
Generative Trajectory Stitching through Diffusion Composition

Yunhao Luo¹ Utkarsh A. Mishra¹ Yilun Du^{2,†} Danfei Xu^{1,†}

¹ Georgia Tech ² Harvard University

Abstract

Effective trajectory stitching for long-horizon planning is a significant challenge in robotic decision-making. While diffusion models have shown promise in planning, they are limited to solving tasks similar to those seen in their training data. We propose CompDiffuser, a novel generative approach that can solve new tasks by learning to compositionally stitch together shorter trajectory chunks from previously seen tasks. Our key insight is modeling the trajectory distribution by subdividing it into overlapping chunks and learning their conditional relationships through a single bidirectional diffusion model. This allows information to propagate between segments during generation, ensuring physically consistent connections. We conduct experiments on benchmark tasks of various difficulties, covering different environment sizes, agent state dimension, trajectory types, training data quality, and show that CompDiffuser significantly outperforms existing methods. Project website at <https://comp-diffuser.github.io/>.

1 Introduction

Generative models have demonstrated remarkable capabilities in modeling complex distributions across domains like images, videos, and 3D shapes. In robot planning, these models offer a promising approach by modeling distributions over plan sequences, which allows amortizing the computational cost of traditional search and optimization methods. This effectively transforms planning into sampling likely solutions given start and goal conditions. Recent works like Diffuser [26] and Decision Diffuser [1] have shown how diffusion models can learn to generate entire plans for long-horizon robotics tasks. However, exhaustively modeling joint distributions over entire plan sequences for all possible start and goal states remains extremely sample-inefficient, as it requires collecting long-horizon plan data covering all possible combinations of initial states and goals.

The concept of trajectory stitching [88] from Reinforcement Learning literature [62] presents a potential solution by combining chunks of different trajectories to create new, potentially better policies. The methods work by identifying high-reward trajectory chunks and stitching them together at states where they overlap or are similar enough, creating composite trajectories that can inform better policy learning. This effectively enables compositional generalization since collecting long consecutive trajectories is costly, and these short chunks can be flexibly assembled to complete new tasks. The key challenge lies in finding appropriate stitching points where trajectories can be combined while maintaining dynamic consistency and feasibility. Our goal is to enable generative planners to solve long-horizon tasks without requiring long-horizon training data, while retaining their ability to generate physically feasible, goal-directed plans.

We propose a novel diffusion-based approach, Compositional Diffuser (CompDiffuser), that enables effective trajectory stitching through goal-conditioned causal trajectory generation.

[†]Equal advising.

Our key insight is that we can model the trajectory distribution compositionally by subdividing it into distributions of overlapping chunks and learning their conditional relationships. Rather than learning separate models for each chunk, we train a single diffusion model that can generate trajectory chunks conditioned on neighboring chunks' states (Figure 1). This allows information to propagate bidirectionally during the reverse diffusion process as illustrated in Figure 2: each chunk's generation is influenced by both past and future chunks. This architecture naturally enables both parallel generation of chunks and causal autoregressive generation, each with different trade-offs in computational cost and planning quality.

We conduct extensive experiments across benchmark tasks of varying difficulty levels, including different environment sizes (from simple U-mazes to complex giant mazes), agent state dimensions (from 2D point agents to 50D humanoid robots), trajectory types (from maze navigation trajectories to ball dribbling trajectories), and training data quality (from clean demonstrations to noisy exploration data). Our results demonstrate that CompDiffuser significantly outperforms multiple imitation learning and offline reinforcement learning baselines across all settings. We show that our approach can effectively solve long-horizon tasks while maintaining plan feasibility and goal-reaching behavior. We validate the importance of our key technical components including the bidirectional conditioning mechanism, the autoregressive sampling process, and the flexible replanning capability.

In summary, the key contributions of this work are:

- A noisy-sample conditioned diffusion planning framework that enables learning compositional trajectory distributions by decomposing the trajectory generation procedure into a sequence of segments each generated by a separate diffusion denoising process.
- A compositional goal-conditioned trajectory planning method that uses bidirectional information propagation during denoising to maintain physical consistency between trajectory chunks.
- A set of empirical results showing significant improvements over existing methods across multiple trajectory stitching benchmarks, with detailed analysis of model capabilities and limitations.

2 Related Work

Diffusion Models For Planning. Many works have studied the applications of diffusion models [58, 24] for generative planning [26, 1, 54, 72, 22, 64, 42, 86, 7]. Diffusion planning has been widely applied in various fields, such as motion planning [5, 43], procedure planning [67], task planning [76, 16], autonomous driving [75, 37, 68], reasoning [79], and reward learning [49]. Many techniques have also been combined with diffusion planning, including hierarchical planning [35, 8, 45], self-evolving planner [36], preference alignment [10], tree branch-pruning [17], refining [31], replanning [84], uncertainty-aware planning [60], equivariance [4]. However, these works are usually constrained to plan within similar horizons as training data. Our work instead proposes a compositional diffusion planning approach that generalizes to much longer horizon via generative trajectory stitching.

Trajectory Stitching. A flurry of works have explored the trajectory stitching problem given offline data. One typical category of solution is based on data augmentation or goal relabeling along with various techniques, such as generative models [33, 41, 28, 30, 34, 40], model-based approaches [6, 32, 85, 23], and clustering [21]. Other supervised learning based methods, such as sequence modeling [9, 25, 73, 87, 71], latent space learning [80, 66, 83], and learning with dynamic programming [18, 27, 74], have also demonstrated some extent of stitching capability. In this work, we propose a different trajectory stitching approach based on generative modeling, where the model only learns from plain short trajectory segments while is able to directly perform goal-conditioned trajectory stitching through test-time compositional generation.

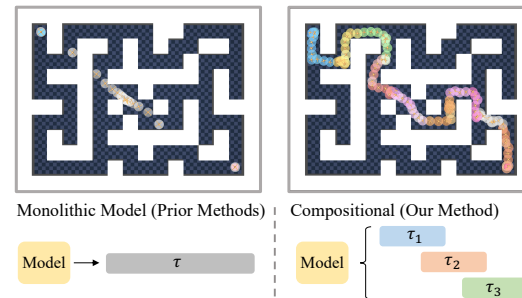


Figure 1: **Compositional Trajectory Generation.** CompDiffuser enables generative trajectory stitching through diffusion composition. Left: Monolithic generative planner fails to generalize to tasks of longer horizon and collapses to the maze center. Right: Our method successfully navigates the ant agent from start to goal by compositionally stitching together shorter trajectories.

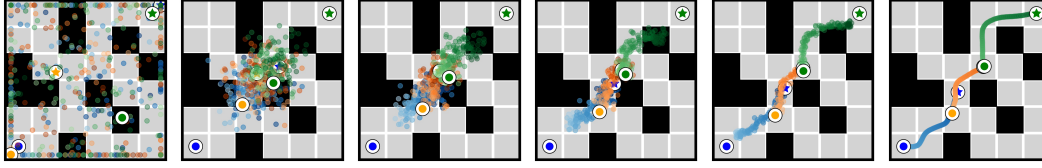


Figure 2: **Illustrating the Trajectory Stitching Process.** Given an unseen start (blue circle) and goal (green star), CompDiffuser generates a long-horizon plan by progressively denoising three trajectory chunks in parallel, with each chunk conditioning on its neighbors to ensure smooth transitions.

Compositional Generative Models. Compositional Generative Models [13, 11, 19, 12, 46, 3, 50, 63] are widely studied in various domains, including visual content generation [39, 11, 77, 82, 59, 57], human motion generation [56, 61, 81], traffic generation [38], robotic planning [78, 2, 43], and policy learning [69, 53]. Most existing work on compositionality focuses on sampling under the conjunction of several given conditions. While several studies [48, 47] have explored sequential compositional models, they are restricted to planning within a given skeleton and hence unable to generalize to longer sequences or new tasks. In contrast, we propose a compositional planning framework that scales to generating much longer sequences and completing new tasks via directly stitching short trajectory segments, without relying on pre-defined task-dependent skeletons.

3 Planning through Compositional Trajectory Generation

We aim to develop a generative planning framework that can generate long-horizon trajectories by composing multiple modular generative models. Our method, Compositional Diffuser (CompDiffuser), trains a single diffusion model on short-horizon trajectories. At inference time, given a start and goal, CompDiffuser runs parallel instances of this model to generate a sequence of overlapping trajectory segments, coordinating their denoising processes to ensure they smoothly connect into a coherent long-horizon plan. This approach enables us to stitch together short-horizon subsequences of training trajectories to form novel long-horizon trajectory plans.

3.1 Compositional Trajectory Modeling

Given a planning problem, consisting of a start state q_s and a goal state q_g , we formulate planning as sampling a trajectory τ from the probability distribution

$$[s^{1:T}, a^{1:T}] \sim p_\theta(\tau | q_s, q_g), \quad (1)$$

where $s^{1:T}$ corresponds to future states to reach the goal state q_g and $a^{1:T}$ corresponds to a set of future actions. To implement this sampling procedure, prior work [1, 26] learns a generative model $p(\tau)$ directly over previous trajectories τ in the environment. However, since the generative model is trained to model the density of previously seen trajectories, it is restricted to generating plans with start and goal that are similar to those seen in the past.

In this paper, we propose to model the generative model over trajectories $p_\theta(\tau | q_s, q_g)$ compositionally [12], where we subdivide trajectory τ into a set of K overlapping sub-chunks τ_k (Figure 1). We then represent the trajectory distribution as

$$p_\theta(\tau | q_s, q_g) \propto p_1(\tau_1 | q_s, \tau_2) p_K(\tau_K | \tau_{K-1}, q_g) \prod_{k=2}^{K-1} p_k(\tau_k | \tau_{k-1}, \tau_{k+1}). \quad (2)$$

In the above expression, each trajectory chunk τ_k is only dependent on nearby trajectory chunks τ_{k-1} and τ_{k+1} . This allows $p_\theta(\tau | q_s, q_g)$ to generate trajectory plans that significantly depart from previously seen trajectories, as long as intermediate trajectory chunks τ_k have been seen. Overall, the goal of our method is to enable long-horizon planning without long-horizon training data.

3.2 Training Compositional Trajectory Models

One approach to represent the composed probability distribution in Equation 2 is to directly learn separate generative models to represent each conditional probability distribution. However, sampling from the composed distribution is challenging, as each individual trajectory chunk depends on the

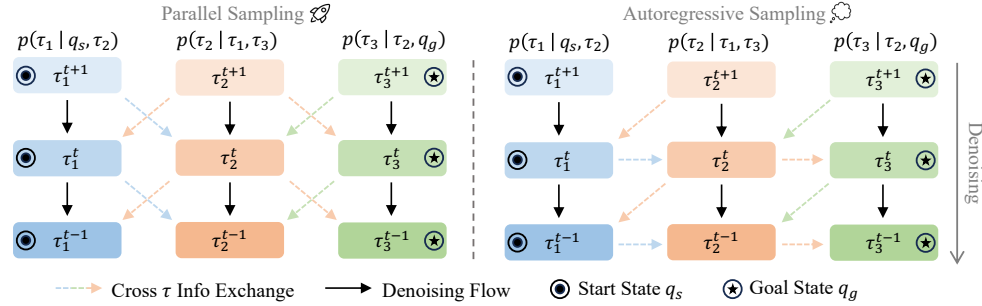


Figure 3: **Compositional Trajectory Planning: Parallel Sampling and Autoregressive Sampling.** We present an illustrative example of sampling three trajectories $\tau_{1:3}$ with the proposed compositional sampling methods. Dashed lines represent cross trajectory information exchange between adjacent trajectories and black lines represent the denoising flow of each trajectory. In parallel sampling, $\tau_{1:3}$ can be denoised concurrently; while in autoregressive sampling, denoising τ_k depends on the previous trajectory τ_{k-1} , e.g., the denoising of τ_2 depends on τ_1 (as shown in the blue horizontal dashed arrows). Additionally, start state q_s and goal state q_g conditioning are applied to the trajectories in the two ends, τ_1 and τ_3 , which enables goal-conditioned planning. Trajectories $\tau_{1:3}$ will be merged to form a longer plan τ_{comp} after the full diffusion denoising process.

values of neighboring chunks. As a result, to sample from the composed distribution, one would need a blocked Gibbs sampling procedure, where the value of each individual trajectory chunk is iteratively sampled given decoded values of neighboring trajectory chunks. This sampling procedure is slow, and passing information across chunks to form a consistent plan is challenging.

Information Propagation Through Noisy-Sample Conditioning. We propose a more efficient approach that addresses these challenges by leveraging the progressive denoising process of diffusion models. The key challenge in composing trajectories is ensuring feasible transitions between each pair of neighboring chunks, i.e., maintaining physical constraints and dynamic consistency at connection points where segments overlap. Our key insight is that we can achieve this by having trajectory segments guide each other’s generation during the diffusion process: as one segment takes shape through denoising, it helps shape its neighbors into compatible configurations.

