
From Experiments to Discovery: A Principled Approach to Measuring How Well LLMs Do Science

Kanishk Gandhi* Michael Y. Li* Lyle Goodyear Agam Bhatia

Louise Li Aditi Bhaskar Mohammed Zaman Noah D. Goodman

Stanford University

Abstract

Understanding the world and explaining it with scientific theories is a central aspiration of artificial intelligence research. Proposing theories, designing experiments to test them, and then revising them based on data are key to scientific discovery. Despite the promise of LLM-based scientific agents, no benchmarks systematically test their ability to propose scientific models, collect experimental data, and revise them in light of new data. We introduce `BoxingGym`, a benchmark with 10 environments for evaluating experimental design (*e.g.*, collecting data to test a scientific theory) and model discovery (*e.g.*, proposing and revising scientific theories). To enable quantitative and principled evaluation, we implement each environment as a generative probabilistic model with which a scientific agent can run interactive experiments. These probabilistic models are drawn from various real-world scientific domains ranging from psychology to ecology. To evaluate a scientific agent’s ability to collect informative experimental data, we compute the expected information gain (EIG), an information-theoretic quantity which measures how much an experiment reduces uncertainty about the parameters of a generative model. A good scientific theory is a concise and predictive explanation. To quantitatively evaluate model discovery, we ask a scientific agent to explain their model and evaluate whether this explanation helps another scientific agent make more accurate predictions. We evaluate several open and closed-source language models of varying sizes. We find that larger models (32B) consistently outperform smaller variants (7B), and that closed-source models generally achieve better results than open-source alternatives. However, all current approaches struggle with both experimental design and model discovery, highlighting these as promising directions for future research. ²

“To understand a system, you must perturb it.”
– George Box (*ad sensum*)

1 Introduction

Helping humans understand the world (and themselves) by discovering scientific theories is a foundational goal of artificial intelligence research [30]. Proposing theories about the world, conducting experiments to test them, and revising them based on data is central to this process [9]. Recent advances in large language models (LLMs), have shown promising potential for accelerating scientific discovery. LLMs have extensive scientific knowledge [2], strong inductive reasoning capabilities [52, 42], and the ability to propose models of data [26, 27, 11]. These promising results suggest that LLMs, functioning as autonomous agents, could be well-suited for experimental design (*i.e.*, collecting informative experiments to test scientific theories) and model discovery (*i.e.*, developing interpretable models based on experimental data).

*Equal Contribution. Corresponding author: kanishk.gandhi@stanford.edu

²Project: <https://github.com/kanishkg/boxing-gym/>

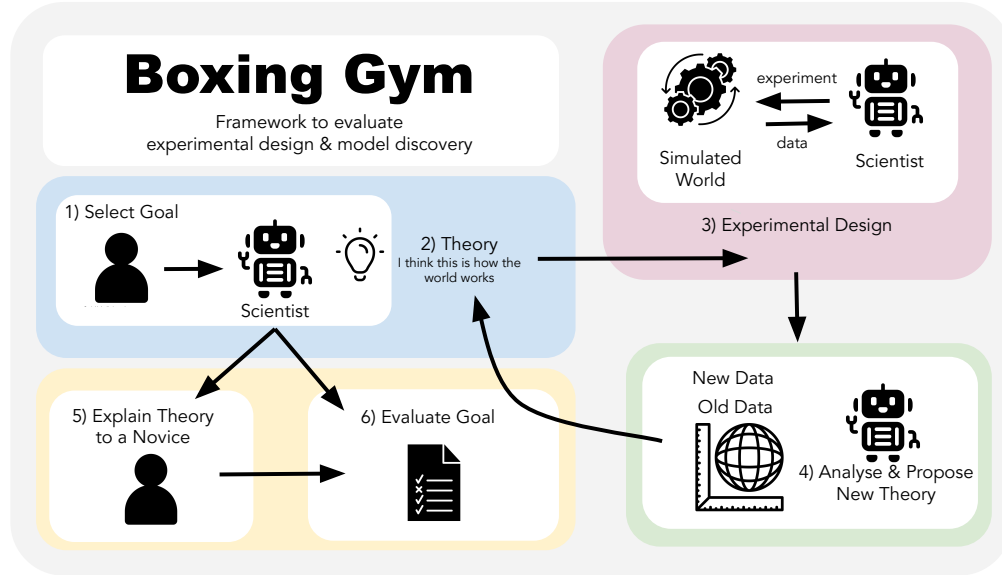


Figure 1: **Overview of BoxingGym.** The BoxingGym Framework is designed to holistically evaluate experimental design and model discovery capabilities in the spirit of George Box [9]. 1) The process starts with a user defining a goal for the scientist agent. 2) The scientist formulates a theory. 3) This theory guides the experimental design, where the scientist interacts with a simulated world to gather new data. 4) The scientist then analyzes the new and old data to propose and refine theories. This iterative process continues for several iterations. 5) The scientist is then asked to explain the findings to a novice. 6) We evaluate the novice and the scientist by casting the goal as a prediction problem.

Previous work has evaluated automated experimental design and model discovery in isolation [16, 17, 15, 26]. However, they are fundamentally coupled in real-world settings: scientists collect experimental data to build better models and better models inform better experiments. While scientific agents are promising, there is currently no systematic way to evaluate an agent’s ability to propose scientific models, collect experimental data, and revise them in light of new data. This motivates the need for a benchmark that evaluates an agent’s capabilities holistically in an integrated scientific discovery pipeline.

We outline the key desiderata for a framework that evaluates experimental design and model discovery: (1) The framework should enable the agent to *actively experiment* with the environment without requiring the agent to perform time-consuming and resource-intensive real-world lab experiments. (2) Since scientific theories come in different forms, the framework should flexibly accommodate *different representations of scientific theories*. (3) The framework should evaluate experimental design and model discovery in an *integrated* way. (4) Science is often *goal-directed* or driven by an inquiry. For example, a biologist might perform experiments with the goal of identifying cellular mechanisms underlying circadian rhythm in mammals. Our framework should allow users to specify high-level goals to guide the agent’s discovery process. Our desiderata are inspired by the framework for scientific modeling introduced by George Box [7, 8], which emphasizes an iterative process of building models, designing experiments to test them, and revising them accordingly.

