

Zen AI Model Family

Zen-Director

Video Scene Generation, Storyboarding, and Cinematic Direction

Technical Report v2025.01

Hanzo AI Research Team* Zoo Labs Foundation[†]

January 2025

Abstract

We present **Zen-Director**, a 7-billion parameter vision-language model specialized for video scene generation, cinematographic planning, and multi-shot storyboard synthesis. Built on the Zen MoDE (Mixture of Distilled Experts) architecture with a temporal transformer extension, Zen-Director understands and generates structured cinematic descriptions—shot lists, scene compositions, camera movements, and narrative arcs—that downstream video generation models can consume directly. The model achieves BLEU-4 of 0.423 on the VideoCaption benchmark, 81.2% accuracy on CinematicQA (our novel evaluation of cinematographic knowledge), and produces storyboards rated 4.1/5.0 by professional cinematographers in blind evaluation. Zen-Director bridges the gap between natural language creative intent and the precise technical specifications required to direct AI video generation systems at professional quality.

Contents

1	Introduction	3
1.1	Model Overview	3
1.2	Key Capabilities	3
2	Architecture	3
2.1	Zen MoDE 7B Language Backbone	3
2.2	Visual Encoder	4
2.3	Temporal Transformer Extension	4
2.4	Shot Plan Generation	4
3	Training	5
3.1	Dataset	5
3.2	Training Protocol	5
4	Evaluation	5
4.1	VideoCaption Benchmark	5
4.2	CinematicQA	5
4.3	Professional Storyboard Evaluation	6
4.4	Downstream Video Generation Quality	6

*research@hanzo.ai

[†]foundation@zoo.ngo

5	Applications	6
5.1	AI Film Production Pipeline	6
5.2	Film Education	7
5.3	Video Game Cutscene Direction	7
6	Related Work	7
7	Conclusion	7

1 Introduction

The rapid maturation of text-to-video generation systems has exposed a critical bottleneck: the gap between a creator’s high-level creative intent and the precise technical specifications these systems require to produce coherent, professional-quality video. Current generation workflows demand that users specify camera angles, focal lengths, lighting conditions, shot durations, and scene transitions—knowledge traditionally held by professional cinematographers and directors.

Zen-Director addresses this bottleneck as a **directorial intelligence**: a model that understands narrative intent and translates it into actionable cinematic specifications. Rather than generating video pixels directly, Zen-Director generates structured **shot plans**: hierarchical scene descriptions that video generation systems consume as structured prompts.

1.1 Model Overview

Property	Value
Parameters	7B
Architecture	Zen MoDE 7B + Temporal Transformer
Context Length	32K tokens + 512 video frames
Visual Encoder	ViT-L/14 (336px)
Temporal Depth	12 temporal transformer layers
Shot Plan Format	Structured JSON + prose description
Training Data	8.2M scene-annotation pairs, 1.4M film scripts

Table 1: Zen-Director Model Specifications

1.2 Key Capabilities

- **Scene storyboarding**: Convert natural language scene descriptions into complete shot-by-shot storyboards with camera specifications.
- **Shot composition**: Recommend and generate compositional guidelines (rule of thirds, leading lines, depth of field) for each shot.
- **Narrative arc planning**: Structure multi-scene video narratives with consistent pacing, tension arcs, and visual motifs.
- **Cinematic vocabulary**: Understand and generate industry-standard cinematographic terminology (establishing shot, dolly zoom, rack focus, etc.).
- **Video comprehension**: Analyze existing video clips and generate directorial notes describing their cinematographic techniques.

2 Architecture

2.1 Zen MoDE 7B Language Backbone

The language backbone is Zen MoDE at 7B scale: 28 transformer layers, 28 attention heads with grouped-query attention (4 KV heads), and MoE feed-forward networks (4 experts, top-2 routing). This provides strong natural language understanding for parsing creative briefs and generating rich cinematic descriptions.

2.2 Visual Encoder

A ViT-L/14 vision encoder operating at 336px resolution encodes reference images and video frames into 256 visual tokens per frame. A two-layer MLP projection maps visual tokens into the language model’s embedding space.

2.3 Temporal Transformer Extension

The key architectural innovation is a 12-layer temporal transformer that operates *across* frames rather than within them. Given T frames, each encoded to 256 tokens, the temporal transformer attends over the time dimension to build a coherent representation of motion, continuity, and scene evolution:

$$H_{\text{temporal}} = \text{TransformerEncoder}([h_1^{CLS}, h_2^{CLS}, \dots, h_T^{CLS}]) \quad (1)$$

where h_t^{CLS} is the CLS token representation of frame t from the visual encoder. The temporal representation is concatenated with text tokens before the language model’s cross-attention layers.

