

Consistent Video-to-Video Translation via Explicit Correspondences

Gaurav Parmar^{1,2} Zhengqi Li² Richard Zhang² Jun-Yan Zhu¹
Srinivasa Narasimhan¹ Eli Shechtman² Yotam Nitzan²

Carnegie Mellon University¹ Adobe Research²

Abstract. Interactive video-to-video applications require real-time generation while maintaining long-range temporal consistency. However, recent methods achieve speed by restricting attention to a fixed-length causal window. This limits long-range consistency whenever relevant history falls outside the context window. To address this limitation, we introduce **vid2vid-long**, a correspondence-based approach grounded in a simple observation. Video-to-video translation provides an *explicit* cross-frame reference: the regions that match in the input video should match in the output video too. Concretely, we compute patch-level correspondences in the input, guiding retrieval of previously generated outputs in the corresponding regions, and adding them to the context window. Our experiments demonstrate that **vid2vid-long** consistently improves long-range temporal consistency across recent video models while preserving generation quality and real-time performance.

1 Introduction

Real-time video-to-video translation unlocks a large class of interactive applications, including streaming, gaming, and augmented and virtual reality. Recent advances in large scale text-to-video diffusion models [1, 6, 16, 30, 40, 56], and subsequently distilling them into few-step, autoregressive models [12, 20, 22, 46, 50], have enabled high-quality generation at interactive speeds.

Autoregressive video generation relies on a history of past frames to ensure temporal continuity. However, determining an efficient mechanism for maintaining this history remains a fundamental challenge. Naïvely retaining all past outputs causes memory to scale linearly and runtime quadratically with sequence length, both of which are prohibitive for real-time applications. To maintain real-time latency, many methods resort to a fixed-length sliding window. While this ensures a fixed memory and runtime per frame, it fundamentally sacrifices long-range consistency: once an object leaves the window, often after only a few seconds, the model “forgets” it forever.

Several lines of work have attempted to bridge this gap. First, attention sinks [22, 37, 44, 46] heuristically retain the initial frames as a global anchor.

However, this approach inherently cannot address novel content that first appears in a later frame, even in cases like simple disocclusions. Another approach involves models with learned hidden states, such as Mamba [4, 8], which

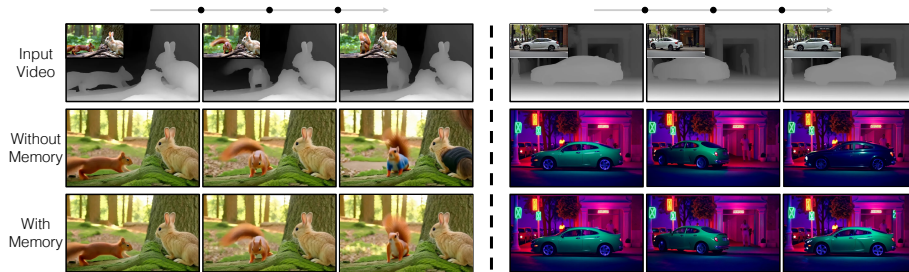


Fig. 1: Long-range consistency without long context windows. Standard autoregressive video translation forgets content once it falls outside its context window (middle row). Our proposed method `vid2vid-long` (bottom row) explicitly retrieves matching past outputs when the inputs recur, maintaining consistency over long videos. Please see the supplement for a large grid of qualitative result videos.

compress history into fixed-size bottlenecks. While efficient, these models often forgo the expressive power of transformers and remain largely limited to scenes with lower temporal complexity [31].

Finally, some methods use pose conditioning [51] or explicit world representation (e.g., 3D point clouds) to maintain a consistent world [41, 45]. These methods struggle to model dynamic scenes and require additional privileged inputs that are hard to obtain.

In this paper, we address the challenge of long-range consistency in video-to-video translation. Our insight is that long-range consistency cues are provided in the input video itself. That is, if two regions match in the input video, their generated outputs should match as well. The input space can therefore serve as an index for retrieving previously generated content, allowing consistency to be enforced, without requiring access to the full output history.

As such, we introduce `vid2vid-long`, a correspondence-based dynamic memory mechanism that enables long-term consistency while preserving real-time performance. Concretely, we compute patch-level correspondences in the input video and use them to index a compressed global representation of prior outputs. This representation grows only when novel content appears, rather than with video length, making the memory growth sub-linear. The approach is complementary to existing autoregressive distillation backbones.

We evaluate our approach in combination with multiple autoregressive distillation frameworks, including Self-Forcing, LongLive, and Rolling Forcing. In all settings, incorporating our technique substantially improves long-range consistency, while maintaining real-time inference speed and high visual quality. Our main contributions are:

- We identify a structural property of video-to-video translation: correspondences in input space induce correspondences in output space, enabling input-driven retrieval for long-range consistency.

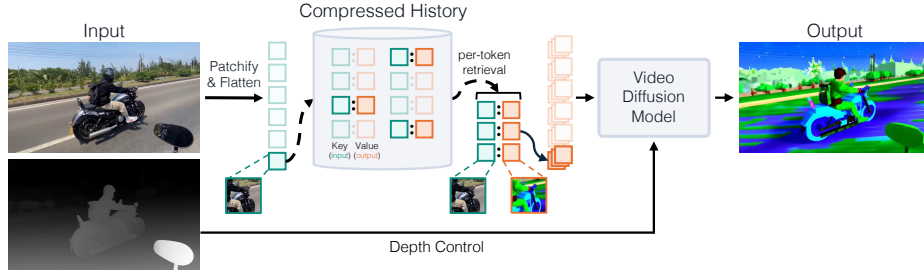


Fig. 2: Correspondence-based video translation. The memory is populated with input-output correspondence pairs from previous generations. To support the consistent translation of the next block of frames, we first retrieve from memory the output patches whose corresponding input is most similar to the current. These output patches, together with the normal input signal, is the input to our video diffusion model. Based on the new induced input-output pairs, we write to memory only those that are sufficiently different from existing.

- We propose a correspondence-based mechanism that maintains a compressed global representation of prior outputs whose size scales with visual novelty rather than sequence length.
- We demonstrate consistent improvements in long-range temporal coherence across multiple real-time autoregressive backbones.

2 Related Works

Controllable Generation. The rapid advancements in text-to-image [18, 33, 43] and text-to-video [1, 10, 40] has been accompanied by controllable adapters that allow a more fine-grained control over the generation. In the image domain, works like *ControlNet* [53], *Sketch-Guided-Diffusion* [39], *T2I adapter* [25], *img2img-turbo* [27] show the ability to add structure control like edge maps, depth maps, and hand drawn sketches. Models such as *IP-adapter* [47], *BLIP-diffusion* [19], *SynCD* [17], *VisualComposer* [28] show ability to condition the generation of subject identity. These methods have also been extended to videos in recent methods like *ControlVideo* [55], *Control-A-Video* [3], *VACE* [14], and *COSMOS-T* [26]. However all of these methods are designed for offline generation; they rely on global bidirectional attention and need multiple diffusion steps, making them both computationally expensive, non-causal, and unsuitable for interactive applications. We aim to address this gap in this paper.

Diffusion Distillation. Several prior works have explored speeding up diffusion models by reducing the number of denoising steps. Initial image domain approaches like *consistency models* [38], *progressive distillation* [34], *adversarial distillation* [36], and *distribution matching distillation* [48, 49], *moment matching distillation* [35] enabled high-quality inference in 1–4 steps. Recent methods

like *CausVid* [50] and *Diffusion Forcing* [2] extend the distillation to video generative models but often struggle with error accumulation during auto-regressive inference. To address this, *Self-Forcing* trains the student model on its own generated rollout, rather than ground truth video. We build on this to ensure our rendering remains stable for longer sessions.

