

Zen-Video-I2V: Photorealistic Image-to-Video Generation

with Controllable Motion Dynamics

Technical Whitepaper v2025.04

Zach Kelling
Zen LM Research Team
research@zenlm.org
papers.zenlm.org

April 2025

Abstract

Zen-Video-I2V is a 14 billion parameter image-to-video generation model that animates static input images into fluid video sequences up to 10 seconds at 1080p resolution, with controllable motion direction, intensity, and style. The model achieves a Fréchet Video Distance (FVD) of 289 on UCF-101, a Warping Error of 0.023 (measuring temporal warp consistency), and 91.3% temporal consistency on the Video Quality Assessment benchmark. Zen-Video-I2V supports both single-image animation (producing plausible dynamic motion from scene content) and image-sequence interpolation (generating smooth transitions between provided keyframes). This paper describes the model architecture, motion conditioning mechanisms, training methodology, and comprehensive video quality evaluations.

Contents

1	Introduction	3
1.1	Model Overview	3
2	Architecture	3
2.1	Diffusion Transformer with Temporal Attention	3
2.2	Dual-Path Motion Conditioning	4
2.3	Image Conditioning	4
2.4	Physics-Informed Regularization	4
3	Training Methodology	5
3.1	Training Data	5
3.2	Training Curriculum	5
3.3	Noise Schedule	5
4	Evaluation	6
4.1	Video Quality Benchmarks	6
4.2	Human Evaluation	6
4.3	Motion Control Evaluation	6
4.4	Resolution and Duration Scaling	6
4.5	Image-Sequence Interpolation	7
4.6	Per-Domain Performance	7

5	Deployment	7
5.1	Inference Requirements	7
5.2	API Interface	7
6	Safety Considerations	8
6.1	Identity and Consent	8
6.2	Deepfake Mitigation	8
7	Related Work	8
8	Conclusion	8

1 Introduction

The transition from image generation to video generation requires solving two qualitatively different problems. Image generation learns the distribution of plausible spatial configurations. Video generation must additionally learn the distribution of plausible temporal dynamics: how objects in a scene move, deform, and interact over time while maintaining physical and visual coherence across frames.

Image-to-video generation constrains this problem by conditioning on a real input image, fixing the spatial configuration and shifting the task to learning what temporal evolutions are consistent with that fixed starting state. This constraint is practically valuable: it allows users to start from real photographs, generated images, or concept art and animate them into video without requiring full video synthesis from text prompts.

Zen-Video-I2V advances image-to-video generation through three contributions:

1. **Dual-path motion conditioning:** Separate conditioning pathways for camera motion (dolly, pan, tilt, zoom, orbit) and scene motion (object-level dynamics), enabling independent control of each.
2. **Physical plausibility constraints:** A physics-informed regularizer trained on simulation data penalizes animations that violate basic physical constraints (rigid body motion, fluid dynamics, cloth behavior).
3. **Temporal attention at scale:** A full-resolution temporal attention mechanism (not limited to downsampled feature maps) operating across all frames simultaneously, producing high temporal coherence even in long-form (10 second) generations.

1.1 Model Overview

Table 1: Zen-Video-I2V Model Specification

Parameter	Value
Architecture	Diffusion Transformer (DiT) with temporal attention
Total Parameters	14B
Output Resolution	720p (1280×720) or 1080p (1920×1080)
Output Duration	2–10 seconds (configurable)
Frame Rate	24 fps (configurable: 8, 16, 24 fps)
FVD on UCF-101	289
Warping Error	0.023
Temporal Consistency	91.3%
Version	v2025.04
Release Date	April 2025

2 Architecture

2.1 Diffusion Transformer with Temporal Attention

Zen-Video-I2V is built on a Diffusion Transformer (DiT) [1] backbone extended with full-resolution temporal attention. The model operates in the latent space of a pretrained video VAE [2] that encodes video frames at $8\times$ spatial compression with 4-frame temporal compression, reducing a 1080p 240-frame video to a $135\times 240\times 60$ -frame latent volume.

