

Persistent Safety Set Guided Offline Safe Reinforcement Learning

Ayan Choudhury¹, Janaka Brahmanage², Akshat Kumar² and Praveen Paruchuri¹

¹International Institute of Information Technology, Hyderabad

²Singapore Management University

{ayan.choudhury@research, praveen.p}@iiit.ac.in, {janakat.2022@phdcs, akshatkumar}@smu.edu.sg

Abstract

Offline safe reinforcement learning learns high-return policies that satisfy hard safety constraints using only a pre-collected dataset. This setting is challenging due to the inability to explore, and the risk of propagating value errors through unsafe state-space regions. To address this, *first*, we characterize the safe state region by developing a framework for learning control barrier functions (CBFs) using a novel *generalized Bellman* operator, yielding a *persistent safety set*, from which the agent can remain safe indefinitely. *Second*, we show that several existing safety set estimation methods (e.g., reachability-constrained RL) can be formulated within our CBF learning framework, highlighting its generality. We further propose a new CBF that ensures safety under environment dynamics uncertainty, unlike standard CBFs designed for deterministic settings. *Third*, we propose a new reward maximization algorithm that effectively exploits our learned persistent safety set for reward critic estimation. Empirical results on standard benchmarks show that our approach achieves state-of-the-art safety with fewer constraint violations while maintaining competitive returns.

1 Introduction

Safety-critical tasks such as autonomous driving [Gulino *et al.*, 2023; Osinski *et al.*, 2019], robotics [Gu *et al.*, 2023], and healthcare [Fang *et al.*, 2024] require RL algorithms [Sutton and Barto, 2020] that maximize performance while guaranteeing constraint satisfaction [Wachi *et al.*, 2024]. A single violation can cause irrecoverable harm, making online exploration risky. In offline RL, where learning is from static datasets [Fujimoto *et al.*, 2019], ensuring safety is even more challenging due to risks of unsafe generalization, value overestimation, and reward propagation through unsafe regions. Prior work on safe RL has mainly used constrained optimization, introducing Lagrange multipliers [Gu *et al.*, 2024; Wachi *et al.*, 2024; Achiam *et al.*, 2017; Tessler *et al.*, 2018]. However, these approaches only encourage safety in expectation. An alternative line formulates safety via reachability:

RCRL [Yu *et al.*, 2022] analyzes worst-case constraint evolution, which BCRL [Brahmanage and Kumar, 2026] extends to cumulative constraints, while FISOR [Zheng *et al.*, 2024] uses diffusion models to estimate safe trajectories. Inspired by control theory, these methods use control barrier functions (CBFs) to enforce forward invariance [Ames *et al.*, 2019]. Recent work further shows that value functions can act as barrier certificates [Tan *et al.*, 2023].

Recent surveys on learning CBFs [Wachi *et al.*, 2024; Ganai *et al.*, 2024] emphasize their importance for connecting verification, optimization, and learning, but also highlight challenges: (a) learning barrier certificates from limited data, (b) integrating barrier constraints into value learning, and (c) preventing safety information “leakage” through unsafe regions. Most offline safe RL algorithms [Chemingui *et al.*, 2025b; Koirala *et al.*, 2025] enforce safety only at policy extraction, allowing reward propagation through unsafe regions in the value critic and weakening safety guarantees. Our work addresses these issues by learning CBFs from offline data via a generalized Bellman operator and introducing principled methods to prevent safety leakage.

Related work. Recent work has connected control barrier functions (CBFs) and reinforcement learning (RL), especially in data-driven and offline settings. Neural CBFs trained from demonstrations or safe trajectories enable safety without explicit models [Harms *et al.*, 2024; Lindemann *et al.*, 2024; So *et al.*, 2024]. Some approaches estimate CBFs from offline data, but often rely on conservative assumptions or are limited to specific systems [Yu *et al.*, 2025; Tabbara and Sibai, 2025]. In offline RL, strategies like boundary-to-region supervision, diffusion regularization, and online optimization layers [Su *et al.*, 2025; Guo *et al.*, 2025a; Chemingui *et al.*, 2025a] mitigate unsafe generalization, but do not provide a unified framework for learning persistent safety sets via Bellman updates. Our work addresses this gap by formulating safety as a global, trajectory-level property, propagating worst-case violations through Bellman updates, and integrating safety critics with offline value-based RL to prevent unsafe value propagation during training.

Contributions. Our main contributions are as follows:

- **Generalized Bellman operator for CBF learning:** We develop a framework for learning discrete time CBFs via a generalized Bellman operator, and show that sev-

eral existing CBFs are special cases. The fixed point of this operator corresponds to the largest forward-invariant safe set, providing a data-driven safety certificate.

- **Forward-invariant safety enforcement:** We integrate learned CBFs as an explicit safety critic in our offline safe RL algorithm, embedding them directly into actor-critic learning so that both value estimation and policy improvement are restricted to barrier-certified safe regions, preventing value propagation via unsafe regions, while still achieving high returns.
- **Empirical validation:** We demonstrate state-of-the-art safety constraint satisfaction with competitive returns on standard offline safe RL benchmarks.

Additional results, code and full paper with proofs are available at <https://ayanroy24.github.io/GDCBF/>.

2 Preliminaries

2.1 Offline Safe Reinforcement Learning

We formulate safe sequential decision making as a constrained MDP (CMDP), and subsequently address the challenges of learning from an offline dataset.

Constrained MDP. A CMDP is defined by the tuple $\mathcal{M} := \langle \mathcal{S}, \mathcal{A}, \mathcal{T}, r, c, \gamma, \mu_0, C_{\text{lim}} \rangle$. Here, $\mathcal{S}$ and $\mathcal{A}$ denote the state and action spaces, respectively [Altman, 2021]. The transition distribution $\mathcal{T}(s' | s, a)$ denotes the probability of transitioning to state s' given state s and action a . Reward function is defined as $r : \mathcal{S} \times \mathcal{A} \rightarrow \mathbb{R}$ and the cost function as $c : \mathcal{S} \times \mathcal{A} \rightarrow \mathbb{R}^+$. The reward discount factor is $\gamma \in [0, 1)$, and μ_0 denotes the initial state distribution. The scalar C_{lim} represents the safety budget, defined as the maximum allowable expected cumulative cost.

The objective is to learn a policy $\pi : \mathcal{S} \rightarrow \mathcal{A}$ that maximizes the expected discounted return $J_R(\pi)$ subject to the safety constraint $J_C(\pi) \leq C_{\text{lim}}$, formally stated as:

$$\begin{aligned} \underset{\pi}{\text{maximize}} \quad & J_R(\pi) = \mathbb{E}_{\pi} \left[\sum_{t=0}^{\infty} \gamma^t r(s_t, a_t) \right] \\ \text{subject to} \quad & J_C(\pi) = \mathbb{E}_{\pi} \left[\sum_{t=0}^{\infty} \gamma^t c(s_t, a_t) \right] \leq C_{\text{lim}} \end{aligned} \quad (1)$$

where $\mathbb{E}_{\pi}$ denotes the expectation over trajectories induced by π and $\mathcal{T}$. In this work, we set the safety budget $C_{\text{lim}} = 0$, addressing the challenging *hard constraint* setting where any constraint violation is inadmissible.

Learning from Offline Data. In complex real-world domains, obtaining an accurate analytical model of the environment (e.g., transition dynamics $\mathcal{T}$) is often intractable. Consequently, we adopt the *offline* safe RL setting, where the agent must learn solely from a fixed dataset $\mathcal{D} = \{(s_i, a_i, r_i, c_i, s'_i)\}_{i=1}^N$ collected by one or more behavioral policies π_{β} , without additional environment interactions.

2.2 Discrete Control Barrier Value Functions

Expected-cost constraints in CMDPs do not prevent transient or worst-case violations. In safety-critical domains, a single violation may be unacceptable regardless of the expected

cost. Control barrier functions (CBFs) provide a stronger safety mechanism by characterizing safety through forward invariance of a feasible set. Next, we adapt CBFs [Agrawal and Sreenath, 2017] to discrete-time offline safe RL.

Safety is enforced through inequalities on value functions, enabling barrier-based reasoning directly from offline datasets. We learn CBFs via a Bellman-style update rule to learn a safety-critic $Q_h(s, a)$ that serves as a safety certificate. We define the safe set $\mathcal{C}$ based on state-action pairs, where the function $h : \mathcal{S} \times \mathcal{A} \rightarrow \mathbb{R}$ serves as a safety indicator.

Definition 2.1 (Safe Set). The safe set $\mathcal{C}$ is defined as the set of state-action pairs satisfying the condition:

$$\mathcal{C} = \{(s, a) \in \mathcal{S} \times \mathcal{A} \mid h(s, a) \leq 0\} \quad (2)$$

where $h(s, a)$ represents the immediate safety indicator (e.g., distance to a safety boundary), and positive values indicate a violation.

For CMDPs, the safety indicator $h(s, a)$ is directly derived from the cost function $c(s, a)$. Following [Zheng *et al.*, 2024], let $h(s, a) = c(s, a)$ if the transition is unsafe ($c(s, a) > 0$), and $h(s, a) = -1$ otherwise. This mapping ensures that $h(s, a) > 0$ corresponds to constraint violations, while $h(s, a) \leq 0$ indicates safety compliance.

While $\mathcal{C}$ captures instantaneous safety, relying solely on it is myopic, as it does not ensure safety can be maintained over time. Guaranteeing long-term safety requires identifying regions where the dynamics allow indefinite constraint satisfaction, leading to two notions: forward invariance under a given policy, and the existence of at least one safe policy, which defines the persistent safety set. For exposition ease, we assume below a deterministic MDP; stochastic setting is addressed later.

Definition 2.2. A set $\mathcal{S}_{\pi} \subseteq \mathcal{S}$ is **forward invariant** with respect to a policy π if, for any initial state $s_0 \in \mathcal{S}_{\pi}$, the subsequent state trajectories generated by the deterministic transition function $f : \mathcal{S} \times \mathcal{A} \rightarrow \mathcal{S}$, $s_{t+1} = f(s_t, \pi(s_t))$, satisfy:

$$s_t \in \mathcal{S}_{\pi} \quad \forall t \geq 0. \quad (3)$$

If $\mathcal{S}_{\pi}$ is a subset of the safe states (where $(s, \pi(s)) \in \mathcal{C}$), then the policy π guarantees safety for all time steps when initialized within $\mathcal{S}_{\pi}$, and thus π satisfies the CMDP safety constraint in (1).

Definition 2.3. The **persistent safety set** $\mathcal{S}_P \subseteq \mathcal{S}$ is defined as the set of all states for which there exists a policy π such that the agent remains safe indefinitely. Formally:

$$\mathcal{S}_P = \{s \in \mathcal{S} \mid \exists \pi, \forall t \geq 0, (s_t, a_t) \in \mathcal{C} \text{ given } s_0 = s\} \quad (4)$$

This set represents the maximal controlled invariant set within the safety constraints. Staying in this set enforces CMDP constraint in (1).

Explicitly computing set $\mathcal{S}_P$ requires solving a global reachability problem [Bansal *et al.*, 2017], which is computationally intractable [Mitchell *et al.*, 2005]. CBFs provide a tractable alternative. By satisfying a specific rate-of-change condition, the CBF function estimates the persistent safety set, effectively identifying the boundary of safe operation [Ames *et al.*, 2019].

Specifically, we learn a safety-critic $Q_h(s, a)$ that serves as a certificate for persistent safety. The indicator $h(s, a)$ captures whether an action is *immediately* safe, $Q_h(s, a)$ accounts for the worst-case future safety violations that may occur along subsequent trajectories. Intuitively, if $Q_h(s, a) \leq 0$, the state-action pair is *persistently* safe, meaning there exists a policy that can maintain safety indefinitely. To ensure forward invariance, we define the following condition:

Definition 2.4 (Deterministic Discrete CBF). Consider a deterministic MDP with transition function f (Definition 2.2). A function $Q_h(s, a)$ is a discrete CBF if there exists a policy π such that for all (s, a) with $Q_h(s, a) \leq 0$:

$$Q_h(f(s, a), \pi(f(s, a))) \leq Q_h(s, a) - \alpha(Q_h(s, a)) \quad (5)$$

where $\alpha(\cdot)$ is an extended class $\mathcal{K}$ function [Agrawal and Sreenath, 2017].

