

Introduction & Setup

- WAMs jointly predict future states & actions, but remain brittle to camera, gripper, and noise perturbations.
- Activation steering** is a powerful approach to influence model behavior at inference time by directly modifying activations during generation.
- We seek to **improve robustness** by steering features related to **environmental perturbations**.

Mechanistic Analysis of WAMs

- Project contrastive activations onto their top-3 PCs and fit a linear SVM; **hinge loss** measures separability (0=separable, 1=random).
- Key result: separability is architecture-dependent:** strong in Cosmos-Policy/DiT4DiT, weak in LingBot-VA.

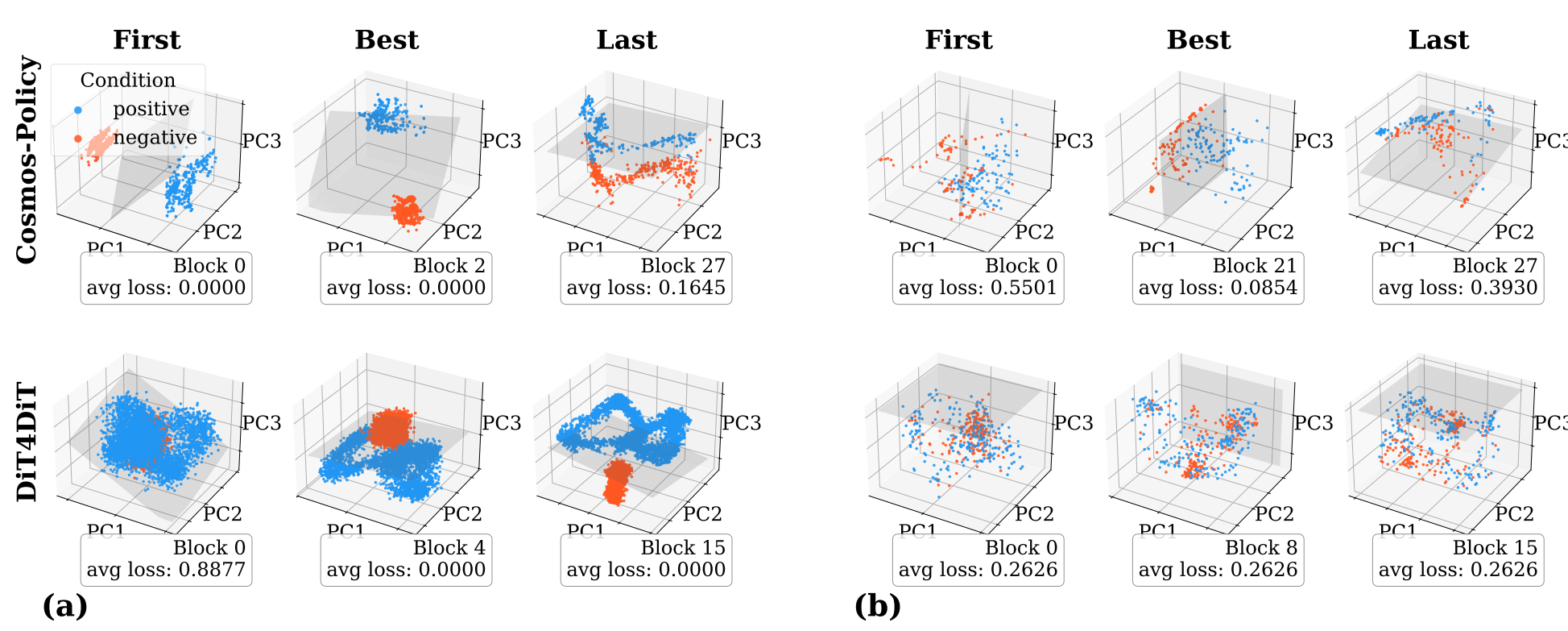


Fig. 2 — Cosmos-Policy & DiT4DiT activations separate cleanly under (a) noise and (b) camera perturbations.