The transformer architecture alternates between spatial attention (across spatial positions within a frame) and temporal attention (across frames at fixed spatial positions):

$$h^{\text{spatial}} = \text{SelfAttn}_{\text{spatial}}(h) \quad (1)$$

$$h^{\text{temporal}} = \text{SelfAttn}_{\text{temporal}}(h^{\text{spatial}}) \quad (2)$$

Temporal attention is applied at full latent resolution, not downsampled, which enables the model to maintain fine-grained temporal coherence in high-frequency details (hair, fabric texture, water surface). The computational cost of full-resolution temporal attention is managed through FlashAttention-3 [3] and temporal sparse attention patterns for long (10s) generations.

2.2 Dual-Path Motion Conditioning

Motion control is provided through two independent conditioning pathways injected via cross-attention:

Camera Motion Pathway: Encodes camera trajectory as a sequence of 6-DoF camera poses $\{(R_t, T_t)\}_{t=1}^F$ where $R_t \in SO(3)$ is rotation and $T_t \in \mathbb{R}^3$ is translation. The camera trajectory is encoded by a lightweight 6-layer MLP producing a 512-dimensional embedding per frame.

Scene Motion Pathway: Encodes object-level dynamics as an optical flow field $\{u_t\}_{t=1}^F$ where $u_t \in \mathbb{R}^{H \times W \times 2}$ is the dense flow from frame $t - 1$ to frame t . The flow field is encoded by a small 4-layer U-Net producing a 256-dimensional embedding at the latent spatial resolution.

The two pathways are combined through learned gating:

$$c_{\text{motion}} = \sigma(g) \cdot c_{\text{camera}} + (1 - \sigma(g)) \cdot c_{\text{scene}} \quad (3)$$

where g is a learned gating scalar conditioned on the generation context. In practice, the gate learns to balance camera and scene motion conditioning based on the input image content: outdoor scenes with clear horizon lines favor camera conditioning, while close-up portraits or product shots favor scene motion conditioning.

2.3 Image Conditioning

The input image is encoded by the video VAE and injected into the generation process as:

1. **First-frame conditioning:** The encoded input image is fixed as the first latent frame, with the model trained to denoise subsequent frames conditioned on this fixed start state.
2. **Global image embedding:** A CLIP image encoder extracts a semantic embedding of the input image that conditions all transformer layers via cross-attention, maintaining semantic consistency throughout the generated video.
3. **Pixel-level conditioning:** For high-fidelity reconstruction of the input image’s specific textures and details, ControlNet-style feature injection [4] copies early encoder features directly into the decoder path.

2.4 Physics-Informed Regularization

Physical plausibility is enforced through an auxiliary physics regularizer trained on simulation data from three domains:

1. **Rigid body:** Motion from a physics engine simulator covering 50K object types.
2. **Fluid:** Water surface simulation data from a fluid dynamics solver.
3. **Cloth and soft body:** Garment simulation covering 20K clothing styles.

The regularizer R_ψ takes a generated optical flow sequence and predicts a physics plausibility score. This score is used as an auxiliary loss during training:

$$\mathcal{L}_{\text{phys}} = -\log R_\psi(u_{1:F}) \quad (4)$$

The physics regularizer is frozen during video model training (trained separately on simulation data only), serving as a learned prior that penalizes motions that violate physical constraints without requiring differentiable physics simulation.

3 Training Methodology

3.1 Training Data

Table 2: Training Data Composition

Source	Video Hours	Fraction
Licensed stock footage	2,400,000	48.0%
Documentary and nature video	800,000	16.0%
Film and television (licensed)	600,000	12.0%
User-generated video (consented)	500,000	10.0%
Physics simulation renders	400,000	8.0%
3D rendering and animation	300,000	6.0%
Total	5,000,000	100%

All training data was filtered for quality (resolution $\geq 720\text{p}$, motion blur $<$ threshold, compression artifact score $<$ threshold) and processed to remove frames with faces not covered by consent agreements.

3.2 Training Curriculum

Training proceeds through three stages:

Stage 1 – Short clip pretraining (2 seconds, 48 frames): The model learns basic temporal coherence and motion dynamics on short clips, where the temporal attention computation is manageable at full resolution.