Long Context Generation. Maintaining long-term consistency in autoregressive generation is computationally challenging due to the quadratic complexity of transformer attention. This is particularly severe for videos: five seconds of 480p video can exceed 32,000 tokens in modern architectures like WAN2.1 [40]. Existing methods such as *Streaming-T2V* [11] or *LongLive* [46] typically use sliding windows or hierarchical latent memory. However, these often suffer from recency bias, where the model forgets the appearance of objects or environments during loop closures or long-duration occlusions. While the language modeling literature has explored KV cache pruning and compression [7, 23], these methods often ignore the spatiotemporal continuity of video.

Rather than deciding which history to retain solely based on temporal position, learned compression, or token importance, we use input-space correspondences to retrieve the specific past outputs that are relevant to the current frame. This makes the memory mechanism task-aware and particularly well-suited to video-to-video translation, where repeated input regions should induce repeated output appearances.

3 Method

Our objective is to design an autoregressive video-to-video diffusion model equipped with a correspondence-based memory mechanism. This architecture enables long-range temporal consistency while maintaining the low-latency performance required for real-time deployment. We begin by reviewing the autoregressive formulation in Section 3.1. In Section 3.2, we introduce our explicit correspondence-based memory component. We then detail our strategy for conditioning on this historical context in Section 3.3. Finally, we describe the specific training objectives and real-time inference procedures in Sections 3.4 and 3.5, respectively.

3.1 Preliminaries

Bidirectional Video Generation. Our method builds upon pre-trained flow-matching models designed for video-to-video translation. Specifically, we adopt VACE [14], which is a latent flow-matching model, operating within the latent space of a pre-trained Variational Autoencoder (VAE) [33]. Given a block of input frames $\mathbf{I}^{(b)} \in \mathbb{R}^{T \times H \times W \times 3}$, the VAE encoder $\mathcal{E}$ maps the pixel data to a latent representation $\mathbf{z}^{(b)} \in \mathbb{R}^{t \times h \times w \times c}$. The latent dimensions are determined by the spatial and temporal compression factors, which for VACE are given as follows: $h = H/8$, $w = W/8$, and $t = 1 + T/4$. Following common practice in vision transformers [29], the latent block is further “patchified” into a sequence of tokens. Each latent frame is partitioned spatially into non-overlapping patches of size $1 \times p \times p$,

where the temporal dimension of a patch remains singular. This process yields a total of $N = t \times (h/p) \times (w/p)$ tokens per block. Specifically for VACE, $p = 2$.

For the generative model itself, VACE utilizes a bidirectional spatiotemporal transformer to generate an entire video at once. This design implies that every generated token depends on the full context of the video, including both past and future tokens. This design is expensive in computation and memory, limiting their operation to short, fixed-length sequences, typically up to five seconds.

Autoregressive Video Generation. To enable long-form generation, recent methods [12, 22, 46, 50] replace bidirectional attention with causal attention. These models generate videos autoregressively in temporal blocks, where each block $\mathbf{z}^{(b)}$ is generated by conditioning on the previously generated blocks $\mathbf{z}^{(<b)}$. Access to previously generated blocks is typically implemented via a causal KV cache. While theoretically sound, maintaining the entire history $\mathbf{z}^{(<b)}$ causes the runtime to grow quadratically with the video length, and the memory usage linearly. To ensure a constant per-frame runtime necessary for streaming applications, the causal attention is typically restricted to a smaller set of past frames, implemented with a sliding window of size W . Under this constraint, the generation of block b conditions only on the most recent context $\mathbf{z}^{(b-W:b-1)}$. Although this approach bounds the computational cost, it fundamentally limits temporal consistency to the window size. Any information that falls outside this window becomes inaccessible to the model, leading to identity drift and “forgetting” over long horizons. Our method addresses this limitation by providing an explicit correspondence-based memory that allows any autoregressive model to access the history arbitrarily far beyond its fixed window.

3.2 Compressed Patch History

In video-to-video translation, the availability of the input stream allows us to move beyond temporal proximity as the sole proxy for relevance. While a sliding window approach utilizes only recent frames for context, we can use explicit correspondences in the input space to identify relevant history from any point in the past. This is a reliable strategy because regions that “match” in the input video should “match” in the generated output (see Figure 2).

By explicitly hard-coding these correspondences, we offload the burden of discovery from the transformer’s expensive attention layers. Rather than forcing the model to rediscover these relationships through expensive high-dimensional matching, we provide them directly via an external index. We represent this memory as an evolving set of input-output pairs. Each entry consists of a descriptor for an input patch, which acts as a retrieval key, and the corresponding generated output patch, which serves as the value for future conditioning. To ensure the system remains scalable for long sequences, we only update the memory when the input video presents novel information. If an incoming patch is already sufficiently represented in the history, we do not save it. This ensures that the memory footprint scales with the visual novelty of the scene, rather than the raw sequence length. (See history growth analysis in the appendix.)

To implement this mechanism, we define three concrete operations that govern the memory lifecycle. We first specify the patch-level representation for the keys and values, followed by the retrieval protocol for reading from history, and finally the novelty-based policy for writing to the global store.

Patch Descriptor. Our memory mechanism leverages the spatial discretization already present in the diffusion transformer (see Section 3.1). In this framework, every input video spatiotemporal patch is already mapped to a small latent patch of shape $(1, 2, 2)$ that corresponds to a single token. From here, we use term “latent patch” for these token-corresponding patches. We assign each latent patch a global index g that uniquely identifies its spatiotemporal coordinates. For every input latent patch at index g , denoted as $\mathbf{z}_g$, we define a descriptor $\mathbf{k}_g$ to serve as a retrieval key. We evaluated several feature extractors ϕ to produce these keys, including external vision backbones applied in pixel space. However, our experiments show that the most effective representation is the latent patch itself (further analysis is included in the appendix). Consequently, we define the feature extraction function as the identity mapping, where $\mathbf{k}_g = \mathbf{z}_g$. This approach also avoids the computational cost of decoding latents back to pixels and the inference run of additional feature extraction backbones. Under this formulation, the memory stores pairs $(\mathbf{k}_g, \mathbf{v}_g)$, where the key is the input latent patch and the value $\mathbf{v}_g$ is the corresponding to the generated output latent patch.

Memory Management. To maintain temporal coherence throughout the generation process, we define the memory as a dynamic set of key-value pairs that evolves alongside the autoregressive blocks. Let $\mathcal{H}_b$ represent the memory state after the generation of block b , initialized as an empty set $\mathcal{H}_0 = \emptyset$. We similarly define $\mathcal{G}_b$ as the set of all global indices g associated with the latent tokens generated within block b . Under this formulation, the generation of each subsequent block $b + 1$ utilizes the existing history $\mathcal{H}_b$ for retrieval and concludes by updating the state to $\mathcal{H}_{b+1}$ using the newly generated tokens in $\mathcal{G}_{b+1}$.

Memory Write Policy. The memory state is updated after the generation of each block by evaluating the novelty of every new latent patch relative to the existing stored patches. For each generated latent patch at index $g \in \mathcal{G}_{b+1}$, we pair the input descriptor $\mathbf{k}_g$ with the corresponding output latent patch $\mathbf{v}_g$. The similarity between a query latent patch $\mathbf{k}_g$ and the current memory $\mathcal{H}_b$ is defined as the maximum cosine similarity:

$$\text{sim}(\mathbf{k}_g, \mathcal{H}_b) = \max_{(\mathbf{k}_{g'}, \mathbf{v}_{g'}) \in \mathcal{H}_b} \text{sim}(\mathbf{k}_g, \mathbf{k}_{g'}) = \max_{(\mathbf{k}_{g'}, \mathbf{v}_{g'}) \in \mathcal{H}_b} \frac{\mathbf{k}_g \cdot \mathbf{k}_{g'}}{\|\mathbf{k}_g\| \|\mathbf{k}_{g'}\|}. \quad (1)$$

The latent patch and its corresponding generated output $\mathbf{v}_g$ are committed to the memory only if this similarity falls below a novelty threshold τ_h . The updated memory $\mathcal{H}_{b+1}$ is thus defined as:

$$\mathcal{H}_{b+1} = \mathcal{H}_b \cup \{(\mathbf{k}_g, \mathbf{v}_g) : \text{sim}(\mathbf{k}_g, \mathcal{H}_b) < \tau_h\}_{g \in \mathcal{G}_{b+1}}. \quad (2)$$

This policy ensures that memory footprint and inference cost scales with discovery of unique visual content rather than total sequence length.

