

# From Articulated Kinematics to Routed Visual Control for Action-Conditioned Surgical Video Generation

Anonymous Author(s)

Affiliation

Address

email

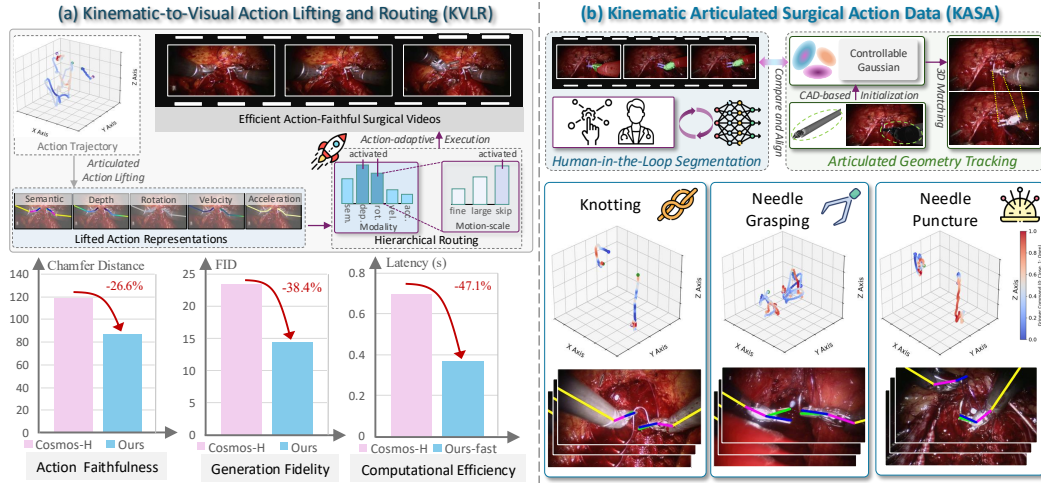


Figure 1: (a) KVLR transforms low-dimensional articulated kinematics into five image-aligned control modalities. A hierarchical routing mechanism selectively activates the most relevant modalities and motion scales, improving action faithfulness, visual fidelity, and computational efficiency. (b) KASA provides articulated action annotations curated from real robotic surgical videos.

## Abstract

1 Action-conditioned surgical video generation is a critical yet highly challenging  
 2 problem for robotic surgery. The core difficulty is that low-dimensional control  
 3 vectors must precisely govern complex image-space evolution. In this work, we  
 4 propose a **kinematic-to-visual** lifting paradigm that converts articulated kinematics  
 5 into a unified set of five image-aligned control modalities. Building on this repre-  
 6 sentation, we introduce a **hierarchically routed visual control** framework that  
 7 selectively activates the most relevant control modalities and motion scales. Instead  
 8 of uniformly applying all control signals, our model performs hierarchical routing  
 9 to dynamically allocate conditioning capacity. We further design kinematic-prior-  
 10 guided routing loss functions to ensure physically meaningful, temporally stable,  
 11 and efficient expert utilization. **To improve efficiency**, we propose a budgeted  
 12 training and inference scheme that leverages routing-induced sparsity. By selec-  
 13 tively discarding low-significance control pathways during training and execution,  
 14 our approach enables adaptive computation that is complementary to standard  
 15 distillation. We additionally construct a **new benchmark** with curated articulated  
 16 annotations, obtained through human-in-the-loop semantic labeling and differen-  
 17 tiable pose tracking, providing realistic supervision for action-conditioned surgical  
 18 video generation. **Extensive experiments** demonstrate that our method consistently  
 19 improves action faithfulness, visual fidelity, and cross-domain generalization over  
 20 diverse baselines. Moreover, our efficient variant achieves substantial reductions in  
 21 latency while maintaining strong control accuracy.

# 22 1 Introduction

23 Robotic surgery is entering an era where perception and control must be tightly integrated, making  
 24 action-conditioned surgical video generation a promising enabling tool [8, 25, 6, 3, 48, 34, 45]. Given  
 25 a desired surgical action sequence, the ability to synthesize realistic and physically consistent visual  
 26 outcomes has broad applications [3, 39, 17]. For instance, surgeons or learning systems can simulate  
 27 “what-if” scenarios (see Fig. 6) under different manipulation strategies. However, as illustrated in  
 28 Fig. 1 (a) bottom, this problem is fundamentally challenging: it requires simultaneously achieving  
 29 high visual realism and precise action faithfulness, where low-dimensional articulated control signals  
 30 must govern complex, high-dimensional image-space evolution.

