
Hierarchical Denoising For Multi-Step Visual Reasoning

Anonymous Author(s)

Affiliation

Address

email

Abstract

Video models are recently evolving into vision foundation models, but they still lack human-like, multi-step reasoning. Existing streaming autoregressive diffusion models are efficient but lack the reasoning ability, whereas bidirectional diffusion allows for global revision but incurs high inference cost due to the dense frames in fixed-sequence denoising. Consequently, both paradigms struggle to maintain logical consistency with low-latency streaming in complex reasoning tasks. Bridging this gap, we propose **HDR** (*Hierarchical Denoising for Visual Reasoning*), a unified framework for multi-step reasoning by integrating hierarchical latents into the causal video generation process. HDR organizes video latents into a tree-structured hierarchy to perform coarse-to-fine reasoning before streaming output. Coarse denoising layers maintain uncertain hypotheses for global planning, while finer denoising layers progressively refine them into concrete visual states. A sparse hierarchical attention pattern (SHAP) further reduces temporal attention cost. We construct a level-stratified multi-step video reasoning benchmark with out-of-distribution cases, covering six tasks: *maze navigation*, *Tower of Hanoi*, *one-line drawing*, *sliding puzzle*, *Sokoban*, and *water pouring*. Compared with the streaming autoregressive diffusion baseline, HDR improves overall success from 34.22 to 60.29 (76.2% relative gain) in multi-step reasoning accuracy, and improves average progress from 76.00 to 89.56, indicating more consistent intermediate reasoning trajectories. For deployment efficiency, HDR maintains low-latency streaming at 0.70s per latent, $54.2\times$ faster than bidirectional diffusion during streaming. HDR also demonstrates strong data efficiency, retaining 82.9% of its full-data success score using only 2% of the training data, compared with 52.0% for bidirectional diffusion. Further experiments on real-world robots showcase the potential of HDR in physical interaction, providing a new paradigm for physical world modeling. Project demo page is available at <https://hierarchical-diffusion-reasoning.github.io/>.

1 Introduction

Video models are evolving from realistic video synthesizers into generalist visual foundation models capable of visual reasoning [22, 21, 20, 23, 32, 31, 10]. Recent work such as Veo3 [22] showed that large video generation models can solve tasks such as maze navigation, symmetry completion, physical reasoning, and tool-use simulation through generated visual trajectories. Complementarily, recent analysis of diffusion-based video reasoning suggests that such reasoning is closely tied to intermediate denoising states, where the model can maintain and refine candidate solutions over multiple steps [21]. These findings raise a key question: how can video models support reliable multi-step reasoning while still enabling low-latency streaming generation?

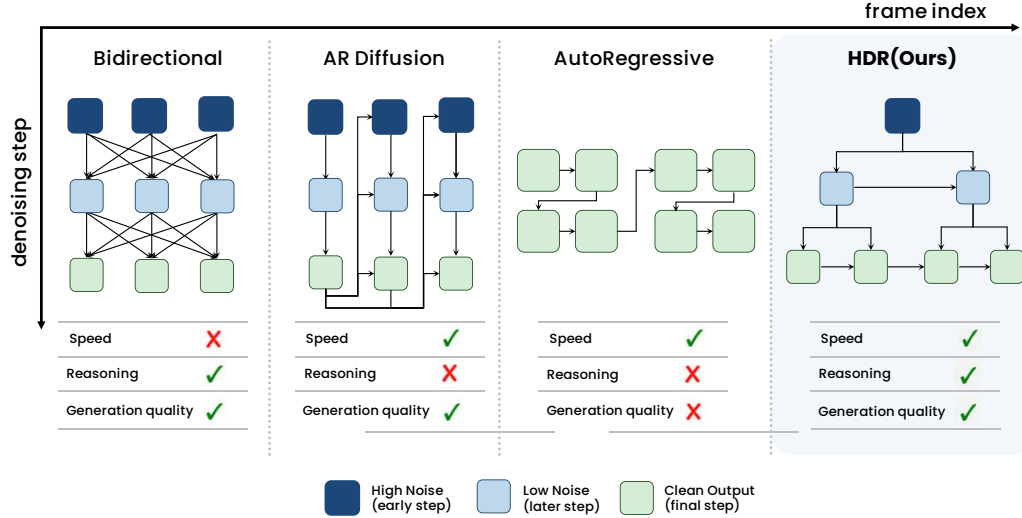


Figure 1: Comparison of video generation paradigms. Streaming autoregressive diffusion models are efficient but commit too early for reliable multi-step reasoning, while bidirectional diffusion supports global revision but requires costly dense fixed-sequence denoising. HDR performs hierarchical denoising before streaming output: coarse layers maintain high-level hypotheses, finer layers refine them into visual states, and sparse hierarchical attention keeps generation efficient.

Existing video generation paradigms struggle to satisfy both multi-step reasoning and low-latency streaming. Streaming autoregressive diffusion models, such as CausVid [28] and CausalForcing [33] generate efficiently by conditioning each step only on past context, but this left-to-right commitment limits their ability to revise previous decisions and perform multi-step reasoning [25, 2, 6, 13]. In contrast, bidirectional diffusion models jointly denoise a fixed video sequence, allowing global information flow and revision across time [15, 19, 22]. However, this requires dense full-sequence updates at every denoising step, leading to high deployment cost and poor compatibility with streaming generation [28, 3, 1, 16]. These limitations reveal a fundamental tension: current models either generate efficiently but struggle with logical consistency over multiple steps, or reason globally at the cost of high-latency fixed-sequence denoising.

Motivated by this observation, we propose **HDR** (*Hierarchical Denoising for Visual Reasoning*), a unified framework that integrates hierarchical latents into the streaming autoregressive diffusion process. As illustrated in Figure 1, HDR organizes video latents into a tree-structured hierarchy and performs coarse-to-fine multi-step reasoning before streaming output. This hierarchy gives the model an explicit intermediate space for high-level planning before it commits to frame-level generation.

A key design of HDR is to match denoising strength to hierarchy level. Instead of fully denoising every layer, HDR keeps coarse layers at higher noise levels so they can preserve multiple possible global plans, while finer layers receive stronger denoising and lower residual noise to instantiate these plans into concrete visual states. HDR further introduces **SHAP** (*Sparse Hierarchical Attention Pattern*), which lets each token communicate only with fixed local and parent-level contexts for fast retrieval. This enables multi-scale information flow without dense full-sequence attention, reducing temporal attention cost while preserving streaming generation.

We construct a level-stratified multi-step video reasoning benchmark with out-of-distribution cases, covering six tasks: *maze navigation*, *Tower of Hanoi*, *one-line drawing*, *sliding puzzle*, *Sokoban*, and *water pouring*. These tasks require models to maintain logical consistency across multi-step reasoning trajectories rather than merely generating locally plausible motion. Experiments show that HDR substantially improves both final task completion and intermediate reasoning consistency: it improves the overall success score from 34.22 to 60.29 (76.2% relative gain), and increases the overall average progress score from 76.00 to 89.56. HDR also preserves efficient streaming behavior, achieving 0.70s per latent during streaming, 54.2× faster than bidirectional diffusion. In addition, HDR demonstrates strong data efficiency, retaining 82.9% of its full-data success score when trained with only 2% of the data, compared with 52.0% for bidirectional diffusion. Finally, real-world robot maze experiments show that HDR can transfer its hierarchical reasoning ability to physical interaction, suggesting its potential as a new paradigm for physical world modeling.

71 Our contributions can be summarized as follows:

- 72 • We identify the core tension in multi-step video reasoning: models must maintain logical consis-
73 tency across long trajectories while also supporting low-latency streaming generation.
- 74 • We propose **HDR** (*Hierarchical Denoising for Visual Reasoning*), a unified framework that
75 integrates hierarchical latents into streaming autoregressive diffusion and performs coarse-to-fine
76 reasoning before streaming output.
- 77 • We introduce a hierarchy-matched denoising schedule and **SHAP** (*Sparse Hierarchical Attention*
78 *Pattern*). The former preserves high-level hypotheses at coarse layers and refines them at finer
79 layers, while the latter reduces temporal attention cost through local and parent-level contexts.
- 80 • We construct a level-stratified multi-step video reasoning benchmark with OOD cases, and
81 demonstrate HDR’s advantages in reasoning accuracy, data efficiency, low-latency streaming, and
82 physical-world robot interaction.

83 2 Related Work

84 **Video Model Reasoning.** Recent studies show that video generation models can exhibit reasoning
85 behaviors beyond visual realism. Benchmarks such as VBVR, V-ReasonBench, and VR-Bench evalu-
86 ate these capabilities across structured problem solving, spatial cognition, physical dynamics, maze
87 navigation, and multi-step planning [20, 14, 29]. Unlike video understanding, where the input video is
88 fixed, video generation requires the model to construct a coherent future trajectory that satisfies both
89 local visual plausibility and global task constraints. This makes generated videos a natural interface
90 for world-model-style reasoning, but also makes early visualized mistakes difficult to correct [23].
91 Recent analyses further suggest that reasoning in video diffusion models is closely tied to iterative
92 denoising dynamics, where intermediate latents maintain and refine uncertain hypotheses, rather than
93 arising solely from frame-by-frame prediction [21]. However, existing work mainly benchmarks
94 or analyzes emergent reasoning in bidirectional generators, rather than designing architectures that
95 preserve reasoning ability under streaming generation constraints [20, 14, 21].

96 **Autoregressive Video Diffusion Models.** Streaming autoregressive diffusion models replace dense
97 full-sequence denoising with sequential temporal computation, enabling low-latency video generation
98 and efficient KV-cache reuse. Recent work improves this paradigm through chunk-wise rollout,
99 queue-based denoising, AR-guided diffusion, causal attention, training–inference alignment, distilla-
100 tion, cache sharing, and sliding-window KV-cache acceleration [4, 9, 12, 28, 3, 6, 33, 26, 17, 8, 7, 24].
101 These properties make autoregressive video models attractive for closed-loop robot interaction
102 [11, 27], where low latency is critical, and their context-as-memory structure allows previously gen-
103 erated visual states to serve as a persistent temporal memory [30, 5]. Moreover, because generation
104 proceeds autoregressively, sliding-window KV-cache mechanisms can extend rollout length and sup-
105 port effectively unbounded video generation [26, 13, 3]. However, the same sequential commitment
106 makes earlier decisions difficult to revise: once a frame or latent chunk has been produced, later
107 predictions condition on this committed history. HDR preserves the low-latency structure of stream-
108 ing autoregressive diffusion while introducing a tree-structured latent hierarchy for coarse-to-fine
109 denoising, enabling revisable high-level planning before committing to fine-grained visual details.

110 3 Method

111 We present HDR (*Hierarchical Denoising for Visual Reasoning*), a hierarchical framework for multi-
112 step video reasoning. HDR preserves the low-latency streaming behavior of streaming autoregressive
113 diffusion while introducing a structured intermediate process for global planning and revision.
114 We first revisit why streaming autoregressive generation struggles with multi-step reasoning, then
115 introduce the hierarchical latent representation, layer-wise flow-matching objective, and **SHAP**
116 (*Sparse Hierarchical Attention Pattern*), which enables efficient inference over flattened tree tokens.

117 3.1 Rethinking Streaming Autoregressive Generation for Multi-step Reasoning

118 We start by comparing bidirectional diffusion and streaming autoregressive diffusion. Let $\mathbf{z}^t =$
119 $\{z_1^t, \dots, z_N^t\}$ denote a video latent sequence at denoising step t , where N is the number of temporal
120 latent tokens and z_i^t is the token at temporal position i . Let c denote the conditioning signal, such

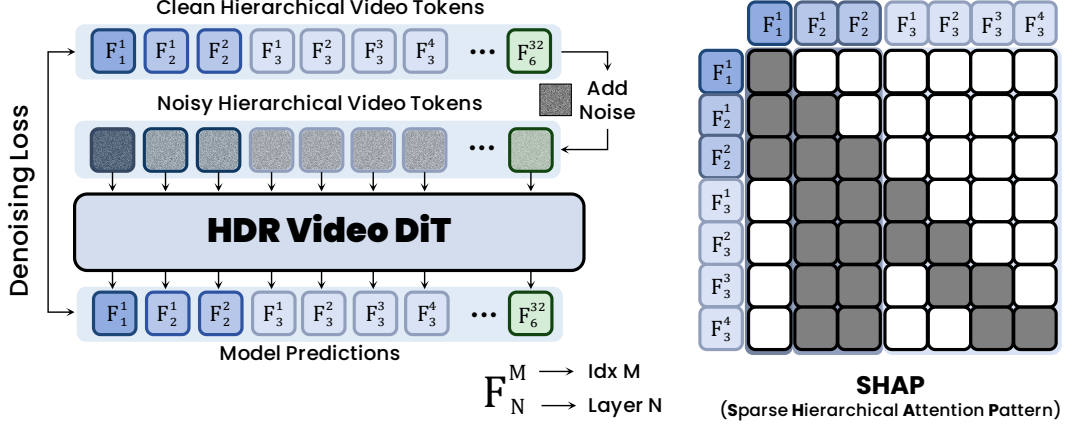


Figure 2: Overview of HDR. Video latents are organized into a tree-structured hierarchy across multiple temporal resolutions. During training, each hierarchy token is corrupted along a flow-matching path and HDR Video DiT is optimized with a layer-wise objective. During inference, all tree tokens are flattened into a coarse-to-fine autoregressive order. **SHAP** (*Sparse Hierarchical Attention Pattern*) defines a structured attention mask over this flattened sequence: each token attends only to fixed local, parent-level, and first-frame contexts rather than the full video sequence. Generated tokens are written into a shared KV cache and reused by later tokens across hierarchy levels, enabling multi-scale information propagation with low temporal attention cost.

121 as text, image, or the first frame. A bidirectional video diffusion model updates the entire sequence
 122 jointly at every denoising step [19, 15]:

$$\mathbf{z}^{t-1} = D_\theta(\mathbf{z}^t, t, c), \quad (1)$$

123 where D_θ is the denoising network. Because all temporal tokens remain noisy during intermediate
 124 denoising steps, information can propagate across the whole sequence before video is committed.
 125 This allows uncertain hypotheses to be maintained and refined over multiple denoising steps [21].
 126 However, the same global revision ability comes with high cost: each step repeatedly updates a dense
 127 fixed-length sequence, making bidirectional diffusion poorly aligned with low-latency streaming.

