
Text-to-Decision Agent: Offline Meta-Reinforcement Learning from Natural Language Supervision

Shilin Zhang^{1*} Zican Hu^{1*} Wenhao Wu^{1*} Xinyi Xie¹ Jianxiang Tang¹
Chunlin Chen¹ Daoyi Dong² Yu Cheng³ Zhenhong Sun^{4,5✉} Zhi Wang^{1✉}

¹ Nanjing University ² University of Technology Sydney ³ The Chinese University of Hong Kong

⁴ Australian National University ⁵ University of New South Wales

{shilinzhang, zicanhu, wenhaowu, xinyixie, jianxiangtang}@smail.nju.edu.cn
{clchen, zhiwang}@nju.edu.cn daoyi.dong@uts.edu.au
chengyu@cse.cuhk.edu.hk zhenhong.sun@anu.edu.au

Abstract

Offline meta-RL usually tackles generalization by inferring task beliefs from high-quality samples or warmup explorations. The restricted form limits their generality and usability since these supervision signals are expensive and even infeasible to acquire in advance for unseen tasks. Learning directly from the raw text about decision tasks is a promising alternative to leverage a much broader source of supervision. In the paper, we propose **Text-to-Decision Agent (T2DA)**, a simple and scalable framework that supervises offline meta-RL with natural language. We first introduce a generalized world model to encode multi-task decision data into a dynamics-aware embedding space. Then, inspired by CLIP, we predict which textual description goes with which decision embedding, effectively bridging their semantic gap via contrastive language-decision pre-training and aligning the text embeddings to comprehend the environment dynamics. After training the text-conditioned generalist policy, the agent can directly realize zero-shot text-to-decision generation in response to language instructions. Comprehensive experiments on MuJoCo and Meta-World benchmarks show that T2DA facilitates high-capacity zero-shot generalization and outperforms various types of baselines. Our code is available at <https://github.com/NJU-RL/T2DA>.

1 Introduction

Reinforcement learning (RL) has emerged as an effective mechanism for training autonomous agents to perform complex tasks in interactive environments [1, 2], unleashing its potential across frontier problems including preference optimization [3], diffusion model training [4], and reasoning [5, 6] such as in OpenAI o1 [7] and DeepSeek-R1 [8]. Despite its achievements, one of the grand challenges is *generalization*: building a general-purpose agent capable of handling multiple tasks in response to diverse user commands [9]. Due to distribution shift and the lack of self-supervised pre-training techniques, RL agents typically struggle with poor generalization to unseen tasks [10]. Offline meta-RL tackles generalization via training on a distribution of offline tasks [11]. However, they usually rely on high-quality samples [12] or warmup explorations [13] to infer task beliefs for generalization at test time, while these supervision signals are expensive and even infeasible to acquire in advance for unseen tasks [14]. *This inspires us to explore a much broader source of supervision.*

Recently, large language models (LLMs), pre-trained on extensive text corpora, can encode a wealth of semantic knowledge with remarkable representation power and transferability [15]. The typical “text-

*Equal contributions. ✉Corresponding authors.

to-text" interface enables task-agnostic architectures to achieve broad generalization with minimal or no domain-specific data [16]. The availability of large-scale text collections for aggregate supervision during model pre-training has revolutionized the language and multi-modal communities in recent years [17]. Naturally, the above advancements highlight a promising question: *Could scalable pre-training methods that learn perception from the supervision embedded in natural language lead to a similar leap forward in developing generalist agents for offline meta-RL?*

RL typically learns from active interactions with the outer world tied to specific environment dynamics, with unique decision-level representation distinct from the unbounded perception-level representation of natural language [18]. Recent studies facilitate language grounding to decision domains from various aspects, such as LLMs as policies [19, 20, 21], LLMs as rewards [22, 23, 24], and language-conditioned policy learning [25, 26, 27]. Despite these efforts, how to harness unbounded representations of natural language knowledge for building generalist decision agents remains the following challenges: i) LLMs, trained on text corpora, typically lack grounding in the physical world and fail to capture any environment dynamics; ii) Leveraging LLMs for decision-making is prone to knowledge misalignment due to the semantic gap between the text and decision modalities; iii) Scalable implementations are necessary to fully harness language knowledge to train generalist decision models.

To tackle these challenges, we propose **Text-to-Decision Agent (T2DA)**, a simple and scalable pre-training framework for offline meta-RL via aligning language knowledge with environment dynamics of decision tasks. First, we pre-train a generalized world model to encode multi-task data into dynamics-aware decision embeddings, effectively capturing task-specific environment dynamics. Second, inspired by CLIP [16], we predict which textual description goes with which decision embedding, efficiently bridging their semantic gap via contrastive language-decision pre-training. It distills the world model structure to the text modality and aligns text embeddings to comprehend the environment dynamics. Finally, we deploy T2DA on two mainstream conditional generation architectures, training generalist policies conditioned on aligned text embeddings. During evaluation, natural language is used to reference learned decision perceptions or describe new ones, and the agent can directly realize zero-shot text-to-decision generation according to textual instructions at hand. Extensive experiments show that T2DA transfers nontrivially to downstream tasks, achieves high-capacity zero-shot generalization, and outperforms various types of baselines.

In summary, our main contributions are as follows:

- We present a generalized world model designed to capture the environment dynamics, facilitating effective grounding of linguistic knowledge in decision domains.
- We introduce contrastive language-decision pre-training to bridge their semantic gap, aligning the text embeddings to comprehend the environment dynamics.
- We propose a simple and scalable framework that supervises offline meta-RL with natural language, and develop scalable implementations with the potential to train decision models at scale: Text-to-Decision Diffuser and Text-to-Decision Transformer.

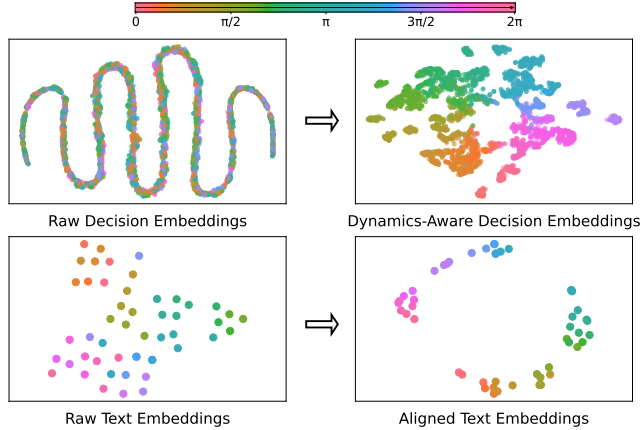


Figure 1: t-SNE visualization of Ant-Dir where tasks with target directions in $[0, 2\pi]$ are mapped to rainbow-colored points. Top: we encode multi-task data into dynamics-aware decision embeddings to capture task-specific environment dynamics. Bottom: we bridge the semantic gap between text and decision via contrastive pre-training. The aligned text embeddings follow a cyclic spectrum that exactly matches the periodicity of angular directions *in a physical sense*. This interesting finding shows that we effectively align text embeddings to comprehend environment dynamics and facilitate convincing language grounding in decision domains.

2 Related Work

Offline Meta-RL learns adaptable policies from pre-collected datasets without requiring online environment interaction [11, 28]. It addresses two fundamental challenges: the distributional shift between behavior and learned policies in offline RL [29, 30], and the quick adaptation to unseen tasks with minimal data. Existing approaches can be broadly categorized into the memory-based (e.g., RL² [31] and LLIRL [32, 33]), the optimization-based (MAML [34] and MACAW [11]), and the context-based methods (e.g., PEARL [35], VariBAD [36], etc. [28, 13, 37]). Recent studies of in-context RL attempt to leverage the in-context learning capability of transformers to improve RL’s generalization via casting RL as an across-episodic sequential prediction problem with few-shot prompting at test time [38, 39, 40, 41], such as AD [42] and DPT [43]. In general, the above approaches rely on high-quality samples or domain knowledge to infer task beliefs during evaluation, limiting their generality with this restricted form of supervision. This inspires us to explore a much broader source of supervision from natural language.

LLMs for RL. Existing studies on leveraging LLMs for RL domains can be categorized into three main threads: LLMs as rewards, LLMs as policies, and language-conditioned policy learning. The first kind uses LLMs or VLMs (vision-language models) to automate the generation and shaping of reward functions [44, 22, 23, 24] or state representations [45], providing denser information to facilitate RL training. The second kind utilizes LLMs or VLMs as the policy backbone [18, 46, 47, 48]. Vision-language-action (VLA) models provide a direct instantiation of using pre-trained VLMs for robotics, fine-tuning visually-conditioned VLMs to generate robot control actions, such as RT-2 [49], OpenVLA [50], and $\pi_{0.5}$ [51]. In contrast, we leverage LLMs to obtain task representations for offline meta-RL and train conventional RL policies based on decision transformer or diffuser architectures.

Language-conditioned policy. We focus on the third kind that learns a language-conditioned policy, extracting the world knowledge encoded in LLMs to help train offline meta-RL agents. Previous studies acquire language embeddings using pre-trained LLMs like BERT and GPT [52], while these embeddings are learned independently from decision tasks and may fail to capture domain-related information. Early methods develop rule-based [53] or task-specific [54] intermediate representations to capture task-related semantics, which can require a laborious design and lack scalability. Under an imitation learning paradigm, BC-Z [25] trains a vision-based robotic manipulation policy conditioned on pre-trained language embeddings, and BAKU [27] produces a simple transformer architecture that fuses multi-modal vision-language information and temporal context for multi-task decision-making. To improve language grounding, [55] translates natural language to task language with a referential game, and [26] uses supervised fine-tuning to enable bidirectional translation between language-described skills and rollout data. Overall, the challenges include the inability to capture environment dynamics, knowledge misalignment due to semantic gaps, and limited model scalability. In the paper, we tackle these challenges to fully harness LLMs for building generalist decision agents.

3 Method

In this section, we present Text-to-Decision Agent (T2DA), a simple and scalable framework to train generalist policies. Figure 2 illustrates the overall pipeline, followed by detailed implementations in the subsequent subsections. Corresponding algorithm pseudocodes are given in Appendix A.

3.1 Problem Statement

We consider a language-conditioned offline meta-RL setting. Tasks follow a distribution $M_k = \langle \mathcal{S}, \mathcal{A}, \mathcal{T}_k, \mathcal{R}_k, \gamma \rangle \sim P(M)$, sharing the same state-action spaces while varying in the reward and state transition functions, i.e., environment dynamics. For each training task k , our system receives a user-provided natural language supervision l_k that describes the task (e.g., open the door), paired with an offline dataset $\mathcal{D}_k = \sum_i (s_i^k, a_i^k, r_i^k, s_i'^k)$ collected by arbitrary behavior policies. The agent can only access the offline datasets $\sum_k (l_k, \mathcal{D}_k)$ to train a generalist policy $\pi(a|s, l)$ to follow language instructions. At test time, natural language is used to reference learned decision perceptions or describe new ones. The agent can perform text-to-decision generation in test environments in a zero-shot manner, given any language instruction l_{new} at hand. The primary objective is to obtain a highly generalizable policy that achieves strong zero-shot performance on unseen tasks.