Memory Read Mechanism. Before generating the output latents for block $b + 1$, we query $\mathcal{H}_b$ using the descriptors of the current input latent. For each patch index g , we identify the set of k nearest neighbors $\mathcal{N}_k(g)$ whose similarity exceeds a retrieval threshold τ_m :

$$\mathcal{N}_k(g) = \text{Top}_k(\{(\mathbf{k}_{g'}, \mathbf{v}_{g'}) \in \mathcal{H}_b : \text{sim}(\mathbf{k}_g, \mathbf{k}_{g'}) \geq \tau_m\}). \quad (3)$$

For every query latent patch that satisfies this condition, we retrieve the stored output values $\mathbf{v}_{g'}$. If no entries in $\mathcal{H}_b$ meet the threshold τ_m , the neighborhood $\mathcal{N}_k(g)$ is empty. We track these occurrences with a binary validity mask $\mathbf{m}_g$, where $\mathbf{m}_g = 1$ if $\mathcal{N}_k(g) \neq \emptyset$ and 0 otherwise. This mask and the retrieved values are then used to condition the diffusion transformer as described in Section 3.3.

3.3 Conditioning via Correspondence Cross-Attention

We now describe how these retrieved values are used to condition the video translation model. The retrieved output patches $\{\mathbf{v}_{g'}\}$ are tokenized using the same patch embedding as the input latents, producing history tokens $\mathbf{h}_{g'}$ in a shared representation space.

To incorporate these history tokens into the video-model architecture, we add correspondence cross-attention layers to every transformer block between self-attention and text cross-attention layers. Each token attends only to its own k retrieved neighbors. Concretely, for each token index g , the current hidden state $\mathbf{x}_g$ provides the query. Let $H_g = \{\mathbf{h}_{g'}\}_{(\mathbf{k}_{g'}, \mathbf{v}_{g'}) \in \mathcal{N}_k(g)}$ denote the set of retrieved history tokens for query g . We define

$$\text{CorrXA}(g) = \mathbf{m}_g \cdot W_o \text{Attn}(W_q \mathbf{x}_g, W_k H_g, W_v H_g). \quad (4)$$

where $\mathbf{m}_g$ is the validity mask that zeros out tokens without valid correspondences. The output is added as a residual: $\mathbf{x}_g \leftarrow \mathbf{x}_g + \text{CorrXA}(g)$. The output projection W_o is zero-initialized [53], so the layer acts as an identity at the start of training, preserving the behavior of the pretrained model. Since each token attends only to its k neighbors, the computational overhead of these layers is minimal. We’ve described the cross-attention operation for a single token. To implement this efficiently, we reshape the query and memory tensors to collapse the sequence length into the batch dimension, effectively treating each token’s cross-attention as an independent $1 \times k$ operation. Please see the appendix for runtime analysis and architecture diagram.

3.4 Training

We incorporate our correspondence-memory with several different autoregressive video model. Unless stated otherwise, we use Rolling Forcing [22] distillation with 5 steps as the default configuration. The student model is initialized

Table 1: Quantitative evaluation on depth-to-video translation. We report long-context reconstruction metrics (LPIPS, DreamSim, CLIP-Sim), quality scores (V-Bench), and inference efficiency. Offline methods are shown in gray on the top and distillation based autoregressive methods are shown at the bottom. The best method in each column is marked in bold and the second best method is marked in underline.

Family	Method	Base Model	Reconstruction			Quality Score $_{\uparrow}$	Efficiency	
			LPIPS $_{\downarrow}$	DreamSim $_{\downarrow}$	CLIP-Sim $_{\uparrow}$		Latency(s) $_{\downarrow}$	FPS $_{\uparrow}$
Offline	VACE 1.3B	N/A	0.463	0.195	0.901	0.835	79	1.03
	Control-A-Video	SD 1.5	0.288	0.169	0.907	0.722	11	0.89
	VideoComposer	N/A	0.664	0.333	0.847	0.724	11	1.36
	ControlVideo	SD 1.5	0.174	0.065	0.954	0.778	550	0.59
Auto-regressive	Self-Forcing (4 steps)	VACE 1.3B	<u>0.548</u>	<u>0.203</u>	<u>0.881</u>	<u>0.789</u>	<u>0.67</u>	<u>13.74</u>
	LongLive (4 steps)	VACE 1.3B	0.511	0.218	0.874	<u>0.817</u>	0.68	<u>13.83</u>
	Rolling Forcing (5 steps)	VACE 1.3B	0.427	0.172	0.905	0.799	2.58	10.99
	Ours (2 step)	VACE 1.3B	<u>0.302</u>	<u>0.079</u>	<u>0.941</u>	0.814	1.09	13.88
	Ours (5 step)	VACE 1.3B	0.240	0.066	0.955	0.828	3.72	8.10

from pretrained VACE 1.3B and all model parameters are finetuned jointly, including the newly introduced correspondence cross-attention layers. We adopt the standard Rolling Forcing objective without modification or extra loss terms. The correspondence memory is active during training, and memory read and write operations are performed as described in Section 3.2. The student conditions on retrieved correspondences from its own previously generated blocks, and newly generated patches are written to memory according to the novelty threshold. The teacher model is a bidirectional VACE 14B model and it does not use correspondence memory. Please see the supplementary for full training details.

3.5 Inference

At inference time, the model generates the video autoregressively, processing fixed-size blocks sequentially. At the start of each video, we set $\mathcal{H}_0 = \emptyset$. For each block b , we compute the input latent tokens $\mathbf{z}_g$, $g \in \mathcal{G}_b$. Using $\mathbf{z}_g$, we query $\mathcal{H}_b$ to compute neighbors $\mathcal{N}_k(g)$ for each $g \in \mathcal{G}_b$. The retrieved neighbors are used to condition the translation network through correspondence cross-attention layers during each denoising step of block b . Retrieval is performed once per block and the resulting neighbors are reused across all denoising steps. After block b is fully denoised, we obtain the output latents $\mathbf{v}_g$, $g \in \mathcal{G}_b$, and update the memory according to Eq. 2. The updated memory $\mathcal{H}_{b+1}$ is used to generate the next block $b + 1$. A complete inference pseudocode is provided in the supplementary.

4 Experiments

Here, we first evaluate our method on depth-to-video translation, focusing on long-range temporal consistency on long videos. Our mechanism is agnostic to the conditioning modality. Experiments are conducted on two benchmarks:

Table 2: Quantitative evaluation on depth-to-video translation (DL3DV). We report long-context reconstruction metrics (LPIPS, DreamSim, CLIP-Sim), quality scores (V-Bench), and inference efficiency. Offline methods (gray) are shown on the top and distillation based autoregressive methods are in the bottom. The best method in each column is marked in bold and the second best method is marked in underline.