We implement this insight using a diffusion model that generates trajectory chunks conditioning on their neighbors’ *noisy samples*. Given a dataset D of trajectories τ , we train a denoising network ϵ_θ to learn the trajectory distribution $p_\theta(\tau_k | \tau_{k-1}, \tau_{k+1})$ with the training objective

$$\mathcal{L}_{\text{nbr}} = \mathbb{E}_{\tau \in D, t, k} [\|\epsilon - \epsilon_\theta(\tau_k^t, t | \tau_{k-1}^t, \tau_{k+1}^t)\|^2], \quad (3)$$

where k identifies a trajectory segment, t is the noise level, and τ_k^t represents segment k corrupted with noise level t . Crucially, when denoising each segment, the network conditions on noisy versions of neighboring segments $\tau_{k-1}^t, \tau_{k+1}^t$ at the same noise level. This allows each segment to influence its neighbors’ denoising process, ensuring their final configurations are dynamically compatible. In addition, further training the network to condition on τ_{k-1}^{t-1} can enable autoregressive compositional sampling, which we will discuss in Section 3.3. In practice, we only need to condition on the small overlapping regions between consecutive trajectories, making the generation process efficient while maintaining consistency across connection points.

Representing Initial States and Goals. In addition, we further train the same denoising network to represent the distributions $p_\theta(\tau_1 | q_s, \tau_2)$ and $p_\theta(\tau_K | \tau_{K-1}, q_g)$. This corresponds to the training objective

$$\mathcal{L}_{\text{start}} = \mathbb{E}_{\tau \in D, t, k} [\|\epsilon - \epsilon_\theta(\tau_1^t, t | q_s, \tau_2^t)\|^2], \quad (4)$$

with an analogous objective for the conditioned goal state q_g . We train the same denoising network ϵ_θ with both conditioning. Please see Appendix E for implementation details. We provide an overview of our proposed training strategy in Algorithm 1.

3.3 Compositional Trajectory Planning

Our compositional framework enables flexible sampling strategies for generating long-horizon plans with Equation 2. The basic sampling process starts by initializing each trajectory chunk τ_k with Gaussian noise. Then, through iterative denoising, each chunk is denoised while being conditioned

Algorithm 1 Training CompDiffuser

```

1: Require: train dataset  $D$ , diffusion denoiser model  $\epsilon_\theta(\tau^t, t \mid \text{st\_cond}, \text{end\_cond})$ , number of training steps  $N$ , diffusion timestep  $T$ 
2: for  $i = 1 \rightarrow N$  do
3:    $\tau^0 \sim D$ , sample clean trajectory from dataset
4:    $\tau_{1:K}^0 \leftarrow$  divide  $\tau^0$  to  $K$  overlapping chunks
5:    $\tau_{1:K}^t \leftarrow$  add noise to  $\tau_{1:K}^0$ ,  $t \sim T$ , following DDPM
6:   # Objective for noisy chunk condition
7:    $\mathcal{L}_{\text{nbr}} = \|\epsilon - \epsilon_\theta(\tau_k^t, t \mid \tau_{k-1}^t, \tau_{k+1}^t)\|^2$ ,  $k \sim [2, K-1]$ 
8:   # Objective for start / end state condition
9:    $\mathcal{L}_{\text{start}} = \|\epsilon - \epsilon_\theta(\tau_1^t, t \mid q_s, \tau_2^t)\|^2$ ,  $q_s = \tau_1^0[0]$ 
10:   $\mathcal{L}_{\text{end}} = \|\epsilon - \epsilon_\theta(\tau_K^t, t \mid \tau_{K-1}^t, q_g)\|^2$ ,  $q_g = \tau_K^0[-1]$ 
11:   $\mathcal{L}_{\text{all}} = \mathcal{L}_{\text{nbr}} + \mathcal{L}_{\text{start}} + \mathcal{L}_{\text{end}}$ 
12:  Backprop to update  $\epsilon_\theta(\cdot)$  using  $\mathcal{L}_{\text{all}}$ 
13: end for
14: return  $\epsilon_\theta(\cdot)$ 

```

Algorithm 2 Autoregressive Trajectory Sampling

```

1: Models: trained diffusion denoiser model  $\epsilon_\theta(\tau, t \mid \text{st\_cond}, \text{g\_cond})$ 
2: Input: start state  $q_s$ , goal state  $q_g$ , number of composed trajectories  $K$ 
3: Initialize  $K$  trajectories  $\tau_{1:K} \sim \mathcal{N}(0, I)$ 
4: for  $t = T \rightarrow 1$  do
5:   # Denoise  $\tau_1$  conditioned on  $q_s$  and  $\tau_2^t$ 
6:    $\tau_1^{t-1} = \epsilon_\theta(\tau_1^t, t \mid q_s, \tau_2^t)$ 
7:   # Denoise intermediate trajectories  $\tau_2$  to  $\tau_{K-1}$ 
8:   for  $k = 2 \rightarrow K-1$  do
9:      $\tau_k^{t-1} = \epsilon_\theta(\tau_k^t, t \mid \tau_{k-1}^{t-1}, \tau_{k+1}^t)$ 
10:  end for
11:  # Denoise  $\tau_K$  conditioned on  $\tau_{K-1}^{t-1}$  and  $q_g$ 
12:   $\tau_K^{t-1} = \epsilon_\theta(\tau_K^t, t \mid \tau_{K-1}^{t-1}, q_g)$ 
13: end for
14:  $\tau_{\text{comp}} =$  Merge the denoised trajectories  $\tau_{1:K}^0$ 
15: return  $\tau_{\text{comp}}$ 

```

on its neighbors. This structure allows for different ways of coordinating the denoising process across trajectory chunks, each offering different tradeoffs between information propagation and computational efficiency. We present two sampling schemes (illustrated in Figure 3):

Parallel Sampling. Our first sampling approach conditions denoising on the values of the noisy adjacent trajectory chunks from the previous denoising timestep, where the update rule is

$$\tau_k^{t-1} = \alpha^t(\tau_k^t - \epsilon_\theta(\tau_k^t \mid \tau_{k-1}^t, \tau_{k+1}^t) + \beta^t \xi), \quad \xi \sim \mathcal{N}(0, 1), \quad (5)$$

where α^t and β^t are diffusion specific hyperparameters. This approach allows us to run denoising on each trajectory chunk in parallel, as each denoising update only requires the values of the adjacent trajectory chunks at a previous noise level. However, information propagation between the values of adjacent trajectory chunks is limited at each denoising timestep, as each trajectory chunk is denoised independently of the denoising updates of other trajectory chunks.

Autoregressive Sampling. To better couple the values of adjacent trajectory chunks, we propose to denoise each trajectory chunk autoregressively dependent on the values of neighboring chunks at each denoising timestep. In particular, we iteratively denoise each trajectory $\tau_{1:K}^t$ starting from the τ_1^t , and condition the denoising of τ_k^t on the previously decoded chunk τ_{k-1}^{t-1} at the current noise level $t-1$ and the future chunk τ_{k+1}^t at the previous noise level t , giving us the sampling equation

$$\tau_k^{t-1} = \alpha^t(\tau_k^t - \epsilon_\theta(\tau_k^t \mid \tau_{k-1}^{t-1}, \tau_{k+1}^t) + \beta^t \xi), \quad \xi \sim \mathcal{N}(0, 1). \quad (6)$$

This sequential generation process enables stronger coordination among the chunks since each chunk is conditioned on the less noisy version of its previous chunk. However, it requires generating chunks one at a time rather than simultaneously, making it computationally less efficient than parallel sampling. We compare the two sampling schemes in Table 6, where we empirically find autoregressive sampling leads to improved performance. Additionally, we provide sampling time comparison in Table 10. We use this autoregressive sampling procedure throughout the experiments in the paper and illustrate pseudocode for sampling in Algorithm 2. Given such final set of generated chunks $\tau_{1:K}$, we then merge the chunks together to construct a final trajectory τ_{comp} by applying exponential trajectory blending to areas where subchunks τ_k overlap (See Appendix E.2 for implementation details).

4 Experiments

In this section, our objective is to (1) validate that our method enforces coherent trajectory stitching on multiple benchmarks, with varying state space dimensions, task design, and training data collection policies (2) understand how planning with higher state dimensions, varying numbers of composed trajectories, different sampling schemes, and replanning affect the performance of the proposed method. Additional results are provided in Appendix B and C with failure analysis in Appendix D.

Baselines. We compare CompDiffuser with three categories of existing methods: (1) for generative planning methods, we include Decision Diffuser (DD) [1] for monolithic trajectory sampling

Env	Size	RvS	RvS (SA)	RvS (GA)	DT	DT (SA)	DT (GA)	DD	GSC	Ours
PointMaze [21]	U-Maze	17±7	97±5	76±5	17±5	65±4	54±4	0±0	100±0	100±0
	Medium	1±2	55±3	21±3	20±2	55±3	62±2	30±1	93±1	100±0
	Large	3±4	38±5	31±5	22±2	35±2	39±5	0±0	99±2	100±0
Env	Type	Size	GCBC	GCIVL	GCIQL	QRL	CRL	HIQL	GSC	Ours
PointMaze [51]	stitch	Medium	23 ±18	70 ±14	21 ±9	80 ±12	0 ±1	74 ±6	100±0	100±0
		Large	7 ±5	12 ±6	31 ±2	84 ±15	0 ±0	13 ±6	100±0	100±0
		Giant	0 ±0	0 ±0	0 ±0	50 ±8	0 ±0	0 ±0	29±3	68±3

Table 1: **Quantitative Results on PointMaze Stitching Datasets in Ghugare et al. [21] and OGBench [51].** We compare CompDiffuser to baselines of multiple categories, including diffusion, data augmentation, and offline reinforcement learning. SA and GA stand for state augmentation and goal augmentation respectively, as described in Ghugare et al. [21]. Our results are averaged over 5 seeds and standard deviations are shown after the $\pm$ sign.

and Generative Skill Chaining (GSC) [48] for compositional sampling; (2) for data augmentation based methods, we include stitching specific data augmentation [21] with RvS [14] and Decision Transformer (DT) [9]; (3) for offline reinforcement learning methods, we include goal-conditioned behavioral cloning (GCBC) [44, 20], goal-conditioned implicit V-learning (GCIVL) and Q-learning (GCIQL) [29], Quasimetric RL (QRL) [70], Contrastive RL (CRL) [15], and Hierarchical implicit Q-learning (HIQL) [52]. See Appendix F for more details of baselines.

Evaluation Setup. For each environment, we report the success rate over all evaluation episodes, where the success criterion is that the agent or target object is close to the goal within a small threshold. We evaluate all methods with 5 random seeds for each experiment and report the mean and standard deviation. Specifically, in Ghugare et al. [21] datasets, we evaluate on 2 tasks in U-Maze, 6 tasks in Medium, and 7 tasks in Large with 10 episodes per task; in OGBench [51], we evaluate on 5 tasks in each environment with 20 episodes per task. Each task is introduced in the respective papers and is defined by a base start and goal state that require trajectory stitching to complete. A random noise is added to the base start and goal state for each evaluation episode.

4.1 PointMaze

We present experiment results on two types of trajectory-stitching datasets in point maze environments, which features different dataset collection strategies. Our method is trained on short trajectories of x - y positions of the point agent and is able to directly generate much longer trajectories from start and goal by composing multiple trajectories (detailed numbers of the composed trajectories in each experiment are provided in Appendix Table 13). Note that our method only requires training one model and we can use the proposed compositional inference method to construct coherent long-horizon trajectories.

PointMaze [21]. Following the original framework, the training data are curated by dividing each environment (here, maze) into several small regions, and feasible trajectories constrained within their respective regions are sampled. Possible stitching between segments is facilitated by a small overlap (one block) between different regions, which can be used to join trajectories across regions. More dataset details are provided in Appendix A.1. We present the quantitative results on these datasets in Table 1, where we compare to goal-conditioned behavior cloning methods trained with data augmentation, Decision Diffuser, and GSC. Most baselines perform suboptimally, likely because they are unable to autonomously identify the small overlapping regions needed for trajectory stitching. In contrast, CompDiffuser successfully resolves all tasks across various maze sizes, demonstrating its ability to stitch trajectories even when the connecting regions are small.

PointMaze in OGBench [51]. In this particular setup, each trajectory in the training dataset is constrained to navigate no more than four blocks in the environment. The start and goal of each trajectory can be sampled from the entire environment provided that the travel distance between the start and goal is within four blocks. These trajectories are much shorter than the ones required for a feasible plan between a given start and goal at inference. More details about these datasets are provided in Appendix A.2. We present the quantitative results in Table 1, where we compared our method to GSC and multiple offline RL baselines following the benchmark established in [51]. We observe that while GSC is able to perform on par with CompDiffuser in Medium and Large mazes, it struggles as the planning horizon further increases, as illustrated in the qualitative results in Figure 4.

