



ShadowDancer: Teaching Video World Models Any Action by Learning Unified Dynamics Representations from a Video and Its Shadow

Jin Cao¹ Zian Meng^{1,2} Kaipeng Zhang^{1†}
¹Alaya Lab ²Shanghai Innovation Institute
<https://ShadowDancer-1.github.io>

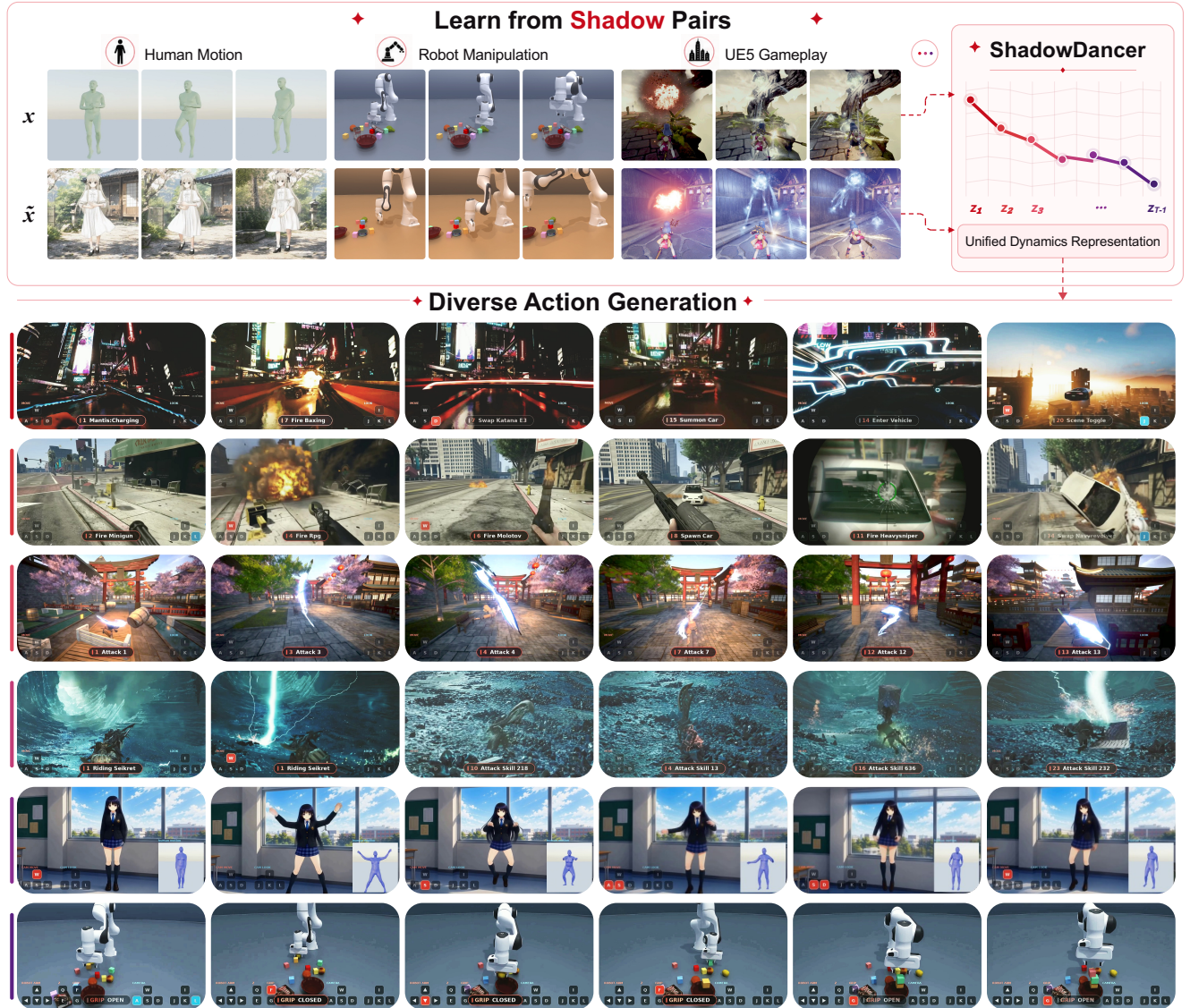


Figure 1. **ShadowDancer** learns any action from a video and its shadow. Top: shadow pairs $(x, \tilde{x})$ replay one dynamics under independently resampled appearance, across heterogeneous sources such as human motion, robot manipulation, and open-world gameplay, and distill it into a unified dynamics representation $z_{1:T-1}$. Bottom: a single model driven through this one latent interface executes diverse commanded actions, spanning first- and third-person combat, driving, locomotion, and robot manipulation.

Abstract

We present ShadowDancer, a novel approach to any-action, frame-level control of interactive video world models. The obstacle is representational: existing interfaces either encode an action loosely, leaving how it unfolds for the model to improvise, or encode it exactly through structured signals that serve one family and are hard to acquire, so precise control across diverse dynamics remains impractical. Demonstration videos are the natural remedy, specifying any dynamics frame by frame; yet a video shows its dynamics only through one particular appearance, a single shadow of the underlying dynamics, so actions learned from demonstrations transfer poorly to new scenes. ShadowDancer addresses this with two key innovations: (1) shadow pairs, video pairs that replay the same dynamics under independently resampled appearance, constructed at scale by our Shadow Library, so that a dynamics family becomes controllable exactly when such pairs can be constructed for it; and (2) cross-shadow prediction, which learns actions by predicting one shadow from the other, so that whatever the pairing resamples is discarded by construction and whatever it preserves becomes the action, yielding a unified dynamics representation that drives a block-causal world model. Any demonstrated clip thus becomes a reusable action asset, replayed in new environments without action labels, motion estimators, or fine-tuning. Experiments demonstrate improved action transfer and long action rollout over strong latent-action and interactive world model baselines across diverse dynamics families. We show video results at <https://ShadowDancer-1.github.io>.

1. Introduction

Interactive world models have emerged as a transformative capability in video generation, synthesizing persistent virtual worlds that users can explore in real time [14, 19, 26, 41, 43, 49, 52]. However, despite the boundless variety of dynamics a world can host, existing approaches lack a unified means of specifying how those dynamics should unfold. The result is a stubborn trade-off between precision and generality: a behavior can be specified loosely for arbitrary dynamics, or exactly for one narrow family, but rarely both at once. This leaves worlds that can be explored but not yet directed, restricting their applicability for interactive entertainment and world simulation. We therefore target *any-action, frame-level control*: precision, so that a command specifies how an action unfolds frame by frame, and generality, so that the same interface extends to any dynamics family.

The bottleneck is not generative capacity, as video diffusion backbones already synthesize complex articulated

motion well [5, 9, 36, 55]; it lies in the *representation* of actions: how the user communicates which dynamics should unfold. Any-action, frame-level control is at its core a representation learning problem, and each existing interface can be read as one candidate representation. Symbolic commands [14, 26, 63] were designed around logged interactions such as keystrokes: they name an event from a finite vocabulary and leave its trajectory for the generator to improvise, so holding a key longer repeats an action yet hardly reshapes it. Free-form text [19, 41, 65] removes the vocabulary but not the ambiguity, since language describes dynamics rather than supplying it. Structured motion inputs such as 3D human motion [8, 16], point tracks [37, 47], and camera pose trajectories do reach frame-level precision, but each format serves a single family, demands its own conditioning machinery, and depends on specialized yet often brittle estimators, so the precise signal it promises is hard to obtain in practice. What remains is the demonstration itself, a natural specification that is temporally dense and format-free, and the way humans are taught new actions. Latent action models already exploit it, distilling the transitions of a reference clip into continuous latent actions z that condition generation [6, 11, 15, 46, 61]. But critically, such models are trained to reconstruct the very clip they read, and self-reconstruction fundamentally conflicts with *control*: reconstruction asks z merely to explain the observed appearance, whereas control asks it to transfer the underlying dynamics to new appearances [20, 32, 34, 40, 60]. Trained without ever observing the same dynamics under two different appearances, the latent cannot tell which traits belong to the action and which merely co-occur with it: it entangles what moves with how it looks indiscriminately, and shrinking it [18] or regularizing it toward semantic features [32] operates within the same supervision: invariance can only be encouraged as a penalty, and the unchosen residue can resurface as spurious motion in the generation. What the task calls for is to learn the representation from data in which dynamics is separated from appearance by construction.

To address this challenge, we propose **ShadowDancer**, built on the idea of observing the same dynamics *twice*. Our approach introduces two key innovations. (1) *Shadow pairs*: a Shadow of a video is a second render of the same dynamics under a different appearance. The two videos are frame-synchronized in motion, while subject, scene, materials, and lighting are independently resampled. Constructing shadows is a protocol rather than an engine: replay the dynamics and resample everything else. Our Shadow Library implements this protocol across animation suites, open-world games, and robotic simulators. (2) *Cross-shadow prediction*: an encoder extracts latent actions from one video, and a decoder predicts the other shadow from those latents under the shadow’s own appearance context. Since the target appearance is already supplied, the prediction requires transferring dynamics rather

[†]Corresponding author.

than reconstructing appearance, making invariance a property of the supervision rather than a penalty. The pairing defines the action itself: whatever the pair shares becomes the controllable factor. What the representation couples to is thus no longer an accident of the data but a choice of the protocol. Pairing on a body motion while resampling the camera yields camera-free control, whereas pairing on a camera trajectory across scenes yields pure camera control. More generally, any dynamics family becomes controllable once a shadow pair can be constructed for it, with no vocabulary to design and no labels to collect. The resulting latent stream is a unified dynamics representation, driving every action family through one interface. It conditions a pretrained video diffusion backbone [9, 55], which we convert into a block-causal autoregressive generator [23, 30] for interactive rollout.

In use, directing an action takes one demonstration: the dynamics of the clip is extracted in a single frozen-encoder pass, stored as a variable-length action asset, composed with other assets along the condition stream, and replayed in a new environment. Unlike symbolic interfaces, which specify what should happen and leave how it happens to the generator, ShadowDancer supplies a dense control signal read from the trajectory itself, so the timing, amplitude, and style of an action are inherited from the demonstration. We evaluate this interface in two settings: action transfer, which tests how precisely a demonstration from one environment is re-enacted in another, and long action rollout, which tests how well stored assets compose over time. In summary, our contributions are threefold:

- We introduce the Shadow learning paradigm, which learns, from a video and its Shadow, dynamics representations invariant to everything the pairing resamples, resolves latent-action non-identifiability by construction, and provides a formal analysis of when the shared dynamics is identifiable.
- We build the Shadow Library, a source-agnostic pairing protocol implemented through paired-rendering scripts across animation suites, open-world games, and robotic simulators, supplying shadow supervision for families including human, camera, and object dynamics without a single action label.
- We present ShadowDancer, an any-action interactive world model driven by dense, variable-length action assets, and show improved action transfer and long action rollout over strong latent-action and interactive world model baselines across the evaluated dynamics families.