Family	Method	Base Model	Reconstruction			Quality Score $\uparrow$	Efficiency	
			LPIPS $\downarrow$	DreamSim $\downarrow$	CLIP-Sim $\uparrow$		Latency(s) $\downarrow$	FPS $\uparrow$
Offline	VACE 1.3B	N/A	0.592	0.240	0.881	0.819	79	1.01
	Control-A-Video	SD 1.5	0.437	0.188	0.895	0.727	14	0.89
	VideoComposer	N/A	0.672	0.303	0.854	0.718	12	1.39
	ControlVideo	SD 1.5	0.205	0.057	0.957	0.794	2518	0.38
Auto-regressive	Self-Forcing	VACE 1.3B	0.616	0.238	0.870	0.790	0.67	13.74
	LongLive	VACE 1.3B	0.534	0.170	0.896	0.806	0.68	13.83
	Rolling Forcing	VACE 1.3B	0.445	0.163	0.907	0.806	2.58	10.99
	Ours (2 step)	VACE 1.3B	0.379	0.083	0.932	0.799	1.09	13.88
	Ours (5 step)	VACE 1.3B	0.305	0.058	0.952	0.821	3.72	8.10

Table 3: Sensitivity to the distillation method. Our correspondence mechanism consistently improves long-context reconstruction across all three autoregressive distillation methods, with minimal impact on quality. The base model used is VACE 1.3B.

	VACE-Bench				DL3DV			
	LPIPS $\downarrow$	DreamSim $\downarrow$	CLIP-Sim $\uparrow$	Quality Score $\uparrow$	LPIPS $\downarrow$	DreamSim $\downarrow$	CLIP-Sim $\uparrow$	Quality Score $\uparrow$
Self-Forcing	0.548	0.203	0.881	0.789	0.616	0.238	0.870	0.790
+ Correspondence	0.343 $\downarrow$ 37%	0.097 $\downarrow$ 52%	0.930 $\uparrow$ 6%	0.794	0.451 $\downarrow$ 27%	0.151 $\downarrow$ 37%	0.897 $\uparrow$ 3%	0.775
LongLive	0.511	0.218	0.874	0.817	0.534	0.170	0.896	0.806
+ Correspondence	0.272 $\downarrow$ 47%	0.057 $\downarrow$ 74%	0.953 $\uparrow$ 9%	0.807	0.345 $\downarrow$ 35%	0.064 $\downarrow$ 63%	0.945 $\uparrow$ 5%	0.786
Rolling Forcing	0.427	0.172	0.905	0.799	0.445	0.163	0.907	0.806
+ Correspondence	0.240 $\downarrow$ 44%	0.066 $\downarrow$ 62%	0.955 $\uparrow$ 6%	0.828	0.305 $\downarrow$ 32%	0.058 $\downarrow$ 65%	0.952 $\uparrow$ 5%	0.821

VACE-Benchmark [14], covering diverse real-world scenes, and DL3DV [21], a geometry-grounded 3D dataset enabling controlled pose revisits. We compare against bidirectional diffusion models, prior video translation methods, and recent autoregressive distillation approaches, and ablate the key components of our correspondence mechanism. Unless stated otherwise, all experiments use VACE 1.3B at 480p resolution.

Datasets. We evaluate on two benchmarks. VACE-Benchmark [14] contains diverse real-world video sequences with varying motion and scene complexity. DL3DV [21] is a geometry-grounded 3D dataset with calibrated camera poses, enabling controlled evaluation of viewpoint revisits and long-range consistency. Additional dataset details and preprocessing are provided in supplement.

Evaluation Protocol. We evaluate our method in terms of video quality, long-range consistency under revisits, and inference efficiency. Video quality is measured on full generated videos using V-Bench [13] metrics (subject consistency, background consistency, temporal flickering, motion smoothness, aesthetic quality, and imaging quality). Experiments report average of scores, individual metrics values are provided in supp. Long-range consistency is evaluated using con-

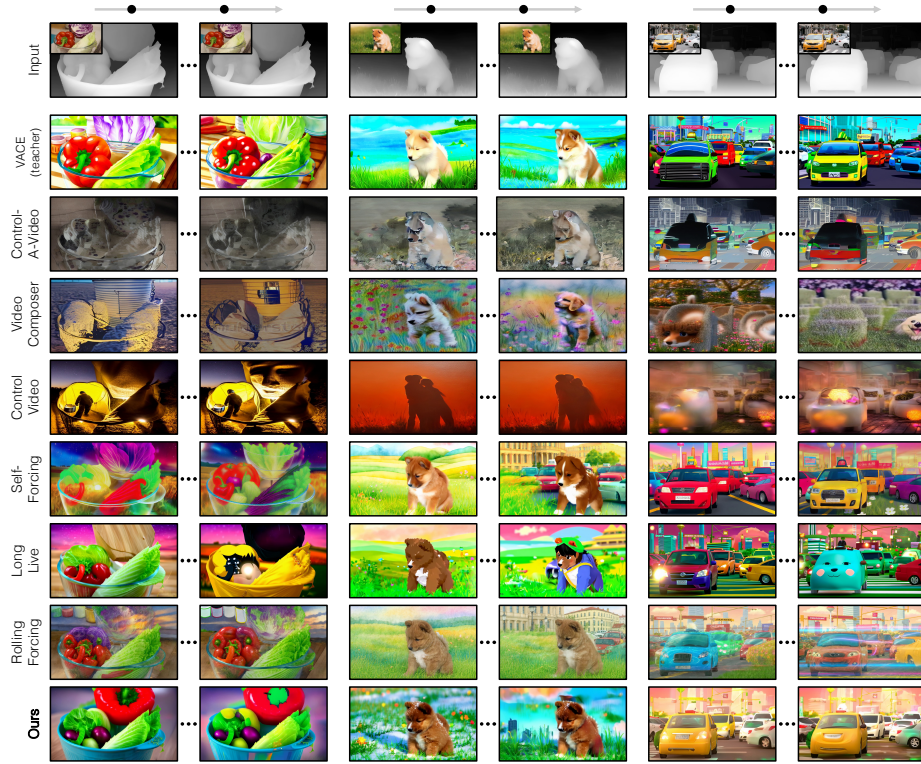


Fig. 3: Qualitative comparisons. We show two frames sampled far apart in time for three example videos (left and right). Bidirectional methods (VACE 1.3B) and Autoregressive distillation methods produce high-quality individual frames but fail to maintain object identity over long durations. Adding our correspondence mechanism preserves consistent appearance throughout the video.

trolled revisit protocols that require model to reproduce previously observed content. On VACE-Benchmark, following [9], we construct synthetic long sequences by concatenating two distinct 5-second clips in an interleaved order ($A \rightarrow B \rightarrow A \rightarrow B$), and measure reconstruction consistency between first and final occurrences. Consistency is quantified using LPIPS [54], DreamSim [5], and CLIP [32] image similarity. On DL3DV, we train a 3D Gaussian Splatting model [15] using calibrated camera poses to render exact pose revisits and measure consistency between first-visit and revisit frames. We measure inference latency and throughput. Inference latency is measured as time to produce first frame and throughput is reported in frames per second (FPS). All measurements are performed on a single NVIDIA H100 GPU.