Env	Type	Size	GCBC	GCIVL	GCIQL	QRL	CRL	HIQL	GSC	Ours
antmaze	stitch	Medium	45 $\pm$ 11	44 $\pm$ 6	29 $\pm$ 6	59 $\pm$ 7	53 $\pm$ 6	94 $\pm$ 1	97 $\pm$ 2	96 $\pm$ 2
		Large	3 $\pm$ 3	18 $\pm$ 2	7 $\pm$ 2	18 $\pm$ 2	11 $\pm$ 2	67 $\pm$ 5	66 $\pm$ 2	86 $\pm$ 2
		Giant	0 $\pm$ 0	0 $\pm$ 0	0 $\pm$ 0	0 $\pm$ 0	0 $\pm$ 0	21 $\pm$ 2	20 $\pm$ 1	65 $\pm$ 3
	explore	Medium	2 $\pm$ 1	19 $\pm$ 3	13 $\pm$ 2	1 $\pm$ 1	3 $\pm$ 2	37 $\pm$ 10	90 $\pm$ 2	81 $\pm$ 2
		Large	0 $\pm$ 0	10 $\pm$ 3	0 $\pm$ 0	0 $\pm$ 0	0 $\pm$ 0	4 $\pm$ 5	21 $\pm$ 3	27 $\pm$ 1
	humanoid maze	Medium	29 $\pm$ 5	12 $\pm$ 2	12 $\pm$ 3	18 $\pm$ 2	71 $\pm$ 3	96 $\pm$ 4	92 $\pm$ 1	91 $\pm$ 1
		Large	6 $\pm$ 3	1 $\pm$ 1	0 $\pm$ 0	3 $\pm$ 1	6 $\pm$ 1	31 $\pm$ 3	70 $\pm$ 3	72 $\pm$ 3
		Giant	0 $\pm$ 0	0 $\pm$ 0	0 $\pm$ 0	0 $\pm$ 0	0 $\pm$ 0	12 $\pm$ 2	5 $\pm$ 1	67 $\pm$ 4

Table 2: **Quantitative Results on AntMaze and HumanoidMaze in OGBench.** We benchmark our method on the 5 test-time tasks defined in OGBench with 20 episodes per task. Our results are averaged over 5 seeds and standard deviations are shown after the $\pm$ sign.

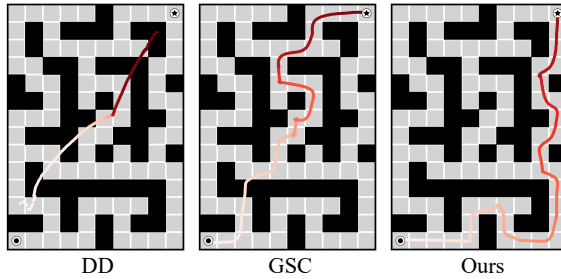


Figure 4: **Qualitative Comparison of DD, GSC and CompDiffuser on OGBench PointMaze Giant.** Effective bidirection information propagation enables CompDiffuser to successfully synthesize trajectories from start (bottom left) to goal (upper right), while other methods generate *o.o.d* trajectories that are disconnected with the start/goal or passing through walls. See Figure 10 for per-segment visualization.

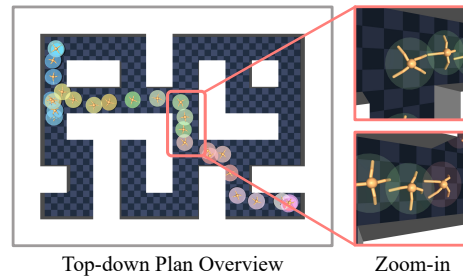


Figure 5: **Qualitative Results of Planning in High Dimension on OGBench AntMaze Large.** Original plan is sub-sampled for clearer view. Our method is able to synthesize plans of valid dynamics while reaching the goal position (bottom right). Note that our method is only trained on trajectory segments of much shorter length.

CompDiffuser significantly outperforms all baselines in Giant maze, demonstrating its efficacy in complex environments. Additional results are provided in Appendix C.1, C.2, and C.3.

4.2 High Dimension Tasks

We present the evaluation results of our method on various trajectory stitching tasks involving higher-dimensional state spaces within OGBench: AntMaze, HumanoidMaze, and AntSoccer.

AntMaze and HumanoidMaze. We conduct maze navigation experiments of multiple agents, ant and humanoid, using the pre-collected Stitch datasets provided in OGBench. The data collection strategy is identical to OGBench PointMaze, where each episode is constrained to travel at most 4 blocks, while at inference, a successful plan requires the agent to travel up to 30 blocks. We compare our method with the offline RL benchmark established in [51] along with the best-performing diffusion-based compositional stitching baseline GSC, as shown in Table 2. Both GSC and CompDiffuser generate plans in a planar x - y space while the agent follows the plans with a learned inverse dynamics model. We observe a pattern similar to the results in PointMaze, where CompDiffuser can consistently give high success rates as the planning horizon and complexity increase while other baselines start to collapse. To complete the study, we also conduct experiments where the full agent state is used for planning instead of the x - y space and provide the results in Section 4.3.

AntMaze with Low-Quality Data. We evaluate CompDiffuser on OGBench AntMaze with a different data collection strategy, Explore. These datasets consist of extremely low-quality yet high-coverage data, where the data collection policy contains a large amount of action noise and will randomly re-sample a new moving direction after every 10 steps (see Figure 9 in Appendix for qualitative examples). Hence, each demonstration episode typically oscillates within only 2-3 blocks. Our planner needs to learn from these clustered trajectories to construct coherent plans that reach goals in large spatial distances. We present the success rate of each method in Table 2.

AntSoccer. The AntSoccer environment in OGBench requires the ant agent to move a soccer ball to a designated goal in the environment, different from maze tasks that require the agent itself to

Env	Size	GCBC	GCIVL	GCIQL	QRL	CRL	HIQL	GSC (4D)	Ours (4D)	GSC (17D)	Ours (17D)
antsoccer	arena	34 $\pm$ 4	21 $\pm$ 3	5 $\pm$ 2	2 $\pm$ 1	2 $\pm$ 1	23 $\pm$ 2	41 $\pm$ 4	55 $\pm$ 6	65 $\pm$ 3	69 $\pm$ 3
stitch	medium	2 $\pm$ 1	1 $\pm$ 0	0 $\pm$ 0	0 $\pm$ 0	0 $\pm$ 0	8 $\pm$ 2	5 $\pm$ 2	13 $\pm$ 1	12 $\pm$ 2	17 $\pm$ 3

Table 3: **Quantitative Results on OGBench AntSoccer Stitch.** We evaluate two generative planners with different planning state dimensions: a 4D planner that operates on the x - y positions of the ant and the ball, and a 17D planner that additionally generates the 13 joint positions of the ant agent.

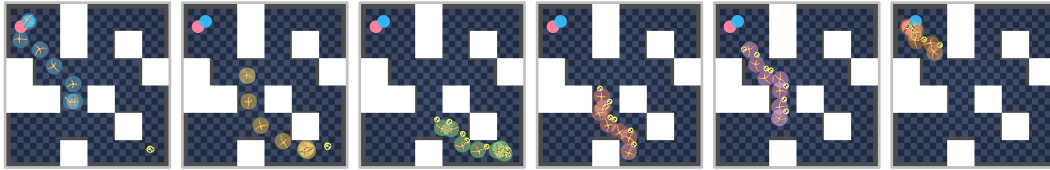


Figure 6: **Qualitative Plan generated by CompDiffuser in OGBench AntSoccer Medium Stitch.** The initial position of the ant is shown by the blue circle and the goal is to move the ball to the pink circle. We plot each individual trajectory $\tau_{1:6}$ separately (as shown from left to right) and mark the ball with a yellow circle for clearer view. These generated trajectories will be merged to form a long-horizon plan τ_{comp} . Note that our model is only trained on two different types of short trajectories: 1) the ant moves without dribbling the ball; 2) the ant moves while dribbling the ball. At test time, the proposed compositional sampling method can stitch these two types of trajectories and construct end-to-end plans of longer horizon to solve the more difficult task – first navigate to the ball from the far side and then dribble the ball to the goal position.

reach the goal. OGBench provides two categories of trajectories in AntSoccer Stitch datasets: (1) ant navigates in the maze without the ball and (2) ant moves and dribbles the ball in the maze. The planner must stitch trajectories of these two skills to complete the goal-reaching objective, because the ant must reach the ball first and then dribble it to the given goal location (see Figure 6 and Figure 14).

We present experimental results in two distinct maze configurations, Arena and Medium, in Table 3 relative to the benchmarking provided in OGBench. Specifically, as an ablation study, we evaluate with two planner state space configurations: (1) a 4D planner consisting of the x - y location of the ant and the ball; (2) a 17D planner consisting of the x - y location of the ant and the ball along with all the joint positions of the ant. We observe that our compositional method outperforms all the baselines in both configurations. Notably, the 17D variant demonstrates slightly higher success rates, likely due to the ability of joint positions to provide more fine-grained information for ball dribbling.

4.3 Ablation Studies

Planning in High Dimension Space. We report experiment results where CompDiffuser synthesizes trajectories in state space of higher dimension. We compare the success rates of our method planning in dimension of 2D, 15D, and 29D in Table 4, and present qualitative plans in Figure 5 and Figure 13. Specifically, the planner operates in the x - y space for 2D, x - y with the joint positions for 15D, x - y with both the joint positions and velocities for 29D. Similar to Section 4.2, we train an MLP inverse dynamics model that takes in the current observation and a goal of the planner dimension to predict an action for the agent to execute.

Our method performs consistently and achieves near-optimal success rates across all planning dimensions in AntMaze Medium. In AntMaze Large and Giant, the success rates decrease when planning in 15D and 29D, which is probably due to the increasing complexity of the trajectory modeling, since the joint positions and velocities of a moving ant are highly dynamic and the tasks require planning hundreds of consecutive future states to reach the goal.

Different Numbers of Composed Trajectories. We study the effect of varying the numbers of trajectories to be composed K . To better study the planning performance with respect to K , we use the challenging PointMaze-Giant-Stitch in OGBench as the testbed. As shown in Figure 7, our method obtains consistent performance when composing 7 to 12 trajectories. We observe that the optimal K is around 9 and 10, which also corresponds to a natural path length to reach the goal. Qualitatively, decreasing K will result in a sparser trajectory while increasing K will cause the final trajectory traveling back and forth to consume the redundant states (See Figure 10).

Size	HIQL	Ours (2D)	Ours (15D)	Ours (29D)
Medium	94 $\pm$ 1	96 $\pm$ 2	95 $\pm$ 0	97$\pm$2
Large	67 $\pm$ 5	86$\pm$2	66 $\pm$ 5	66 $\pm$ 5
Giant	21 $\pm$ 2	65$\pm$3	41 $\pm$ 3	28 $\pm$ 4

Table 4: **Quantitative Results of Different Planning Dimensions on AntMaze Stitch Datasets in OGBench.** Our method compositionally constructs feasible plans that reach long-distance goals while modeling complex environment dynamics, such as the positions and velocities of the agent’s joints.

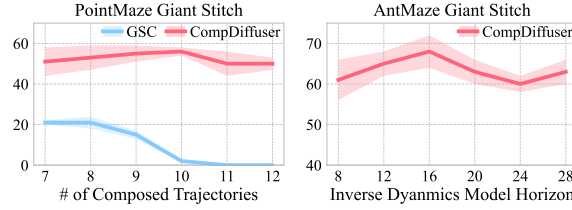


Figure 7: **Quantitative Results of Different Setup.** Left: success rate versus different numbers of composed trajectories K in OGBench PointMaze Giant (w/o replan). Right: success rate versus different inverse dynamics model horizons in OGBench AntMaze Giant.

Env	Size	QRL	HIQL	Ours w/o Replan	Ours w/ Replan
PointMaze	Medium	80 $\pm$ 12	74 $\pm$ 6	100$\pm$0	100$\pm$0
	Large	84 $\pm$ 15	13 $\pm$ 6	100$\pm$0	100$\pm$0
	Giant	50 $\pm$ 8	0 $\pm$ 0	53$\pm$6	68$\pm$3
AntMaze	Medium	59 $\pm$ 7	94$\pm$1	92$\pm$2	96$\pm$2
	Large	18 $\pm$ 2	67 $\pm$ 5	76$\pm$2	86$\pm$2
	Giant	0 $\pm$ 0	21 $\pm$ 2	27$\pm$4	65$\pm$3

Table 5: **Quantitative Results of CompDiffuser with and without Replanning.** We report the success rates on OGBench PointMaze and AntMaze Stitch datasets. For w/o replan, CompDiffuser only synthesizes one trajectory and executes the trajectory in a close-loop manner; for w/ replan, CompDiffuser will synthesize a new trajectory if the agent loses track of the current plan.

Env	Size	Replan	Parallel	AR
PointMaze	Large	✓	100$\pm$0	100$\pm$0
	Giant	✗	45 $\pm$ 1	53$\pm$6
	Giant	✓	66 $\pm$ 2	68$\pm$3
AntMaze	Large	✓	84 $\pm$ 5	86$\pm$2
	Giant	✗	18 $\pm$ 4	27$\pm$4
	Giant	✓	48 $\pm$ 1	65$\pm$3

Table 6: **Quantitative Comparison of two Sampling Schemes: Parallel and Autoregressive (AR).** Parallel sampling performs on par with AR sampling in the easier Large maze, while AR sampling can generate trajectories of higher quality when constructing longer plans (i.e., composing more trajectories), likely due to its causal denoising strategy.

Replanning with CompDiffuser. Our method can also flexibly replan during a rollout, which enables the agent to recover from failure, such as when the agent fails to track the planned trajectory due to sub-optimal inverse dynamic actions. In practice, we replan if the distance between the agent’s current observation and the synthesized subgoal is larger than a threshold. Please see Appendix E.3 for details. In Table 5, we present ablation studies of CompDiffuser with and without replanning in PointMaze and AntMaze Stitch datasets in OGBench. CompDiffuser outperforms the best performing baselines QRL and HIQL even without replanning in 5 out of 6 tasks. We also observe that w/ and w/o replanning yield similar performance in maze size Medium and Large, while offering significant performance boost in the more complex Giant maze.

Parallel vs. Autoregressive Sampling. We compare the performance of the two proposed compositional sampling schemes, parallel and autoregressive, as presented in Table 6. Autoregressive sampling consistently outperforms the parallel sampling across various tasks in plan quality, showing that the causal information flow, where each trajectory chunk is conditioned on the already-denoised (less noisy) version of previous chunk, leads to more coherent and physically consistent transitions between chunks. We include additional discussion and sampling time comparison in Appendix B.3.