Stage 2 – Mid-length extension (5 seconds, 120 frames): The context length is extended through curriculum continuation training, where the model is initialized from Stage 1 weights and fine-tuned on longer clips with sparse temporal attention for positions beyond 48 frames.

Stage 3 – Long-form fine-tuning (10 seconds, 240 frames): Final stage fine-tuning on high-quality long clips (filtered for production quality), with the sparse temporal attention pattern expanded to cover the full temporal extent.

3.3 Noise Schedule

Zen-Video-I2V uses a continuous-time diffusion process with a learned noise schedule that is adaptive per-frame: the noise level at frame t is conditioned on the frame index, applying higher signal-to-noise ratios for early frames (closer to the conditioning image) and allowing more variance for later frames (more uncertain motion trajectories):

$$\sigma_t^{(f)} = \sigma_{\min} + \frac{f}{F}(\sigma_{\max} - \sigma_{\min}), \quad f \in \{1, \dots, F\} \quad (5)$$

This schedule encourages the model to generate temporally consistent short-range dynamics while maintaining temporal diversity over longer horizons.

4 Evaluation

4.1 Video Quality Benchmarks

Table 3: Video Quality Benchmark Results

Model	FVD (UCF-101) ↓	Warping Error ↓	Temporal Consistency ↑	PSNR ↑
Zen-Video-I2V (14B)	289	0.023	91.3%	
Comparable I2V model A	341	0.031	87.6%	
Comparable I2V model B	387	0.038	84.2%	
Comparable I2V model C	312	0.027	89.1%	

4.2 Human Evaluation

Table 4: Human Evaluation MOS (1–5 scale, 500 clips rated by 3 raters each)

Model	Realism	Temporal Coherence	Motion Quality	Overall
Zen-Video-I2V (14B)	4.1	4.0	3.9	4.0
Comparable model A	3.7	3.5	3.4	3.5
Comparable model B	3.4	3.3	3.1	3.3

4.3 Motion Control Evaluation

We evaluate the precision of the dual-path motion control on a benchmark of 200 clips with specified camera trajectories and object motion targets:

Table 5: Motion Control Accuracy

Motion Type	Direction Accuracy	Intensity Correlation
Camera pan / tilt	94.2%	0.91
Camera dolly (zoom equivalent)	91.7%	0.88
Camera orbit	88.3%	0.84
Object translation	87.4%	0.82
Object deformation (cloth, fluid)	79.1%	0.76
Average	88.1%	0.84

4.4 Resolution and Duration Scaling

Table 6: Quality vs. Resolution and Duration

Resolution	Duration	FVD	Temporal Consistency	Gen Time
720p	2s	241	93.7%	18s
720p	5s	267	92.1%	44s
720p	10s	289	91.3%	87s
1080p	2s	253	93.1%	31s
1080p	5s	278	91.4%	74s
1080p	10s	302	90.7%	148s

Generation times measured on 4×H100 80GB SXM5 (FP8, 50 DDIM steps).

4.5 Image-Sequence Interpolation

When provided with multiple keyframe images, Zen-Video-I2V generates smooth interpolated video transitions:

Table 7: Keyframe Interpolation Quality

Metric	Zen-Video-I2V	Linear Interpolation	Morph baseline
Perceptual quality MOS	4.2/5	2.1/5	3.1/5
Boundary frame fidelity (PSNR)	41.3 dB	44.2 dB	38.7 dB
Intermediate plausibility MOS	4.0/5	1.8/5	2.9/5

4.6 Per-Domain Performance

Table 8: FVD by Content Domain

Domain	Zen-Video-I2V FVD	Best Comparable FVD
Nature / landscape	243	298
Human portrait	318	394
Animal motion	271	337
Urban / architecture	261	319
Product / object	234	287
Abstract / artistic	312	381

Human portrait is the most challenging domain due to the perceptual sensitivity to uncanny motion artifacts in faces; Zen-Video-I2V’s physics-informed regularizer is less effective for biological motion than for physical objects.