2.4 Shot Plan Generation

Zen-Director generates shot plans as structured JSON objects:

```
1 {
2   "scene": {
3     "id": "s01",
4     "location": "rain-soaked_rooftop, _night",
5     "mood": "tense, _noir",
6     "duration_sec": 45
7   },
8   "shots": [
9     {
10      "id": "s01_001",
11      "type": "establishing",
12      "camera": {"position": "wide", "angle": "high_angle", "move": "
13        slow_push_in"},
14      "focal_length": "24mm",
15      "subject": "city_skyline_with_protagonist_silhouette",
16      "duration_sec": 8,
17      "lighting": "practical_neon, _rain_reflections",
18      "notes": "City_should_feel_overwhelming_relative_to_protagonist"
19    }
20  ]
21 }
```

Listing 1: Shot Plan Schema

3 Training

3.1 Dataset

Source	Samples	Proportion	Content
Film scripts + frames	2,400,000	29.3%	Script-to-scene alignment
Cinematography textbooks	180,000	2.2%	Technical knowledge
Film criticism corpus	820,000	10.0%	Aesthetic analysis
Video-caption pairs	3,100,000	37.8%	Visual understanding
Storyboard collections	1,700,000	20.7%	Shot plan examples
Total	8,200,000	100%	

Table 2: Zen-Director Training Data

3.2 Training Protocol

Stage 1 – Visual encoder alignment (20K steps): The ViT encoder and MLP projection are trained to align visual representations with cinematic language descriptions.

Stage 2 – Temporal pretraining (40K steps): The temporal transformer is pretrained on video sequences with a masked-frame prediction objective.

Stage 3 – Directorial SFT (60K steps): The full model is fine-tuned on storyboard generation, scene description, and cinematographic Q&A tasks jointly.

Stage 4 – RLHF from cinematographers (15K steps): A reward model trained on 50,000 pairwise comparisons from professional cinematographers (recruited via film school partnerships) is used to further refine shot plan quality via PPO.

4 Evaluation

4.1 VideoCaption Benchmark

We evaluate video description quality on a held-out set of 5,000 film clips spanning 20 genres.

Model	BLEU-4	METEOR	CIDEr	ROUGE-L
BLIP-2	0.281	0.314	0.872	0.512
InstructBLIP	0.308	0.341	0.941	0.534
VideoChat2	0.352	0.374	1.021	0.561
Video-LLaMA2	0.387	0.402	1.103	0.584
Zen-Director	0.423	0.441	1.187	0.612

Table 3: VideoCaption Benchmark Results (higher is better)

4.2 CinematicQA

CinematicQA is a novel benchmark we introduce comprising 2,000 multiple-choice questions testing cinematographic knowledge: shot types, camera movements, lighting techniques, editing principles, and genre conventions. Questions were authored by five professional cinematographers.

Model	Accuracy (%)
GPT-4o (zero-shot)	68.3
Claude 3.5 Sonnet	71.4
Gemini 1.5 Pro	66.8
Specialist fine-tuned 7B	74.2
Zen-Director 7B	81.2

Table 4: CinematicQA Accuracy (%)

4.3 Professional Storyboard Evaluation

Twenty professional cinematographers evaluated storyboards generated from 100 scene descriptions, rating each on a 5-point scale across five dimensions.

Dimension	Mean Score (1–5)
Technical accuracy of shot specifications	4.3
Narrative coherence across shots	4.1
Creative quality / originality	3.9
Pacing appropriateness	4.0
Overall directorial vision	4.1
Overall Mean	4.1

Table 5: Professional Cinematographer Evaluation (N=20 evaluators, 100 storyboards)

4.4 Downstream Video Generation Quality

We evaluate whether Zen-Director shot plans improve final video quality when used as structured prompts for a text-to-video generation system. Using the same creative brief:

Prompt Method	Human Preference (%)	FVD ↓
Raw creative brief (baseline)	18.4%	412
Manual cinematographer spec	42.1%	287
Zen-Director shot plan	39.5%	298

Table 6: Downstream Video Quality with Zen-Director Shot Plans

Zen-Director reaches 94% of manual cinematographer performance at a fraction of the cost and time, validating its role as an effective creative intermediary.

5 Applications

5.1 AI Film Production Pipeline

Zen-Director is integrated into the Hanzo AI film production pipeline as the directorial layer between a human creative brief and the Zen-Video generation system. A typical workflow:

1. Human writes a scene description in natural language.
2. Zen-Director generates a shot plan JSON with full cinematographic specs.
3. Human reviews and optionally edits the shot plan (typically 3–5 minutes).

4. Zen-Video consumes the shot plan and generates video clips per shot.
5. Zen-Director evaluates temporal consistency across clips and suggests retakes.

5.2 Film Education

Zen-Director serves as an interactive cinematography tutor: students can submit scene descriptions and receive expert-level directorial notes explaining the reasoning behind each shot choice.

5.3 Video Game Cutscene Direction

Game studios use Zen-Director to generate cinematic specifications for in-engine cutscene directors, reducing the time from narrative script to playable sequence by an estimated 60%.

6 Related Work

Video understanding models (VideoChat, Video-LLaMA) focus primarily on description generation. Text-to-video generation systems (Sora, CogVideo, ModelScope) focus on pixel synthesis. Zen-Director uniquely occupies the directorial planning layer between these two stages, drawing on the rich tradition of computational narrative research and cinematography theory.

7 Conclusion

Zen-Director establishes a new model category: the AI cinematographer. By training on film scripts, storyboards, and cinematographic knowledge at 7B scale, the model achieves 81.2% on CinematicQA and generates shot plans rated 4.1/5 by professional cinematographers. Integration with the Zen-Video generation backbone creates a complete AI film production pipeline from creative brief to rendered footage.

References

- [1] J. Li et al., “BLIP-2: Bootstrapping Language-Image Pre-training with Frozen Image Encoders and Large Language Models,” ICML, 2023.
- [2] K. Li et al., “VideoChat: Chat-Centric Video Understanding,” arXiv:2305.06355, 2023.
- [3] A. Makarov et al., “Computational Cinematography: A Survey,” IEEE TPAMI, 2022.