5 Discussion and Conclusion

Limitations. Stitching short trajectories to solve for unseen longer-horizon plans is a challenging problem. While our method excels in our empirical evaluation, there still remain a number of future venues of research. When composing a large sequence of trajectories, the error accumulation in the long chain of bidirectional information propagation might lead to infeasible plans. This issue can be potentially mitigated by using domain/task specific rejection sampling techniques to select from multiple candidate plans or leveraging more expensive MCMC sampling. Moreover, the optimal number of the test-time composed chunks K is task-dependent. Our future research will explore ways to automatically identify a suitable number of chunks K , potentially by incrementally increasing the number of chunks dependent on the quality of the generated plan.

Conclusion. We introduce CompDiffuser, a generative trajectory stitching method that leverages the compositionality of diffusion models. We introduce a noise-conditioned score function formulation that helps in performing autoregressive sampling of multiple short-horizon trajectory diffusion models and eventually stitching them to form a longer-horizon goal-conditioned trajectory. Our method demonstrates effective trajectory stitching capabilities as evident from the extensive experiments on tasks of various difficulty, including different environment sizes, planning state dimensions, trajectory types, and training data quality.

Acknowledgments

We would like to thank Zilai Zeng for helpful early discussion on trajectory stitching and feedback of the manuscript.

References

- [1] Ajay Anurag, Du Yilun, Gupta Abhi, Tenenbaum Joshua B, Jaakkola Tommi S, Agrawal Pulkit. Is Conditional Generative Modeling all you need for Decision Making? // The Eleventh International Conference on Learning Representations. 2022. [1](#), [2](#), [3](#), [5](#), [33](#), [35](#)
- [2] Ajay Anurag, Han Seungwook, Du Yilun, Li Shuang, Gupta Abhi, Jaakkola Tommi S., Tenenbaum Joshua B., Kaelbling Leslie Pack, Srivastava Akash, Agrawal Pulkit. Compositional Foundation Models for Hierarchical Planning // Thirty-seventh Conference on Neural Information Processing Systems. 2023. [3](#)
- [3] Bradley Arwen, Nakkiran Preetum, Berthelot David, Thornton James, Susskind Joshua M. Mechanisms of Projective Composition of Diffusion Models // arXiv preprint arXiv:2502.04549. 2025. [3](#)
- [4] Brehmer Johann, Bose Joey, De Haan Pim, Cohen Taco S. EDGI: Equivariant diffusion for planning with embodied agents // Advances in Neural Information Processing Systems. 2023. 36. 63818–63834. [2](#)
- [5] Carvalho Joao, Le An T, Baierl Mark, Koert Dorothea, Peters Jan. Motion planning diffusion: Learning and planning of robot motions with diffusion models // 2023 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS). 2023. 1916–1923. [2](#)
- [6] Char Ian, Mehta Viraj, Villafior Adam, Dolan John M, Schneider Jeff. Bats: Best action trajectory stitching // arXiv preprint arXiv:2204.12026. 2022. [2](#)
- [7] Chen Boyuan, Martí Monsó Diego, Du Yilun, Simchowicz Max, Tedrake Russ, Sitzmann Vincent. Diffusion forcing: Next-token prediction meets full-sequence diffusion // Advances in Neural Information Processing Systems. 2024. 37. 24081–24125. [2](#)
- [8] Chen Chang, Deng Fei, Kawaguchi Kenji, Gulcehre Caglar, Ahn Sungjin. Simple hierarchical planning with diffusion // arXiv preprint arXiv:2401.02644. 2024. [2](#)
- [9] Chen Lili, Lu Kevin, Rajeswaran Aravind, Lee Kimin, Grover Aditya, Laskin Misha, Abbeel Pieter, Srinivas Aravind, Mordatch Igor. Decision transformer: Reinforcement learning via sequence modeling // Advances in neural information processing systems. 2021. 34. 15084–15097. [2](#), [6](#), [35](#)
- [10] Dong Zibin, Yuan Yifu, HAO Jianye, Ni Fei, Mu Yao, ZHENG YAN, Hu Yujing, Lv Tangjie, Fan Changjie, Hu Zhipeng. AlignDiff: Aligning Diverse Human Preferences via Behavior-Customisable Diffusion Model // The Twelfth International Conference on Learning Representations. 2024. [2](#)
- [11] Du Yilun, Durkan Conor, Strudel Robin, Tenenbaum Joshua B, Dieleman Sander, Fergus Rob, Sohl-Dickstein Jascha, Doucet Arnaud, Grathwohl Will Sussman. Reduce, reuse, recycle: Compositional generation with energy-based diffusion models and mcmc // International conference on machine learning. 2023. 8489–8510. [3](#)

- [12] *Du Yilun, Kaelbling Leslie*. Compositional generative modeling: A single model is not all you need // arXiv preprint arXiv:2402.01103. 2024. [3](#)
- [13] *Du Yilun, Li Shuang, Mordatch Igor*. Compositional visual generation with energy based models // Advances in Neural Information Processing Systems. 2020. 33. 6637–6647. [3](#)
- [14] *Emmons Scott, Eysenbach Benjamin, Kostrikov Ilya, Levine Sergey*. Rvs: What is essential for offline rl via supervised learning? // arXiv preprint arXiv:2112.10751. 2021. [6](#), [35](#)
- [15] *Eysenbach Benjamin, Zhang Tianjun, Levine Sergey, Salakhutdinov Russ R*. Contrastive learning as goal-conditioned reinforcement learning // Advances in Neural Information Processing Systems. 2022. 35. 35603–35620. [6](#), [35](#)
- [16] *Fang Xiaolin, Garrett Caelan Reed, Eppner Clemens, Lozano-Pérez Tomás, Kaelbling Leslie Pack, Fox Dieter*. Dimsam: Diffusion models as samplers for task and motion planning under partial observability // 2024 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS). 2024. 1412–1419. [2](#)
- [17] *Feng Lang, Gu Pengjie, An Bo, Pan Gang*. Resisting Stochastic Risks in Diffusion Planners with the Trajectory Aggregation Tree // Forty-first International Conference on Machine Learning. 2024. [2](#)
- [18] *Gao Chen-Xiao, Wu Chenyang, Cao Mingjun, Kong Rui, Zhang Zongzhang, Yu Yang*. Act: Empowering decision transformer with dynamic programming via advantage conditioning // Proceedings of the AAAI Conference on Artificial Intelligence. 38, 11. 2024. 12127–12135. [2](#)
- [19] *Garipov Timur, De Peuter Sebastiaan, Yang Ge, Garg Vikas, Kaski Samuel, Jaakkola Tommi*. Compositional sculpting of iterative generative processes // arXiv preprint arXiv:2309.16115. 2023. [3](#)
- [20] *Ghosh Dibya, Gupta Abhishek, Reddy Ashwin, Fu Justin, Devin Coline, Eysenbach Benjamin, Levine Sergey*. Learning to reach goals via iterated supervised learning // arXiv preprint arXiv:1912.06088. 2019. [6](#), [35](#)
- [21] *Ghugare Raj, Geist Matthieu, Berseth Glen, Eysenbach Benjamin*. Closing the Gap between TD Learning and Supervised Learning-A Generalisation Point of View. // The Twelfth International Conference on Learning Representations. 2024. [2](#), [6](#), [23](#), [24](#), [34](#), [35](#)
- [22] *He Haoran, Bai Chenjia, Xu Kang, Yang Zhuoran, Zhang Weinan, Wang Dong, Zhao Bin, Li Xuelong*. Diffusion model is an effective planner and data synthesizer for multi-task reinforcement learning // Advances in neural information processing systems. 2023. 36. 64896–64917. [2](#)
- [23] *Hepburn Charles Alexander, Montana Giovanni*. Model-based Trajectory Stitching for Improved Offline Reinforcement Learning // 3rd Offline RL Workshop: Offline RL as a “Launchpad”. 2022. [2](#)
- [24] *Ho Jonathan, Jain Ajay, Abbeel Pieter*. Denoising diffusion probabilistic models // Advances in neural information processing systems. 2020. 33. 6840–6851. [2](#)
- [25] *Huang Xingshuai, Wu Di, Boulet Benoit*. DRDT3: Diffusion-Refined Decision Test-Time Training Model // arXiv preprint arXiv:2501.06718. 2025. [2](#)
- [26] *Janner Michael, Du Yilun, Tenenbaum Joshua, Levine Sergey*. Planning with Diffusion for Flexible Behavior Synthesis // International Conference on Machine Learning. 2022. 9902–9915. [1](#), [2](#), [3](#)
- [27] *Kim Jeonghye, Lee Suyoung, Kim Woojun, Sung Youngchul*. Adaptive Q -Aid for Conditional Supervised Learning in Offline Reinforcement Learning // Advances in Neural Information Processing Systems. 2024. 37. 87104–87135. [2](#)
- [28] *Kim Sungyoon, Choi Yunseon, Matsunaga Daiki E, Kim Kee-Eung*. Stitching Sub-trajectories with Conditional Diffusion Model for Goal-Conditioned Offline RL // Proceedings of the AAAI Conference on Artificial Intelligence. 38, 12. 2024. 13160–13167. [2](#)

- [29] *Kostrikov Ilya, Nair Ashvin, Levine Sergey*. Offline reinforcement learning with implicit q-learning // arXiv preprint arXiv:2110.06169. 2021. 6, 35
- [30] *Lee Jaewoo, Yun Sujin, Yun Taeyoung, Park Jinkyoo*. GTA: Generative Trajectory Augmentation with Guidance for Offline Reinforcement Learning // The Thirty-eighth Annual Conference on Neural Information Processing Systems. 2024. 2
- [31] *Lee Kyowoon, Kim Seongun, Choi Jaesik*. Refining diffusion planner for reliable behavior synthesis by automatic detection of infeasible plans // Advances in Neural Information Processing Systems. 2024. 36. 2
- [32] *Lei Xing, Zhang Xuetao, Wang Donglin*. MGDA: Model-based Goal Data Augmentation for Offline Goal-conditioned Weighted Supervised Learning // arXiv preprint arXiv:2412.11410. 2024. 2
- [33] *Li Guanghe, Shan Yixiang, Zhu Zhengbang, Long Ting, Zhang Weinan*. DiffStitch: Boosting Offline Reinforcement Learning with Diffusion-based Trajectory Stitching // Forty-first International Conference on Machine Learning. 2024. 2
- [34] *Li Shangzhe, Zhang Xinhua*. Augmenting Offline Reinforcement Learning with State-only Interactions // arXiv preprint arXiv:2402.00807. 2024. 2
- [35] *Li Wenhao, Wang Xiangfeng, Jin Bo, Zha Hongyuan*. Hierarchical diffusion for offline decision making // International Conference on Machine Learning. 2023. 20035–20064. 2
- [36] *Liang Zhixuan, Mu Yao, Ding Mingyu, Ni Fei, Tomizuka Masayoshi, Luo Ping*. AdaptDiffuser: Diffusion Models as Adaptive Self-evolving Planners // International Conference on Machine Learning. 2023. 20725–20745. 2
- [37] *Liao Bencheng, Chen Shaoyu, Yin Haoran, Jiang Bo, Wang Cheng, Yan Sixu, Zhang Xinbang, Li Xiangyu, Zhang Ying, Zhang Qian, others*. DiffusionDrive: Truncated Diffusion Model for End-to-End Autonomous Driving // arXiv preprint arXiv:2411.15139. 2024. 2
- [38] *Lin Haohong, Huang Xin, Phan-Minh Tung, Hayden David S, Zhang Huan, Zhao Ding, Srinivasa Siddhartha, Wolff Eric M, Chen Hongge*. Causal Composition Diffusion Model for Closed-loop Traffic Generation // arXiv preprint arXiv:2412.17920. 2024. 3
- [39] *Liu Nan, Li Shuang, Du Yilun, Torralba Antonio, Tenenbaum Joshua B*. Compositional visual generation with composable diffusion models // European Conference on Computer Vision. 2022. 423–439. 3
- [40] *Liu Zhihong, Qian Long, Liu Zeyang, Wan Lipeng, Chen Xingyu, Lan Xuguang*. Enhancing Decision Transformer with Diffusion-Based Trajectory Branch Generation // arXiv preprint arXiv:2411.11327. 2024. 2
- [41] *Lu Cong, Ball Philip, Teh Yee Whye, Parker-Holder Jack*. Synthetic experience replay // Advances in Neural Information Processing Systems. 2024. 36. 2
- [42] *Lu Haoifei, Han Dongqi, Shen Yifei, Li Dongsheng*. What Makes a Good Diffusion Planner for Decision Making? // The Thirteenth International Conference on Learning Representations. 2025. 2
- [43] *Luo Yunhao, Sun Chen, Tenenbaum Joshua B, Du Yilun*. Potential Based Diffusion Motion Planning // International Conference on Machine Learning. 2024. 33486–33510. 2, 3
- [44] *Lynch Corey, Khansari Mohi, Xiao Ted, Kumar Vikash, Tompson Jonathan, Levine Sergey, Sermanet Pierre*. Learning latent plans from play // Conference on robot learning. 2020. 1113–1132. 6, 35
- [45] *Ma Xiao, Patidar Sumit, Haughton Iain, James Stephen*. Hierarchical Diffusion Policy for Kinematics-Aware Multi-Task Robotic Manipulation // Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. 2024. 18081–18090. 2
- [46] *Mahajan Divyat, Pezeshki Mohammad, Mitliagkas Ioannis, Ahuja Kartik, Vincent Pascal*. Compositional Risk Minimization // arXiv preprint arXiv:2410.06303. 2024. 3