2. Related Work

Interactive Video World Models. World models predict future observations and support planning or interactive simulation in games, robotics, and driving [1, 17, 24, 25], and recent systems generate explorable worlds at impressive fi-

delity and horizon [14, 19, 26, 41, 43, 49, 52]. Their control interfaces, however, are symbolic or textual: frame-level keyboard and mouse states logged from game engines [3, 14, 59, 63], or free-form text prompts [19, 41, 65]. These commands specify what happens but not how (Sec. 1), and telemetry-logged interfaces further bind each model to the action schema of its data-collection pipeline [26, 28, 50, 62]; structured motion inputs [8, 16, 37, 47] are precise instead, but per-family and hard to acquire accurately. ShadowDancer instead conditions on a dense per-frame latent extracted from a reference clip, so that the speed, amplitude, and style of an action are supplied by the reference, and any dynamics family with constructible shadow pairs is controlled through one interface.

Latent Actions as Control Interfaces. Latent action models (LAMs) infer controls directly from unlabeled video, serving as interfaces for interactive generation [6, 18, 31], cross-embodiment policy learning [7, 11–13, 35, 61], and world-model pretraining [15, 20, 46, 57]. Unlike imitation learning, which aims to execute a demonstrated behavior on a target embodiment, our focus is the representation itself: a reusable dynamics signal, extracted from a demonstration, that transfers across appearances and environments. The core difficulty is that inverse-dynamics latents couple control to whatever appearance happens to co-occur with it: clip-local reconstruction admits shortcut solutions through context cues [20, 60] and is non-identifiable across contexts [34, 40], so the same underlying dynamics need not map to a consistent latent [32]. Existing remedies constrain the latent (information bottlenecks, quantization) [6, 18] or regularize it toward motion-centric references [32]; in practice, regularization reduces but does not eliminate the entanglement. All of these operate on a single observation of each transition, so the desired invariance is imposed as a penalty rather than exhibited by the data. Shadow pairing changes the data instead: each dynamics is observed twice under independently resampled appearance, so invariance holds by construction and the objective merely has to read it off.

Invariance by Construction. Our supervision transplants a proven recipe from self-supervised learning: a representation is defined by the pairs it is trained on, since what a pair varies is discarded and what it preserves is kept. Contrastive learning pairs two augmentations of one image, so the representation drops crops and color while retaining visual identity [10]; CLIP pairs an image with its caption, so it retains the semantics that language can express [45]; masked and predictive video objectives extend the same logic to spatiotemporal structure [4, 53]. In each case the objective does not create the invariance; it reads off the invariance that the pairing already wrote into the data. Dynamics, however, resists pixel-space augmentation: a post-hoc transform of one video can hardly resample its appearance while preserving its motion frame by frame. The shadow pair moves the

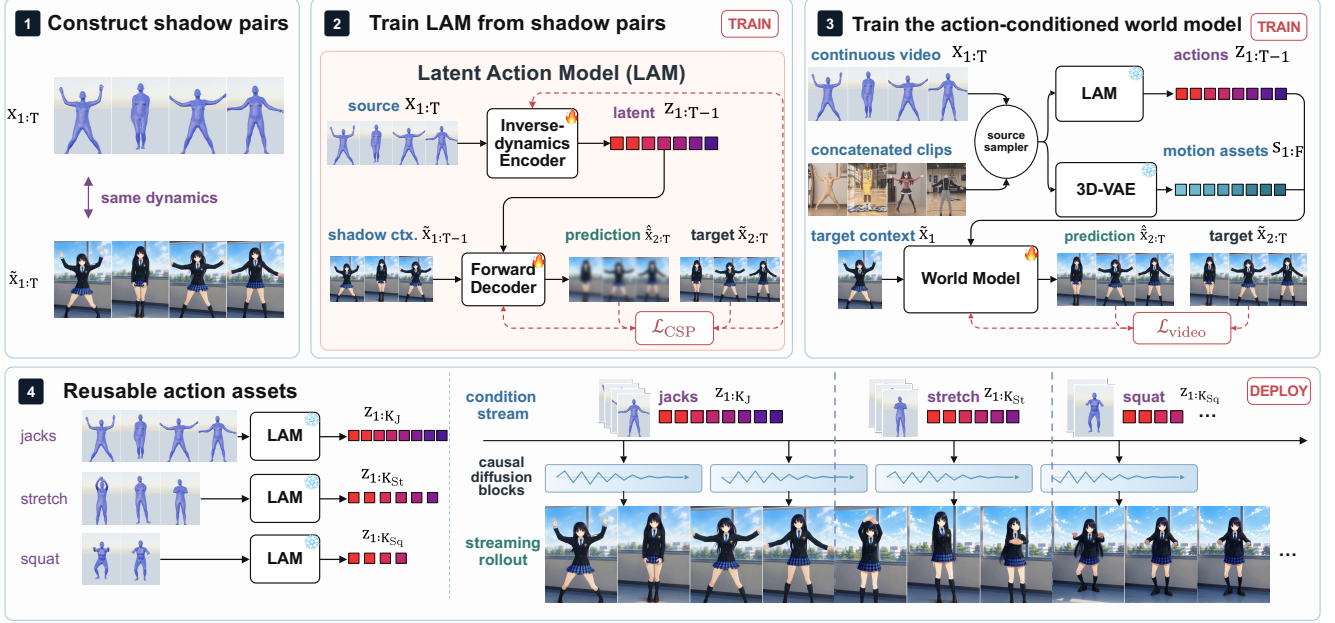


Figure 2. **Overview of ShadowDancer**, taking human motion as the running example; the pipeline is identical for camera, object, and every other dynamics family. (1) A shadow pair renders one dynamics d twice, $x=R(d, c)$ and $\tilde{x}=R(d, \tilde{c})$, with the remaining factors independently resampled. (2) Cross-shadow prediction trains the LAM: the encoder reads $z_{1:T-1}$ from the source x , and the decoder predicts the next shadow frame from $(\tilde{x}_t, z_t)$, so only what the pair shares survives in z . (3) With the LAM and 3D-VAE frozen, a video diffusion backbone is fine-tuned into a block-causal world model, conditioned on the action trajectory z and on source motion assets s , and trained on mixed continuous and concatenated sources. (4) At deployment, variable-length action assets (z, s) are stored, composed along the condition stream, and streamed through causal blocks for long-horizon rollout; a new action costs one pass through the frozen encoders.

augmentation into the renderer, replaying the trajectory and resampling everything else, which grounds cross-shadow prediction (Sec. 3.2) as a sound way to recover the preserved dynamics, invariant to everything resampled.

3. Method

Overview. Fig. 2 summarizes ShadowDancer. The method has three parts, in pipeline order: shadow pairs make dynamics identifiable in the data (Sec. 3.1); *cross-shadow prediction* distills it into the unified representation z (Sec. 3.2); and a video diffusion backbone, conditioned on z for control and on source assets for motion detail, becomes a block-causal interactive world model whose deployment we describe last (Sec. 3.3). The Shadow Library (Sec. 3.4) implements the pairing protocol at scale. Three objects recur throughout: the LAM extracts a latent trajectory z invariant to everything its pairing resamples (*appearance*, Sec. 3.1); a frozen 3D-VAE encodes the aligned source-detail stream s ; and together $a = (z, s)$ is the reusable *action asset* the world model consumes, while the *unified dynamics representation* always denotes z alone, unified in the sense of one encoder, one latent space, and one interface for every dynamics family.

3.1. The Shadow Formulation

We model a video $x = (x_1, \dots, x_T)$ as the output of a rendering process

$$x = R(d, c), \quad (1)$$

where $d = (d_1, \dots, d_T)$ is a time-varying *dynamics* trajectory, the factor that drives change between frames, and c collects the remaining generative factors, such as subject identity, scene, materials, and lighting, which we call *appearance*. In this notation the metaphor of our title is literal: a video is a *shadow* of its dynamics, one rendering of d under one particular c . Observing x alone, d and c are entangled: many different (d, c) pairs can explain the same pixels, which is the root cause of the shortcut and non-identifiability failures of latent action learning [32, 40, 60].

A **Shadow pair** comprises two independent renders of the same dynamics,

$$(x = R(d, c), \tilde{x} = R(d, \tilde{c})), \quad \tilde{c} \sim p(c), \quad (2)$$

where the two videos are frame-synchronized: frame t of x and frame t of $\tilde{x}$ share the identical d_t . Both videos are shadows of the same d , and each is therefore the Shadow of the other: the relation is symmetric, and “the Shadow of a video” refers to a fellow projection of its dynamics rather

than a projection of the video itself. Because $\tilde{c}$ is resampled independently of c , the only factor deliberately preserved across the two rendering processes is d ; source appearance cannot systematically predict target appearance.

The split between d and c in Eq. (1) is not fixed a priori; it is chosen by the pairing protocol. Replaying a body motion while resampling the camera assigns the camera to c ; replaying a camera trajectory while resampling the scene assigns it to d . This operational definition is what makes the paradigm any-action: designating a factor as controllable requires only constructing pairs that preserve it. Shadow pairing therefore differs from data augmentation: augmentation encourages invariance to nuisances chosen within a fixed task, whereas the pairing chooses the task variable itself, by deciding what is preserved. This choice cuts both ways: a factor the pairing does not resample is shared by both shadows and is therefore retained in the representation. Such retention is often wanted rather than tolerated, because many actions owe part of their identity to how they look; a fire spell stripped of its flames is a different action, so the protocol keeps the trait inside the action by simply not resampling it. In our combat families the acting instrument is likewise bound to its motion: the same weapon appears in both renders, its identity belongs to the action rather than to the appearance, and a stored action carries its weapon into the new world, which is exactly the behavior the action-fidelity evaluation of Sec. 4.3 demands. The representation is thus invariant to what a pairing resamples and faithful to what it preserves; decoupling a factor is a capability of the protocol, not an obligation. Accordingly, *appearance* in this paper always denotes the factors a pairing chooses to resample, and *dynamics* denotes what it preserves: motion at its core, together with any trait bound to that motion.

3.2. Cross-Shadow Prediction

Objective. Cross-shadow prediction turns the pairing constraint into a learning objective. We instantiate it with a latent action model (LAM): an inverse-dynamics encoder produces a variational posterior $q_\phi(z_t | x)$ over per-frame latent actions $z_t \in \mathbb{R}^{d_z}$ for $t = 1, \dots, T-1$, and a forward decoder p_θ performs next-frame prediction. Standard LAMs [18, 32] read and predict the same clip. Our decisive difference is the cross-shadow arrangement: the encoder reads the source video x , while the decoder predicts the other shadow, reconstructing $\tilde{x}_{t+1}$ from the target’s own previous frame $\tilde{x}_t$ and the source-derived action z_t . Training minimizes the β -VAE objective [2, 27]:

$$\mathcal{L}_{\theta, \phi}^{\text{CSP}} = \frac{1}{T-1} \sum_{t=1}^{T-1} \left(-\mathbb{E}_{q_\phi(z_t | x)} [\log p_\theta(\tilde{x}_{t+1} | \tilde{x}_t, z_t)] + \beta \text{KL}(q_\phi(z_t | x) \| p(z_t)) \right), \quad (3)$$

where $p(z_t)$ is a fixed prior $\mathcal{N}(0, I)$.

Why this works. Source-context information in z_t beyond what is needed to infer d cannot improve population prediction of $\tilde{x}$, whose appearance is supplied by its own context $\tilde{x}_t$ and independently resampled from c : given the pairing, appearance ceases to help the prediction, and the bottleneck, pressured by the KL term, has room only for what does help, the shared dynamics. Unlike feature-alignment regularizers [32], the predictive supervision itself supplies this cross-context constraint. With probability p we also let a video pair with itself ($\tilde{x} = x$), recovering the standard single-video LAM objective as a degenerate case; these self-pairs admit unpaired real video into the same training stream (Sec. 3.4) as an empirical realism augmentation, outside the guarantee below.