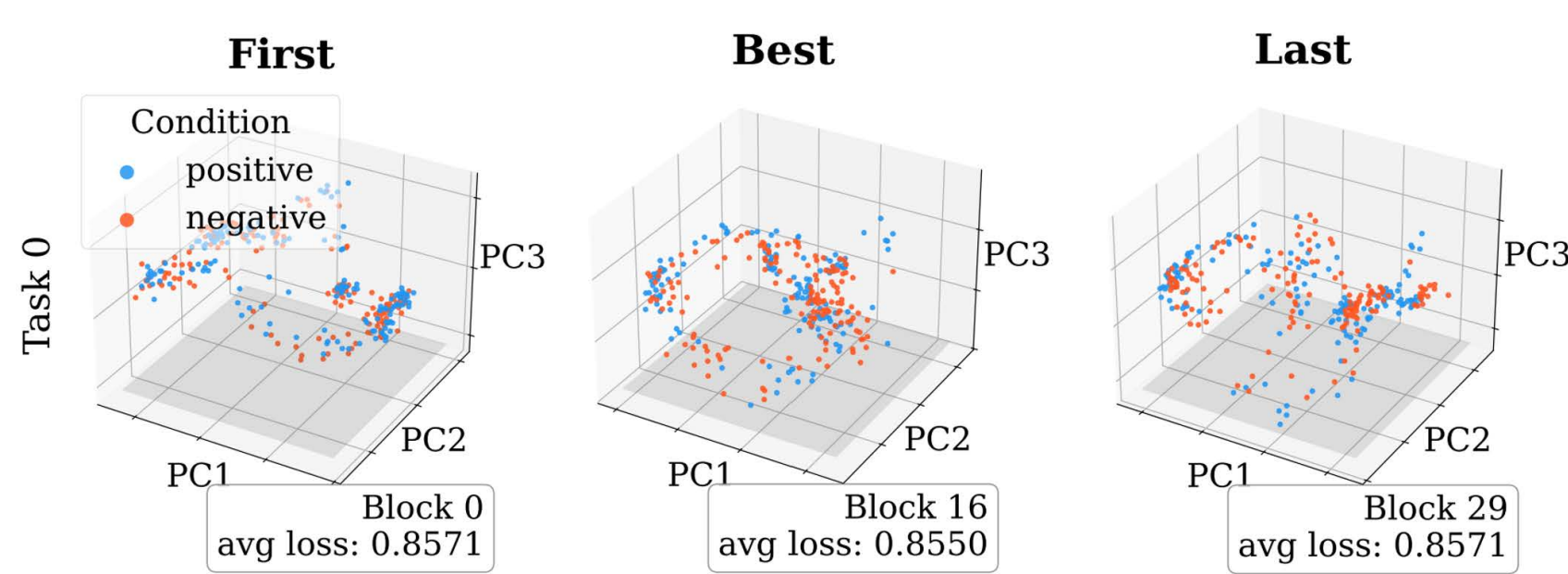


Fig. 4 — LingBot-VA activations overlap heavily — much weaker separability.