An extended class- $\mathcal{K}$ function $\alpha : \mathbb{R} \rightarrow \mathbb{R}$ is continuous, strictly increasing, with $\alpha(0) = 0$; it modulates the decay rate in Eq.5, permitting residency in the safe set for $Q_h \leq 0$ while enforcing a corrective decrease toward safety for $Q_h > 0$. Hence $\alpha(q) > 0$ for $q > 0$ and $\alpha(q) < 0$ for $q < 0$, so $\alpha(Q_h)$ has the same sign as Q_h . A canonical choice is the linear form $\alpha(q) = \lambda q$ with $\lambda \in (0, 1]$. Intuitively, this inequality guarantees forward invariance by restricting the evolution of Q_h so that the agent is strictly prevented from crossing the safety boundary. The function $\alpha(\cdot)$ modulates this restriction, effectively dictating how aggressively the policy must steer away from the unsafe region based on the current safety margin. In stochastic environments, relying on a single next state or an expectation is insufficient for strict safety guarantees. We therefore adopt a robust formulation.

Definition 2.5 (Robust Discrete CBF Condition). For a stochastic MDP defined by $\mathcal{T}(\cdot|s, a)$, a function $Q_h(s, a)$ is a robust discrete CBF if there exists a policy π such that for all (s, a) with $Q_h(s, a) \leq 0$:

$$\max_{s' \in \text{supp}(\mathcal{T}(\cdot|s, a))} Q_h(s', \pi(s')) \leq Q_h(s, a) - \alpha(Q_h(s, a)) \quad (6)$$

where $\text{supp}(\mathcal{T}(\cdot|s, a))$ denotes the set of possible next states (or support), and $\alpha(\cdot)$ is an extended class $\mathcal{K}$ function.

This ensures that, even under worst-case transitions, the safety value remains decreasing (or sufficiently negative) over time. Small increases in Q_h may occur due to stochasticity, but α ensures forward invariance is maintained without requiring strict decrease.

3 Learning CBFs via Bellman Updates

3.1 Generalized Bellman Operator

Learning control barrier functions from data presents several challenges in the offline reinforcement learning setting. Classical CBF formulations typically assume access to system dynamics or require evaluating all possible successor states [Ames *et al.*, 2014], which is infeasible without an explicit model. Existing reinforcement learning approaches instantiate barrier functions using specific Bellman operators, such as hard maximum-based reachability updates [Yu *et al.*, 2022] or additive cost-style backups [Tan *et al.*, 2023],

each inducing different trade-offs between safety and conservatism. Moreover, offline datasets provide limited transition coverage [Fujimoto *et al.*, 2019], requiring safety propagation mechanisms that reason about worst-case future violations rather than expected behavior. These challenges motivate the need for a generalized Bellman backup framework that unifies existing formulations while ensuring convergence and forward-invariant safety properties.

To learn this function, we define a generalized Bellman-like operator that propagates the worst-case safety constraint violation. This effectively enforces a reachability constraint: if a future state is unsafe, the current value must reflect that potential violation.

Assumption 3.1. The function $\Psi : \mathbb{R} \times \mathbb{R} \rightarrow \mathbb{R}$ satisfies the following properties: it is strictly increasing in its second argument. Furthermore, $\Psi(h, \cdot)$ is a contraction in the second argument, i.e., $|\Psi(h, v_1) - \Psi(h, v_2)| \leq \gamma |v_1 - v_2|$ with $\gamma < 1$ for all h .

Assumption 3.2. The function Ψ satisfies a consistency condition: for any h, v , if $q = \Psi(h, v)$, then $v \leq q - \alpha(q)$ for some extended class $\mathcal{K}$ function α . This ensures that the fixed point induces a valid decay for the safety value.

Intuitively, this condition ensures that any fixed point of the operator $\mathcal{B}$ automatically respects the required safety decay, guaranteeing that the learned value function acts as a valid barrier certificate.

Definition 3.1 (Generalized Reachability Bellman Operator). We define the operator $\mathcal{B}$ acting on the state-action value function:

$$\mathcal{B}Q_h(s, a) = \Psi \left(h(s, a), \max_{s' \in \text{supp}(\mathcal{T}(\cdot|s, a))} \min_{a' \in \mathcal{A}} Q_h(s', a') \right) \quad (7)$$

where Ψ mixes the immediate safety cost with the future cost.

3.2 Theoretical Analysis

Having defined the generalized operator, we now establish its theoretical properties. Specifically, we show that under the imposed assumptions on the mixing function Ψ , the operator defines a contraction mapping in the value function space. Consequently, its iterative application converges to a unique fixed point. Crucially, we prove that this fixed point satisfies the robust discrete CBF condition, thereby providing a valid certificate for the forward-invariant safety set.

Theorem 3.1. Consider a CMDP with continuous safety indicator $h(s, a)$. The iterative application of the operator $\mathcal{B}$ defined in (7) converges to a unique fixed point Q_h^* . Furthermore, Q_h^* satisfies the Robust Discrete CBF condition defined in (6), provided Assumptions 3.1 and 3.2 hold.

The proof of Theorem 3.1 is available in Appendix A.1.

As a consequence of the fixed-point characterization, the boundary of the persistent safety set is given by the zero level-set of the learned barrier value function. In particular, the set

$$\mathcal{C}^* = \{(s, a) \mid Q_h^*(s, a) \leq 0\}$$

defines a forward-invariant safety region, while its boundary is implicitly induced by the condition $Q_h^*(s, a) = 0$. Thus, the safety boundary is not learned explicitly or parameterized a priori, but emerges naturally from the generalized Bellman updates.

Maximal Forward-Invariant Safe Set. As safety in sequential decision-making is inherently global and set-valued, validity requires guarantees over entire trajectories, robustness to worst-case dynamics, and compatibility with value-based policy optimization. Moreover, in the offline setting, safety guarantees must be grounded in dataset-supported transitions rather than extrapolated dynamics. Accordingly, we establish additional properties of the fixed point Q_h^* . These results show that the zero sublevel set of Q_h^* is not merely forward invariant but constitutes the *maximal* set of state-action pairs from which safety can be maintained indefinitely, ruling out unnecessary conservatism. We use the standard pointwise order on value functions, i.e., $Q_1 \preceq Q_2$ if $Q_1(s, a) \leq Q_2(s, a)$ for all (s, a) , and define the safe set induced by Q as $\mathcal{C}_Q = \{(s, a) \mid Q(s, a) \leq 0\}$.

Lemma 3.2. A set $\mathcal{C} \subseteq \mathcal{S} \times \mathcal{A}$ is forward invariant if there exists a policy π such that, for all $(s, a) \in \mathcal{C}$ and all $s' \in \text{supp}(\mathcal{T}(\cdot|s, a))$, we have $(s', \pi(s')) \in \mathcal{C}$. A forward-invariant safe set $\mathcal{C}^*$ is maximal if, for any other forward-invariant safe set $\tilde{\mathcal{C}}$, it holds that $\tilde{\mathcal{C}} \subseteq \mathcal{C}^*$. Let Q_h^* be the fixed point of the operator $\mathcal{B}$ defined in (7). Then the set

$$\mathcal{C}^* := \{(s, a) \in \mathcal{S} \times \mathcal{A} \mid Q_h^*(s, a) \leq 0\}$$

is the maximal forward-invariant safe set induced by $\mathcal{B}$.

The proof of Lemma 3.2 is available in Appendix A.2.

4 Generalized Discrete CBF Variants

The theoretical framework established in Section 3 relies on a generalized mixing function Ψ . In practice, the choice of this function determines the nature of the learned safety certificate and its optimization landscape. In this section, we discuss what are the potential instantiations of the framework.

4.1 Deterministic Case

Maximum. The simplest instantiation of our framework aligns with the Reachability Constrained RL (RCRL) approach [Yu *et al.*, 2022]. Here, the operator propagates the maximum of the immediate safety cost and the discounted future safety value:

$$\mathcal{B}Q_h(s, a) = \max\{h(s, a), \gamma V_h(s')\} \quad (8)$$

where, $V_h(s') = \min_{a'} Q_h(s', a')$. This corresponds to $\Psi(h, v) = \max(h, \gamma v)$. The discount factor γ serves to decay the safety value over the horizon; specifically, it causes the safety margin of safe states (negative values) to degrade towards zero over time, while diminishing the magnitude of distant unsafe violations (positive values). While this operator strictly enforces the safety condition, it can suffer from value explosion if the horizon is long and costs are additive.

Smooth. The FISOR [Zheng *et al.*, 2024] algorithm employs a convex combination to smooth the transition between the immediate cost and the future value:

$$\mathcal{B}Q_h(s, a) = (1 - \gamma)h(s, a) + \gamma \max\{h(s, a), V_h(s')\} \quad (9)$$

This corresponds to $\Psi(h, v) = (1 - \gamma)h + \gamma \max(h, v)$. By incorporating the term $(1 - \gamma)h(s, a)$, this operator effectively anchors the value estimate to the immediate safety cost. This

anchoring ensures that the safety critic remains responsive to local constraint violations and stabilizes the learning process by balancing the immediate feedback with the bootstrapped lookahead.

Additive. The ‘‘Your Value Function is a Barrier Function’’ framework [Tan *et al.*, 2023] proposes an additive formulation where the Bellman operator aggregates the immediate safety cost and the discounted future value:

$$\mathcal{B}Q_h(s, a) = h(s, a) + \gamma V_h(s') \quad (10)$$

This corresponds to the linear aggregator $\Psi(h, v) = h + \gamma v$. Unlike the maximum-based operators, this formulation resembles the standard Bellman equation for cumulative cost. By interpreting the value function itself (shifted by a constant) as the barrier function, this approach ensures that the safety certificate is learned via standard value iteration techniques, provided the immediate cost $h(s, a)$ is designed to reflect the distance to the safety boundary.

Proposition 4.1. The Maximum, Smooth, and Additive update rules all satisfy Assumption 3.1 (Monotonicity and Contraction) with contraction factor γ . Furthermore, they satisfy the robust CBF condition (Assumption 3.2) with appropriate choice of α .

The proof of Proposition 4.1 is available in Appendix A.5.

4.2 Stochastic Case

The deterministic operators discussed above assume a single possible outcome for each action, which is often violated in real-world systems with aleatoric uncertainty or unmodeled dynamics. In such non-deterministic environments, the agent must satisfy the Robust Discrete CBF condition (Definition 2.5), which requires bounding the maximum violation over all possible next states: $\max_{s'} V_h(s')$.

In the offline setting, we lack an explicit transition model to query the set of reachable states. To overcome this, we propose to employ an *asymmetric aggregation* [Hansen-Estruch *et al.*, 2023] to approximate this worst-case future violation directly from the data distribution. We denote this asymmetric aggregation as E^τ , where the parameter τ can be adjusted to control the level of conservativeness when approximating the maximum. Common choices for such objectives include expectile regression, percentile (pinball) loss, or Gumbel regression [Garg *et al.*, 2023]. By tuning the parameter τ , we can estimate the upper tail of the future safety values observed in the dataset.

In practice, we apply the asymmetric aggregation *outside* the aggregator function Ψ to handle stochasticity in the transitions (s, a, s') . The update rule becomes:

$$\mathcal{B}Q_h(s, a) = \mathbb{E}_{s' \sim \mathcal{D}(s, a)}^\tau [\Psi(h(s, a), V_h(s'))] \quad (11)$$

As $\tau \rightarrow 1$, the expectile $\mathbb{E}^\tau$ approaches the maximum operator, it also approaches the worst-case max over the transition support as seen in Fig.2, thereby recovering the Robust Discrete CBF condition (Definition 2.5) which requires bounding the worst-case future violation. Note that Eq.11 adapts Eq.7 to the stochastic case: a direct max over transitions is impractical in offline data (identical (s, a) pairs rarely repeat), so asymmetric regression estimates the upper tail as a

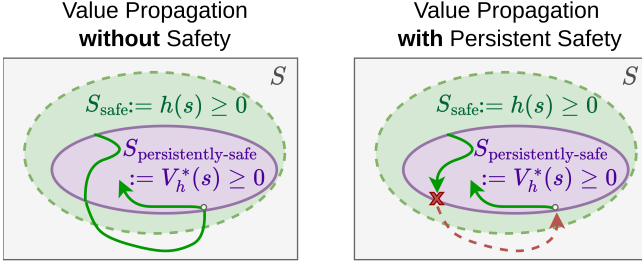


Figure 1: Value propagation with vs. without persistent safety. The figure shows a trajectory that briefly leaves the feasible region. Prior methods allow value propagation through unsafe states by enforcing safety only at policy extraction. Our approach enforces safety within the Bellman update, blocking value propagation through unsafe transitions.

data-efficient surrogate for worst-case estimation under limited coverage. In many practical implementations, Q-learning employs a Huber or L_1 loss [Mnih *et al.*, 2013] to handle outliers and stabilize training. The corresponding asymmetric variant for these robust losses is the percentile loss. Therefore, we utilize the percentile loss to implicitly approximate the robust target $\max_{s'} V_h(s')$ over the transition distribution, thereby enforcing safety against stochasticity in dynamics.