To achieve these desiderata, we introduce BoxingGym (Fig. 1) a flexible framework for evaluating experimental design and model discovery with autonomous agents. Our benchmark consists of 10 *environments* grounded in real-world scientific models. To enable agents to actively experiment, we implement each environment as a generative model. This key design choice makes simulating active experimentation tractable because it corresponds to sampling from the underlying generative model, conditioned on the experimental interventions. To accommodate various representations of scientific theories, all environments are designed with a flexible language based interface (Fig. 2). Finally, our environments can be instantiated with different goals, or intents for inquiry, that encourage the agent to adapt their experimentation towards accomplishing the goal (*e.g.*, understand the parameters underlying participant behavior in a psychology study) by specifying the goal in language.

We introduce principled evaluation metrics that measure the quality of experiments and discovered models. To evaluate experimental design, we draw from *Bayesian optimal experimental* (BOED)

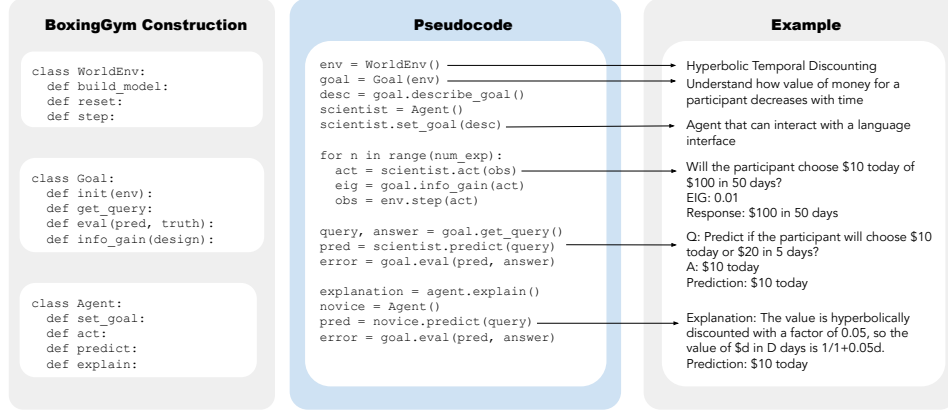


Figure 2: **Python pseudocode examples.** (left) BoxingGym is instantiated as modular classes and methods for the environment (WorldEnv), goals (Goal), and agents (Agent). (center) Pseudocode illustrating the workflow of setting goals, performing experiments, predicting outcomes, and providing explanations. (right) An example, hyperbolic temporal discounting, where the agent predicts a participant’s choice between immediate and delayed rewards and explains the concept to a novice.

design [43] and use *expected information gain* (EIG) to measure the informativeness of an experiment. EIG captures how much an experiment reduces uncertainty in the parameters of a generative model and, importantly, this measure complements our decision to implement environments as generative models. To evaluate model discovery, we take inspiration from the fact that science is a communicative endeavor. We propose a *communication-based* evaluation strategy: we ask a scientist agent to distill their experiments into a natural language explanation and evaluate how much that explanation empowers a novice agent, who does not have access to the experiments conducted by the scientist, to make accurate predictions about the environment.

We evaluate several open and closed-source language models ranging from 7B to 32B parameters. We find that larger models consistently outperform smaller variants, and closed-source models generally achieve better results than open-source alternatives. We also evaluate Box’s Apprentice [26], which augments language models with statistical modeling capabilities, but find that this augmentation does not reliably improve performance. Notably, we observe substantial variation in difficulty across environments, which remaining challenging even for the strongest models. Promisingly, some environments show clear performance improvements with model scale. These results highlight significant opportunities for improving automated scientific reasoning.

2 Related Works

Optimal Experimental Design. Bayesian optimal experimental design (BOED) is a principled framework for designing maximally informative experiments across various disciplines [48, 12, 34]. While theoretically appealing, BOED’s practical implementation is challenging due to the intractability of information gain metrics like expected information gain (EIG). Although several methods [43, 16, 17] exist to approximate EIG, they assume the data follows a fixed generative model—limiting their utility when model revision is needed as new data is collected.

Automated Model Discovery. Automated model discovery from data has been a long-standing goal in AI, aiming to build interpretable models that capture underlying patterns in data—from physical laws [6, 31] to nonparametric regression [15]. Recent work [26, 27] has integrated language models into this process, leveraging their ability to both propose and critique candidate models, demonstrating their potential as tools for automated model discovery. This work highlights the potential of using language models as a powerful tool for model discovery.

Reasoning and Exploration with LLMs. Language models have shown promising capabilities in both deductive reasoning (deriving consequences from hypotheses) Saparov et al. [46], Saparov and He [45], Poesia et al. [41] and inductive reasoning (inferring hypotheses from observations) [52, 42]. While reinforcement learning has improved LLMs’ reasoning abilities [23, 21, 20, 22], these advances have primarily focused on deterministic, verifiable systems rather than the stochastic data typical in scientific discovery. Efficient exploration and information-seeking are crucial for experimental design and model building. Recent work [36, 32, 19, 18, 47, 25] has investigated in-context exploration

strategies and shown how language models can learn how to search and explore directly through sequence modeling, developing effective search strategies in language.

Interactive Environments. Drawing inspiration from established reinforcement learning principles [10, 33], BoxingGym adopts the modularity and simplicity of classic environments like OpenAI Gym while shifting focus to evaluation rather than agent training. While recent work has expanded interactive benchmarks to language agents —spanning tasks from software debugging [24] to automated scientific research[35, 28], our work advances this direction by introducing a principled framework for evaluating language agents’ capabilities in iterative experimental design and model discovery.

3 Boxing Gym

3.1 Problem Formulation.

We formalize experimental design and model discovery using probabilistic modeling and Bayesian optimal experimental design (BOED). In BoxingGym, each environment is implemented as a generative model defining a joint distribution over the experimental outcome y , experimental design d , and unobserved parameters θ . This joint distribution is defined in terms of a prior distribution over θ , $p(\theta)$ and a *simulator* $p(y|\theta, d)$ which is a model of the experimental outcome y given parameters θ and design d . For example, in a psychology experiment, θ could be the parameters of a behavioral model of participants, d could be the questions posed to participants, and y could be the participant’s response to d . Running an experiment corresponds to choosing a design d and observing a sample y from the marginal predictive distribution conditioned on that design, i.e., $y \sim p(y|d) = E_{p(\theta)}[p(y|\theta, d)]$ ³.