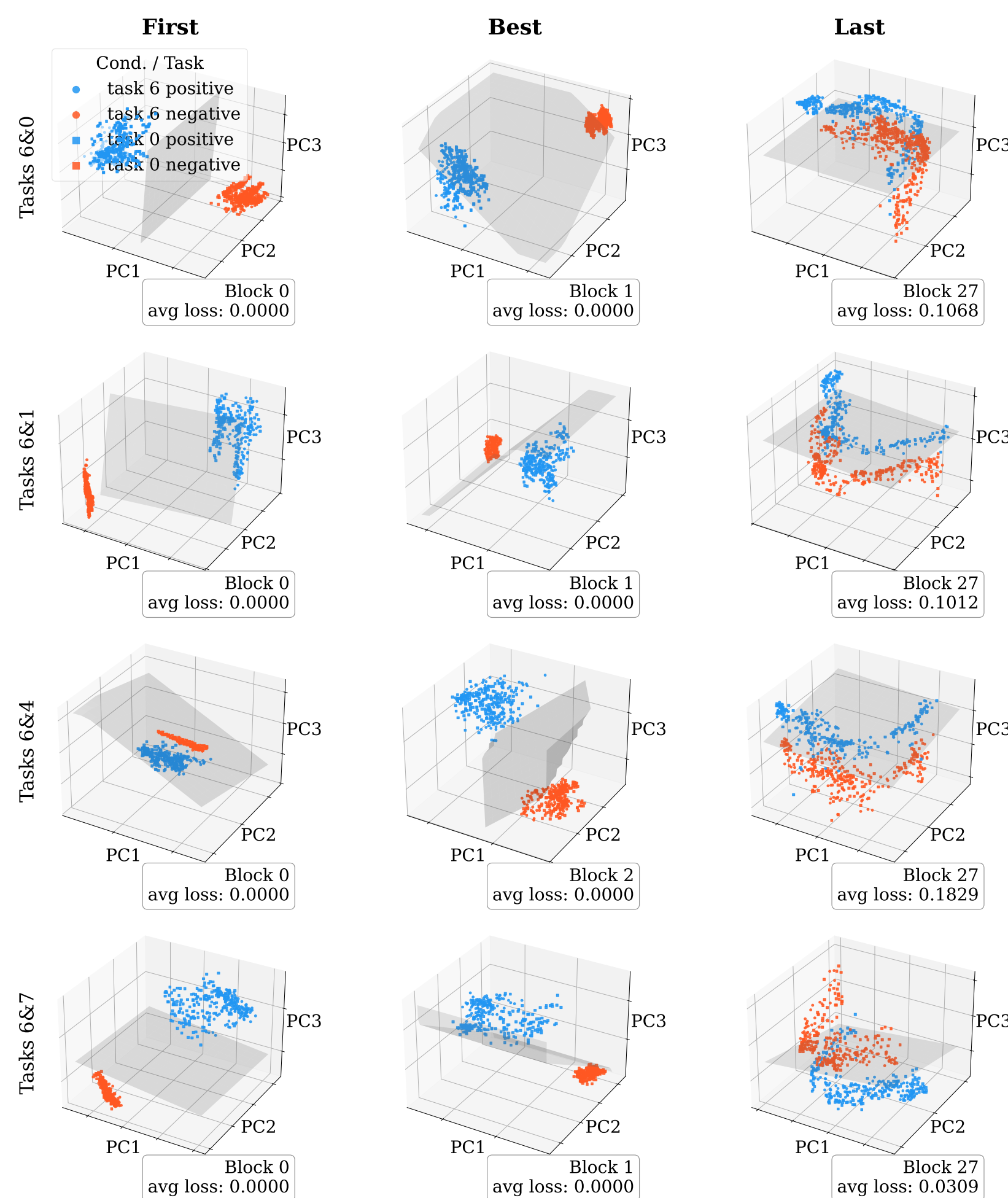


Fig. 3 — Some task pairs share separable directions, enabling cross-task steering.