5 Utilizing Learned Persistent Safety Set for the Value Function Estimation

As illustrated in Figure 1, enforcing safety at the value function level prevents reward propagation through unsafe intermediate states—a common issue in prior works such as FISOR [Zheng *et al.*, 2024] that apply safety constraints only during policy extraction. We adopt an Implicit Q-Learning (IQL) [Kostrikov *et al.*, 2022] framework for offline RL to learn the optimal value function from the static dataset $\mathcal{D}$, while enforcing the persistent safety.

Safety-Aware Joint Bellman Target. A core challenge in offline safe RL is ensuring value estimates reflect safety constraints, even when the dataset contains unsafe trajectories. Discarding unsafe transitions is undesirable because they often contain useful information, including safe segments. To address this, we first learn the safety critic Q_h using the chosen CBF instantiation and then use it to define a modified Bellman target. We propose a joint Bellman target for the reward Q -function, Q_R , which explicitly handles safe and unsafe transitions based on learned persistent safety. The target $y_R(s, a, s')$ is defined based on the safety status predicted by our learned CBF Q_h^* :

$$y_R(s, a, s') = \begin{cases} r + \gamma V_R(s') & \text{if } Q_h^*(s, a) \leq 0 \\ \frac{r_{\min}}{1-\gamma} - Q_h^*(s, a) & \text{if } Q_h^*(s, a) > 0 \end{cases} \quad (12)$$

where $r_{\min} = \min_{(s,a,r,c,s') \in \mathcal{D}} r$ is the minimum reward observed across the entire offline dataset (a single dataset-wide constant, not a per- (s, a) minimum). For safe actions ($Q_h^*(s, a) \leq 0$), we update using the standard Bellman backup with the value function V_R . For unsafe actions

Algorithm 1 Barrier-Guided Offline Actor-Critic (CBF)

Require: Offline dataset $\mathcal{D} = \{(s, a, r, c, s')\}$, discount γ , expectiles (τ_r, τ_h) , temperature β , target rate η , samples N

- 1: Initialize actor π_θ , reward critic Q_r^ω , value V_r^ψ , safety critic Q_h^ϕ , safety value V_h^ν , and target networks
- 2: **for** each gradient step **do**
- 3: Sample a minibatch $\mathcal{M} = \{(s, a, r, c, s')\} \sim \mathcal{D}$
- 4: **Barrier target:** $y_h \leftarrow \Psi(c, V_h(s'))$ for $(s, a, s') \in \mathcal{M}$
- 5: **Safety value update:**
 $\nu \leftarrow \arg \min_\nu \mathbb{E}_{\mathcal{M}} [\mathcal{L}_{\tau_h} (Q_h^\phi(s, a) - V_h^\nu(s))]$
- 6: **Safety critic update:**
 $\phi \leftarrow \arg \min_\phi \mathbb{E}_{\mathcal{M}} [Q_h^\phi(s, a) - y_h]$
- 7: **Safety-aware reward target (Eq.(12):**

$$y_r = \begin{cases} r + \gamma V_r(s') & Q_h^\phi(s, a) \leq 0 \\ \frac{r_{\min}}{1-\gamma} - Q_h^\phi(s, a) & \text{otherwise} \end{cases}$$
- 8: **Reward critic update:**
 $\omega \leftarrow \arg \min_\omega \mathbb{E}_{\mathcal{M}} [(Q_r^\omega(s, a) - y_r)^2]$
- 9: **Reward value update:**

$$\psi \leftarrow \arg \min_\psi \mathbb{E}_{\mathcal{M}} [\mathcal{L}_{\tau_r} (Q_r^\omega(s, a) - V_r^\psi(s))]$$
- 10: **Policy update:** $w \leftarrow \exp(\beta(Q_r^\omega(s, a) - V_r^\psi(s)))$
 $\theta \leftarrow \arg \min_\theta \mathbb{E}_{\mathcal{M}} [-w \log \pi_\theta(a|s)]$
- 11: Update target networks
- 12: **end for**
- 13: **Evaluation:** sample $\{a_i\}_{i=1}^N \sim \pi_\theta(\cdot|s)$,
- 14: return $a^* = \arg \min_{a_i} Q_h(s, a_i)$

($Q_h^*(s, a) > 0$), we assign a penalty term. This penalty consists of the worst-case discounted return $\frac{r_{\min}}{1-\gamma}$ minus the magnitude of the safety violation $Q_h^*(s, a)$. This ensures that unsafe actions have strictly lower Q -values than safe actions, effectively creating a “soft” barrier in the value landscape.

We define *reward leakage* as the propagation of reward through unsafe transitions during Bellman updates, causing the reward critic to overestimate values for policies that violate safety. Our safety-aware target (Eq. 12) directly penalizes such transitions during training, preventing leakage. In contrast, FISOR [Zheng *et al.*, 2024] enforces safety only at policy extraction, allowing value propagation through unsafe regions. Enforcing safety at the value level ensures the Bellman operator converges to a safety-consistent fixed point as shown in figure 3 and the toy example in Appendix D. The reward Q -function is updated by minimizing the Mean Squared Error (MSE) loss:

$$L_Q(\theta) = \mathbb{E}_{(s,a,r,s') \sim \mathcal{D}} [(y_R(s, a, s') - Q_R(s, a))^2] \quad (13)$$

Value Function Estimation via Expectile Regression. We then learn the state-value function $V_R(s)$ via expectile regression over the learned Q_R values:

$$L_V(\psi) = \mathbb{E}_{(s,a) \sim \mathcal{D}} [L_2^\tau(Q_R(s, a) - V_R(s))] \quad (14)$$

By utilizing the safety-aware Q_R , the expectile regression (with $\tau > 0.5$) naturally approximates the value of the best *safe* actions in the dataset distribution. In states where both

	CAPS		CCAC		LSPC		FISOR		CBF _{smooth}	
	reward $\uparrow$	cost $\downarrow$	reward $\uparrow$	cost $\downarrow$	reward $\uparrow$	cost $\downarrow$	reward $\uparrow$	cost $\downarrow$	reward $\uparrow$	cost $\downarrow$
CarButton1	-0.04 ± 0.04	0.98 ± 0.53	0.45 ± 0.06	11.31 ± 0.06	-0.03 ± 0.03	0.59 ± 0.11	-0.02 ± 0.05	0.26 ± 0.24	0.12 ± 0.04	0.47 ± 0.42
CarButton2	-0.15 ± 0.02	0.56 ± 0.01	0.53 ± 0.01	9.82 ± 0.01	-0.13 ± 0.06	0.93 ± 0.09	0.01 ± 0.00	0.58 ± 0.33	0.05 ± 0.06	0.13 ± 0.15
CarGoal1	0.23 ± 0.04	1.09 ± 0.57	0.84 ± 0.13	1.81 ± 0.14	0.27 ± 0.22	0.24 ± 0.47	0.49 ± 0.03	0.83 ± 0.16	0.55 ± 0.33	0.27 ± 0.41
CarGoal2	0.14 ± 0.08	0.52 ± 0.01	0.90 ± 0.05	5.53 ± 0.05	0.16 ± 0.41	0.41 ± 0.36	0.06 ± 0.03	0.33 ± 0.31	0.18 ± 0.12	0.20 ± 0.35
CarPush1	0.20 ± 0.00	0.32 ± 0.00	-0.12 ± 0.00	0.59 ± 0.09	0.18 ± 0.07	0.41 ± 0.12	0.28 ± 0.04	0.28 ± 0.15	0.30 ± 0.16	0.13 ± 0.28
CarPush2	0.05 ± 0.03	2.31 ± 1.03	0.12 ± 0.03	6.06 ± 0.03	0.07 ± 0.05	0.80 ± 0.02	0.14 ± 0.07	0.89 ± 0.51	0.20 ± 0.12	0.23 ± 0.25
AntVelocity	0.87 ± 0.00	0.22 ± 0.01	0.51 ± 0.09	0.59 ± 0.11	0.97 ± 0.28	0.14 ± 0.04	0.89 ± 0.00	0.00 ± 0.00	0.88 ± 0.02	0.17 ± 0.30
HalfCheetahVelocity	0.87 ± 0.00	0.29 ± 0.01	0.94 ± 0.15	0.94 ± 0.02	0.97 ± 0.37	1.00 ± 0.81	0.89 ± 0.01	0.00 ± 0.00	0.93 ± 0.01	0.00 ± 0.00
SwimmerVelocity	0.41 ± 0.10	2.64 ± 2.85	-0.08 ± 0.01	0.00 ± 0.00	0.46 ± 0.19	1.28 ± 0.15	-0.04 ± 0.03	0.00 ± 0.00	0.56 ± 0.03	0.21 ± 0.36
Average (9)	0.29 ± 0.37	0.99 ± 0.89	0.41 ± 0.33	3.18 ± 3.06	0.34 ± 0.31	1.10 ± 1.08	0.30 ± 0.37	0.35 ± 0.35	0.42 ± 0.33	0.24 ± 0.33
AntCircle	0.32 ± 0.00	0.00 ± 0.00	0.54 ± 0.05	0.00 ± 0.00	0.35 ± 0.14	1.66 ± 0.18	0.20 ± 0.00	0.00 ± 0.00	0.43 ± 0.13	0.12 ± 0.21
AntRun	0.53 ± 0.00	1.83 ± 0.00	0.02 ± 0.00	0.00 ± 0.00	0.75 ± 0.03	15.58 ± 0.04	0.45 ± 0.08	0.03 ± 0.01	0.63 ± 0.09	0.53 ± 0.31
BallCircle	0.33 ± 0.00	0.00 ± 0.00	0.67 ± 0.14	0.00 ± 0.00	0.50 ± 0.31	0.14 ± 0.07	0.34 ± 0.01	0.00 ± 0.00	0.44 ± 0.02	0.22 ± 0.21
BallRun	0.10 ± 0.02	0.00 ± 0.00	0.31 ± 0.04	0.00 ± 0.00	0.17 ± 0.09	0.00 ± 0.00	0.18 ± 0.01	0.00 ± 0.00	0.27 ± 0.02	0.00 ± 0.00
CarCircle	0.42 ± 0.02	0.01 ± 0.01	0.68 ± 0.08	0.74 ± 0.06	0.70 ± 0.47	0.36 ± 0.11	0.40 ± 0.02	0.11 ± 0.20	0.58 ± 0.03	0.12 ± 0.21
CarRun	0.96 ± 0.00	0.12 ± 0.00	0.92 ± 0.09	0.29 ± 0.02	0.97 ± 0.42	1.52 ± 0.39	0.73 ± 0.02	0.14 ± 0.22	0.84 ± 0.07	0.23 ± 0.34
DroneCircle	0.34 ± 0.01	0.00 ± 0.00	0.14 ± 0.01	0.17 ± 0.06	0.57 ± 0.04	2.74 ± 0.06	0.48 ± 0.00	0.00 ± 0.00	0.37 ± 0.02	0.38 ± 0.44
DroneRun	0.42 ± 0.04	17.69 ± 8.48	0.40 ± 0.13	15.86 ± 0.16	0.56 ± 0.16	0.00 ± 0.00	0.30 ± 0.03	0.55 ± 0.55	0.55 ± 0.01	0.00 ± 0.00
Average (8)	0.43 ± 0.25	2.46 ± 6.19	0.46 ± 0.30	2.13 ± 5.55	0.57 ± 0.25	2.75 ± 5.28	0.39 ± 0.18	0.10 ± 0.19	0.59 ± 0.19	0.18 ± 0.16
easysparse	0.20 ± 0.15	1.05 ± 1.07	-0.05 ± 0.00	0.32 ± 0.00	0.85 ± 0.16	5.76 ± 3.9	0.38 ± 0.06	0.53 ± 0.30	0.52 ± 0.01	0.46 ± 0.15
easymean	0.15 ± 0.10	0.61 ± 0.34	-0.06 ± 0.01	0.14 ± 0.10	0.78 ± 0.05	3.59 ± 1.27	0.38 ± 0.05	0.25 ± 0.16	0.38 ± 0.07	0.36 ± 0.37
easydense	0.10 ± 0.01	0.57 ± 0.31	-0.05 ± 0.00	0.24 ± 0.06	0.61 ± 0.39	5.20 ± 3.26	0.36 ± 0.03	0.25 ± 0.13	0.55 ± 0.21	0.61 ± 0.18
mediumsparse	0.46 ± 0.12	0.40 ± 0.63	-0.08 ± 0.00	0.24 ± 0.06	0.95 ± 0.07	4.24 ± 1.11	0.42 ± 0.05	0.22 ± 0.36	0.54 ± 0.04	0.23 ± 0.46
mediummean	0.39 ± 0.03	0.33 ± 0.25	-0.06 ± 0.01	0.06 ± 0.09	0.97 ± 0.03	4.14 ± 0.93	0.39 ± 0.04	0.08 ± 0.07	0.54 ± 0.15	0.30 ± 0.36
mediumdense	0.50 ± 0.13	0.36 ± 0.20	-0.07 ± 0.00	0.20 ± 0.00	0.70 ± 0.39	2.77 ± 1.73	0.49 ± 0.05	0.44 ± 0.32	0.53 ± 0.16	0.26 ± 0.15
hardsparse	0.32 ± 0.04	1.76 ± 0.32	-0.05 ± 0.00	0.24 ± 0.06	0.51 ± 0.11	1.99 ± 0.66	0.30 ± 0.02	0.01 ± 0.01	0.42 ± 0.02	0.77 ± 0.07
hardmean	0.18 ± 0.01	0.17 ± 0.15	-0.05 ± 0.00	0.20 ± 0.00	0.53 ± 0.12	4.48 ± 1.63	0.26 ± 0.02	0.09 ± 0.07	0.36 ± 0.06	0.28 ± 0.16
harddense	0.25 ± 0.02	1.03 ± 0.35	-0.03 ± 0.01	0.34 ± 0.12	0.49 ± 0.12	3.56 ± 1.87	0.30 ± 0.03	0.34 ± 0.31	0.36 ± 0.09	0.87 ± 0.42
Average (9)	0.28 ± 0.14	0.70 ± 0.50	-0.06 ± 0.02	0.22 ± 0.10	0.71 ± 0.16	3.97 ± 1.82	0.36 ± 0.07	0.25 ± 0.17	0.47 ± 0.08	0.46 ± 0.24