3.2 Evaluation

3.2.1 Evaluating experimental design via Expected Information Gain

To evaluate experimental design, we take inspiration from the Bayesian OED literature [16, 17]. Crucially, our choice to implement environments as generative models enables us to leverage this literature. For each domain, we have an underlying predictive model $p(y|\theta, d)$. We quantify the *informativeness* of a design d through the expected information gain (EIG), that measures the reduction in posterior uncertainty about the model parameters θ after running an experiment d . Below, H is the Shannon entropy.

$$\text{EIG}(d) = \mathbb{E}_{p(y|d)} [H[p(\theta)] - H[p(\theta|y, d)]]$$

Since the EIG is typically not available in closed-form, we use a Nested Monte Carlo estimator

$$\hat{\mu}_{\text{NMC}}(d) = \frac{1}{N} \sum_{n=1}^N \log \left(\frac{p(y_n|\theta_{n,0}, d)}{\frac{1}{M} \sum_{m=1}^M p(y_n|\theta_{n,m}, d)} \right) \quad \text{where} \quad \theta_{n,m} \stackrel{\text{i.i.d.}}{\sim} p(\theta), \quad y_n \sim p(y|\theta = \theta_{n,0}, d)$$

We chose this estimator because it is a consistent estimator of the true EIG [43] and is straightforward to implement. EIG measures the value of an experiment under the assumption that the true distribution of experimental outcomes is modeled by $p(y|d)$. In general, this assumption is not true, but EIG is still a useful measure since we generate data from an underlying model in our benchmarks.

3.2.2 Evaluating model discovery via communication

To evaluate the quality of a model, we use standard model evaluation metrics (e.g., prediction MSE) and a communication-based metric that takes advantage of the natural language interface. In particular, a *scientist agent* interacts with an environment through experiments. After these experiments, we ask the scientist agent to synthesize their findings through an *explanation*. We then evaluate how much that explanation enables a *novice agent* to make more accurate predictions about the environment without any additional experiments. Since a good explanation is both *predictive* and *parsimonious*, we set a token limit on the explanation. Crucially, this evaluation method can accommodate different forms of scientific theories. In our experiments, we ask the scientist agent to produce a statistical model and then distill the model into a natural language explanation to guide the novice agent.

³In the sequential setting, we replace the prior $p(\theta)$ with the posterior $p(\theta|y, d)$.

3.2.3 Evaluating goals via prediction

To evaluate success at achieving a specific goal (*e.g.*, how do the populations of predator and prey change with time) we employ a prediction target (*e.g.*, predict the population of predators at a particular time) and calculate a standardized prediction error. First, we compute the error between the predicted and true values. Then, we standardize this error with respect to the prior predictive mean, which is obtained by assuming a uniform prior over the design space. Specifically, for each domain, we sample a design d uniformly from the design space and a parameter θ from the prior distribution $p(\theta)$. We then generate samples from the predictive model $p(y|\theta, d)$ and average over multiple d and θ to obtain the prior predictive mean μ_0 and variance σ_0 . Let $\{y_i\}_{i=1}^n$ be the ground truth outputs for inputs $\{x_i\}_{i=1}^n$, and let $\{\hat{y}_i\}_{i=1}^n$ be the predictions of the agent. The standardized prediction error is then calculated using these quantities, providing a measure of the agent’s performance relative to the prior predictive mean. We use a domain-specific function f computing the discrepancy between a prediction $\hat{y}_i$ and ground truth value y_i (*e.g.*, MSE). We compute the errors $\epsilon_i = f(\hat{y}_i, y_i)$ and $\epsilon_{\mu_0} = f(\mu_0, y_i)$. Finally, we compute the standardized error as $\frac{\epsilon_i - \epsilon_{\mu_0}}{\sigma_0}$. Crucially, since this metric is computed with respect to the prior predictive, this metric can be negative.

3.3 Design Decisions in Constructing BoxingGym

We outline the key design decisions of BoxingGym that allow it to capture key aspects of scientific discovery within a flexible, simulated, and extensible environment.

Discovery via active experimentation. The agent actively interacts with the environment by conducting experiments, reflecting the real-world coupling of experimentation and model discovery. This approach assesses the agent’s ability to gather relevant data and refine its models based on experimental results.

Real-world scientific models. Our environments are grounded in real-world scientific models from several domains, ensuring the benchmark tests the agent’s ability to handle realistic scenarios. We implement these environment as pymc generative models to make active experimentation an automatic and tractable process.

Goal-driven discovery. Each environment has a specific goal, mirroring the inquiry-driven nature of scientific research. This encourages the agent to engage in targeted experimentation.

Language-based interface for experiments. We use a language-based interface for our experiments because it’s flexible (*i.e.*, scientific domains can generally be described in language), easily integrates with LLMs, and interpretable to humans.

Emphasis on Measuring Discovery with Explanations. BoxingGym places a strong emphasis on measuring the quality of the agent’s discoveries through the explanations it can provide after experimentation (§3.2.2). This design decision is motivated by two considerations. From a theoretical perspective, science is fundamentally about developing better theories, and scientific theories are explanations of observed phenomena. From a practical perspective, communicating findings to the broader scientific community is an essential aspect of scientific research. By using language, we do not have to commit to a particular representation of a scientific theory. We illustrate this flexibility, by showing how different representations can be easily integrated within our method for measuring natural language explanations.

Extensible/modular environments for benchmarking agents. BoxingGym is easily extensible and modular, enabling researchers to integrate new environments and test different agents with minimal effort. We illustrate this in Fig. 2 which provides a pseudo-code example of how to implement a new environment and goal in BoxingGym.

3.4 Domains

BoxingGym consists of 10 environments (see App. D for full details) that cover a range of scientific domains and test different aspects of experimental design and model discovery. Some environments are designed to test optimal experiment design, while others focus on model discovery or involve simulated neuro-symbolic human participants.