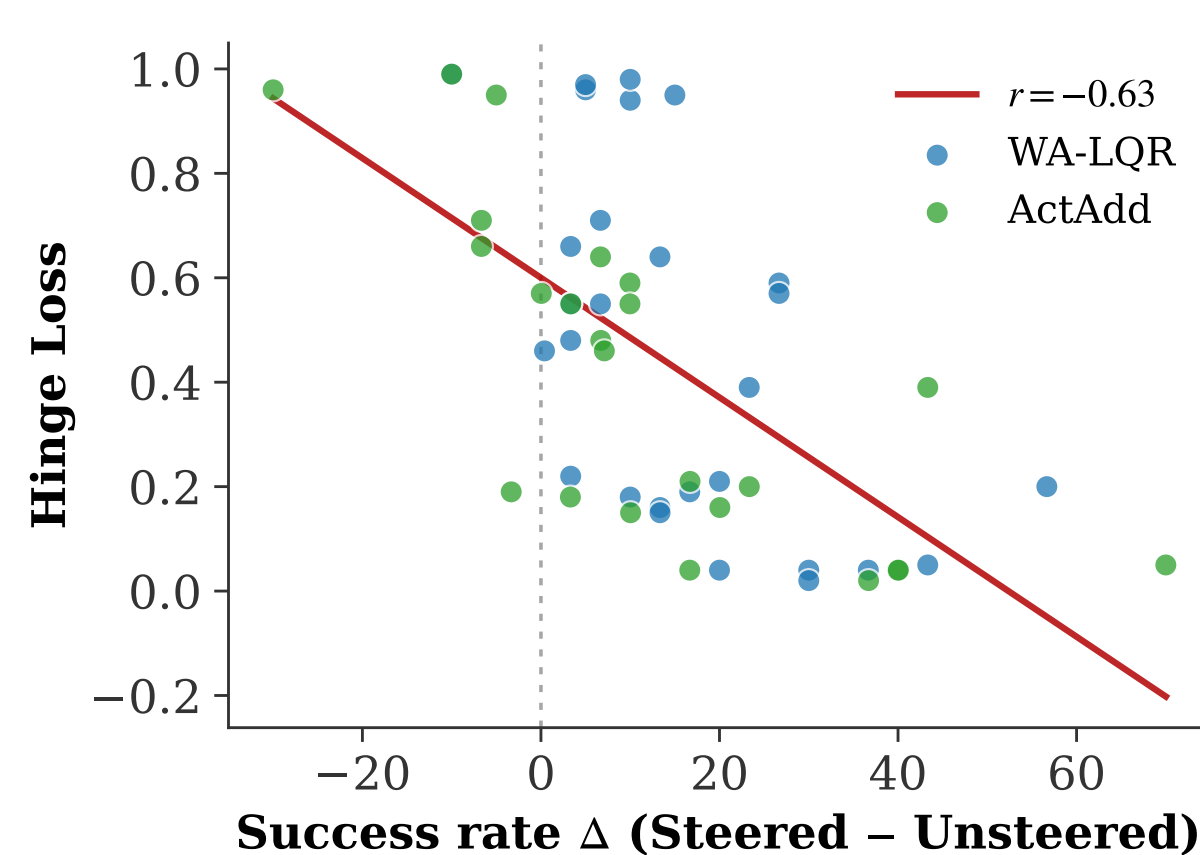


Fig. 7 — Lower hinge loss (better separability) predicts larger steering gains.

Key Takeaway

World Action Models can be steered to be more robust — **but only when contrastive activations are linearly separable**.

We propose an **inference-time alignment** approach that

- is *training-free and weight-preserving*
- adapts *online* (WA-LQR), minimizing *invasive edits*
- improves *Cosmos-Policy & DiT4DiT*, but not *LingBot-VA* — *exactly as predicted*

Separability $\Rightarrow$ *Steerability*, but separability is architecture dependent