Table 1: Results for normalized cost and normalized reward where $\uparrow$ indicates that a higher value is better, while $\downarrow$ indicates that a lower value is preferable. **Bold** text denotes safe policies, **Red** indicates unsafe policies, and **Blue** highlights the highest reward among the safe policies for each task. Each result is obtained as an average by running over three random seeds and evaluating over 20 episodes.

safe and unsafe actions are present, the high penalty on unsafe actions ensures they fall below the expectile, causing $V_R(s)$ to converge towards the value of the safe actions. In states where only unsafe actions are available, $V_R(s)$ correctly reflects the low, penalized value, signaling the agent to avoid such states. In practice, we can either learn the CBF as a pre-training step and then learn the reward critic Q_R , or learn both in parallel. The latter approach aligns more closely with standard offline RL, which typically extracts the policy while learning the Q-value. Following this approach, we learn all three components—the CBF Q_h , the reward critic Q_R , and the policy π —concurrently, as presented in the algorithm pseudocode.

Pseudocode and Computational cost. Algorithm 1 summarizes our method where action-value functions (Q_h, Q_r) are parameterized to support Bellman updates and policy improvement, while value functions (V_h, V_r) are also learned via function approximation but optimized using expectile regression to provide stable targets. At each step, we fit the safety critic Q_h using a Bellman-like update, then use it to filter transitions for the reward critic Q_r : safe transitions use standard bootstrapped targets, while unsafe ones receive a pessimistic penalty based on $r_{\min}$. The policy is improved via advantage-weighted regression. This approach prevents reward propagation through unsafe states and ensures efficient, stable offline

training. On a single RTX 2080 Ti GPU, 500k gradient steps require about 20 minutes.

6 Experiments

Goal of our experimental evaluation is:

1. Can a learned discrete control barrier function (CBF) effectively enforce hard safety constraints in the offline setting while obtaining high rewards?
2. Assess how different deterministic CBFs—CBF_{smooth} Eq.(9), CBF_{additive} Eq.(10), and CBF_{maximum} Eq.(8)—perform in the offline setting.
3. Does the robust smooth barrier update (section 4.2), utilizing an asymmetric loss, provides a safety advantage over the standard version in the presence of uncertainty?
4. Ablation studies to analyze the algorithm’s sensitivity to hyperparameters (see Appendix B.2).

Evaluation Setup. We evaluate on diverse safety-critical continuous control tasks (Safety Gym, Bullet Gym, MetaDrive) using the same experimental setup as FISOR [Zheng *et al.*, 2024]. Rewards are normalized by the dataset’s min and max, and costs by a predefined limit. A policy is considered safe if its normalized cost does not exceed one. Further details are in Appendix C.

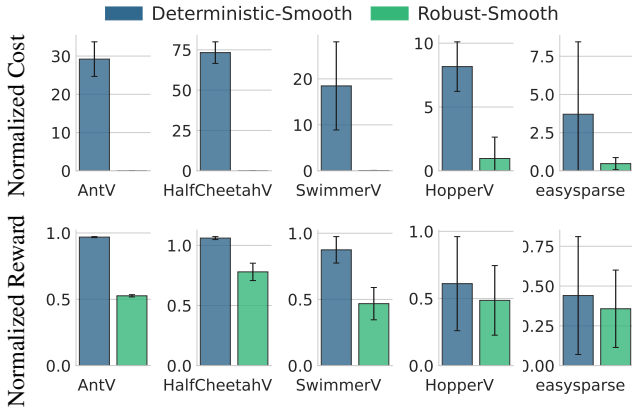


Figure 2: Effect of asymmetric aggregation on safety and reward propagation. This experiment evaluates whether the robust CBF variant in Sec. 4.2 improves performance under uncertainty. We add outliers to DSRL and compare Standard-Smooth with Robust-Smooth, which uses an asymmetric loss to approximate the maximum cost over the transition function. Across tasks, Robust-Smooth achieves lower normalized costs with comparable rewards.

Comparison with Baselines. We compare our method to safe offline RL baselines: CAPS [Chemingui *et al.*, 2025b], LSPC [Koirala *et al.*, 2025] (soft constraints), and CCAC [Guo *et al.*, 2025b], FISOR [Zheng *et al.*, 2024] (hard constraints). As shown in Table 1, methods enforcing safety only at policy extraction (e.g., CAPS, CCAC, LSPC) often achieve high rewards but violate constraints. In contrast, our approach (Algorithm 1) maintains near-zero cost across tasks, enforcing hard safety while achieving top rewards among safe policies. This is especially clear in recovery-heavy tasks, where restricting value backups to the barrier-certified set preserves high returns and prevents unsafe transitions. Thus, integrating a learned discrete CBF into value critic learning (Section 5) enables strict safety without sacrificing performance. FISOR and our method both satisfy the cost criterion, but our safety-aware Bellman target yields more accurate value estimates, supporting both high reward, safety.

Comparison of CBFs. A comparison of different barrier formulations is available in Appendix B.1, which reveals distinct trade-offs. For instance, $\text{CBF}_{\text{maximum}}$ tends to optimize for reward, while $\text{CBF}_{\text{smooth}}$ yields the most stable results. This distinction highlights that the propagation of safety information is a decisive factor in offline safe RL.

Robustness via Asymmetric Losses. Figure 2 compares the standard smooth (deterministic) barrier update (Eq. (9)) with its robust variant (Eq. (11)). To test robustness, we modify the DSRL dataset by mislabeling 70% of high-cost trajectories as safe using the DSRL’s built-in functionality. We evaluate both deterministic and robust versions of Smooth CBF across velocity control and MetaDrive tasks. The robust version utilizes an asymmetric loss (percentile) with $\tau = 0.6$, while the standard version uses the standard L1 loss. Robust-Smooth consistently reduces normalized cost compared to the Deterministic-Smooth version. Particularly in velocity tasks, Robust-Smooth achieves nearly zero cost whereas the Deterministic-Smooth version incurs high costs. This im-

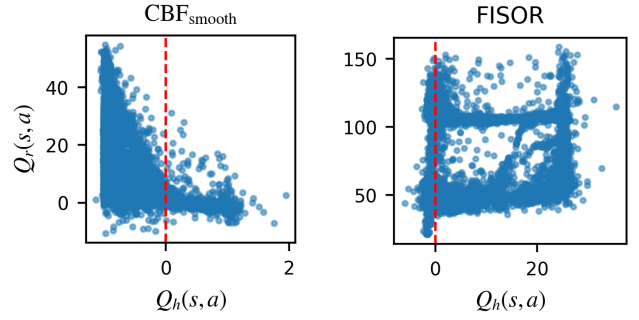


Figure 3: Reward critic $Q_r(s, a)$ vs. safety critic $Q_h(s, a)$ distributions for Swimmer Velocity. The red dashed line at $Q_h(s, a) = 0$ marks the safety boundary: points to the left are safe, and those to the right are unsafe. **Left:** $\text{CBF}_{\text{smooth}}$ prevents reward leakage as $Q_r(s, a)$ is low when $Q_h(s, a) > 0$ while **Right:** FISOR shows reward leakage as $Q_r(s, a)$ is high even for unsafe pairs.

provement is achieved with only a modest reduction in normalized reward, demonstrating that asymmetric losses effectively estimate a persistent safety set under uncertainty. Importantly, these trends are consistent across all tasks, suggesting that improvements arise from stochasticity-aware safety estimation rather than task-specific tuning. Although, the robust variant uses asymmetric losses to approximate worst-case transitions, reducing costs on average in stochastic environments, but for *hardenselhardspare*, limited data coverage causes $\mathbb{E}^\tau$ to underestimate extreme violations, leading to less conservative behavior and slightly higher costs than the deterministic formulation. Complete results across all velocity control, MetaDrive environments are in Appendix B.3.

Reward Leakage Analysis. To empirically evaluate reward leakage, we save the fully trained model and randomly sample 10K state-action pairs from the offline dataset. For each pair, we compute the safety critic $Q_h(s, a)$ and reward critic $Q_r(s, a)$ using the saved model, and then plot $Q_r(s, a)$ versus $Q_h(s, a)$ to verify that, for $Q_h(s, a) > 0$, the reward critic remains low, indicating that reward does not propagate through unsafe transitions. Figure 3 confirms that $\text{CBF}_{\text{smooth}}$ approach effectively blocks reward propagation through unsafe regions, ensuring that high rewards are only assigned to safe state-action pairs and providing strong safety guarantees during value learning.

7 Conclusions

We propose an offline safe RL framework that learns discrete CBFs via a generalized Bellman backup, unifying prior CBF methods with convergence guarantees. We further introduce a robust asymmetric variant for worst-case transitions and a safety-aware Bellman target to prevent reward leakage, together with a value-based offline RL method for safe policy extraction. Across continuous control and autonomous driving benchmarks, our method satisfies constraints while achieving competitive or superior reward.

Ethical Statement

There are no ethical issues.

Contribution Statement

Ayan Choudhury and Janaka Brahmanage contributed equally to this work.