195 **Location finding.** [17] In an n -dimensional space with k signal-emitting sources, the scientist
196 measure signals at any grid location. Goals include predicting the signal at any point or locating the
197 sources.

198 **Hyperbolic temporal discounting.** [17] The scientist observes a participant’s choices for different
199 immediate rewards (ir), delayed rewards (dr), and delay periods (D days) Fig. 2 (right). Goals
200 include predicting choices of a participant or discount factors.

201 **Death process.** [17] A disease spreads at an infection rate. The scientist can measure the number
202 of infected individuals at different points of time to predict future infections or the infection rate.

203 **Item Response Theory (IRT).** [44] In this environment, there is a set of students and a set of
204 questions. The experimenter can observe the correctness of a student’s response to a particular
205 question. The goal is to discover the underlying model that relates student ability and question
206 difficulty to the probability of a correct response.

207 **Animal growth curves.** [29] An experimenter can observe the length of a dugong at a particular
208 age. The goal is to discover the underlying growth model of dugongs.

209 **Population growth dynamics.** [29] An experimenter can observe the population of peregrines at a
210 particular point in time. The goal is to discover the underlying population dynamics model. This is
211 tested by asking the experimenter to predict population dynamics at a particular point in time.

212 **Mastectomy Survival analysis.** [13] The experimenter can observe if a patient is alive after a
213 mastectomy, including metastasis status and time since surgery. The goal is to predict survival
214 probabilities for new patients.

215 **Predator-Prey dynamics.** [51] This simulates predator-prey populations over time. The goal is to
216 discover models like the Lotka-Volterra equations to predict future populations.

217 **Emotion from outcome.** [37] Participants guess a player’s emotions after a gambling game’s outcome.
218 The experimenter designs games with varied probabilities and prizes to model how participants judge
219 the emotions of a player from outcomes. Human participants are simulated using a probabilistic
220 model translated into natural language by a language model.

221 **Moral Machines.** [5] Participants face moral dilemmas, choosing which group an autonomous car
222 should save. Experimenters manipulate group compositions and required actions to model moral
223 decision-making. Human participants are simulated with a probabilistic model, and their actions are
224 translated into natural language by a language model.

225 4 Experiments

226 We conduct experiments to evaluate the performance of two baseline agents on BoxingGym . Our
227 goal is to assess their ability to perform experimental design and theory building across a diverse set
228 of environments. We benchmark two types of agents: a standard language model (GPT-4o, OpenAI
229 [38]) and a language model augmented with symbolic reasoning capabilities (Box’s Apprentice).

230 **LLM Agent.** We consider 6 LLMs, GPT-4o [38], Claude-3.7-sonnet [3], Qwen-2.5-32b-instruct,
231 Qwen-2.5-7b-instruct [54], and reasoning variants OpenThinker-32b, and OpenThinker-7b [50]; the
232 reasoning variants are finetuned on math and coding task. We prompt these models to interact with
233 our environment, purely through natural language, without additional tools (see Fig. 2, see App. B
234 for details).

235 **Box’s Apprentice.** The apprentice agent augments language models by enabling them to implement
236 generative models of observed data. For model discovery, the agent writes a pymc program [26] after
237 10 experiments, which is then fit and provided to the scientist explaining findings to the novice. For
238 experimental design, the agent creates and uses these models to guide subsequent experiments.

239 **Experiment Setup.** For each environment, we run the agents for 5 independent trials. At each
240 step, the agent chooses to perform an experiment, by specifying a design, and observes the outcome.

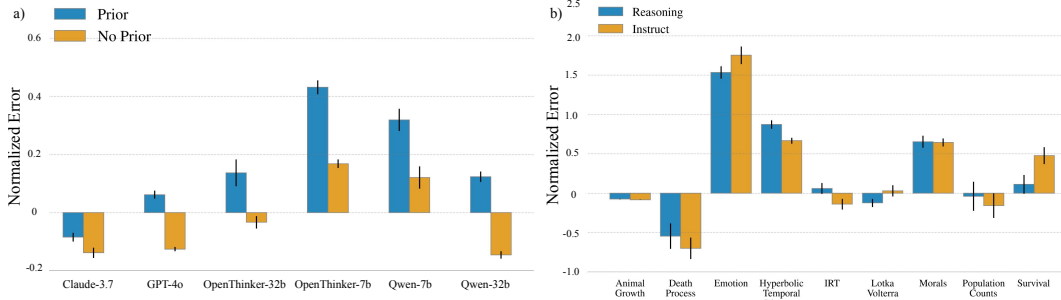


Figure 3: **Normalized Error Compared across Models.** (a) Comparison of the normalized errors for different LLMs with or without prior information included in the prompt. (b) Comparison of reasoning models (OpenThinker) and instruct models (Qwen) across environments. Error bars are the standard error across 5 runs.

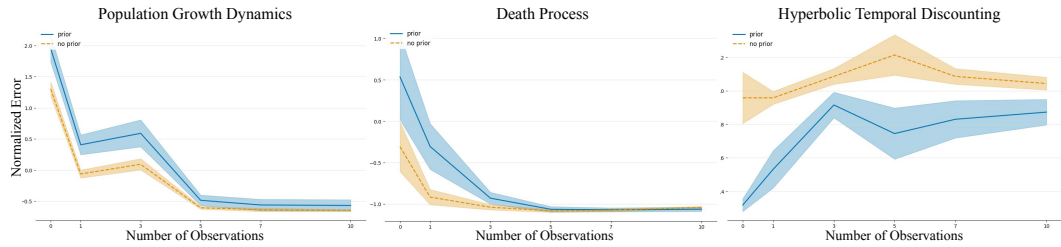


Figure 4: **Normalized Errors Over Number of Observations.** Normalized errors for the LLM agent with gpt-4o with prior information (solid blue) and without prior information (dotted yellow) across three domains: Population Growth Dynamics (left), IRT (center) and Hyperbolic Discounting (right). Error bars are the standard error across 5 runs.