- [47] *Mishra Utkarsh Aashu, Chen Yongxin, Xu Danfei*. Generative Factor Chaining: Coordinated Manipulation with Diffusion-based Factor Graph // 8th Annual Conference on Robot Learning. 2024. [3](#)
- [48] *Mishra Utkarsh Aashu, Xue Shangjie, Chen Yongxin, Xu Danfei*. Generative Skill Chaining: Long-Horizon Skill Planning with Diffusion Models // 7th Annual Conference on Robot Learning. 2023. [3](#), [6](#), [35](#)
- [49] *Nuti Felipe, Franzmeyer Tim, Henriques João F*. Extracting reward functions from diffusion models // Advances in Neural Information Processing Systems. 2024. [36](#). [2](#)
- [50] *Okawa Maya, Lubana Ekdeep S, Dick Robert, Tanaka Hidenori*. Compositional abilities emerge multiplicatively: Exploring diffusion models on a synthetic task // Advances in Neural Information Processing Systems. 2024. [36](#). [3](#)
- [51] *Park Seohong, Frans Kevin, Eysenbach Benjamin, Levine Sergey*. OGBench: Benchmarking Offline Goal-Conditioned RL // International Conference on Learning Representations (ICLR). 2025. [6](#), [7](#), [24](#), [35](#)
- [52] *Park Seohong, Ghosh Dibya, Eysenbach Benjamin, Levine Sergey*. Hiql: Offline goal-conditioned rl with latent states as actions // Advances in Neural Information Processing Systems. 2024. [36](#). [6](#), [35](#)
- [53] *Patil Omkar, Sah Anant, Gopalan Nakul*. Composing Diffusion Policies for Few-shot Learning of Movement Trajectories // arXiv preprint arXiv:2410.17479. 2024. [3](#)
- [54] *Pearce Tim, Rashid Tabish, Kanervisto Anssi, Bignell Dave, Sun Mingfei, Georgescu Raluca, Macua Sergio Valcarcel, Tan Shan Zheng, Momennejad Ida, Hofmann Katja, others*. Imitating human behaviour with diffusion models // The Eleventh International Conference on Learning Representations. 2023. [2](#)
- [55] *Peebles William, Xie Saining*. Scalable diffusion models with transformers // Proceedings of the IEEE/CVF International Conference on Computer Vision. 2023. 4195–4205. [33](#)
- [56] *Shafir Yonatan, Tevet Guy, Kapon Roy, Bermano Amit H*. Human motion diffusion as a generative prior // arXiv preprint arXiv:2303.01418. 2023. [3](#)
- [57] *Skreta Marta, Atanackovic Lazar, Bose Avishek Joey, Tong Alexander, Neklyudov Kirill*. The Superposition of Diffusion Models Using the Itô Density Estimator // arXiv preprint arXiv:2412.17762. 2024. [3](#)
- [58] *Sohl-Dickstein Jascha, Weiss Eric, Maheswaranathan Niru, Ganguli Surya*. Deep unsupervised learning using nonequilibrium thermodynamics // International conference on machine learning. 2015. 2256–2265. [2](#)
- [59] *Su Jocelin, Liu Nan, Wang Yanbo, Tenenbaum Joshua B, Du Yilun*. Compositional image decomposition with diffusion models // arXiv preprint arXiv:2406.19298. 2024. [3](#)
- [60] *Sun Jiankai, Jiang Yiqi, Qiu Jianing, Nobel Parth, Kochenderfer Mykel J, Schwager Mac*. Conformal prediction for uncertainty-aware planning with diffusion dynamics model // Advances in Neural Information Processing Systems. 2024. [36](#). [2](#)
- [61] *Sun Shanlin, De Araujo Gabriel, Xu Jiaqi, Zhou Shenghan, Zhang Hanwen, Huang Ziheng, You Chenyu, Xie Xiaohui*. CoMA: Compositional Human Motion Generation with Multi-modal Agents // arXiv preprint arXiv:2412.07320. 2024. [3](#)
- [62] *Sutton Richard S, Barto Andrew G*. Reinforcement learning: An introduction. 2018. [1](#)
- [63] *Thornton James, Bethune Louis, Zhang Ruixiang, Bradley Arwen, Nakkiran Preetum, Zhai Shuangfei*. Composition and Control with Distilled Energy Diffusion Models and Sequential Monte Carlo // arXiv preprint arXiv:2502.12786. 2025. [3](#)
- [64] *Ubukata Toshihide, Li Jialong, Tei Kenji*. Diffusion model for planning: A systematic literature review // arXiv preprint arXiv:2408.10266. 2024. [2](#)

- [65] Vaswani Ashish, Shazeer Noam, Parmar Niki, Uszkoreit Jakob, Jones Llion, Gomez Aidan N, Kaiser Łukasz, Polosukhin Illia. Attention is all you need // Advances in neural information processing systems. 2017. 30. [33](#)
- [66] Venkatraman Siddarth, Khaitan Shivesh, Akella Ravi Tej, Dolan John, Schneider Jeff, Berseth Glen. Reasoning with latent diffusion in offline reinforcement learning // arXiv preprint arXiv:2309.06599. 2023. [2](#)
- [67] Wang Hanlin, Wu Yilu, Guo Sheng, Wang Limin. Pdpp: Projected diffusion for procedure planning in instructional videos // Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. 2023. 14836–14845. [2](#)
- [68] Wang Junming, Zhang Xingyu, Xing Zebin, Gu Songen, Guo Xiaoyang, Hu Yang, Song Ziyang, Zhang Qian, Long Xiaoxiao, Yin Wei. HE-Drive: Human-Like End-to-End Driving with Vision Language Models // arXiv preprint arXiv:2410.05051. 2024. [2](#)
- [69] Wang Lirui, Zhao Jialiang, Du Yilun, Adelson Edward H, Tedrake Russ. Poco: Policy composition from and for heterogeneous robot learning // arXiv preprint arXiv:2402.02511. 2024. [3](#)
- [70] Wang Tongzhou, Torralba Antonio, Isola Phillip, Zhang Amy. Optimal goal-reaching reinforcement learning via quasimetric learning // International Conference on Machine Learning. 2023. 36411–36430. [6](#), [35](#)
- [71] Wang Yuanfu, Yang Chao, Wen Ying, Liu Yu, Qiao Yu. Critic-guided decision transformer for offline reinforcement learning // Proceedings of the AAAI Conference on Artificial Intelligence. 38, 14. 2024. 15706–15714. [2](#)
- [72] Wang Zhendong, Hunt Jonathan J, Zhou Mingyuan. Diffusion policies as an expressive policy class for offline reinforcement learning // arXiv preprint arXiv:2208.06193. 2022. [2](#)
- [73] Wu Yueh-Hua, Wang Xiaolong, Hamaya Masashi. Elastic decision transformer // Advances in neural information processing systems. 2023. 36. 18532–18550. [2](#)
- [74] Yamagata Taku, Khalil Ahmed, Santos-Rodriguez Raul. Q-learning decision transformer: Leveraging dynamic programming for conditional sequence modelling in offline rl // International Conference on Machine Learning. 2023. 38989–39007. [2](#)
- [75] Yang Brian, Su Huangyuan, Gkanatsios Nikolaos, Ke Tsung-Wei, Jain Ayush, Schneider Jeff, Fragkiadaki Katerina. Diffusion-es: Gradient-free planning with diffusion for autonomous driving and zero-shot instruction following // arXiv preprint arXiv:2402.06559. 2024. [2](#)
- [76] Yang Cheng-Fu, Xu Haoyang, Wu Te-Lin, Gao Xiaofeng, Chang Kai-Wei, Gao Feng. Planning as In-Painting: A Diffusion-Based Embodied Task Planning Framework for Environments under Uncertainty // arXiv preprint arXiv:2312.01097. 2023. [2](#)
- [77] Yang Mengjiao, Du Yilun, Dai Bo, Schuurmans Dale, Tenenbaum Joshua B, Abbeel Pieter. Probabilistic adaptation of text-to-video models // arXiv preprint arXiv:2306.01872. 2023. [3](#)
- [78] Yang Zhutian, Mao Jiayuan, Du Yilun, Wu Jiajun, Tenenbaum Joshua B, Lozano-Pérez Tomás, Kaelbling Leslie Pack. Compositional diffusion-based continuous constraint solvers // arXiv preprint arXiv:2309.00966. 2023. [3](#)
- [79] Ye Jiacheng, Gao Jiahui, Gong Shansan, Zheng Lin, Jiang Xin, Li Zhenguo, Kong Lingpeng. Beyond autoregression: Discrete diffusion for complex reasoning and planning // arXiv preprint arXiv:2410.14157. 2024. [2](#)
- [80] Zeng Zilai, Zhang Ce, Wang Shijie, Sun Chen. Goal-conditioned predictive coding for offline reinforcement learning // Advances in Neural Information Processing Systems. 2023. 36. 25528–25548. [2](#)
- [81] Zhang Jianrong, Fan Hehe, Yang Yi. EnergyMoGen: Compositional Human Motion Generation with Energy-Based Diffusion Model in Latent Space // arXiv preprint arXiv:2412.14706. 2024. [3](#)

- [82] *Zhang Qinsheng, Song Jiaming, Huang Xun, Chen Yongxin, Liu Ming-Yu*. Diffcollage: Parallel generation of large content with diffusion models // 2023 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). 2023. 10188–10198. [3](#)
- [83] *Zhang Ziqi, Xu Jingzehua, Liu Jinxin, Zhuang Zifeng, Wang Donglin, Liu Miao, Zhang Shuai*. Context-former: Stitching via latent conditioned sequence modeling // arXiv preprint arXiv:2401.16452. 2024. [2](#)
- [84] *Zhou Siyuan, Du Yilun, Zhang Shun, Xu Mengdi, Shen Yikang, Xiao Wei, Yeung Dit-Yan, Gan Chuang*. Adaptive online replanning with diffusion models // Advances in Neural Information Processing Systems. 2024. 36. [2](#)
- [85] *Zhou Zhaoyi, Zhu Chuning, Zhou Runlong, Cui Qiwen, Gupta Abhishek, Du Simon Shaolei*. Free from Bellman Completeness: Trajectory Stitching via Model-based Return-conditioned Supervised Learning // The Twelfth International Conference on Learning Representations. 2024. [2](#)
- [86] *Zhu Zhengbang, Liu Minghuan, Mao Liyuan, Kang Bingyi, Xu Minkai, Yu Yong, Ermon Stefano, Zhang Weinan*. Madiff: Offline multi-agent learning with diffusion models // Advances in Neural Information Processing Systems. 2024. 37. 4177–4206. [2](#)
- [87] *Zhuang Zifeng, Peng Dengyun, Liu Jinxin, Zhang Ziqi, Wang Donglin*. Reinformer: Max-Return Sequence Modeling for Offline RL // International Conference on Machine Learning. 2024. 62707–62722. [2](#)
- [88] *Ziebart Brian D, Maas Andrew L, Bagnell J Andrew, Dey Anind K, others*. Maximum entropy inverse reinforcement learning. // Aaai. 8. 2008. 1433–1438. [1](#)

NeurIPS Paper Checklist

1. Claims

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

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

Justification: We clearly and accurately outlined our scope and contribution in the Abstract and the Introduction Section.

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: We provided description of limitations of the work in Section 5. We also provide additional analysis to the proposed method, such as the factors that might influence the performance of our method (Appendix D) and computation efficiency (Appendix B.3).

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: This paper does not include theoretical results.

Guidelines:

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

4. Experimental result reproducibility

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

Answer: [Yes]

Justification: We have described our experiment settings in Section 4 and Appendix E. In addition, we will release our code upon acceptance.

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: We will release our code implementation and corresponding instructions upon acceptance.

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: The training and evaluation details are included in Section 4 and Appendix E.

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: As shown in our quantitative results, we report the mean and standard deviation based on 5 random seeds for each of our 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: We provide description of the computation platform of our experiments along with resource requirements in 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: This paper conforms with the NeurIPS Code of Ethics in every respect.

Guidelines:

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

10. Broader impacts

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

Answer: [NA]

Justification: We believe the societal impact of this work should be similar to a generic machine learning paper. No other issues we feel must be specifically highlighted.

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: This paper poses no such risks.

Guidelines:

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

12. Licenses for existing assets

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

Answer: [Yes]

Justification: The assets used in our paper are properly cited and related URLs are also provided.

Guidelines:

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

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

13. **New assets**

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

Answer: [Yes]

Justification: The new assets in our paper will be the corresponding code implementation of the proposed method. We will release the code along with documentation upon acceptance.

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

Guidelines:

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

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

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

Answer: [NA]

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

Guidelines:

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

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

16. **Declaration of LLM usage**

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

Answer: [NA]

Justification: This research 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.