References

- [Achiam *et al.*, 2017] Joshua Achiam, David Held, Aviv Tamar, and Pieter Abbeel. Constrained policy optimization. In *Proceedings of the 34th International Conference on Machine Learning - Volume 70*, page 22–31, 2017.
- [Agrawal and Sreenath, 2017] Ayush Agrawal and Koushil Sreenath. Discrete Control Barrier Functions for Safety-Critical Control of Discrete Systems with Application to Bipedal Robot Navigation. In *Robotics: Science and Systems XIII*. Robotics: Science and Systems Foundation, 2017.
- [Altman, 2021] Eitan Altman. *Constrained Markov Decision Processes: Stochastic Modeling*. Routledge, Boca Raton, 2021.
- [Ames *et al.*, 2014] Aaron D. Ames, Jessy W. Grizzle, and Paulo Tabuada. Control barrier function based quadratic programs with application to adaptive cruise control. In *53rd IEEE Conference on Decision and Control*, pages 6271–6278, 2014.
- [Ames *et al.*, 2019] Aaron D. Ames, Samuel Coogan, Magnus Egerstedt, Gennaro Notomista, Koushil Sreenath, and Paulo Tabuada. Control Barrier Functions: Theory and Applications. In *European Control Conference*, pages 3420–3431, 2019.
- [Bansal *et al.*, 2017] Somil Bansal, Mo Chen, Sylvia Herbert, and Claire J. Tomlin. Hamilton-Jacobi Reachability: A Brief Overview and Recent Advances, September 2017. arXiv:1709.07523 [cs, math].
- [Brahmanage and Kumar, 2026] Janaka Brahmanage and Akshat Kumar. Beyond hard constraints: Budget-conditioned reachability for safe offline reinforcement learning. In *International Conference on Automated Planning and Scheduling (ICAPS)*, 2026.
- [Chemingui *et al.*, 2025a] Yassine Chemingui, Aryan Deshwal, Alan Fern, Thanh Nguyen-Tang, and Jana Doppa. Online optimization for offline safe reinforcement learning. In *The Thirty-ninth Annual Conference on Neural Information Processing Systems*, 2025.
- [Chemingui *et al.*, 2025b] Yassine Chemingui, Aryan Deshwal, Honghao Wei, Alan Fern, and Janardhan Rao Doppa. Constraint-Adaptive Policy Switching for Offline Safe Reinforcement Learning. In *Proceedings of the Thirty-Ninth AAAI Conference on Artificial Intelligence and Thirty-Seventh Conference on Innovative Applications of Artificial Intelligence and Fifteenth Symposium on Educational Advances in Artificial Intelligence*, pages 15722–15730, 2025.
- [Fang *et al.*, 2024] Nan Fang, Guiliang Liu, and Wei Gong. Offline Inverse Constrained Reinforcement Learning for Safe-Critical Decision Making in Healthcare, October 2024. arXiv:2410.07525 [cs].
- [Fujimoto *et al.*, 2019] Scott Fujimoto, David Meger, and Doina Precup. Off-Policy Deep Reinforcement Learning without Exploration, August 2019. arXiv:1812.02900 [cs, stat].
- [Ganai *et al.*, 2024] Milan Ganai, Sicun Gao, and Sylvia L. Herbert. Hamilton-Jacobi Reachability in Reinforcement Learning: A Survey. *IEEE Open Journal of Control Systems*, 3:310–324, 2024.
- [Garg *et al.*, 2023] Divyansh Garg, Joey Hejna, Matthieu Geist, and Stefano Ermon. Extreme Q-Learning: Max-Ent RL without Entropy. In *The Eleventh International Conference on Learning Representations*, February 2023.
- [Gu *et al.*, 2023] Shangding Gu, Jakub Grudzien Kuba, Yuanpei Chen, Yali Du, Long Yang, Alois Knoll, and Yaodong Yang. Safe multi-agent reinforcement learning for multi-robot control. *Artificial Intelligence*, 319:103905, 2023.
- [Gu *et al.*, 2024] Shangding Gu, Long Yang, Yali Du, Guang Chen, Florian Walter, Jun Wang, and Alois Knoll. A Review of Safe Reinforcement Learning: Methods, Theory and Applications. In *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2024.
- [Gulino *et al.*, 2023] Cole Gulino, Justin Fu, Wenjie Luo, George Tucker, Eli Bronstein, Yiren Lu, Jean Harb, Xinlei Pan, Yan Wang, Xiangyu Chen, John D Co-Reyes, Rishabh Agarwal, Rebecca Roelofs, Yao Lu, Nico Montali, Paul Mouglin, Zoey Yang, Brandyn White, Aleksandra Faust, Rowan McAllister, Dragomir Anguelov, and Benjamin Sapp. Waymax: An Accelerated, Data-Driven Simulator for Large-Scale Autonomous Driving Research. In *Proceedings of Thirty-seventh Conference on Neural Information Processing Systems*, 2023.
- [Guo *et al.*, 2025a] Junyu Guo, Zhi Zheng, Donghao Ying, Ming Jin, Shangding Gu, Costas Spanos, and Javad Lavaei. Don’t trade off safety: Diffusion regularization for constrained offline RL. In *The Thirty-ninth Annual Conference on Neural Information Processing Systems*, 2025.
- [Guo *et al.*, 2025b] Zijian Guo, Weichao Zhou, Shengao Wang, and Wenchao Li. Constraint-Conditioned Actor-Critic for Offline Safe Reinforcement Learning. In *The Thirteenth International Conference on Learning Representations*, 2025.
- [Hansen-Estruch *et al.*, 2023] Philippe Hansen-Estruch, Ilya Kostrikov, Michael Janner, Jakub Grudzien Kuba, and Sergey Levine. IDQL: Implicit Q-Learning as an Actor-Critic Method with Diffusion Policies. *Artificial Intelligence (cs.AI), FOS: Computer and information sciences, Machine Learning (cs.LG)*, 2023. arxiv:2304.10573.
- [Harms *et al.*, 2024] Marvin Harms, Mihir Kulkarni, Nikhil Khedekar, Martin Jacquet, and Kostas Alexis. Neural control barrier functions for safe navigation. In *IEEE/RSJ In-*

- ternational Conference on Intelligent Robots and Systems*, pages 10415–10422, 2024.
- [Koirala *et al.*, 2025] Prajwal Koirala, Zhanhong Jiang, Soumik Sarkar, and Cody Fleming. Latent Safety-Constrained Policy Approach for Safe Offline Reinforcement Learning. In *The Thirteenth International Conference on Learning Representations*, 2025.
- [Kostrikov *et al.*, 2022] Ilya Kostrikov, Ashvin Nair, and Sergey Levine. Offline Reinforcement Learning with Implicit Q-Learning. In *The Tenth International Conference on Learning Representations*, 2022.
- [Lindemann *et al.*, 2024] Lars Lindemann, Alexander Robey, Lejun Jiang, Satyajeet Das, Stephen Tu, and Nikolai Matni. Learning robust output control barrier functions from safe expert demonstrations. *IEEE Open Journal of Control Systems*, 3:158–172, 2024.
- [Littman and Szepesvári, 1996] Michael L Littman and Csaba Szepesvári. A generalized reinforcement-learning model: Convergence and applications. In *Proceedings of the Thirteenth International Conference on International Conference on Machine Learning*, pages 310–318, 1996.
- [Mitchell *et al.*, 2005] I.M. Mitchell, A.M. Bayen, and C.J. Tomlin. A time-dependent Hamilton-Jacobi formulation of reachable sets for continuous dynamic games. *IEEE Transactions on Automatic Control*, 50(7):947–957, 2005.
- [Mnih *et al.*, 2013] Volodymyr Mnih, Koray Kavukcuoglu, David Silver, Alex Graves, Ioannis Antonoglou, Daan Wierstra, and Martin Riedmiller. Playing Atari with Deep Reinforcement Learning, 2013.
- [Osinski *et al.*, 2019] Blazej Osinski, Adam Jakubowski, Piotr Milos, Pawel Ziecina, Christopher Galias, Silviu Homocceanu, and Henryk Michalewski. Simulation-based reinforcement learning for real-world autonomous driving. *CoRR*, abs/1911.12905, 2019.
- [So *et al.*, 2024] Oswin So, Zachary Serlin, Makai Mann, Jake Gonzales, Kwesi Rutledge, Nicholas Roy, and Chuchu Fan. How to train your neural control barrier function: Learning safety filters for complex input-constrained systems. In *IEEE International Conference on Robotics and Automation*, pages 11532–11539, 2024.
- [Su *et al.*, 2025] Huikang Su, Dengyun Peng, Zifeng Zhuang, YuHan Liu, Qiguang Chen, Donglin Wang, and Qinghe Liu. Boundary-to-region supervision for offline safe reinforcement learning. In *The Thirty-ninth Annual Conference on Neural Information Processing Systems*, 2025.
- [Sutton and Barto, 2020] Richard S. Sutton and Andrew Barto. *Reinforcement learning: an introduction*. Adaptive computation and machine learning. The MIT Press, Cambridge, Massachusetts London, England, second edition, 2020.
- [Tabbara and Sibai, 2025] Ihab Tabbara and Hussein Sibai. Learning conservative neural control barrier functions from offline data, 2025. arXiv:2505.00908 [cs].
- [Tan *et al.*, 2023] Daniel C H Tan, Fernando Acero, Robert McCarthy, Dimitrios Kanoulas, and Zhibin Li. Your Value Function is a Control Barrier Function. *International Conference on Machine Learning: 2nd Workshop on Formal Verification of Machine Learning (WVVML)*, 2023.
- [Tessler *et al.*, 2018] Chen Tessler, Daniel J. Mankowitz, and Shie Mannor. Reward Constrained Policy Optimization. In *Proceedings of The Seventh International Conference on Learning Representations*, 2018.
- [Wachi *et al.*, 2024] Akifumi Wachi, Xun Shen, and Yanan Sui. A Survey of Constraint Formulations in Safe Reinforcement Learning, May 2024. arXiv:2402.02025 [cs].
- [Yu *et al.*, 2022] Dongjie Yu, Haitong Ma, Shengbo Eben Li, and Jianyu Chen. Reachability Constrained Reinforcement Learning. In *Proceedings of the 39th International Conference on Machine Learning*, 2022.
- [Yu *et al.*, 2025] Hongzhan Yu, Seth Farrell, Ryo Yoshimitsu, Zhizhen Qin, Henrik I. Christensen, and Sicun Gao. Estimating control barriers from offline data, 2025. arxiv:2503.10641.
- [Zheng *et al.*, 2024] Yinan Zheng, Jianxiong Li, Dongjie Yu, Yujie Yang, Shengbo Eben Li, Xianyu Zhan, and Jingjing Liu. Safe Offline Reinforcement Learning with Feasibility-Guided Diffusion Model. In *The Twelfth International Conference on Learning Representations*, 2024.

A Theoretical Results

A.1 Convergence and CBF property

Theorem 3.1. Consider a CMDP with continuous safety indicator $h(s, a)$. The iterative application of the operator $\mathcal{B}$ defined in (7) converges to a unique fixed point Q_h^* . Furthermore, Q_h^* satisfies the Robust Discrete CBF condition defined in (6), provided Assumptions 3.1 and 3.2 hold.

Proof. The proof consists of two parts: proving convergence via contraction mapping and verifying the CBF property of the fixed point.

Part 1: Convergence. We show that $\mathcal{B}$ is a contraction mapping under the sup-norm $\|\cdot\|_\infty$. Let Q_1, Q_2 be two arbitrary functions.

$$\begin{aligned} \|\mathcal{B}Q_1 - \mathcal{B}Q_2\|_\infty &= \sup_{s,a} |\Psi(h, \max_{s'} \min_{a'} Q_1') - \Psi(h, \max_{s'} \min_{a'} Q_2')| \\ &\leq \gamma \sup_{s,a} |\max_{s'} \min_{a'} Q_1' - \max_{s'} \min_{a'} Q_2'| \\ &\leq \gamma \|Q_1 - Q_2\|_\infty \end{aligned}$$

The first inequality follows directly from Assumption 3.1 (contraction of Ψ). The second inequality holds due to the non-expansiveness of the max and min operators [Littman and Szepesvári, 1996]. Specifically, we use the property $|\max_x f(x) - \max_x g(x)| \leq \sup_x |f(x) - g(x)|$ sequentially:

$$\begin{aligned} &\left| \max_{s'} \min_{a'} Q_1(s', a') - \max_{s'} \min_{a'} Q_2(s', a') \right| \\ &\leq \sup_{s'} \left| \min_{a'} Q_1(s', a') - \min_{a'} Q_2(s', a') \right| \\ &\leq \sup_{s'} \left| \max_{a'} (-Q_1(s', a')) - \max_{a'} (-Q_2(s', a')) \right| \\ &\leq \sup_{s', a'} |Q_1(s', a') - Q_2(s', a')| = \|Q_1 - Q_2\|_\infty. \end{aligned}$$

Since $\gamma < 1$, $\mathcal{B}$ is a contraction. By the Banach Fixed Point Theorem, it converges to a unique fixed point Q_h^* .