241 After a fixed number of steps (0, 1, 3, 5, 7, 10), we evaluate the agent’s performance using the
 242 metrics described earlier §3.2. The performance of models is averaged across 5 runs and over 10
 243 evaluation points. We also explore a *prior* vs *no prior* condition to investigate whether domain
 244 knowledge helps or hinders scientific discovery. In the prior condition, we give the LM full context
 245 about the problem domain (e.g., “you are observing how participants balance delayed vs immediate
 246 rewards”), simulating scientists with background knowledge. In the no prior condition, we remove
 247 this context and describe the setting in a domain-agnostic way (e.g., “you receive a tuple of three
 248 values”), resembling reasoning from raw observations without preconceptions. This tests whether
 249 prior knowledge scaffolds discovery or creates biases that constrain exploration.

250 4.1 Experimental Design Evaluation

251 **Setup.** To evaluate the agents’ performance, we first assess their ability to gather valuable informa-
 252 tion through their experiment selection and then measure how effectively they use this information
 253 to predict the environment. The Expected Information Regret (EI Regret) compares the Expected
 254 Information Gain (EIG) (§3.2.1) of the agent’s chosen experiments to the maximum EIG achievable
 255 from 100 random experiments. Lower EI Regret indicates more informative experiment selection.

256 **Prior information does not improve performance.** We find that models often perform better
 257 when given no prior information after 10 experiments (Fig. 3a). In some cases, this is because the
 258 LLM makes an overly strong assumption about the environment (e.g., the signal decay is symmetric
 259 around the origin) and does not revise the assumption after more experiments; this is consistent with
 260 findings reported by Li et al. [26]. In other cases, such as the hyperbolic discounting environment
 261 (Fig. 4, right), the model overfits to limited observations.

262 **More experiments generally lead to better predictions.** We plot the learning trajectories for
 263 three environments in (Fig. 4). The agent’s average prediction error decreases as it performs more

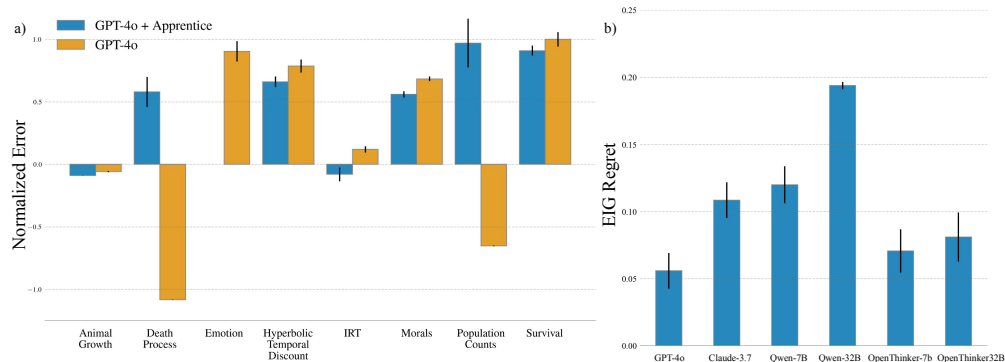


Figure 5: (a) Comparison of the Box’s Apprentice with an LLM agent. (b) EIG Regret scores for six large language models, with lower values indicating better performance.

experiments. The Hyperbolic Temporal Discounting environments shows an unexpected trends where more experiments actually increases error. This may again be related to how prior knowledge interferes with effective learning from data.

Models Improve with Scale. Larger models consistently outperform their smaller counterparts within the same model family. Both OpenThinker-32B and Qwen2.5-32B demonstrate significantly better performance than their respective 7B variants across environments (Fig. 3a), highlighting the benefits of scale for experimental design tasks.

Instruction-Tuned Models outperform Reasoning Models. Surprisingly, the instruction-tuned Qwen2.5 models outperform the reasoning-focused OpenThinker models (Fig. 3b). This may be because OpenThinker models are finetuned to perform well on a relatively narrow set of verifiable problems in math and code, while instruction-tuned models retain broader capabilities that could be useful for experimental design.

Models performance varies substantially across environments. Models show varying performance across different environments (Fig. 3b). Performance is strongest on environments like population growth dynamics and death process, where the LM agent achieves negative standardized error, indicating that the LM successfully leveraged information gained through experimentation. However, in environments like hyperbolic discounting, performance is low even after experimentation, suggesting that some domains are inherently more challenging for current models.

EIG Regret reveals relationship between experimental design and prediction. Our EIG regret analysis (Fig. 5b) provides insight into the relationship between two key components of scientific reasoning: designing informative experiments and making accurate predictions from collected data. GPT-4o achieves both the lowest EIG regret and strong predictive performance across several environments, suggesting these capabilities can be aligned. However, the varying performance of other models is informative — for instance, Qwen-32B shows higher EIG regret despite good predictive performance in some domains, indicating that while these abilities may be related, excellence in prediction doesn’t automatically translate to optimal experimental design.

LLMs cannot always optimally leverage statistical models. While Box’s Apprentice can propose and fit explicit statistical models to observed data, it does not consistently improve over the non-augmented LLM (GPT-4o) (Fig. 5a) From qualitative analysis of the models, we find that Box’s Apprentice tends to favor overly simple functional forms due to limited data, such as using linear approximations for inherently nonlinear phenomena.