128 AR diffusion improves deployment efficiency by factorizing generation from left to right:

$$p(\mathbf{z} | c) = \prod_{i=1}^N p(z_i | z_{<i}, c), \quad (2)$$

129 where $\mathbf{z} = \{z_1, \dots, z_N\}$ is the clean latent sequence and $z_{<i}$ denotes all previously generated
 130 temporal tokens. With an autoregressive attention mask, token z_i attends only to past tokens and the
 131 condition c , enabling streaming inference and KV-cache reuse [28, 33, 6]. However, this structure
 132 also creates an irreversible rollout: once z_i is generated, future predictions follow $z_i \rightarrow p(z_{i+1} |$
 133 $z_{\leq i}, c) \rightarrow p(z_{i+2} | z_{\leq i+1}, c) \rightarrow \dots$. If an early token encodes an incorrect decision, later tokens
 134 must condition on this committed history and cannot revise it, which weakens logical consistency
 135 across multi-step reasoning trajectories.

136 Thus, the central challenge is not simply choosing between bidirectional and streaming autoregres-
 137 sive generation. Bidirectional diffusion supports global revision but incurs high deployment cost,
 138 while streaming autoregressive diffusion is efficient but lacks a mechanism for revisable multi-step
 139 reasoning. HDR addresses this tension by introducing a hierarchy of latent variables: coarse levels
 140 preserve noisy high-level hypotheses for global planning, while finer levels progressively refine them
 141 into concrete visual states before streaming output.

142 3.2 Hierarchical Denoising for Visual Reasoning

143 HDR represents a video using a tree-structured hierarchy of latent tokens:

$$\mathcal{T} = \{\mathcal{V}^1, \mathcal{V}^2, \dots, \mathcal{V}^L\}, \quad \mathcal{V}^\ell = \{v_{\ell,1}, \dots, v_{\ell,N_\ell}\}. \quad (3)$$

144 Here, L is the number of hierarchy levels, $\mathcal{V}^\ell$ is the set of latent tokens at level ℓ , N_ℓ is the number of
 145 tokens at that level, and $v_{\ell,i}$ denotes the i -th token. We use $\ell = 1$ for the coarsest level and $\ell = L$

for the finest level. Coarse tokens summarize global temporal structure and represent high-level plans, while fine tokens encode local visual details and instantiate the final video dynamics. Each non-root token has a parent in the previous coarser level, denoted as $\pi(\ell, i) = (\ell - 1, p_\ell(i))$ for $\ell > 1$, where $p_\ell(i)$ maps token $v_{\ell, i}$ to its parent index at level $\ell - 1$. The parent token represents the coarse temporal segment that the current token refines.

As shown in Figure 2, HDR constructs hierarchical tokens from the input video latent sequence and performs coarse-to-fine reasoning before streaming output. During training, each hierarchy token is corrupted along a flow-matching path, and the model learns to predict the corresponding velocity target under hierarchical context. During inference, HDR generates tokens from coarse to fine: upper layers form noisy but revisable global hypotheses, while lower layers refine these hypotheses into concrete visual states. Within each level, generation proceeds autoregressively from left to right, preserving the streaming structure of autoregressive diffusion.

3.3 Layer-wise Flow-Matching Objective

We now describe the training objective over hierarchical tokens. For a clean hierarchy token $v_{\ell, i}^0$, we sample a noise token $\epsilon \sim \mathcal{N}(0, I)$ and construct an interpolated token at continuous time $t \in [0, 1]$ as $v_{\ell, i}^t = (1 - t)v_{\ell, i}^0 + t\epsilon$. Under this linear probability path, the target velocity field is $u_{\ell, i}^t = \epsilon - v_{\ell, i}^0$. The HDR network is trained to predict this flow velocity from the interpolated token, the timestep, the hierarchical context $h_{\ell, i}$, and the condition c . The layer-wise flow-matching objective sums the velocity regression loss over all hierarchy levels and tokens:

$$\mathcal{L}_{\text{HDR}} = \sum_{\ell=1}^L \lambda_\ell \frac{1}{N_\ell} \sum_{i=1}^{N_\ell} \mathbb{E}_{t, \epsilon} \left[\left\| (\epsilon - v_{\ell, i}^0) - u_\theta(v_{\ell, i}^t, t, h_{\ell, i}, c) \right\|_2^2 \right]. \quad (4)$$

Here, $h_{\ell, i}$ denotes the sparse hierarchical context available to token $v_{\ell, i}$, and λ_ℓ balances the contribution of different hierarchy levels. This objective encourages each level to learn a velocity field appropriate to its temporal abstraction: coarse levels model global planning structure, while fine levels model concrete visual states and motion details.

A key design of HDR is to match denoising strength to hierarchy level. Instead of assigning the same inference budget to every layer, HDR uses a level-dependent sampling budget K_ℓ , with $K_1 < K_2 < \dots < K_L$. Equivalently, the residual noise level decreases from coarse to fine, i.e., $\rho_1 > \rho_2 > \dots > \rho_L$. Coarse layers are intentionally stopped at higher noise levels, so their predictions remain partially stochastic and can preserve multiple possible global plans. Finer layers receive stronger denoising and lower residual noise, progressively instantiating these hypotheses into visual states. We provide the entropy-matched derivation of K_ℓ and its ablation in Appendix B.

3.4 Sparse Hierarchical Attention Pattern

HDR implements hierarchical reasoning with **SHAP** (*Sparse Hierarchical Attention Pattern*), a structured attention mask over flattened tree tokens. Each hierarchy token is indexed by (ℓ, i) , where ℓ is the hierarchy level and i is the temporal position within that level. To perform inference autoregressively, we flatten all tree tokens into a coarse-to-fine order using $\phi(\ell, i) = i + \sum_{r < \ell} N_r$, with $s = \phi(\ell, i)$ denoting the flattened token index. Thus, all tokens in $\mathcal{V}^1$ are generated first, followed by $\mathcal{V}^2$, and so on until the finest level $\mathcal{V}^L$. This ordering matches the inference path shown in Figure 2: higher-level tokens are generated before the lower-level tokens that refine them.

For token $v_{\ell, i}$, SHAP first defines its sparse context in tree coordinates:

$$\mathcal{A}_{\text{SHAP}}(\ell, i) = \mathbf{1}[i > 1]\{(\ell, i - 1)\} \cup \mathbf{1}[i = 1]\{x_{\text{ref}}\} \cup \mathbf{1}[\ell > 1]\{\pi(\ell, i), \text{left}(\pi(\ell, i)), \text{right}(\pi(\ell, i))\}. \quad (5)$$

Here, x_{ref} is the clean first-frame condition, $(\ell, i - 1)$ provides same-level autoregressive continuity, and $\pi(\ell, i)$ is the parent token in the coarser level. The neighboring parent-level tokens $\text{left}(\pi(\ell, i))$ and $\text{right}(\pi(\ell, i))$ provide boundary information from adjacent coarse segments. Invalid boundary indices are omitted. For root-level tokens, which have no parent, the hierarchical context reduces to the first-frame condition and same-level autoregressive context.

The tree context is then converted into a binary attention mask over the flattened sequence. Let $S = \sum_{\ell=1}^L N_\ell$ be the total number of hierarchy tokens and $M_{\text{SHAP}} \in \{0, 1\}^{S \times S}$ be the SHAP mask.

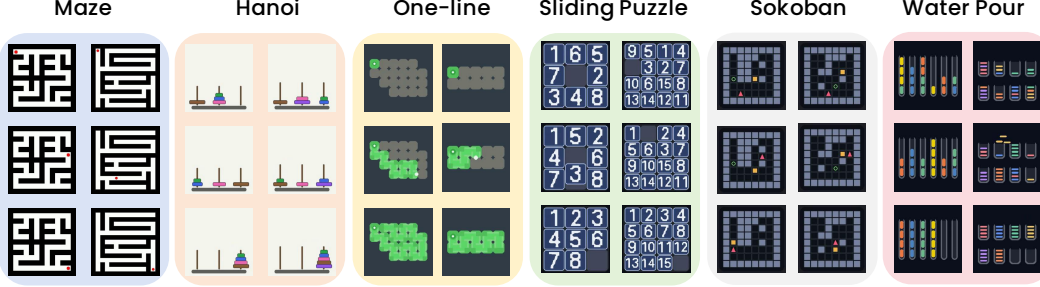


Figure 3: Visualization of the six multi-step video reasoning benchmarks: maze navigation, Tower of Hanoi, one-line drawing, sliding puzzle, Sokoban, and water pouring. Each task requires logical consistency across multiple reasoning steps rather than only local visual plausibility.

192 For $s = \phi(\ell, i)$ and $r = \phi(\ell', j)$, we define:

$$M_{\text{SHAP}}[s, r] = \mathbf{1}[(\ell', j) \in \mathcal{A}_{\text{SHAP}}(\ell, i)]. \quad (6)$$

193 This mask is sparse by construction: each row contains only a constant number of valid attention
 194 targets, independent of video length. It is also autoregressive under the flattened order, because every
 195 valid context token is either a previous same-level token, a coarser-level token that has already been
 196 generated, or the first-frame condition.

197 SHAP naturally induces cross-level KV-cache sharing. During inference, once token $v_{\ell, i}$ is generated,
 198 its key-value state is written into a global hierarchy cache, denoted by $\mathcal{C}_s = \mathcal{C}_{s-1} \cup \{(k_{\ell, i}, v_{\ell, i}^{\text{KV}})\}$,
 199 where $s = \phi(\ell, i)$. When generating a later token $v_{\ell', j}$, HDR retrieves only the cached states selected
 200 by the SHAP mask, i.e., $h_{\ell', j} = \text{Attn}(q_{\ell', j}, \{(k_{\ell, i}, v_{\ell, i}^{\text{KV}}) : M_{\text{SHAP}}[\phi(\ell', j), \phi(\ell, i)] = 1\})$. Thus,
 201 information generated at coarse levels is reused by lower levels through the shared KV cache, while
 202 local temporal continuity is preserved by same-level autoregressive links. Because each token attends
 203 to a fixed-size SHAP context rather than dense full-sequence tokens, HDR reduces temporal attention
 204 cost while maintaining multi-scale information flow before streaming output.

205 4 Experiments

206 We evaluate HDR along two axes: multi-step reasoning ability and low-latency streaming generation
 207 efficiency. Section 4.1 introduces the benchmark, metrics, baselines, and implementation setup.
 208 Section 4.2 compares HDR with full-attention and streaming baselines. Section 4.3 analyzes the
 209 hierarchy layers impact, Section 4.4 tests robustness under reduced denoising budgets and limited
 210 data, and Section 4.5 evaluates transfer to real-world robot interaction.

211 4.1 Benchmarks, Baselines, and Implementation Details

212 **Benchmarks and metrics.** Existing video reasoning benchmarks are not fully suitable for evaluating
 213 streaming multi-step generation: many do not release complete evaluation scripts, and most emphasize
 214 short clips or perceptual consistency rather than logical consistency across multiple reasoning steps.
 215 We therefore construct a controlled benchmark suite with six tasks: Tower of Hanoi, maze navigation,
 216 one-line drawing, sliding puzzle, Sokoban, and water pouring. The benchmark is stratified by
 217 difficulty and includes OOD cases to test rule transfer beyond frequent training patterns. The
 218 scored evaluation set contains 370 held-out videos. Each task is evaluated with *success*, which
 219 measures exact task completion, and *average progress*, which measures partial progress and reflects
 220 intermediate logical consistency. We report task results as *Success / Avg. Progress* pairs and compute
 221 overall performance as an unweighted average across the six tasks. Full benchmark construction and
 222 task-specific evaluation details are provided in Appendix A.

223 **Baselines and implementation.** We compare HDR with full-attention baselines, including bidirectional
 224 diffusion and VideoMAE [18], and streaming baselines, including CausalForcing [33] and
 225 VideoGPT [25]. For a controlled comparison, the main diffusion baselines and HDR are built on
 226 Wan2.2-5B-TI2V [19]; HDR uses six latent hierarchy levels with the entropy-matched denoising
 227 schedule [5, 8, 13, 20, 32, 50]. All methods are trained on the same 18,000-video reasoning dataset,

Table 1: Main comparison on multi-step video reasoning benchmarks. Scores are reported as mean $\pm$ standard deviation for success and average progress. The best and second-best means are shown in **bold** and underlined, respectively. Full-attention baselines are grayed out. Compared with CausalForcing, HDR improves overall success from 34.22 to 60.29 and average progress from 76.00 to 89.56.

Method	Full Attn.	Metric	Hanoi	Maze	One-line	Sliding	Sokoban	Water	Overall
VideoMAE [18]	✓	Success	22.50 $\pm$ 6.18	52.00 $\pm$ 8.96	58.33 $\pm$ 9.67	1.67 $\pm$ 0.00	63.33 $\pm$ 9.42	43.33$\pm$4.94	40.53 $\pm$ 6.53
		Avg. Progress	58.23 $\pm$ 3.80	93.40 $\pm$ 1.32	91.92 $\pm$ 1.71	49.93 $\pm$ 0.00	100.00$\pm$0.00	73.26$\pm$2.48	77.79 $\pm$ 1.55
Bidirectional [19]	✓	Success	<u>45.00$\pm$6.78</u>	90.00$\pm$3.12	81.67$\pm$7.60	<u>31.67$\pm$6.98</u>	90.00$\pm$4.21	21.67 $\pm$ 6.83	<u>60.00$\pm$5.92</u>
		Avg. Progress	<u>73.58$\pm$2.95</u>	99.89$\pm$0.11	99.26$\pm$0.27	<u>93.00$\pm$1.56</u>	100.00$\pm$0.00	57.08 $\pm$ 3.66	<u>87.13$\pm$1.42</u>
VideoGPT [25]	✗	Success	18.75 $\pm$ 5.70	22.00 $\pm$ 12.83	31.67 $\pm$ 7.73	10.00 $\pm$ 5.57	15.00 $\pm$ 5.46	8.33 $\pm$ 3.00	17.57 $\pm$ 6.71
		Avg. Progress	25.82 $\pm$ 5.46	26.32 $\pm$ 12.30	81.74 $\pm$ 2.34	27.81 $\pm$ 4.01	23.50 $\pm$ 4.91	38.03 $\pm$ 3.56	36.88 $\pm$ 5.43
CausalForcing [33]	✗	Success	<u>45.00$\pm$6.73</u>	12.00 $\pm$ 8.08	48.33 $\pm$ 7.66	21.67 $\pm$ 5.16	40.00 $\pm$ 10.69	<u>38.33$\pm$5.25</u>	34.22 $\pm$ 7.26
		Avg. Progress	70.47 $\pm$ 2.76	55.04 $\pm$ 8.36	95.15 $\pm$ 1.28	86.13 $\pm$ 2.62	82.25 $\pm$ 5.79	<u>66.94$\pm$3.03</u>	76.00 $\pm$ 3.97
HDR	✗	Success	58.75$\pm$4.50	<u>78.00$\pm$5.25</u>	<u>70.00$\pm$10.12</u>	33.33$\pm$7.93	<u>78.33$\pm$9.67</u>	43.33$\pm$4.79	60.29$\pm$7.04
		Avg. Progress	79.62$\pm$3.04	<u>97.18$\pm$1.07</u>	<u>97.84$\pm$0.61</u>	93.69$\pm$1.47	<u>99.69$\pm$0.31</u>	<u>69.34$\pm$2.84</u>	89.56$\pm$1.56

conditioned on the first frame, and optimized with the same flow-matching training setup. Complete baseline definitions, training details, and implementation settings are provided in Appendix H.