Contents

1	Introduction	1
2	Related Work	2
3	Planning through Compositional Trajectory Generation	3
3.1	Compositional Trajectory Modeling	3
3.2	Training Compositional Trajectory Models	3
3.3	Compositional Trajectory Planning	4
4	Experiments	5
4.1	PointMaze	6
4.2	High Dimension Tasks	7
4.3	Ablation Studies	8
5	Discussion and Conclusion	9
A	Environment and Stitching Datasets	24
A.1	Stitching Datasets in Ghugare et al. [21]	24
A.2	OGBench Datasets	24
B	Additional Quantitative Results	26
B.1	Number of Composed Trajectories	26
B.2	Inverse Dynamics Model	26
B.3	Sampling Time Comparison: Parallel vs. Autoregressive	26
B.4	Starting Direction of Autoregressive Compositional Sampling	27
C	Additional Qualitative Results	28
C.1	Composing Different Numbers of Trajectories	28
C.2	Diverse Trajectory Morphology	28
C.3	Planning Results on Different Tasks.	29
C.4	Compositional Planning in High Dimensional Space	29
C.5	Compositional Planning in AntSoccer Arena	30
D	Failure Mode Analysis	30
D.1	Infeasible Generated Plan	30
D.2	Suboptimal Inverse Dynamics Model	31
D.3	Suboptimal Number of Composed Trajectories K	31
E	Implementation Details	32
E.1	Our Conditional Diffusion Model	32
E.2	Trajectory Merging	33
E.3	Replanning	34
F	Baselines	35

Appendix

In this appendix, we first introduce the evaluation environments and stitching datasets in Section A. Next, we provide additional quantitative results in Appendix B and additional qualitative results in Appendix C. We provide failure mode analysis in Section D. Following this, we provide implementation details in Section E, including model architecture, training and evaluation setups, trajectory merging, and replanning. Lastly, we introduce notable baselines in Section F.

A Environment and Stitching Datasets

In this paper, we directly evaluate our method on public stitching datasets introduced in two recent papers Ghugare et al. [21] and OGBench [51]. In this section, we provide detailed descriptions of each dataset along with qualitative examples of trajectories in these datasets.

A.1 Stitching Datasets in Ghugare et al. [21]

This paper [21] divides each evaluation environment into several small regions and each demonstration trajectory in the training datasets can only navigate within a specific region. There is a small overlap (one block) between each region, which can be used to stitch trajectories across regions. Therefore, to complete test-time goals, the agent needs to conduct effective reasoning based on the given start state and goal state and identify the corresponding overlap joints. The division of regions is visualized in the original paper [21]. We use the environments and datasets from their official implementation release at <https://github.com/RajGhugare19/stitching-is-combinatorial-generalisation>.

A.2 OGBench Datasets

OGBench is a comprehensive benchmark designed for offline goal-conditioned RL. Since our focus is to evaluate the trajectory stitching ability of CompDiffuser, we use the *Stitch* and *Explore* dataset types in OGBench.

In *Stitch* datasets, trajectories are constrained to navigate no more than 4 blocks in the environment. The start and goal state of each trajectory can be sampled from the entire environment provided that the travel distance between the start and goal is within 4 blocks. Qualitative examples of the trajectories in the *Stitch* dataset are shown in Figure 8.

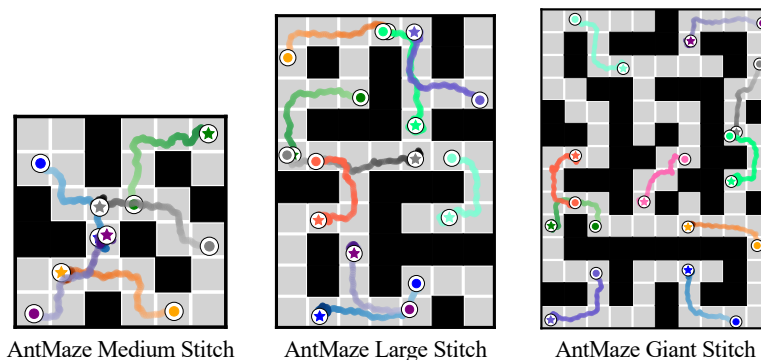


Figure 8: **Trajectory Examples in OGBench AntMaze Stitch Datasets.** Each Trajectory is limited to travel at most 4 blocks for dataset type *Stitch*, while at inference, the distance between the start and goal can be up to 30 in the *Giant* Maze.

In *Explore* datasets, trajectories are of extremely low-quality though high-coverage. The data collection policy contains a large amount of action noise and will randomly re-sample a new moving direction after every 10 steps. Hence, each demonstration trajectory in the training dataset typically moves within only 2-3 blocks due to the random moving direction. These datasets might be even more challenging due to the large noisy and cluster-like trajectory pattern. Qualitative examples of the trajectories in the *Explore* dataset are shown in Figure 9.

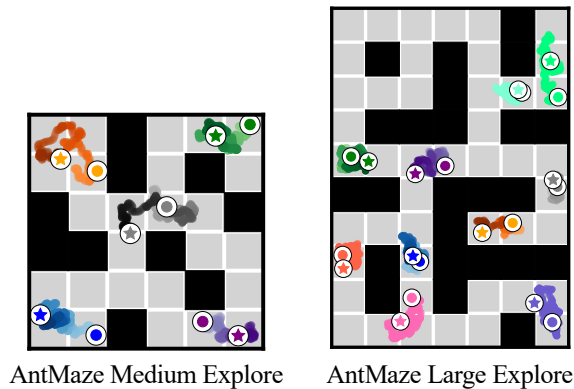


Figure 9: **Trajectory Examples in OGBench AntMaze Explore Datasets.** Trajectories in Explore datasets are of extremely low-quality though high-coverage. The data collection policy contains a large amount of action noise and will randomly re-sample a new moving direction after every 10 steps.

We show the environment names, corresponding datasets, and the maximum environment steps for each evaluation episode in Table 7. All methods are trained on the OGBench public release datasets. We slightly increase the environment steps for some environments due to the task difficulty (e.g., Giant Maze). For these environments, we follow the implementation in <https://github.com/seohongpark/ogbench>. and rerun all baselines with the increased maximum environment steps; for other environments, we directly adopt the reported success rates in the original paper.

Environment	Type	Size	Dataset Name	Env Steps
pointmaze	stitch	Medium	pointmaze-medium-stitch-v0	1000
		Large	pointmaze-large-stitch-v0	1000
		Giant	pointmaze-giant-stitch-v0	1000
AntMaze	Stitch	Medium	antmaze-medium-stitch-v0	1000
		Large	antmaze-large-stitch-v0	2000
		Giant	antmaze-giant-stitch-v0	2000
AntSoccer	stitch	Arena	antsoccer-arena-stitch-v0	5000
		Medium	antsoccer-medium-stitch-v0	5000
HumanoidMaze	Stitch	Medium	humanoid-medium-stitch-v0	5000
		Large	humanoid-large-stitch-v0	5000
		Giant	humanoid-giant-stitch-v0	8000

Table 7: **Our Evaluation Environments and Datasets in OGBench.**

B Additional Quantitative Results

In this section, we provide additional quantitative experiment results. In Section B.1, we study how varying the number of composed trajectories K affects the performance of our method. In Section B.2, we investigate the effect of inverse dynamic models of different horizons. Next, we provide a sampling time comparison of Decision Diffuser and the proposed parallel and autoregressive sampling schemes in Section B.3. Following this, in Section B.4, we analyze how the proposed autoregressive sampling scheme performs when using different denoising starting directions.

B.1 Number of Composed Trajectories

In Table 8, we compare CompDiffuser with GSC over composing different numbers of trajectories (results are also shown in Figure 7). Our method performs steadily when composing different numbers of trajectories while GSC collapses.

# Comp	PointMaze Giant Stitch					
	7	8	9	10	11	12
GSC	21 $\pm$ 1	21 $\pm$ 3	15 $\pm$ 2	2 $\pm$ 1	0 $\pm$ 0	0 $\pm$ 0
CompDiffuser	51 $\pm$ 7	53 $\pm$ 6	55 $\pm$ 4	56 $\pm$ 2	50 $\pm$ 6	50 $\pm$ 3

Table 8: **Quantitative Results over Different Numbers of Composed Trajectories.** We report success rates (w/o replanning) of composing 7 to 12 trajectories in the OGBench PointMaze Giant Stitch dataset. CompDiffuser can consistently construct feasible trajectories over various numbers of composed trajectories while GSC gradually collapses.

B.2 Inverse Dynamics Model

In this section, we evaluate our planner with inverse dynamics models of different horizons (results are also shown in Figure 7). Specifically, we use an MLP to implement the inverse dynamics model, which takes as input the start and goal state and outputs an action. We use the same training dataset (as to train the planner) to train the corresponding inverse dynamics model. As shown in Table 9, our method performs steadily across different inverse dynamics model horizons, showing that the subgoals generated by our planners are of high feasibility and are robust to various inverse dynamics models' configurations.

Env	Type	Size	8	12	16	20	24	28
AntMaze	Stitch	Large	77 $\pm$ 4	86 $\pm$ 2	80 $\pm$ 2	85 $\pm$ 3	76 $\pm$ 2	77 $\pm$ 3
		Giant	61 $\pm$ 5	65 $\pm$ 3	68 $\pm$ 4	63 $\pm$ 3	60 $\pm$ 2	63 $\pm$ 3

Table 9: **Quantitative Results of CompDiffuser with Inverse Dynamics Models of Different Horizons.** We present the success rates of CompDiffuser with 6 different inverse dynamics model horizons. In both environments, Large and Giant, our method performs consistently across all configurations, showing that the synthesized plans adhere to the transition dynamics and are easy to follow.

B.3 Sampling Time Comparison: Parallel vs. Autoregressive

In Table 10, we compare the diffusion denoising sampling time of three methods: (1) monolithic model method, Decision Diffuser (DD), where we directly sample a long trajectory with the same horizon as the proposed compositional sampling method; (2) parallel compositional sampling, as shown in the left of Figure 3, where we denoise all trajectories $\tau_{1:K}$ in one batch; (3) autoregressive compositional sampling, as shown in the right of Figure 3, where we sequentially denoise each τ_k .

We observe that both parallel and AR sampling methods require more sampling time than DD probably due to (1) the simpler and smaller denoiser network of DD; (2) time for condition

encoding (our method will first encode the noisy adjacent trajectories τ_{k-1} and τ_{k+1} and feed the resulting latents to the denoiser ϵ_θ) and (3) the overhead of trajectory merging.

In addition, in our parallel sampling scheme, we stack all trajectories $\tau_{1:K}$ to one batch and feed it to the denoiser network ϵ_θ . While it indeed requires only one model forward, the batch size increases implicitly, which is probably the major reason that the sampling time of Ours (Parallel) does not proportionally decrease as the number of composed trajectories K , in comparison to Ours (AR).

Env	Type	Size	DD (Monolithic)	Ours (Parallel)	Ours (AR)
PointMaze	Stitch	Medium	0.23	1.02	1.54
		Large	0.23	1.67	3.39
		Giant	0.23	2.73	5.00

Table 10: **Quantitative Comparison of Sampling Time** We report the time for sampling one (compositional) trajectory in three different PointMaze environments using one Nvidia L40S GPU (unit: second). The reported results are averaged over 20 sampling. In Parallel and AR (Autoregressive) mode, we use the proposed compositional sampling scheme as shown in Figure 3. Specifically, we compose 3, 6, and 9 trajectories in Maze Medium, Large, and Giant, respectively. In Decision Diffuser (DD), we directly sample one trajectory with identical length as the compositional counterparts.

B.4 Starting Direction of Autoregressive Compositional Sampling

In the proposed autoregressive sampling scheme described in Figure 3, inside each diffusion denoising timestep, the denoising starts from τ_1 and sequentially proceeds to τ_K (from left to right), which we denote as forward passing. Another implementation variant is to denoise in the reverse order, that is, first denoise τ_K and sequentially proceed to τ_1 (from right to left), which we denote as backward passing.

We provide a quantitative comparison of these two starting directions in Table 11. We observe that these two methods yield similar performance, demonstrating that either autoregressive sampling direction can enable effective information propagation and exchange.

Env	Type	Size	Ours (Forward)	Ours (Backward)
PointMaze	Stitch	Medium	100 $\pm$ 0	100 $\pm$ 0
		Large	100 $\pm$ 0	100 $\pm$ 0
		Giant	55 $\pm$ 4	56 $\pm$ 2

Table 11: **Quantitative Comparison of Forward and Backward Information Propagation.** We study the effect of the starting direction of the autoregressive sampling, from τ_1 vs. from τ_K . To directly study the trajectory quality, we report the results w/o replanning for both methods. Either Forward or Backward achieves similar performance, suggesting that our sampling method is robust to different sampling configurations.

C Additional Qualitative Results

In this section, we present additional qualitative results of CompDiffuser. Videos and interactive demos are provided at our project website (see *Abstract* section for the link).

C.1 Composing Different Numbers of Trajectories

We present qualitative results of composing 8 to 11 trajectories in OGBench PointMaze Giant Stitch environments in Figure 10. We compositionally sample multiple trajectories to construct a long-horizon plan where the given start is at the bottom-left corner and the given goal is at the top-right corner. For clearer view, we present the results before applying trajectory merging (i.e., we show each individual trajectory of $\tau_{1:K}$) and use different colors to highlight different trajectories.