4.2 Evaluating Model Discovery via Communication

Setup. Next, we evaluate the agents’ ability to build and communicate models that capture the underlying phenomena in each environment. To test this, we have the agents interact with the environment for 10 steps (scientist phase) and then generate a natural language explanation of their

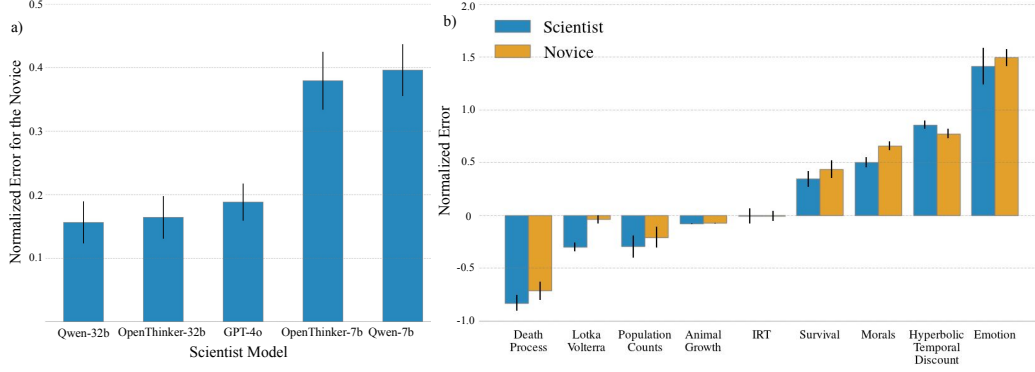


Figure 6: **Evaluation of Model Discovery via Communication.** (a) Comparison of the standardized error of the Novice (gpt-4o) with different Scientist models. (b) Comparison of errors made by the Novice and the Scientist (both models are gpt-4o). Error bars are standard error.

findings. We then provide this explanation to a *novice* agent, which must make predictions about the environment without any direct interaction (novice phase by using the explanation from the scientist; §3.2.2). The novice agent is always gpt-4o. The scientist’s prediction after 10 observations (Error After Experiments) acts as a weak positive control. Ideally, if the scientist’s explanation is effective, the novice’s error should approach the positive control.

Explanations improve with scale. Larger models generally produce more effective explanations, as evidenced by better novice performance when using explanations from 32B variants compared to 7B models (Fig. 6a). This suggests that increased model scale improves not just experimentation but also the ability to distill and communicate findings.

Explanations are not as good as experiments As expected, novice agents perform worse than scientists who directly interacted with the environment (Fig. 6b). The gap suggests that current explanation methods do not fully capture the knowledge gained through experimentation.

Explanations are more helpful for some environments. However, the effectiveness of explanations varies substantially across domains (Fig. 6b). For instance, explanations are helpful for animal growth, but struggle with complex domains like moral judgments. This variation likely reflects the complexity of different domains and the current limitations of language models in capturing and communicating certain types of patterns.

5 Discussion

We introduced *BoxingGym*, a benchmark measuring language-based agents’ capabilities in experimental design and model discovery across 10 real-world-based environments. We evaluated experimental design using information gain metrics and developed a novel model discovery metric based on an agent’s ability to explain its model to a novice agent. Our evaluation across multiple model scales (7B-32B parameters) shows that while larger and closed-source models generally perform better, fundamental challenges persist. Neither domain-specific prior knowledge nor statistical modeling capabilities consistently improved performance. Some environments yielded strong results with larger models, while others remained challenging for all approaches. *BoxingGym* has limitations: it uses pre-defined experimental paradigms rather than requiring design from scratch [14], ignores resource constraints, and covers limited scientific domains. Future work should address these limitations by incorporating experiment design from scratch, resource constraints, and more diverse fields. We could also expand the human behavior environments (Moral Machines, Emotions) with more sophisticated participant simulations [4, 1, 49, 39, 40]. While our experiments demonstrated potential for interfaces that augment language models’ scientific reasoning capabilities, future research should explore data visualization, strategic simulations [27], model validation, and web-based research strategies to enhance experimental guidance and discovery.

References

- [1] Gati V Aher, Rosa I Arriaga, and Adam Tauman Kalai. Using large language models to simulate multiple humans and replicate human subject studies. In *International Conference on Machine Learning*, pages 337–371. PMLR, 2023.
- [2] Microsoft Research AI4Science and Microsoft Azure Quantum. The impact of large language models on scientific discovery: a preliminary study using gpt-4, 2023.
- [3] Anthropic. Claude 3.7 sonnet. <https://www.anthropic.com>, 2024.
- [4] Lisa P Argyle, Ethan C Busby, Nancy Fulda, Joshua R Gubler, Christopher Rytting, and David Wingate. Out of one, many: Using language models to simulate human samples. *Political Analysis*, 31(3):337–351, 2023.
- [5] Edmond Awad, Sohan Dsouza, Richard Kim, Jonathan Schulz, Joseph Henrich, Azim Shariff, Jean-François Bonnefon, and Iyad Rahwan. The moral machine experiment. *Nature*, 563(7729): 59–64, 2018.
- [6] Josh C. Bongard and Hod Lipson. Automated reverse engineering of nonlinear dynamical systems. *Proceedings of the National Academy of Sciences*, 104:9943 – 9948, 2007.
- [7] G. E. P. Box and William G. Hunter. A Useful Method for Model-Building. *Technometrics*, 4: 301–318, 1962.
- [8] George E. P. Box. Sampling and Bayes’ Inference in Scientific Modelling and Robustness. *Journal of the Royal Statistical Society. Series A (General)*, 143(4):383–430, 1980. ISSN 00359238.
- [9] George EP Box. Science and statistics. *Journal of the American Statistical Association*, 71 (356):791–799, 1976.
- [10] Greg Brockman, Vicki Cheung, Ludwig Pettersson, Jonas Schneider, John Schulman, Jie Tang, and Wojciech Zaremba. Openai gym. *arXiv preprint arXiv:1606.01540*, 2016.
- [11] Pablo Samuel Castro, Nenad Tomasev, Ankit Anand, Navodita Sharma, Rishika Mohanta, Aparna Dev, Kuba Perlin, Siddhant Jain, Kyle Levin, Noémi Éltető, Will Dabney, Alexander Novikov, Glenn C Turner, Maria K Eckstein, Nathaniel D Daw, Kevin J Miller, and Kimberly L Stachenfeld. Discovering symbolic cognitive models from human and animal behavior. *bioRxiv*, 2025. doi: 10.1101/2025.02.05.636732.
- [12] Kathryn Chaloner and Isabella Verdinelli. Bayesian Experimental Design: A Review. *Statistical Science*, 10(3):273 – 304, 1995. doi: 10.1214/ss/1177009939. URL <https://doi.org/10.1214/ss/1177009939>.
- [13] David Roxbee Cox. *Analysis of survival data*. Chapman and Hall/CRC, 2018.
- [14] Michael Dennis, Natasha Jaques, Eugene Vinitsky, Alexandre Bayen, Stuart Russell, Andrew Critch, and Sergey Levine. Emergent complexity and zero-shot transfer via unsupervised environment design. *Advances in neural information processing systems*, 33:13049–13061, 2020.
- [15] David Duvenaud, James Lloyd, Roger Grosse, Joshua Tenenbaum, and Ghahramani Zoubin. Structure discovery in nonparametric regression through compositional kernel search. In Sanjoy Dasgupta and David McAllester, editors, *Proceedings of the 30th International Conference on Machine Learning*, volume 28 of *Proceedings of Machine Learning Research*, pages 1166–1174, Atlanta, Georgia, USA, 17–19 Jun 2013. PMLR.
- [16] Adam Foster, Martin Jankowiak, Elias Bingham, Paul Horsfall, Yee Whye Teh, Thomas Rainforth, and Noah Goodman. Variational bayesian optimal experimental design. In H. Wallach, H. Larochelle, A. Beygelzimer, F. d’Alché-Buc, E. Fox, and R. Garnett, editors, *Advances in Neural Information Processing Systems*, volume 32. Curran Associates, Inc., 2019.