Part 2: CBF Property. The fixed point satisfies $Q_h^*(s, a) = \Psi(h(s, a), \max_{s'} \min_{a'} Q_h^*(s', a'))$. By Assumption 3.2, this implies:

$$\max_{s' \in \text{supp}(\mathcal{T}(\cdot|s, a))} \min_{a' \in \mathcal{A}} Q_h^*(s', a') \leq Q_h^*(s, a) - \alpha(Q_h^*(s, a)) \quad (15)$$

Let $\pi(s) = \arg \min_{a \in \mathcal{A}} Q_h^*(s, a)$, be the optimal safety policy. Then:

$$\max_{s' \in \text{supp}(\mathcal{T}(\cdot|s, a))} Q_h^*(s', \pi(s')) \leq Q_h^*(s, a) - \alpha(Q_h^*(s, a)) \quad (16)$$

This recovers the exact condition required for a Robust Discrete CBF in (6). $\square$

A.2 Maximality of the Forward-Invariant Safe Set

Lemma 3.2. A set $\mathcal{C} \subseteq \mathcal{S} \times \mathcal{A}$ is forward invariant if there exists a policy π such that, for all $(s, a) \in \mathcal{C}$ and all $s' \in \text{supp}(\mathcal{T}(\cdot|s, a))$, we have $(s', \pi(s')) \in \mathcal{C}$. A forward-invariant safe set $\mathcal{C}^*$ is maximal if, for any other forward-invariant safe set $\tilde{\mathcal{C}}$, it holds that $\tilde{\mathcal{C}} \subseteq \mathcal{C}^*$. Let Q_h^* be the fixed point of the operator $\mathcal{B}$ defined in (7). Then the set

$$\mathcal{C}^* := \{(s, a) \in \mathcal{S} \times \mathcal{A} \mid Q_h^*(s, a) \leq 0\}$$

is the maximal forward-invariant safe set induced by $\mathcal{B}$.

Proof. We equip the space of bounded functions $Q : \mathcal{S} \times \mathcal{A} \rightarrow \mathbb{R}$ with the pointwise partial order

$$Q_1 \preceq Q_2 \iff Q_1(s, a) \leq Q_2(s, a), \forall (s, a).$$

We first show that the Bellman operator $\mathcal{B}$ is order-preserving. Let $Q_1 \preceq Q_2$. By monotonicity of the min and max operators,

$$\max_{s'} \min_{a'} Q_1(s', a') \leq \max_{s'} \min_{a'} Q_2(s', a').$$

By Assumption 3.1, the aggregator $\Psi(h, \cdot)$ is strictly increasing in its second argument. Hence, for any fixed $h(s, a)$,

$$\Psi\left(h(s, a), \max_{s'} \min_{a'} Q_1(s', a')\right) \leq \Psi\left(h(s, a), \max_{s'} \min_{a'} Q_2(s', a')\right). \quad (17)$$

which implies

$$\mathcal{B}Q_1(s, a) \leq \mathcal{B}Q_2(s, a), \quad \forall (s, a).$$

Thus, $\mathcal{B}$ is monotone with respect to $\preceq$. Assumption 3.2 implies that for any $q = \Psi(h, v)$, there exists an extended class $\mathcal{K}$ function α such that

$$v \leq q - \alpha(q).$$

Since $\Psi(h, \cdot)$ is strictly increasing and continuous (Assumption 3.1), it admits a well-defined inverse with respect to its second argument, which we denote by $k^{-1}(\cdot)$. Strict monotonicity of $\Psi(h, \cdot)$ directly implies that k^{-1} is also strictly increasing. This inverse mapping is used implicitly when interpreting the Bellman update as propagating future safety violations backward to the current state-action pair, and its monotonicity ensures that order relations are preserved under Bellman iteration.

Let $Q_0 = h$ and define $Q_{k+1} = \mathcal{B}Q_k$. By monotonicity of $\mathcal{B}$, the sequence $\{Q_k\}$ is non-decreasing, and by Theorem 3.1, it converges to the unique fixed point Q_h^* .

Each iterate Q_k corresponds to the worst-case safety violation over trajectories of length at most k . Therefore, the induced sets

$$\mathcal{C}_k := \{(s, a) \mid Q_k(s, a) \leq 0\}$$

form a nested sequence of forward-invariant sets, whose limit is $\mathcal{C}^*$.

Let $\tilde{\mathcal{C}}$ be any other forward-invariant safe set induced by some function $\tilde{Q}$. Forward invariance implies that $\tilde{Q}$ satisfies the robust discrete CBF condition, which in turn implies $\mathcal{B}\tilde{Q} \preceq \tilde{Q}$. By monotonicity of $\mathcal{B}$,

$$Q_h^* = \lim_{k \rightarrow \infty} \mathcal{B}^k \tilde{Q} \preceq \tilde{Q}.$$

Hence, $\tilde{\mathcal{C}} \subseteq \mathcal{C}^*$, and $\mathcal{C}^*$ is the maximal forward-invariant safe set. □

A.3 Minimality of the Barrier Value Function

Beyond establishing the maximal safe set, we further characterize Q_h^* in terms of its minimality and consistency with the offline dataset. Specifically, we show that Q_h^* is the least restrictive barrier function satisfying the safety conditions, ensuring that high-reward actions are not artificially suppressed during learning. Here, Q_h^* is shown to be *minimal* among all functions satisfying the discrete CBF condition, ensuring that safety constraints do not artificially suppress high-reward actions during value learning and policy optimization.

Lemma A.1. Let $\tilde{Q}$ be any function satisfying the Robust Discrete CBF condition (6). Then

$$Q_h^*(s, a) \leq \tilde{Q}(s, a), \quad \forall (s, a).$$

We emphasize that the persistent safety set $\mathcal{C}^* = \{(s, a) \mid Q_h^*(s, a) \leq 0\}$ is not restricted to a single connected region. Since Q_h^* is learned pointwise via Bellman iteration over dataset-supported transitions, the induced safe set may consist of multiple disconnected components. Each connected component corresponds to a region from which safety can be maintained indefinitely under some policy supported by the data. No assumptions on convexity or connectedness of the safe set are required.

Proof. Since $\tilde{Q}$ satisfies the Robust Discrete CBF condition, there exists a policy $\tilde{\pi}$ such that

$$\max_{s'} \tilde{Q}(s', \tilde{\pi}(s')) \leq \tilde{Q}(s, a) - \alpha(\tilde{Q}(s, a)).$$

By Assumption 3.2, this implies

$$\max_{s'} \tilde{Q}(s', \tilde{\pi}(s')) \leq v \quad \Rightarrow \quad \Psi(h(s, a), v) \leq \tilde{Q}(s, a).$$

Since $\min_{a'} \tilde{Q}(s', a') \leq \tilde{Q}(s', \tilde{\pi}(s'))$, we obtain

$$\mathcal{B}\tilde{Q}(s, a) \leq \tilde{Q}(s, a),$$

so $\tilde{Q}$ is a post-fixed point of $\mathcal{B}$.

By contraction and monotonicity of $\mathcal{B}$, iterating from $\tilde{Q}$ yields

$$Q_h^* = \lim_{k \rightarrow \infty} \mathcal{B}^k \tilde{Q} \preceq \tilde{Q}.$$

Thus, Q_h^* is minimal among all valid discrete CBFs. □

A.4 Dataset Consistency

In the offline setting, a transition (s, a, s') is said to be dataset-supported if it appears in the offline dataset $\mathcal{D}$. A trajectory is dataset-supported if all its transitions are dataset-supported. Assume Bellman updates are computed only over transitions supported by an offline dataset $\mathcal{D}$. Then Q_h^* depends solely on dataset-supported transitions and does not rely on unseen dynamics.

Proof. Let $\mathcal{D} = \{(s, a, s')\}$ denote the offline dataset and restrict the operator $\mathcal{B}$ to

$$\max_{s' \in \text{supp}(\mathcal{T}(\cdot|s,a)) \cap \mathcal{D}} \min_{a'} Q(s', a').$$

Define $Q_{k+1} = \mathcal{B}Q_k$ with $Q_0 = h$. Each iterate satisfies

$$Q_k(s, a) = \max_{0 \leq t \leq k} h(s_t, a_t),$$

over trajectories entirely supported by $\mathcal{D}$.

Taking the limit yields

$$Q_h^*(s, a) = \sup_{\tau \in \mathcal{T}_{\mathcal{D}}(s, a)} \max_{t \geq 0} h(s_t, a_t),$$

where $\mathcal{T}_{\mathcal{D}}(s, a)$ denotes dataset-supported trajectories.

Thus, Q_h^* encodes unavoidable safety violations provable from the data alone, ensuring dataset consistency. $\square$

A.5 Proof of Proposition 4.1

Proof. We verify Assumptions 3.1 and 3.2 for each variant.

Maximum $\Psi(h, v) = \max(h, \gamma v)$.

- **Contraction:** Since the max operator is non-expansive, we have:

$$|\max(h, \gamma v_1) - \max(h, \gamma v_2)| \leq |\gamma v_1 - \gamma v_2| = \gamma |v_1 - v_2|.$$

Thus, Ψ is a contraction in v with factor $\gamma < 1$.

- **CBF Condition:** Let $q = \max(h, \gamma v)$. By definition, $q \geq \gamma v$, which implies $v \leq q/\gamma$. To satisfy the discrete CBF condition $v \leq q - \alpha(q)$, it suffices to show that:

$$\frac{q}{\gamma} \leq q - \alpha(q) \iff \alpha(q) \leq q \left(1 - \frac{1}{\gamma}\right).$$

Consider the safe set where $q \leq 0$ and the discount factor $\gamma \in (0, 1)$:

1. The term $(1 - 1/\gamma)$ is negative.
2. Since q is negative (safe), the term $q(1 - 1/\gamma)$ is **positive**.
3. By definition, for any extended class $\mathcal{K}$ function, $\alpha(q) \leq 0$ when $q \leq 0$.

Therefore, the condition requires a negative value ($\alpha(q)$) to be less than or equal to a positive value ($q(1 - 1/\gamma)$), which holds for **any** valid extended class $\mathcal{K}$ function α .

Smooth $\Psi(h, v) = (1 - \gamma)h + \gamma \max(h, v)$.

- **Contraction:** The operator is a contraction with modulus γ . For any v_1, v_2 :

$$\begin{aligned} |\Psi(h, v_1) - \Psi(h, v_2)| &= |((1 - \gamma)h + \gamma \max(h, v_1)) - ((1 - \gamma)h + \gamma \max(h, v_2))| \\ &= \gamma |\max(h, v_1) - \max(h, v_2)| \\ &\leq \gamma |v_1 - v_2|. \end{aligned}$$

Thus, Ψ is a contraction in v with factor $\gamma < 1$.

- **CBF Condition:** Let $q = \Psi(h, v)$. We analyze the discrete CBF condition $v \leq q - \alpha(q)$ for safe states ($q \leq 0$):

1. **Case $v \leq h$:** In this case, $\max(h, v) = h$, so $q = (1 - \gamma)h + \gamma h = h$. The condition becomes $v \leq h - \alpha(h)$. Since $v \leq h$ and $-\alpha(h) \geq 0$ (because α is an extended class $\mathcal{K}$ function and h is safe), the inequality $v \leq h + \text{positive}$ holds trivially for **any** valid α .

