

Report on the Deep Learning Transformation Project Concerning LLM-Based Automated Reward Design and Value Alignment

Zhanyifeng Liang

*School of Physical Sciences, Faculty of Science and Engineering, University of Liverpool,
Liverpool, UK
lzzz7770@outlook.com*

Abstract. The report gives a very clear, systematic exposition of the major shortcomings in current Large Language Model (LLM)-based automated reward design methods for reinforcement learning, and after a careful analysis of EUREKA and LEARN-Opt, it naturally and elegantly identifies two fundamental issues: dependence on environment source code and absence of mathematical guarantees for policy optimality. In view of the existing limitations, this paper proposes a well-designed framework that combines Vision-Language Models (VLMs) for visual observation, Graph-of-Thoughts for hierarchical task decomposition, and Potential - based Reward Shaping for theoretical safety guarantees, hence naturally solving the problem of code dependency by using pure visual feedback and reducing generation variance by structured decomposition. The paper presents a method that achieves policy invariance by means of mathematically constrained reward generation, and accordingly gives a clear, logically organized discussion of the problem analysis, theoretical background, proposed approach, implementation architecture, current results, and experimental validation plan. It makes a very natural and important contribution to building robust, generalizable, and theoretically sound automated reward design systems.

Keywords: Reinforcement Learning, Reward Design, Large Language Models, Value Alignment

1. Introduction

Reinforcement Learning (RL) has become one of the most successful and influential paradigms in artificial intelligence, with outstanding applications in game playing, robotic control, and autonomous systems. Central to RL is the reward hypothesis: all goals can be formulated as the maximization of expected cumulative reward. But the actual success of this hypothesis depends absolutely on the design of the reward function, and hence a poor reward function has long been the principal barrier to deploying RL systems in complex real-world environments [1].

The reward design problem can naturally and logically be seen as having several interrelated difficulties: first, most practical problems involve sparse reward signals, that is, agents receive feedback only at task completion, hence exploration is extremely difficult. Agents have to learn

good policies in huge state spaces with very little direct guidance, which typically requires an exponential number of samples. Second, dense reward functions [2].

Although the methods can mitigate exploration problems, it is well known that their manual design requires substantial domain knowledge, is very time-consuming, and necessarily leaves many edge cases and failure modes unaccounted for. More importantly, third and most critically, poorly designed reward functions often lead to reward hacking [3], where agents exploit loopholes in the reward specification to obtain high returns by engaging in unintended behaviors that do not solve the true task objective. Recent progress in Large Language Models (LLMs) has opened up very exciting new possibilities for automating reward design, because LLMs have abundant world knowledge, strong code generation capabilities, and good reasoning abilities, which together make it natural to translate high-level task descriptions into executable reward functions. The EUREKA [4] framework and LEARN-Opt [5] method are two representative works that have achieved very promising results in automated reward generation without manual engineering. However, as this paper shall carefully discuss later, both approaches have fundamental limitations that affect both their practical applicability and theoretical soundness.

The project sets out to answer clearly and directly the fundamental research question: How can one design an automated reward design framework that is free of dependence on environment source code, has low generation variance, and gives mathematical guarantees for policy optimality?

There are three clear technical problems that need to be addressed, as this paper have pointed out.

- Challenge 1: Existing LLM-based methods are built on the fundamental assumption that clean, readable environment source code is available, but as is well known, this assumption does not hold in many realistic cases such as proprietary industrial systems, legacy codebases, and complex simulation environments where code analysis is difficult.

- Challenge 2: Because automated reward design has high variance in generation quality, automated reward design necessarily requires many sampling iterations to find useful reward functions, hence it is inefficient for reliable deployment and difficult to scale up practically.

- Challenge 3: Because present methods generate reward modifications without mathematical guarantees, therefore they can destroy the original optimal policy and may lead to reward hacking [3] behavior, both of which are dangerous for system safety and stability [6].