- [17] Adam Foster, Desi R Ivanova, Ilyas Malik, and Tom Rainforth. Deep adaptive design: Amortizing sequential bayesian experimental design. In *Proceedings of the 38th International Conference on Machine Learning*, Proceedings of Machine Learning Research. PMLR, 18–24 Jul 2021.
- [18] Kanishk Gandhi, Dorsa Sadigh, and Noah D Goodman. Strategic reasoning with language models. *arXiv preprint arXiv:2305.19165*, 2023.
- [19] Kanishk Gandhi, Denise Lee, Gabriel Grand, Muxin Liu, Winson Cheng, Archit Sharma, and Noah D Goodman. Stream of search (sos): Learning to search in language. *arXiv preprint arXiv:2404.03683*, 2024.
- [20] Kanishk Gandhi, Ayush Chakravarthy, Anikait Singh, Nathan Lile, and Noah D Goodman. Cognitive behaviors that enable self-improving reasoners, or, four habits of highly effective stars. *arXiv preprint arXiv:2503.01307*, 2025.
- [21] Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Ruoyu Zhang, Runxin Xu, Qihao Zhu, Shirong Ma, Peiyi Wang, Xiao Bi, et al. Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning. *arXiv preprint arXiv:2501.12948*, 2025.
- [22] Alex Havrilla, Yuqing Du, Sharath Chandra Raparthy, Christoforos Nalmpantis, Jane Dwivedi-Yu, Maksym Zhuravinskyi, Eric Hambro, Sainbayer Sukhbaatar, and Roberta Raileanu. Teaching large language models to reason with reinforcement learning. *arXiv preprint arXiv:2403.04642*, 2024.
- [23] Aaron Jaech, Adam Kalai, Adam Lerer, Adam Richardson, Ahmed El-Kishky, Aiden Low, Alec Helyar, Aleksander Madry, Alex Beutel, Alex Carney, et al. Openai o1 system card. *arXiv preprint arXiv:2412.16720*, 2024.
- [24] Carlos E Jimenez, John Yang, Alexander Wettig, Shunyu Yao, Kexin Pei, Ofir Press, and Karthik R Narasimhan. SWE-bench: Can language models resolve real-world github issues? In *The Twelfth International Conference on Learning Representations*, 2024. URL <https://openreview.net/forum?id=VTF8yNQM66>.
- [25] Lucas Lehnert, Sainbayer Sukhbaatar, DiJia Su, Qinqing Zheng, Paul Mccvay, Michael Rabbat, and Yuandong Tian. Beyond a*: Better planning with transformers via search dynamics bootstrapping. *arXiv preprint arXiv:2402.14083*, 2024.
- [26] Michael Y Li, Emily B Fox, and Noah D Goodman. Automated Statistical Model Discovery with Language Models. In *International Conference on Machine Learning (ICML)*, 2024.
- [27] Michael Y. Li, Vivek Vajipey, Noah D. Goodman, and Emily B. Fox. Critical: Critic automation with language models, 2024. URL <https://arxiv.org/abs/2411.06590>.
- [28] Chris Lu, Cong Lu, Robert Tjarko Lange, Jakob Foerster, Jeff Clune, and David Ha. The ai scientist: Towards fully automated open-ended scientific discovery. *arXiv preprint arXiv:2408.06292*, 2024.
- [29] Måns Magnusson, Paul Bürkner, and Aki Vehtari. posteriordb: a set of posteriors for Bayesian inference and probabilistic programming, October 2023.
- [30] John McCarthy, Marvin L Minsky, Nathaniel Rochester, and Claude E Shannon. A proposal for the dartmouth summer research project on artificial intelligence, august 31, 1955. *AI magazine*, 27(4):12–12, 1955.
- [31] B. A. McKinney, J. E. Crowe, H. U. Voss, P. S. Crooke, N. Barney, and J. H. Moore. Hybrid grammar-based approach to nonlinear dynamical system identification from biological time series. *Phys. Rev. E*, 73:021912, Feb 2006. doi: 10.1103/PhysRevE.73.021912.
- [32] Sewon Min, Mike Lewis, Luke Zettlemoyer, and Hannaneh Hajishirzi. Metaicl: Learning to learn in context. *arXiv preprint arXiv:2110.15943*, 2021.
- [33] Volodymyr Mnih, Koray Kavukcuoglu, David Silver, Alex Graves, Ioannis Antonoglou, Daan Wierstra, and Martin Riedmiller. Playing atari with deep reinforcement learning. *arXiv preprint arXiv:1312.5602*, 2013.