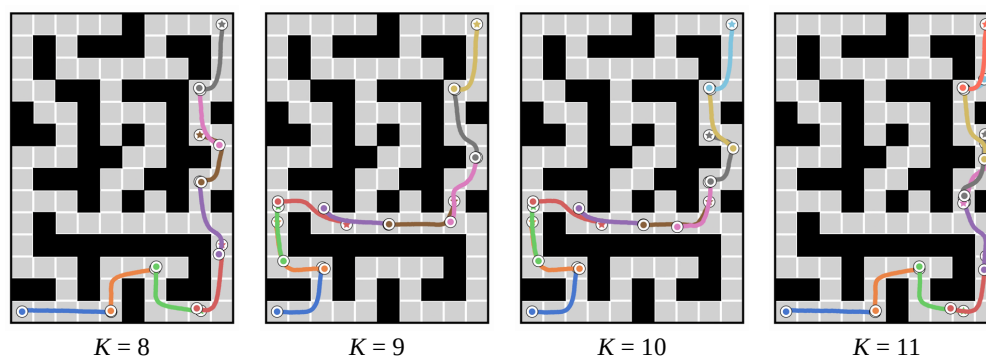


Figure 10: **Composing Different Numbers of Trajectories at Inference.** Our method is only trained on short trajectory segments that travel at most 4 blocks, while is able to compositionally generate coherent long trajectories given the start (circle) and goal (star). When K is smaller and barely sufficient to reach the goal (e.g., 8), the length of the overlapping part between segments decreases so as to extend the travel distance of the overall compositional plan. In contrast, if given a larger K (e.g., 11), some parts of the compositional plan might travel back and forth to consume the extra length.

C.2 Diverse Trajectory Morphology

The proposed compositional sampling method allows direct generalization to long-horizon planning tasks at test time through its noisy-sample conditioning and bidirectional information propagation design. Meanwhile, this sampling approach also preserves the multi-modal nature of the diffusion model, enabling a diverse range of trajectory morphology. As shown in Figure 11, given a similar start and goal pair, our method can construct trajectories that reach the goal via various possible paths. With such multi-modal flexibility, the proposed sampling process can be further customized with specific preferences by integrating additional test-time steering techniques.

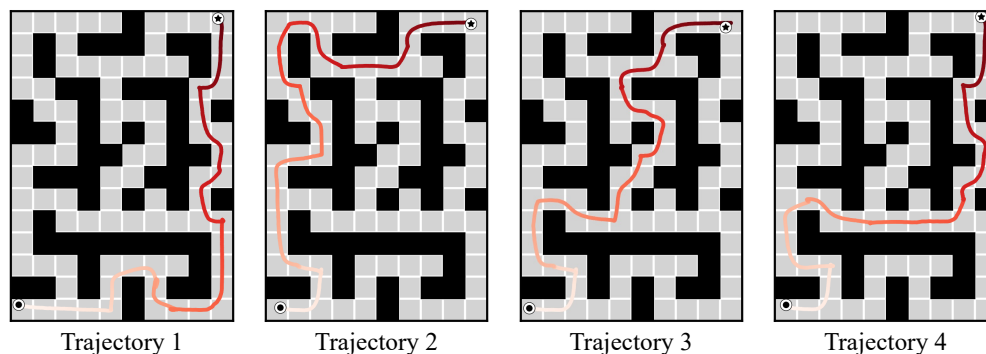


Figure 11: **Diverse Trajectory Morphology.** We present four trajectories with similar start and goal in OGBench PointMaze Giant Stitch. All trajectories are generated by CompDiffuser, composing 9 trajectories. CompDiffuser preserves the multi-modal nature of diffusion models and is able to flexibly sample trajectories of diverse morphology.

C.3 Planning Results on Different Tasks.

In Section C.2 above, we present multiple trajectories of Task 1 in OGBench PointMaze Giant. In this section, we present qualitative results of the following Task 2 to Task 5, as in Figure 12. We set the number of composed trajectories K to 9 in all tasks. The state state is shown by the black circle and the goal state is shown by the black star. Our method successfully constructs feasible plans for various start-goal configurations across different spatial distances.

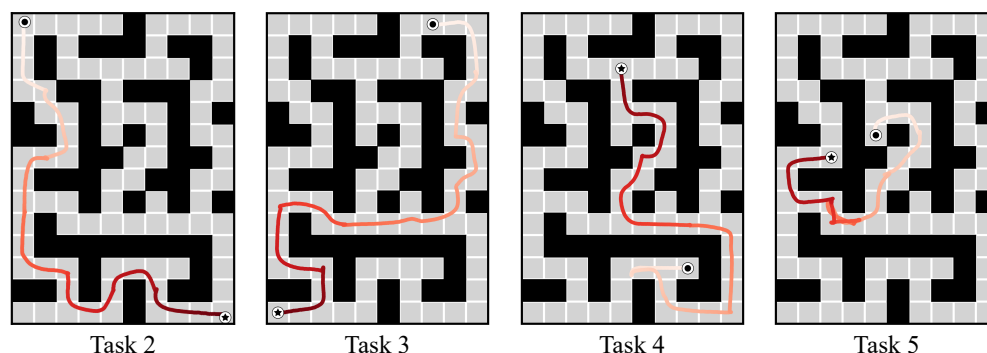


Figure 12: **Different Tasks in OGBench PointMaze Giant Stitch.** We present qualitative plans by CompDiffuser on OGBench Task 2-5 (See Figure 11 for results of Task 1). We set the number of composed trajectories to 9 for all the tasks above. The black circle denotes the start state and the black star denotes the goal state. In task 2 and 3, our method effectively constructs long horizon plans that reach the goals on the opposite side of the environment (despite being trained only on trajectories that travel at most 4 blocks). In comparison, Task 4 and 5 feature a relatively smaller spatial gap between the start and goal, thus requiring a shorter planning horizon. Our method still generalizes to these tasks by generating plans that traverse additional distance or leveraging back-and-forth movements to consume the extra plan length.

C.4 Compositional Planning in High Dimensional Space

In this section, we present additional high dimensional trajectories generated by CompDiffuser in OGBench AntMaze Large Stitch environment. Similar to other experiments in the paper, the models are trained on the corresponding OGBench public release datasets.

AntMaze Large Stitch 29D. We train CompDiffuser on the state space of x - y position along with the ant’s joint positions and velocities, resulting in a 29D planning task. Note that CompDiffuser is only trained on short trajectory segments (we set its horizon to 160), while at test time, we compose 5 trajectories to directly construct trajectory plans of horizon 584. We present additional qualitative results in Figure 13. Note that we uniformly sub-sample the length of the trajectory to 50 for clearer view. Corresponding quantitative results are reported in Table 4.

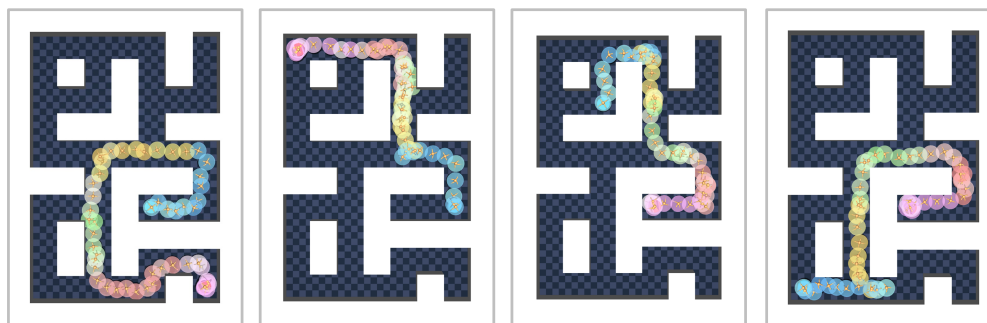


Figure 13: **Compositional Planning in 29D on OGBench AntMaze Large Stitch Tasks.** Original trajectory plans are much denser and we uniformly sub-sample 50 states from the original generated trajectories for better view. Five trajectories are composed as shown by different colors: the blue one indicates the first trajectory τ_1 and the purple one indicates the fifth trajectory τ_5 .

C.5 Compositional Planning in AntSoccer Arena

We provide additional qualitative results in OGBench AntSoccer Arena Stitch environment in Figure 14. In this task, the ant is initialized to the location of the blue circle and is tasked to move the ball to the goal location indicated by the pink circle. Hence, the ant needs to first reach the ball from the far side of the environment and dribble the ball to the goal.

However, such long-horizon trajectories (the ant reaches the ball and then dribbles the ball to the goal) do not exist in the training dataset. The training dataset only contains two distinct types of trajectories: (1) the ant moves in the environment without the ball, (2) the ant moves while dribbling the ball. Therefore, the planner needs to generalize and stitch in a zero-shot manner – constructing an end-to-end trajectory that first approaches the ball and then dribbles the ball to the goal.

We train CompDiffuser on the state space of the ant's x - y position and joint positions along with the x - y position of the ball, resulting in a 17D planning task. Similar to Figure 6, we present each individual trajectory $\tau_{1:K}$ in Figure 14 for clearer view. Corresponding quantitative results are provided in Table 3.

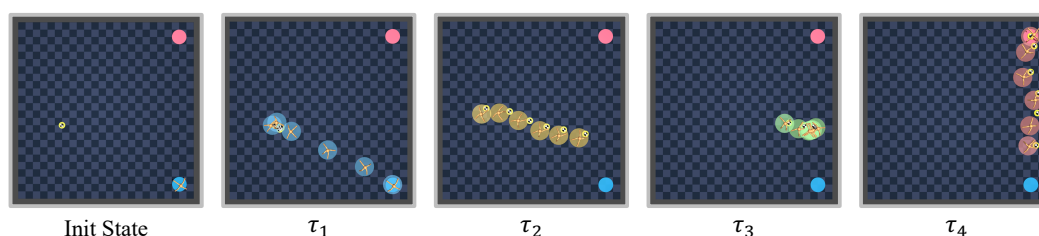


Figure 14: **Compositional Planning on OGBench AntSoccer Arena Stitch.** We present the initial state for planning and each individual trajectories $\tau_{1:4}$ above. The start position of the ant is shown by the blue circle (bottom right) and the goal is to move the ball to the pink circle (upper right). We compositionally sample 4 trajectories (as shown from left to right), which will then be merged to form a long-horizon plan. The ball is highlighted with a small yellow circle. Our compositional sampling method effectively stitches two different types of trajectories and generalizes to more difficult tasks unseen in the training data.

D Failure Mode Analysis

In this section, we present several failure cases of our method along with analyses and qualitative examples of these failure modes. Our proposed framework consists of two components: a generative planner and an inverse dynamics model. We outline and discuss several failure modes below.

D.1 Infeasible Generated Plan

As the number of composed trajectories K increases, our method may show inconsistencies in the trajectory plan, especially when planning in high dimensional state space. We observe that although the start and goal conditioning can consistently ensure that the composed plan begins at the initial state and terminates at the provided goal, the intermediate trajectories may include implausible transitions, such as trajectories that jump over walls, making the overall plan infeasible. However, empirically these failures usually occur at considerably higher K as compared to baseline methods (e.g., GSC, see Table 1 and Table 2).

In Figure 15, we present a qualitative example of infeasible plan in OGBench AntMaze Giant Stitch when planning in 29D space (ant's x - y position, joint positions, and joint velocity) and composing 9 trajectories. The original white walls are rendered transparent for better view of infeasible states. Though the generated compositional plan is mostly valid, the second trajectory τ_2 , as shown in yellow, is infeasible as it passes through walls. This is probably due to the suboptimal coordination among trajectories in the compositional sampling process.

As illustrated in the figure, certain trajectories (e.g., the neighboring blue, green, and red ones) barely proceeds, moving only 2-3 blocks. To bridge the resulting spatial gap, the intermediate yellow trajectory must span a much longer distance. However, the training data does not contain such long-horizon trajectories, making it difficult for a single trajectory to extend to such length, which

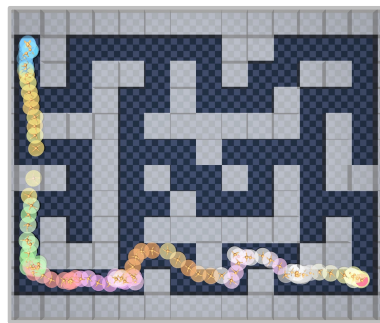
finally results in an o.o.d sample that goes through walls. We could potentially mitigate this issue with rejection or guided sampling techniques to select feasible candidate plans.

D.2 Suboptimal Inverse Dynamics Model

In our experiments, we observe that even if CompDiffuser synthesizes a feasible trajectory for the agent to follow, the agent might not successfully reach the goal due to the error by the inverse dynamics model. For example, the agent might bump into walls due to unstable locomotion or get stuck in some local region, yet the synthesized plan is collision-free and coherent.

In implementation, we use a simple MLP to parameterize the inverse dynamics model and train it with a regression MSE loss. We believe that more optimized model architecture or specific finetuning might further boost the performance of the inverse dynamics model, hence boosting the overall performance of our method. We deem that out of the scope of this work.

In Figure 16, we present a qualitative example of failure due to the inverse dynamics model in OGBench AntMaze Giant Stitch. The planning is in 15D space (ant's x - y position and joint positions) and composes 3 trajectories. While we observe that the synthesized plan is feasible and successfully reaches the goal, the environment execution rollout fails during the yellow trajectory. In this instance, the ant agent fails to execute the right turn, loses track of subsequent subgoals, and becomes trapped in a local region. To address this issue, incorporating effective replanning strategies could help the agent recover (since the new plan will start from the current trapped state). In addition, we deem that employing more robust or specialized inverse dynamics models may further mitigate such failure scenarios.