4.2 Main Results: Bridging Reasoning Precision and Streaming Efficiency

HDR improves both the final success of multi-step reasoning trajectories and their intermediate logical consistency. We support this conclusion through quantitative comparisons in Table 1 and qualitative examples in Figure 4. We further evaluate streaming efficiency in Table 2, showing that HDR maintains comparable latency to AR diffusion baseline [33].

Overall comparison. Table 1 presents the main comparison on our multi-step video reasoning benchmark. Compared with CausalForcing, the streaming autoregressive diffusion baseline, HDR improves overall success from 34.22 to 60.29, corresponding to a 76.2% relative gain. Average progress also increases from 76.00 to 89.56, indicating stronger intermediate logical consistency. Full-attention baselines, shown in gray, enable dense global interaction but are less aligned with low-latency streaming. Despite not using full temporal attention, HDR achieves the best overall success and average progress, showing that hierarchical denoising can recover strong reasoning precision while preserving the streaming structure of autoregressive diffusion.

Qualitative evidence of revisable planning. Figure 4 compares CausalForcing and HDR on Maze and One-line cases. CausalForcing commits to an early local decision that leads to failure: taking the wrong branch in Maze or missing the top block in One-line. HDR instead resolves the ambiguous decision point through hierarchical planning before committing to fine-grained outputs, leading to successful task completion.

Streaming efficiency. We further report inference speed in Table 2. Latency measures the time required for each streaming generation step after KV-cache initialization. HDR achieves comparable latency to CausalForcing, 0.70s versus 0.72s, while being substantially faster than bidirectional diffusion at 37.92s. This shows that HDR improves multi-step reasoning while preserving low-latency streaming behavior.

Table 2: Inference speed comparison. Latency is the average time required for each streaming generation step after KV-cache initialization.

Method	Latency
Bidirectional	37.92s
CausalForcing [33]	0.72s
HDR	0.70s

4.3 Mechanism Analysis: Layer-wise Analysis

Hierarchical layer importance. We study how performance changes as the number of active hierarchical layers increases. Figure 5 reorganizes the variants by layer count, from a single-layer setting equivalent to CausalForcing to the full six-layer HDR hierarchy. The one-layer setting lacks the coarse-to-fine reasoning process enabled by the hierarchy. As additional layers are introduced, the model gains progressively richer high-level planning and stronger multi-step reasoning behavior.

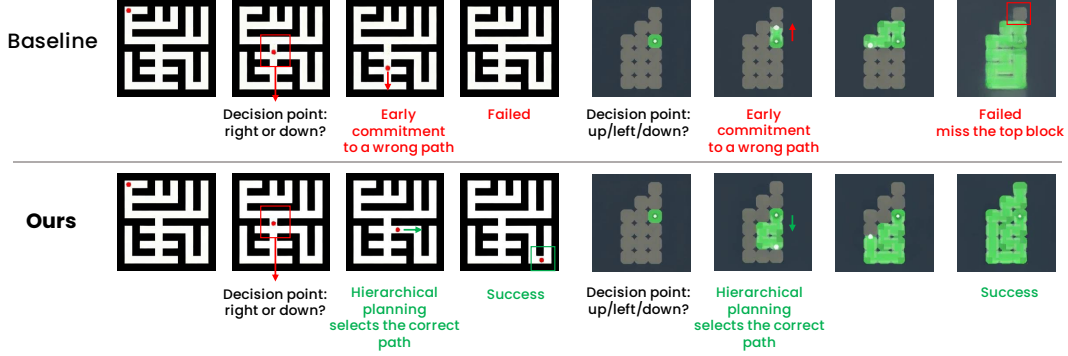


Figure 4: Qualitative comparison between Baseline (CausalForcing) and HDR on Maze and One-line tasks. CausalForcing makes an early local commitment that leads to failure, while HDR performs hierarchical planning before committing to the final trajectory.

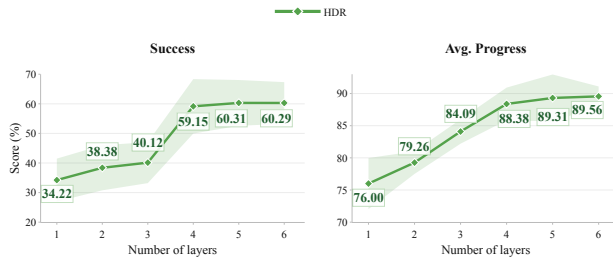


Figure 5: Hierarchical layer importance. Curves show mean performance with one-standard-deviation bands. Increasing active hierarchy layers from 1 to 6 consistently improves HDR over CausalForcing [33], showing that each layer contributes to the full coarse-to-fine reasoning process.

262 This trend shows that HDR’s advantage does not come only from the final frame-level refinement:
 263 each hierarchy level contributes useful structure, and deeper hierarchies lead to better performance.

264 4.4 Robustness: Performance under Reduced Budgets and Limited Data

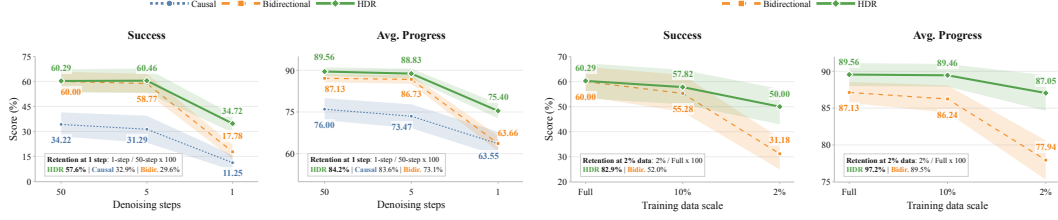
265 **Denosing-step reduction.** We first study robustness under reduced denoising budgets. Figure 6(a)
 266 compares CausalForcing, bidirectional diffusion, and HDR as the number of inference denoising
 267 steps decreases. Both baselines degrade substantially under aggressive step reduction: bidirectional
 268 diffusion drops from 60.00 to 17.78 in overall success with one step, while CausalForcing drops from
 269 34.22 to 11.25. In contrast, HDR remains substantially more robust, achieving 34.72 success with
 270 one denoising step and preserving 57.6% of its full-step performance, compared with 29.6% and
 271 32.9% for the bidirectional and CausalForcing baselines, respectively.

272 **Data reduction.** We test whether HDR can learn reasoning rules from limited data. Figure 6(b)
 273 compares HDR and bidirectional diffusion using the full training set, 10%, and 2% of the data, with
 274 task-level results in Appendix Table 9. As data decreases, HDR degrades more gracefully: with
 275 only 2% of the training data, it retains 82.9% of its full-data success and 97.2% of its full-data
 276 average progress, compared with 52.0% and 89.5% for bidirectional diffusion. This suggests that
 277 the hierarchy provides a useful inductive bias for learning transferable task rules rather than simply
 278 memorizing training examples.

279 Together, these ablations show that HDR’s robustness comes from its hierarchical architecture. The
 280 model remains effective under limited denoising budgets because coarse layers provide stable global
 281 guidance, while lower layers refine this guidance into detailed visual states. Conversely, when coarse
 282 levels are absent or training data is limited, the hierarchical reasoning process becomes the key factor
 283 that determines whether the model can preserve multi-step logical consistency.

284 4.5 Physical World Modeling: Transfer to Robot Interaction

285 To evaluate whether HDR transfers beyond synthetic benchmarks, we conduct a physical-world robot
 286 maze experiment. The task requires a robot arm to pick up a plush toy and move it through a maze
 287 built from toy blocks. This setting introduces visual domain shifts, imperfect object localization,
 288 occlusion, lighting variation, and irregular maze boundaries. We pretrain HDR on 3,000 virtual maze



(a) Denoising-step reduction. HDR remains more robust than both CausalForcing [33], the streaming AR diffusion baseline, and bidirectional diffusion. (b) Data reduction. HDR degrades more gracefully than bidirectional diffusion as the training data scale decreases from the full set to 10% and 2%.

Figure 6: Robustness ablations of HDR. Curves show mean performance, and shaded bands denote one standard deviation. The left plot evaluates reduced denoising budgets, and the right plot evaluates reduced training-data scales. Full task-level data-reduction results are provided in Appendix Table 9.

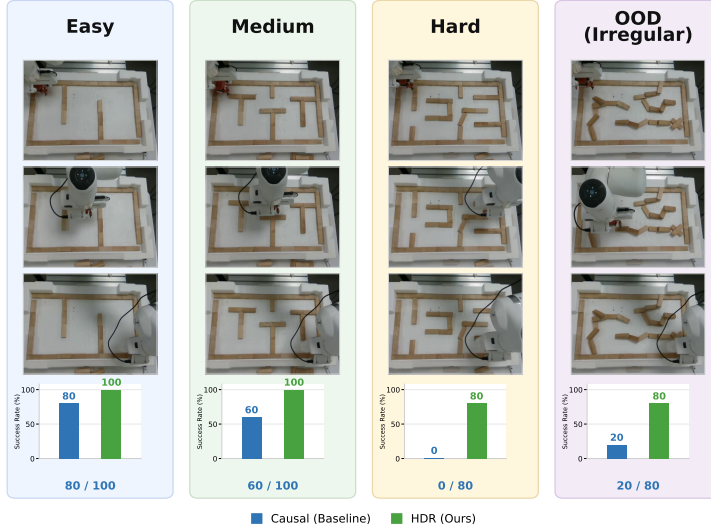


Figure 7: Physical-world robot maze experiment. We fine-tune both CausalForcing [33] and HDR using only 50 real-world robot maze videos, then use an inverse dynamics model (IDM) to convert generated videos into executable robot actions. HDR maintains strong success rates across easy, medium, hard, and OOD mazes, where OOD mazes contain diagonal or stacked wooden blocks.

289 videos, fine-tune both HDR and CausalForcing using only 50 real-world robot videos, and convert
 290 generated videos into executable robot actions using an inverse dynamics model (IDM).

291 We categorize physical-world mazes by difficulty. Easy, medium, and hard mazes are defined by
 292 maze complexity, including the number of branches and turns required by the solution path. We also
 293 include an OOD setting, where wooden blocks are placed diagonally or stacked together, producing
 294 layouts that differ from regular grid-like training mazes. As shown in Figure 7, HDR maintains
 295 strong success rates across all settings, indicating that hierarchical denoising can transfer multi-step
 296 reasoning ability to physical interaction.

297 5 Conclusion

298 We introduced HDR, a framework for long-horizon video reasoning. By organizing video latents
 299 into a coarse-to-fine temporal tree, equipping the hierarchy with sparse structured attention, and
 300 allocating denoising budgets according to an entropy-matched principle, the method recovers much
 301 of the reasoning ability of bidirectional diffusion without relying on full temporal attention.

302 Empirically, HDR outperforms the causal baseline and remains competitive with bidirectional
 303 diffusion across six benchmarks. Ablations show both ingredients matter: removing coarse hierarchy
 304 levels hurts, and fully denoising every level is inferior to preserving uncertainty at upper levels.

305 More broadly, our results suggest that global reasoning in generative video modeling does not
 306 necessarily require dense all-to-all temporal computation. Structured multi-scale latent planning
 307 provides a promising alternative that preserves causal efficiency while preserving reasoning ability.

References

- [1] Andreas Blattmann, Tim Dockhorn, Sumith Kulal, Daniel Mendelevitch, Maciej Kilian, Dominik Lorenz, Yam Levi, Zion English, Vikram Voleti, Adam Letts, Varun Jampani, and Robin Rombach. Stable video diffusion: Scaling latent video diffusion models to large datasets, 2023.
- [2] Haoge Deng, Ting Pan, Haiwen Diao, Zhengxiong Luo, Yufeng Cui, Huchuan Lu, Shiguang Shan, Yonggang Qi, and Xinlong Wang. Autoregressive video generation without vector quantization, 2025.
- [3] Kaifeng Gao, Jiaxin Shi, Hanwang Zhang, Chunping Wang, Jun Xiao, and Long Chen. Ca2-vdm: Efficient autoregressive video diffusion model with causal generation and cache sharing, 2025.
- [4] Roberto Henschel, Levon Khachatryan, Hayk Poghosyan, Daniil Hayrapetyan, Vahram Tadevosyan, Zhangyang Wang, Shant Navasardyan, and Humphrey Shi. Streamingt2v: Consistent, dynamic, and extendable long video generation from text, 2025.
- [5] Yicong Hong, Yiqun Mei, Chongjian Ge, Yiran Xu, Yang Zhou, Sai Bi, Yannick Hold-Geoffroy, Mike Roberts, Matthew Fisher, Eli Shechtman, Kalyan Sunkavalli, Feng Liu, Zhengqi Li, and Hao Tan. Relic: Interactive video world model with long-horizon memory, 2025.
- [6] Xun Huang, Zhengqi Li, Guande He, Mingyuan Zhou, and Eli Shechtman. Self forcing: Bridging the train-test gap in autoregressive video diffusion, 2025.
- [7] Dongya Jia, Zhuo Chen, Jiawei Chen, Chenpeng Du, Jian Wu, Jian Cong, Xiaobin Zhuang, Chumin Li, Zhen Wei, Yuping Wang, et al. Ditar: Diffusion transformer autoregressive modeling for speech generation. In *International Conference on Machine Learning*, pages 27255–27270. PMLR, 2025.
- [8] Yang Jin, Zhicheng Sun, Ningyuan Li, Kun Xu, Kun Xu, Hao Jiang, Nan Zhuang, Quzhe Huang, Yang Song, Yadong MU, and Zhouchen Lin. Pyramidal flow matching for efficient video generative modeling. In *The Thirteenth International Conference on Learning Representations*, 2025.
- [9] Jihwan Kim, Junoh Kang, Jinyoung Choi, and Bohyung Han. Fifo-diffusion: Generating infinite videos from text without training, 2024.
- [10] Dan Kondratyuk, Lijun Yu, Xiuye Gu, José Lezama, Jonathan Huang, Grant Schindler, Rachel Hornung, Vighnesh Birodkar, Jimmy Yan, Ming-Chang Chiu, Krishna Somandepalli, Hassan Akbari, Yair Alon, Yong Cheng, Josh Dillon, Agrim Gupta, Meera Hahn, Anja Hauth, David Hendon, Alonso Martinez, David Minnen, Mikhail Sirotenko, Kihyuk Sohn, Xuan Yang, Hartwig Adam, Ming-Hsuan Yang, Irfan Essa, Huisheng Wang, David A. Ross, Bryan Seybold, and Lu Jiang. Videopoet: A large language model for zero-shot video generation, 2024.
- [11] Lin Li, Qihang Zhang, Yiming Luo, Shuai Yang, Ruilin Wang, Fei Han, Mingrui Yu, Zelin Gao, Nan Xue, Xing Zhu, Yujun Shen, and Yinghao Xu. Causal world modeling for robot control, 2026.
- [12] Zongyi Li, Shujie Hu, Shujie Liu, Long Zhou, Jeongsoo Choi, Lingwei Meng, Xun Guo, Jinyu Li, Hefei Ling, and Furu Wei. Arlon: Boosting diffusion transformers with autoregressive models for long video generation, 2025.
- [13] Kunhao Liu, Wenbo Hu, Jiale Xu, Ying Shan, and Shijian Lu. Rolling forcing: Autoregressive long video diffusion in real time, 2025.
- [14] Yang Luo, Xuanlei Zhao, Baijiong Lin, Lingting Zhu, Liyao Tang, Yuqi Liu, Ying-Cong Chen, Shengju Qian, Xin Wang, and Yang You. V-reasonbench: Toward unified reasoning benchmark suite for video generation models, 2025.
- [15] NVIDIA, :, Niket Agarwal, Arslan Ali, Maciej Bala, Yogesh Balaji, Erik Barker, Tiffany Cai, Prithvijit Chattopadhyay, Yongxin Chen, Yin Cui, Yifan Ding, Daniel Dworakowski, Jiaojiao Fan, Michele Fenzi, Francesco Ferroni, Sanja Fidler, Dieter Fox, Songwei Ge, Yunhao Ge, Jinwei Gu, Siddharth Gururani, Ethan He, Jiahui Huang, Jacob Huffman, Pooya Jannaty,

