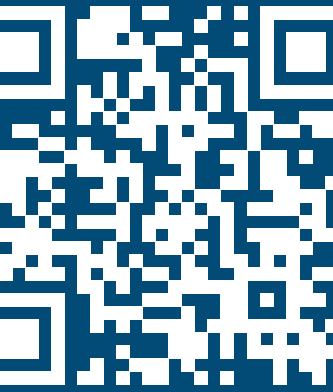


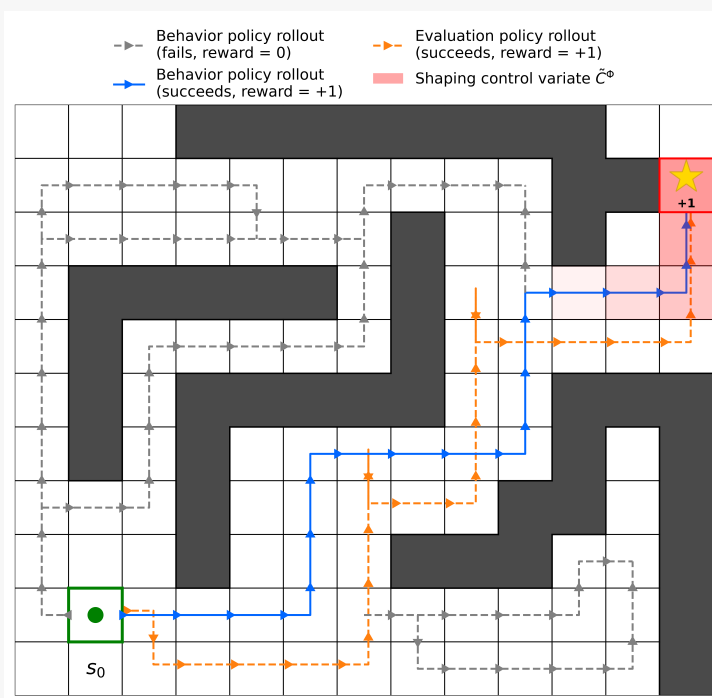


Aim of the Paper

Develop **unbiased** OPE estimators for **sparse-reward** settings using **potential-based reward shaping** as zero-mean control variates.

- Reduce variance **without adding bias**.
- Extract signal from **every** trajectory not just rare successes.
- Practical Applications: Healthcare, Education, Finance, where sparse and delayed rewards are prevalent.

The Sparse Reward OPE Problem



- Most trajectories give **zero reward**.
- OPE dominated by rare successes $\Rightarrow$ **high variance**.
- Shaped** estimators use **dense state signal** as control variates.

Why Existing Methods Fail

- IS**: Exploding IS-weights on rare successes
- DR**: Must learn Q from $\sim \emptyset$ TD targets.
- MIS**: Density-ratio objective ill-conditioned when reward support is sparse

Reward Shaping Control Variates: Intuition

- Rewards are sparse, but **state visitations are dense**.
- Learn $\Phi(s)$: a **progress score** over states.
- Use Φ to build a **zero-mean control variate** for the IS return.
- Trained *only* from how well it correlates with returns, no reward labels needed.
- Signal from **all** trajectories, not just successes.

Proposed Shaping Estimator

$$\hat{V}_{\text{Shaped-PDIS}}^{\pi_e}(\lambda) = \frac{1}{N} \sum_{n=1}^N \left[\underbrace{Y^{(n)}}_{\text{IS return}} + \lambda^\top \underbrace{\tilde{C}^{\Phi, (n)}}_{\text{shaping control variate}} \right]$$

Building blocks

- PDIS return**: $Y^{(n)} = \sum_{t=0}^{T-1} \gamma^t W_t r_t$
Standard PDIS estimator. High-variance under sparse rewards.
- Shaping feature**:
 $C_t^\Phi = \gamma^t W_t (\gamma \Phi(s_{t+1}) - \Phi(s_t))$
Potential-based shaping, responsible for variance reduction.
- Boundary**: $\Phi(s_T) = 0$
Anchors the potential and forces telescoping.