Formal guarantee. This preference for dynamics is not only an intuition; we prove it. Write D for the shared dynamics, X for the source video, (B, Y) for the target-side context and prediction target, and K_d for the law of Y given B under dynamics d . Whenever the pairing protocol delivers

$$\underbrace{D = h(X)}_{\text{observable}}, \quad \underbrace{X \perp\!\!\!\perp (B, Y) \mid D}_{\text{independent resampling}}, \quad \underbrace{K_d = K_{d'} \Rightarrow d = d'}_{\text{effect separation}}, \quad (4)$$

the minimal representation sufficient for predicting one shadow from the other is the shared dynamics itself, up to an invertible reparameterization (Theorem C.2); formal statements and proofs are in Supplementary Sec. C.

Factor-selective control. Dynamics itself is composite: the observer’s own motion (ego, *i.e.* the camera) and the motion of everything observed (scene) are superimposed in every video. Pair construction separates them within one latent space and specifies the readout: pairs preserving only the camera trajectory supervise a `cam` head, pairs preserving only the scene dynamics supervise a `dyn` head, and pairs preserving both supervise a `full` head; what the `dyn` factor contains is itself defined by each family’s pairing protocol, the arm trajectory alone in our robot manipulation pairs, the body motion in our human pairs. At inference the same reference video can therefore be read as a pure camera action, a pure scene action, or their joint; the mechanism is detailed in Supplementary Sec. B.1.

3.3. From Representation to World Model

The conditioning design follows two principles: (i) the action latent z carries the unified, transferable control, and (ii) the source stream s supplies the high-frequency motion detail that an invariance-constrained bottleneck is deliberately too small to carry.

Action-conditioned generation. We fine-tune a pretrained video diffusion transformer (DiT) [9, 44, 55] on Shadow Library clips with flow matching [39]: the frozen LAM’s actions are fused into the timestep embedding for AdaLN modulation and attended through a dedicated cross-attention

arm, while the source video is encoded by the 3D-VAE and injected as *motion assets*, through channel concatenation at the input convolution and as additional cross-attention context. Because source and target form a shadow pair, copying source appearance is not rewarded, so the generator is trained to read the source for motion detail rather than appearance; all injection routes are zero-initialized, with details in Supplementary Sec. B.2.

Block-causal conversion for interactive rollout. An interactive world model must generate forward-only over long horizons, while diffusion backbones are bidirectional and fixed-length. Following recent causal conversions of video diffusion [23, 30], we fine-tune the bidirectional model into a block-causal generator: latent frames are grouped into blocks, attention is causal across blocks, and denoising is conditioned on the clean history of previous blocks. At inference the model rolls out block by block with a key-value cache, consuming a stream of latent actions and emitting a stream of frames.

Actions as reusable assets. At inference, the model receives a world seed $\mathbf{I}_0$ and a demonstration, given as reference clips or stored action assets: the frozen encoder reads the demonstration into per-frame latent actions z , and the block-causal generator rolls out a video in which the seeded world re-enacts the demonstrated dynamics. Because control is a latent trajectory rather than a per-frame key state, an “action” in ShadowDancer is a *variable-length dynamics segment*, stored as an action asset: extract $a = (z_{1:K}, s)$ from a clip of any duration, store it, replay it in any world, and compose segments by concatenation along the rollout, so teaching the model a new action takes a single pass through the frozen encoders, with no vocabulary design, labels, or fine-tuning. Blocks group latent frames only for causal denoising, while the action stream stays per-frame, so a single action may span many blocks and its temporal detail is not tied to the block rate. At deployment, discrete commands retrieve and concatenate short stored assets rather than supplying one continuous demonstration; we therefore expose the generator to the same discontinuous conditioning during training, replacing the continuous source, with some probability, by a concatenation of *canonical action chunks* drawn from a fixed asset library, so a command at inference is a library lookup that matches the input statistics the generator was trained on.

3.4. The Shadow Library

The shadow protocol, which replays the dynamics and resamples everything else, admits many implementations, and we deliberately build a library rather than a single engine. Animation suites replay articulated human motion across characters, scenes, lighting, and cameras; open-world game environments replay trajectories across districts, weather, and time of day; robotic simulators replay manipulation across arms, objects, and tabletops; camera-trajectory sources re-

play the same path through different scenes. All sources emit the same artifact, frame-synchronized clip sets that share one dynamics, so a new action family integrates by contributing one more shadow script.

Real-world video, which lacks exact shadows, enters the very same stream as degenerate self-pairs ($\tilde{x} = x$): it contributes visual diversity and realism priors, while the invariance pressure is supplied by the synthetic pairs. Genuine shadow pairs supply the identifying signal; self-pairs broaden the visual support.

4. Experiments

4.1. Implementation Details

Following [18, 32], the LAM has latent dimension $d_z=32$ and is trained on the Shadow Library at half the world model’s spatial resolution ($H/2 \times W/2$) with $\beta=0.01$; real videos enter as degenerate self-pairs. The world model is built on the SkyReels-V2-1.3B I2V DiT backbone [9], fine-tuned with flow matching and converted to a block-causal generator. All experiments are conducted on NVIDIA H200 GPUs. Additional implementation details are provided in Supplementary Sec. D.

Datasets The Shadow Library spans articulated human motion (SMPL-X sequences re-rendered in Blender across characters, scenes, lighting, and cameras, together with paired renders driven by [8] that replay one motion under a resampled character and environment), robotic manipulation (ManiSkill [51]), first-person open-world games (GTA- and Cyberpunk-style urban scenes with first-person weapon fire and driving), third-person games (Unreal Engine scenes; Monster Hunter-style action), and camera trajectories, both around articulated subjects and through static scenes (DL3DV [38]); unpaired real video (OpenX robot corpora [42], Internet clips [33]) joins as self-pairs. Held-out shadow pairs from every source form the evaluation sets; the full composition of the Shadow Library is tabulated in Supplementary Sec. D.

4.2. Action Transfer Across Dynamics Families

Our first experiment tests the central claim in one protocol: a single model with one latent interface re-enacts every family of dynamics in a new environment. For every family in the Shadow Library (human motion, first-person combat, third-person action, robot manipulation, and camera control), we hold out shadow pairs, extract the action from video x , and generate from the first frame of its Shadow $\tilde{x}$, which realizes the identical, frame-synchronized dynamics. Four families are scored by frame-aligned reconstruction against $\tilde{x}$ (PSNR, LPIPS [64]); camera control is scored by trajectory error (ATE/RPE [48]) between the target trajectory and poses recovered from the generations with VGGT [56]. The baseline is Olaf-World [32], the state-of-the-art self-reconstruction

Table 1. **Quantitative comparison of action transfer across dynamics families.** Best in **bold**.

Method	Human motion		First-person combat		Third-person action		Camera control		Robot manipulation	
	PSNR $\uparrow$	LPIPS $\downarrow$	PSNR $\uparrow$	LPIPS $\downarrow$	PSNR $\uparrow$	LPIPS $\downarrow$	ATE $\downarrow$	RPE $\downarrow$	PSNR $\uparrow$	LPIPS $\downarrow$
Olaf-World [32]	18.2	0.288	13.0	0.532	12.6	0.506	0.072	0.021	14.0	0.478
ShadowDancer (ours)	22.4	0.184	17.0	0.354	17.4	0.309	0.005	0.003	22.6	0.116

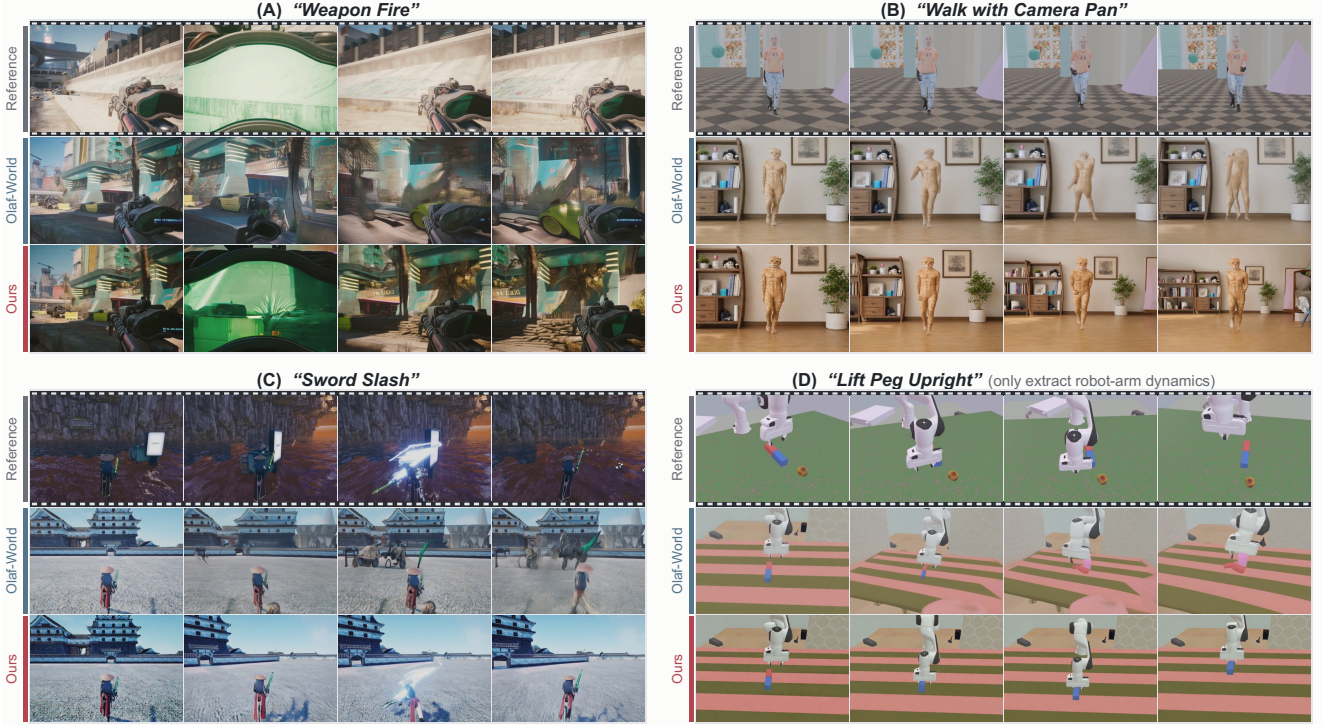


Figure 3. **Qualitative action transfer against Olaf-World.** Top: reference videos providing the action across four families; below: frame-synchronized generations in the new environment. Olaf-World warps subjects and leaks spurious motion; ShadowDancer re-enacts the reference dynamics faithfully.

latent-action model: we match the training data and world-model backbone, inject its latent through the same routes, and retain its original single-head, z -only design, which has no source-asset stream. Table 1 reports the results: ShadowDancer leads on every family by a wide margin. The gap is visible in Fig. 3: Olaf-World routinely produces warped, twisting subjects and spurious camera motion, because its entangled latent cannot cleanly extract the intended dynamics from the source, so what it transfers is the action mixed with whatever appearance happened to accompany it; with the shadow-trained latent, the same interface re-enacts the demonstrated dynamics faithfully in the new environment. Latent-level probes corroborate this before any generation is involved: linear probes on z show that cross-shadow training preserves action content while reducing leakage of the resampled scene (Supplementary Sec. D).