- Jingyi Jin, Seung Wook Kim, Gergely Klár, Grace Lam, Shiyi Lan, Laura Leal-Taixe, Anqi Li, Zhaoshuo Li, Chen-Hsuan Lin, Tsung-Yi Lin, Huan Ling, Ming-Yu Liu, Xian Liu, Alice Luo, Qianli Ma, Hanzi Mao, Kaichun Mo, Arsalan Mousavian, Seungjun Nah, Sriharsha Niverty, David Page, Despoina Paschalidou, Zeeshan Patel, Lindsey Pavao, Morteza Ramezanali, Fitsum Reda, Xiaowei Ren, Vasanth Rao Naik Sabavat, Ed Schmerling, Stella Shi, Bartosz Stefaniak, Shitao Tang, Lyne Tchapmi, Przemek Tredak, Wei-Cheng Tseng, Jibin Varghese, Hao Wang, Haoxiang Wang, Heng Wang, Ting-Chun Wang, Fangyin Wei, Xinyue Wei, Jay Zhangjie Wu, Jiashu Xu, Wei Yang, Lin Yen-Chen, Xiaohui Zeng, Yu Zeng, Jing Zhang, Qinsheng Zhang, Yuxuan Zhang, Qingqing Zhao, and Artur Zolkowski. Cosmos world foundation model platform for physical ai, 2025.
- [16] Zezhong Qian, Xiaowei Chi, Yuming Li, Shizun Wang, Zhiyuan Qin, Xiaozhu Ju, Sirui Han, and Shanghang Zhang. Wristworld: Generating wrist-views via 4d world models for robotic manipulation, 2025.
- [17] Robbyant Team, Zelin Gao, Qiuyu Wang, Yanhong Zeng, Jiapeng Zhu, Ka Leong Cheng, Yixuan Li, Hanlin Wang, Yinghao Xu, Shuailei Ma, Yihang Chen, Jie Liu, Yansong Cheng, Yao Yao, Jiayi Zhu, Yihao Meng, Kecheng Zheng, Qingyan Bai, Jingye Chen, Zehong Shen, Yue Yu, Xing Zhu, Yujun Shen, and Hao Ouyang. Advancing open-source world models, 2026.
- [18] Zhan Tong, Yibing Song, Jue Wang, and Limin Wang. Videomae: Masked autoencoders are data-efficient learners for self-supervised video pre-training, 2022.
- [19] Team Wan, Ang Wang, Baole Ai, Bin Wen, Chaojie Mao, Chen-Wei Xie, Di Chen, Feiwei Yu, Haiming Zhao, Jianxiao Yang, Jianyuan Zeng, Jiayu Wang, Jingfeng Zhang, Jingren Zhou, Jinkai Wang, Jixuan Chen, Kai Zhu, Kang Zhao, Keyu Yan, Lianghua Huang, Mengyang Feng, Ningyi Zhang, Pandeng Li, Pingyu Wu, Ruihang Chu, Ruili Feng, Shiwei Zhang, Siyang Sun, Tao Fang, Tianxing Wang, Tianyi Gui, Tingyu Weng, Tong Shen, Wei Lin, Wei Wang, Wei Wang, Wenmeng Zhou, Wenten Wang, Wenting Shen, Wenyuan Yu, Xianzhong Shi, Xiaoming Huang, Xin Xu, Yan Kou, Yangyu Lv, Yifei Li, Yijing Liu, Yiming Wang, Yingya Zhang, Yitong Huang, Yong Li, You Wu, Yu Liu, Yulin Pan, Yun Zheng, Yuntao Hong, Yupeng Shi, Yutong Feng, Zeyinzi Jiang, Zhen Han, Zhi-Fan Wu, and Ziyu Liu. Wan: Open and advanced large-scale video generative models, 2025.
- [20] Maijunxian Wang, Ruisi Wang, Juyi Lin, Ran Ji, Thaddäus Wiedemer, Qingying Gao, Dezhi Luo, Yaoyao Qian, Lianyu Huang, Zelong Hong, Jiahui Ge, Qianli Ma, Hang He, Yifan Zhou, Lingzi Guo, Lantao Mei, Jiachen Li, Hanwen Xing, Tianqi Zhao, Fengyuan Yu, Weihang Xiao, Yizheng Jiao, Jianheng Hou, Danyang Zhang, Pengcheng Xu, Boyang Zhong, Zehong Zhao, Gaoyun Fang, John Kitaoka, Yile Xu, Hua Xu, Kenton Blacutt, Tin Nguyen, Siyuan Song, Haoran Sun, Shaoyue Wen, Linyang He, Runming Wang, Yanzhi Wang, Mengyue Yang, Ziqiao Ma, Raphaël Millière, Freda Shi, Nuno Vasconcelos, Daniel Khashabi, Alan Yuille, Yilun Du, Ziming Liu, Bo Li, Dahua Lin, Ziwei Liu, Vikash Kumar, Yijiang Li, Lei Yang, Zhongang Cai, and Hokin Deng. A very big video reasoning suite, 2026.
- [21] Ruisi Wang, Zhongang Cai, Fanyi Pu, Junxiang Xu, Wanqi Yin, Maijunxian Wang, Ran Ji, Chenyang Gu, Bo Li, Ziqi Huang, Hokin Deng, Dahua Lin, Ziwei Liu, and Lei Yang. Demystifying video reasoning, 2026.
- [22] Thaddäus Wiedemer, Yuxuan Li, Paul Vicol, Shixiang Shane Gu, Nick Matarese, Kevin Swersky, Been Kim, Priyank Jaini, and Robert Geirhos. Video models are zero-shot learners and reasoners, 2025.
- [23] Jialong Wu, Xiaoying Zhang, Hongyi Yuan, Xiangcheng Zhang, Tianhao Huang, Changjing He, Chaoyi Deng, Renrui Zhang, Youbin Wu, and Mingsheng Long. Visual generation unlocks human-like reasoning through multimodal world models, 2026.
- [24] Guangxuan Xiao, Yuandong Tian, Beidi Chen, Song Han, and Mike Lewis. Efficient streaming language models with attention sinks. In *The Twelfth International Conference on Learning Representations*, 2024.
- [25] Wilson Yan, Yunzhi Zhang, Pieter Abbeel, and Aravind Srinivas. Videogpt: Video generation using vq-vae and transformers, 2021.

- 409 [26] Shuai Yang, Wei Huang, Ruihang Chu, Yicheng Xiao, Yuyang Zhao, Xianbang Wang, Muyang
410 Li, Enze Xie, Yingcong Chen, Yao Lu, Song Han, and Yukang Chen. Longlive: Real-time
411 interactive long video generation, 2025.
- 412 [27] Seonghyeon Ye, Yunhao Ge, Kaiyuan Zheng, Shenyuan Gao, Sihyun Yu, George Kurian, Suneel
413 Indupuru, You Liang Tan, Chuning Zhu, Jiannan Xiang, Ayaan Malik, Kyungmin Lee, William
414 Liang, Nadun Ranawaka, Jiasheng Gu, Yinzhen Xu, Guanzhi Wang, Fengyuan Hu, Avnish
415 Narayan, Johan Bjorck, Jing Wang, Gwanghyun Kim, Dantong Niu, Ruijie Zheng, Yuqi Xie,
416 Jimmy Wu, Qi Wang, Ryan Julian, Danfei Xu, Yilun Du, Yevgen Chebotar, Scott Reed, Jan
417 Kautz, Yuke Zhu, Linxi "Jim" Fan, and Joel Jang. World action models are zero-shot policies,
418 2026.
- 419 [28] Tianwei Yin, Qiang Zhang, Richard Zhang, William T. Freeman, Fredo Durand, Eli Shechtman,
420 and Xun Huang. From slow bidirectional to fast autoregressive video diffusion models, 2025.
- 421 [29] Jiashuo Yu, Yue Wu, Meng Chu, Zhifei Ren, Zizheng Huang, Pei Chu, Ruijie Zhang, Yinan He,
422 Qirui Li, Songze Li, Zhenxiang Li, Zhongying Tu, Conghui He, Yu Qiao, Yali Wang, Yi Wang,
423 and Limin Wang. Vrbench: A benchmark for multi-step reasoning in long narrative videos,
424 2025.
- 425 [30] Jiwen Yu, Jianhong Bai, Yiran Qin, Quande Liu, Xintao Wang, Pengfei Wan, Di Zhang, and
426 Xihui Liu. Context as memory: Scene-consistent interactive long video generation with memory
427 retrieval, 2025.
- 428 [31] Lijun Yu, José Lezama, Nitesh B. Gundavarapu, Luca Versari, Kihyuk Sohn, David Minnen,
429 Yong Cheng, Vighnesh Birodkar, Agrim Gupta, Xiuye Gu, Alexander G. Hauptmann, Boqing
430 Gong, Ming-Hsuan Yang, Irfan Essa, David A. Ross, and Lu Jiang. Language model beats
431 diffusion – tokenizer is key to visual generation, 2024.
- 432 [32] Kevin Zhang, Kuangzhi Ge, Xiaowei Chi, Renrui Zhang, Shaojun Shi, Zhen Dong, Sirui Han,
433 and Shanghang Zhang. Can world models benefit vlms for world dynamics?, 2025.
- 434 [33] Hongzhou Zhu, Min Zhao, Guande He, Hang Su, Chongxuan Li, and Jun Zhu. Causal forcing:
435 Autoregressive diffusion distillation done right for high-quality real-time interactive video
436 generation, 2026.

A Benchmark and Eval Details

We evaluate all methods on a mixed benchmark containing six long-horizon reasoning tasks: maze navigation, Tower of Hanoi, one-line drawing, sliding puzzle, Sokoban, and water pouring. These tasks cover diverse reasoning patterns, including spatial planning, state transition, object manipulation, and trajectory consistency.

Each generated video is evaluated with task-specific success criteria and reported using two metrics. The success score $S \in \{0, 1\}$ measures exact task completion under the benchmark-specific success criterion. The average progress score $A \in [0, 1]$ measures partial progress toward the target outcome and is therefore more tolerant of minor visual deviations that do not alter the core semantic trajectory. Results in the main paper are reported as Success / Avg. Progress pairs, with both values multiplied by 100.

For a scored set $\mathcal{D}$, the reported task score is computed as the mean over evaluation samples:

$$\text{Success}(\mathcal{D}) = \frac{1}{|\mathcal{D}|} \sum_{i \in \mathcal{D}} S_i, \quad \text{Avg. Progress}(\mathcal{D}) = \frac{1}{|\mathcal{D}|} \sum_{i \in \mathcal{D}} A_i.$$

The overall score is obtained by averaging the task-level scores across the six benchmarks. This gives equal weight to each reasoning task and reflects the model’s robustness across different forms of long-horizon video reasoning.

The benchmark uses task-specific difficulty levels. Maze and Hanoi are grouped by native problem size, while the remaining four tasks use level directories. The evaluation split is summarized in Table 3.

Table 3: Difficulty grouping of the mixed reasoning benchmark.

Task	Levels / Sizes	Eval Split	Main Difficulty Factor
Maze	$N = 6, 7, 8, 9, 10$	10 per size	Grid size, path length
Hanoi	2, 3, 4, 5 disks	10 ID + 10 OOD per size	Disk count, initial rods
One-line	levels 2, 3, 4	25 ID + 25 OOD per level	Path length, shape topology
Sliding	levels 2, 3, 4	30 ID + 20 OOD per level	Board size, optimal moves
Sokoban	levels 2, 3, 4	30 ID + 20 OOD per level	Layout grammar, action horizon
Water	levels 2, 3, 4	30 ID + 20 OOD per level	Tubes, capacity, solution length

A.1 Maze Navigation

Maze samples are generated on square grids with fixed start $(0, 0)$ and goal $(N - 1, N - 1)$. The generator constructs a recursive-division maze graph, solves it with breadth-first search, stores the graph edges and solution path, and renders a red ball moving along the path. The default evaluation set contains 50 samples, with 10 samples for each $N \in \{6, 7, 8, 9, 10\}$.