Why it works

- Centered CV**: $\tilde{C}^\Phi = C^\Phi - \mathbb{E}_{\pi_b}[C^\Phi]$
Zero-mean by construction $\Rightarrow$ adding $\lambda^\top \tilde{C}^\Phi$ never changes the expectation.
- Optimal weight**: $\lambda^* = -\Sigma_{CC}^{-1} \Sigma_{CY}$
Closed-form: projection of Y onto the span of shaping columns.
- Variance reduction**:
 $\text{Var}(Y + \lambda^{*\top} \tilde{C}^\Phi) = \text{Var}(Y) - \Sigma_{CY}^\top \Sigma_{CC}^{-1} \Sigma_{CY}$
Strict reduction whenever $\tilde{C}^\Phi$ correlates with Y. Never larger than $\text{Var}(Y)$.

Training Objective: Learning Φ_β for Variance Reduction: $\max_{\beta} J(\beta) = \Sigma_{CY}(\beta)^\top \Sigma_{CC}(\beta)^{-1} \Sigma_{CY}(\beta)$

No dense reward labels required: Φ_β is trained *only* to make its shaping columns correlate with the noisy IS return Y.

Same construction extends to **Shaped-DR** and **Shaped-MRDR** by replacing Y with Y_{DR} or Y_{MRDR} .

Advantages over DR / MRDR

- Easier to learn.**
 - DR needs Q(s, a) satisfying *Bellman consistency* across long horizons, under sparse TD targets $\sim \emptyset$.
 - $\Phi(s)$ only needs to *correlate* with Y with sparse reward labels.
- Interpretable.**
 - $\Phi(s)$ is a scalar over states
 - Learned λ_t reveal *which timesteps* drive variance reduction.

Algorithm

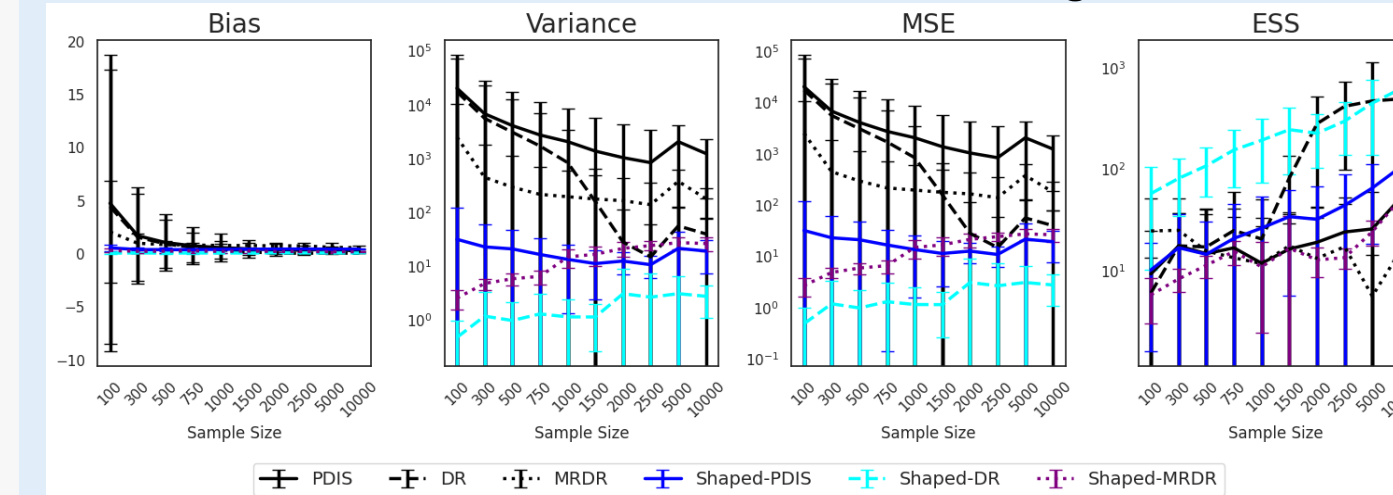
Input: trajectories from π_b , target π_e , γ , ridge α
For $k = 1, \dots, K$:

- Split into fit set ξ_{fit} and eval set ξ_k .
- On ξ_{fit} : build IS returns $Y^{(n)}$ and shaping columns $C^{\Phi, (n)}$.
- Estimate covariances $\hat{\Sigma}_{CC}, \hat{\Sigma}_{CY}$.
- Update β by gradient ascent on $J(\beta) = \Sigma_{CY}^\top \Sigma_{CC}^{-1} \Sigma_{CY}$.
- Compute $\lambda_k = -(\hat{\Sigma}_{CC} + \alpha I)^{-1} \hat{\Sigma}_{CY}$.
- Evaluate on ξ_k :
 $\hat{V}_k^{\pi_e} = \frac{1}{|\xi_k|} \sum_{n \in \xi_k} (Y^{(n)} + \lambda_k^\top C^{\Phi, (n)})$.

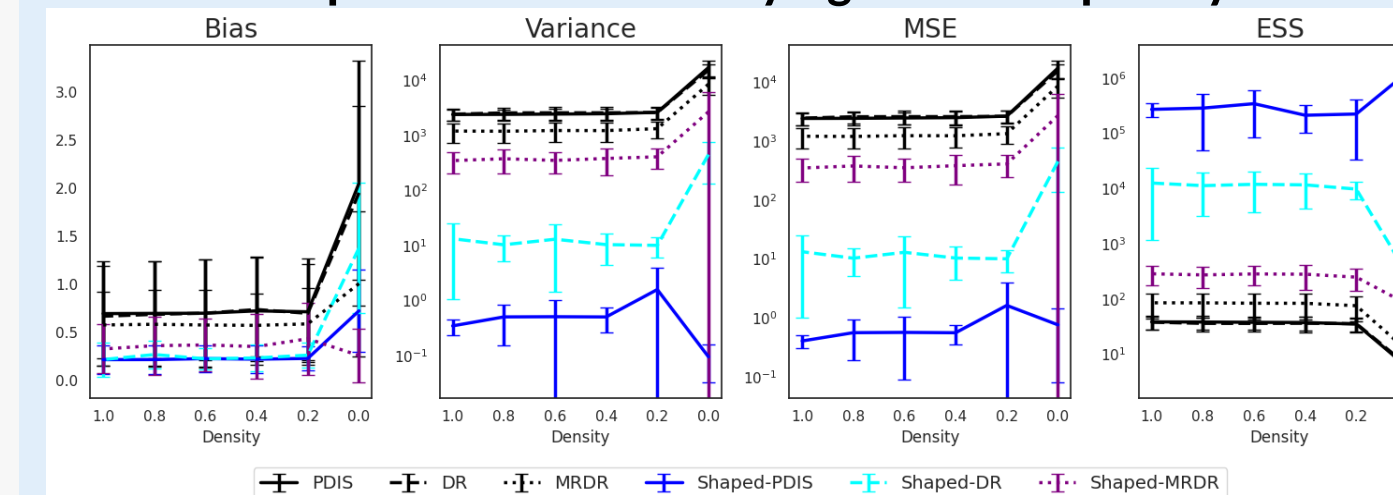
Output: $\hat{V}_{\text{Shaped-PDIS}}^{\pi_e} = \frac{1}{K} \sum_k \hat{V}_k^{\pi_e}$.