Implementation Details. We use VACE 1.3B [14] as the student model and VACE 14B as the teacher. The student is initialized from the pretrained VACE 1.3B checkpoint and trained using Rolling Forcing [22]. For correspondence retrieval,

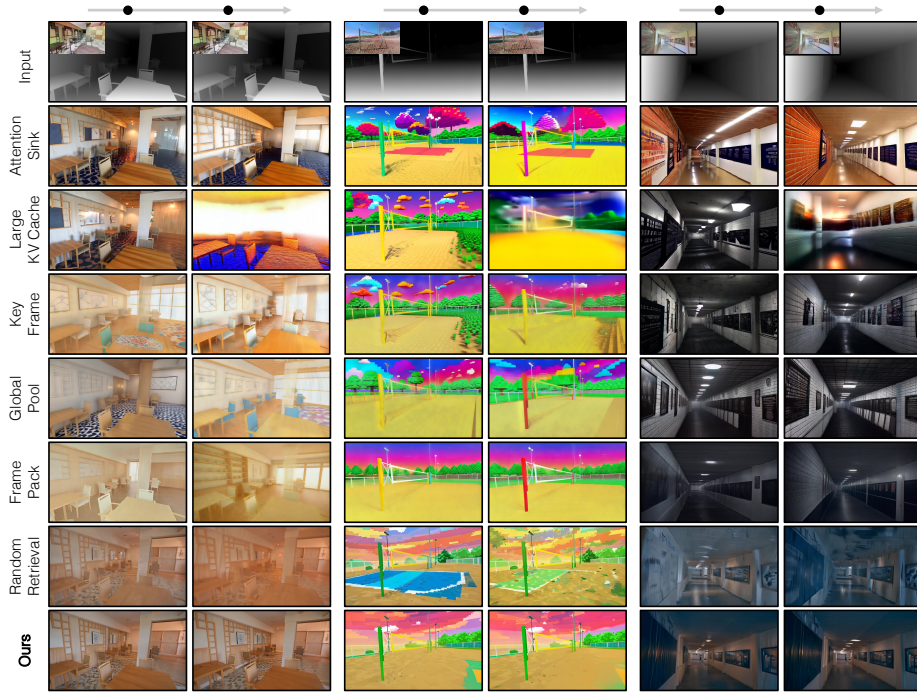


Fig. 4: Comparing to alternate consistency mechanisms. We compare different strategies for extending the temporal context beyond sliding window, including enlarging the KV cache, keyframe retention, global pooling, FramePack [52] compression, and random retrieval. While these approaches provide partial improvements, they suffer for loss of details or quality degradation under long-range revisits. In contrast, our correspondence based method (bottom row) enables accurate reuse of historical content and preserving consistency over long time horizons.

we use VAE encoder’s latent features as keys, with $k = 3$ retrieved neighbors, history threshold $\tau_h = 0.75$, and memory threshold $\tau_m = 0.75$. Training is performed using AdamW [24] with batch size 64, a learning rate of 2×10^{-6} for the student network and 4×10^{-7} for the critic.

4.1 Baseline Comparisons

We evaluate our method against several video-to-video translation frameworks. Quantitative and qualitative results are in Table 1 and Figure 3. First, we compare to Control-A-Video [3], VideoComposer [42], and ControlVideo [55].

They operate in an offline setting with multi-step diffusion and global attention over fixed-length sequences. In our experiments, they produce lower visual quality and struggle to maintain long-range consistency (Figure 3). ControlVideo achieves relatively strong reconstruction scores due to its hierarchical sampling strategy over long sequences. However, this comes at the cost of high latency

Table 4: Comparison of alternate consistency mechanisms. We compare several strategies for extending temporal context beyond the default sliding window (Self-Forcing, $W=21$): increasing the KV cache window size, retaining every N -th keyframe, maintaining a global average-pooled KV cache, Attention Sinks, and FramePack. While larger windows and keyframe retention provide modest reconstruction gains, our correspondence-based retrieval achieves substantially better long-range consistency.

Method	VACE-Bench				DL3DV			
	LPIPS $\downarrow$	DreamSim $\downarrow$	CLIP-Sim $\uparrow$	Quality $\uparrow$	LPIPS $\downarrow$	DreamSim $\downarrow$	CLIP-Sim $\uparrow$	Quality $\uparrow$
KV cache (2 $\times$)	0.350	0.139	0.910	0.795	0.710	0.364	0.814	0.759
KV cache (3 $\times$)	0.308	0.120	0.920	0.793	0.770	0.496	0.744	0.752
KV cache (4 $\times$)	0.330	0.134	0.913	0.795	0.759	0.543	0.714	0.753
Keyframe (every 3rd)	0.329	0.115	0.923	0.790	0.607	0.257	0.857	0.776
Keyframe (every 5th)	0.350	0.128	0.914	0.789	0.581	0.240	0.857	0.776
Keyframe (every 10th)	0.373	0.134	0.912	0.789	0.523	0.190	0.880	0.777
Global avg. pool	0.450	0.176	0.891	0.753	0.457	0.145	0.909	0.772
Attention Sink	0.436	0.159	0.899	0.793	0.502	0.156	0.901	0.788
FramePack	0.304	0.095	0.938	0.775	0.344	0.105	0.932	0.802
Random retrieval	0.396	0.144	0.916	0.807	0.451	0.156	0.911	0.801
Correspondence (Ours)	0.240	0.066	0.955	0.828	0.305	0.058	0.952	0.821

Table 5: Importance of correspondence patches. We measure the effect of removing the correspondence patches. Without correspondence, reconstruction metrics degrade substantially while quality scores remain similar, demonstrating the effectiveness of correspondence tokens for maintaining long term consistency.

Method	VACE-Bench				DL3DV			
	LPIPS $\downarrow$	DreamSim $\downarrow$	CLIP-Sim $\uparrow$	Quality $\uparrow$	LPIPS $\downarrow$	DreamSim $\downarrow$	CLIP-Sim $\uparrow$	Quality $\uparrow$
Ours w/o Corr.	0.379	0.160	0.910	0.829	0.437	0.124	0.921	0.822
Ours	0.240 $\downarrow$ 37%	0.066 $\downarrow$ 59%	0.955 $\uparrow$ 5%	0.828	0.305 $\downarrow$ 30%	0.058 $\downarrow$ 53%	0.952 $\uparrow$ 3%	0.821

(2518 seconds in our setting) and reduced perceptual quality, making it impractical for real-time video-to-video translation.

We next compare to newer, state-of-the-art video translation models, such as VACE [14]. As discussed before, VACE is a bidirectional model which is limited to a fixed-length, short video sequences. Unfortunately, there is no publicly available adaption of VACE to an unbounded-length, streaming use case. In order to still apply VACE on our setting, we adapt it in two manners.

First, we directly evaluate VACE by applying it on non-overlapping video segments. We find that this naive adaptation of VACE produces high-quality results, but unsurprisingly has poor long-range consistency. For example, note row 2 in Figure 3 where each individual frame is of high-quality, but objects change their colors drastically over the length of the video.

Second, we adapt VACE to a streaming use-case with three recent causal, autoregressive distillation methods: Self-Forcing [12], LongLive [46], and Rolling Forcing [22]. These methods were originally designed for text-to-video generation and aren’t applied to video-to-video yet. Here, we implement and adapt the three methods to video-to-video with VACE. We find that all methods achieve competitive quality scores and substantially lower latency than the bidirectional

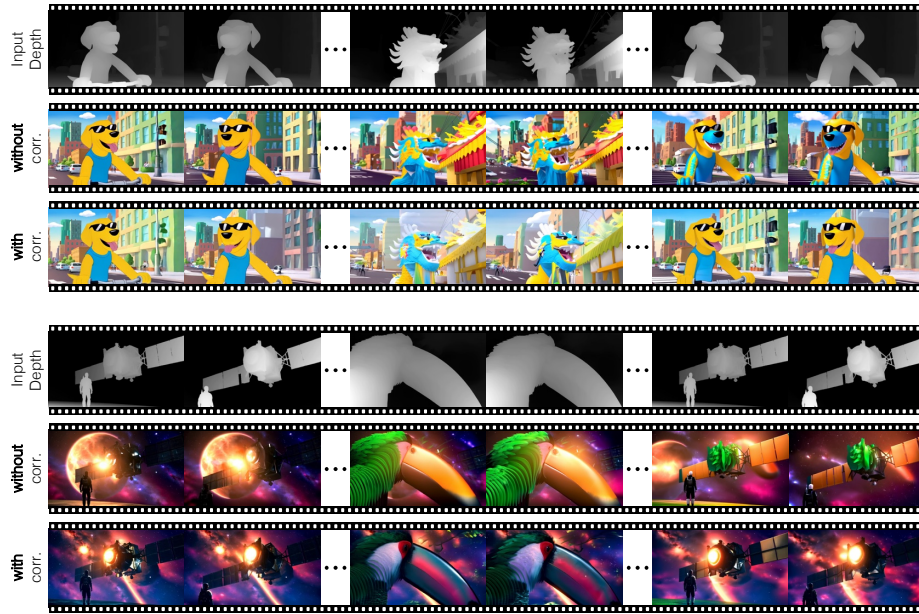


Fig. 5: Effect of correspondence pointers on long-range consistency. Without correspondence pointers, the model generates plausible but inconsistent frames. Object appearance drifts over time. Consider the top example, the dog starts off having a yellow snout, but it changes to blue when it re-appears. With correspondence pointers, the model retrieves and reuses visual context from earlier frames, maintaining consistent identity throughout the video. For example in the top video, the dog snout is yellow at the beginning of the video. However it turns into a blue snout when reappearing if correspondence is not used. However it stays consistently yellow with our method.