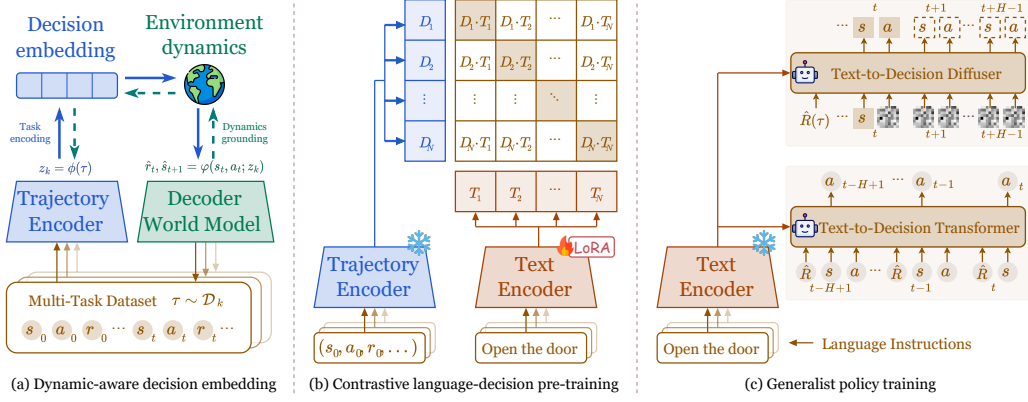


Figure 2: The overall pipeline of T2DA. (a) We encode the multi-task trajectories into dynamics-aware decision embeddings and decode the generalized world model conditioned on that embedding, effectively capturing the environment dynamics. (b) We bridge the semantic gap between decision and text by fine-tuning the text encoder (initialized from popular language models such as CLIP or T5) to align the produced text embeddings with dynamics-aware decision embeddings using contrastive loss. It distills the world model structure from decision embeddings to the text modality, aligning text embeddings to comprehend the environment dynamics. (c) We condition the generalist policy on the aligned text embeddings, and develop scalable implementations with the potential to train decision models at scale: Text-to-Decision Diffuser and Text-to-Decision Transformer. During evaluation, the agent can directly realize zero-shot text-to-decision generation according to textual instructions at hand, enabling high-capacity zero-shot generalization to downstream tasks.

3.2 Dynamics-Aware Decision Embedding

Natural language is a complex and unbounded representation of human instruction, and training policies to harness language supervision is nontrivial. A standard approach is to convert language instructions to simpler latent embeddings using popular pre-trained language models that can draw on a wealth of knowledge learned from copious amounts of text. However, directly applying these models to decision domains could fail to solve user-specific problems, as they are typically learned independently of decision-making tasks. A major challenge is that LLMs trained on text corpora typically lack grounding in the interacting environment and fail to capture its dynamics.

For effective grounding, the agent should comprehend environment dynamics of underlying decision tasks. RL typically learns from active interactions with the outer environment. The world model, i.e., the reward and state transition functions $p(s', r|s, a)$, fully characterizes the environment and presents a promising alternative to capturing environment dynamics. Hence, we introduce a generalized world model to encode the multi-task decision data into a dynamics-aware embedding space. We separate the reasoning about the world model into two parts: i) encoding the dynamics-specific information into a latent embedding, and ii) decoding environment dynamics conditioned on that embedding.

First, we use a trajectory encoder ϕ to abstract multi-task decision trajectories into a compact embedding z that captures task-specific dynamics. Specifically, we “tokenize” different elements (i.e., state, action, and reward) in the raw sequence $\tau = (s_0, a_0, r_0, \dots, s_L, a_L, r_L) \sim \mathcal{D}_k$, by lifting them to a common representation space using element-specific tokenizers $f_\phi^s, f_\phi^a, f_\phi^r$ as

$$\tau_e = (e_0^s, e_0^a, e_0^r, \dots, e_L^s, e_L^a, e_L^r), \text{ where } e_t^s = f_\phi^s(s_t), e_t^a = f_\phi^a(a_t), e_t^r = f_\phi^r(r_t), \forall t \in [0, L]. \quad (1)$$

We use a bi-directional transformer E_ϕ to extract the dynamics-aware embedding from input tokens as $z_k = E_\phi(\tau_e)$. The bi-directional structure allows for capturing the forward and inverse dynamics.

Second, we introduce a decoder φ containing a reward model $\mathcal{R}_\varphi$ and state transition model $\mathcal{T}_\varphi$. The latent embedding is augmented into the input to predict the instant reward $\hat{r}_t = \mathcal{R}_\varphi(s_t, a_t; z_k)$ and next state $\hat{s}_{t+1} = \mathcal{T}_\varphi(s_t, a_t; z_k)$. The encoder-decoder pipeline is jointly trained by minimizing the reward and state transition prediction error conditioned on the decision embedding as

$$\mathcal{L}(\phi, \varphi) = \mathbb{E}_{\tau \sim \mathcal{D}_k} \left[\mathbb{E}_{z_k \sim \phi(\tau)} \left[\mathbb{E}_t \left[(r_t - \mathcal{R}_\varphi(s_t, a_t; z_k))^2 + (s_{t+1} - \mathcal{T}_\varphi(s_t, a_t; z_k))^2 \right] \right] \right], \forall k. \quad (2)$$

In practice, we randomly sample a trajectory from the dataset as $\tau \sim \mathcal{D}_k$ and use its decision embedding to decode the dynamics of other trajectories in the same dataset as $\tau^* \in \mathcal{D}_k \setminus \tau$. The reason is that, since the embedding has access to the entire trajectory’s information, using it to decode the same trajectory can lead to deceptive traps during training. After proper training, we freeze the world model for the subsequent learning phases.

3.3 Contrastive Language-Decision Pre-training

To successfully instruct generalist training with human language, we must bridge the semantic gap between text and the “decision modality”. Inspired by CLIP [16], we predict which textual description goes with which training task, bridging their semantic gap via contrastive learning. The core idea is to connect the language supervision with corresponding decision embeddings, distilling the world model structure to the text modality and aligning text embeddings to comprehend environment dynamics.

Formally, given a batch of N (trajectory τ , textual task description l) pairs, we predict which of the $N \times N$ possible (trajectory, text) pairings across a batch actually occurred. We use the pre-trained trajectory encoder in Sec. 3.2 to extract the dynamics-aware decision embedding from the trajectory as $z = \phi(\tau)$. The text encoder is a transformer-based model ψ that could be initialized from one of a wide variety of popular language models, such as CLIP [16] or T5 [56]. Then, the text embedding z_T is easily derived by tokenizing the textual instruction l and feeding it to the encoder as $z_T = \psi(l)$.

We model the pre-training task as a text-decision multi-modal learning problem by fine-tuning the text encoder while keeping the trajectory encoder fixed. The objective is to maximize the cosine similarity of the decision and text embeddings of the N real pairs in the batch while minimizing the cosine similarity of the embeddings of the $N^2 - N$ incorrect pairings. The two directional similarities between the modalities, which are transposes of each other, are defined as

$$\text{sim}(\tau, l) = \exp(\alpha) \cdot \frac{\phi(\tau)W_D \cdot \psi(l)W_T}{\|\phi(\tau)W_D\| \cdot \|\psi(l)W_T\|}, \quad \text{sim}(l, \tau) = \exp(\alpha) \cdot \frac{\psi(l)W_T \cdot \phi(\tau)W_D}{\|\psi(l)W_T\| \cdot \|\phi(\tau)W_D\|}, \quad (3)$$

where α is a learnable temperature controlling the range of similarities, and W_D and W_T are learnable projections that map feature representations of each modality to a joint multi-modal embedding space. The decision-to-text $p(\tau)$ and text-to-decision $p(l)$ similarity scores in the batch are calculated as

$$p(\tau_k) = \frac{e^{\text{sim}(\tau_k, l_k)}}{\sum_{i=1}^N e^{\text{sim}(\tau_k, l_i)}}, \quad p(l_k) = \frac{e^{\text{sim}(l_k, \tau_k)}}{\sum_{i=1}^N e^{\text{sim}(l_k, \tau_i)}}. \quad (4)$$

Let $q(\tau)$ and $q(l)$ denote the ground-truth similarity scores, where negative pairs have a probability of 0 and positive pairs have 1. We optimize a symmetric cross-entropy loss over similarity scores as

$$\mathcal{L}(\psi) = 0.5 \cdot \mathbb{E}_{\tau, l} [\text{CE}(p(\tau), q(\tau)) + \text{CE}(p(l), q(l))], \quad (5)$$

where CE is the cross-entropy between two distributions. To effectively adapt the text encoder to the semantics of decision modality at a lightweight computing cost, we employ low-rank adaptation (LoRA) [57] to fine-tune the text encoder with the above contrastive loss. Experiments in Appendix E present the analysis of its efficiency and superiority compared to full-parameter fine-tuning.

This alignment process distills the environment dynamics from the decision embedding to the text modality. The aligned text embeddings not only serve as mere linguistic representations but also establish stronger comprehension of underlying decision tasks, effectively bridging the semantic gap between language supervision and decision dynamics.

3.4 Generalist Policy Learning

A widely adopted manner to tackle generalization in RL domains is to condition the policy on some representation that can capture task-relevant information, such as context-based meta-learning [35, 36] and prompt-based methods [12, 14]. Under this unified framework, we formulate the generalist agent as a task-conditioned policy model. There exists a true variable that represents the task identity, and we use a latent representation h to approximate that variable. Then, a base policy $\pi(\cdot, h)$ can be shared across tasks by conditioning on the task representation as $\pi_k(a|s) = \pi(a|s; h_k), \forall k$. Previous studies usually approximate the task representation by acquiring expert decision data [12] or exploring the unseen task first [13]. This restricted form of supervision limits their universality and practicability since expensive high-quality samples or explorations are needed to specify any new task.

Motivated by the recent success of language models, we study a promising alternative that leverages a much broader source of supervision. Natural language is a flexible representation for transferring a variety of human ideas and intentions. Learning from its supervision can not only enhance the representation power but also enable efficient knowledge transfer. We map the textual task description l to a latent embedding $\psi(l)$ via the fine-tuned text encoder in Sec. 3.3, and use it to approximate the task representation as $h \approx \psi(l)$. Then, the task-agnostic generalist policy is approximated by

$$\pi(a \mid s; h_k) \approx \pi(a \mid s; \psi(l_k)), \quad \forall k. \quad (6)$$

At test time, the agent can directly perform text-to-decision policy generation according to any textual instruction l_{new} at hand, enabling high-capacity zero-shot generalization to downstream tasks without the need for expensive decision demonstrations or warmup explorations.

3.5 Scalable Implementations

We develop scalable implementations of T2DA using two mainstream conditional generation architectures that hold the promise to train RL models at scale: the decision diffuser [58] and the decision transformer [59], yielding the following T2DA variants.