31 The core difficulty lies in how to effectively introduce control signals into video generation [47, 21, 61].  
 32 Text-based conditioning [8, 17, 25], while flexible, is insufficient for surgical settings due to its  
 33 ambiguity and lack of geometric grounding. Prior efforts have explored dense visual control signals  
 34 (e.g., depth, segmentation, or flow) [6, 3, 39, 48], which improve controllability but rely on expensive  
 35 annotations or unavailable inputs in real robotic systems. This motivates the need for a unified  
 36 representation that is both physically grounded and visually aligned.

37 We propose KVLR, a **K**inematic-to-**V**isual **L**ifting and **R**outing method (Sec. 2.2) that transforms  
 38 articulated kinematics into visual control. Specifically, we lift low-dimensional articulated kinematics  
 39 into a pixel-aligned representation consisting of five interpretable modalities (semantics, depth,  
 40 rotation, velocity, and acceleration), forming a unified visual control basis (Fig. 1 (a) top). This  
 41 kinematic-to-visual action field enables the model to explicitly encode where the tool is, how it is  
 42 oriented, and how it moves. Compared to prior works [6, 39, 17], this representation is lightweight  
 43 and can be directly derived from articulated actions commonly available in robotic settings.

44 Building on this representation, KVLR allows a hierarchical routing mechanism (Sec. 2.3) that  
 45 dynamically allocates control capacity. As illustrated in Fig. 1 (a) top, instead of uniformly applying  
 46 all control modalities [57, 52, 11, 59, 29], our model performs modality-level routing to select the  
 47 most relevant control modalities, and further applies motion-scale routing to specialize computation  
 48 for fine manipulation, large transport motion, or static regions. To ensure that routing reflects physical  
 49 structure rather than arbitrary patterns, we design kinematic-prior-guided loss functions (Sec. 2.4)  
 50 that enforce physically meaningful expert utilization and temporal consistency.

51 An important product of hierarchical routing is the emergence of controllable sparsity, which we  
 52 leverage for efficient generation. As shown in Fig. 1 (a) top and Sec. 2.5, routing naturally identifies  
 53 low-significance control pathways that can be skipped without degrading performance. We exploit this  
 54 property through a budgeted training and inference scheme, combining action-aware distillation with  
 55 significance-driven execution. This enables the model to selectively discard unnecessary computation,  
 56 achieving substantial efficiency gains. Empirically, our efficient variant reduces latency by 47.1%  
 57 while maintaining strong control accuracy (Fig. 1 (a) bottom).

58 To support this task, we further construct KASA, a new benchmark for **K**inematic **A**rticulated  
 59 **S**urgical **A**ction. As illustrated in Fig. 1 (b) and Sec. 3, KASA provides curated annotations that  
 60 bridge real surgical videos and articulated control signals. Our pipeline combines human-in-the-loop  
 61 segmentation with differentiable pose tracking, enabling the extraction of part-aware geometry and  
 62 temporally consistent trajectories without requiring external robot logs. The dataset covers diverse  
 63 surgical actions such as knotting, needle grasping, and puncture, capturing both large-scale motion  
 64 and fine-grained manipulation, and serves as a realistic testbed for evaluating structured control.

65 Extensive experiments (Sec. 4) demonstrate that KVLR consistently improves action faithfulness,  
 66 visual fidelity, and generalization. As summarized in Fig. 1 (a) bottom, our method achieves  
 67 26.6% Chamfer Distance decrease and 38.4% FID decrease compared to strong baselines. KVLR  
 68 consistently outperforms text-conditioned, raw-action-conditioned, and dense-visual-conditioned  
 69 methods. Furthermore, KVLR shows favorable zero-shot transfer on unseen scenes.

70 In summary, this work makes the following contributions: (1) A kinematic-to-visual lifting paradigm  
 71 that transforms articulated kinematics into unified, image-aligned control modalities; (2) A hierar-  
 72 chical routing method that dynamically allocates control across modalities and motion scales; (3)  
 73 An action-adaptive efficient generation scheme that leverages routing-induced sparsity for signif-  
 74 icant computational acceleration; (4) KASA, a new benchmark with articulated annotations from  
 75 real surgical videos for evaluating action-conditioned generation; (5) Comprehensive experiments  
 76 demonstrating improved action faithfulness, visual fidelity, efficiency, and generalization.