VACE baseline (Tables 1, 2). Nevertheless, these methods still perform poorly on long-range consistency, where Self-Forcing and LongLive under-perform bidirectional VACE and Rolling Forcing only slightly over-performs it.

Our main contribution, correspondence mechanism, can be readily applied with any autoregressive method. In Table 1, 2 and Figure 3, we report results of our best configuration, where our method is applied with Rolling Forcing. Our models achieve the strongest long-range performance and visual quality while maintaining real-time performance. In Table 3, we report the results of our method when applied on top of all autoregressive methods. Adding our correspondence mechanism consistently improves long-term consistency across all autoregressive backbones while preserving visual quality. Table 3 results demonstrate that our method is orthogonal to choice of autoregressive backbone. We also compare different conditioning mechanisms in the supplement.

4.2 Alternate Consistency Mechanisms

We compare our correspondence mechanism against alternative strategies for extending temporal context beyond the default sliding window. Qualitative and quantitative results are shown in Figure 4 and Table 4.

Larger KV Cache window. We increase the sliding window size from 21 latent frames to 42 ($\times 2$), 63 ($\times 3$), and 84 ($\times 4$). On VACE-Benchmark, this yields only marginal consistency improvements (LPIPS decreases from 0.35 to 0.33 and DreamSim from 0.139 to 0.134).

On VACE-Bench, increasing the KV cache size results in slight consistency improvement (LPIPS decreases from 0.35 to 0.33 and DreamSim from 0.139 to 0.134). However, on DL3DV—where sequences are substantially longer, performance degrades sharply in both quality and long-range consistency.

Keyframe retention. Next, instead of sliding window, we retain KV cache of every N -th block ($N=3, 5, 10$), providing a sparse but more complete coverage of the history. While this improves over simply enlarging the window, it remains worse than correspondence-based retrieval in long-range reconstruction (Table 4).

Global Pooling. We maintain a single block of KV tokens as a running cumulative average of all previous blocks, providing constant memory access to the entire history. Although this improves long-range consistency on DL3DV, it degrades visual quality on both datasets (quality scores drop to 0.753 and 0.772).

Random Retrieval. To isolate the role of correspondence matching from the cross-attention layers, we replace the similarity-based retrieval with uniform random sampling from the memory bank. Each query patch receives k randomly selected history patches instead of the top- k most similar, the history creation process remains the same. This significantly degrades long-range reconstruction performance, demonstrating that accurate correspondence matching is essential.

Attention Sink. Recent methods such as LongLive, MotionStream, and Rolling Forcing retain initial frames as an attention sink to provide global context. However, this mechanism cannot capture novel content that appears later in the sequence and is insufficient to prevent identity drift in long videos.

FramePack. FramePack [52] is applied to the task of video to video translation, storing past blocks at geometrically decreasing spatial resolution. Recent blocks maintain full resolution while older blocks are progressively downsampled. While this retains the entire history, aggressive spatial compression reduces detail, leading to weaker long-range reconstruction.

Overall, correspondence-based retrieval outperforms all alternative mechanisms for extending temporal context, achieving stronger long-range consistency without sacrificing visual quality. Please see supplement for additional results.

4.3 Analysis and Ablations

Please see supplement for additional ablation studies, analyses, and results with other conditioning signals.

Importance of Correspondence. First, we show the importance of using correspondence patches through an ablation where we remove the correspondences and disable the correspondence cross-attention layer (Equation 4). As in Table 5, without using correspondence the long context consistency gets significantly worse. The LPIPS reconstruction improves by 37% and 40% on the two datasets. DreamSim and CLIP similarity scores are also similarly effected. Figure 5 shows this qualitatively. The top row shows the generated video when the correspondence is disabled. In this video, the dog starts with a yellow snout, but as the video progresses and the dog re-enters the scene, the snout turns blue. In contrast, our full method with correspondence is shown in the bottom row. Here the dog retains the identity and re-enters the scene with the same yellow snout.

Sensitivity to different distillation methods. Next, we show that our method is not specific to any single distillation method. We apply our method to three different distillation methods: SelfForcing [12], LongLive [46], and RollingForcing [22]. Add our correspondence method improves the consistency across each of the base distillation methods as seen in Table 3.

5 Conclusion and Limitations

We introduce `vid2vid-long`, a correspondence-based memory mechanism for long-range consistent video-to-video translation. By exploiting a structural property of the task, matching regions in the input should match in the output, we replace implicit long-range attention with explicit input-driven retrieval. Our dynamic memory grows with visual novelty rather than length of the video, enabling sub-linear scaling while preserving real-time performance. Across multiple autoregressive distillation backbones, our method consistently improves long-range reconstruction and prevents identity drift without sacrificing visual quality or latency. Compared to alternative context extension strategies, explicit correspondence retrieval proves both more effective and more efficient.

Limitations. Our model struggles on fast videos with large motion, where reliable cross-frame correspondences become harder to establish. It can also fail when the text prompt is inconsistent with the input depth video, since the model must reconcile conflicting semantic and structural conditioning signals.

Acknowledgment. We thank Sheng-Yu Wang, Nupur Kumari, Maxwell Jones, and Xun Huang for their helpful comments and discussion. The project was partially done during an internship with Adobe Research. The project is partially funded by NSF IIS-2239076, NSF ISS-2403303, the IITP grant funded by the Korean Government (MSIT) (No. RS-2024-00457882, National AI Research Lab Project), and the Packard Fellowship.