Table 2. **Long action rollout.** Blinded 2AFC favor rate of ShadowDancer over each baseline; $>50\%$ favors ours.

ShadowDancer (ours) over	Action Control	Action Fidelity	Long-horizon Consist.
Olaf-World [32]	94%	95%	94%
Yume-1.5 [41]	88%	86%	91%
LingBot-World 2.0 [19]	78%	83%	64%

4.3. Long Action Rollout

Our second experiment evaluates what the system is ultimately for: sustained, action-driven generation over long horizons. Discrete commands are mapped to canonical action assets by lookup and streamed into long rollouts that chain several commands per video, spanning navigation, aiming and firing, and skill casts across first- and third-person worlds, together with human-motion and robot-manipulation streams. Unlike the transfer setting, a chained command stream has no ground-truth realization, so control must be *judged* rather than measured: following the

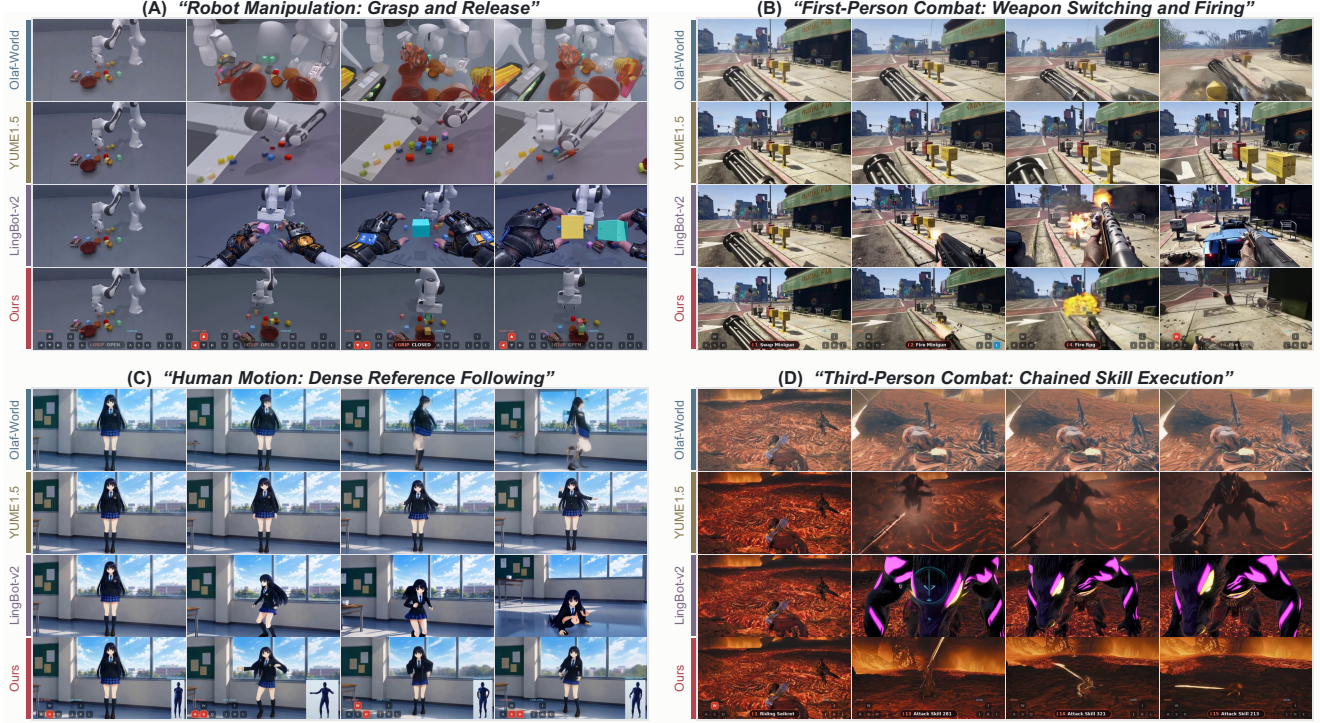


Figure 4. **Qualitative long action rollouts.** Each comparison starts from a shared first frame and follows the same command stream; rows are methods, columns are time steps along the rollout. Where the baselines drift or degrade, ShadowDancer stays aligned.

preference protocol of [65], each ShadowDancer rollout is compared head-to-head against a baseline’s in a blinded two-alternative forced choice (2AFC), judged by a VLM (Table 5) over all pairwise comparisons. The comparison runs along three action-centric axes: *action control* (is the commanded action performed at all), *action fidelity* (is the specific weapon/object/motion preserved), and *long-horizon consistency* (no drift or degradation as the rollout extends); visual fidelity is excluded as orthogonal to control. We compare against the state-of-the-art open interactive world models LingBot-World 2.0 [19], driven by a camera-pose trajectory with text-triggered events, and Yume-1.5 [41], driven by text and keystrokes, alongside Olaf-World’s latent interface on our backbone; each receives the same command stream through its native interface. The comparison is deliberately system-level: the interfaces receive different information by design, and how faithfully a user’s commands survive each interface is precisely what is being measured. Table 2 reports the favor rates, and Fig. 4 shows representative rollouts; where the baselines drift, hallucinate, or degrade, ShadowDancer stays aligned to the commanded action.

4.4. Ablation Studies

Effectiveness of shadow pairing and source assets As detailed in Table 3, we study the impact of the proposed components in four settings: (a) the prior z -only latent inter-

Table 3. **Impact of key components.** From the z -only latent (Olaf-World), adding *pairing*, then *assets* (=ShadowDancer).

Setting	z	Pairing	Assets	1st-person combat		3rd-person action	
				PSNR $\uparrow$	LPIPS $\downarrow$	PSNR $\uparrow$	LPIPS $\downarrow$
(a) Olaf-World (z only)	✓	×	×	12.70	0.524	12.90	0.532
(b) + pairing	✓	✓	×	14.92	0.445	15.56	0.412
(c) assets, no action	×	×	✓	14.82	0.442	16.41	0.376
(d) ShadowDancer (ours)	✓	✓	✓	16.00	0.397	17.13	0.351

face, a separately trained model following the Olaf-World recipe [32] on our data (its single-head latent is architecturally distinct, so it cannot be recovered by ablation); (b) the same interface with the LAM trained on shadow pairs instead of self-reconstruction; (c) source assets without the action z ; and (d) the full model, where (b)–(d) probe the same ShadowDancer checkpoint with the corresponding routes ablated at inference. The results demonstrate that both components matter, with pairing as the foundation: once the latent is constrained to carry dynamics alone, the same interface becomes transferable ((a) vs. (b)). Assets are a strong carrier in their own right, since the source video itself contains the action and its spatial detail, and (c) accordingly matches or surpasses the bottlenecked z alone; but without z the model must infer the intended action implicitly from the raw source, whereas the explicit, disentangled z tells it what to extract, so the combination (d) is uniformly best.

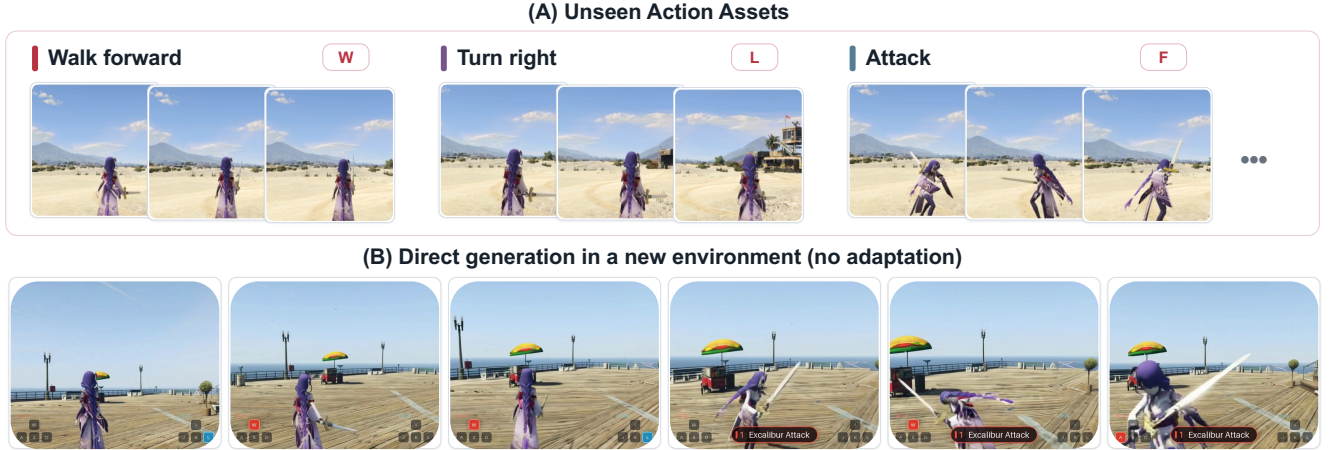


Figure 5. **Transferring unseen actions:** assets recorded from a modded, unseen character (A) drive generation in a new environment (B).

4.5. Transferring Unseen Actions

A latent interface is only as general as its behavior on dynamics it has never seen, so we test transfer on actions introduced through game mods after training: an unseen player character and a two-handed sword whose attack is a swing rather than a shot; neither asset appears in the Shadow Library, nor does the motion either one produces. We record walking, turning, and sword attacks from the modded character as action assets. The source clips are recorded on one map and the extracted actions replayed on held-out maps, so transfer to a new environment is built into the protocol; Fig. 5 shows the results. The modded dynamics survive this replay: this character’s gait and this sword’s swing, motion the model never observed, are reproduced on a held-out map, indicating that the encoder reads dynamics rather than memorized action categories. This generalization is where the scaling potential of the paradigm lies: a behavior demonstrated after training joins the action library at the cost of one pass through the frozen encoders, so coverage grows with new demonstrations rather than with new training runs.

5. Conclusion

ShadowDancer enables any-action, frame-level control of interactive video world models by demonstration. Through shadow pairs and cross-shadow prediction, it learns a unified dynamics representation that discards, by construction, whatever the pairing resamples and keeps whatever it preserves, so users can extract the dynamics of any clip as a reusable action asset and replay it in new environments without labels, estimators, or fine-tuning. Our experiments demonstrate clear improvements over latent-action and interactive world model baselines in both action transfer and long action rollout, across dynamics families ranging from human motion and camera control to gameplay and robot manipulation. ShadowDancer opens new possibilities for interactive world models, which can now be driven

by showing rather than telling, from entertainment to simulation-based training of embodied agents. We discuss limitations and future directions in Supplementary Sec. A.