Text-to-Decision Diffuser (T2DA-D). Inspired by the recent success of text-to-image diffusion models [60, 61], we implement the text-to-decision counterpart by modeling the policy as a return-conditioned diffusion model with action planning. For model simplicity and scalability, we choose to diffuse over the H -step state-action trajectory as

$$x_c(\tau) = \begin{bmatrix} s_t & s_{t+1} & \dots & s_{t+H-1} \\ a_t & a_{t+1} & \dots & a_{t+H-1} \end{bmatrix}_c, \quad (7)$$

where c denotes the timestep in the forward diffusion process, and $c=0$ corresponds to the unperturbed data.¹ To enable a generalist policy with natural language supervision, we include the fine-tuned text embedding $\psi(l)$ as an additional condition to the diffusion model. To this end, the generalist policy training is formulated as the standard problem of conditional generative modeling as

$$\max_{\theta} \mathbb{E}_{\tau \sim \mathcal{D}} \left[\log p_{\theta} \left(x_0(\tau) \mid \hat{R}(\tau); \psi(l) \right) \right], \quad (8)$$

where θ denotes diffusion parameters and $\hat{R}$ is the return-to-go. During evaluation, we can cast planning in RL as sampling from the diffuser. At each timestep, we observe a state s_t in the environment, sample an H -step trajectory $x_0(\tau)$ with the diffusion process conditioned on a target return-to-go and task prompt $\psi(l)$, execute the first action a_t in $x_0(\tau)$, and transition to the next state s_{t+1} . More details can be referred to as in Algorithms 3 and 5 in Appendix A.

Text-to-Decision Transformer (T2DA-T). We also implement T2DA with the causal transformer architecture, drawing upon the simplicity, scalability, and associated advances in language modeling such as BERT and GPT. We prepend a task prompt to the input of the causal transformer to realize generalization across tasks, akin to Prompt-DT [12]. For each task, we take the prompt-augmented trajectory τ^+ as the input, which contains both the fine-tuned text embeddings $\psi(l)$ and the most recent H -step history sampled from offline datasets as

$$\tau^+ = \left(\psi(l); \hat{R}_{t-H+1}, s_{t-H+1}, a_{t-H+1}, \dots, \hat{R}_t, s_t, a_t \right), \quad (9)$$

where the return-to-go $\hat{R}_t$ is the cumulative rewards from the current time step till the end of the episode. Then, the model predicts actions autoregressively using a causal self-attention mask. This architecture introduces only one additional token compared to the standard decision transformer, allowing minimal architecture change and lightweight cost for generalist training. More details can be referred to as in Algorithms 4 and 6 in Appendix A.

¹Decision diffuser diffuses over only states and trains another inverse dynamics model to generate actions. In contrast, we only keep a unified diffusion model to i) ensure simplicity with a single end-to-end model in a multi-task setting, and ii) unlock the potential of diffusion models at scale, rather than depending on additional MLPs to generate part of the information, e.g., actions.

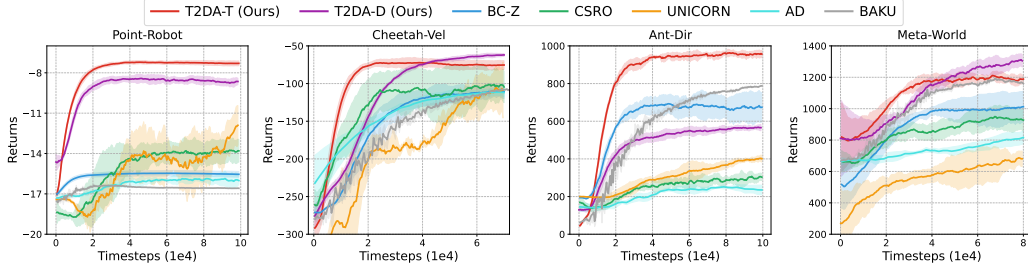


Figure 3: Zero-shot test return curves of T2DA against baselines using Mixed datasets.

Table 1: Zero-shot test returns of T2DA against baselines using Mixed datasets, i.e., numerical results of converged performance from Figure 3. Best results in **bold** and second best underlined.

Environment	T2DA-T	T2DA-D	BC-Z	CSRO	UNICORN	AD	BAKU
Point-Robot	-7.2 ± 0.1	<u>-8.4</u> ± 0.2	-15.5 ± 0.2	-13.7 ± 0.5	-11.4 ± 1.3	-15.8 ± 0.2	-16.4 ± 0.1
Cheetah-Vel	-70.3 ± 3.8	-60.4 ± 1.8	-107.4 ± 8.5	-92.6 ± 5.9	-94.1 ± 20.0	-107.6 ± 2.3	-104.5 ± 4.8
Ant-Dir	970.3 ± 8.7	570.6 ± 13.1	700.0 ± 43.3	317.1 ± 27.6	407.3 ± 21.6	268.9 ± 3.4	<u>786.9</u> ± 17.3
Meta-World	<u>1274.5</u> ± 48.4	1376.4 ± 39.6	1053.8 ± 33.8	1005.2 ± 69.6	754.6 ± 62.7	921.0 ± 66.9	1187.6 ± 18.2

4 Experiments

We comprehensively evaluate and analyze our method on popular benchmarking domains across datasets of varying qualities, aiming to answer the following research questions:

- Can T2DA achieve consistent performance gain on zero-shot generalization capacity to unseen tasks? We compare it to various types of strong baselines, including offline meta-RL, in-context RL, and language-conditioned policy learning approaches. (Sec. 4.1)
- What is the contribution of each component to T2DA’s performance? We ablate both the T2DA-D and T2DA-T architectures to analyze the respective impact of world model pre-training, contrastive language-decision pre-training, and language supervision. (Sec. 4.2)
- How robust is T2DA across diverse settings? We evaluate T2DA against baselines using offline datasets of varying qualities, and assess T2DA’s performance when initializing the text encoder from different LLMs. (Sec. 4.3)

Environments. We evaluate T2DA on three benchmarks that are widely adopted to assess generalization capacities of RL algorithms: i) the 2D navigation Point-Robot; ii) the multi-task MuJoCo locomotion control, containing Cheetah-Vel and Ant-Dir; and iii) the Meta-World platform for robotic manipulation, where a robotic arm is designed to perform a wide range of manipulation tasks, such as close faucet, lock door, open door, and press button. For each domain, we randomly sample a distribution of tasks captioned with text descriptions, and split them into training and test sets. We employ SAC [62] to independently train a single-task policy on each training task for offline dataset collection. We develop three types of offline datasets: Mixed, Medium, and Expert, following the common practice in offline RL [63]. Appendix B presents more details of environments and datasets.

Baselines. We compare T2DA to five competitive baselines that cover three representative paradigms in tackling RL generalization: the context-based offline meta-RL approaches, including 1) CSRO [13] and 2) UNICORN [37]; the in-context RL method of 3) AD [42]; and the language-conditioned policy learning methods including 4) BC-Z [25] and 5) BAKU [27].

For evaluation, all experiments are conducted using five different random seeds. The mean of received return is plotted as the bold line, with 95% bootstrapped confidence intervals of the mean indicated by the shaded region. Appendix D gives detailed implementations of T2DA. Appendix E gives analysis on fine-tuning of the text encoder.

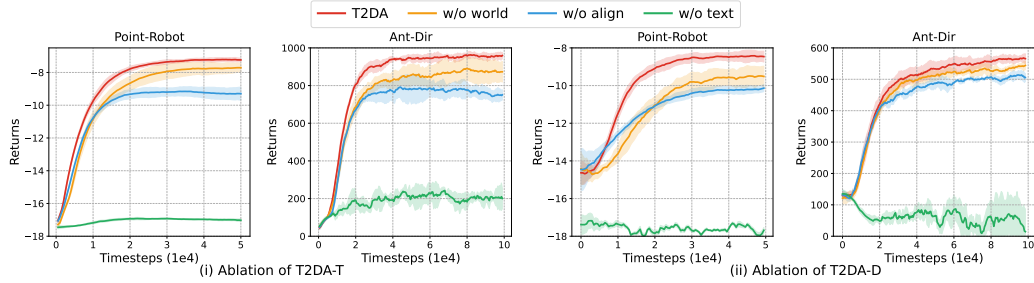


Figure 4: Ablation results using Mixed datasets. w/o world omits pre-training the trajectory encoder, w/o align omits contrastive pre-training, and w/o text omits the language supervision.

Table 2: Numerical results of ablation study on both T2DA-D and T2DA-T using Mixed datasets, i.e., the final test returns after learning convergence from Figure 4. Best results in **bold**.

Environment	T2DA-T	w/o world	w/o align	w/o text	T2DA-D	w/o world	w/o align	w/o text
Point-Robot	-7.2$\pm$0.1	-7.7 $\pm$ 0.3	-9.2 $\pm$ 0.2	-16.9 $\pm$ 0.1	-8.4$\pm$0.3	-9.5 $\pm$ 0.4	-10.1 $\pm$ 0.2	-17.2 $\pm$ 0.4
Ant-Dir	963.4$\pm$15.3	891.0 $\pm$ 54.4	792.3 $\pm$ 48.3	242.6 $\pm$ 13.7	568.0$\pm$10.8	544.1 $\pm$ 11.1	513.8 $\pm$ 11.0	135.1 $\pm$ 10.3

4.1 Main Results

We compare our method against various baselines under an aligned zero-shot setting. Figure 3 and Table 1 present test return curves and numerical results of converged performance on various benchmark environments using Mixed datasets. A noteworthy point is that language-conditioned baselines such as BC-Z and BAKU generally achieve better performance than offline meta-RL and in-context RL baselines, especially in harder environments like Ant-Dir and Meta-World. This again validates our motivation of exploring a much broader source of supervision from natural language, harnessing the representation power and knowledge transferability embodied in pre-trained LLMs.

In these diversified environments, both T2DA-D and T2DA-T consistently achieve significantly superior performance regarding learning speed and final asymptotic results compared to the three types of baselines. In most cases, T2DA-D and T2DA-T take the top two rankings for the best and second-best performance, and showcase comparably strong zero-shot generalization capacities on average across all evaluated environments. It highlights the effectiveness of both T2DA implementations using the two mainstream conditional generation architectures. Furthermore, our method typically demonstrates lower variance during generalist policy learning, signifying not only enhanced learning efficiency but also improved training stability.

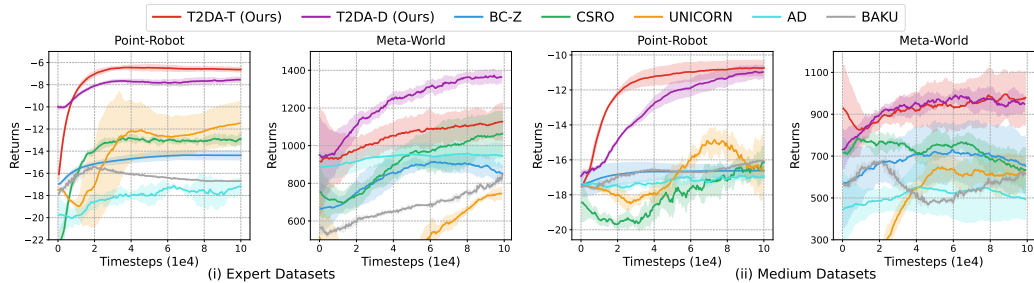


Figure 5: Results of robustness to data quality, where T2DA is compared to baselines using Expert and Medium datasets. T2DA-D and T2DA-T achieve consistent superiority across various datasets.