2. **Case $v > h$:** In this case, $\max(h, v) = v$, so $q = (1 - \gamma)h + \gamma v$. Solving for v , we get $v = \frac{q - (1 - \gamma)h}{\gamma}$. Substituting this into the CBF condition $v \leq q - \alpha(q)$:

$$\frac{q - (1 - \gamma)h}{\gamma} \leq q - \alpha(q) \iff \alpha(q) \leq q - \frac{q - (1 - \gamma)h}{\gamma}.$$

Simplifying the RHS:

$$\alpha(q) \leq \frac{\gamma q - q + (1 - \gamma)h}{\gamma} = \frac{(1 - \gamma)(h - q)}{\gamma}.$$

Since $v > h$, we have $q > h$, which implies $(h - q)$ is negative. Thus, the condition holds if $\alpha(q)$ is sufficiently negative (strong decay):

$$\alpha(q) \leq -\frac{1 - \gamma}{\gamma} |h - q|.$$

Additive $\Psi(h, v) = h + \gamma v$.

- **Contraction:** The operator is a contraction with modulus γ :

$$|\Psi(h, v_1) - \Psi(h, v_2)| = |(h + \gamma v_1) - (h + \gamma v_2)| = \gamma |v_1 - v_2|.$$

- **CBF Condition:** Let $q = \Psi(h, v) = h + \gamma v$. This form mirrors the Bellman update for a value function. Solving for v , we have $v = \frac{q - h}{\gamma}$. Directly satisfying the standard condition $v \leq q - \alpha(q)$ for $q \leq 0$ is generally not possible without invalid dependence on h .

However, as established in [Tan *et al.*, 2023], this additive form satisfies the CBF property if we define the barrier function as the value shifted by a safety threshold R . Let $\mathcal{H}(q) = q - R$. The condition becomes:

$$v - R \leq (q - R) - \alpha(q - R).$$

Substituting $v = \frac{q - h}{\gamma}$:

$$\frac{q - h}{\gamma} - R \leq (q - R) - \alpha(q - R).$$

Rearranging guarantees that a valid class $\mathcal{K}$ function α exists provided that R is chosen within specific bounds (sufficiently large to separate safe and unsafe states). Thus, the Additive mixer generates a valid CBF if the zero-level set is learned as a threshold R , rather than being fixed at 0.

□

B Additional Experiments and Ablations

B.1 Comparison of Barrier-Value Formulations

Table 2 isolates the effect of the barrier-value formulation by comparing $\text{CBF}_{\text{smooth}}$, $\text{CBF}_{\text{additive}}$, and $\text{CBF}_{\text{maximum}}$ functions on MetaDrive environments. $\text{CBF}_{\text{maximum}}$ achieves best reward in Metadrive medium settings but has higher cost deviation than $\text{CBF}_{\text{smooth}}$ function which is used for the main results. In contrast, the $\text{CBF}_{\text{additive}}$ barrier formulation consistently achieves lower cost than $\text{CBF}_{\text{smooth}}$ while attaining higher reward in some cases. The empirical trend of the average values aligns with our theoretical analysis that the learned barrier value function is minimal among all valid discrete CBFs, avoiding unnecessary conservatism while preserving forward invariance. Overall, these results demonstrate that how safety information is propagated through the Bellman operator critically affects the safety-performance trade-off in offline RL.

B.2 Sensitivity to Reward and Cost Expectile

We next evaluate the sensitivity of the choice of reward and cost thresholds in the asymmetric Bellman backup. Table 3 varies the reward expectile from 0.7 to 0.8 at a fixed cost expectile of 0.15, while Table 4 reports results across a range of cost expectile from 0.1 to 0.2 at a fixed reward expectile of 0.75. The cost and reward expectile here refers to the values of τ_r and τ_h used to determine the expectile loss during training, while the cost and reward values in table are the results from the evaluation corresponding to those trained models. Across tasks, the pair of reward and cost expectile at 0.75 and 0.15 respectively, consistently maintains zero or near-zero cost while achieving stable reward. Moreover, the expectile values at lower or higher than the optimum exhibit better returns with safe costs, indicating stability across hyperparameter tuning. This suggests that safety is enforced structurally by the learned barrier value function rather than being delicately balanced by penalty coefficients.

	CBF _{smooth}		CBF _{additive}		CBF _{maximum}	
	reward ↑	cost ↓	reward ↑	cost ↓	reward ↑	cost ↓
easysparse	0.52 ±0.01	0.46 ±0.15	0.53 ±0.04	0.41 ±0.22	0.30 ±0.23	0.03 ±0.02
easymean	0.38 ±0.07	0.36 ±0.37	0.42 ±0.12	1.82 ±2.20	0.13 ±0.20	0.05 ±0.05
easydense	0.55 ±0.21	0.61 ±0.18	0.31 ±0.12	0.16 ±0.23	0.47 ±0.11	0.29 ±0.16
mediumsparse	0.54 ±0.04	0.23 ±0.46	0.67 ±0.08	0.04 ±0.03	0.74 ±0.06	0.74 ±0.57
mediummean	0.54 ±0.15	0.30 ±0.36	0.31 ±0.33	0.02 ±0.02	0.61 ±0.03	0.45 ±0.60
mediumdense	0.53 ±0.16	0.26 ±0.15	0.55 ±0.09	0.06 ±0.05	0.73 ±0.09	0.32 ±0.41
hardsparse	0.42 ±0.02	0.77 ±0.07	0.34 ±0.04	0.05 ±0.07	0.16 ±0.22	0.15 ±0.25
hardmean	0.36 ±0.06	0.28 ±0.16	0.29 ±0.06	0.03 ±0.03	0.37 ±0.03	0.10 ±0.06
harddense	0.36 ±0.09	0.87 ±0.42	0.20 ±0.19	0.04 ±0.04	0.42 ±0.04	1.08 ±0.91
Average (9)	0.47 ±0.08	0.46 ±0.24	0.40 ±0.15	0.29 ±0.59	0.44 ±0.22	0.36 ±0.35

Table 2: Results for comparison of normalized cost and normalized reward using Maximum and Additive barrier value functions with the Smooth barrier value function on Metadrive where ↑ indicates that a higher value is better, while ↓ indicates that a lower value is preferable. **Bold** text denotes safe policies, **Red** indicates unsafe policies, and **Blue** highlights the highest reward among the safe policies for each task. Each result is obtained as an average by running three random seeds and evaluating over 20 episodes.

	0.700		0.725		0.750		0.775		0.800	
	reward ↑	cost ↓	reward ↑	cost ↓	reward ↑	cost ↓	reward ↑	cost ↓	reward ↑	cost ↓
AntVelocity	0.87 ±0.01	0.00 ±0.00	0.87 ±0.01	0.00 ±0.00	0.85 ±0.02	0.00 ±0.00	0.88 ±0.03	0.00 ±0.00	0.85 ±0.02	0.00 ±0.00
HalfCheetahVelocity	0.82 ±0.05	0.00 ±0.00	0.80 ±0.01	0.00 ±0.00	0.80 ±0.02	0.00 ±0.00	0.82 ±0.02	0.00 ±0.00	0.80 ±0.03	0.00 ±0.00
SwimmerVelocity	0.53 ±0.08	0.00 ±0.00	0.50 ±0.05	0.00 ±0.00	0.48 ±0.08	0.00 ±0.00	0.47 ±0.12	0.00 ±0.00	0.50 ±0.06	0.07 ±0.12
Average (3)	0.74 ±0.18	0.00 ±0.00	0.72 ±0.20	0.00 ±0.00	0.71 ±0.20	0.00 ±0.00	0.72 ±0.22	0.00 ±0.00	0.72 ±0.19	0.02 ±0.04
AntCircle	0.50 ±0.03	0.00 ±0.00	0.51 ±0.05	0.00 ±0.00	0.50 ±0.08	0.20 ±0.35	0.49 ±0.04	0.00 ±0.00	0.44 ±0.04	0.87 ±1.50
BallCircle	0.45 ±0.02	0.00 ±0.00	0.43 ±0.03	0.00 ±0.00	0.43 ±0.02	0.00 ±0.00	0.41 ±0.01	0.00 ±0.00	0.40 ±0.03	0.00 ±0.00
CarCircle	0.56 ±0.06	0.00 ±0.00	0.55 ±0.03	0.00 ±0.00	0.51 ±0.03	0.00 ±0.00	0.53 ±0.05	0.00 ±0.00	0.55 ±0.05	1.13 ±1.96
DroneCircle	0.31 ±0.05	0.00 ±0.00	0.34 ±0.07	0.00 ±0.00	0.30 ±0.08	0.00 ±0.00	0.31 ±0.07	0.00 ±0.00	0.26 ±0.09	0.07 ±0.12
Average (4)	0.46 ±0.11	0.08 ±0.16	0.46 ±0.09	0.08 ±0.17	0.43 ±0.09	0.13 ±0.15	0.43 ±0.10	0.08 ±0.16	0.41 ±0.12	0.57 ±0.52
easysparse	0.14 ±0.30	0.03 ±0.05	0.41 ±0.17	0.09 ±0.02	0.32 ±0.11	0.04 ±0.02	0.30 ±0.28	0.08 ±0.11	0.08 ±0.19	0.01 ±0.02
mediumsparse	0.70 ±0.04	0.75 ±0.62	0.73 ±0.09	0.73 ±0.63	0.76 ±0.07	0.48 ±0.32	0.76 ±0.07	0.47 ±0.41	0.63 ±0.10	0.06 ±0.06
hardsparse	0.37 ±0.04	0.44 ±0.26	0.39 ±0.02	0.38 ±0.33	0.39 ±0.03	0.76 ±0.57	0.38 ±0.04	0.45 ±0.39	0.24 ±0.24	0.08 ±0.11
Average (3)	0.40 ±0.28	0.41 ±0.36	0.51 ±0.19	0.40 ±0.32	0.49 ±0.24	0.43 ±0.36	0.48 ±0.25	0.33 ±0.22	0.32 ±0.28	0.05 ±0.04

Table 3: Results for comparison of normalized cost and normalized reward across range of reward-tau at a fixed cost-tau where ↑ indicates that a higher value is better, while ↓ indicates that a lower value is preferable. **Bold** text denotes safe policies, **Red** indicates unsafe policies, and **Blue** highlights the highest reward among the safe policies for each task. Each result is obtained as an average by running three seeds of (0,10,20) and evaluating over 20 episodes.

B.3 Complete Results on Effectiveness of The Asymmetric Aggregation

Figure 4 reports the full set of results for the experiment summarized in Figure 2 of the main paper, covering all velocity control tasks (Ant, HalfCheetah, Swimmer, Hopper) and the complete MetaDrive suite (*easydense*, *easymean*, *easysparse*, *mediumdense*, *mediummean*, *mediumsparse*, *harddense*, *hardmean*, *hardsparse*). Using the same protocol—70% of high-cost trajectories mislabeled as safe via DSRL—we compare the deterministic Smooth CBF against its robust variant with asymmetric loss ($\tau = 0.6$). Across all tasks, Robust-Smooth consistently lowers normalized cost while retaining comparable normalized reward, confirming the trend observed in the main paper. The exceptions are the *harddense* and *hardsparse* settings, where limited data coverage of high-cost transitions causes the expected $\mathbb{E}^T$ to underestimate extreme violations, yielding slightly higher cost than the deterministic baseline. Overall, the broader evaluation supports the conclusion that asymmetric aggregation provides a tighter approximation of worst-case future safety violations under uncertainty.