References

- [1] Niket Agarwal, Arslan Ali, Maciej Bala, Yogesh Balaji, Erik Barker, Tiffany Cai, Prithvijit Chattopadhyay, Yongxin Chen, Yin Cui, Yifan Ding, et al. Cosmos world foundation model platform for physical ai. *arXiv preprint arXiv:2501.03575*, 2025. 3
- [2] Alexander A. Alemi, Ian Fischer, Joshua V. Dillon, and Kevin Murphy. Deep variational information bottleneck. In *ICLR*, 2017. 5
- [3] Eloi Alonso, Adam Jelley, Vincent Micheli, Anssi Kanervisto, Amos J Storkey, Tim Pearce, and François Fleuret. Diffusion for world modeling: Visual details matter in atari. *NeurIPS*, 2024. 3
- [4] Mido Assran, Adrien Bardes, David Fan, Quentin Garrido, Russell Howes, Matthew Muckley, Ammar Rizvi, Claire Roberts, Koustuv Sinha, Artem Zhohus, et al. V-JEPA 2: Self-supervised video models enable understanding, prediction and planning. *arXiv preprint arXiv:2506.09985*, 2025. 3
- [5] Tim Brooks, Bill Peebles, Connor Holmes, Will DePue, Yufei Guo, Li Jing, David Schnurr, Joe Taylor, Troy Luhman, Eric Luhman, Clarence Ng, Ricky Wang, and Aditya Ramesh. Video generation models as world simulators. 2024. 2
- [6] Jake Bruce, Michael D Dennis, Ashley Edwards, Jack Parker-Holder, Yuge Shi, Edward Hughes, Matthew Lai, Aditi Mavalankar, Richie Steigerwald, Chris Apps, et al. Genie: Generative interactive environments. In *ICML*, 2024. 2, 3
- [7] Qingwen Bu, Yanting Yang, Jisong Cai, Shenyuan Gao, Guanghui Ren, Maoqing Yao, Ping Luo, and Hongyang Li. UniVLA: Learning to act anywhere with task-centric latent actions. *arXiv preprint arXiv:2505.06111*, 2025. 3
- [8] Chenjie Cao, Jingkai Zhou, Shikai Li, Jingyun Liang, Chao-hui Yu, Fan Wang, Xiangyang Xue, and Yanwei Fu. Unifying

- precisely 3D-enhanced camera and human motion controls for video generation. In *SIGGRAPH Asia*, 2025. 2, 3, 6
- [9] Guibin Chen, Dixuan Lin, Jiangping Yang, Chunze Lin, Junchen Zhu, Mingyuan Fan, Hao Zhang, Sheng Chen, Zheng Chen, Chengcheng Ma, et al. SkyReels-V2: Infinite-length film generative model. *arXiv preprint arXiv:2504.13074*, 2025. 2, 3, 5, 6
- [10] Ting Chen, Simon Kornblith, Mohammad Norouzi, and Geoffrey Hinton. A simple framework for contrastive learning of visual representations. In *ICML*, 2020. 3
- [11] Xiaoyu Chen, Junliang Guo, Tianyu He, Chuheng Zhang, Pushi Zhang, Derek Cathera Yang, Li Zhao, and Jiang Bian. IGOR: Image-GOal Representations are the Atomic Control Units for Foundation Models in Embodied AI. *arXiv preprint arXiv:2411.00785*, 2024. 2, 3
- [12] Xiaoyu Chen, Hangxing Wei, Pushi Zhang, Chuheng Zhang, Kaixin Wang, Yanjiang Guo, Rushuai Yang, Yucen Wang, Xinquan Xiao, Li Zhao, et al. villa-X: enhancing latent action modeling in vision-language-action models. *arXiv preprint arXiv:2507.23682*, 2025.
- [13] Yi Chen, Yuying Ge, Weiliang Tang, Yizhuo Li, Yixiao Ge, Mingyu Ding, Ying Shan, and Xihui Liu. Moto: Latent motion token as the bridging language for learning robot manipulation from videos. In *ICCV*, 2025. 3
- [14] Decart, Julian Quevedo, Quinn McIntyre, Spruce Campbell, Xinlei Chen, and Robert Wachen. Oasis: A universe in a transformer. 2024. 2, 3
- [15] Ashley Edwards, Himanshu Sahni, Yannick Schroecker, and Charles Isbell. Imitating latent policies from observation. In *ICML*, 2019. 2, 3
- [16] Zhixue Fang, Xu He, Songlin Tang, Haoxian Zhang, Qingfeng Li, Xiaoqiang Liu, Pengfei Wan, and Kun Gai. 3D-aware implicit motion control for view-adaptive human video generation. *arXiv preprint arXiv:2602.03796*, 2026. 2, 3
- [17] Shenyuan Gao, Jiazhi Yang, Li Chen, Kashyap Chitta, Yihang Qiu, Andreas Geiger, Jun Zhang, and Hongyang Li. Vista: A generalizable driving world model with high fidelity and versatile controllability. *NeurIPS*, 2024. 3
- [18] Shenyuan Gao, Siyuan Zhou, Yilun Du, Jun Zhang, and Chuang Gan. AdaWorld: Learning adaptable world models with latent actions. In *ICML*, 2025. 2, 3, 5, 6
- [19] Zelin Gao, Qiuyu Wang, Jiapeng Zhu, Jingye Chen, Zichen Liu, Qingyan Bai, Jiahao Wang, Yufeng Yuan, Hanlin Wang, Yichong Lu, Ka Leong Cheng, Haojie Zhang, Jian Gao, Tianrui Feng, Yuzheng Liu, Yao Yao, Yinghao Xu, Xing Zhu, Yujun Shen, and Hao Ouyang. Infinite worlds with versatile interactions. *arXiv preprint arXiv:2607.07534*, 2026. 2, 3, 7, 8
- [20] Quentin Garrido, Tushar Nagarajan, Basile Terver, Nicolas Ballas, Yann LeCun, and Michael Rabbat. Learning latent action world models in the wild. *arXiv preprint arXiv:2601.05230*, 2026. 2, 3
- [21] Google DeepMind. Veo. Model page. 1
- [22] Luigi Gresele, Paul K. Rubenstein, Arash Mehrjou, Francesco Locatello, and Bernhard Schölkopf. The incomplete rosetta stone problem: Identifiability results for multi-view nonlinear ICA. In *Proceedings of the 35th Conference on Uncertainty in Artificial Intelligence*, pages 217–227, 2020. 2
- [23] Yuchao Gu, Weijia Mao, and Mike Zheng Shou. Long-context autoregressive video modeling with next-frame prediction. *arXiv preprint arXiv:2503.19325*, 2025. 3, 6
- [24] David Ha and Jürgen Schmidhuber. World models. *arXiv preprint arXiv:1803.10122*, 2018. 3
- [25] Danijar Hafner, Jurgis Pasukonis, Jimmy Ba, and Timothy Lillicrap. Mastering diverse domains through world models. *arXiv preprint arXiv:2301.04104*, 2023. 3
- [26] Xianglong He, Chunli Peng, Zexiang Liu, Boyang Wang, Yifan Zhang, Qi Cui, Fei Kang, Biao Jiang, Mengyin An, Yangyang Ren, Baixin Xu, Hao-Xiang Guo, Kaixiong Gong, Cyrus Wu, Wei Li, Xuchen Song, Yang Liu, Eric Li, and Yahui Zhou. Matrix-Game 2.0: An open-source, real-time, and streaming interactive world model. *arXiv preprint arXiv:2508.13009*, 2025. 2, 3
- [27] Irina Higgins, Loic Matthey, Arka Pal, Christopher Burgess, Xavier Glorot, Matthew Botvinick, Shakir Mohamed, and Alexander Lerchner. beta-VAE: Learning basic visual concepts with a constrained variational framework. In *ICLR*, 2017. 5
- [28] Yicong Hong, Yiqun Mei, Chongjian Ge, Yiran Xu, Yang Zhou, Sai Bi, Yannick Hold-Geoffroy, Mike Roberts, Matthew Fisher, Eli Shechtman, et al. RELIC: Interactive video world model with long-horizon memory. *arXiv preprint arXiv:2512.04040*, 2025. 3
- [29] Kurt Hornik. Approximation capabilities of multilayer feed-forward networks. *Neural Networks*, 4(2):251–257, 1991. 4
- [30] Xun Huang, Zhengqi Li, Guande He, Mingyuan Zhou, and Eli Shechtman. Self forcing: Bridging the train-test gap in autoregressive video diffusion. *arXiv preprint arXiv:2506.08009*, 2025. 3, 6
- [31] Joel Jang, Seonghyeon Ye, Zongyu Lin, Jiannan Xiang, Johan Bjorck, Yu Fang, Fengyuan Hu, Spencer Huang, Kaushil Kundalia, Yen-Chen Lin, et al. DreamGen: Unlocking generalization in robot learning through video world models. *arXiv preprint arXiv:2505.12705*, 2025. 3
- [32] Yuxin Jiang, Yuchao Gu, Ivor W. Tsang, and Mike Zheng Shou. Olaf-world: Orienting latent actions for video world modeling. *arXiv preprint arXiv:2602.10104*, 2026. 2, 3, 4, 5, 6, 7, 8
- [33] Xuan Ju, Yiming Gao, Zhaoyang Zhang, Ziyang Yuan, Xintao Wang, Ailing Zeng, Yu Xiong, Qiang Xu, and Ying Shan. Miradata: A large-scale video dataset with long durations and structured captions. *NeurIPS*, 2024. 6
- [34] Ilyes Khemakhem, Diederik Kingma, Ricardo Monti, and Aapo Hyvarinen. Variational autoencoders and nonlinear ica: A unifying framework. In *AISTATS*, 2020. 2, 3
- [35] Hanjung Kim, Jaehyun Kang, Hyolim Kang, Meedeum Cho, Seon Joo Kim, and Youngwoon Lee. UniSkill: Imitating human videos via cross-embodiment skill representations. *arXiv preprint arXiv:2505.08787*, 2025. 3
- [36] Weijie Kong, Qi Tian, Zijian Zhang, Rox Min, Zuozhuo Dai, Jin Zhou, Jiangfeng Xiong, Xin Li, Bo Wu, Jianwei Zhang, et al. HunyuanVideo: A systematic framework for large video generative models. *arXiv preprint arXiv:2412.03603*, 2024. 2
- [37] Yao-Chih Lee, Zhoutong Zhang, Jiahui Huang, Jui-Hsien Wang, Joon-Young Lee, Jia-Bin Huang, Eli Shechtman, and