Infeasible Plan: Passing through Walls

Figure 15: Failure Mode: Infeasible Plan. We increase the transparency of the original white walls for better view of infeasible states. The start state is at the top left corner and the goal state is at the bottom right corner. Nine trajectories are composed, highlighted by different colors. The second trajectory τ_2 (shown in yellow) is infeasible and passes through walls. This is probably due to the suboptimal coordination between trajectories in the compositional sampling process: the neighboring blue and green trajectories barely progress, leaving a long o.o.d gap for the yellow trajectory to fill.



Feasible Plan



Rollout: Suboptimal Actions

Figure 16: Failure Mode: Suboptimal Inverse Dynamics Actions. We present a failure case of planning in 15D space (ant's x - y position and joint positions) in AntMaze Medium Stitch. Left: the compositional plan where three trajectories (represented in blue, yellow, and green) are composed. We sub-sample the plan to 36 states for visualization (the original plan is dense with over 300 states). Right: the corresponding environment execution rollout of the compositional plan. Though the synthesized plan is valid and successfully reaches the goal, the inverse dynamics model may fail, as illustrated on the right which gets stuck in a right turn and is unable to proceed. Further incorporating some effective replanning strategies or employing more robust inverse dynamics models could potentially mitigate these failure scenarios.

D.3 Suboptimal Number of Composed Trajectories K

In our current implementation, the number of composed trajectories K needs to be manually specified at test time. As shown by the quantitative results Table 8 and qualitative results Figure 10, a relatively larger K does not significantly affect the model performance as the extra plan length will be consumed by staying or circulating within certain valid regions in the environment. However, an aggressively smaller K might lead to failure plans – since the overall planning horizon becomes insufficient to cover the substantial spatial gap between the start and the goal.

In this section, we provide qualitative examples of infeasible plans due to impractical K . In Figure 17, we show each individual trajectory τ_k generated by our compositional sampling method when K ranges from 3 to 6. Note that a feasible horizon to reach the goal from the start (bottom left) to the goal (upper right) requires composing at least 8 trajectories, i.e., $K = 8$.

Therefore, while the generated trajectories can begin at the start state and terminate at the goal state, the intermediate segments become disconnected since there are too few trajectories to bridge the significant spatial gap (given that the training data contains only short trajectories). Nonetheless, we observe that the overall flow of the trajectories is directed toward the goal, and as K increases, the plan's structure gradually becomes more feasible.

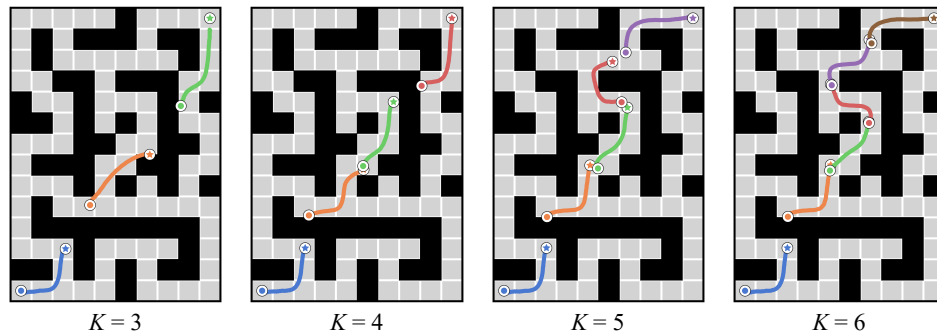


Figure 17: **Failure Model: Suboptimal Number of Composed Trajectories K .** Our method may generate infeasible plans if K , the number of composed trajectories, is set significantly below the required minimum. For example, in the planning task illustrated above, the start state is located in the bottom left corner while the goal state is at the upper right corner. A feasible plan for this task typically requires composing at least 8 trajectories (i.e., $K = 8$). We present qualitative examples of plans when K is smaller than the minimum threshold. Though the first and last trajectory segments correctly attach to the start and goal states, the intermediate trajectories are disconnected due to the insufficient plan length. However, the overall plan still progresses towards the goal, which may provide useful guidance signals to the agent.

E Implementation Details

Software: The computation platform is installed with Ubuntu 20.04.6, Python 3.9.20, PyTorch 2.5.0.

Hardware: We use 1 NVIDIA GPU for each experiment. A GPU with 24GB memory is sufficient to train our models and it takes 1-2 days to train a model using a recent mid-level NVIDIA GPU.

E.1 Our Conditional Diffusion Model

In this section, we introduce detailed implementation of our conditional diffusion model $\epsilon_\theta(\tau^t, t \mid \text{st_cond}, \text{end_cond})$.

Model Inputs and Outputs. Our diffusion model takes as input the noisy trajectory τ^t , diffusion noise timestep t , and the start condition st_cond and end condition end_cond for τ^t . st_cond can be either noisy chunk or start state q_s and end_cond can be either noisy chunk or goal state q_g . We use the predicting τ^0 formulation to implement this network, i.e., the network directly predicts the clean sample τ^0 .

Implementation of Noisy Chunk Conditioning. As described in Section 3, we only train *one* diffusion model ϵ_θ that is able to condition on both the start state q_s , goal state q_g , and noisy chunk $\tau_{k-1}^t, \tau_{k+1}^t$, such that in test time we can use the proposed compositional sampling approach to construct long-horizon plans. This unified one model design eliminate the need for manual task subdivision (across multiple models) or reliance on predefined planning skeleton, thereby enabling holistic end-to-end planning.

In practice, we can implement the noisy neighbor conditioning $\mathcal{L}_{\text{nbr}}$ simply using the overlapping part between the two chunks instead of a full chunk τ_{k-1}/τ_{k+1} conditioning, because the overlapping part is sufficient to ensure the connectivity between two adjacent trajectories. Such design further simplifies the training procedure: instead of dividing a training trajectory τ to K chunks, we only need to sample a noisy sub-chunk from τ itself.

Specifically, assume τ^t is the noisy trajectory to be denoised and $\hat{\tau}^t$ is another independent noisy version of τ also at noisy timestep t . We denote the length of τ as h and the length of the overlapped part as h_o , then we can use $\hat{\tau}^t[0 : h_o]$ and $\hat{\tau}^t[h - h_o : h]$ as the noisy sub-chunks for `st_cond` and `end_cond` during training, respectively. Hence, when inference, the `end_cond` for τ_1 can be set to $\hat{\tau}_2^t[0 : h_o]$ (the front chunk of the second trajectory τ_2) and the `st_cond` for τ_2 can be set to $\hat{\tau}_1^t[h - h_o : h]$ (the tail chunk of the first trajectory τ_1).

Model Architecture. For planning in 2D x - y space, we follow Decision Diffuser [1] (<https://github.com/anuragajay/decision-diffuser/>) and use a conditional U-Net as the denoiser network. For planning in high dimension state space, we use a DiT [55] based transformer [65] as the denoiser network (<https://github.com/facebookresearch/DiT/>).

Training Pipeline. We provide detailed hyperparameters for training our model on PointMaze Giant Stitch environment in Table 12. We do not apply any hyperparameter search or learning rate scheduler. Please refer to our codebase for more implementation details.

Hyperparameters	Value
Horizon	160
Diffusion Time Step	512
Probability of Condition Dropout	0.2
Iterations	1.2M
Batch Size	128
Optimizer	Adam
Learning Rate	2e-4
U-Net Base Dim	128
U-Net Encoder Dims	(128, 256, 512, 1024)

Table 12: Hyperparameters for Training on PointMaze Giant Stitch environment.

Inference Pipeline. In Table 13, we present the single model horizon (length of individual τ_k) and the inference-time number of composed trajectories K corresponding to the reported results in Table 1, Table 2 and Table 3.

For each evaluation problem, we generate B samples in a batch and we use a simple heuristic to select one sample as the output plan. Specifically, we compute the L2-distance of each overlapping parts in the generated trajectory segments $\tau_{1:K}$. The one with the smallest average distance will be adopted as the output plan, in the sense that a small distance in the overlapping parts indicates better coherency between adjacent trajectories. we deem that developing some more advanced inference-time methods with CompDiffuser may be an interesting future research direction, such as probability or density based plan selection methods or compositional sampling with flexible preference steering.

E.2 Trajectory Merging

As introduced in Method Section 3 and Algorithm 2, the generated trajectories $\tau_{1:K}$ are mutually overlapped and we merge these K trajectories to form a long-horizon compositional trajectory τ_{comp} . In this section, we describe the implementation of the exponential trajectory blending technique which we use for merging.

We directly leverage the classic exponential trajectory blending formulation. For simplicity, let τ_1 and τ_2 denote the trajectories to blend, $\tau_1[t]$ denote the t -th state in τ_1 , t_{start} and t_{end} denote the start and end index of the region to apply blending. Note that, in practice, we *only* blend the overlapped part between two adjacent trajectories. We provide the equation for exponential blending below,

$$\tau_{\text{comp}}[t] = w(t) * \tau_1[t] + (1 - w(t)) * \tau_2[t], \quad \text{where } w(t) = \frac{e^{-\beta \left(\frac{t - t_{\text{start}}}{t_{\text{end}} - t_{\text{start}}} \right)} - e^{-\beta}}{1 - e^{-\beta}} \quad (7)$$

We set $\beta = 2$ across all our experiments. In practice, various other trajectory blending techniques can also be directly applied, such as cosine blending and linear blending.

Environment	Type	Size	Single Model Horizon	# of Composed Trajectories
PointMaze [21]	-	U-maze	40	5
		Medium	144	5
		Large	192	5
pointmaze	Stitch	Medium	160	3
		Large	160	5
		Giant	160	8
AntMaze	Stitch	Medium	160	3
		Large	160	6
		Giant	160	9
AntMaze	Explore	Medium	192	5
		Large	192	6
AntSoccer	Stitch	Arena	160	5
		Medium	160	6
HumanoidMaze	Stitch	Medium	336	4
		Large	336	6
		Giant	336	11

Table 13: **Number of Composed Trajectories for Each Evaluation Environment.** Our diffusion models are trained with a short horizon as listed in the *Single Model Horizon* column. In test time, we compositionally generate multiple such short trajectories to enable trajectory stitching and construct plans of much longer horizon.

E.3 Replanning

In this section, we describe the detailed implementation of replanning. While our method is designed to directly generate an end-to-end trajectory from the given start state to the goal state, replanning can be performed at any given timesteps during a rollout. Specifically, we initiate replanning if the agent loses track of the current subgoal, i.e., the L2 distance between the agent and the subgoal is larger than a threshold.

In a larger maze, the required number of composed trajectories, denoted as K , is usually large due to the distance between the start state and the goal state. However, if we keep replanning with a similar large K even when the agent is already close to the goal, the generated trajectory might travel back and forth to consume the unnecessary intermediate length (see (2) and (3) in Figure 12), thus delaying the agent’s progress toward the goal.

To address this, we use a receding scheme for K to encourage faster convergence to the goal. Let K denote the number of composed trajectories of the current plan (the plan that the agent is following), h_{comp} denote the length of the current plan, h_{exe} denote the length of the current plan that the agent already executes in the environment. The number of composed trajectories used for replanning K_{replan} is given by

$$K_{\text{replan}} = \text{ceil} \left(\left(1 - \frac{\max(h_{\text{exe}} - \delta, 0)}{h_{\text{comp}}} \right) * K \right) \quad (8)$$

where δ is a hyper-parameter that controls the convergence speed to the goal, for example, K_{replan} will decrease faster if δ is set to a small (or negative) number while K_{replan} will decrease very slowly if δ is a large positive number.

F Baselines

We compare our approach with a wide variety of baselines, including diffusion-based trajectory planning algorithms, data augmentation based stitching algorithms, and goal-conditioned offline RL algorithms.

Particularly, we include the following methods:

- For generative planning methods, we include Decision Diffuser (DD) [1] for monolithic trajectory sampling and Generative Skill Chaining (GSC) [48] for trajectory stitching;
- For data augmentation based methods, we include stitching specific data augmentation [21] with RvS [14] and Decision Transformer (DT) [9]. Since the evaluation setting is identical, we directly adopt the reported numbers in the original paper [21].
- For offline reinforcement learning methods, we include goal-conditioned behavioral cloning (GCBC) [44, 20], goal-conditioned implicit V-learning (GCIVL) and Q-learning (GCIQL) [29], Quasimetric RL (QRL) [70], Contrastive RL (CRL) [15], and Hierarchical implicit Q-learning (HIQL) [52]. For these baselines, we follow the implementation setup established by OGBench [51] throughout our experiments.

We describe the implementation details of a few notable ones below.

Generative Skill Chaining (GSC) [48]. GSC is a recent diffusion model-based skill stitching method. We directly adopt its original score-averaging based stitching algorithm and apply it in our tasks. Specifically, when composing K trajectories $\tau_{1:K}$, GSC averages the scores of the overlapping segments between adjacent trajectories (in total $K - 1$ overlapping segments in this case) prior to each denoising timestep. For a fair comparison, we re-use the same diffusion denoiser networks in CompDiffuser for every GSC experiment.

Hierarchical Implicit Q-Learning (HIQL) [52]. HIQL is a recently proposed Q-learning based method that employs a hierarchical framework for training goal-conditioned RL agents. It learns a goal conditioned value function and uses it to learn feature representations, high-level policy, and low-level policy. We follow the original implementation of the method in OGBench [51].

Goal-Conditioned Behavioral Cloning (GCBC) [44]. GCBC is a classic imitation learning-based method. In our experiments, GCBC trains an MLP that takes as input the observation state and a future goal state from the same offline trajectory and outputs a corresponding action for the agent. We use the same implementation as in OGBench [51].