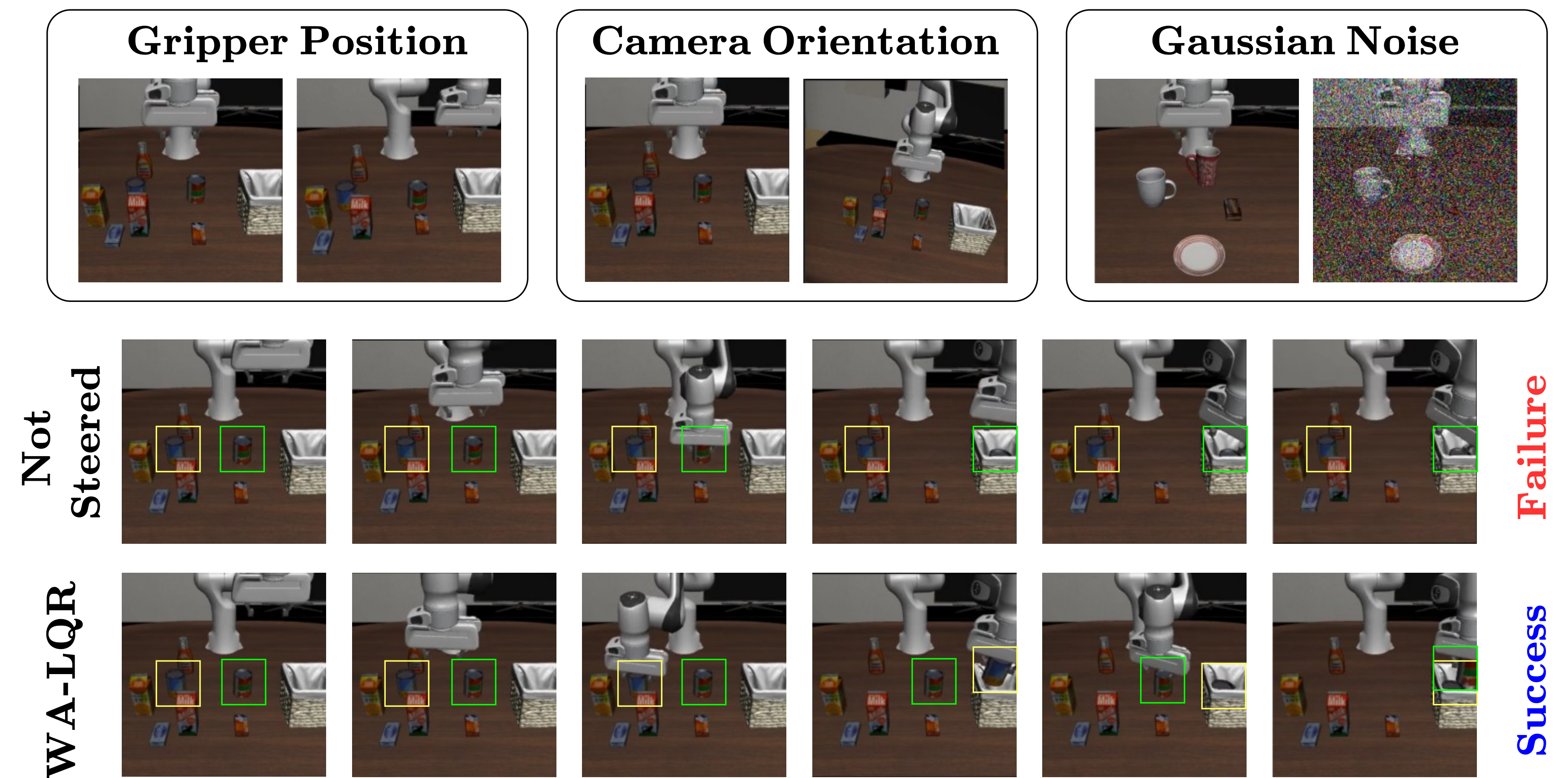


Fig. 1 — Overview: WA-LQR steers WAM activations toward robust behavior under camera, gripper, and noise perturbations.

Method 1: Open-Loop Steering (ActAdd)

Average contrastive directions into a steering vector, then add it at fixed strength γ :

$$\mathbf{a}_{l,t} = \frac{1}{NH_s} \sum_{n,\tau} \mathbf{d}_{l,t,\tau}^{(n)}, \quad \mathbf{x}_{l,t,\tau} \leftarrow \mathbf{x}_{l,t,\tau} + \gamma \mathbf{a}_{l,t}.$$

Simple and training-free, but *open-loop*: it ignores the current activation state and can oversteer.

Method 2: Closed-Loop Steering (WA-LQR)

- Project into a low-dim contrastive subspace via SVD ($d_z \ll d_x$):

$$\mathbf{z}_{l,t,\tau} = \mathbf{P}_{l,t} \mathbf{x}_{l,t,\tau} \in \mathbb{R}^{d_z}.$$

- Linearize latent dynamics across transformer blocks:

$$\delta \mathbf{z}_{l+1,t,\tau} \approx \tilde{\mathbf{A}}_{l,t} \delta \mathbf{z}_{l,t,\tau} + \tilde{\mathbf{B}}_{l,t} \delta \mathbf{u}_{l,t,\tau}.$$

- Track a feature setpoint $\beta_{l,t}^{z,*} = \lambda \| \mathbf{e}_{l,t}^z \|_2$ with online error $\alpha_{l,t,\tau} = \beta_{l,t}^{z,*} - (\mathbf{v}_{l,t}^z)^\top \mathbf{z}_{l,t,\tau}$.

- Solve the LQR tracking problem over the L transformer blocks (Riccati recursion, precomputed offline):

$$\min_{\{\delta \mathbf{u}_{l,t}\}} \sum_l (\delta \bar{\mathbf{z}}_{l,t}^\top \mathbf{Q}_{l,t} \delta \bar{\mathbf{z}}_{l,t} + \delta \mathbf{u}_{l,t}^\top \mathbf{R}_{l,t}^{\text{chunk}} \delta \mathbf{u}_{l,t}) \quad \text{s.t.} \quad \delta \bar{\mathbf{z}}_{l+1,t} = \tilde{\mathbf{A}}_{l,t} \delta \bar{\mathbf{z}}_{l,t} + \tilde{\mathbf{B}}_{l,t} \delta \mathbf{u}_{l,t}.$$

- Closed-form control law — apply the precomputed gains $\mathbf{K}_{l,t}$ online:

$$\mathbf{u}_{l,t,\tau}^* = \bar{\mathbf{u}}_{l,t,\tau} + \mathbf{K}_{l,t,\tau} \alpha_{l,t,\tau}$$

Intervention strength scales automatically with the *online tracking error* — steering only when, and as much as, needed. $\mathbf{R}_{l,t}^{\text{chunk}}$ decays across the action chunk, so early timesteps are steered most.

Why Is WA-LQR Tractable?