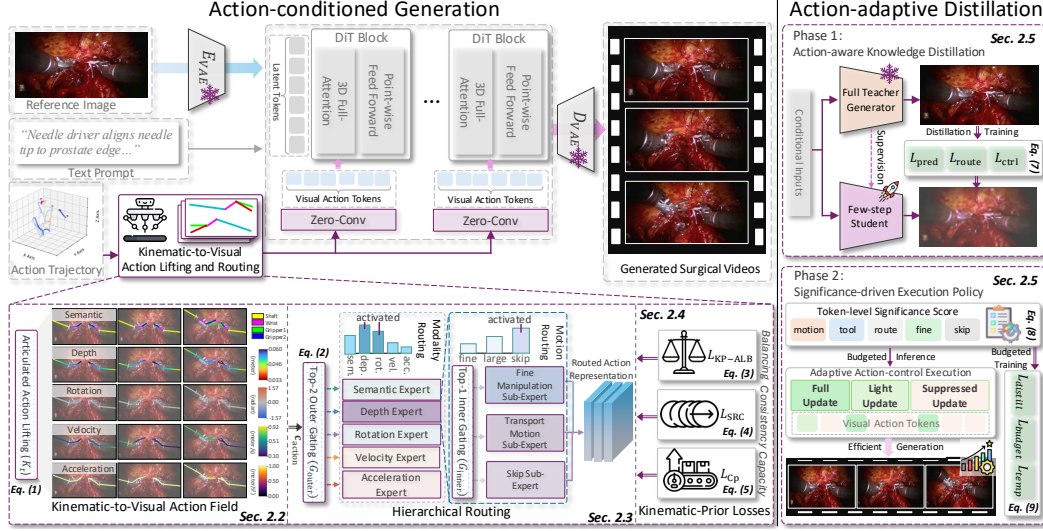


Figure 2: Architecture overview. Articulated kinematics are lifted into a pixel-aligned KVA-Field, then hierarchically routed across physical modalities and motion scales. This paradigm improves action faithfulness under heterogeneous dynamics and supports action-adaptive efficient generation.

## 2 Methodology

### 2.1 Framework Overview

As shown in Fig. 2 top left, given a reference image  $I^{\text{ref}}$ , a text prompt  $y$ , and an articulated action sequence  $a_{1:T}$  from the surgical robot, our goal is to generate a video  $\hat{x}_{1:T}$  conditioned on the specified surgical actions:  $\hat{x}_{1:T} = \mathcal{F}(I^{\text{ref}}, y, a_{1:T})$ . As illustrated in Fig. 2 bottom left, we first lift articulated kinematics into a pixel-aligned *Kinematic-to-Visual Action Field (KVA-Field)*, which encodes multi-modal cues directly in image space (Sec. 2.2). We then apply *hierarchical routing* to allocate conditional computation according to physical modalities and motion scales (Sec. 2.3). The kinematic-prior-guided routing loss functions are further designed to ensure physically meaningful, temporally stable, and efficient expert utilization (Sec. 2.4). This design reveals intrinsic conditional sparsity, which is later leveraged for action-adaptive efficient generation (Sec. 2.5).

### 2.2 Kinematic-to-Visual Action Field

To bridge the gap between low-dimensional control and dense visual synthesis, we lift articulated instrument actions into a pixel-aligned *Kinematic-to-Visual Action Field (KVA-Field)*. We model the surgical tool with a part-aware articulated structure composed of a shaft, a wrist, and two gripper components (as shown in Fig. 2 bottom left). For each frame  $t$ , the articulated action is represented as  $a_t = [p_t, r_t, q_t^{\text{sw}}, q_t^{\text{lg}}, q_t^{\text{rg}}] \in \mathbb{R}^9$ . Here,  $p_t \in \mathbb{R}^3$  and  $r_t \in \mathbb{R}^3$  denote the wrist translation and wrist orientation, respectively, with  $r_t$  expressed in axis-angle form. The scalar variables  $q_t^{\text{sw}}$ ,  $q_t^{\text{lg}}$ , and  $q_t^{\text{rg}}$  denote the shaft-to-wrist, wrist-to-left-gripper, and wrist-to-right-gripper joint states. Together, these variables are turned into a compact articulated action representation  $K_t \in \mathbb{R}^{H \times W \times 9}$ , with channels corresponding to part semantics, depth, orientation, velocity, and acceleration:

$$K_t(h, w) = [s_t(h, w), d_t(h, w), \rho_t(h, w), v_t(h, w), \alpha_t(h, w)]. \quad (1)$$