- [34] Jay I. Myung, Daniel R. Cavagnaro, and Mark A. Pitt. A tutorial on adaptive design optimization. *Journal of Mathematical Psychology*, 57(3):53–67, 2013. ISSN 0022-2496. doi: <https://doi.org/10.1016/j.jmp.2013.05.005>. URL <https://www.sciencedirect.com/science/article/pii/S0022249613000503>.
- [35] Deepak Nathani, Lovish Madaan, Nicholas Roberts, Nikolay Bashlykov, Ajay Menon, Vincent Moens, Amar Budhiraja, Despoina Magka, Vladislav Vorotilov, Gaurav Chaurasia, et al. Mlgym: A new framework and benchmark for advancing ai research agents. *arXiv preprint arXiv:2502.14499*, 2025.
- [36] Allen Nie, Yi Su, Bo Chang, Jonathan N Lee, Ed H Chi, Quoc V Le, and Minmin Chen. Evolve: Evaluating and optimizing llms for exploration. *arXiv preprint arXiv:2410.06238*, 2024.
- [37] Desmond C Ong, Jamil Zaki, and Noah D Goodman. Affective cognition: Exploring lay theories of emotion. *Cognition*, 143:141–162, 2015.
- [38] OpenAI. Hello, GPT-4. <https://openai.com/index/hello-gpt-4o/>, 2024. Accessed: 2024-06-04.
- [39] Joon Sung Park, Lindsay Popowski, Carrie Cai, Meredith Ringel Morris, Percy Liang, and Michael S Bernstein. Social simulacra: Creating populated prototypes for social computing systems. In *Proceedings of the 35th Annual ACM Symposium on User Interface Software and Technology*, pages 1–18, 2022.
- [40] Joon Sung Park, Joseph O’Brien, Carrie Jun Cai, Meredith Ringel Morris, Percy Liang, and Michael S Bernstein. Generative agents: Interactive simulacra of human behavior. In *Proceedings of the 36th Annual ACM Symposium on User Interface Software and Technology*, pages 1–22, 2023.
- [41] Gabriel Poesia, Kanishk Gandhi, Eric Zelikman, and Noah D Goodman. Certified deductive reasoning with language models. *arXiv preprint arXiv:2306.04031*, 2023.
- [42] 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. In *The Twelfth International Conference on Learning Representations*, 2024.
- [43] Tom Rainforth, Robert Cornish, Hongseok Yang, Andrew Warrington, and Frank Wood. On Nesting Monte Carlo Estimators. *International Conference on Machine Learning (ICML)*, 2018.
- [44] Georg Rasch. *Probabilistic models for some intelligence and attainment tests*. ERIC, 1993.
- [45] Abulhair Saparov and He He. Language models are greedy reasoners: A systematic formal analysis of chain-of-thought. *arXiv preprint arXiv:2210.01240*, 2022.
- [46] Abulhair Saparov, Richard Yuanzhe Pang, Vishakh Padmakumar, Nitish Joshi, Mehran Kazemi, Najoung Kim, and He He. Testing the general deductive reasoning capacity of large language models using ood examples. *Advances in Neural Information Processing Systems*, 36, 2024.
- [47] John Schultz, Jakub Adamek, Matej Jusup, Marc Lanctot, Michael Kaisers, Sarah Perrin, Daniel Hennes, Jeremy Shar, Cannada Lewis, Anian Ruoss, et al. Mastering board games by external and internal planning with language models. *arXiv preprint arXiv:2412.12119*, 2024.
- [48] Ben Shababo, Brooks Paige, Ari Pakman, and Liam Paninski. Bayesian inference and online experimental design for mapping neural microcircuits. In C.J. Burges, L. Bottou, M. Welling, Z. Ghahramani, and K.Q. Weinberger, editors, *Advances in Neural Information Processing Systems*, volume 26. Curran Associates, Inc., 2013. URL https://proceedings.neurips.cc/paper_files/paper/2013/file/17c276c8e723eb46aef576537e9d56d0-Paper.pdf.
- [49] Omar Shaikh, Valentino Chai, Michele J Gelfand, Diyi Yang, and Michael S Bernstein. Rehearsal: Simulating conflict to teach conflict resolution. *arXiv preprint arXiv:2309.12309*, 2023.
- [50] Open Thoughts Team. Open Thoughts, January 2025.

- 476 [51] Vito Volterra. Variations and fluctuations of the number of individuals in animal species living
477 together. *ICES Journal of Marine Science*, 3(1):3–51, 1928.
- 478 [52] Ruocheng Wang, Eric Zelikman, Gabriel Poesia, Yewen Pu, Nick Haber, and Noah D. Goodman.
479 Hypothesis search: Inductive reasoning with language models. In *The Twelfth International*
480 *Conference on Learning Representations*, 2024.
- 481 [53] Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Fei Xia, Ed Chi, Quoc V Le,
482 Denny Zhou, et al. Chain-of-thought prompting elicits reasoning in large language models.
483 *Advances in neural information processing systems*, 35:24824–24837, 2022.
- 484 [54] An Yang, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chengyuan
485 Li, Dayiheng Liu, Fei Huang, Haoran Wei, Huan Lin, Jian Yang, Jianhong Tu, Jianwei Zhang,
486 Jianxin Yang, Jiaxi Yang, Jingren Zhou, Junyang Lin, Kai Dang, Keming Lu, Keqin Bao, Kexin
487 Yang, Le Yu, Mei Li, Mingfeng Xue, Pei Zhang, Qin Zhu, Rui Men, Runji Lin, Tianhao Li,
488 Tingyu Xia, Xingzhang Ren, Xuancheng Ren, Yang Fan, Yang Su, Yichang Zhang, Yu Wan,
489 Yuqiong Liu, Zeyu Cui, Zhenru Zhang, and Zihan Qiu. Qwen2.5 technical report. *arXiv preprint*
490 *arXiv:2412.15115*, 2024.

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: We describe the design of our benchmark accurately, summarize results with different models.

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: See discussion.

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: No proofs or new theoretical result.

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: Yes, further, all our code, results and scripts are available on github.

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: All the code is accessible on the github.

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: We describe this in detail in experimental setup and have the full specification 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: We report statistical significance in all our results...

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: See appendix section B.

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: Single blind submission and we follow the code.

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: [No]

Justification: We don't discuss these as there are no direct negative societal impacts.

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: Not relevant for the paper.

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: All models have been cited appropriately. The papers that inspired the environments have been credited too.

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: [\[Yes\]](#)

Justification: We add documentation to the BoxingGym code.

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: No human participants were recruited.

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: Paper does not use human participants.

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

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

805 Answer: [NA]

806 Justification: None of the core methods used LLMs.

807 Guidelines:

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