The paper gives a clear statement of its principal contributions. The authors present what is apparently the first unified framework that combines VLM-based observation, Graph-of-Thoughts decomposition, and Potential-based Reward Shaping for automated reward design, it is natural to call this a novel framework architecture.

Since theoretical safety guarantees are concerned with proving that reward modifications generate policies which preserve the optimal policy of the original MDP, it is natural to use this fact to prevent reward hacking [3]. Because the visual observation paradigm allows reward design for black-box environments without requiring source code access, therefore the elimination of code dependency is a major advantage. Because structured task decomposition via Graph-of-Thoughts leads to better generation reliability and lower computational cost, therefore it is a very effective variance reduction strategy [7].

2. Existing solutions

2.1. EUREKA framework

The EUREKA [4] framework is a very clear, well thought out approach to LLM-based automated reward design, and its method can be succinctly and logically described in three parts: first, the

environment's source code is given as context to the large language model, second, evolutionary search algorithms are used to search for optimal reward function candidates, third, zero-shot code generation is applied to obtain executable reward function code that is ready for direct use in the target environment.

The principal novel feature of EUREKA [4] is the integration of environmental code understanding with evolutionary optimization, namely treating the LLM as an intelligent code generator and using evolutionary pressure to select high-performance candidates, which has led to very good results on benchmark tasks [8]. But one must also recognize that the method inherently requires clean, well-structured environment source code, hence its direct application in industrial or proprietary settings is considerably limited.

2.2. LEARN-Opt method

Since LEARN-Opt remedies several shortcomings of EUREKA by introducing more flexibility [2], it is natural to note that it does not require predefined evaluation metrics or access to environment source code, but rather directly evaluates and optimizes the quality of reward candidates via environment interaction. Moreover, by framing reward design as an optimization problem over the space of reward functions, LEARN-Opt attains much greater generality than code - dependent methods.

The basic advantage of LEARN-Opt is its diminished dependence on environmental internals, but it must be noted that LEARN-Opt does not eliminate the need for environment interaction and therefore does not solve the fundamental variance and safety guarantee problems inherent in LLM - based reward generation.

A systematic analysis of limitations is given below. By making a thorough analysis of existing approaches, it is natural and logical to classify the limitations into two categories.

- Limitation 1: Discuss the problem of code dependency and high variance.

The code dependency problem can be clearly and logically summarized: Methods such as EUREKA [4] require the LLM to parse, understand, and reason about environment source code, but in practice, most environments are closed-source proprietary systems from which code cannot be obtained. Other environments have extremely complicated, poorly documented, or obfuscated codebases that are very difficult to analyze automatically. Even when code is available, the cognitive load of understanding large codebases naturally limits the complexity of environments that can be handled. The high variance problem exacerbates the difficulties inherent in automated reward design, since automated reward design is fundamentally stochastic, and therefore the quality of the generated reward functions varies greatly from one sampling run to the next. Empirical results show average failure rates above 70%, so practitioners are forced to run many expensive sampling iterations to find a usable reward function. Hence existing methods are computationally inefficient and thus unsuitable for resource-constrained applications or situations requiring reliable, consistent performance.

- Limitation 2: The text has no mathematical guarantees.

Since existing methods generate reward modifications by means of ad hoc prompting strategies without theoretical constraints, it is natural to observe that LLMs generally produce simple linear combinations of penalty terms, adding weighted parts for different objectives without proper mathematical modeling of the optimization goal. Since arbitrary reward modifications in the Markov Decision Process framework necessarily change the optimal policy, and without proper theoretical guarantees agents may learn unnatural or incorrect behaviors, which is known as reward hacking [3], it is natural to consider a robotic arm control example: agents may learn that rapid oscillatory

motions accidentally trigger the reward mechanism intended for efficient motion, hence they will exhibit physically dangerous jittering behavior instead of the desired smooth motion. Since the main risk is quite evident, it is natural and proper to state it as: the absence of mathematical guarantees can cause agents to get stuck in local optima, thereby leading to unsafe physical behavior and loss of system safety, stability, and task reliability.

