

Counterfactual Data Augmentation for Offline Reinforcement Learning via Conditional Diffusion Models

Rishi Ojha¹, Shivam Rajput², Shivam Dhakad³, Vikash Narveriya⁴, Dr. Shiv Kumar Sharma⁵

^{1, 2, 3, 4} CS-Data Science, ITM GOI, Gwalior, India ^{5} Associate Professor

HOD Department of Mechanical Engineering, ITM GOI, Gwalior, India

¹rjha7221@gmail.com, ²shivamrajput6643@gmail.com, ³Shivamdhakad0450@gmail.com

⁴vikashnarveriya444@gmail.com, ⁵sksme830@gmail.com

Abstract—

Offline reinforcement learning (RL) is limited by the coverage of the static training dataset: regions of the state-action space absent from the data lead to conservative, brittle policies. Data augmentation addresses this by adding synthetic transitions, but existing methods either apply random perturbations that ignore dynamics or generate model-based rollouts that invite model exploitation—neither accounts for the causal structure of the environment. We propose Counterfactual Diffusion Augmentation (CDA), which uses conditional diffusion models to generate counterfactual next states for actions not present in the offline dataset. A dual plausibility filter—combining forward dynamics uncertainty and inverse dynamics consistency—discards unrealistic synthetic transitions before they reach the training buffer, blocking model exploitation. The filtered counterfactuals augment any standard offline RL algorithm without modifying its training objective. On D4RL MuJoCo and Adroit tasks, CDA raises the returns of CQL, IQL, and TD3+BC by up to 25%, with the largest gains in low-data regimes. The same diffusion model also supports offline-to-online fine-tuning without catastrophic forgetting. CDA shows that counterfactual generation with diffusion models is a practical augmentation strategy for offline RL.

Keywords—Offline reinforcement learning, data augmentation, diffusion models, counterfactual reasoning, plausibility filtering.

I. INTRODUCTION

Offline reinforcement learning aims to learn effective policies from pre-collected static datasets without further environment interaction [1]. Real-world datasets typically cover only a fraction of the state-action space, so learned policies tend to be conservative and fail at deployment. Data augmentation widens training coverage by inserting additional synthetic transitions drawn from the true environment manifold.

Existing augmentation strategies have two weaknesses. Random perturbation methods add noise to states or actions [2], [3] but are dynamics-agnostic and can produce physically impossible transitions. Model-based approaches train a world model and roll out synthetic trajectories [4], [5], but the policy can learn to exploit model inaccuracies in underrepresented regions.

Neither family uses counterfactual reasoning—asking what would have happened under a different action. Counterfactual transitions respect the causal structure of the environment and carry precise information about the effect of unexplored actions. Conditional diffusion models can generate realistic next states for arbitrary state-action pairs, including those absent from the training data [6], [7]. Using such samples without verification, however, risks polluting the training buffer with out-of-distribution (OOD) artefacts.

We introduce Counterfactual Diffusion Augmentation (CDA), which generates counterfactual transitions with a diffusion model and filters them before use. Our contributions are:

1) A procedure to train a conditional diffusion model on the offline dataset and sample counterfactual next states for alternative actions.

2) A dual plausibility filter combining forward uncertainty and inverse dynamics consistency that removes unrealistic counterfactuals before they reach the buffer.

3) A drop-in augmentation protocol compatible with CQL, IQL, and TD3+BC, with gains of 8–27% across D4RL tasks and nearly doubled returns under 10% data.

4) An extension to offline-to-online fine-tuning that reuses the same diffusion model to generate counterfactuals for novel online states, avoiding catastrophic forgetting.

II. RELATED WORK

A. Data Augmentation in Offline RL

RAD [2] and DrQ [15] apply random image augmentations in online RL. S4RL [3] adapts state perturbation to offline RL but does not perform causal counterfactual generation. MOPO [4] and COMBO [16] generate model-based rollouts and use uncertainty penalties to limit exploitation, though this can be overly conservative. CDA generates targeted counterfactuals and removes implausible ones via a dedicated two-stage filter.

B. Diffusion Models in RL

Diffusion models have been used as world models (DIAMOND [8]) and for trajectory-level planning (Diffuser [9]). CDA uses diffusion as a data augmentation module, not for planning or model-based rollouts, preserving the stability guarantees of the underlying model-free offline algorithms.