Here,  $s_t \in \mathbb{R}^3$  denotes part-aware semantic channels for shaft, wrist, and gripper;  $d_t \in \mathbb{R}$  is rendered depth;  $\rho_t \in \mathbb{R}$  is a local orientation descriptor;  $v_t \in \mathbb{R}^3$  is projected velocity; and  $\alpha_t \in \mathbb{R}$  is acceleration magnitude.  $v_t$  and  $\alpha_t$  are computed from consecutive articulated states:  $v_t = \frac{\Phi(M_t) - \Phi(M_{t-1})}{\Delta t}$ ,  $\alpha_t = \frac{\|v_t - v_{t-1}\|_2}{\Delta t}$ . Here,  $M_t$  is the articulated tool state and  $\Phi(\cdot)$  denotes its camera-projected feature representation on the image plane. **The resulting field encodes not only where the instrument is, but also how it is oriented and how it moves.** By explicitly separating semantic, geometric, and dynamic factors in image space, the KVA-Field provides the structured representation on top of which our hierarchical routing allocates conditional computation.

### 2.3 Hierarchical Routing

To account for the spatial heterogeneity and heavy-tailed distribution of surgical motion, we introduce the *Hierarchical Routing strategy*. Unlike standard Mixture-of-Experts (MoE) [36, 62] architectures

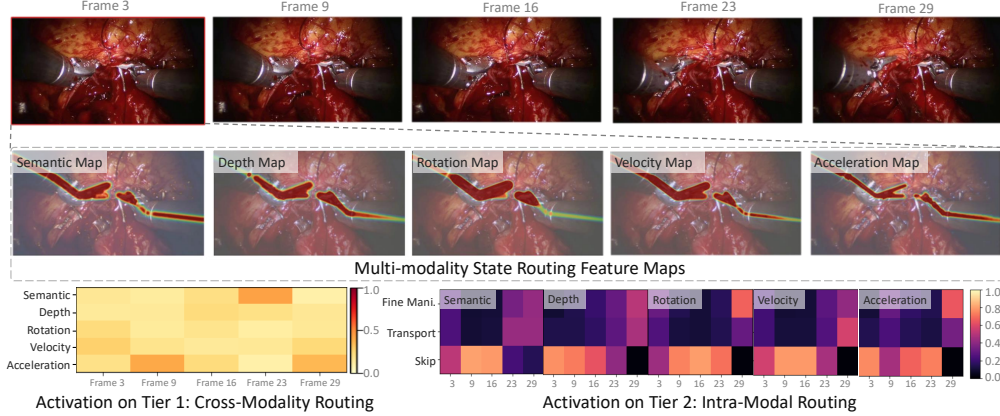


Figure 3: Visualization of hierarchical routing. We render the lifted action cues and show the learned Tier-1 modality and Tier-2 motion-scale activations. The routing patterns vary across motion phases, indicating action-dependent allocation of expert computation.

that route based solely on token content, our design employs a two-tier hierarchy. As shown in Fig. 2 bottom left, this architecture decouples the routing decision into two complementary axes: (1) *Which physical modality is dominant at this spatial location?* (Tier 1), and (2) *What scale of motion dynamics should be extracted within that modality?* (Tier 2). This hierarchical disentanglement allows the model to adaptively allocate computational resources to complex instrument dynamics.