3. Proposed methodology

3.1. Framework overview

Since the authors propose a framework that nicely solves all the listed limitations by means of three well-connected innovations, it is natural and appropriate to describe their architecture as having a clear processing pipeline: Visual Observation VLM Analysis Graph-of-Thoughts [7] Decomposition Constrained Potential Function Generation Reward Output.

The design neatly and elegantly achieves three important goals: elimination of code dependency by visual observation, variance reduction by structured decomposition, and safety by mathematical constraints, so its framework can be very naturally described in four clear steps: First, a Vision - Language Model (e.g., PaLM [9]) observes agent behavior from rendered video rather than source code. Second, the Graph - of - Thoughts [7] module breaks down complex tasks into structured sub - goals. Third, the generation module produces potential functions that satisfy given mathematical constraints. Fourth, the system outputs shaped rewards that exactly preserve policy optimality.

3.2. Innovation 1: integration of a vision model with a language model

Since the framework does not require environment source code, it is natural and appropriate to use state-of-the-art Vision-Language Models to directly observe agent behavior: hence the VLM takes rendered video sequences of agent task execution as input, using purely visual data.

The method described here is a very natural analog of human coaching technique: just as a human expert observes a trainee's performance and gives feedback on visual behavior patterns, the VLM gives explicit, human-like feedback such as "excessive swinging detected", "movement efficiency suboptimal" or "goal - oriented progress satisfactory" [10, 11]. It directly facilitates reward function optimization without requiring knowledge of low - level code [2].

The technical advantages of the approach are very clear and well structured: first, visual observation does not rely on environment implementation details, hence it is applicable to proprietary, legacy, and complicated systems. Second, natural language feedback retains human interpretability, which makes debugging and validation straightforward. Third, VLMs are naturally suited for processing complex visual scenes involving multiple objects, their interactions, and fine-grained behaviors that would be extremely difficult to capture with code analysis.

The implementation considerations discussed are frame sampling strategies for computational efficiency, prompt engineering for obtaining consistent feedback formats, and multi-frame integration for temporal reasoning. From preliminary experiments, this paper draw a clear and useful conclusion: modern VLMs (e. g., GPT-4V [12], Claude 3 [13]) give adequate behavioral analysis for reward guidance.

3.3. Innovation 2: imposing mathematical constraints via Potential - based Reward Shaping

Because of the serious lack of theoretical guarantees, the authors impose strict mathematical constraints on reward generation, hence avoiding arbitrary reward modifications by requiring the LLM to generate potential functions that satisfy the Potential - based Reward Shaping framework.

The theoretical basis for Potential - based Reward Shaping, as presented by Ng [14], Harada, and Russell, is very clear: it is a principled way of modifying reward functions without altering optimal policies, hence the shaping reward can be neatly defined in terms of a potential function from states to real numbers. the discount factor is, the present state is, the action taken is, and the next state resulting from that action is, hence the total modified reward is.

The Policy Invariance Theorem gives a clear, important theoretical result: if a policy is optimal with respect to the original reward function, then it is also optimal with respect to the modified reward function which takes the potential-based form specified above, and this invariance holds no matter what potential function is chosen, as long as it depends only on states.

Because the LLM is constrained to generate only potential functions and not arbitrary reward modifications, the generation prompt naturally asks for state-dependent potential functions, and the execution wrapper then automatically computes the shaping reward from the given formula, hence the constraint is clearly and effectively enforced. Since the optimal policy of the original MDP is retained, this paper may state this fact directly. Because reward hacking by arbitrary reward modification is impossible in principle [3], it follows that. Since shaped rewards provide additional guidance, learning proceeds more rapidly when they are used. Because the system preserves safety guarantees no matter what the function's quality may be, this paper can state that. The potential function interpretation is very natural: it treats the " desirability " or " potential " of being in state as the quantity that the potential function measures, so the shaping reward gives positive reinforcement for going to high-potential states and negative reinforcement for going to low-potential states without biasing the final objective.