- Zhengqi Li. Generative video motion editing with 3D point tracks. In *CVPR*, 2026. 2, 3
- [38] Lu Ling, Yichen Sheng, Zhi Tu, Wentian Zhao, Cheng Xin, Kun Wan, Lantao Yu, Qianyu Guo, Zixun Yu, Yawen Lu, et al. D13dv-10k: A large-scale scene dataset for deep learning-based 3d vision. *arXiv preprint arXiv:2312.16256*, 2023. 6
- [39] Xingchao Liu, Chengyue Gong, and qiang liu. Flow straight and fast: Learning to generate and transfer data with rectified flow. In *ICLR*, 2023. 5
- [40] Francesco Locatello, Stefan Bauer, Mario Lucic, Gunnar Raetsch, Sylvain Gelly, Bernhard Schölkopf, and Olivier Bachem. Challenging common assumptions in the unsupervised learning of disentangled representations. In *ICML*, 2019. 2, 3, 4
- [41] Xiaofeng Mao, Zhen Li, Chuanhao Li, Xiaojie Xu, Kaining Ying, Tong He, Jiangmiao Pang, Yu Qiao, and Kaipeng Zhang. Yume-1.5: A text-controlled interactive world generation model. *arXiv preprint arXiv:2512.22096*, 2025. 2, 3, 7, 8
- [42] Abby O’Neill, Abdul Rehman, Abhiram Maddukuri, Abhishek Gupta, Abhishek Padalkar, Abraham Lee, Acorn Pooley, Agrim Gupta, et al. Open X-Embodiment: Robotic learning datasets and RT-X models. In *IEEE International Conference on Robotics and Automation (ICRA)*, pages 6892–6903, 2024. 6
- [43] Jack Parker-Holder, Shlomi Fruchter, et al. Genie 3: A new frontier for world models. <https://deepmind.google/models/genie/>. Blog post. 2, 3
- [44] William Peebles and Saining Xie. Scalable diffusion models with transformers. In *ICCV*, 2023. 5
- [45] Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, et al. Learning transferable visual models from natural language supervision. In *ICML*, 2021. 3
- [46] Oleh Rybkin, Karl Pertsch, Andrew Jaegle, Konstantinos G. Derpanis, and Kostas Daniilidis. Learning what you can do before doing anything. In *ICLR*, 2019. 2, 3
- [47] Joonghyuk Shin, Zhengqi Li, Richard Zhang, Jun-Yan Zhu, Jaesik Park, Eli Shechtman, and Xun Huang. MotionStream: Real-time video generation with interactive motion controls. In *ICLR*, 2026. 2, 3
- [48] Jürgen Sturm, Nikolas Engelhard, Felix Endres, Wolfram Burgard, and Daniel Cremers. A benchmark for the evaluation of RGB-D SLAM systems. In *IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, pages 573–580, 2012. 6
- [49] Wenqiang Sun, Haiyu Zhang, Haoyuan Wang, Junta Wu, Zehan Wang, Zhenwei Wang, Yunhong Wang, Jun Zhang, Tengfei Wang, and Chunchao Guo. WorldPlay: towards long-term geometric consistency for real-time interactive world modeling. *arXiv preprint arXiv:2512.14614*, 2025. 2, 3
- [50] Junshu Tang, Jiacheng Liu, Jiaqi Li, Longhuang Wu, Haoyu Yang, Penghao Zhao, Siruis Gong, Xiang Yuan, Shuai Shao, and Qinglin Lu. Hunyuan-GameCraft-2: Instruction-following interactive game world model. *arXiv preprint arXiv:2511.23429*, 2025. 3
- [51] Stone Tao, Fanbo Xiang, Arth Shukla, Yuzhe Qin, Xander Hinrichsen, Xiaodi Yuan, Chen Bao, Xinsong Lin, Yulin Liu, Tse-kai Chan, Yuan Gao, Xuanlin Li, Tongzhou Mu, Nan Xiao, Arnav Gurha, Viswesh Nagaswamy Rajesh, Yong Woo Choi, Yen-Ru Chen, Zhiao Huang, Roberto Calandra, Rui Chen, Shan Luo, and Hao Su. Maniskill3: Gpu parallelized robotics simulation and rendering for generalizable embodied ai. *arXiv preprint arXiv:2410.00425*, 2024. 6
- [52] Robbyant Team, Zelin Gao, Qiuyu Wang, Yanhong Zeng, Jiapeng Zhu, Ka Leong Cheng, Yixuan Li, Hanlin Wang, Yinghao Xu, Shuailei Ma, et al. Advancing open-source world models. *arXiv preprint arXiv:2601.20540*, 2026. 2, 3
- [53] Zhan Tong, Yibing Song, Jue Wang, and Limin Wang. VideoMAE: Masked autoencoders are data-efficient learners for self-supervised video pre-training. *NeurIPS*, 2022. 3
- [54] Julius von Kügelgen, Yash Sharma, Luigi Gresele, Wieland Brendel, Bernhard Schölkopf, Michel Besserve, and Francesco Locatello. Self-supervised learning with data augmentations provably isolates content from style. In *Advances in Neural Information Processing Systems*, pages 16451–16467, 2021. 2
- [55] Team Wan, Ang Wang, Baole Ai, Bin Wen, Chaojie Mao, Chen-Wei Xie, Di Chen, Feiwu Yu, Haiming Zhao, Jianxiao Yang, et al. Wan: Open and advanced large-scale video generative models. *arXiv preprint arXiv:2503.20314*, 2025. 2, 3, 5, 1
- [56] Jianyuan Wang, Minghao Chen, Nikita Karaev, Andrea Vedaldi, Christian Rupprecht, and David Novotny. VGGT: Visual geometry grounded transformer. In *CVPR*, 2025. 6
- [57] Yucen Wang, Fengming Zhang, De-Chuan Zhan, Li Zhao, Kaixin Wang, and Jiang Bian. Co-Evolving latent action world models. *arXiv preprint arXiv:2510.26433*, 2025. 3
- [58] Halbert White. Connectionist nonparametric regression: Multilayer feedforward networks can learn arbitrary mappings. *Neural Networks*, 3(5):535–549, 1990. 4
- [59] Zeqi Xiao, Yushi Lan, Yifan Zhou, Wenqi Ouyang, Shuai Yang, Yanhong Zeng, and Xingang Pan. WorldMem: Long-term consistent world simulation with memory. *arXiv preprint arXiv:2504.12369*, 2025. 3
- [60] Jiange Yang, Yansong Shi, Haoyi Zhu, Mingyu Liu, Kaijing Ma, Yating Wang, Gangshan Wu, Tong He, and Limin Wang. CoMo: Learning continuous latent motion from internet videos for scalable robot learning. *arXiv preprint arXiv:2505.17006*, 2025. 2, 3, 4
- [61] Seonghyeon Ye, Joel Jang, Byeongguk Jeon, Sejeun Joo, Jianwei Yang, Baolin Peng, Ajay Mandlekar, Reuben Tan, Yu-Wei Chao, Bill Yuchen Lin, et al. Latent Action Pretraining from Videos. In *ICLR*, 2025. 2, 3
- [62] Yixuan Ye, Xuanyu Lu, Yuxin Jiang, Yuchao Gu, Rui Zhao, Qiwei Liang, Jiachun Pan, Fengda Zhang, Weijia Wu, and Alex Jinpeng Wang. MIND: Benchmarking memory consistency and action control in world models. *arXiv preprint arXiv:2602.08025*, 2026. 3
- [63] Jiwen Yu, Yiran Qin, Xintao Wang, Pengfei Wan, Di Zhang, and Xihui Liu. GameFactory: Creating new games with generative interactive videos. *arXiv preprint arXiv:2501.08325*, 2025. 2, 3, 1

- [64] Richard Zhang, Phillip Isola, Alexei A Efros, Eli Shechtman, and Oliver Wang. The unreasonable effectiveness of deep features as a perceptual metric. In *CVPR*, 2018. [6](#)
- [65] Shangwen Zhu, Qianyu Peng, Zhao Pu, Zhilei Shu, Xiangrui Ke, Zhaohu Xing, Zizhao Tong, Zeqing Wang, Xinyu Cui, Zian Zheng, Huangji Wang, Jian Zhao, Yeying Jin, Fan Cheng, and Ruili Feng. Incantation: Natural language as the action interface for multi-entity video world models. *arXiv preprint arXiv:2605.18601*, 2026. [2](#), [3](#), [8](#)



ShadowDancer: Teaching Video World Models Any Action by Learning Unified Dynamics Representations from a Video and Its Shadow

Supplementary Material

A. Limitations and Future Directions

Limitation ShadowDancer depends on two preconditions. First, action assets must be prepared in advance: at deployment the model consumes reference clips or stored assets, so a dynamics for which no demonstration exists must first be captured or authored. Second, shadow pairs are easy to construct in games and simulators, which replay the same dynamics under resampled appearance on demand, but are much harder to collect in the real world, where such re-runs are rarely possible; real video therefore enters training only as self-pairs, contributing visual realism rather than the identifying signal.

Future directions Both preconditions point to generative remedies. First, asset generation: because the encoder is appearance-invariant by construction, a demonstration does not have to be filmed. Off-the-shelf video generation [21, 55] is not sufficient by itself, however: a usable asset set consists of many clips of the same subject that stay strictly consistent in appearance and starting state and differ only in the performed action, whereas current models generate each video independently. Designing an *asset-generation* video model with this cross-clip consistency would let the action library scale with the generative ecosystem rather than with engine coverage. Second, shadows for real video: since the pairing is a protocol rather than an engine, video editing models that change a clip’s appearance while keeping its motion could manufacture shadows of real footage, extending the identifying supervision from synthetic worlds to the real domain.

B. More Model Details

B.1. Factor-Selective Readout of the Unified Representation

The main text (Sec. 3.2) states that ego and scene dynamics are separated within one latent space; this section gives the mechanism. The inverse-dynamics encoder is a spatiotemporal transformer with causal temporal attention: each per-frame action z_t is inferred from the frames up to the transition it describes, with no access to later frames, so clips of any length are encoded in a single pass. We use *three prompt slots*: three learnable tokens (`cam`, `dyn`, and `full`) are prepended to every frame’s patch sequence, attend jointly through the shared encoder, and each owns a variational readout head. The pairing protocol routes supervision among them: pairs preserving only the camera trajectory (scene and subject resampled or frozen) supervise the `cam` head; pairs

Table 4. **Pairing taxonomy.** Per dynamics family: what the shadow pairs preserve (this becomes the action in z), what they independently resample, and which readout head they supervise.

Family	Preserved (= action)	Resampled	Head
Human motion	body trajectory	character, scene, lighting, camera	<code>dyn</code>
Camera control	camera trajectory	scene, subject	<code>cam</code>
Robot manipulation	arm trajectory	camera, scene	<code>dyn</code>
Open-world games	camera + scene trajectory, bound weapon	district, weather, time of day	<code>full</code>
Paired character renders	body + camera trajectory	character, environment	<code>full</code>

preserving only scene motion (camera resampled) supervise the `dyn` head; pairs preserving both supervise the `full` head. A per-sample mask $(m_{\text{cam}}, m_{\text{dyn}}) \in \{0, 1\}^2$ selects one head, so each head is supervised only by the factor its pairs preserve. What constitutes the `dyn` factor is not fixed a priori but defined by each family’s pairing protocol: our robot manipulation pairs preserve the arm trajectory alone, so for this family `dyn` reads out arm dynamics and nothing else in the scene; our human pairs preserve the body motion. This is the protocol nature of the Shadow Library: designating a factor as `dyn` requires only constructing pairs that preserve it. Table 4 summarizes, per family, what the pairs preserve, what they resample, and which head they supervise.

All heads share one encoder and one latent space $\mathcal{Z}$, so downstream consumption is identical; yet at inference they are independent dials: the same reference video yields a pure camera action, a pure scene action, or their joint, and a body motion, a camera move, and an object interaction are simply different trajectories z .

B.2. Conditioning Architecture

This section details the conditioning routes summarized in Sec. 3.3 of the main text. The backbone operates on latents of a causal 3D video VAE with temporal compression $r=4$; we therefore group the per-frame actions by latent frame, $\mathcal{Z}_f = (z_{r(f-1)+1}, \dots, z_{rf})$, following [63]. The modulation route fuses the group-averaged action $\bar{z}_f$ into the per-latent-frame timestep embedding,

$$e_f = \psi(\tau) + \alpha_z W_z \bar{z}_f, \quad (\text{B.1})$$

where $\psi(\tau)$ embeds the diffusion timestep and e_f drives the adaptive layer norm of every DiT block, providing global, low-frequency control. The cross-attention route preserves all r action tokens: the spatial hidden states h_f of latent frame f attend to a per-frame context C_f through a dedicated arm,

$$h_f \leftarrow h_f + W_o \text{Attn}(h_f, C_f), \quad (\text{B.2})$$

with $\mathcal{C}_f = [W_a \mathcal{Z}_f]$ initially holding the projected action tokens. The gates α_z and W_o are zero-initialized, so at step 0 the conditioning is an exact no-op.