**Tier 1: Multi-modality routing.** The outer layer consists of 5 modality-specific expert branches  $\mathcal{E}_{\text{mod}} = \{E_i\}_{i=1}^5$ . Each modality expert  $E_i$  is structurally constrained to process only its corresponding physical channels from the 9-channel input KVA-Field  $K_t$ . As shown in Fig. 2 bottom left, routing at this level is governed by an outer gating network  $G_{\text{outer}}$ . We encode  $K_t$  with a lightweight shared action module  $\mathcal{A}_{\text{act}}$  to obtain a holistic gating feature  $\mathbf{c}_{\text{action}} = \mathcal{A}_{\text{act}}(K_t)$ .  $G_{\text{outer}}$  is then computed as:

$$G_{\text{outer}}(\mathbf{c}_{\text{action}}, t) = \text{Softmax}(\text{Conv}_{1 \times 1}(\text{Concat}[\text{Pool}(\mathbf{c}_{\text{action}}), \tau(t)])), \quad (2)$$

where  $\text{Pool}$  and  $\tau(t)$  denote spatial average pooling and the diffusion timestep embedding, respectively. To ensure training stability, we employ a scheduling strategy [36, 62] that transitions from dense softmax fusion to sparse top- $k$  routing (default  $k = 2$ ).

**Tier 2: Motion-scale routing.** As illustrated in Fig. 2 bottom left, within each modality expert, we introduce a secondary layer to handle diverse motion dynamics. Each modality branch is composed of three sub-experts controlled by a top-1 inner gating mechanism  $G_{\text{inner}}$ : (1) **Fine Manipulation Sub-Expert**: Utilizes  $1 \times 1$  and dilated  $3 \times 3$  convolutions to capture high-frequency details required for precision tasks like knotting or grasping. (2) **Transport Motion Sub-Expert**: Employs large-kernel ( $7 \times 7$ ) depthwise convolutions and global context pooling to model fast tool sweeps and large displacements. (3) **Skip Sub-Expert**: An identity mapping branch that allows zero-cost computation for static instrument regions, preventing unnecessary feature transformation.

**Routing fusion.** The inner router uses a  $1 \times 1$  convolution to map the modality-specific feature vector at each spatial location to motion-scale logits. This produces an independent fine, large, or skip decision for each token, allowing the selected operator to adapt to the local kinematic state. The output of each modality branch is the selected sub-expert feature weighted by its inner-routing confidence. The final routed action representation is obtained by fusing the selected expert outputs.

**Hierarchical injection.** We then inject the fused control features into the video generator through zero-initialized additive layers on top of the DiT backbone, as shown in Fig. 2 top left. This preserves the pretrained generative prior at initialization while allowing the model to progressively learn action conditioning. In this way, the routing strategy allocates conditional computation according to both physical modality and motion complexity. Fig. 3 shows that the learned routing is motion-dependent: low-motion regions favor semantic and skip pathways, while larger motions increasingly activate dynamic modalities and transport-oriented computation, matching the design of hierarchical routing.

## 2.4 Kinematic-Prior Learning Loss Functions

The routing module in Sec. 2.3 should reflect the physical structure of surgical motion rather than purely statistical expert usage. However, as illustrated in Tab. 7, previous MoE [43, 10, 60] objectives, such as uniform load balancing, are poorly matched to the heavy-tailed and heterogeneous dynamics of articulated manipulation. We therefore introduce *kinematic-prior learning* loss functions to make hierarchical routing physically meaningful, temporally stable, and compatible with efficient inference.

**Kinematic-prior adaptive load balancing.** Instead of enforcing uniform expert utilization, we guide the modality router with a dynamic physical signal derived from the KVA-Field. The **motivation** is that different frames are dominated by different physical factors: (1) static tool regions may be dominated by semantic and depth cues, (2) wrist articulation by rotation cues, and (3) rapid tool motion by velocity or acceleration cues. Thus, the target expert usage should depend on the physical signal present in the current action field rather than being fixed across all samples. Let  $K_t^{\text{sem}} \in \mathbb{R}^{H \times W \times 3}$ ,  $K_t^{\text{dep}} \in \mathbb{R}^{H \times W}$ ,  $K_t^{\text{rot}} \in \mathbb{R}^{H \times W}$ ,  $K_t^{\text{vel}} \in \mathbb{R}^{H \times W \times 3}$ , and  $K_t^{\text{acc}} \in \mathbb{R}^{H \times W}$  denote the feature fields of depth, semantic, rotation, velocity and acceleration channels of the KVA-Field  $K_t$  at frame  $t$ . For each spatial token  $u \in \Omega$ , where  $\Omega$  is the image-token domain, we compute modality-wise physical magnitudes as  $e_t(u) = [\|K_t^{\text{sem}}(u)\|_2, |K_t^{\text{dep}}(u)|, |K_t^{\text{rot}}(u)|, \|K_t^{\text{vel}}(u)\|_2, |K_t^{\text{acc}}(u)|]^\top$ . These terms respectively measure part-presence strength, camera-depth magnitude, articulated rotation magnitude, projected motion magnitude, and acceleration magnitude. We aggregate them over spatial tokens to obtain the frame-level physical energy vector  $\mathcal{E}_i(K_t) = \frac{1}{|\Omega|} \sum_{u \in \Omega} e_{t,i}(u)$ ,  $i \in \mathcal{M}$ , where  $\mathcal{M} = \{\text{sem}, \text{dep}, \text{rot}, \text{vel}, \text{acc}\}$  is the set of modality experts.