The evaluator tracks the red ball using color segmentation and connected components. The success score requires the decoded route to start correctly, remain in open cells, move only through valid graph edges, and reach the goal. The average progress score is the normalized longest-common-subsequence overlap between the relaxed decoded route $\hat{r}$ and the reference path r^* :

$$A_{\text{maze}} = \frac{\text{LCS}(\hat{r}, r^*)}{\max(1, |r^*|)}.$$

A.2 Tower of Hanoi

Hanoi samples use 2–5 disks and three rods. In-domain samples initialize disks on rods $\{0, 1\}$, while OOD samples may initialize disks on $\{0, 1, 2\}$ and require at least one disk to already appear on the goal rod. The goal is always rod 2. For each initial assignment, the generator solves the shortest legal plan with breadth-first search and renders one lifted disk move at a time.

The evaluator decodes disk tracks, infers moves of the form (disk, source, target), and simulates them from the ground-truth initial state. The success score requires all inferred moves to be legal and

471 the final state to place every disk on the goal rod. The average progress score measures plan overlap
 472 with the reference shortest plan:

$$A_{\text{hanoi}} = \frac{2 \text{LCS}(\hat{m}, m^*)}{|\hat{m}| + |m^*|}.$$

473 A.3 One-line Drawing

474 One-line drawing levels 2, 3, 4 use increasingly larger boards and longer self-avoiding paths. The
 475 generator samples a topology family, searches for an adjacent-cell path without revisits, applies
 476 random geometric transforms, checks uniqueness and shape constraints, and treats the resulting path
 477 as both the occupied shape and the required solution. In the evaluation set, level 2 uses a 9×9 board,
 478 level 3 uses 10×10 , and level 4 uses 12×12 .

479 The evaluator extracts the drawing head, visited cells, white trace cells, and whether the trace remains
 480 on the required shape. The success score requires starting from the highlighted cell, covering every
 481 occupied cell, avoiding revisits, avoiding non-adjacent jumps, and never leaving the shape. A small
 482 visual bridge tolerance is allowed when the white trace already connects an apparent short jump. The
 483 average progress score combines coverage, legal-transition ratio, and on-shape ratio:

$$A_{\text{one}} = 0.5 C_{\text{cover}} + 0.3 R_{\text{legal}} + 0.2 R_{\text{shape}}.$$

484 A.4 Sliding Puzzle

485 Sliding puzzle levels 2 and 3 use 3×3 boards, while level 4 uses a 4×4 board. The goal state is the
 486 ordered board with the blank tile in the final position. For 3×3 puzzles, states are sampled from
 487 an exact solvable pool under bucket constraints; for 4×4 puzzles, states are sampled by backward
 488 scrambling from the goal. Each accepted sample is solved optimally and filtered by solution length,
 489 family, and diversity constraints.

490 The evaluator decodes board states from video frames and selects the best legal subsequence starting
 491 at the initial state. The success score requires the subsequence to reach the goal, preserve tile
 492 inventory, contain no illegal blank moves, and have an unresolved-frame ratio below 0.8. The average
 493 progress score combines normalized Manhattan progress, final-state quality, rule adherence, and
 494 observability:

$$A_{\text{slide}} = 0.45 P + 0.20 Q_{\text{final}} + 0.20 R_{\text{rule}} + 0.15 R_{\text{obs}}.$$

495 A.5 Sokoban

496 Sokoban samples use an 8×8 grid with walls, one player, one box, and one target. The generator
 497 builds layouts from grammar operators, samples valid player and box positions, solves each candidate
 498 with shortest-path search, and filters by action length, push count, box-target distance, deadlock-like
 499 cells, and diversity. Higher levels use longer and more constrained layouts.

500 The evaluator decodes player and box positions and checks whether the recovered state sequence can
 501 be explained by legal walking and pushing. The success score requires the player and box to start
 502 correctly, the box to end on the target, no wall crossing or overlap, no illegal pulling or pushing, and
 503 an unresolved-frame ratio below 0.8. The average progress score emphasizes box progress toward
 504 the target:

$$A_{\text{sokoban}} = 0.50 P_{\text{box}} + 0.20 Q_{\text{final}} + 0.15 R_{\text{rule}} + 0.15 R_{\text{obs}}.$$

505 A.6 Water Pouring

506 Water pouring samples contain vertical tubes with fixed capacity and colored blocks. The generator
 507 constructs a solved goal state, samples a backward chain of legal reverse pours to obtain an initial
 508 state, solves or validates the resulting puzzle, and filters candidates by solution length, buried blocks,
 509 color runs, branching factor, and state diversity. Higher levels increase tube count, color count,
 510 capacity, and solution horizon.

511 The evaluator extracts tube states from stationary frames. A legal move transfers the complete
 512 contiguous top run of one color into an empty tube or onto the same color, limited by remaining
 513 destination capacity. The success score requires the final tube state to be solved, all decoded transitions

514 to be legal, block conservation and capacity constraints to hold, and unresolved stationary frames to
 515 remain below 0.35. The average progress score combines rule adherence, progress, final-state quality,
 516 and observability:

$$A_{\text{water}} = 0.45 R_{\text{rule}} + 0.35 P + 0.15 Q_{\text{final}} + 0.05 R_{\text{obs}}.$$

517 B Entropy-Matched versus Fully Denoised Hierarchies

518 Our final ablation studies how denoising budgets should be allocated across the hierarchy. A naive
 519 strategy is to fully denoise every hierarchy level before moving to the next one. In our implementation,
 520 this corresponds to assigning 50 denoising steps to all layers. Although intuitive, this strategy ignores
 521 the different entropy scales of different hierarchy levels. Coarse layers contain fewer temporal tokens
 522 and represent lower-resolution hypotheses, while fine layers contain more tokens and carry more
 523 detailed visual entropy. Therefore, forcing all levels to remove the same amount of uncertainty is not
 524 well matched to the hierarchical representation.

We instead use an entropy-matched denoising schedule. Let N_ℓ denote the number of tokens at layer ℓ , and let $\tilde{N}_\ell$ denote its effective temporal support. Since a finer layer represents a larger number of temporal degrees of freedom, we assume that the entropy to be removed at layer ℓ should grow with its effective support:

$$\Delta H_\ell \propto \tilde{N}_\ell^\beta,$$

525 where β is a tunable parameter controlling how aggressively denoising budget increases from coarse
 526 to fine layers. Larger β allocates more denoising steps to fine layers, while smaller β makes the
 527 schedule more uniform.

We convert this entropy allocation into a layer-wise denoising budget:

$$K_\ell = \left\lceil K_{\max} \left(\frac{\tilde{N}_\ell}{\tilde{N}_L} \right)^\beta \right\rceil,$$

528 where $K_{\max}$ is the maximum denoising budget used by the finest layer. This rule gives fewer
 529 denoising steps to coarse layers and more steps to fine layers, matching the intuition that coarse layers
 530 should preserve uncertainty while fine layers should resolve it into concrete visual states.

For the 21-frame setting, our hierarchy has actual layer sizes:

$$N_\ell = [1, 2, 4, 8, 16, 21].$$

The last layer is truncated by the video length, but the underlying binary hierarchy has effective temporal supports:

$$\tilde{N}_\ell = [1, 2, 4, 8, 16, 32].$$

Using $K_{\max} = 50$ and $\beta = 0.66$, the entropy-matched rule gives:

$$K_\ell = \left\lceil 50 \left(\frac{[1, 2, 4, 8, 16, 32]}{32} \right)^{0.66} \right\rceil = [5, 8, 13, 20, 32, 50].$$

531 This is the default denoising schedule used by HDR in the main experiments.

This schedule can also be interpreted through the lens of uncertainty preservation. Coarse layers remove only a small fraction of their uncertainty, so they act as flexible high-level hypotheses. Fine layers remove much more entropy, allowing them to instantiate these hypotheses into detailed frame-level latents. In contrast, the All-50 strategy sets

$$K_1 = K_2 = \dots = K_L = 50,$$

532 which forces every level to remove the same amount of uncertainty regardless of its temporal scale.
 533 This prematurely collapses high-level hypotheses and reduces the ability of lower layers to correct or
 534 refine them.

535 Table 4 compares the entropy-matched schedule with fully denoised and alternative schedules. The
 536 All-50 variant performs worse overall than the entropy-matched schedule: the overall success score
 537 drops from 60.29 to 58.38, and the average progress score drops from 89.56 to 88.21. We also include
 538 two additional schedules for completeness: a sparse-to-full schedule $[5, 5, 5, 5, 5, 50]$, which keeps
 539 coarse layers minimally denoised before fully denoising the finest layer, and an exponential schedule
 540 $[2, 4, 8, 16, 32, 50]$, which increases the denoising budget more aggressively. Their benchmark results
 541 are left blank.

Table 4: Entropy-matched versus alternative denoising schedules. The entropy-matched schedule allocates fewer denoising steps to coarse layers and more steps to fine layers according to their entropy scale, achieving the best overall average progress and remaining competitive in overall success. Compared with sparse-to-full and exponential schedules, entropy matching provides a better balance between preserving high-level uncertainty and refining fine-grained video states.

Denoising Strategy	Schedule	Hanoi	Maze	One-line	Sliding	Sokoban	Water	Overall
Entropy-matched	[5, 8, 13, 20, 32, 50]	58.75/79.62	78.00/97.18	70.00/97.84	33.33/93.69	78.33/99.69	43.33/69.34	60.29/89.56
All-50	[50, 50, 50, 50, 50, 50]	60.00/79.70	76.00/97.18	65.00/97.45	30.00/94.34	80.00/99.60	41.67/65.32	58.38/88.21
Sparse-to-full	[5, 5, 5, 5, 5, 50]	55.00/74.58	42.00/92.86	63.33/97.25	41.67/93.76	76.67/99.33	23.33/58.99	50.33/86.13
Exponential	[2, 4, 8, 16, 32, 50]	53.75/77.36	58.00/95.53	73.33/98.03	41.67/94.30	86.67/98.73	46.67/68.26	60.02/88.70

C Time Complexity Analysis

We analyze the temporal attention complexity of bidirectional diffusion, Streaming AR Diffusion, and HDR. Let N denote the number of frame-level video tokens and K denote the number of denoising steps. We omit constant factors such as hidden dimension, number of heads, model depth, and the fixed number of tokens attended by HDR.

Bidirectional diffusion. Bidirectional diffusion updates the full video sequence at every denoising step. Each token attends to all N tokens, so one denoising step costs $\mathcal{O}(N^2)$ temporal attention. With K denoising steps, the total complexity is:

$$\mathcal{O}(KN^2).$$

This dense all-to-all interaction supports global reasoning, but it repeatedly recomputes the full temporal attention map throughout inference.

Streaming AR Diffusion. Streaming AR Diffusion generates tokens from left to right. When generating token i , it attends to all previously generated tokens. Although KV caching avoids recomputing previous hidden states, the attention length still grows with time. Thus, the total temporal attention cost is:

$$\mathcal{O}\left(K \sum_{i=1}^N i\right) = \mathcal{O}(KN^2).$$

Therefore, Streaming AR Diffusion improves streaming efficiency compared with bidirectional diffusion, but its attention complexity with respect to the number of tokens remains quadratic.

HDR. HDR generates a hierarchical latent tree instead of a flat sequence. For a video with N frame-level tokens, the total number of hierarchy tokens is linear in N ; for example, a binary tree contains approximately $2N - 1$ nodes. More importantly, each token attends only to a fixed-size context, including its previous same-layer token, its parent, and neighboring parent tokens. Therefore, the temporal attention cost is linear in the number of hierarchy tokens:

$$\mathcal{O}(K_{\text{avg}}N),$$

where K_{avg} denotes the average denoising budget across hierarchy tokens. Since the hierarchy contains approximately $2N - 1$ tokens for N frame-level tokens, the full attention cost is $\mathcal{O}(K_{\text{avg}}(2N - 1))$, which simplifies to $\mathcal{O}(K_{\text{avg}}N)$. This linear-size hierarchy introduces a longer one-time prefill stage, 16.19s for HDR, compared with 1.48s for bidirectional diffusion and 2.44s for CausalForcing; however, the subsequent streaming latency remains 0.70s per latent and much faster than bidirectional diffusion at 37.92s. In practice, HDR assigns fewer denoising steps to coarse layers, so K_{avg} is smaller than the full denoising budget used by standard diffusion.

This analysis highlights the main computational advantage of HDR. Although HDR introduces additional hierarchy tokens, their number grows only linearly with the video length. Since each hierarchy token attends to a constant-size context, the overall temporal attention cost remains linear in N . In contrast, both bidirectional diffusion and Streaming AR Diffusion have quadratic attention complexity with respect to the number of video tokens.

Table 5: Temporal attention complexity comparison. Bidirectional diffusion uses all-to-all attention, Streaming AR Diffusion attends to the full prefix, while HDR attends to a fixed-size context over a linear-size hierarchy.

Method	Attention Pattern	Token Count	Complexity
Bidirectional	all-to-all	N	$\mathcal{O}(KN^2)$
Streaming AR Diffusion	prefix attention	N	$\mathcal{O}(KN^2)$
HDR	fixed sparse attention	$\approx 2N$	$\mathcal{O}(K_{\text{avg}}N)$

Table 6: Comparison with state-of-the-art closed-source video generation models on video reasoning benchmarks. Scores are reported using the success metric and the average progress metric.

Method	Full Attn.	Fine-tuned	Metric	Hanoi	Maze	One-line	Sliding	Sokoban	Water	Overall
Wan2.6	✓	✗	Success	3.75	6.00	1.67	0.00	0.00	1.67	2.16
			Avg. Progress	20.15	46.93	72.13	61.38	67.56	35.50	49.06
Veo	✓	✗	Success	0.00	0.00	0.00	0.00	0.00	0.00	0.00
			Avg. Progress	0.09	24.42	66.46	0.00	0.33	0.82	14.28
HDR	✗	✓	Success	58.75	78.00	70.00	33.33	78.33	43.33	60.29
			Avg. Progress	79.62	97.18	97.84	93.69	99.69	69.34	89.56

D Comparison with SOTA Closed-source Models

We further compare HDR with state-of-the-art closed-source video generation models, including Wan2.6 and Veo. Since these models are only accessible through closed-source inference APIs or released inference interfaces, we cannot modify their architectures, apply our hierarchical training strategy, or fine-tune them on the proposed benchmark data. Therefore, we evaluate them in an off-the-shelf setting using the same prompts and evaluation protocol as our method.