The bottleneck $d_z=32$ is a feature rather than a limitation: invariance demands that z be too small to smuggle appearance, but it is then also too small to carry the high-frequency content of a motion, such as the precise articulation of limbs and the exact timing of contacts. We therefore feed the source video itself to the generator as assets. Let $s = \mathcal{E}_{\text{vae}}(x^{\text{src}})$ be the source’s 3D-VAE latents. On the input side, s is channel-concatenated with the noisy latents u and a binary availability mask m , then embedded by the input 3D convolution:

$$h^{(0)} = \text{Conv3D}([u; s; m]); \quad (\text{B.3})$$

on the context side, each latent frame s_f is patchified into tokens that join the cross-attention context of Eq. (B.2),

$$\mathcal{C}_f = [W_a \mathcal{Z}_f; W_s \text{patchify}(s_f)], \quad (\text{B.4})$$

alongside a per-frame semantic embedding of the source in a third, parallel arm.

C. Formal Guarantee for Shadow Learning

This section makes precise the sense in which a shadow pair makes an arbitrary dynamics family learnable. An unconditional statement is impossible: if two dynamics produce exactly the same observations, no algorithm can distinguish them; if a finite-dimensional code is too small for the intrinsic dynamics dimension, no continuous encoder can store them. We therefore state the necessary observability, separation, and capacity conditions explicitly. The result is agnostic to the renderer and to whether the dynamics describe a human, camera, robot, object, or an entire multi-action trajectory.

Our argument is related in spirit to results showing that paired views can isolate invariant content from independently changing style [54] and can restore identifiability in multi-view nonlinear models [22]. Here we do not assume independent latent components or attempt to recover a privileged coordinate system. We prove the task-specific statement needed by cross-shadow prediction: the coarsest statistic sufficient for predicting one shadow from the other is exactly the shared dynamics, up to an invertible reparameterization.

C.1. Setup and Assumptions

Let D denote the dynamics to be represented. It may be one transition d_t , a variable-length segment $d_{1:T}$, or a full rollout. Let X be a source video, let B collect the target-side information already available to the decoder (for example, the previous target frame or clean block history), and let Y be the target next frame or block. All variables take values

in standard Borel spaces, so the conditional distributions below exist. Define the target transition kernel

$$K_d(b) := \mathbb{P}(Y \in \cdot \mid B = b, D = d). \quad (\text{C.1})$$

Three explicit conditions capture the shadow construction.

Assumption A1 (Independent context resampling). *The source is conditionally independent of the target rendering once the shared dynamics is fixed:*

$$X \perp\!\!\!\perp (B, Y) \mid D. \quad (\text{C.2})$$

This is the probabilistic statement of “replay the dynamics, resample everything else.” It allows both views to depend arbitrarily on D .

Assumption A2 (Source observability). *There is a measurable map h such that*

$$D = h(X) \quad \text{almost surely.} \quad (\text{C.3})$$

Thus the requested dynamics is actually visible in the source clip. This is necessary for any deterministic encoder of one video to recover it.

Assumption A3 (Target overlap and separation). *There is a probability measure ν on the target-context space such that $\mathbb{P}(B \in \cdot \mid D = d)$ is equivalent to ν for almost every d ; that is, every dynamics is rendered over the same target contexts up to null sets. Moreover, distinct dynamics induce distinct transition kernels:*

$$K_d = K_{d'} \implies d = d'. \quad (\text{C.4})$$

Here equality means equality ν -almost everywhere. We also assume that $d \mapsto [K_d]_\nu$ is a measurable map into a standard Borel coding of these kernels. Equivalently, every pair $d \neq d'$ differs in its effect in a set of target contexts of nonzero probability. Separation is necessary. The stated mutual absolute continuity is a transparent sufficient coverage condition that prevents the decoder-side context B from perfectly revealing D .

Definition C.1 (Cross-shadow representation). A representation $Z = s(X)$ is *cross-shadow sufficient* when

$$Y \perp\!\!\!\perp X \mid (B, Z), \quad (\text{C.5})$$

and is called *minimal* when every other cross-shadow sufficient statistic determines it. Minimality does not impose a particular coordinate system; two minimal statistics may differ by a bijection.

C.2. Identification Theorem

Theorem C.2 (Shadow identification). *Under Assumptions A1–A3, define the conditional-law statistic*

$$\Gamma(x) := [K_{h(x)}]_\nu. \quad (\text{C.6})$$

Then: (i) $\Gamma(X)$ is cross-shadow sufficient and minimal; (ii) it is invariant to every source-context factor resampled by the shadow construction; and (iii) every cross-shadow sufficient representation $Z = s(X)$ determines D . Consequently, every minimal sufficient Z and D determine one another and are equivalent up to a one-to-one reparameterization.

Proof. Because $D = h(X)$ and $X \perp\!\!\!\perp (B, Y) \mid D$, for almost every x and target context b ,

$$\mathbb{P}(Y \in \cdot \mid B = b, X = x) \quad (\text{C.7})$$

$$= \mathbb{P}(Y \in \cdot \mid B = b, D = h(x)) \quad (\text{C.8})$$

$$= K_{h(x)}(b). \quad (\text{C.9})$$

The conditional law of Y given (B, X) therefore depends on X only through $\Gamma(X)$, which proves sufficiency. The same identity shows invariance: source videos with the same D have the same Γ regardless of source context.

It remains to prove minimality without allowing the target-side variable B to do the identification. Let $Z = s(X)$ be any cross-shadow sufficient statistic and write $L_z(b) = \mathbb{P}(Y \in \cdot \mid B = b, Z = z)$. Sufficiency and Eq. (C.9) give

$$K_D(B) = L_Z(B) \quad \text{almost surely.} \quad (\text{C.10})$$

Because Z is a function of X , Assumption A1 implies $B \perp\!\!\!\perp Z \mid D$. Therefore, for almost every pair (d, z) in the support of (D, Z) , Eq. (C.10) holds under $\mathbb{P}(B \in \cdot \mid D = d)$. Equivalence of this measure to ν in Assumption A3 upgrades this to the kernel-class identity

$$[K_D]_\nu = [L_Z]_\nu \quad \text{almost surely.} \quad (\text{C.11})$$

The Borel injection $d \mapsto [K_d]_\nu$ has a measurable inverse on its image, so

$$D = (d \mapsto [K_d]_\nu)^{-1}([L_Z]_\nu) =: r(Z) \quad \text{almost surely.} \quad (\text{C.12})$$

Consequently, $\Gamma(X) = [K_{r(Z)}]_\nu$ is a measurable function of every sufficient Z , proving minimality. Conversely, the same inverse recovers D from $\Gamma(X)$. Any other minimal sufficient Z is also a function of Γ , so Z and D are bijective up to null sets. $\square$

Corollary C.3 (Cross-shadow likelihood learns a dynamics representation). *Assume the relevant conditional laws admit densities with respect to a common dominating measure and have finite population log loss. Let $Z = s(X)$ and allow the decoder to range over all conditional distributions*

$q(Y \mid B, Z)$. Under Assumptions A1–A3, Z attains the same population negative log-likelihood as a decoder given the full source X if and only if Z is cross-shadow sufficient. Therefore every Bayes-optimal code contains D , and every minimal Bayes-optimal code is a one-to-one reparameterization of D .

Proof. For a fixed Z , the optimal decoder is the true conditional law $\mathbb{P}(Y \mid B, Z)$. Since Z is a function of X , the gap from the full-source Bayes risk is

$$\mathcal{R}^*(Z) - \mathcal{R}^*(X) = \mathbb{E}[\text{KL}(\mathbb{P}(Y \mid B, X) \parallel \mathbb{P}(Y \mid B, Z))] \quad (\text{C.13})$$

$$\geq 0. \quad (\text{C.14})$$

Equality holds exactly when $Y \perp\!\!\!\perp X \mid (B, Z)$. The conclusion then follows from Theorem C.2. $\square$

The theorem identifies a representation, not a preferred axis system: if φ is any bijection, $\varphi(D)$ is equally valid. This is the strongest identifiability one can request of an unlabeled latent. A nonminimal Bayes-optimal encoder may additionally store source appearance, but that information is provably unnecessary for cross-shadow prediction and is removed by passing to the minimal sufficient statistic. In our model, the finite variational bottleneck is the mechanism that favors this minimal solution; the theorem does not claim that a KL penalty or a particular SGD run uniquely guarantees it.

Remark C.4 (Deterministic representation versus variational channel). Theorems C.2 and Corollary C.3 concern a deterministic statistic $Z = s(X)$, such as encoder parameters or a posterior mean. The sampled Gaussian channel $q_\phi(z \mid x)$ and the β -weighted KL term in our implementation form a rate–distortion approximation to this ideal. They favor a compact code but do not, by themselves, guarantee exact sufficiency or minimality at finite rate.

C.3. Realization and Learning with Differentiable Networks

The previous theorem is nonparametric. We next show that, under standard regularity conditions, its representation and predictor can be realized by differentiable neural networks.

Theorem C.5 (Differentiable neural representation). *Assume Assumptions A1–A3. In addition, suppose that the supports of X , B , and D are compact subsets of finite-dimensional Euclidean spaces; h is continuous; D admits a continuous embedding $e : \mathcal{D} \rightarrow \mathbb{R}^m$ into the chosen latent dimension; and the target rendering is*

$$Y = F(B, D) \quad (\text{C.15})$$

for a continuous F . Then there exist sequences of differentiable neural networks (E_n, G_n) such that

$$\sup_x \|E_n(x) - e(h(x))\|_2 \rightarrow 0 \quad (\text{C.16})$$

and

$$\sup_{b,x} \|G_n(b, E_n(x)) - F(b, h(x))\|_2 \longrightarrow 0. \quad (\text{C.17})$$

Thus a differentiable neural encoder can represent D up to an invertible coordinate change while discarding source context, and its decoder can attain arbitrarily small cross-shadow reconstruction error.

Proof. Set $E_0 = e \circ h$. Since e is a continuous injection from a compact space into a Hausdorff space, its inverse is continuous on $e(\mathcal{D})$. Hence

$$G_0(b, z) = F(b, e^{-1}(z)), \quad z \in e(\mathcal{D}), \quad (\text{C.18})$$

is continuous and satisfies $G_0(B, E_0(X)) = F(B, D) = Y$. Extend G_0 continuously from the compact set $\mathcal{B} \times e(\mathcal{D})$ to a compact Euclidean neighborhood. Universal approximation with a smooth, bounded, nonconstant sigmoidal activation (such as logistic or tanh) then gives sequences of differentiable networks that approximate E_0 and this extension of G_0 uniformly [29]. Uniform continuity of the extension of G_0 turns both approximations into Eqs. (C.16) and (C.17). $\square$

Corollary C.6 (Exact reconstruction identifies the representation). *In the deterministic setting of Theorem C.5, suppose a code $Z = s(X)$ and decoder G attain*

$$\mathbb{E}\|Y - G(B, Z)\|_2^2 = 0. \quad (\text{C.19})$$

Then Z determines D almost surely. If Z is minimal among exact reconstruction codes, it is a one-to-one reparameterization of D .