References

1. Blattmann, A., Dockhorn, T., Kulal, S., Mendelevitch, D., Kilian, M., Lorenz, D., Levi, Y., English, Z., Voleti, V., Letts, A., et al.: Stable video diffusion: Scaling latent video diffusion models to large datasets. arXiv preprint arXiv:2311.15127 (2023) [1](#), [3](#)
2. Chen, B., Martí Monsó, D., Du, Y., Simchowitz, M., Tedrake, R., Sitzmann, V.: Diffusion forcing: Next-token prediction meets full-sequence diffusion. *Advances in Neural Information Processing Systems* **37**, 24081–24125 (2024) [4](#)
3. Chen, W., Ji, Y., Wu, J., Wu, H., Xie, P., Li, J., Xia, X., Xiao, X., Lin, L.: Control-a-video: Controllable text-to-video diffusion models with motion prior and reward feedback learning. arXiv preprint arXiv:2305.13840 (2023) [3](#), [11](#)
4. Dao, T., Gu, A.: Transformers are ssms: Generalized models and efficient algorithms through structured state space duality (2024), <https://arxiv.org/abs/2405.21060> [1](#)
5. Fu, S., Tamir, N., Sundaram, S., Chai, L., Zhang, R., Dekel, T., Isola, P.: Dreamsim: Learning new dimensions of human visual similarity using synthetic data. arXiv preprint arXiv:2306.09344 (2023) [10](#)
6. Gao, Y., Guo, H., Hoang, T., Huang, W., Jiang, L., Kong, F., Li, H., Li, J., Li, L., Li, X., Li, X., Li, Y., Lin, S., Lin, Z., Liu, J., Liu, S., Nie, X., Qing, Z., Ren, Y., Sun, L., Tian, Z., Wang, R., Wang, S., Wei, G., Wu, G., Wu, J., Xia, R., Xiao, F., Xiao, X., Yan, J., Yang, C., Yang, J., Yang, R., Yang, T., Yang, Y., Ye, Z., Zeng, X., Zeng, Y., Zhang, H., Zhao, Y., Zheng, X., Zhu, P., Zou, J., Zuo, F.: Seedance 1.0: Exploring the boundaries of video generation models (2025), <https://arxiv.org/abs/2506.09113> [1](#)
7. Ge, S., Zhang, Y., Liu, L., Zhang, M., Han, J., Gao, J.: Model tells you what to discard: Adaptive kv cache compression for llms. arXiv preprint arXiv:2310.01801 (2023) [4](#)
8. Gu, A., Dao, T.: Mamba: Linear-time sequence modeling with selective state spaces (2024), <https://arxiv.org/abs/2312.00752> [1](#)
9. Guo, H., Jia, Z., Li, J., Li, B., Cai, Y., Wang, J., Li, Y., Lu, Y.: Efficient autoregressive video diffusion with dummy head. arXiv preprint arXiv:2601.20499 (2026) [10](#)
10. HaCohen, Y., Chiprut, N., Brazowski, B., Shalem, D., Moshe, D., Richardson, E., Levin, E., Shiran, G., Zabari, N., Gordon, O., Panet, P., Weissbuch, S., Kulikov, V., Bitterman, Y., Melumian, Z., Bibi, O.: Ltx-video: Realtime video latent diffusion. arXiv preprint arXiv:2501.00103 (2024) [3](#)
11. Henschel, R., Khachatryan, L., Poghosyan, H., Hayrapetyan, D., Tadevosyan, V., Wang, Z., Navasardyan, S., Shi, H.: Streamingt2v: Consistent, dynamic, and extendable long video generation from text. In: *Proceedings of the Computer Vision and Pattern Recognition Conference*. pp. 2568–2577 (2025) [4](#)
12. Huang, X., Li, Z., He, G., Zhou, M., Shechtman, E.: Self forcing: Bridging the train-test gap in autoregressive video diffusion. arXiv preprint arXiv:2506.08009 (2025) [1](#), [5](#), [12](#), [15](#)
13. Huang, Z., He, Y., Yu, J., Zhang, F., Si, C., Jiang, Y., Zhang, Y., Wu, T., Jin, Q., Chanpaisit, N., Wang, Y., Chen, X., Wang, L., Lin, D., Qiao, Y., Liu, Z.: VBench: Comprehensive benchmark suite for video generative models. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition* (2024) [9](#)

14. Jiang, Z., Han, Z., Mao, C., Zhang, J., Pan, Y., Liu, Y.: Vace: All-in-one video creation and editing. In: Proceedings of the IEEE/CVF International Conference on Computer Vision. pp. 17191–17202 (2025) [3](#), [4](#), [9](#), [10](#), [12](#)
15. Kerbl, B., Kopanas, G., Leimkühler, T., Drettakis, G.: 3d gaussian splatting for real-time radiance field rendering. *ACM Transactions on Graphics* **42**(4) (July 2023), <https://repo-sam.inria.fr/fungraph/3d-gaussian-splatting/> [10](#)
16. Kong, W., Tian, Q., Zhang, Z., Min, R., Dai, Z., Zhou, J., Xiong, J., Li, X., Wu, B., Zhang, J., et al.: Hunyuanvideo: A systematic framework for large video generative models. *arXiv preprint arXiv:2412.03603* (2024) [1](#)
17. Kumari, N., Yin, X., Zhu, J.Y., Misra, I., Azadi, S.: Generating multi-image synthetic data for text-to-image customization. In: IEEE International Conference on Computer Vision (ICCV) (2025) [3](#)
18. Labs, B.F.: Flux. <https://github.com/black-forest-labs/flux> (2024) [3](#)
19. Li, D., Li, J., Hoi, S.: Blip-diffusion: Pre-trained subject representation for controllable text-to-image generation and editing. *Advances in Neural Information Processing Systems* **36**, 30146–30166 (2023) [3](#)
20. Lin, S., Xia, X., Ren, Y., Yang, C., Xiao, X., Jiang, L.: Diffusion adversarial post-training for one-step video generation (2025), <https://arxiv.org/abs/2501.08316> [1](#)
21. Ling, L., Sheng, Y., Tu, Z., Zhao, W., Xin, C., Wan, K., Yu, L., Guo, Q., Yu, Z., Lu, Y., et al.: D13dv-10k: A large-scale scene dataset for deep learning-based 3d vision. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 22160–22169 (2024) [9](#)
22. Liu, K., Hu, W., Xu, J., Shan, Y., Lu, S.: Rolling forcing: Autoregressive long video diffusion in real time. *arXiv preprint arXiv:2509.25161* (2025) [1](#), [5](#), [7](#), [10](#), [12](#), [15](#)
23. Liu, X., Tang, Z., Dong, P., Li, Z., Liu, Y., Li, B., Hu, X., Chu, X.: Chunkkv: Semantic-preserving kv cache compression for efficient long-context llm inference. *arXiv preprint arXiv:2502.00299* (2025) [4](#)
24. Loshchilov, I., Hutter, F.: Decoupled weight decay regularization. *arXiv preprint arXiv:1711.05101* (2017) [11](#)
25. Mou, C., Wang, X., Xie, L., Wu, Y., Zhang, J., Qi, Z., Shan, Y., Qie, X.: T2i-adapter: Learning adapters to dig out more controllable ability for text-to-image diffusion models. *arXiv preprint arXiv:2302.08453* (2023) [3](#)
26. NVIDIA, Ali, A., Bai, J., Bala, M., Balaji, Y., Blakeman, A., Cai, T., Cao, J., Cao, T., Cha, E., Chao, Y.W., Chattopadhyay, P., Chen, M., Chen, Y., Chen, Y., Cheng, S., Cui, Y., Diamond, J., Ding, Y., Fan, J., Fan, L., Feng, L., Ferroni, F., Fidler, S., Fu, X., Gao, R., Ge, Y., Gu, J., Gupta, A., Gururani, S., El Hanafi, I., Hassani, A., Hao, Z., Huffman, J., Jang, J., Jannaty, P., Kautz, J., Lam, G., Li, X., Li, Z., Liao, M., Lin, C.H., Lin, T.Y., Lin, Y.C., Ling, H., Liu, M.Y., Liu, X., Lu, Y., Luo, A., Ma, Q., Mao, H., Mo, K., Nah, S., Narang, Y., Panaskar, A., Pavao, L., Pham, T., Ramezanali, M., Reda, F., Reed, S., Ren, X., Shao, H., Shen, Y., Shi, S., Song, S., Stefaniak, B., Sun, S., Tang, S., Tasmeen, S., Tchapmi, L., Tseng, W.C., Varghese, J., Wang, A.Z., Wang, H., Wang, H., Wang, H., Wang, T.C., Wei, F., Xu, J., Yang, D., Yang, X., Ye, H., Ye, S., Zeng, X., Zhang, J., Zhang, Q., Zheng, K., Zhu, A., Zhu, Y.: World simulation with video foundation models for physical ai. *arXiv preprint arXiv:2511.00062* (2025) [3](#)
27. Parmar, G., Park, T., Narasimhan, S., Zhu, J.Y.: One-step image translation with text-to-image models. *arXiv preprint arXiv:2403.12036* (2024) [3](#)
28. Parmar, G., Patashnik, O., Wang, K.C., Ostashev, D., Narasimhan, S., Zhu, J.Y., Cohen-Or, D., Aberman, K.: Object-level visual prompts for compositional image