As shown in Table 6, off-the-shelf closed-source models achieve limited success under the exact-completion criterion on our reasoning benchmarks. Although they can often generate visually plausible videos, they frequently fail to satisfy the exact multi-step constraints required by these tasks. In contrast, HDR is explicitly trained for hierarchical diffusion reasoning and achieves substantially higher scores across all tasks. This comparison highlights that strong general-purpose video generation capability does not necessarily translate into reliable multi-step reasoning under irreversible decision-making.

E Full Table of Ablations

The main paper visualizes the key trends of the denoising-step, hierarchical-layer, and data-reduction ablations in Figure 6a, Figure 5, and Figure 6b. For completeness, we provide the corresponding full numerical results in Table 7, Table 8, and Table 9. Table 7 contains the task-level scores behind Figure 6a, including the full success and average-progress results for the causal baseline, bidirectional baseline, and HDR under different denoising budgets. Table 8 provides the task-level breakdown for Figure 5, organized by the number of active hierarchical layers from 1 to 6, where the 1-layer setting corresponds to the causal baseline. Table 9 reports the task-level results behind Figure 6b, comparing the bidirectional baseline and HDR at different training data scales. All entries are reported as mean $\pm$ standard deviation. For each task-level metric computed over N evaluation samples, we estimate the standard deviation by repeatedly sampling $\lfloor N/2 \rfloor$ examples without replacement, recomputing the mean for 10 trials, and then reporting the population standard deviation across the resulting sample means.

These tables support the same conclusions as the figures while exposing the per-task details. First, the denoising-step ablation shows that HDR retains stronger performance under reduced denoising budgets, especially at the one-step setting. Second, the hierarchical-layer ablation shows that shallower variants stay closer to the causal baseline, while adding more layers generally improves performance

Table 7: Denoising step ablation. Values are reported as mean $\pm$ std. HDR remains substantially more robust than the causal baseline under reduced inference budgets.

Method	Step	Metric	Hanoi	Maze	One-line	Sliding	Sokoban	Water	Overall
CausalForcing [33]	50	Success	45.00 $\pm$ 6.73	12.00 $\pm$ 8.08	48.33 $\pm$ 7.66	21.67 $\pm$ 5.16	40.00 $\pm$ 10.69	38.33 $\pm$ 5.25	34.22 $\pm$ 7.26
		Avg. Progress	70.47 $\pm$ 2.76	55.04 $\pm$ 8.36	95.15 $\pm$ 1.28	86.13 $\pm$ 2.62	82.25 $\pm$ 5.79	66.94 $\pm$ 3.03	76.00 $\pm$ 3.97
	5	Success	38.75 $\pm$ 5.14	14.00 $\pm$ 11.31	43.33 $\pm$ 9.26	20.00 $\pm$ 6.00	35.00 $\pm$ 10.19	36.67 $\pm$ 6.83	31.29 $\pm$ 8.12
		Avg. Progress	66.31 $\pm$ 2.78	49.67 $\pm$ 9.17	94.30 $\pm$ 2.00	83.73 $\pm$ 2.79	79.81 $\pm$ 5.85	66.97 $\pm$ 3.26	73.47 $\pm$ 4.31
	1	Success	7.50 $\pm$ 4.36	10.00 $\pm$ 5.63	0.00 $\pm$ 0.00	18.33 $\pm$ 6.75	28.33 $\pm$ 8.92	3.33 $\pm$ 1.80	11.25 $\pm$ 4.57
		Avg. Progress	39.24 $\pm$ 4.57	60.29 $\pm$ 9.31	92.82 $\pm$ 1.00	79.13 $\pm$ 3.56	77.45 $\pm$ 5.79	32.40 $\pm$ 1.70	63.55 $\pm$ 4.32
Bidirectional [19]	50	Success	45.00 $\pm$ 6.78	90.00 $\pm$ 3.12	81.67 $\pm$ 7.60	31.67 $\pm$ 6.98	90.00 $\pm$ 4.21	21.67 $\pm$ 6.83	60.00 $\pm$ 5.92
		Avg. Progress	73.58 $\pm$ 2.95	99.89 $\pm$ 0.11	99.26 $\pm$ 0.27	93.00 $\pm$ 1.56	100.00 $\pm$ 0.00	57.08 $\pm$ 3.66	87.13 $\pm$ 1.42
	5	Success	41.25 $\pm$ 6.91	88.00 $\pm$ 2.56	75.00 $\pm$ 8.79	35.00 $\pm$ 5.74	86.67 $\pm$ 4.81	26.67 $\pm$ 5.85	58.77 $\pm$ 5.78
		Avg. Progress	75.01 $\pm$ 3.04	99.75 $\pm$ 0.15	98.70 $\pm$ 0.72	92.45 $\pm$ 1.66	99.11 $\pm$ 0.89	55.34 $\pm$ 3.16	86.73 $\pm$ 1.60
	1	Success	5.00 $\pm$ 2.78	0.00 $\pm$ 0.00	31.67 $\pm$ 9.08	28.33 $\pm$ 4.35	41.67 $\pm$ 9.17	0.00 $\pm$ 0.00	17.78 $\pm$ 4.23
		Avg. Progress	43.86 $\pm$ 3.13	43.87 $\pm$ 4.43	91.32 $\pm$ 1.42	94.26 $\pm$ 0.62	84.79 $\pm$ 5.62	23.86 $\pm$ 1.53	63.66 $\pm$ 2.79
HDR	Full	Success	58.75 $\pm$ 4.50	78.00 $\pm$ 5.25	70.00 $\pm$ 10.12	33.33 $\pm$ 7.93	78.33 $\pm$ 9.67	43.33 $\pm$ 4.79	60.29 $\pm$ 7.04
		Avg. Progress	79.62 $\pm$ 3.04	97.18 $\pm$ 1.07	97.84 $\pm$ 0.61	93.69 $\pm$ 1.47	99.69 $\pm$ 0.31	69.34 $\pm$ 2.84	89.56 $\pm$ 1.56
	5	Success	58.75 $\pm$ 5.25	74.00 $\pm$ 6.93	75.00 $\pm$ 8.96	36.67 $\pm$ 8.26	75.00 $\pm$ 10.62	43.33 $\pm$ 3.54	60.46 $\pm$ 7.26
		Avg. Progress	79.38 $\pm$ 3.56	96.88 $\pm$ 1.20	97.88 $\pm$ 0.61	94.02 $\pm$ 1.59	97.61 $\pm$ 1.19	67.18 $\pm$ 2.45	88.83 $\pm$ 1.77
	1	Success	50.00 $\pm$ 6.05	0.00 $\pm$ 0.00	45.00 $\pm$ 12.03	35.00 $\pm$ 6.15	76.67 $\pm$ 7.23	1.67 $\pm$ 1.00	34.72 $\pm$ 5.41
		Avg. Progress	74.03 $\pm$ 2.74	57.10 $\pm$ 9.21	93.80 $\pm$ 1.11	95.30 $\pm$ 0.79	99.81 $\pm$ 0.19	32.35 $\pm$ 1.89	75.40 $\pm$ 2.66

Table 8: Hierarchical layer ablation. Values are reported as mean $\pm$ std. Results are organized by the number of active hierarchical layers, where the 1-layer setting corresponds to the causal baseline. Adding more layers generally improves reasoning performance, showing that each hierarchical level contributes to HDR’s coarse-to-fine reasoning.

Layers	Metric	Hanoi	Maze	One-line	Sliding	Sokoban	Water	Overall
1	Success	45.00 $\pm$ 6.73	12.00 $\pm$ 8.08	48.33 $\pm$ 7.66	21.67 $\pm$ 5.16	40.00 $\pm$ 10.69	38.33 $\pm$ 5.25	34.22 $\pm$ 7.26
	Avg. Progress	70.47 $\pm$ 2.76	55.04 $\pm$ 8.36	95.15 $\pm$ 1.28	86.13 $\pm$ 2.62	82.25 $\pm$ 5.79	66.94 $\pm$ 3.03	76.00 $\pm$ 3.97
2	Success	31.25 $\pm$ 4.19	44.00 $\pm$ 7.69	53.33 $\pm$ 11.81	20.00 $\pm$ 9.03	51.67 $\pm$ 6.16	30.00 $\pm$ 6.83	38.38 $\pm$ 7.62
	Avg. Progress	68.05 $\pm$ 2.86	80.60 $\pm$ 1.57	97.17 $\pm$ 0.57	86.83 $\pm$ 1.95	83.72 $\pm$ 0.71	59.17 $\pm$ 2.91	79.26 $\pm$ 1.76
3	Success	38.75 $\pm$ 4.58	22.00 $\pm$ 8.03	75.00 $\pm$ 8.24	30.00 $\pm$ 9.75	45.00 $\pm$ 7.24	30.00 $\pm$ 3.67	40.12 $\pm$ 6.92
	Avg. Progress	73.38 $\pm$ 3.27	84.76 $\pm$ 1.84	98.55 $\pm$ 0.46	87.79 $\pm$ 2.37	93.72 $\pm$ 0.66	66.35 $\pm$ 2.72	84.09 $\pm$ 1.89
4	Success	61.25 $\pm$ 7.69	72.00 $\pm$ 10.64	80.00 $\pm$ 11.62	33.33 $\pm$ 9.57	73.33 $\pm$ 10.71	35.00 $\pm$ 4.98	59.15 $\pm$ 9.20
	Avg. Progress	81.59 $\pm$ 3.57	93.33 $\pm$ 3.16	98.72 $\pm$ 0.60	91.23 $\pm$ 2.09	98.71 $\pm$ 3.18	66.71 $\pm$ 2.65	88.38 $\pm$ 2.54
5	Success	57.50 $\pm$ 8.22	76.00 $\pm$ 9.40	73.33 $\pm$ 8.10	33.33 $\pm$ 6.57	80.00 $\pm$ 8.91	41.67 $\pm$ 5.27	60.31 $\pm$ 7.74
	Avg. Progress	80.70 $\pm$ 5.10	95.40 $\pm$ 5.90	98.57 $\pm$ 0.93	92.65 $\pm$ 1.99	99.25 $\pm$ 4.23	69.32 $\pm$ 3.73	89.31 $\pm$ 3.65
6	Success	58.75 $\pm$ 4.50	78.00 $\pm$ 5.25	70.00 $\pm$ 10.08	33.33 $\pm$ 7.93	78.33 $\pm$ 9.67	43.33 $\pm$ 4.79	60.29 $\pm$ 7.04
	Avg. Progress	79.62 $\pm$ 3.04	97.18 $\pm$ 1.07	97.84 $\pm$ 0.46	93.69 $\pm$ 1.48	99.69 $\pm$ 0.31	69.34 $\pm$ 2.90	89.56 $\pm$ 1.54

and makes HDR’s reasoning stronger. Third, the data-reduction ablation shows that HDR degrades more gracefully than the bidirectional baseline as the amount of training data decreases. Together, the full tables confirm that HDR’s robustness comes from both its hierarchical coarse-to-fine structure and its ability to preserve useful reasoning behavior under limited denoising computation and limited data.

F Visualization of Hierarchical Intermediate Predictions

We provide qualitative visualizations of the intermediate predictions produced by HDR during hierarchical reasoning. As shown in Figure 8, the columns labeled L1, L2, and L3 correspond to the model’s clean x_0 predictions before noise is added at different hierarchy levels. These intermediate outputs reveal the coarse-to-fine inference process of HDR: higher layers first capture coarse task structure and global plans, while lower layers progressively refine them into concrete video states.

Table 9: Data reduction ablation. Values are reported as mean $\pm$ std. We train the bidirectional baseline and HDR using the full training set, 10% of the training set, and 2% of the training set. This experiment evaluates whether the models can learn reasoning rules from limited data.

Method	Data Scale	Metric	Hanoi	Maze	One-line	Sliding	Sokoban	Water	Overall
Bidirectional [19]	Full	Success	45.00 $\pm$ 6.78	90.00 $\pm$ 3.12	81.67 $\pm$ 7.60	31.67 $\pm$ 6.98	90.00 $\pm$ 4.21	21.67 $\pm$ 6.83	60.00 $\pm$ 5.92
		Avg. Progress	73.58 $\pm$ 2.95	99.89 $\pm$ 0.11	99.26 $\pm$ 0.27	93.00 $\pm$ 1.56	100.00 $\pm$ 0.00	57.08 $\pm$ 3.66	87.13 $\pm$ 1.42
	10%	Success	55.00 $\pm$ 5.36	80.00 $\pm$ 9.94	80.00 $\pm$ 9.13	26.67 $\pm$ 7.24	70.00 $\pm$ 6.67	20.00 $\pm$ 6.15	55.28 $\pm$ 7.42
		Avg. Progress	73.94 $\pm$ 3.00	95.08 $\pm$ 2.03	98.70 $\pm$ 0.59	93.45 $\pm$ 1.08	99.08 $\pm$ 0.90	57.20 $\pm$ 3.55	86.24 $\pm$ 1.86
	2%	Success	43.75 $\pm$ 6.74	20.00 $\pm$ 7.76	76.67 $\pm$ 9.27	16.67 $\pm$ 3.27	23.33 $\pm$ 7.42	6.67 $\pm$ 3.14	31.18 $\pm$ 6.27
		Avg. Progress	72.54 $\pm$ 2.59	67.76 $\pm$ 6.35	97.69 $\pm$ 0.66	92.51 $\pm$ 0.71	90.11 $\pm$ 3.32	47.03 $\pm$ 2.39	77.94 $\pm$ 2.67
HDR	Full	Success	58.75 $\pm$ 4.50	78.00 $\pm$ 5.25	70.00 $\pm$ 10.12	33.33 $\pm$ 7.93	78.33 $\pm$ 9.67	43.33 $\pm$ 4.79	60.29 $\pm$ 7.04
		Avg. Progress	79.62 $\pm$ 3.04	97.18 $\pm$ 1.07	97.84 $\pm$ 0.61	93.69 $\pm$ 1.47	99.69 $\pm$ 0.31	69.34 $\pm$ 2.84	89.56 $\pm$ 1.56
	10%	Success	56.25 $\pm$ 5.78	64.00 $\pm$ 9.95	68.33 $\pm$ 8.34	38.33 $\pm$ 5.86	80.00 $\pm$ 5.83	40.00 $\pm$ 5.57	57.82 $\pm$ 6.89
		Avg. Progress	79.26 $\pm$ 3.55	93.63 $\pm$ 2.05	97.91 $\pm$ 0.55	94.40 $\pm$ 0.82	99.92 $\pm$ 0.08	71.67 $\pm$ 2.17	89.46 $\pm$ 1.54
	2%	Success	60.00 $\pm$ 5.59	50.00 $\pm$ 13.43	50.00 $\pm$ 7.98	26.67 $\pm$ 5.71	70.00 $\pm$ 7.20	43.33 $\pm$ 2.49	50.00 $\pm$ 7.07
		Avg. Progress	77.71 $\pm$ 3.75	86.78 $\pm$ 6.54	96.40 $\pm$ 0.90	93.76 $\pm$ 0.68	98.82 $\pm$ 0.47	68.84 $\pm$ 2.18	87.05 $\pm$ 2.42