4.2 Ablation Study

We compare T2DA to three ablations: 1) w/o world, it eliminates pre-training the trajectory encoder and instead updates it jointly with the text encoder during contrastive pre-training; 2) w/o align, it

Table 3: Numerical results of robustness to data quality where T2DA is compared against baselines using Expert and Medium datasets, i.e., the final test returns after learning convergence from Figure 5. Both T2DA-D and T2DA-T achieve consistent superiority across datasets of varying quality. Best results in **bold** and second best underlined.

Expert	T2DA-T	T2DA-D	BC-Z	CSRO	UNICORN	AD	BAKU
Point-Robot	-6.4 ± 0.2	-7.5 ± 0.2	-14.4 ± 0.4	-12.8 ± 0.2	-11.5 ± 1.7	-17.1 ± 0.2	-15.4 ± 0.1
Meta-World	<u>1142.8</u> ± 106.6	1387.7 ± 21.6	935.5 ± 26.7	1087.0 ± 78.8	787.1 ± 16.0	991.0 ± 37.7	833.6 ± 8.3
Medium	T2DA-T	T2DA-D	BC-Z	CSRO	UNICORN	AD	BAKU
Point-Robot	-10.7 ± 0.4	-11.0 ± 0.3	-16.6 ± 0.2	-16.1 ± 0.8	-14.9 ± 0.5	-16.9 ± 0.2	-16.1 ± 0.1
Meta-World	1033.3 ± 130.6	<u>1025.3</u> ± 43.6	756.9 ± 103.4	854.1 ± 17.9	715.8 ± 93.5	585.7 ± 126.1	669.5 ± 11.4

omits contrastive language-decision pre-training and fixes the text encoder initialized from pre-trained language models; and 3) w/o text, it completely removes the text encoder.

Figure 4 presents ablation results of both architectures. First, the erasure of pre-training the trajectory encoder (T2DA vs. w/o world) results in decreased test returns with greater variances. It indicates that capturing environment dynamics via the world model can enable more precise and stable knowledge alignment between text and decision. Second, omitting the language-decision pre-training (w/o world vs. w/o align) further decreases the performance. It confirms the existence of a semantic gap between text and decision, whereas the gap can be effectively bridged by contrastive knowledge alignment. Finally, removing the language supervision (w/o align vs. w/o text) can produce catastrophic performance with poor zero-shot generalization, highlighting the necessity of leveraging a wealth of knowledge from natural language. In summary, the performance of T2DA declines when any component is omitted, verifying that all components are essential to its final capabilities.

4.3 Robustness Study

A reliable framework with scalability should be robust across diverse settings. We assess T2DA’s robustness in two critical aspects: data quality and choice of text encoders.

Robustness to Data Quality. We compare T2DA against baselines using Expert and Medium offline datasets. As shown in Figure 5 and Table 3, both T2DA-D and T2DA-T consistently outperform baselines on these datasets. Notably, while most baselines suffer significant performance degradation when trained on Medium datasets, T2DA maintains high-capacity generalization despite the lower data quality. It showcases the promising applicability of T2DA in real-world scenarios where the agent often needs to learn from sub-optimal data. In summary, T2DA achieves consistent superiority across datasets of varying quality, yielding excellent robustness.

Robustness to Text Encoders. We also investigate T2DA’s sensitivity to different representations of language supervision by initializing the text encoder from three distinct language models: CLIP [16], BERT [64], and T5 [56]. As shown in Figure 6, the near-identical performance across all language models highlights that T2DA’s effectiveness is independent of the specific text encoder architecture. It suggests that T2DA can effectively leverage the semantic understanding capabilities inherent to different language models while maintaining consistent performance.

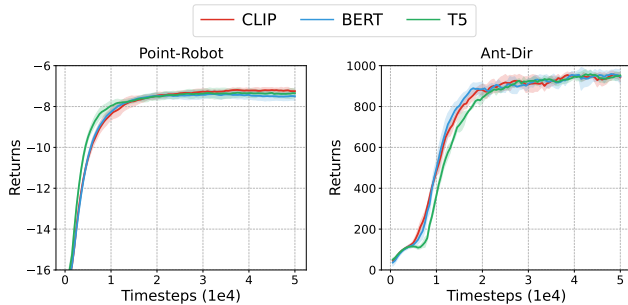


Figure 6: Results of robustness to text encoders where T2DA initializes the text encoder from CLIP, BERT, and T5, showing that T2DA can efficiently harness language knowledge embedded in different language models.

4.4 Visualization Insights

We gain deep insights into the knowledge alignment between text and decision through t-SNE visualization on the example Cheetah-Vel task as shown in Figure 7.

Decision Embeddings. Raw decision embeddings are derived from a randomly initialized trajectory encoder. In the dynamics-aware embedding space, sample points from different tasks are more clearly distinguished, and similar tasks are grouped more closely. The initially entangled trajectories are transformed into well-separated clusters, highlighting the successful capture of task-specific environment dynamics via the generalized world model.

Text Embeddings. Raw text embeddings are derived from pre-trained language models, and are scattered with a vague distribution structure. In the aligned space, text embeddings follow a more distinguished distribution that matches the metric of task similarity, highlighting successful bridging of the semantic gap. Notably, the aligned embeddings form a clear rectilinear distribution from red (low velocity) to blue (high velocity), exactly matching the rectilinear spectrum of target velocities *in a physical sense*. It verifies efficient comprehension of the environment dynamics.

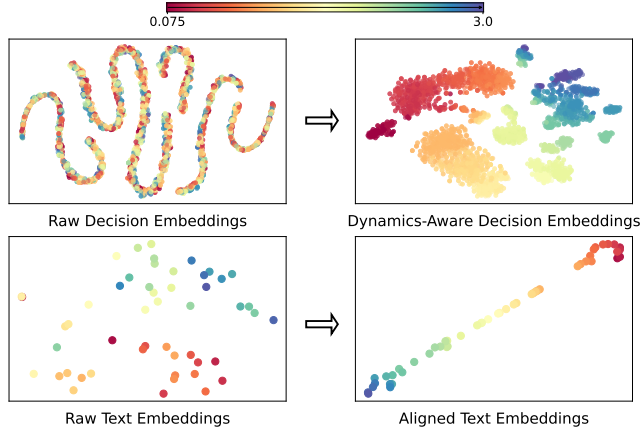


Figure 7: t-SNE visualization of Cheetah-Vel where tasks with target velocities in $[0.075, 3.0]$ are mapped into rainbow-colored points. Top: the evolution from Raw Decision Embeddings to Dynamics-Aware Decision Embeddings, where the initially entangled trajectories are transformed into well-separated clusters that highlight the successful capture of task-specific environment dynamics. Bottom: using dynamics-aware decision embeddings to align text embeddings with contrastive loss. This inspiring finding verifies the effective alignment of text embeddings to comprehend environment dynamics and interpretable language grounding.

5 Conclusions, Limitations, and Future Work

In the paper, we tackle offline meta-RL challenge via leveraging a much broader source of supervision from natural language. Improvements in zero-shot generalization capacities highlight the potential impact of our scalable implementations: text-to-decision diffuser and text-to-decision transformer. These findings open up new avenues for building more flexible and accessible generalist agents that can understand and act upon natural language instructions without extensive domain-specific training.

Though, our agent is trained on lightweight datasets compared to popular large models. An essential step is to implement on vast datasets with diversified domains, unleashing the scaling law with the diffusion or transformer architectures. Another step is to deploy the idea of knowledge alignment in VLA models, or deploy the text-to-decision paradigm on real robots toward embodied intelligence.

Acknowledgements

This work was supported in part by the National Natural Science Foundation of China (Nos. 62376122 and 72394363), and in part by the AI & AI for Science Project of Nanjing University.

References

- [1] Volodymyr Mnih, Koray Kavukcuoglu, David Silver, Andrei A Rusu, Joel Veness, et al. Human-level control through deep reinforcement learning. *Nature*, 518(7540):529–533, 2015.
- [2] Alhussein Fawzi, Matej Balog, Aja Huang, Thomas Hubert, Bernardino Romera-Paredes, et al. Discovering faster matrix multiplication algorithms with reinforcement learning. *Nature*, 610(7930):47–53, 2022.
- [3] Rafael Rafailov, Archit Sharma, Eric Mitchell, Christopher D Manning, Stefano Ermon, and Chelsea Finn. Direct preference optimization: Your language model is secretly a reward model. In *Advances in Neural Information Processing Systems*, volume 36, 2023.
- [4] Kevin Black, Michael Janner, Yilun Du, Ilya Kostrikov, and Sergey Levine. Training diffusion models with reinforcement learning. In *Proceedings of International Conference on Learning Representations*, 2024.
- [5] Ganqu Cui, Yuchen Zhang, Jiacheng Chen, Lifan Yuan, Zhi Wang, Yuxin Zuo, Haozhan Li, Yuchen Fan, Huayu Chen, Weize Chen, et al. The entropy mechanism of reinforcement learning for reasoning language models. *arXiv preprint arXiv:2505.22617*, 2025.
- [6] Jianhao Yan, Yafu Li, Zican Hu, Zhi Wang, Ganqu Cui, Xiaoye Qu, Yu Cheng, and Yue Zhang. Learning to reason under off-policy guidance. In *Advances in Neural Information Processing Systems*, 2025.
- [7] Aaron Jaech, Adam Kalai, Adam Lerer, Adam Richardson, Ahmed El-Kishky, et al. OpenAI o1 system card. *arXiv preprint arXiv:2412.16720*, 2024.
- [8] Daya Guo, Dejian Yang, Haowei Zhang, Junxiao Song, Ruoyu Zhang, et al. DeepSeek-R1: Incentivizing reasoning capability in llms via reinforcement learning. *arXiv preprint arXiv:2501.12948*, 2025.
- [9] Kevin Black, Noah Brown, Danny Driess, Adnan Esmail, Michael Equi, Chelsea Finn, Niccolo Fusai, Lachy Groom, Karol Hausman, Brian Ichter, et al. π_0 : A vision-language-action flow model for general robot control. *arXiv preprint arXiv:2410.24164*, 2024.
- [10] Thomas Schmied, Markus Hofmarcher, Fabian Paischer, Razvan Pascanu, and Sepp Hochreiter. Learning to modulate pre-trained models in RL. In *Advances in Neural Information Processing Systems*, 2023.
- [11] Eric Mitchell, Rafael Rafailov, Xue Bin Peng, Sergey Levine, and Chelsea Finn. Offline meta-reinforcement learning with advantage weighting. In *Proceedings of International Conference on Machine Learning*, pages 7780–7791, 2021.
- [12] Mengdi Xu, Yikang Shen, Shun Zhang, Yuchen Lu, Ding Zhao, Joshua Tenenbaum, and Chuang Gan. Prompting decision transformer for few-shot policy generalization. In *Proceedings of International Conference on Machine Learning*, volume 162, pages 24631–24645, 2022.
- [13] Yunkai Gao, Rui Zhang, Jiaming Guo, Fan Wu, Qi Yi, Shaohui Peng, Siming Lan, Ruizhi Chen, Zidong Du, Xing Hu, et al. Context shift reduction for offline meta-reinforcement learning. In *Advances in Neural Information Processing Systems*, volume 36, 2023.
- [14] Zhi Wang, Li Zhang, Wenhao Wu, Yuanheng Zhu, Dongbin Zhao, and Chunlin Chen. Meta-DT: Offline meta-rl as conditional sequence modeling with world model disentanglement. In *Advances in Neural Information Processing Systems*, 2024.
- [15] Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, et al. GPT-4 technical report. *arXiv preprint arXiv:2303.08774*, 2023.
- [16] Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, et al. Learning transferable visual models from natural language supervision. In *Proceedings of International Conference on Machine Learning*, pages 8748–8763, 2021.