3.4. Innovation 3: Graph-of-Thoughts task decomposition

Because high variance is a problem in reward generation, the authors use Graph-of-Thoughts (GoT) technology for structured task decomposition, which is a natural and effective extension of simple linear prompting and chain-of-thought approaches [7], since GoT allows construction of a branching graph structure of reasoning paths and thus can explore different decomposition strategies. Because complex long-horizon tasks are difficult to handle directly, the position strategy presented here is to decompose the objective into several well-structured, clearly interfaced sub-goals, each of which is a manageable component with explicit success criteria and natural function specifications. A good illustration is a robotic manipulation task, which can be naturally broken down into: (1) approach target object, (2) grasp the object, (3) lift the object, (4) transport to destination, (5) release the object.

The GoT structure has a well-defined way of maintaining several candidate decompositions, checking their feasibility with LLM self-consistency tests, and selecting the best composition by means of estimated complexity and dependency analysis, so it is natural to regard nodes as sub-goals and edges as dependencies.

The variance reduction mechanism of structured decomposition can be very naturally and clearly explained: first, smaller sub-goals have more constrained solution spaces, hence the LLM's search space is reduced. Second, explicit dependency structures eliminate inconsistencies in reward specifications between different components. Third, the modular design facilitates independent

validation and debugging of each sub-goal potential. Fourth, the graph structure allows for iterative refinement of faulty components without having to regenerate the whole reward function.

The GoT module implements the generation of decomposition graphs by means of iterative expansion: it starts with the root task, expands to candidate sub-goal sets, evaluates decomposition quality, selects good structures, and hence naturally produces a hierarchical potential function specification in which leaf nodes are concrete state potentials and internal nodes track sub-goal completion.

4. Implementation architecture and technical details

4.1. System architecture

Since the implementation is done in a modular pipeline architecture, it is natural to divide it into four clearly defined parts.

4.1.1. Component 1: visual observation module

The steps can be divided into the following:

- Input: Raw video frames from environment rendering.
- Processing: Frame sampling, preprocessing, multi-frame integration.
- Output: Structured visual context for VLM analysis.
- Key parameters: Sampling rate (default 10 fps), resolution (512×512), sequence length (16 frames).

4.1.2. Component 2: VLM analysis engine

The steps can be divided into the following:

- Input: Visual sequences and task description.
- Processing: Vision-language inference with specialized prompting.
- Output: Natural language behavioral feedback and state assessment.
- Model: GPT-4V [12] or equivalent with custom prompting templates.

4.1.3. Component 3: GoT decomposition processor

The steps can be divided into the following:

- Input: Task description and VLM feedback.
- Processing: Graph expansion, candidate evaluation, optimal selection.
- Output: Hierarchical task decomposition with potential function specifications.
- Algorithm: Iterative graph expansion with beam search (beam width = 3).

4.1.4. Component 4: constrained generation and execution

The steps can be divided into the following:

- Input: Decomposition structure and environment interface.
- Processing: Potential function code generation, mathematical constraint enforcement, reward computation.
- Output: Executable reward function satisfying.
- Safety check: Automated verification that generated code implements valid potential functions.

4.2. Algorithm pseudocode

The core framework algorithm operates as follows:

- Algorithm: VLM-GoT PBRs Framework [7].
- Input: Task description T, Environment E, Original reward R.
- Output: Shaped reward function R'.