Two empirical properties of WAM activations justify the reduced-order LQR design above.

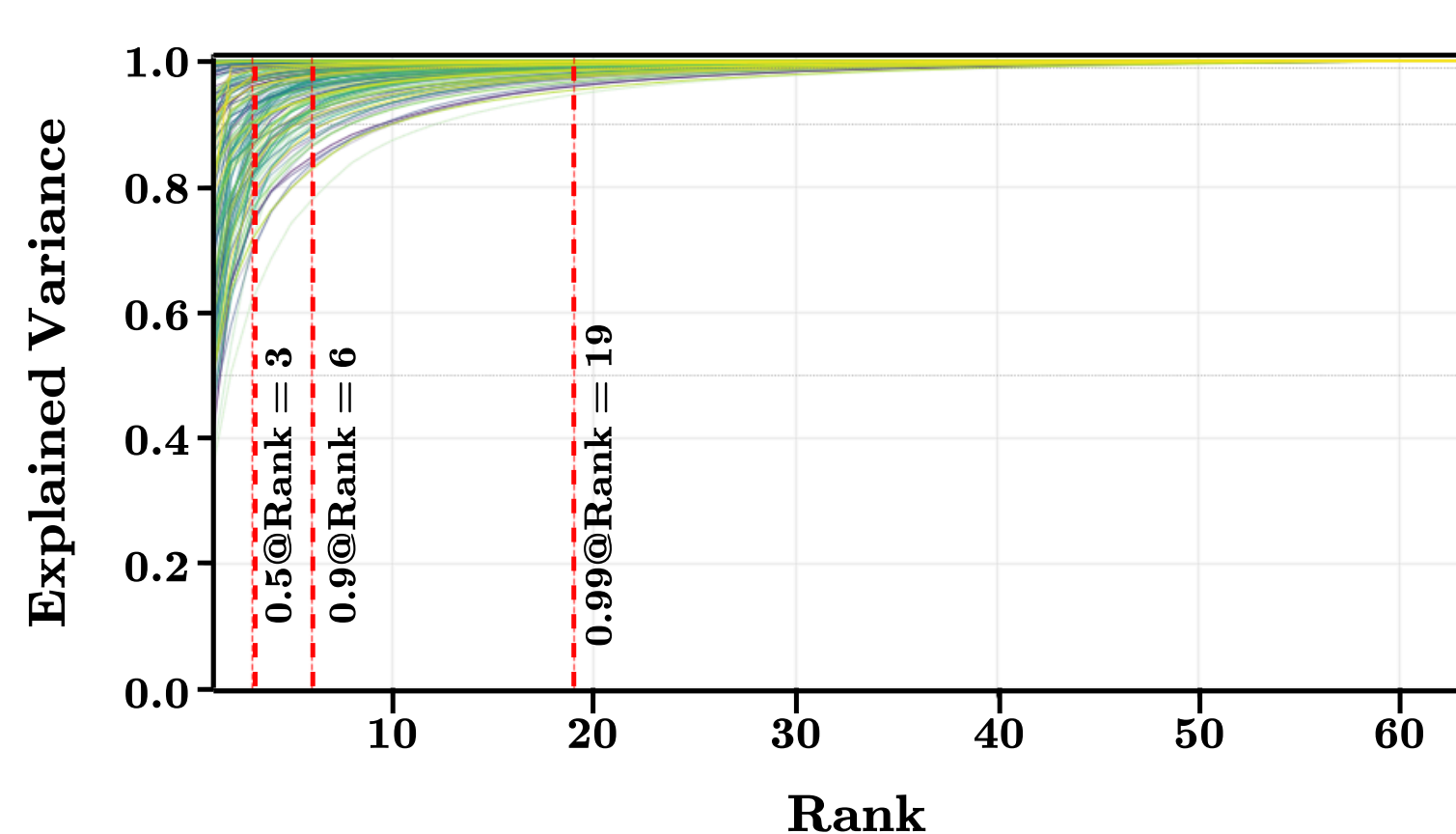


Fig. 5 — Cumulative variance of latent contrastive vectors, explained by the top- k singular vectors — variance concentrates in a low-dimensional subspace.