- [17] Christian Schlarman, Naman Deep Singh, Francesco Croce, and Matthias Hein. Robust CLIP: Unsupervised adversarial fine-tuning of vision embeddings for robust large vision-language models. In *Proceedings of International Conference on Machine Learning*, 2024.
- [18] Michael Ahn, Anthony Brohan, Noah Brown, Yevgen Chebotar, Omar Cortes, Byron David, Chelsea Finn, Chuyuan Fu, Keerthana Gopalakrishnan, Karol Hausman, et al. Do as I can, not as I say: Grounding language in robotic affordances. In *Proceedings of Conference on Robot Learning*, 2022.
- [19] Andrew Szot, Max Schwarzer, Harsh Agrawal, Bogdan Mazouze, Rin Metcalf, Walter Talbott, Natalie Mackraz, R Devon Hjelm, and Alexander T Toshev. Large language models as generalizable policies for embodied tasks. In *Proceedings of International Conference on Learning Representations*, 2024.
- [20] Zican Hu, Wei Liu, Xiaoye Qu, Xiangyu Yue, Chunlin Chen, Zhi Wang, and Yu Cheng. Divide and conquer: Grounding LLMs as efficient decision-making agents via offline hierarchical reinforcement learning. In *Proceedings of International Conference on Machine Learning*, 2025.
- [21] Runzhe Zhan, Yafu Li, Zhi Wang, Xiaoye Qu, Dongrui Liu, Jing Shao, Derek F Wong, and Yu Cheng. ExGRPO: Learning to reason from experience. *arXiv preprint arXiv:2510.02245*, 2025.
- [22] Juan Rocamonde, Victoriano Montesinos, Elvis Nava, Ethan Perez, and David Lindner. Vision-language models are zero-shot reward models for reinforcement learning. In *Proceedings of International Conference on Learning Representations*, 2024.
- [23] Tianbao Xie, Siheng Zhao, Chen Henry Wu, Yitao Liu, Qian Luo, Victor Zhong, Yanchao Yang, and Tao Yu. Text2Reward: Reward shaping with language models for reinforcement learning. In *Proceedings of International Conference on Learning Representations*, 2024.
- [24] Yuwei Fu, Haichao Zhang, Di Wu, Wei Xu, and Benoit Boulet. FuRL: Visual-language models as fuzzy rewards for reinforcement learning. In *Proceedings of International Conference on Machine Learning*, 2024.
- [25] Eric Jang, Alex Irpan, Mohi Khansari, Daniel Kappler, Frederik Ebert, Corey Lynch, Sergey Levine, and Chelsea Finn. BC-Z: Zero-shot task generalization with robotic imitation learning. In *Proceedings of Conference on Robot Learning*, pages 991–1002, 2022.
- [26] Jing-Cheng Pang, Si-Hang Yang, Kaiyuan Li, Jiaji Zhang, Xiong-Hui Chen, Nan Tang, and Yang Yu. Knowledgeable agents by offline reinforcement learning from large language model rollouts. In *Advances in Neural Information Processing Systems*, 2024.
- [27] Siddhant Haldar, Zhuoran Peng, and Lerrel Pinto. BAKU: An efficient transformer for multi-task policy learning. In *Advances in Neural Information Processing Systems*, 2024.
- [28] Haoqi Yuan and Zongqing Lu. Robust task representations for offline meta-reinforcement learning via contrastive learning. In *Proceedings of International Conference on Machine Learning*, pages 25747–25759, 2022.
- [29] Sergey Levine, Aviral Kumar, George Tucker, and Justin Fu. Offline reinforcement learning: Tutorial, review, and perspectives on open problems. *arXiv preprint arXiv:2005.01643*, 2020.
- [30] Aviral Kumar, Aurick Zhou, George Tucker, and Sergey Levine. Conservative Q-learning for offline reinforcement learning. In *Advances in Neural Information Processing Systems*, volume 33, pages 1179–1191, 2020.
- [31] Yan Duan, John Schulman, Xi Chen, Peter L Bartlett, Ilya Sutskever, and Pieter Abbeel. RL²: Fast reinforcement learning via slow reinforcement learning. *arXiv preprint arXiv:1611.02779*, 2016.
- [32] Zhi Wang, Chunlin Chen, and Daoyi Dong. Lifelong incremental reinforcement learning with online Bayesian inference. *IEEE Transactions on Neural Networks and Learning Systems*, 33(8):4003–4016, 2022.

- [33] Zhi Wang, Chunlin Chen, and Daoyi Dong. A Dirichlet process mixture of robust task models for scalable lifelong reinforcement learning. *IEEE Transactions on Cybernetics*, 53(12):7509–7520, 2023.
- [34] Chelsea Finn, Pieter Abbeel, and Sergey Levine. Model-agnostic meta-learning for fast adaptation of deep networks. In *Proceedings of International conference on machine learning*, pages 1126–1135, 2017.
- [35] Kate Rakelly, Aurick Zhou, Chelsea Finn, Sergey Levine, and Deirdre Quillen. Efficient off-policy meta-reinforcement learning via probabilistic context variables. In *Proceedings of International Conference on Machine Learning*, pages 5331–5340, 2019.
- [36] Luisa Zintgraf, Sebastian Schulze, Cong Lu, Leo Feng, Maximilian Igl, Kyriacos Shiarlis, Yarin Gal, Katja Hofmann, and Shimon Whiteson. VariBAD: Variational Bayes-adaptive deep RL via meta-learning. *The Journal of Machine Learning Research*, 22(1):13198–13236, 2021.
- [37] Lanqing Li, Hai Zhang, Xinyu Zhang, Shatong Zhu, Junqiao Zhao, and Pheng-Ann Heng. Towards an information theoretic framework of context-based offline meta-reinforcement learning. In *Advances in Neural Information Processing Systems*, 2024.
- [38] Sili Huang, Jifeng Hu, Hechang Chen, Lichao Sun, and Bo Yang. In-context decision transformer: Reinforcement learning via hierarchical chain-of-thought. In *Proceedings of International Conference on Machine Learning*, 2024.
- [39] Jake Grigsby, Linxi Fan, and Yuke Zhu. AMAGO: Scalable in-context reinforcement learning for adaptive agents. In *Proceedings of International Conference on Learning Representations*, 2024.
- [40] Wenhao Wu, Fuhong Liu, Haoru Li, Zican Hu, Daoyi Dong, Chunlin Chen, and Zhi Wang. Mixture-of-experts meets in-context reinforcement learning. In *Advances in Neural Information Processing Systems*, 2025.
- [41] Jinmei Liu, Fuhong Liu, Jianye Hao, Bo Wang, Huaxiong Li, Chunlin Chen, and Zhi Wang. Scalable in-context Q-learning. *arXiv preprint arXiv:2506.01299*, 2025.
- [42] Michael Laskin, Luyu Wang, Junhyuk Oh, Emilio Parisotto, Stephen Spencer, et al. In-context reinforcement learning with algorithm distillation. In *Proceedings of International Conference on Learning Representations*, 2023.
- [43] Jonathan Lee, Annie Xie, Aldo Pacchiano, Yash Chandak, Chelsea Finn, Ofir Nachum, and Emma Brunskill. Supervised pretraining can learn in-context reinforcement learning. *Advances in Neural Information Processing Systems*, 36, 2023.
- [44] Minae Kwon, Sang Michael Xie, Kalesha Bullard, and Dorsa Sadigh. Reward design with language models. In *Proceedings of International Conference on Learning Representations*, 2023.
- [45] Boyuan Wang, Yun Qu, Yuhang Jiang, Jianzhun Shao, Chang Liu, Wenming Yang, and Xi-angyang Ji. LLM-empowered state representation for reinforcement learning. In *Proceedings of International Conference on Machine Learning*, 2024.
- [46] Yuexiang Zhai, Hao Bai, Zipeng Lin, Jiayi Pan, Shengbang Tong, Yifei Zhou, Alane Suhr, Saining Xie, Yann LeCun, Yi Ma, et al. Fine-tuning large vision-language models as decision-making agents via reinforcement learning. In *Advances in Neural Information Processing Systems*, 2024.
- [47] Ruizhe Shi, Yuyao Liu, Yanjie Ze, Simon Shaolei Du, and Huazhe Xu. Unleashing the power of pre-trained language models for offline reinforcement learning. In *Proceedings of International Conference on Learning Representations*, 2024.
- [48] Zican Hu, Shilin Zhang, Yafu Li, Jianhao Yan, Xuyang Hu, Leyang Cui, Xiaoye Qu, Chunlin Chen, Yu Cheng, and Zhi Wang. Diversity-incentivized exploration for versatile reasoning. *arXiv preprint arXiv:2509.26209*, 2025.