- generation. In: Proceedings of the SIGGRAPH Asia 2025 Conference Papers. pp. 1–12 (2025) [3](#)
29. Peebles, W., Xie, S.: Scalable diffusion models with transformers. In: Proceedings of the IEEE/CVF international conference on computer vision. pp. 4195–4205 (2023) [4](#)
 30. Peng, X., Zheng, Z., Shen, C., Young, T., Guo, X., Wang, B., Xu, H., Liu, H., Jiang, M., Li, W., Wang, Y., Ye, A., Ren, G., Ma, Q., Liang, W., Lian, X., Wu, X., Zhong, Y., Li, Z., Gong, C., Lei, G., Cheng, L., Zhang, L., Li, M., Zhang, R., Hu, S., Huang, S., Wang, X., Zhao, Y., Wang, Y., Wei, Z., You, Y.: Open-sora 2.0: Training a commercial-level video generation model in \$200k. arXiv preprint arXiv:2503.09642 (2025) [1](#)
 31. Po, R., Nitzan, Y., Zhang, R., Chen, B., Dao, T., Shechtman, E., Wetzstein, G., Huang, X.: Long-context state-space video world models (2025), <https://arxiv.org/abs/2505.20171> [2](#)
 32. Radford, A., Kim, J.W., Hallacy, C., Ramesh, A., Goh, G., Agarwal, S., Sastry, G., Askell, A., Mishkin, P., Clark, J., et al.: Learning transferable visual models from natural language supervision. In: International conference on machine learning. pp. 8748–8763. PmLR (2021) [10](#)
 33. Rombach, R., Blattmann, A., Lorenz, D., Esser, P., Ommer, B.: High-resolution image synthesis with latent diffusion models. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 10684–10695 (2022) [3](#), [4](#)
 34. Salimans, T., Ho, J.: Progressive distillation for fast sampling of diffusion models. arXiv preprint arXiv:2202.00512 (2022) [3](#)
 35. Salimans, T., Mensink, T., Heek, J., Hoogeboom, E.: Multistep distillation of diffusion models via moment matching. Advances in Neural Information Processing Systems **37**, 36046–36070 (2024) [3](#)
 36. Sauer, A., Lorenz, D., Blattmann, A., Rombach, R.: Adversarial diffusion distillation. In: European Conference on Computer Vision. pp. 87–103. Springer (2024) [3](#)
 37. Shin, J., Li, Z., Zhang, R., Zhu, J.Y., Park, J., Shechtman, E., Huang, X.: MotionStream: Real-Time Video Generation with Interactive Motion Controls. In: Proceedings of the International Conference on Learning Representations (ICLR (2026) [1](#)
 38. Song, Y., Dhariwal, P., Chen, M., Sutskever, I.: Consistency models. In: International Conference on Machine Learning. pp. 32211–32252. PMLR (2023) [3](#)
 39. Voynov, A., Aberman, K., Cohen-Or, D.: Sketch-guided text-to-image diffusion models. In: ACM SIGGRAPH 2023 conference proceedings. pp. 1–11 (2023) [3](#)
 40. Wan, T., Wang, A., Ai, B., Wen, B., Mao, C., Xie, C.W., Chen, D., Yu, F., Zhao, H., Yang, J., et al.: Wan: Open and advanced large-scale video generative models. arXiv preprint arXiv:2503.20314 (2025) [1](#), [3](#), [4](#)
 41. Wang, J., Ye, L., Lu, T., Xiao, J., Zhang, J., Guo, Y., Liu, X., Chellappa, R., Peng, C., Yuille, A., Chen, J.: Evoworld: Evolving panoramic world generation with explicit 3d memory (2025), <https://arxiv.org/abs/2510.01183> [2](#)
 42. Wang, X., Yuan, H., Zhang, S., Chen, D., Wang, J., Zhang, Y., Shen, Y., Zhao, D., Zhou, J.: Videocomposer: Compositional video synthesis with motion controllability. Advances in Neural Information Processing Systems **36**, 7594–7611 (2023) [11](#)
 43. Wu, C., Li, J., Zhou, J., Lin, J., Gao, K., Yan, K., ming Yin, S., Bai, S., Xu, X., Chen, Y., Chen, Y., Tang, Z., Zhang, Z., Wang, Z., Yang, A., Yu, B., Cheng, C.,

- Liu, D., Li, D., Zhang, H., Meng, H., Wei, H., Ni, J., Chen, K., Cao, K., Peng, L., Qu, L., Wu, M., Wang, P., Yu, S., Wen, T., Feng, W., Xu, X., Wang, Y., Zhang, Y., Zhu, Y., Wu, Y., Cai, Y., Liu, Z.: Qwen-image technical report (2025), <https://arxiv.org/abs/2508.02324> 3
44. Xiao, G., Tian, Y., Chen, B., Han, S., Lewis, M.: Efficient streaming language models with attention sinks. arXiv (2023) 1
 45. Xiao, Z., Lan, Y., Zhou, Y., Ouyang, W., Yang, S., Zeng, Y., Pan, X.: Worldmem: Long-term consistent world simulation with memory (2025), <https://arxiv.org/abs/2504.12369> 2
 46. Yang, S., Huang, W., Chu, R., Xiao, Y., Zhao, Y., Wang, X., Li, M., Xie, E., Chen, Y., Lu, Y., et al.: Longlive: Real-time interactive long video generation. arXiv preprint arXiv:2509.22622 (2025) 1, 4, 5, 12, 15
 47. Ye, H., Zhang, J., Liu, S., Han, X., Yang, W.: Ip-adapter: Text compatible image prompt adapter for text-to-image diffusion models (2023) 3
 48. Yin, T., Gharbi, M., Park, T., Zhang, R., Shechtman, E., Durand, F., Freeman, B.: Improved distribution matching distillation for fast image synthesis. *Advances in neural information processing systems* **37**, 47455–47487 (2024) 3
 49. Yin, T., Gharbi, M., Zhang, R., Shechtman, E., Durand, F., Freeman, W.T., Park, T.: One-step diffusion with distribution matching distillation. In: *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*. pp. 6613–6623 (2024) 3
 50. Yin, T., Zhang, Q., Zhang, R., Freeman, W.T., Durand, F., Shechtman, E., Huang, X.: From slow bidirectional to fast autoregressive video diffusion models. In: *CVPR* (2025) 1, 4, 5
 51. Yu, J., Bai, J., Qin, Y., Liu, Q., Wang, X., Wan, P., Zhang, D., Liu, X.: Context as memory: Scene-consistent interactive long video generation with memory retrieval. arXiv preprint arXiv:2506.03141 (2025) 2
 52. Zhang, L., Cai, S., Li, M., Wetzstein, G., Agrawala, M.: Frame context packing and drift prevention in next-frame-prediction video diffusion models. In: *The Thirty-ninth Annual Conference on Neural Information Processing Systems* (2025) 11, 14
 53. Zhang, L., Rao, A., Agrawala, M.: Adding conditional control to text-to-image diffusion models (2023) 3, 7
 54. Zhang, R., Isola, P., Efros, A.A., Shechtman, E., Wang, O.: The unreasonable effectiveness of deep features as a perceptual metric. In: *CVPR* (2018) 10
 55. Zhang, Y., Wei, Y., Jiang, D., Zhang, X., Zuo, W., Tian, Q.: Controlvideo: Training-free controllable text-to-video generation (2023), <https://arxiv.org/abs/2305.13077> 3, 11
 56. Zheng, Z., Peng, X., Yang, T., Shen, C., Li, S., Liu, H., Zhou, Y., Li, T., You, Y.: Open-sora: Democratizing efficient video production for all. arXiv preprint arXiv:2412.20404 (2024) 1