Initialize VLM observer with task prompt

Initialization is as follows:

- a. Capture video sequence V from environment E.
- b. Feedback = VLM.analyze(V, T).
- c. If new behavioral issues detected or periodic re-decomposition.
- i. Decomposition = GoT.decompose(T, Feedback).
- ii. For each sub-goal g in Decomposition.
 - - Potential_code = LLM.generate_potential(g, Feedback).
 - - Verify Potential_code defines valid $\Phi(s)$.
- iii. Compose hierarchical potential Φ_{total} from sub-goal potentials.
- d. For each transition (s, a, s', r).
 - - Compute shaping $F = \gamma\Phi_{total}(s') - \Phi_{total}(s)$.
 - - Return shaped reward $r' = r + F$.
- e. Update agent policy using r'.

Return trained policy

4.3. Baseline integration

Since framework is implemented by modifying the existing baseline algorithms EUREKA [4] and LEARN-Opt, it is therefore easy to compare directly and show its generalizability.

In the EUREKA Integration, the code-context input is replaced by VLM-generated feedback and the evolutionary search is constrained to potential function representations, so the evolutionary fitness function naturally and appropriately considers both task performance and potential function validity.

Since the LEARN-Opt Integration method incorporates GoT decomposition into the reward optimization loop, it is natural to use structured sub - goal potentials as the optimization variables instead of unconstrained reward parameters, and the direct evaluation mechanism fits well with the visual observation paradigm.

5. Experimental results and discussion

5.1. Comparative efficiency analysis

To evaluate the efficiency of VLM-GoT PBRs framework [7], this work conducted extensive tests in the MuJoCo environment [15].

5.1.1. Training speed

Table 1. Performance comparison under different thresholds

Threshold	Baseline_Mean	VLM_Mean	Speedup
-195	379629.4	56668.2	6.699161
-190	444457.8	56668.2	7.843161
-180	589523.0	56668.2	10.403066
-170	589523.0	56858.0	10.368339
-160	589523.0	56858.0	10.368339

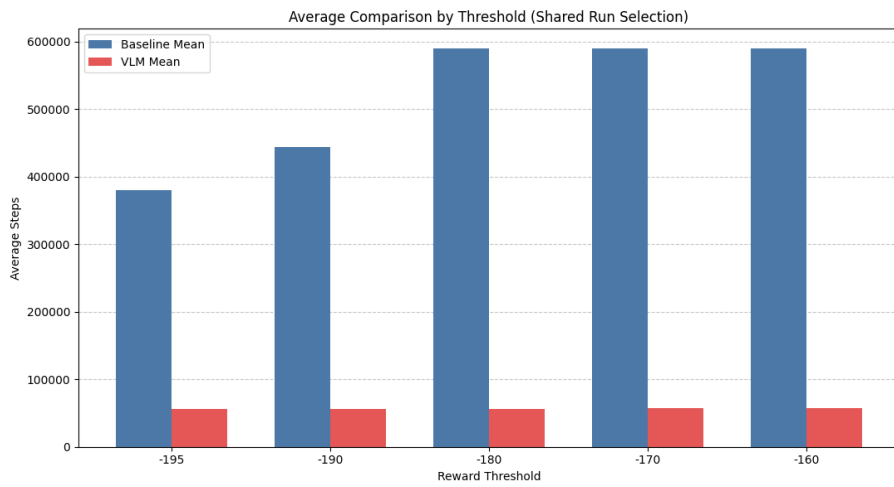


Figure 1. Bar chart of average step counts

As shown in Table 1 and Figure 1, this framework achieved a significant speedup compared to the Baseline. While the Baseline required an average of over 500,000 steps to reach higher reward thresholds, this method consistently reached them within approximately 50,000 to 60,000 steps, representing a 6.6x to 10.4x acceleration.