- [49] Anthony Brohan, Noah Brown, Justice Carbajal, Yevgen Chebotar, Xi Chen, Krzysztof Choro-manski, Tianli Ding, Danny Driess, Avinava Dubey, Chelsea Finn, et al. RT-2: Vision-language-action models transfer web knowledge to robotic control. *arXiv preprint arXiv:2307.15818*, 2023.
- [50] Moo Jin Kim, Karl Pertsch, Siddharth Karamcheti, Ted Xiao, Ashwin Balakrishna, et al. OpenVLA: An open-source vision-language-action model. In *Proceedings of Conference on Robot Learning*, 2024.
- [51] Physical Intelligence, Kevin Black, Noah Brown, James Darpinian, Karan Dhabalia, et al. $\pi_{0.5}$: a vision-language-action model with open-world generalization. *arXiv preprint arXiv:2504.16054*, 2025.
- [52] Felix Hill, Sona Mokra, Nathaniel Wong, and Tim Harley. Human instruction-following with deep reinforcement learning via transfer-learning from text. *arXiv preprint arXiv:2005.09382*, 2020.
- [53] Muhannad Alomari, Paul Duckworth, David Hogg, and Anthony Cohn. Natural language acquisition and grounding for embodied robotic systems. In *Proceedings of AAAI Conference on Artificial Intelligence*, volume 31, 2017.
- [54] Ahmed Akakzia, Cédric Colas, Pierre-Yves Oudeyer, Mohamed Chetouani, and Olivier Sigaud. Grounding language to autonomously-acquired skills via goal generation. In *Proceedings of International Conference on Learning Representations*, 2021.
- [55] Jing-Cheng Pang, Xin-Yu Yang, Si-Hang Yang, Xiong-Hui Chen, and Yang Yu. Natural language instruction-following with task-related language development and translation. In *Advances in Neural Information Processing Systems*, volume 36, 2023.
- [56] Colin Raffel, Noam Shazeer, Adam Roberts, Katherine Lee, Sharan Narang, Michael Matena, Yanqi Zhou, Wei Li, and Peter J Liu. Exploring the limits of transfer learning with a unified text-to-text transformer. *Journal of Machine Learning Research*, 21(140):1–67, 2020.
- [57] Edward J Hu, Phillip Wallis, Zeyuan Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, Weizhu Chen, et al. LoRA: Low-rank adaptation of large language models. In *Proceedings of International Conference on Learning Representations*, 2022.
- [58] Anurag Ajay, Yilun Du, Abhi Gupta, Joshua Tenenbaum, Tommi Jaakkola, and Pulkit Agrawal. Is conditional generative modeling all you need for decision-making? In *Proceedings of International Conference on Learning Representations*, 2023.
- [59] Lili Chen, Kevin Lu, Aravind Rajeswaran, Kimin Lee, Aditya Grover, Misha Laskin, Pieter Abbeel, Aravind Srinivas, and Igor Mordatch. Decision transformer: Reinforcement learning via sequence modeling. In *Advances in Neural Information Processing Systems*, volume 34, pages 15084–15097, 2021.
- [60] Chitwan Saharia, William Chan, Saurabh Saxena, Lala Li, Jay Whang, et al. Photorealistic text-to-image diffusion models with deep language understanding. In *Advances in Neural Information Processing Systems*, volume 35, pages 36479–36494, 2022.
- [61] Lvmin Zhang, Anyi Rao, and Maneesh Agrawala. Adding conditional control to text-to-image diffusion models. In *Proceedings of IEEE/CVF International Conference on Computer Vision*, pages 3836–3847, 2023.
- [62] Tuomas Haarnoja, Aurick Zhou, Pieter Abbeel, and Sergey Levine. Soft actor-critic: Off-policy maximum entropy deep reinforcement learning with a stochastic actor. In *Proceedings of International Conference on Machine Learning*, pages 1861–1870, 2018.
- [63] Justin Fu, Aviral Kumar, Ofir Nachum, George Tucker, and Sergey Levine. D4RL: Datasets for deep data-driven reinforcement learning. *arXiv preprint arXiv:2004.07219*, 2020.
- [64] Jacob Devlin Ming-Wei Chang Kenton and Lee Kristina Toutanova. BERT: Pre-training of deep bidirectional transformers for language understanding. In *Proceedings of the North American Chapter of the Association for Computational Linguistics - Human Language Technologies*, volume 1, 2019.

- [65] Tianhe Yu, Deirdre Quillen, Zhanpeng He, Ryan Julian, Karol Hausman, Chelsea Finn, and Sergey Levine. Meta-World: A benchmark and evaluation for multi-task and meta reinforcement learning. In *Proceedings of Conference on Robot Learning*, pages 1094–1100, 2020.
- [66] Fei Ni, Jianye Hao, Yao Mu, Yifu Yuan, Yan Zheng, Bin Wang, and Zhixuan Liang. MetaDiffuser: Diffusion model as conditional planner for offline meta-RL. In *Proceedings of the 40th International Conference on Machine Learning*, pages 26087–26105, 2023.
- [67] William Peebles and Saining Xie. Scalable diffusion models with transformers. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 4195–4205, 2023.
- [68] Dan Hendrycks and Kevin Gimpel. Gaussian error linear units (GELUS). *arXiv preprint arXiv:1606.08415*, 2016.
- [69] Carmelo Sferrazza, Dun-Ming Huang, Xingyu Lin, Youngwoon Lee, and Pieter Abbeel. Humanoidbench: Simulated humanoid benchmark for whole-body locomotion and manipulation. *arXiv preprint arXiv:2403.10506*, 2024.

Appendix

A	Algorithm Pseudocodes	17
B	Evaluation Environments and Dataset Construction	17
B.1	The Details of Environments	17
B.2	The Details of Datasets	19
C	Baseline Methods	20
D	Implementation Details of T2DA	21
D.1	Network Architecture and Hyperparameters	21
D.2	Computation	22
E	Analysis of Parameter-Efficient Fine-tuning	23
F	Analysis of Training Pipeline	23
G	Meta-RL Baselines with Goal Information	24
H	Performance on Unstructured Language Instructions	24
I	Experimental Results on Humanoid	24
J	Comparison to More OMRL Baselines	25
K	Limitations	25

A Algorithm Pseudocodes

Based on the implementations in Sec. 3, this appendix gives the brief procedures of each component in T2DA. First, Algorithm 1 presents the pre-training of the generalized world model. Based on the dynamics-aware decision embedding encoded by Algorithm 1, Algorithm 2 shows the contrastive language-decision pre-training that bridges the semantic gap between text and decision, and aligns text embeddings to comprehend the environment dynamics of decision tasks.

Then, Algorithm 3 and Algorithm 4 present the model training processes of Text-to-Decision Diffuser and Text-to-Decision Transformer, respectively, where the generalist policy is conditioned on the text embeddings fine-tuned in Algorithm 2. Finally, Algorithm 5 and Algorithm 6 show zero-shot evaluations on test tasks, where the agent can directly realize text-to-decision generation according to textual instructions at hand.

Algorithm 1: Pre-training the generalized world model

Input: Offline datasets $\mathcal{D}_{\text{train}}$; Batch size m
Trajectory encoder ϕ ; Reward decoder $\mathcal{R}_\varphi$ and Transition decoder $\mathcal{T}_\varphi$

```

1 for each iteration do
2   for  $i = 1, \dots, m$  do
3     Sample a task  $M_k$  and obtain the corresponding dataset  $\mathcal{D}_k$  from  $\mathcal{D}_{\text{train}}$ 
4     Sample two trajectories from  $\mathcal{D}_k$  as  $\tau = (s_0, a_0, r_0, s_1, a_1, r_1, \dots)$  and
        $\tau^* = (s_0^*, a_0^*, r_0^*, s_1^*, a_1^*, r_1^*, \dots)$ 
5     Obtain the decision embedding by feeding  $\tau$  to the trajectory encoder as  $z = \phi(\tau)$ 
6     for  $t = 0, 1, \dots$  do
7       Compute the predicted reward  $\hat{r}_t^* = \mathcal{R}_\varphi(s_t^*, a_t^*; z)$  and next state
          $\hat{s}_{t+1}^* = \mathcal{T}_\varphi(s_t^*, a_t^*; z)$  on trajectory  $\tau^*$ 
8     end
9   end
10  Update  $\phi$  and  $\varphi$  jointly using loss
       
$$\mathcal{L}(\phi, \varphi) = \frac{1}{m} \sum_i \sum_t \left[ (r_t^* - \mathcal{R}_\phi(s_t^*, a_t^*; z))^2 + (s_{t+1}^* - \mathcal{T}_\varphi(s_t^*, a_t^*; z))^2 \right]$$

11 end
```

Algorithm 2: Contrastive Language-Decision Pre-training

Input: Offline datasets $\mathcal{D}_{\text{train}}$; Batch size N
Pre-trained trajectory encoder ϕ ; Text encoder ψ

```

1 Freeze the trajectory encoder  $\phi$  pre-trained in Algorithm 1
2 Initialize text encoder  $\psi$  from popular pre-trained language models like CLIP or T5
3 for each iteration do
4   Sample a batch of  $N$  (trajectory  $\tau$ , textual task description  $l$ ) pairs from  $\mathcal{D}_{\text{train}}$ 
5   Compute the symmetric similarities between the two modalities,  $\text{sim}(\tau, l)$  and  $\text{sim}(l, \tau)$ , as in
     Eq. (3)
6   Derive the decision-to-text and text-to-decision similarity scores,  $p(\tau)$  and  $p(l)$ , as in Eq. (4)
7   Update the text encoder  $\psi$  using LoRA by minimizing the symmetric cross-entropy loss  $\mathcal{L}(\psi)$ 
     as in Eq. (5)
8 end
```

B Evaluation Environments and Dataset Construction

B.1 The Details of Environments

We evaluate T2DA and all the baselines on three classical benchmarks that are widely adopted to assess the generalization capacity of RL algorithms, containing the following four environments:

- **Point-Robot** is a 2D navigation environment where an agent starts from a fixed origin and navigates to a target location g . The agent receives its coordinate position as the observation

Algorithm 3: Model Training of Text-to-Decision Diffuser

Input: Offline datasets $\mathcal{D}_{\text{train}} = \sum_k (\mathcal{D}_k, l_k)$; Pre-trained text encoder ψ
Noise model ϵ_θ ; Batch size N

```
1 Freeze the text encoder  $\psi$  pre-trained in Algorithm 2
2 while not converged do
3   Sample a batch of  $N$  (trajectory  $x_0(\tau^k)$ , language supervision  $l_k$ ) pairs from different tasks
   in  $\mathcal{D}_{\text{train}}$ 
4    $c \sim \text{Uniform}(\{1, \dots, C\})$ 
5    $\epsilon \sim \mathcal{N}(0, I)$ 
6    $x_c(\tau_{s_t}^k) \leftarrow s_t^k$  // Constrain the first state of the plan
7   Take gradient descent step on  $\nabla_\theta \sum_{k=1}^N \left\| \epsilon - \epsilon_\theta \left( x_c(\tau^k), \hat{R}(\tau), \psi(l_k), c \right) \right\|^2$ 
8 end
```

Algorithm 4: Model Training of Text-to-Decision Transformer

Input: Offline datasets $\mathcal{D}_{\text{train}} = \sum_k (\mathcal{D}_k, l_k)$; Pre-trained text encoder ψ
Causal transformer F_θ ; Batch size N

```
1 Freeze the text encoder  $\psi$  pre-trained in Algorithm 2
2 while not converged do
3   Sample a batch of  $N$  (trajectory  $\tau_k$ , language supervision  $l_k$ ) pairs from different tasks in
    $\mathcal{D}_{\text{train}}$ 
4   Get a batch of inputs  $\mathcal{B} = \{\tau_k^+\}_{k=1}^N$ , where
    $\tau_k^+ = (\psi(l_k), \hat{R}_{t-H+1}^k, s_{t-H+1}^k, a_{t-H+1}^k, \dots, \hat{R}_t^k, s_t^k, a_t^k)$ 
5    $a^{\text{pred}} = F_\theta(\tau_k^+), \forall \tau_k^+ \in \mathcal{B}$ 
6   Minimize loss  $\mathcal{L}(\theta) = \frac{1}{N} \sum_{\tau_k^+ \in \mathcal{B}} (a - a^{\text{pred}})^2$ 
7 end
```

Algorithm 5: Zero-Shot Evaluation of Text-to-Decision Diffuser

Input: Language instruction l_{new} ; Pre-trained text encoder ψ
Trained noise model ϵ_θ ; Diffusion steps C ; Guidance scale w

```
1 Obtain the text embedding from language instruction as  $\psi(l_{\text{new}})$ 
2 while not done do
3   Observe state  $s_t$ ; Initialize  $x_C(\tau) \sim \mathcal{N}(0, \alpha I)$ 
4   for  $c = C, \dots, 1$  do
5      $x_c(\tau_{s_t}) \leftarrow s_t$  // Constrain the first state of the plan
6      $\hat{\epsilon} \leftarrow \epsilon_\theta(x_c(\tau), c) + w \left[ \epsilon_\theta(x_c(\tau), \hat{R}(\tau), \psi(l_{\text{new}}), c) - \epsilon_\theta(x_c(\tau), c) \right]$  //
     Classifier-free guidance
7      $(\mu_{c-1}, \Sigma_{c-1}) \leftarrow \text{Denoise}(x_c(\tau), \hat{\epsilon})$ 
8      $x_{c-1}(\tau) \sim \mathcal{N}(\mu_{c-1}, \alpha \Sigma_{c-1})$ 
9   end
10  Execute first action of plan as  $a_t \leftarrow x_0(\tau_{a_t})$ 
11   $t \leftarrow t + 1$ 
12 end
```

and outputs actions in the range of $[-0.1, 0.1]^2$, representing displacement in the X - and Y -axis directions. The reward is the negative Euclidean distance to the goal: $r_t = -\|s_t - g\|_2$, where s_t denotes the current position. The maximal episode step is set to 20. Tasks differ in the target location g that is randomly sampled within a unit square of $[-0.5, 0.5]^2$. Each task is captioned by a text description as "Please navigate to the goal position of g ".