Quantitative Improvements

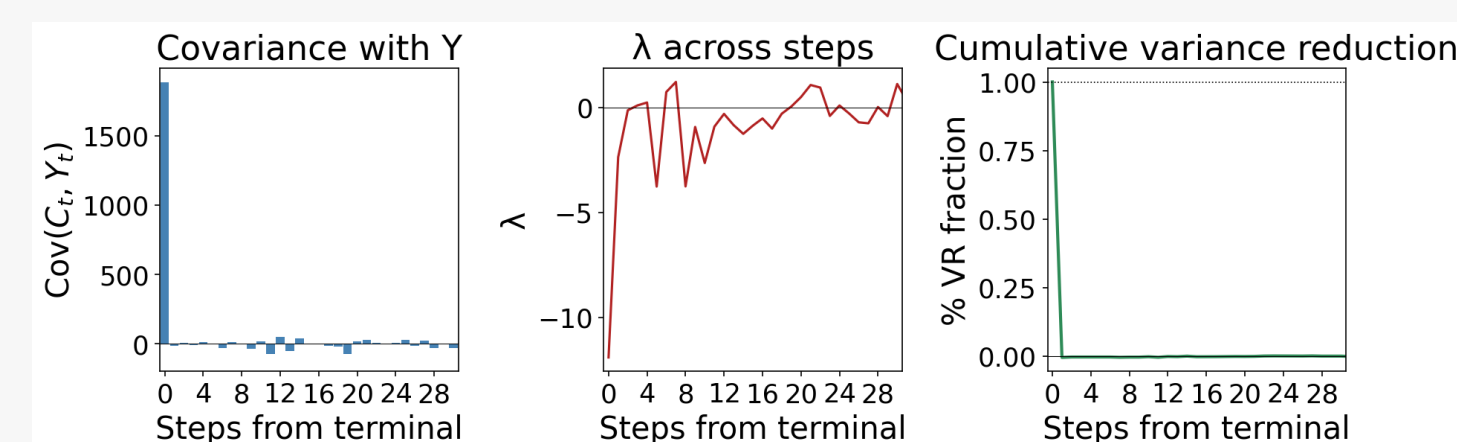
Unbiased, Lower variance, MSE and higher ESS



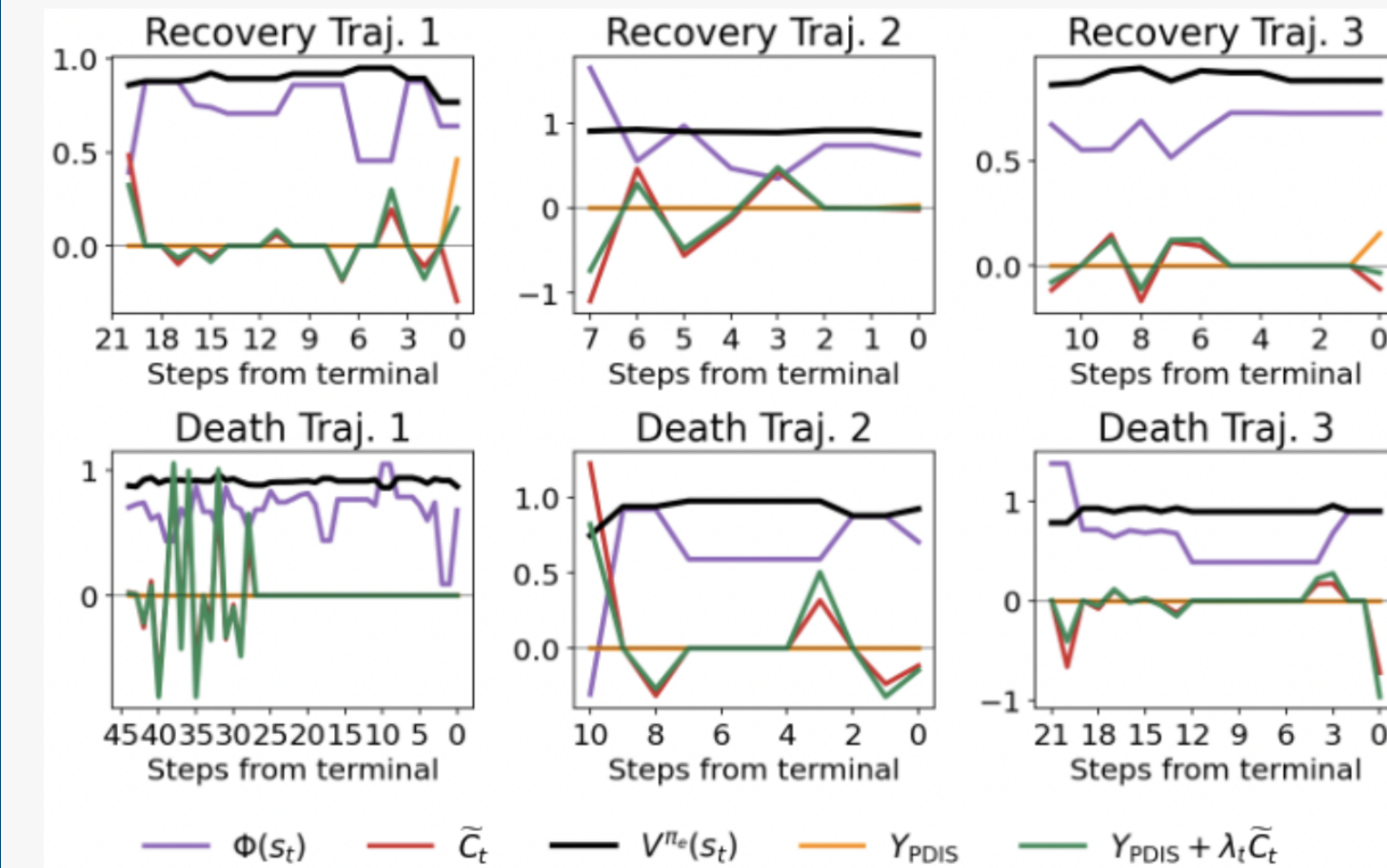
Outperforms across varying levels of sparsity



