_obj.id)
                if player_obj.bottom_side == conveyer_belt_obj.top_side:
                    satisfied_ids.append(conveyer_belt_obj.id)
    return touch_ids, satisfied_ids

# Constraint 3: Player's feet must align with the top of the platform
def c3(obj_list: ObjList, _, touch_side=3, touch_percent=0.5) -> ObjList:
    touch_ids, satisfied_ids = [], []
    player_objs = obj_list.get_objs_by_obj_type('player') # get all Obj of type
'player'
    platform_objs = obj_list.get_objs_by_obj_type('platform') # get all Obj of type
'platform'
    for player_obj in player_objs: # player_obj is of type Obj
        for platform_obj in platform_objs: # platform_obj is of type Obj
            if player_obj.touches(platform_obj, touch_side, touch_percent):
                touch_ids.append(platform_obj.id)
                if player_obj.bottom_side == platform_obj.top_side:
                    satisfied_ids.append(platform_obj.id)
    return touch_ids, satisfied_ids

# Constraint 4: Player's body must align with the center of the ladder
def c4(obj_list: ObjList, _, touch_side=3, touch_percent=1.0) -> ObjList:
    touch_ids, satisfied_ids = [], []
    player_objs = obj_list.get_objs_by_obj_type('player') # get all Obj of type
'player'
    ladder_objs = obj_list.get_objs_by_obj_type('ladder') # get all Obj of type
'ladder'
    for player_obj in player_objs: # player_obj is of type Obj
        for ladder_obj in ladder_objs: # ladder_obj is of type Obj
            if player_obj.touches(ladder_obj, touch_side, touch_percent):
                touch_ids.append(ladder_obj.id)
                if player_obj.center_x == ladder_obj.center_x:
                    satisfied_ids.append(ladder_obj.id)
    return touch_ids, satisfied_ids

```

Table 10: The collection of constraints synthesized for Montezuma's Revenge. Function names have been shortened to save space.