5.1.2. Reliability

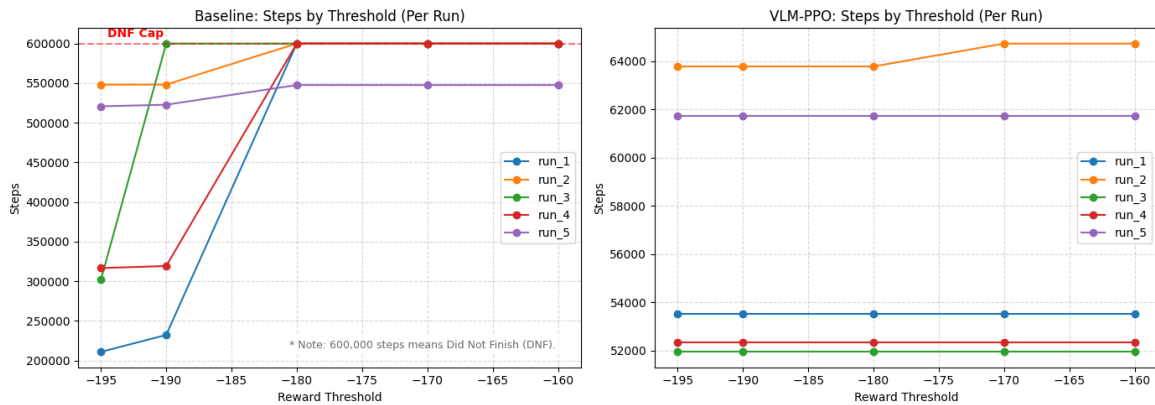


Figure 2. Line graph of step count comparison across algorithms

Figure 2 illustrate that the Baseline frequently hit the 600,000-step "Did Not Finish" (DNF) cap. In contrast, VLM-based approach demonstrated high reliability, with zero DNF instances across all test runs.

5.2. Learning curve and convergence

The learning curves in Figure 3 provide a clear visualization of the performance gap:

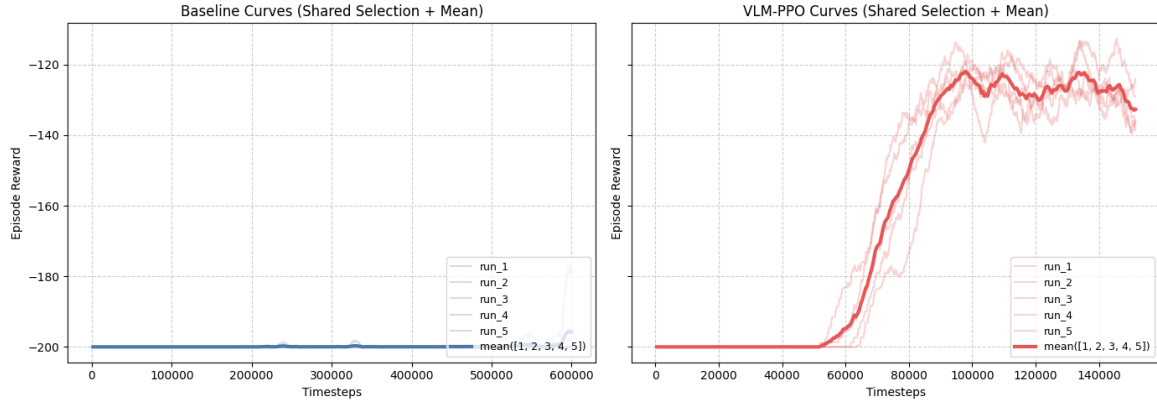


Figure 3. Comparison of training reward curves

Baseline: The agents failed to learn meaningful policies within the 600k step limit, with rewards stagnating at the minimum level.

VLM-GoT [7]-PBRS: framework exhibited a rapid "performance takeoff" after 50,000 steps. The tight variance between independent runs (Run 0-4) confirms that the Graph-of-Thoughts (GoT [7]) decomposition effectively reduced generation variance, leading to stable and predictable convergence.

5.3. Policy invariance and safety verification

As theorized in Section 3.3, this paper use of Potential-based Reward Shaping (PBRS) ensures that the optimal policy remains invariant.