Proof. Equation (C.19) implies $Y = G(B, Z)$ almost surely, hence $Y \perp\!\!\!\perp X \mid (B, Z)$. The claim follows from Theorem C.2. Conversely, in this deterministic setting any sufficient Z admits the exact decoder $G(B, Z) = \mathbb{E}[Y \mid B, Z] = Y$, so minimal exact-reconstruction codes and minimal sufficient codes coincide. $\square$

Theorem C.5 is a realizability statement. For stochastic renderers, the same construction applies if the rendering randomness is included in B and available to the decoder; alternatively one may approximate the full conditional kernel with a neural density decoder and use Corollary C.3. There is also a standard statistical learning counterpart: for i.i.d. shadow pairs, bounded second moments, globally optimized empirical squared loss, and a neural-network sieve whose capacity grows at a controlled rate, neural nonparametric regression is risk-consistent [58]. This consistency result learns the composite Bayes predictor; it does not by itself imply convergence of the internal encoder. The representation conclusion follows separately, at a minimal population optimum, from Theorem C.2 and Corollaries C.3 and C.6. It

does not assert finite-sample recovery, a particular convergence rate, convergence of every encoder parameterization, or that arbitrary nonconvex optimization finds a global solution.

C.4. Scope of the “Any Dynamics” Claim

The identification result places no semantic restriction on D and no component-independence assumption on the renderer. For heterogeneous action families, the guarantee applies family by family: let F index the pairing protocol (the family), and let Assumptions A1–A3 hold conditionally on F with a family-specific context measure ν_F . The overlap condition of Assumption A3 is then required only *within* a family, never across families; a human context need never support a robot dynamics. Since F is itself observable from the source (which family a clip shows is visible in the clip), conditioning on F loses nothing, and the theorem identifies D within every family. What is shared *across* families—one encoder and one latent space $\mathcal{Z}$ —is an architectural choice rather than a consequence of the theorem: it is what makes the per-family identified representations unified in practice, as the cross-family transfer results of the main text support empirically. Therefore, any video dynamics for which one can (i) replay the dynamics while independently resampling the other generative factors, (ii) observe the requested dynamics in the source, and (iii) render overlapping contexts, within its family, in which distinct dynamics have distinct effects admits a context-invariant cross-shadow representation. Under the additional finite-dimensionality, compactness, continuity, and latent-capacity conditions of Theorem C.5, that representation and its decoder admit differentiable neural realizations to arbitrary accuracy.

If separation fails, define

$$d \sim d' \iff K_d = K_{d'}. \quad (\text{C.20})$$

Provided this equivalence class is itself observable from X , the same proof identifies $[D]$ —exactly the part of the dynamics that can affect a shadow—and no predictive method can do better. If source observability fails even for $[D]$, a deterministic representation of one source video is impossible; in general only the conditional-law statistic $x \mapsto \mathbb{P}(Y \in \cdot \mid B = \cdot, X = x)$ is identified (it may be representable as a mixture over equivalence classes). These boundaries prevent “any dynamics” from being read as a promise to recover information that no video contains.

Finally, the guarantee applies to exact shadows satisfying Assumption A1. The degenerate self-pairs used to import unpaired real video are an empirical addition and do not, by themselves, provide the identifiability guarantee.

D. Additional Experimental Details and Ablations

D.1. Composition of the Shadow Library

Table 5 summarizes the sources of the Shadow Library: how each source is paired, which control channel it supervises via the per-sample mask ($m_{\text{cam}}, m_{\text{dyn}}$) of Sec. B.1, and its approximate scale and share of the sampling mixture.

D.2. Training and Inference Details

The LAM is trained with $\beta=0.01$ and a fixed prior $\mathcal{N}(0, I)$, always at half the world model’s spatial resolution; real videos enter the stream as degenerate self-pairs. The self-pair probability is 0.5 for the human body-dynamics source and 0 for the other paired sources; since unpaired real-video sources are always self-paired, roughly one third of training samples are self-pairs. The world model is first fine-tuned bidirectionally with flow matching, then converted to a block-causal generator with 3 latent frames per block and rolled out with a key-value cache at inference. The action-transfer comparisons (Table 1) use the model at 480×720 (LAM at 240×360); the long-rollout model (Table 2) generates at 544×960 (LAM at 272×480), warm-started from a lower-resolution checkpoint and fine-tuned at the target resolution for 10k steps on $8 \times \text{H200}$ GPUs with fully-sharded data parallelism. At inference we sample with 30 flow-matching denoising steps and classifier-free guidance 5.0.

D.3. Evaluation Details

In Table 1, appearance families are scored on 81-frame clips from the first frame, averaging 45–50 held-out pairs each; the camera split uses held-out camera-trajectory pairs. Since monocular pose recovery is scale-ambiguous, the VGGT-recovered trajectory is aligned to the target with a similarity transform (Sim(3)) before computing ATE/RPE. The ablation of Table 3 is scored on a smaller held-out subset under the same protocol, so its absolute numbers are not directly comparable to Table 1. **Long-rollout judging protocol** The comparison of Table 2 uses 16 long command streams. For each stream, ShadowDancer and the baseline generate from the same first frame under the same commands; both rollouts are cut at the same boundaries into 240-frame segments, and corresponding segments are compared, giving roughly 64 comparison pairs per baseline. Each pair is anonymized as A and B with randomized assignment and judged by a VLM (Table 5) in a forced choice on each of the three axes, with no tie option; a human audit reviews a random 20% of the judgments to validate the VLM’s decisions. The judge receives both segments, the command list, and an instruction of the following form:

You are judging two video rollouts, A and B, generated by two different systems from the same first

frame under the same sequence of action commands: `<commands>`. Judge the actions, not the image quality. Answer three independent questions; for each you must answer exactly “A” or “B”, even if the difference is small. (1) *Action control*: in which video are the commanded actions actually performed, at the commanded times? (2) *Action fidelity*: which video better preserves the specific weapon, object, and motion of each commanded action? (3) *Long-horizon consistency*: which video stays coherent for longer, without drifting away from the commanded behavior, freezing, or degrading as the rollout proceeds? Ignore differences in visual style, sharpness, or aesthetics, except where degradation makes the action itself unreadable. Output exactly three lines: `control: A|B, fidelity: A|B, consistency: A|B`.

D.4. Representation-Level Probes

Before touching the world model, we probe the latent z itself on held-out shadow pairs (Table 6). First, the *cross-pair transfer ratio*: with the LAM’s own decoder, we reconstruct the target from its visual context using either the target’s own action ($z(\tilde{x})$, “self”) or the shadow source’s ($z(x)$, “cross”); the ratio $\text{MSE}_{\text{cross}}/\text{MSE}_{\text{self}}$ measures how much a cross-appearance z degrades reconstruction. Second, two linear probes on pooled z over held-out UE clips: decoding the *character* and the *scene*. In this evaluation set each character is given its own distinct action set, so character identity is readable from motion alone and the character probe measures action content (higher is better), while the scene probe measures pure appearance (lower is better). Since characters differ in appearance as well as in motion, character accuracy alone could in principle reflect either; the baseline contrast below rules out the appearance route: the self-reconstruction latent leaks *more* appearance by the scene probe yet decodes character only at chance, so character accuracy tracks the motion set rather than looks. The cross-shadow latent transfers almost losslessly on first-person pairs (ratio 1.02 vs. 1.27) and dominates both probes: it decodes character at $11 \times$ chance while leaking less scene identity, whereas the self-reconstruction latent is at chance on content yet leaks *more* appearance — regularization does not substitute for pairing. Note that the cross/self ratio is informative only for a content-bearing latent: the baseline decodes content at chance, so its near-unity third-person ratio reflects an uninformative z , for which self and cross reconstructions are equally unguided, rather than successful transfer.

D.5. Architecture Ablations

Latent size and self-pair ratio The choice $d_z=32$ is inherited from prior latent-action models [18, 32], where this size was validated as the best operating point for self-reconstruction; since real video enters our training stream as self-pairs under the same reconstruction objective (Sec. 3.4),

Table 5. **Composition of the Shadow Library.** Shadow pairs are two frame-synchronized renders of one dynamics trajectory with the remaining generative factors resampled; unpaired corpora enter as self-pairs. (cam, dyn) is the control channel each source supervises. Counts are approximate; “–” marks sources outside the main sampling mixture.

Category	Data source	Pairing	(cam, dyn)	#Seqs	#Frames	Mixture
Human motion	SMPL-X re-render, body dynamics only	shadow pair	(0, 1)	~1k	~1.1M	13.6%
	SMPL-X re-render, camera only	shadow pair	(1, 0)	~1.1k	~1.2M	10.9%
	SMPL-X re-render, camera + body	shadow pair	(1, 1)	~1k	~1.2M	8.2%
	Paired character renders [8]	shadow pair	(1, 1)	~2k	~0.3M	8.8%
Robot manipulation	ManiSkill [51], arm only	shadow pair	(0, 1)	~1.1k	~1M	2.7%
	ManiSkill [51], camera only	shadow pair	(1, 0)	~1.2k	~1.1M	6.8%
	ManiSkill [51], camera + arm	shadow pair	(1, 1)	~1.2k	~1.1M	5.4%
First-person games	GTA-style urban driving and weapon fire	shadow pair	(1, 1)	~0.9k	~0.4M	8.8%
	Cyberpunk-style night city	shadow pair	(1, 1)	~3.4k	~1.6M	8.2%
Third-person games	Unreal Engine scenes (melee, spellcasting)	shadow pair	(1, 1)	~30 scenes	~0.6M	–
	Monster Hunter-style action	shadow pair	(1, 1)	~1k	~0.5M	–
Static-scene camera	DL3DV [38]	self-pair	(1, 0)	~1k scenes	~0.3M	6.8%
Unpaired real video	MiraData Internet clips [33]	self-pair	(1, 1)	~5k	~6.5M	5.4%
	OpenX BridgeData [42]	self-pair	(0, 1)	~29k	~0.8M	0.1%
	OpenX FurnitureBench [42]	self-pair	(0, 1)	~5k	~2.2M	4.1%
	OpenX Berkeley UR5 [42]	self-pair	(0, 1)	~1k	~0.1M	2.7%
	OpenX Jaco Play [42]	self-pair	(0, 1)	~1.1k	~0.1M	2.7%
	OpenX Stanford HYDRA [42]	self-pair	(0, 1)	~0.6k	~0.3M	2.7%
	OpenX RoboTurk [42]	self-pair	(0, 1)	~2k	~0.1M	2.0%

Table 6. **LAM representation probes** on held-out shadow pairs (24 pairs) and UE clips (60 clips; character chance = 0.04, scene chance = 0.20). The Olaf-recipe row follows [32]: self-reconstruction with feature alignment and a single head.

LAM training	Cross/self MSE ratio↓		Linear probe on z	
	1st-person	3rd-person	Char. (content)↑	Scene (appear.)↓
Self-recon.+reg., single head	1.265	1.067	0.050	0.650
Cross-shadow, 3-prompt (ours)	1.017	1.103	0.450	0.433

their operating point carries over, and it is consistent with the invariance argument of Sec. B.2 that the bottleneck should stay too small to smuggle appearance. A budget-matched sweep confirmed this choice: the representation probes were insensitive to $d_z \in \{16, 64\}$ and to forcing the self-pair ratio, while reconstruction favored $d_z \leq 32$, so we retain $d_z=32$.