C Additional Experimental Details

C.1 Evaluation Metrics

We evaluate performance using the normalized reward return and the normalized cost return. The evaluation metrics are defined as

$$R_{\text{normalized}} = \frac{R_{\pi} - r_{\min}}{r_{\max} - r_{\min}}, \quad C_{\text{normalized}} = \frac{C_{\pi} + \epsilon}{l + \epsilon},$$

	0.100		0.125		0.150		0.175		0.200	
	reward $\uparrow$	cost $\downarrow$	reward $\uparrow$	cost $\downarrow$	reward $\uparrow$	cost $\downarrow$	reward $\uparrow$	cost $\downarrow$	reward $\uparrow$	cost $\downarrow$
AntVelocity	0.86 ± 0.02	0.00 ± 0.00	0.88 ± 0.03	0.00 ± 0.00	0.87 ± 0.01	0.00 ± 0.00	0.88 ± 0.02	0.00 ± 0.00	0.86 ± 0.03	0.00 ± 0.00
HalfCheetahVelocity	0.82 ± 0.05	0.00 ± 0.00	0.85 ± 0.04	0.00 ± 0.00	0.79 ± 0.00	0.00 ± 0.00	0.82 ± 0.05	0.00 ± 0.00	0.80 ± 0.02	0.00 ± 0.00
SwimmerVelocity	0.45 ± 0.11	0.00 ± 0.00	0.49 ± 0.06	0.00 ± 0.00	0.48 ± 0.11	0.00 ± 0.00	0.54 ± 0.06	0.00 ± 0.00	0.50 ± 0.09	0.00 ± 0.00
Average (3)	0.71 ± 0.23	0.00 ± 0.00	0.74 ± 0.22	0.00 ± 0.00	0.71 ± 0.21	0.00 ± 0.00	0.75 ± 0.18	0.00 ± 0.00	0.72 ± 0.20	0.02 ± 0.04
AntCircle	0.45 ± 0.03	0.00 ± 0.00	0.26 ± 0.22	0.00 ± 0.00	0.36 ± 0.31	0.20 ± 0.35	0.39 ± 0.22	0.00 ± 0.00	0.42 ± 0.07	0.87 ± 1.50
BallCircle	0.46 ± 0.09	0.00 ± 0.00	0.42 ± 0.03	0.00 ± 0.00	0.44 ± 0.03	0.00 ± 0.00	0.40 ± 0.05	0.00 ± 0.00	0.35 ± 0.03	0.00 ± 0.00
CarCircle	0.63 ± 0.03	0.00 ± 0.00	0.56 ± 0.05	0.00 ± 0.00	0.54 ± 0.06	0.00 ± 0.00	0.52 ± 0.04	0.00 ± 0.00	0.52 ± 0.04	0.00 ± 0.00
DroneCircle	0.36 ± 0.05	0.00 ± 0.00	0.32 ± 0.04	0.00 ± 0.00	0.29 ± 0.08	0.00 ± 0.00	0.32 ± 0.07	0.00 ± 0.00	0.23 ± 0.16	0.00 ± 0.00
Average (4)	0.47 ± 0.11	0.09 ± 0.18	0.39 ± 0.13	0.08 ± 0.16	0.41 ± 0.11	0.12 ± 0.14	0.41 ± 0.08	0.08 ± 0.16	0.38 ± 0.12	0.06 ± 0.12
easyparse	0.29 ± 0.04	0.08 ± 0.04	0.30 ± 0.31	0.34 ± 0.56	0.11 ± 0.25	0.03 ± 0.05	0.49 ± 0.01	0.29 ± 0.36	0.29 ± 0.26	0.07 ± 0.06
mediumsparse	0.82 ± 0.04	0.39 ± 0.43	0.74 ± 0.08	0.12 ± 0.09	0.51 ± 0.50	0.22 ± 0.26	0.73 ± 0.11	0.28 ± 0.34	0.51 ± 0.12	0.15 ± 0.05
hardspare	0.09 ± 0.21	0.01 ± 0.01	0.30 ± 0.06	0.08 ± 0.07	0.38 ± 0.01	0.23 ± 0.12	0.37 ± 0.06	0.11 ± 0.08	0.30 ± 0.10	0.10 ± 0.08
Average (3)	0.40 ± 0.38	0.16 ± 0.20	0.45 ± 0.25	0.18 ± 0.14	0.33 ± 0.20	0.16 ± 0.11	0.53 ± 0.19	0.22 ± 0.10	0.37 ± 0.13	0.10 ± 0.04

Table 4: Results for comparison of normalized cost and normalized reward across range of cost-tau at a fixed reward-tau where $\uparrow$ indicates that a higher value is better, while $\downarrow$ indicates that a lower value is preferable. **Bold** text denotes safe policies, **Red** indicates unsafe policies, and **Blue** highlights the highest reward among the safe policies for each task. Each result is obtained as an average by running three seeds of (0,10,20) and evaluating over 20 episodes.

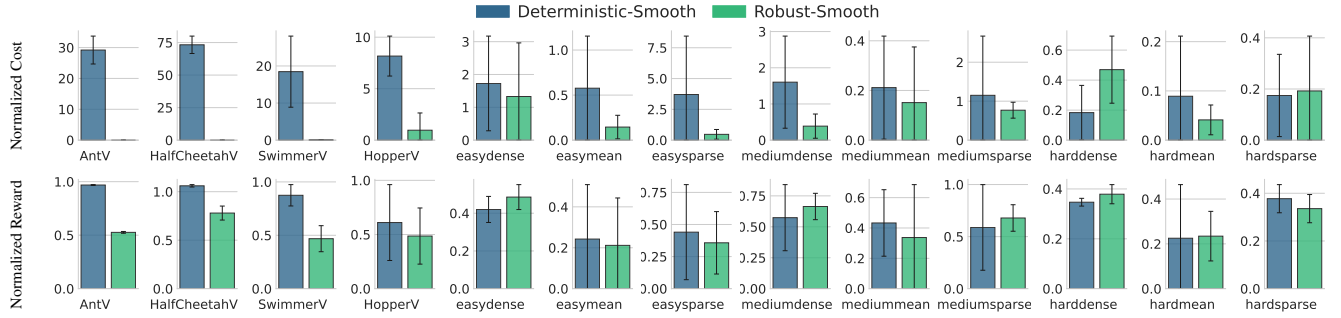


Figure 4: **Effect of asymmetric aggregation on safety and reward propagation.** Experiment investigates whether the robust version of CBF discussed in Sec 4.2 offers an advantage under uncertainty. We modify the DSRL with outliers and use a Smooth barrier update. Figure compares the Standard smooth version against the Robust smooth version, which utilizes an asymmetric loss when learning the Q-value to approximate the maximum cost over the transition function. We report normalized cost (top) and reward (bottom) across tasks. With identical settings, Robust-Smooth achieves lower costs and comparable rewards.

where $R_\pi = \sum_t r_t$ denotes the cumulative reward return of a policy and $C_\pi = \sum_t c_t$ denotes its cumulative cost return. In all benchmark tasks, the per-step cost $c_t = 0$ if the state is safe and $c_t > 0$ for unsafe states. The quantities $r_{\min}$ and $r_{\max}$ correspond to the minimum and maximum reward returns observed in the offline dataset, respectively. The constant l denotes a predefined cost limit, and ϵ is a small positive constant added to ensure numerical stability. Following this protocol, a task is considered safe if the cumulative cost does not exceed the cost limit l , or equivalently, if the normalized cost return is no greater than one.

C.2 Training Details

All value functions are implemented as two-layer multilayer perceptrons with ReLU activation functions and 256 hidden units per layer. We use a Gaussian model as the actor policy in Safety Gym and Bullet Gym while we use a Diffusion model as the actor policy in Metadrive to cater to the high dimensionality (259) of the state space. Expectile regression is employed for value learning, with the expectile parameter set to $\tau_r = 0.75$ for reward value learning and $\tau_h = 0.15$ for cost value learning. During policy optimization, exponential advantages are clipped to $(-\infty, 100]$. We further use reward temperature to control the sharpness of the policy update. We perform function approximation using a constraint violation function $h(s)$. In practice, the benchmark environments used in our experiments do not provide an explicit constraint violation signal, but instead provide a per-state cost function $c(s)$. To address this, when $c(s) = 0$, the state is considered safe and we set $h(s) = -1$, whereas when $c(s) > 0$, the state is considered unsafe and we set $h(s) = M$, where M is the cost scale. The cost scale used to normalize the raw cost values is set to 1 for Safety Gym and Bullet Gym, while for Metadrive it is set to 25 to account for the sparse data.

Table 5: Hyperparameters used for Barrier-Guided Offline Actor–Critic (CBF)

Parameter	Value
Optimizer	Adam
Learning rate	$3e^{-4}$
Discount factor γ	0.99
Soft target update rate α	0.001
Reward expectile τ_r	0.75
Cost expectile τ_h	0.15
Reward temperature β	3.0
Minimum reward $r_{\min}$	$-1e^{-3}$
Cost upper bound c_{ub}	150
Cost limit	10
Hidden layers (value / critic)	2
Neurons per hidden layer	256
Activation function	ReLU
Number of times Gaussian noise is added T	5
Number of action candidates N	16
Training steps	5×10^5

C.3 Asymmetric Losses

Percentile loss, often called Pinball loss, generalizes the absolute error (L_1) to estimate conditional quantiles. By weighting positive and negative errors differently based on $\tau \in (0, 1)$, it provides a robust metric for interval prediction that is less sensitive to extreme outliers than squared methods.

$$L_{\tau}^{\text{pin}}(y, \hat{y}) = \begin{cases} \tau(y - \hat{y}) & \text{if } y \geq \hat{y} \\ (\tau - 1)(y - \hat{y}) & \text{if } y < \hat{y} \end{cases} \quad (18)$$

Expectile loss generalizes the squared error (L_2) to estimate expectiles, which are the squared-loss analogs to quantiles. It applies asymmetric weights to the squared residuals, offering the advantage of being differentiable everywhere while still allowing the model to focus on the tails of the distribution.

$$L_{\tau}^{\text{exp}}(y, \hat{y}) = \begin{cases} \tau(y - \hat{y})^2 & \text{if } y \geq \hat{y} \\ (1 - \tau)(y - \hat{y})^2 & \text{if } y < \hat{y} \end{cases} \quad (19)$$

C.4 Hyperparameters

We use Adam optimizer with a learning rate of $3e^{-4}$. We use a batch size of 512 for training. We report the detailed setup in Table 5. We use a Gaussian model as the actor policy in Safety Gym and Bullet Gym while we use a Diffusion model as the actor policy in Metadrive to cater to the high dimensionality (259) of the state space.

D Example explaining reward leakage

Consider the simple deterministic MDP shown below:

$$s_0 \xrightarrow{a_s} s_1 \rightarrow r = 1, \quad s_0 \xrightarrow{a_u} s_u \rightarrow r = 10,$$

where:

- a_s is a safe action leading to state s_1 ,
- a_u leads to an unsafe state s_u ,
- the safety constraint is violated at s_u ,
- both transitions are present in the offline dataset.

Assume the agent receives reward only at the terminal state and $\gamma = 1$ for simplicity.

Standard Bellman Backup. A conventional offline RL critic learns:

$$Q(s_0, a) = r(s_0, a) + V(s'),$$

without considering whether the transition is safety-feasible.

Thus:

$$Q(s_0, a_s) = 1, \quad Q(s_0, a_u) = 10.$$

Even after convergence, the learned value function assigns:

$$Q(s_0, a_u) > Q(s_0, a_s),$$

because the Bellman operator propagates reward through the unsafe transition.

However, under the safety constraint, the trajectory through s_u is infeasible and should not contribute to the optimal safe value function.

Safety-Aware Bellman Backup. Our method modifies the Bellman target using the safety critic:

$$y_r = \begin{cases} r + \gamma V_r(s'), & Q_h(s, a) \leq 0, \\ \frac{r_{\min}}{1 - \gamma} - Q_h(s, a), & \text{otherwise.} \end{cases}$$

Since (s_0, a_u) is unsafe ($Q_h(s_0, a_u) > 0$), the unsafe transition receives a pessimistic target:

$$Q(s_0, a_u) \ll 1.$$

As a result:

$$Q(s_0, a_s) > Q(s_0, a_u),$$

and the learned value function converges to the optimal *safe* solution.

Thus, the issue is not merely transient optimization behavior before convergence, but the fixed point toward which the Bellman updates converge. Standard Bellman backups optimize expected return over all dataset transitions, including trajectories that violate safety constraints. Consequently, the learned reward critic converges to the optimal unconstrained value function, rather than the optimal value function restricted to safe trajectories. In particular, high future rewards that are only reachable through unsafe intermediate transitions can still propagate backward through bootstrapping, assigning artificially high value to states/actions that are infeasible under the safety constraint. This effect persists even after convergence because the unsafe transitions remain part of the Bellman operator itself. Our safety-aware Bellman target modifies the operator directly by assigning pessimistic targets to unsafe transitions, ensuring convergence toward a value function consistent with the forward-invariant safe set.