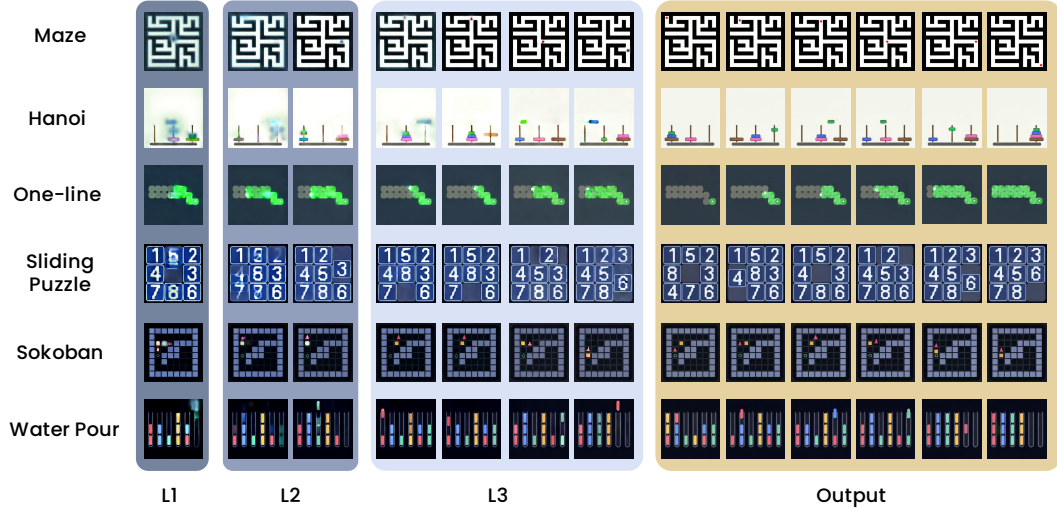


Figure 8: Visualization of HDR’s hierarchical intermediate predictions. Columns labeled L1, L2, and L3 show clean x_0 predictions before noise is added at different hierarchy levels, while the rightmost block shows the final generated output. The progression from coarse layers to fine layers illustrates how HDR first forms high-level plans and then refines them into detailed video states.

604 This supports our claim that hierarchical diffusion reasoning enables the model to form high-level
605 hypotheses before committing to final streaming outputs.

606 G Failure Case Studies

607 This appendix provides qualitative case studies that complement the quantitative results in the main
608 paper. Section G.1 compares HDR with the streaming autoregressive diffusion baseline and illustrates
609 how hierarchical denoising helps avoid early irreversible mistakes. Section G.2 then examines
610 representative residual failures of HDR, focusing on state-consistency drift in maze navigation,
611 Sokoban, and water pouring.

612 G.1 Qualitative Comparison with Streaming AR Diffusion

613 We first provide representative qualitative cases where the streaming autoregressive diffusion baseline,
614 represented by CausalForcing, fails while HDR succeeds. These examples complement the benchmark
615 results in the main paper by showing a recurring failure mode: the baseline makes an early, locally
616 plausible decision, but its irreversible temporal commitment prevents later frames from repairing the



Figure 9: Qualitative comparison across all six reasoning tasks where the streaming autoregressive diffusion baseline fails but HDR succeeds. For each task, the top row shows CausalForcing and the bottom row shows HDR. Red boxes and arrows mark the baseline’s failure points, while green boxes and arrows highlight HDR’s successful coarse-to-fine refinement. Across tasks, the baseline typically commits to an early local mistake, such as leaving the valid maze corridor, missing a required disk transfer, breaking path continuity, misplacing the blank tile, pushing the box away from the goal, or violating water-pouring constraints. In contrast, HDR maintains globally consistent task structure and reaches the correct final state.

trajectory. Once the rollout deviates from the correct plan, subsequent frames inherit the error and the video ends in a globally inconsistent state.

By contrast, HDR maintains revisable high-level hypotheses at coarse hierarchy levels and refines them through coarse-to-fine denoising before committing to the final streaming output. This hierarchical inference process enables implicit backtracking in latent space: the model can preserve uncertainty about multi-step structure, revise an incorrect partial plan, and instantiate a rule-consistent solution at finer levels. Figure 9 illustrates this behavior across maze navigation, Tower of Hanoi, one-line drawing, sliding puzzle, Sokoban, and water pouring.

G.2 Residual Failure Modes of HDR

Although HDR substantially improves multi-step video reasoning, it can still fail when exact state consistency must be preserved until the final frames. Figures 10, 11, and 12 show three representative residual failure modes. In maze navigation, HDR follows a plausible route but a wall disappears near the end of the rollout, making the final maze state invalid. In Sokoban, the box trajectory largely approaches the target, but a late wall-disappearance artifact breaks physical consistency. In water

631 pouring, the grouping process is mostly correct, yet one liquid color collapses into another, producing
632 a visually plausible but semantically invalid final arrangement.

633 These failures suggest that HDR’s remaining errors are less often complete planning failures and more
634 often late-stage state-consistency drift. The coarse hierarchy levels typically encode the intended
635 multi-step structure, and the finer levels refine this intent into nearly correct video states. The residual
636 errors appear near the end of rollout, where the model must preserve exact geometry, object identity,
637 or terminal constraints. This suggests that future improvements should focus on constraint-aware
638 refinement or lightweight verification during the final generation stage.

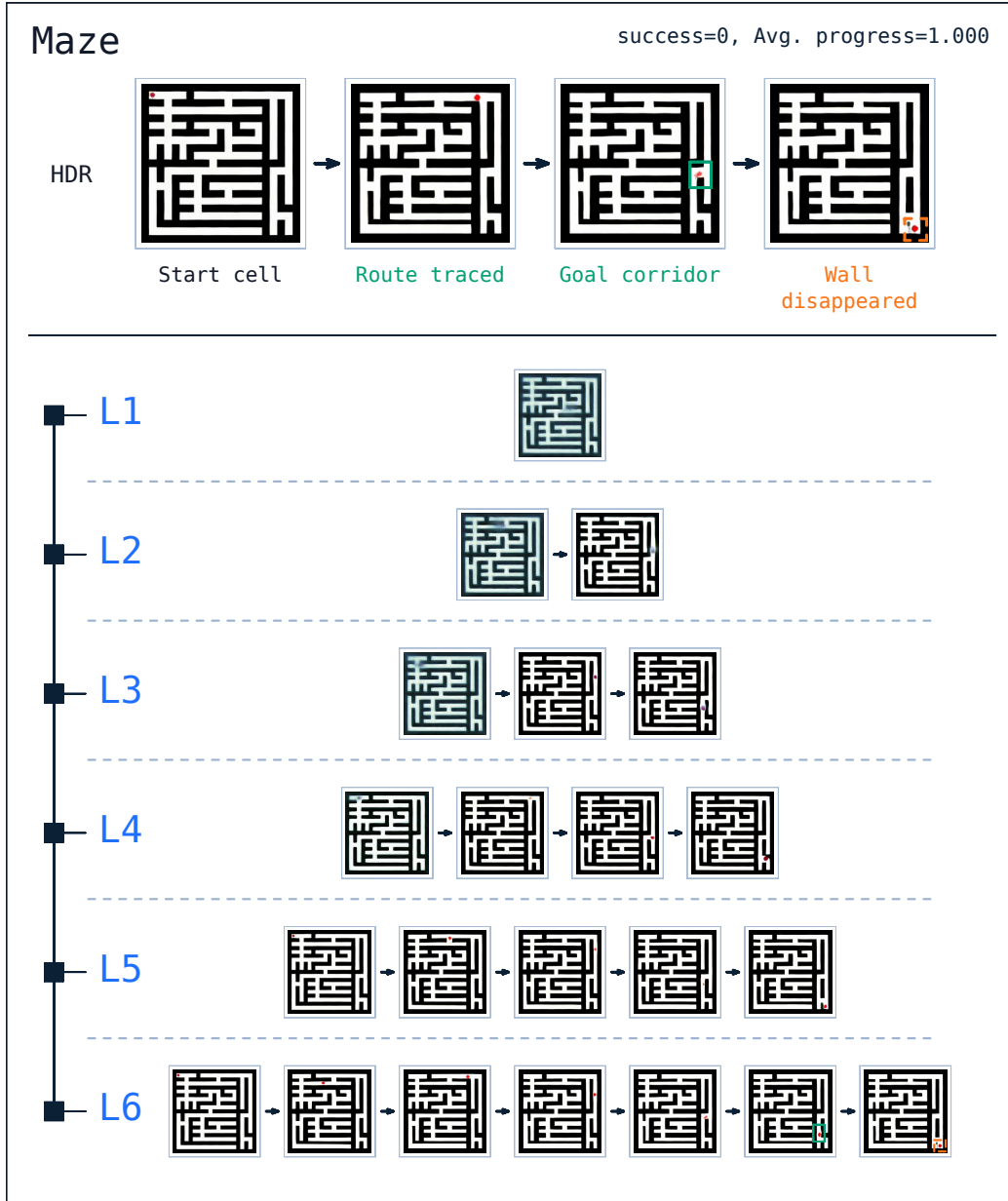


Figure 10: Maze failure case of HDR. The model traces a plausible route and reaches the goal corridor, but a wall disappears near the end of the rollout, producing an invalid maze state despite otherwise correct multi-step planning.

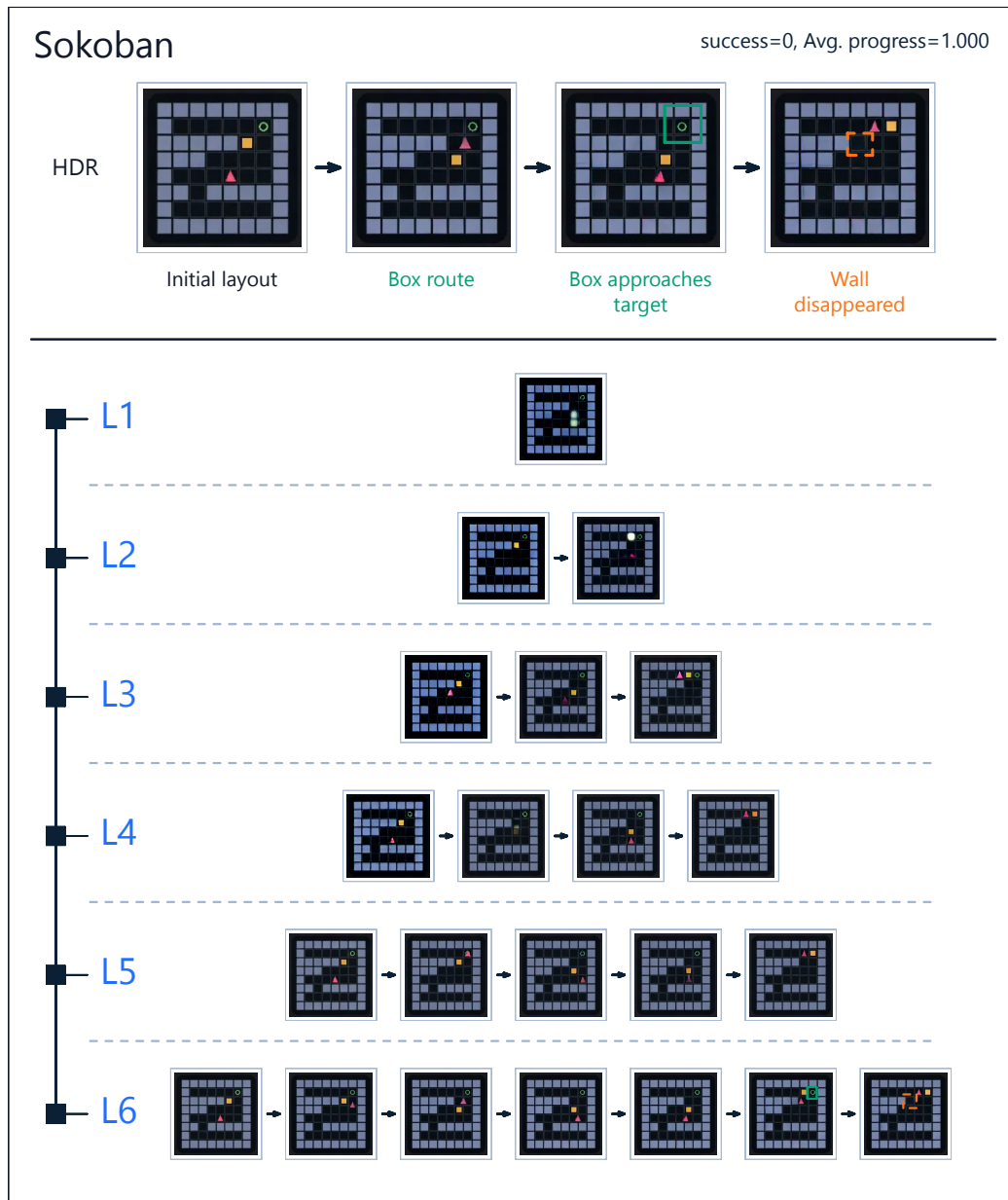


Figure 11: Sokoban failure case of HDR. The box trajectory is largely correct and approaches the target, but a wall disappears near the end of the rollout, making the final scene physically inconsistent.