C. Counterfactual Reasoning in RL

Counterfactual methods have appeared in model-based RL [10] and off-policy evaluation [11]. CDA is the first to combine counterfactual generation via diffusion models with offline RL data augmentation.

III. PRELIMINARIES

A. Offline RL

The environment is a Markov Decision Process (MDP) (S, A, P, r, γ) . The agent learns a policy $\pi(a|s)$ from a fixed dataset $D = \{(s_i, a_i, r_i, s'_i)\}$ collected by unknown behaviour policies, with no further environment interaction.

B. Conditional Diffusion Models

A conditional diffusion model learns $p(x|c)$ by reversing a gradual noising process. The forward process adds Gaussian noise: $q(x_t|x_{t-1}) = N(\sqrt{1-\beta_t}x_{t-1}, \beta_t I)$. A denoising network $\varepsilon\theta(x_t, \tau, c)$ reverses this process by minimising:

$$L = E\| \varepsilon - \varepsilon\theta(\sqrt{\alpha_\tau}x_0 + \sqrt{1-\alpha_\tau}\varepsilon, \tau, c) \|^2 \quad (1)$$

For CDA, $x_0 = s_{t+1}$ and $c = (s_t, a_t)$.

IV. COUNTERFACTUAL DIFFUSION AUGMENTATION (CDA)

CDA proceeds in three stages: (i) train a conditional diffusion model and an inverse dynamics model; (ii) generate and filter counterfactual transitions; (iii) train an offline RL algorithm on the augmented dataset.

A. Model Training

Forward diffusion model. We train $p\theta(s_{t+1}|s_t, a_t)$ on D using Eq. (1). Once trained, it produces next-state samples for any state-action pair, including pairs absent from the dataset.

Inverse dynamics model. We train a network $q\phi(a_t|s_t, s_{t+1})$ via maximum likelihood on D to predict the action that caused a given transition. It serves as the consistency check in the plausibility filter.

B. Counterfactual Generation with Plausibility Filtering

Given a real transition (s_t, a_t, s_{t+1}) from D , we draw a counterfactual action $\tilde{a}_t$ from the current policy $\pi\psi$ (with exploration noise) or a behaviour-cloning model, then sample $\hat{s}_{t+1} \sim p\theta(\cdot|s_t, \tilde{a}_t)$.

Generated samples are checked with a two-stage plausibility filter:

- 1) Forward confidence. We draw M samples per input and compute their variance [17]. Samples with variance above η_{fwd} are rejected.
- 2) Inverse dynamics consistency. We pass $(s_t, \hat{s}_{t+1})$ through $q\phi$ to get $\hat{a}_t$. If $\|\hat{a}_t - \tilde{a}_t\|^2 > \eta_{inv}$, the sample is rejected.

Transitions passing both checks enter the augmented buffer D_{au}^ϕ . Thresholds η_{fwd} and η_{inv} are the 95th percentiles of the corresponding statistics over D .

Reward relabelling. We use the environment reward function where available (as in D4RL), or a reward model trained on D , to assign a reward to each accepted counterfactual.

C. Augmented Offline RL Training

Any model-free offline RL algorithm (CQL [12], IQL [13], TD3+BC [14]) is then trained on $D \cup D_{au}^\phi$ without

modification. Broader action coverage reduces over-conservatism; the plausibility filter blocks OOD exploitation.

D. Offline-to-Online Extension

During online fine-tuning, the same diffusion model generates counterfactuals for novel states encountered online. The fraction of augmented samples is kept low early in fine-tuning and increased as the policy stabilises, preventing catastrophic forgetting.

Algorithm 1 summarises the offline training phase.

Algorithm 1 CDA: Counterfactual Diffusion Augmentation (Offline)

```

Input: Dataset  $D$ , diffusion model  $p\theta$ ,
inverse model  $q\phi$ ,
offline RL algorithm  $\text{Alg}$ ,
thresholds  $\eta_{fwd}$ ,  $\eta_{inv}$ 
Train  $p\theta$  and  $q\phi$  on  $D$ 
 $D_{au}^\phi \leftarrow \emptyset$ 
for each  $(s_t, a_t, s_{t+1}) \in D$  do
    Sample  $\tilde{a}_t$  from proposal distribution
    (policy + noise)
    Generate  $\hat{s}_{t+1} \sim p\theta(\cdot|s_t, \tilde{a}_t)$ 
    Estimate forward uncertainty  $\sigma_{fwd}$ 
    if  $\sigma_{fwd} > \eta_{fwd}$  then continue
    end if
    Compute  $\hat{a}_t = q\phi(s_t, \hat{s}_{t+1})$ 
    if  $\|\hat{a}_t - \tilde{a}_t\|^2 > \eta_{inv}$  then continue
    end if
    Predict reward  $\hat{r}$ ; add  $(s_t, \tilde{a}_t, \hat{r}, \hat{s}_{t+1})$ 
    to  $D_{au}^\phi$ 
end for
Run  $\text{Alg}$  on  $D \cup D_{au}^\phi$  to obtain  $\pi\psi$ 
return  $\pi\psi$ 