5 Deployment

5.1 Inference Requirements

Table 9: Inference Hardware Requirements

Config	Hardware	Resolution	Gen Time (5s)
High quality	4 × H100 80GB	1080p	74s
Standard	4 × H100 80GB	720p	44s
Fast (reduced steps)	2 × H100 80GB	720p	28s

5.2 API Interface

Zen-Video-I2V is accessible through the Zen LM API with the following parameters:

- **image**: Input image (JPEG/PNG, min 512px on shorter edge).
- **duration**: Output duration in seconds (2.0–10.0, default 5.0).
- **resolution**: Output resolution (“720p” or “1080p”).
- **camera_motion**: Camera trajectory type (“static”, “dolly_in”, “dolly_out”, “pan_left”, “pan_right”, “orbit”, or custom 6-DoF trajectory).

- `motion_intensity`: Overall motion scale (0.0–1.0, default 0.5).
- `fps`: Output frame rate (8, 16, or 24).
- `seed`: Reproducibility seed for deterministic generation.

6 Safety Considerations

6.1 Identity and Consent

Zen-Video-I2V includes a face detection classifier that identifies human faces in input images. Generation involving detected faces requires:

- Explicit consent acknowledgment from the API caller that they have rights to animate the depicted individual.
- Invisible watermarking applied to all generated video (C2PA-compatible provenance metadata).
- Rate limiting and audit logging for accounts generating face-containing animations.

6.2 Deepfake Mitigation

Zen-Video-I2V applies Content Authenticity Initiative (CAI) metadata to all generated videos, marking them as AI-generated in a tamper-evident way. The model is also trained to degrade gracefully (reduce quality) on inputs that appear to be composite images designed for deceptive impersonation (face-swapped images, morphed composites).

7 Related Work

Image-to-video generation has advanced through GAN-based [5], flow-based [6], and diffusion-based [7, 8] approaches. Temporal attention mechanisms in video diffusion have been studied in [9, 10]. Motion controllability has been explored through trajectory conditioning [11] and optical flow guidance [12]. Zen-Video-I2V advances this line of work through dual-path motion conditioning, physics-informed regularization, and full-resolution temporal attention at 14B parameter scale.

8 Conclusion

Zen-Video-I2V establishes new state-of-the-art results on image-to-video generation benchmarks (FVD 289, Warping Error 0.023, Temporal Consistency 91.3%) while introducing dual-path motion control and physics-informed regularization that improve the plausibility and controllability of generated animations. At 14B parameters, the model generates photorealistic 10-second 1080p video from single images, suitable for creative production, product visualization, and interactive media applications. The physics-informed regularizer and full-resolution temporal attention architecture provide strong foundations for future extensions to longer-form and higher-resolution video generation.

References

- [1] Peebles, W. & Xie, S. (2023). Scalable Diffusion Models with Transformers. ICCV 2023.
- [2] Kingma, D.P. & Welling, M. (2014). Auto-Encoding Variational Bayes. ICLR 2014.

- [3] Shah, J. et al. (2024). FlashAttention-3: Fast and Accurate Attention with Asynchrony and Low-precision. arXiv:2407.08608.
- [4] Zhang, L. et al. (2023). Adding Conditional Control to Text-to-Image Diffusion Models. ICCV 2023.
- [5] Vondrick, C. et al. (2016). Generating Videos with Scene Dynamics. NeurIPS 2016.
- [6] Holynski, A. et al. (2021). Animating Pictures with Eulerian Motion Fields. CVPR 2021.
- [7] Chai, L. et al. (2023). StableVideo: Text-driven Consistency-aware Diffusion Video Editing. ICCV 2023.
- [8] Blattmann, A. et al. (2023). Stable Video Diffusion: Scaling Latent Video Diffusion Models to Large Datasets. arXiv:2311.15127.
- [9] Wu, J.Z. et al. (2023). Tune-A-Video: One-Shot Tuning of Image Diffusion Models for Text-to-Video Generation. ICCV 2023.
- [10] He, Y. et al. (2023). VideoCrafter1: Open Diffusion Models for High-Quality Video Generation. arXiv:2310.19512.
- [11] Wang, S. et al. (2024). DragNUWA: Fine-grained Control in Video Generation by Integrating Text, Image, and Trajectory. ECCV 2024.
- [12] Wang, Z. et al. (2024). MotionCtrl: A Unified and Flexible Motion Controller for Video Generation. SIGGRAPH 2024.