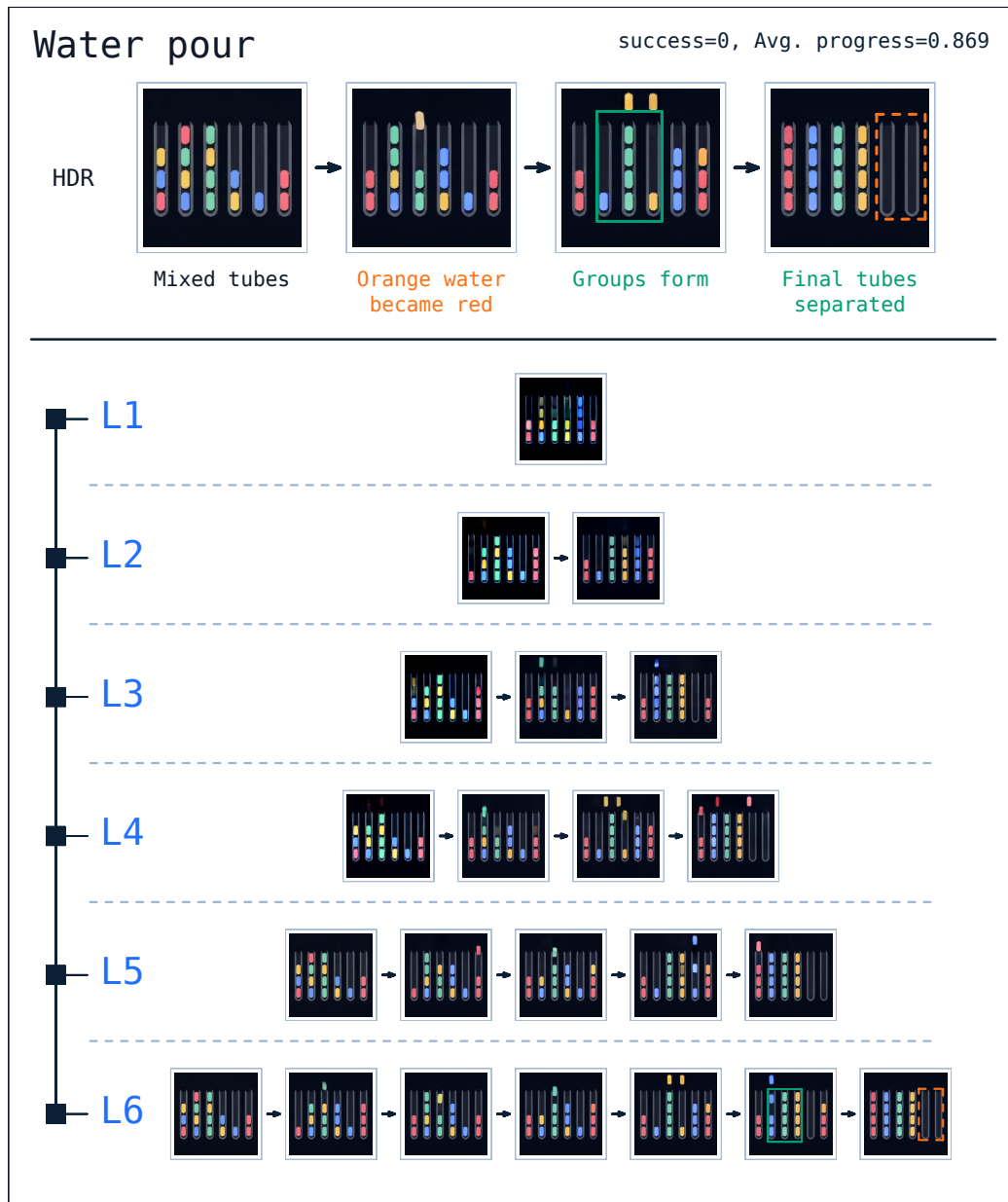


Figure 12: Water-pouring failure case of HDR. The model nearly completes the tube grouping correctly, but orange liquid collapses into red, yielding a visually plausible yet invalid final arrangement.

H Baselines and Implementation Details

Baselines. We compare HDR against three families of baselines. The first family consists of full-attention methods, including bidirectional diffusion and VideoMAE [18]. These methods allow global temporal interactions and therefore provide strong reasoning-oriented comparisons, but they require dense sequence processing. For VideoMAE, we enable image-to-video generation by randomly masking 90%–100% of tokens during training and masking all frames except the first frame at inference. The second family consists of non-full-attention methods, including causal diffusion, autoregressive generation, and VideoGPT [25]. These methods are more efficient and better aligned with streaming generation, but they are more vulnerable to irreversible early mistakes. For a controlled comparison, both the bidirectional diffusion baseline and the causal diffusion baseline are built on Wan2.2-5B-TI2V [19]. The bidirectional baseline follows the original full-attention inference paradigm of the backbone. The causal baseline replaces full temporal attention with a causal attention mask and is trained using the teacher-forcing strategy from Causal Forcing [33]. This setup isolates the effect of temporal attention and causal rollout while keeping the underlying generative backbone comparable.

HDR implementation. Our implementation is also built on Wan2.2-5B-TI2V and augments the backbone with the hierarchical latent tree described in Section 3. For the 125-frame setting used in the main experiments, we employ six hierarchy levels in latent space with sizes 1, 2, 4, 8, 16, 32 and the entropy-matched denoising schedule [5, 8, 13, 20, 32, 50]. Unless otherwise stated, this is the default setting for HDR. We provide the derivation and ablation of this entropy-matched schedule in Appendix B. For training, we use a mixed reasoning dataset containing 18,000 video clips in total, consisting of 3,000 synthetic videos for each of the six tasks. All videos are resized to 224×224 . We train from raw videos rather than precomputed latents: videos are decoded online and encoded by the Wan2.2 TI2V VAE during training. We condition on the first frame and allow this condition to be visible to all hierarchy levels. The model is initialized from a pretrained Wan2.2-5B-TI2V checkpoint and optimized with the flow-matching objective for 100,000 training steps. We use AdamW with learning rate 2×10^{-6} , betas (0, 0.999), weight decay 0.01, bfloat16 mixed-precision training, gradient checkpointing, and hybrid full-sharding FSDP. The per-device batch size is 1 and the total batch size is 8. During inference, we use classifier-free guidance with scale 3.0. Each main training run uses 8 GPUs and takes approximately 48 hours. All baselines are trained and evaluated under the same configuration for fair comparison.

I Existing Assets and Licenses

We use several existing research assets in this work, including Wan2.2-5B-TI2V as the video generation backbone, VideoMAE and VideoGPT as video modeling baselines, and CausalForcing as the streaming autoregressive diffusion baseline. We cite the original papers and repositories for all existing models and methods used in our experiments. Wan2.2-5B-TI2V is released under the Apache-2.0 license. We use all existing assets only for research, training, and evaluation purposes, and follow their corresponding license terms and usage restrictions. We do not redistribute third-party model checkpoints or code beyond what is permitted by their licenses. The supplementary material includes code adapted for our experiments together with attribution to the original projects and license files where applicable.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper’s contributions and scope?

Answer: [Yes].

Justification: The abstract and introduction state the central claims: HDR improves multi-step video reasoning while preserving low-latency streaming generation. These claims are supported by the main benchmark results, inference-speed comparison, data-efficiency experiments, and physical-world robot maze experiment reported in the paper.

Guidelines:

- The answer [N/A] means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [Yes].

Justification: The paper discusses residual failure modes in the appendix, including late-stage state-consistency drift and near-solved but invalid trajectories. The physical-world robot experiment is also framed as a limited transfer evaluation rather than a complete deployment study.

Guidelines:

- The answer [N/A] means that the paper has no limitation while the answer [No] means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate “Limitations” section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [N/A].

Justification: The paper does not present formal theoretical theorems or proofs. It includes algorithmic formulations, flow-matching objectives, attention-mask definitions, and temporal attention complexity analysis, but these are methodological descriptions rather than formal theoretical claims.

Guidelines:

- The answer [N/A] means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.

730 • The proofs can either appear in the main paper or the supplemental material.

731 **4. Experimental result reproducibility**

732 Question: Does the paper fully disclose all the information needed to reproduce the main ex-
733 perimental results of the paper to the extent that it affects the main claims and/or conclusions
734 of the paper?

735 Answer: [Yes].

736 Justification: The paper describes the benchmark tasks, evaluation metrics, baseline construc-
737 tion, model backbone, hierarchy levels, denoising schedule, training dataset size, optimizer,
738 training steps, and inference settings. Additional benchmark construction, task-specific
739 evaluation criteria, and implementation details are provided in the appendix.

740 Guidelines:

741 • The answer [N/A] means that the paper does not include experiments.

742 • If the paper includes experiments, a [No] answer to this question will not be perceived
743 well by the reviewers.

744 • If the contribution is a dataset and/or model, the authors should describe the steps taken
745 to make their results reproducible or verifiable.

746 • Reproducibility can be provided through code, data, model checkpoints, detailed
747 instructions, or hosted access.

748 **5. Open access to data and code**

749 Question: Does the paper provide open access to the data and code, with sufficient instruc-
750 tions to faithfully reproduce the main experimental results, as described in supplemental
751 material?

752 Answer: [Yes].

753 Justification: We provide training and inference code, as well as dataset generation and
754 evaluation code, in the supplementary material. These materials include the necessary
755 components to reproduce the main experimental results and benchmark evaluations.

756 Guidelines:

757 • The answer [N/A] means that paper does not include experiments requiring code.

758 • While release of code and data is encouraged, [No] is acceptable if justified.

759 • The instructions should contain the exact command and environment needed to repro-
760 duce the results if code is released.

761 • The authors should provide instructions on data access and preparation.

762 **6. Experimental setting/details**

763 Question: Does the paper specify all the training and test details necessary to understand the
764 results?

765 Answer: [Yes].

766 Justification: The paper specifies the task split, number of evaluation videos, success and
767 average-progress metrics, model backbone, hierarchy levels, denoising schedule, optimizer,
768 learning rate, batch size, training steps, precision, FSDP configuration, and inference
769 guidance scale. Full baseline and benchmark details are deferred to the appendix.

770 Guidelines:

771 • The answer [N/A] means that the paper does not include experiments.

772 • The experimental setting should be presented in the core paper to a level of detail
773 necessary to appreciate the results.

774 • Full details can be provided either with code, in appendix, or as supplemental material.

775 **7. Experiment statistical significance**

776 Question: Does the paper report error bars suitably and correctly defined or other appropriate
777 information about the statistical significance of the experiments?

778 Answer: [Yes].

Justification: The main tables report mean $\pm$ standard deviation for success and average progress, and the ablation figures show shaded bands denoting one standard deviation. These quantify variability across repeated runs or evaluation conditions used for the reported experiments.

Guidelines:

- The answer [N/A] means that the paper does not include experiments.
- The authors should answer [Yes] if the results are accompanied by error bars, confidence intervals, or statistical significance tests.
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- If error bars are reported, the paper should explain how they were calculated.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources needed to reproduce the experiments?

Answer: [Yes].

Justification: The paper reports the main training configuration, including mixed precision, gradient checkpointing, FSDP, per-device batch size, total batch size, and training duration. Each main training run uses 8 GPUs and takes approximately 48 hours.

Guidelines:

- The answer [N/A] means that the paper does not include experiments.
- The paper should indicate CPU/GPU type, memory, storage, execution time, and total compute.
- The paper should disclose whether the full research project required more compute than the reported experiments.

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics?

Answer: [Yes].

Justification: The work studies synthetic multi-step video reasoning and a controlled physical-world robot maze setup. We do not use human-subject data, personally identifiable information, or sensitive attributes, and the robot experiments are conducted in a constrained tabletop environment.

Guidelines:

- The answer [N/A] means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer [No], they should explain the special circumstances requiring deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity.

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [N/A].

Justification: This work focuses on synthetic multi-step video reasoning benchmarks and a controlled tabletop robot maze experiment. We do not study deployment in open-world settings or applications involving users, sensitive data, or high-stakes decisions; therefore, we regard broader societal impact discussion as not applicable beyond the standard considerations for generative video modeling.

Guidelines:

- The answer [N/A] means that there is no societal impact of the work performed.
- If the authors answer [N/A] or [No], they should explain why.

829 • The authors should consider harms from correct use, incorrect use, and misuse.
830 • If there are negative societal impacts, authors can discuss mitigation strategies.

831 **11. Safeguards**

832 Question: Does the paper describe safeguards that have been put in place for responsible
833 release of data or models that have a high risk for misuse?

834 Answer: [N/A].

835 Justification: The paper does not currently release a general-purpose pretrained video
836 generator or web-scraped dataset. The benchmark data are synthetic reasoning tasks and
837 controlled robot maze videos; if models, data, or code are released later, we will include
838 appropriate usage terms and safeguards.

839 Guidelines:

840 • The answer [N/A] means that the paper poses no such risks.
841 • Released models with high misuse risk should be released with necessary safeguards.
842 • Datasets scraped from the Internet could pose safety risks.

843 **12. Licenses for existing assets**

844 Question: Are the creators or original owners of assets used in the paper properly credited
845 and are the license and terms of use explicitly mentioned and properly respected?

846 Answer: [Yes].

847 Justification: We cite the original papers and repositories for all existing assets used in the
848 paper, including Wan2.2-5B-TI2V, VideoMAE, VideoGPT, and CausalForcing. Appendix I
849 summarizes their usage and license information, and we follow the corresponding license
850 terms and usage restrictions.

851 Guidelines:

852 • The answer [N/A] means that the paper does not use existing assets.
853 • The authors should cite the original paper that produced the code package or dataset.
854 • The name of the license should be included for each asset.
855 • For scraped data, copyright and terms of service should be provided.

856 **13. New assets**

857 Question: Are new assets introduced in the paper well documented and is the documentation
858 provided alongside the assets?

859 Answer: [Yes].

860 Justification: The paper introduces a new level-stratified multi-step video reasoning bench-
861 mark with OOD cases. The appendix documents task generation, difficulty levels, success
862 criteria, and average-progress metrics for each task.

863 Guidelines:

864 • The answer [N/A] means that the paper does not release new assets.
865 • Researchers should communicate the details of the dataset/code/model as part of their
866 submissions via structured templates.
867 • Documentation should include intended use, limitations, licensing, and maintenance
868 plan when assets are released.

869 **14. Crowdsourcing and research with human subjects**

870 Question: For crowdsourcing experiments and research with human subjects, does the paper
871 include the full text of instructions given to participants and screenshots, if applicable, as
872 well as details about compensation and IRB approval?

873 Answer: [N/A].

874 Justification: The paper does not involve crowdsourcing or human-subject experiments.
875 The benchmark videos are synthetic or collected from controlled robot maze experiments
876 without human participants.

877 Guidelines:

878 • The answer [N/A] means that the paper does not involve crowdsourcing nor research
879 with human subjects.
880 • If human subjects are involved, the paper should include instructions, consent, compen-
881 sation, and IRB details.

882 15. **Institutional review board (IRB) approvals or equivalent for research with human**
883 **subjects**

884 Question: Does the paper describe potential risks incurred by study participants, whether
885 such risks were disclosed to the subjects, and whether Institutional Review Board approval
886 was obtained?

887 Answer: [N/A].

888 Justification: The paper does not involve human-subject research or participant data. There-
889 fore, IRB approval for human-subject experiments is not applicable.

890 Guidelines:

891 • The answer [N/A] means that the paper does not involve human subjects.
892 • Depending on local regulation, human-subject research may require IRB approval or
893 exemption.