```

V. EXPERIMENTS

We evaluate CDA on D4RL tasks [18]: MuJoCo locomotion (HalfCheetah, Hopper, Walker2d) and Adroit manipulation (Pen, Door), with "medium" (m), "medium-expert" (m-e), and "10% medium" datasets.

A. Baselines and Setup

CDA is combined with CQL [12], IQL [13], and TD3+BC [14]. Baselines include each algorithm alone, each augmented with S4RL perturbations [3], and MOPO [4] without augmentation. Diffusion models use a U-Net with 3 training steps and DDIM sampling for inference speed. The inverse dynamics model is a 3-layer MLP. Thresholds are set at the 95th percentile of each metric on D . Each original transition generates exactly one counterfactual candidate, doubling the buffer if it passes the filter.

B. Main Results

Table I reports average normalized return over 5 seeds. CDA improves all three base algorithms on every task, with gains of 8–27%. The largest gains are on Adroit Door and the 10% medium Hopper dataset, where dataset coverage is most limited.

TABLE I.
AVERAGE NORMALIZED RETURN AFTER 100K GRADIENT STEPS
(MEAN $\pm$ STD, 5 SEEDS). CDA ADDED TO CQL/IQL/TD3+BC.

Task	CQL	CQL+CDA	IQL	IQL+CDA
HalfCheetah-m	47.0 $\pm$ 0.5	55.4 $\pm$ 0.7	47.4 $\pm$ 0.6	56.0 $\pm$ 0.8
Hopper-m	58.5 $\pm$ 3.2	68.9 $\pm$ 2.1	66.7 $\pm$ 2.5	73.2 $\pm$ 1.9
Walker2d-m	72.5 $\pm$ 2.0	79.3 $\pm$ 1.6	74.0 $\pm$ 1.8	81.1 $\pm$ 1.4
Pen-h	37.5 $\pm$ 5.1	50.2 $\pm$ 3.8	42.3 $\pm$ 4.6	55.7 $\pm$ 3.5
Door-h	9.9 $\pm$ 5.6	21.3 $\pm$ 4.9	15.2 $\pm$ 5.8	28.4 $\pm$ 4.1
Hopper-10% $\hat{m}$	31.2 $\pm$ 4.1	55.0 $\pm$ 3.5	34.7 $\pm$ 3.9	58.2 $\pm$ 3.2

C. Ablation Studies

Table II isolates each component of CDA using CQL on Hopper-m. Removing the plausibility filter and keeping all generated counterfactuals hurts performance relative to CQL alone, confirming that unfiltered diffusion samples introduce harmful OOD transitions. Dropping only the inverse dynamics check while keeping forward uncertainty filtering reduces the gain by 5.8 points, showing that the two criteria capture different failure modes. Using uniformly random counterfactual actions degrades results further, as many random actions lead to states far off the data manifold.

TABLE II.
ABLATION ON HOPPER-M, CQL+CDA (NORMALIZED RETURN, 5 SEEDS).

Variant	Return
CQL (no augmentation)	58.5 $\pm$ 3.2
CDA full	68.9 $\pm$ 2.1
No filter (all counterfactuals)	55.2 $\pm$ 4.3
No inverse dynamics check	63.1 $\pm$ 2.8
Uniform random $\hat{a}$ t	61.5 $\pm$ 2.5

D. Offline-to-Online Fine-Tuning

After offline training on Hopper-m, the agent is given 20k online steps. CDA-augmented CQL reaches a return of 86.2 versus 78.5 for CQL alone, with no drop at the offline-to-online transition. Continued use of filtered counterfactuals during fine-tuning accounts for both the gain and the stability.