- **Cheetah-Vel** is a multi-task MuJoCo environment where a planar cheetah robot aims to run a target velocity v_g along the X -axis. The reward function combines a quadratic control cost and a

Algorithm 6: Zero-Shot Evaluation of Text-to-Decision Transformer

Input: Language instruction l_{new} ; Pre-trained text encoder ψ
 Trained casual transformer F_θ ; Target return G^*

- 1 Obtain the text embedding from language instruction as $\psi(l_{\text{new}})$
- 2 Initialize desired return $\hat{R} = G^*$
- 3 **for** each timestep t **do**
- 4 Observe state s_t
- 5 Obtain the H -step history trajectory as $\tau = (\hat{R}_{t-H+1}, s_{t-H+1}, a_{t-H+1}, \dots, \hat{R}_t, s_t)$
- 6 Augment the trajectory with the text embedding as $\tau^+ = (\psi(l_{\text{new}}), \tau)$
- 7 Get action $a_t = F_\theta(\tau^+)$
- 8 Step env. and get feedback $s, a, r, \hat{R} \leftarrow \hat{R} - r$
- 9 Append $(\hat{R}, s, a)$ to recent history τ
- 10 **end**

velocity matching term as $r_t = -0.05||a_t||_2 - |v_t - v_g|$, where a_t and v_t represent the action and current velocity, respectively. The maximal episode step is set to 200. Tasks differ in the target velocity v_g that is randomly sampled from a uniform distribution as $v_g \sim \text{Uniform}[0.075, 3.0]$. Each task is captioned by a text description as "Please run at the target velocity of v_g ".

- **Ant-Dir** is also a multi-task MuJoCo environment where a 3D ant robot navigates toward a target direction θ . The primary reward component measures directional alignment as $r_{\text{forward}} = v_x \cos \theta + v_y \sin \theta$, where (v_x, v_y) represents the horizontal and vertical velocities. The final reward incorporates additional terms for the control cost, contact penalty, and survival bonus. The maximal episode step is set to 200. Tasks differ in the goal direction θ that is uniformly sampled from the full space as $\theta \sim \text{Uniform}[0, 2\pi]$. Each task is captioned by a text description as "Please walk toward the target direction of θ ".
- **Meta-World** comprises 50 robotic manipulation tasks where a robotic arm interacts with various objects on a tabletop. Each task's reward function (ranging from 0 to 10) combines multiple components for fundamental behaviors (such as reaching, grasping, and placing), with 10 indicating successful completion. The maximal episode step is set to 200. Table 4 illustrates the detailed description of the training and test tasks. More details can be found in [65].

For each domain of Point-Robot, Cheetah-Vel, and Ant-Dir, we sample 50 tasks in total and split them into 45 training tasks and 5 test tasks. For Meta-World, we use 18 training tasks and 4 test tasks as detailed as shown in Table 4.

B.2 The Details of Datasets

We employ the soft actor-critic (SAC) algorithm [62] to independently train a policy for each task and save policy checkpoints at different training steps. The hyperparameters used for SAC training across all environments are listed in Table 5. We consider three types of datasets as

- **Medium:** We load the checkpoint of a medium policy that achieves approximately half the performance of expert policies, and use the medium policy to generate a number of trajectories to construct the Medium dataset.
- **Expert:** We load the checkpoint of an expert policy that achieves the highest return after the convergence of training, and use the expert policy to collect several trajectories to construct the Expert dataset.
- **Mixed:** We load all saved checkpoints of policies at different training steps, and use these policies to generate a diverse set of trajectories to construct the Mixed dataset.

Specifically, we collect 200 trajectories for each kind of dataset in all evaluated environments.

Table 4: Details of the training and test tasks of Meta-World.

Training Task	Description
faucet-close	rotate faucet clockwise
door-lock	lock door by rotating clockwise
door-unlock	unlock door by rotating counter-clockwise
window-close	push and close window
window-open	push and open window
coffee-button	push button on coffee machine
drawer-open	open drawer
door-open	open door with revolving joint
button-press	press button
button-press-topdown	press button from top
button-press-topdown-wall	bypass wall and press button from top
button-press-wall	bypass wall and press button
handle-press	press handle down
handle-pull	pull handle up
plate-slide-back	slide plate back
plate-slide-side	slide plate side
plate-slide	slide plate
plate-slide-back-side	slide plate back side
Test Task	Description
faucet-open	rotate faucet counter-clockwise
drawer-close	push and close drawer
reach-wall	bypass wall and reach goal
handle-press-side	press handle down sideways
handle-pull-side	Pull a handle up sideways
hand-insert	Insert the gripper into a hole

Table 5: Hyperparameters of SAC used to collect multi-task offline datasets.

Environments	Training steps	Warmup steps	Save frequency	Learning rate	Soft update	Discount factor	Entropy ratio
Point-Robot	2000	100	40	3e-4	0.005	0.99	0.2
Cheetah-Vel	500000	2000	10000	3e-4	0.005	0.99	0.2
Ant-Dir	500000	2000	10000	3e-4	0.005	0.99	0.2
Meta-World	1000000	5000	100	1e-3	0.005	0.99	0.2

C Baseline Methods

This section gives details of the competitive baselines, including context-based offline meta-RL (OMRL), in-context RL, and language-conditioned policy learning approaches. These baselines are thoughtfully selected to cover the main studies that tackle the generalization problem of RL algorithms. The baselines are introduced as follows:

- **CSRO** [13], **Context Shift Reduction for OMRL**. It addresses the context shift problem in offline meta-RL through two key components: 1) During meta-training, it employs a max-min mutual information representation learning mechanism to minimize mutual information between the task representation and the behavior policy while maximizing mutual information between the task representation and the task information; 2) During meta-test, it introduces a non-prior context collection strategy to first randomly explore the environment and then gradually update the task representation.
- **UNICORN** [37], **Unified Information Theoretic Framework of Context-Based Offline Meta-Reinforcement Learning**. It provides a general framework to unify several context-based offline meta-RL algorithms by proving that they optimize different bounds of mutual information between

the task variable M and latent representation Z . Based on this theoretical insight, it proposes a supervised and a self-supervised implementation of the mutual information derivation $I(Z; M)$.

- **AD** [42], **Algorithm Distillation**. It distills RL algorithms into neural networks by modeling their training histories with a causal sequence model. Using a dataset of learning histories generated by a source RL algorithm, AD trains a causal transformer to autoregressively predict actions given their preceding learning histories as the context, enabling in-context policy improvement without updating network parameters.
- **BC-Z** [25], **Behavior Cloning Z**. It is an imitation learning system designed for zero-shot generalization to novel vision-based manipulation tasks. It flexibly conditions the policy on pre-trained embeddings of natural language or video of human-performing tasks. To align BC-Z to the benchmarks investigated in this paper, we replace the video data with RL datasets $\mathcal{D} = \sum_i (s_i, a_i, r_i, s'_i)$ and modify its video encoder to a trajectory encoder ϕ identical to that used in T2DA. This adjustment for a fair comparison results in the optimization objective: $\min \sum_{\text{task } k} \sum_{(s,a) \sim \mathcal{D}} [-\log \pi(a|s, \psi(l_k)) + D_{\cos}(\psi(l_k), \phi(\tau_k))]$, where $D_{\cos}$ denotes the cosine distance.
- **BAKU** [27], an efficient transformer architecture for multi-task policy learning that improves upon prior work in offline imitation learning. It carefully integrates multiple key components: observation trunks for multi-modal fusion, action chunking for smoother control, multi-sensory observation processing, and modular action heads for flexible prediction. The model is optimized through behavior cloning with the objective: $\min \sum_{\text{task } k} \sum_{(o,a) \sim \mathcal{D}} |a - \pi(a|o, T_k)|^2$, where o represents multi-modal observations and T_k denotes the task instruction. For a fair comparison, we adjust BAKU to the language-conditioned zero-shot setting.
- **Prompt-DT** [12], **Prompt-based Decision Transformer**. It is a DT-based offline meta-RL method that leverages the transformer architecture and prompting mechanisms to enable few-shot adaptation. It introduces trajectory prompts, which consist of short demonstration segments from the target task, to encode task-specific context and guide policy generation. Without requiring any fine-tuning on unseen tasks, Prompt-DT achieves strong few-shot performance through expert demonstrations in constructed prompts.
- **MetaDiffuser** [66], a diffusion-based offline meta-RL method that considers the generalization problem as conditional trajectory generation task with contextual representation. It trains a context conditioned diffusion model to generate task-oriented trajectories for planning. To further enhance the dynamics consistency of the generated trajectories while encouraging trajectories to achieve high returns, it introduces a dual-guided module in the sampling process of the diffusion model.
- **Meta-DT** [14], **Meta Decision Transformer**. It harnesses the sequential modeling capability of the transformer architecture and enables robust task representation learning through a disentangled world model, thereby achieving efficient generalization in offline meta-RL. It learns a compact task representation via context-aware world model, used as a contextual condition to guide task-oriented sequence generation. Complementary to the task representation, it selects a self-guided trajectory segment from online interactions to further exploit the architectural inductive bias. Thereby, Meta-DT exhibits strong few- and zero-shot performance without expert demonstrations.

For a fair comparison, all baselines are adjusted to an aligned zero-shot setting, where no warm-start data is available before policy evaluation. All these baselines can only use samples generated by the trained meta-policy during evaluation to infer task representations.

D Implementation Details of T2DA

D.1 Network Architecture and Hyperparameters

- **World Model**. We implement the generalized world model using simple architecture: a trajectory encoder and two decoders for the reward and state transition prediction. Specifically, the trajectory encoder consists of a transformer encoder with 6 layers, where each layer contains a multi-head self-attention module (8 heads) and a feed-forward network with ReLU activation. States, actions, and rewards in the raw trajectory are first tokenized into 256-dimensional vectors through separate linear layers. These embeddings are concatenated and processed by the transformer encoder to produce a 256-dimensional decision embedding z via mean pooling and linear projection. Two MLPs serve as decoders: a reward decoder that predicts the instant reward r given the tuple

$(s, a, s'; z)$, and a state transition decoder that predicts the next state s' given the tuple $(s, a, r; z)$. Both MLPs utilize two hidden layers of 256 dimensions each.