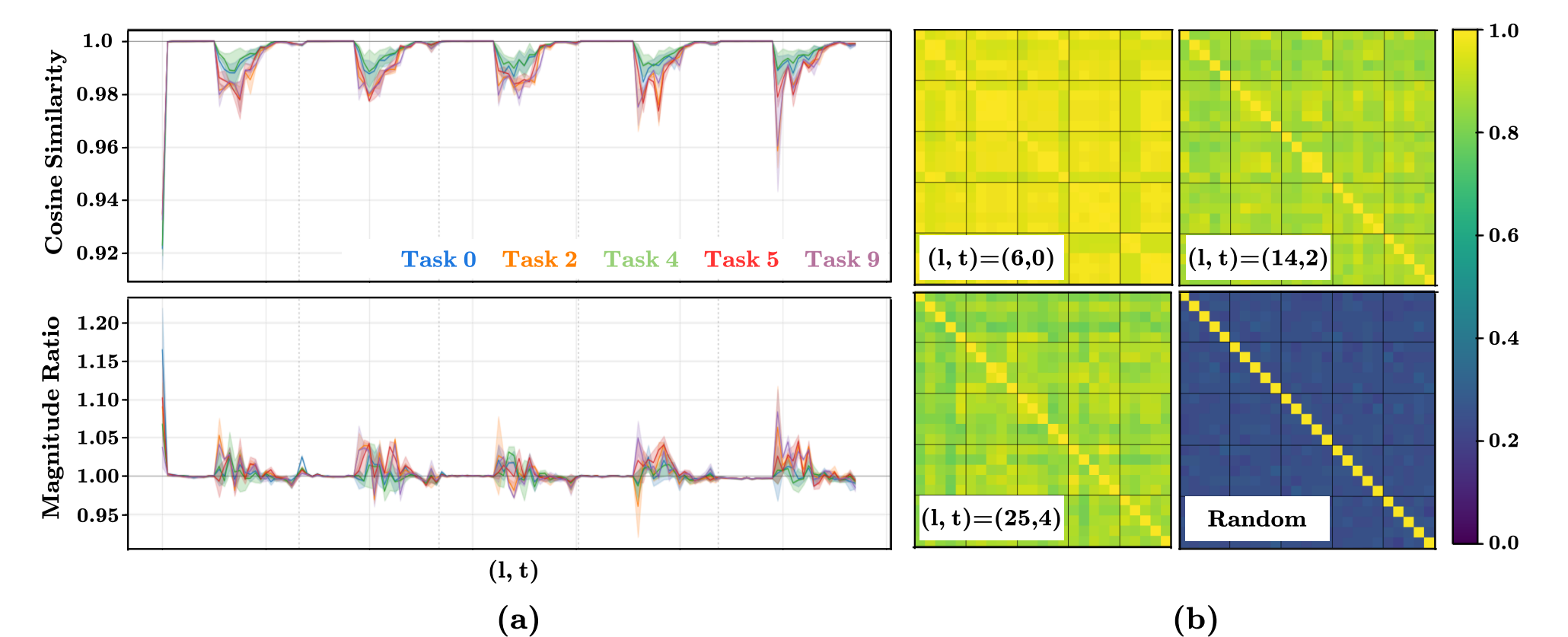


Fig. 6 — (a) Cosine similarity & magnitude ratio between the true update and its linear approximation: activation dynamics are **locally linear**. (b) Subspace overlap across tasks — one controller **transfers**.

Results

Table 1. Success rate on LIBERO-10 under perturbations.

Model	Perturbation	No Steer	ActAdd	WA-LQR
Cosmos-Policy	Camera Orientation	46.0%	49.3%	59.3%
	Init. Gripper Pos.	61.3%	63.3%	72.7%
	Gaussian Noise	26.7%	67.3%	58.7%
DiT4DiT	Init. Gripper Pos.	65.7%	65.8%	71.7%
	Gaussian Noise	15.6%	43.3%	48.9%
LingBot-VA	Camera Orientation	45.0%	43.3%	48.3%

- WA-LQR improves Cosmos-Policy & DiT4DiT robustness across perturbation types.
- LingBot-VA gains are small/inconsistent — as predicted by its weak separability.

Summary & Limitations

- We study **mechanistic interpretability** and **activation steering** for robustifying World Action Models (WAMs).
- Linear **separability** of robustness-relevant features **predicts steerability** — and both are architecture-dependent.
- WA-LQR**, exploits locally-linear activation dynamics for training-free robustness gains.
- No method yet predicts *a priori* which tasks share transferable directions.
- Steerability depends strongly on the backbone architecture.