VI. DISCUSSION

Why does counterfactual augmentation help? Offline RL policies become uncertain in regions absent from the dataset. CDA populates those regions with transitions that stay near the data manifold, checked by both the forward diffusion model and the inverse dynamics model. This broadens the training distribution without the OOD exploitation that unconstrained rollouts produce.

Limitations. Training the diffusion model adds upfront cost, though DDIM inference with 3 steps keeps per-sample cost low. Reward relabelling by a learned reward model introduces bias; tasks with known reward functions (e.g.,

D4RL) avoid this. Currently, CDA generates one counterfactual per real transition; an adaptive augmentation ratio could capture more of the benefit in settings where the filter acceptance rate is high.

VII. CONCLUSION

CDA generates counterfactual transitions with a conditional diffusion model, discards implausible ones with a two-stage plausibility filter, and uses the remainder to augment standard offline RL algorithms. On D4RL MuJoCo and Adroit tasks, this raises returns by 8–27% and nearly doubles performance under 10% of the data, with no modification to the underlying RL objective. The same model supports offline-to-online transfer without catastrophic forgetting. Extending CDA to image observations and combining it with uncertainty-aware offline algorithms are natural directions for future work.

ACKNOWLEDGMENT

The authors would like to thank the Department of Computer Science and Engineering and the Department of Mechanical Engineering at the Institute of Technology and Management, Gwalior, for providing the academic support that made this work possible.

REFERENCES

- [1] S. Levine, A. Kumar, G. Tucker, and J. Fu, "Offline reinforcement learning: Tutorial, review, and perspectives on open problems," arXiv:2005.01643, 2020.
- [2] M. Laskin, K. Lee, A. Stooke, L. Pinto, P. Abbeel, and A. Srinivas, "Reinforcement learning with augmented data," NeurIPS, 2020.
- [3] S. Sinha, A. Mandlekar, and A. Garg, "S4RL: Surprisingly simple self-supervision for offline reinforcement learning," CoRL, 2022.
- [4] T. Yu, G. Thomas, L. Yu, S. Ermon, J. Y. Zou, S. Levine, C. Finn, and T. Ma, "MOPO: Model-based offline policy optimization," NeurIPS, 2020.
- [5] M. Janner, J. Fu, M. Zhang, and S. Levine, "When to trust your model: Model-based policy optimization," NeurIPS, 2019.
- [6] J. Ho, A. Jain, and P. Abbeel, "Denoising diffusion probabilistic models," NeurIPS, 2020.
- [7] Y. Song, J. Sohl-Dickstein, D. P. Kingma, A. Kumar, S. Ermon, and B. Poole, "Score-based generative modeling through stochastic differential equations," ICLR, 2021.
- [8] L. Buesing et al., "DIAMOND: Diffusion world models," arXiv preprint, 2024.
- [9] M. Janner, Y. Du, J. B. Tenenbaum, and S. Levine, "Planning with diffusion for flexible behavior synthesis," ICML, 2022.
- [10] C. Lu, B. Scholkopf, and J. M. Hernandez-Lobato, "Counterfactual data augmentation for machine learning," ICML Workshop, 2020.
- [11] M. Oberst and D. Sontag, "Counterfactual off-policy evaluation with G-estimation," NeurIPS, 2021.
- [12] A. Kumar, A. Zhou, G. Tucker, and S. Levine, "Conservative Q-learning for offline reinforcement learning," NeurIPS, 2020.
- [13] I. Kostrikov, A. Nair, and S. Levine, "Offline reinforcement learning with implicit Q-learning," ICML, 2021.
- [14] S. Fujimoto and S. S. Gu, "A minimalist approach to offline reinforcement learning," NeurIPS, 2021.

- [15] D. Yarats, R. Fergus, A. Lazaric, and L. Pinto, "Image augmentation is all you need: Regularizing deep reinforcement learning from pixels," ICLR, 2021.
- [16] T. Yu, A. Kumar, R. Rafailov, A. Rajeswaran, S. Levine, and C. Finn, "COMBO: Conservative offline model-based policy optimization," NeurIPS, 2021.
- [17] B. Lakshminarayanan, A. Pritzel, and C. Blundell, "Simple and scalable predictive uncertainty estimation using deep ensembles," NeurIPS, 2017.
- [18] J. Fu, A. Kumar, O. Nachum, G. Tucker, and S. Levine, "D4RL: Datasets for deep data-driven reinforcement learning," arXiv:2004.07219, 2020.