- **Text Encoder.** We initialize text encoders from pre-trained CLIP [16], BERT [64], and T5 [56] in this paper. Specifically, we use `openai/clip-vit-base-patch32`, `google-bert/bert-base-uncased`, and `google-t5/t5-small`, respectively.
- **Diffusion Architecture of T2DA.** We implement the diffusion architecture of T2DA-D based on the decision diffuser [58] framework². Following DiT [67], we replace U-Net backbone with a transformer to represent the noise model ϵ_θ . Each block employs a multi-head self-attention module followed by a feed-forward network with GELU activation [68], using adaptive layer normalization (adaLN) for conditioning. The model processes state-action trajectories through linear embeddings, which are combined with positional embeddings derived from RL timesteps. The diffusion time embeddings are encoded via sinusoidal embeddings processed through an MLP. For conditioning, text embeddings and returns-to-go are concatenated with diffusion time embeddings. The detailed configurations and hyperparameters of T2DA-D are presented in Table 7.
- **Transformer Architecture of T2DA.** We implement the transformer architecture of T2DA-T based on the decision transformer [59] framework³. Specifically, we employ modality-specific embeddings for states, actions, returns-to-go, and timesteps. The timestep embeddings serve as positional encodings and are added to each token embedding. These tokens are concatenated with text embedding of task description and fed into a GPT architecture which predicts actions autoregressively using a causal self-attention mask. The detailed configurations and hyperparameters of T2DA-T are presented in Table 6.

Table 6: Configurations and hyperparameters in the training process of T2DA-T.

Configuration	Point-Robot	Cheetah-Vel	Ant-Dir	Meta-World
layers num	3	3	3	3
attention head num	1	1	1	1
embedding dim	128	128	128	128
horizon	20	200	40	50
training steps	100000	70000	100000	100000
learning rate	1e-4	2e-4	1e-4	1e-4

Table 7: Configurations and hyperparameters in the training process of T2DA-D.

Configuration	Point-Robot	Cheetah-Vel	Ant-Dir	Meta-World
DiT layers num	4	4	8	4
DiT attention head num	8	8	8	8
embedding dim	128	128	256	128
horizon	10	50	40	50
diffusion steps	20	20	20	20
training steps	100000	70000	100000	100000
learning rate	5e-5	1e-4	1e-4	1e-4

D.2 Computation

We train our models on one Nvidia RTX4080 GPU with the Intel Core i9-10900X CPU and 256G RAM. Dynamic-aware decision embedding and constrastive language-decision pre-training cost about 0.5 hours. The generalist policy training cost about 0.5-3 hours depending on the complexity of the environment.

²<https://github.com/anuragajay/decision-diffuser>

³<https://github.com/kzl/decision-transformer>

E Analysis of Parameter-Efficient Fine-tuning

As stated in Sec. 3.3, we employ LoRA, a classical parameter-efficient fine-tuning method, to fine-tune the text encoder at a lightweight cost. As shown in Table 8, the number of trainable parameters is substantially smaller than full-parameter fine-tuning (less than 1%), enabling much more efficient computation and memory usage. To investigate whether a limited number of trainable parameters constrains the capacity for knowledge alignment, we evaluate T2DA’s performance by comparing LoRA fine-tuning with full-parameter fine-tuning during contrastive language-decision pre-training. Figure 8 illustrates the results on representative environments with the text encoder initialized from CLIP. Obviously, the two fine-tuning approaches yield nearly identical performance, highlighting T2DA’s superior parameter efficiency of T2DA without sacrificing effectiveness.

Table 8: The number and memory usage of parameters for LoRA fine-tuning and full-parameter fine-tuning.

Method	Trainable parameters (Memory)	Total parameters (Memory)	Percentage
LoRA Tuning	589,824 (2.25MB)	63,755,776 (243.21MB)	0.93%
Full Tuning	63,165,952 (240.96MB)	63,165,952 (240.96MB)	100%

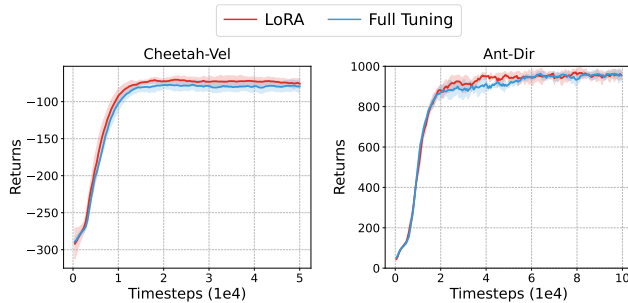


Figure 8: Test return curves of T2DA using LoRA-tuned and fully-tuned CLIP text encoders. With less than 1% parameter count, T2DA using LoRA fine-tuning yields nearly identical performance compared to full-parameter fine-tuning, highlighting T2DA’s superior parameter efficiency without sacrificing effectiveness.

F Analysis of Training Pipeline

T2DA separately trains three key components: (a) decision embedding, (b) knowledge alignment, and (c) generalist policy. To demonstrate the effectiveness of this modular approach, we conducted ablation studies comparing our method against two alternatives: jointly training (a)+(b), and jointly training all three components (a)+(b)+(c) simultaneously. As shown in Table 9, our separate training pipeline achieves substantial performance gains of 20%-56% across different environments. This modular architecture not only provides a more stable training process but also offers superior scalability, as individual components can be easily substituted or upgraded independently without affecting the entire system.

Table 9: Ablation study results on modular architecture. Best results in **bold**.

Environment	T2DA	(a)+(b)	(a)+(b)+(c)
Point-Robot	-7.2	-8.9	-16.3
Meta-World	1274.5	1008	817.3

G Meta-RL Baselines with Goal Information

For comprehensive evaluation, we enhanced meta-RL baselines by incorporating goal information into their observation space. This approach provides a standardized framework where all compared methods have access to equivalent task information.

The results in Table 10 demonstrate T2DA’s performance advantages over the goal-enhanced baselines, further validating the effectiveness of leveraging natural language supervision. This goal-integrated setting was maintained consistently for all meta-RL baseline comparisons throughout the paper.

Table 10: Performance comparison with goal-enhanced meta-RL baselines. Best results in **bold**.

Environment	T2DA	CSRO (w/ goal)	UNICORN (w/ goal)
Cheetah-Vel	-60.4 ± 1.8	-84.0 ± 4.3	-82.7 ± 9.2
Ant-Dir	970.3 ± 8.7	384.4 ± 30.1	427.1 ± 25.3

H Performance on Unstructured Language Instructions

This section presents an evaluation of the method’s generalization capabilities when confronted with less structured, ambiguous, or noisy natural language instructions. Experiments were conducted where tasks were captioned using two additional types of textual descriptions.

Task descriptions were reconstructed in two distinct linguistic styles:

- **Noisy style:** *"Um, there's this location I think... somewhere around $\{round(number=task[0], ndigits=2)\}$ -ish in the x direction? And maybe $\{round(number=task[1], ndigits=2)\}$ or so up?"*
- **Conversational style:** *"Could you move over to the spot that's roughly $\{round(number=task[0], ndigits=2)\}$ units to the right and $\{round(number=task[1], ndigits=2)\}$ units up? Thanks!"*

Results demonstrate T2DA’s consistent performance across these different textual description formats:

Table 11: Performance comparison across different instruction styles.

Environment	T2DA	T2DA (Noisy)	T2DA (conversational)
Point-Robot	-7.2 ± 0.1	-7.6 ± 0.1	-8.0 ± 0.1

These findings highlight the robustness of leveraging natural language supervision, as the language representations effectively capture task-relevant information and successfully extrapolate meta-level knowledge across tasks presented with varying textual descriptions.

I Experimental Results on Humanoid

This section presents experimental results on stimulated humanoid robot benchmark [69]. Table 12 shows the converged performance of T2DA and baselines using Mixed datasets under aligned zero-shot setting. The results demonstrate that T2DA consistently outperforms baselines in this more challenging environment.

Table 12: Performance comparison of different methods on Humanoid. Best results in **bold** and second best underlined.

Method	T2DA-T	T2DA-D	BC-Z	DT	DD	CSRO	UNICORN	AD
Returns	610.4	<u>575.4</u>	531.7	520.3	381.3	520.3	505.6	427.4

J Comparison to More OMRL Baselines

In this section, we compare T2DA to three additional OMRL baselines: Meta-DT [14], MetaDiffuser [66], and Prompt-DT [12]. More details about these baselines can be found in Appendix C. The results in Table 13 demonstrate T2DA’s consistent superiority over these OMRL baselines.

Table 13: Performance comparison with OMRL baselines. Best results in **bold**.

Environment	T2DA-T	T2DA-D	Meta-DT	MetaDiffuser	Prompt-DT
Point-Robot	-7.2 ± 0.1	-8.4 ± 0.2	-8.4 ± 0.4	-11.7 ± 0.6	-14.3 ± 1.7
Cheetah-Vel	-70.3 ± 3.8	-60.4 ± 1.8	-78.2 ± 5.3	-86.5 ± 4.8	-94.6 ± 8.1
Ant-Dir	970.3 ± 8.7	570.6 ± 13.1	922.1 ± 11.9	858.5 ± 42.5	895.2 ± 17.3

K Limitations

We tackle the offline meta-RL challenge via leveraging the supervision from natural language. Our scalable implementations demonstrate that the generalist RL agent can understand and act upon natural language instructions in a zero-shot manner. While our approach uses relatively lightweight datasets compared to large models, it demonstrates promising results. An essential step of future work is to implement on vast datasets with diversified domains and deploy the text-to-decision paradigm on real robots toward embodied intelligence.

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: Please refer to the Abstract and Section 1: Introduction.

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: Please refer to Section 5: Conclusions, Limitations, and Future Work, and Appendix K.

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: Our work does not include theory assumptions and proofs.

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: Please refer to Appendix A for algorithm pseudocodes. Appendix B for the details of environments and dataset construction, and Appendix D for the implementation details and hyperparameters of our method. Our source is included in the supplementary materials.

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: Please refer to Appendix A, Appendix B and Appendix D for details of our method and dataset construction. We also provide our source code in the supplementary materials

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: Please refer to Section 4: Experiments, Appendix B, Appendix D.

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: Please refer to Section 4: Experiments.

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: Please refer to Appendix D and Appendix E.

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: Please refer to Appendix A, Appendix B and Appendix D for details of our implementation and dataset construction. We also provide our source code in the supplementary materials.

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: [Yes]

Justification: Please refer to Section 5: Conclusions, Limitations, and Future Work.

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: Our work poses no such risks.

Guidelines:

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

12. Licenses for existing assets

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

Answer: [NA]

Justification: We construct the datasets used in this work ourselves, and not use existing assets, please refer to Appendix B for details.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.

- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. **New assets**

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

Answer: [NA]

Justification: Our work does not release new assets.

Guidelines:

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

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

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

Answer: [NA]

Justification: Our work does not involve crowdsourcing and research with human subjects.

Guidelines:

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

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

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

Answer: [NA]

Justification: Our work does not involve crowdsourcing and research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.

- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. **Declaration of LLM usage**

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

Answer: [NA]

Justification: The core method development in our work does not involve LLMs as any important, original, or non-standard components.

Guidelines:

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