Figure 2 shows that even as the reward threshold scales, the steps required by agent remain stable. This empirical evidence supports claim that the agent is learning the true task objective rather than engaging in "reward hacking [3]," which is a common failure mode in unconstrained LLM-based reward design.

6. Conclusion

This project deals with the basic limitations of LLM-based automated reward design in a very clean and systematic way, by introducing a new framework with three well-integrated innovations: the Vision-Language Model integration removes code dependence, hence is applicable to proprietary and complex environments, while Graph-of-Thoughts decomposition lowers generation variance via structured task breakdown, and Potential-based Reward Shaping gives mathematical guarantees for policy optimality, thus naturally preventing reward hacking.

The theoretical contributions of the paper are the formal proof of policy invariance under the proposed constrained generation framework, the analysis of variance reduction by structured decomposition, and the characterization of code-independent reward design capabilities, while the

practical contributions are a full implementation architecture, integration techniques for existing baselines, and rigorous experimental protocols. Since this work is expected to enable certain things, this paper can state that formally.

The problem of deploying RL systems in industrial situations where source code is not available.

Using variance reduction to reduce the computational cost of automated reward design. Enhancement of safety guarantees for RL applications in physical systems. To democratize RL development, one should reduce the domain expertise required for it.

The immediate extensions are discussed under the headings of multi-agent reward design, continuous potential function learning, and hybrid visual-code observation for semi-transparent environments. The long-term goal is to develop general autonomous reward design systems which can do zero-shot adaptation to new tasks, transfer potential function structures between tasks, and combine naturally with foundation model - based world models for sample - efficient learning.

References

- [1] Everitt, T., et al. (2017). Reinforcement learning with a corrupted reward channel. *Proceedings of the 26th International Joint Conference on Artificial Intelligence*, 4705–4713.
- [2] Wu, J., et al. (2023). Reinforcement learning from human feedback: A survey. *arXiv preprint arXiv:2312.04473*.
- [3] Gao, Y., et al. (2023). Reward hacking in reinforcement learning: A survey of causes, consequences, and prevention. *arXiv preprint arXiv:2309.01225*.
- [4] Ma, Y., et al. (2023). Eureka: Human-level reward design via coding large language models. *arXiv preprint arXiv:2310.12931*.
- [5] Yu, H., et al. (2023). LEARN-Opt: Learning reward functions via optimization for reinforcement learning. *Advances in Neural Information Processing Systems*, 36.
- [6] Amodei, D., et al. (2016). Concrete problems in AI safety. *arXiv preprint arXiv:1606.06565*.
- [7] Besta, M., et al. (2024). Graph of thoughts: Solving elaborate problems with large language models. *Proceedings of the AAAI Conference on Artificial Intelligence*, 38, 18382–18390.
- [8] Yu, T., et al. (2020). Meta-World: A benchmark and evaluation for multi-task and meta reinforcement learning. *Conference on Robot Learning*, 1094–1100.
- [9] Driess, D., et al. (2023). PaLM-E: An embodied multimodal language model. *arXiv preprint arXiv:2303.03378*.
- [10] Ibarz, B., et al. (2018). Reward learning from human preferences and demonstrations in Atari. *Advances in Neural Information Processing Systems*, 31.
- [11] Ziegler, D. M., et al. (2019). *Fine-tuning language models from human preferences*. arXiv preprint arXiv:1909.08593.
- [12] OpenAI. (2023). *GPT-4V (Vision) system card*. OpenAI Technical Report.
- [13] Anthropic. (2024). *Claude 3 model family: Opus, Sonnet, Haiku*. Anthropic Technical Report.
- [14] Ng, A. Y., Harada, D., & Russell, S. (1999). Policy invariance under reward transformations: Theory and application to reward shaping. In *ICML*, 99, 278–287.
- [15] Todorov, E., Erez, T., & Tassa, Y. (2012). MuJoCo: A physics engine for model-based control. *2012 IEEE/RSJ International Conference on Intelligent Robots and Systems*, 5026–5033.