Variance Reduction Concentrates Near Terminal Reward

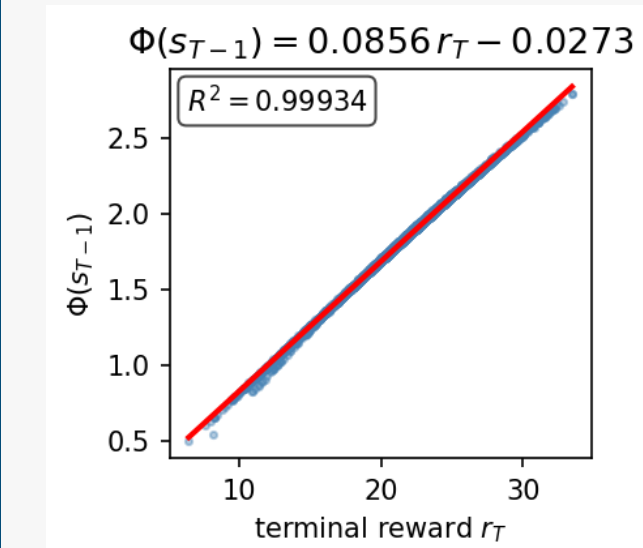


Learnt Control Variate $\tilde{C}^\Phi$ is not the same as value function V^{π_e}



Φ satisfies **no Bellman constraint**, it is optimised purely for variance reduction.

Shaping potential Φ can reveal meaningful structure on state quality



- Slope (~ 0.08) scales returns down and drives variance reduction.
- Sign encodes **state quality**: $\Phi > 0$ high-portfolio, $\Phi < 0$ low.
- Holds when IS-weight dispersion doesn't dominate.

Limitations & Future Work

- Limited shaping of earlier timesteps.
- Long-horizon IS-weights **limit** Φ interpretability.
- State-density **not** explicitly modelled in variance reduction.

Conclusions

- Reward-shaping control variates are **unbiased** by construction.
- Shaped estimators achieve **strict variance reduction** over traditional counterparts.
- State-of-the-art** on sparse, noisy, and misspecified OPE benchmarks across healthcare, finance, and tabular domains.
- Open a path to **interpretable** OPE: Φ and λ expose *where* and *when* variance reduction happens.