The normalized physical signal is then  $\pi_i(K_t) = \frac{\mathcal{E}_i(K_t)}{\sum_{j \in \mathcal{M}} \mathcal{E}_j(K_t)}$ . This signal defines the desired routing mass for each modality expert according to the physical content of the current action field. Let  $P_i$  be the mean soft routing probability assigned to expert  $i$ , and let  $f_i$  be the fraction of tokens whose top-ranked routing assignment selects expert  $i$ . Following sparse MOE load estimation, we define the realized routing load as  $\ell_i = f_i P_i$  and minimize:

$$\mathcal{L}_{\text{KP-ALB}} = \sum_{i \in \mathcal{M}} (\ell_i - \pi_i(K_t))^2. \quad (3)$$

This objective encourages the router to allocate capacity according to physically meaningful action cues, rather than forcing all modality experts to be used uniformly.

**Spatiotemporal routing consistency.** Existing MoE diffusion models process tokens within a batch-level global pool, ignoring the temporal correlation inherent in sequential video frames [10, 60]. We introduce a Spatiotemporal Routing Consistency (SRC) loss to enforce temporal coherence in the continuous routing distributions, localized to moving instruments:

$$\mathcal{L}_{\text{SRC}} = \frac{1}{T-1} \sum_{t=1}^{T-1} \sum_{h,w} M_{\text{tool}}^{(t)}(h,w) \|R_t(h,w) - R_{t-1}(h,w)\|_2^2, \quad (4)$$

where  $R \in \mathbb{R}^{H \times W \times N}$  denotes the continuous routing probabilities at frame  $t$ , and  $M_{\text{tool}}^{(t)} \in \{0, 1\}^{H \times W}$  is the binary spatial mask indicating surgical tool presence (derived from  $K_t^{\text{sem}}$ ). By masking the loss to instrument regions, we avoid penalizing routing changes in static background areas while stabilizing expert selection for dynamic tools.

**Capacity predictor loss.** To enable efficient inference without computing full top- $k$  routing at runtime, we train a lightweight capacity predictor [36, 62] to anticipate routing decisions. The predictor receives detached action features and outputs raw logits  $\hat{R}$  for each expert. We optimize via Binary Cross-Entropy against the actual routing mask  $R^*$  generated by the gating networks:

$$\mathcal{L}_{\text{CP}} = \text{BCE}(\sigma(\hat{R}), \text{sg}(R^*)), \quad (5)$$

where  $\sigma(\cdot)$  is the sigmoid function. During inference, the predictor’s outputs are thresholded using EMA-updated per-expert quantiles, enabling dynamic routing with  $\mathcal{O}(1)$  complexity per spatial token. The thresholds are updated every training step to track the evolving routing distribution, preventing expert starvation from positive-feedback loops. More implementation details are provided in Sec. D.5.

**Overall objective.** The final objective combines the generation loss with the above routing priors:

$$\mathcal{L}_{\text{total}} = \mathcal{L}_{\text{Flow}} + \lambda_1 \mathcal{L}_{\text{KP-ALB}} + \lambda_2 \mathcal{L}_{\text{SRC}} + \lambda_3 \mathcal{L}_{\text{CP}}, \quad (6)$$

where  $\mathcal{L}_{\text{Flow}}$  is the standard flow-matching loss [30], and  $\lambda_{1:3}$  are weighting coefficients tuned as  $\lambda_1 = 0.01$ ,  $\lambda_2 = 0.005$ ,  $\lambda_3 = 0.01$ . More details on each loss function are presented in Sec. D.5.

## 2.5 Action-adaptive Efficient Generation

The structured control pathway improves action faithfulness and reveals where dense conditional computation is necessary. We exploit this property to build an *action-adaptive efficient generation* scheme.

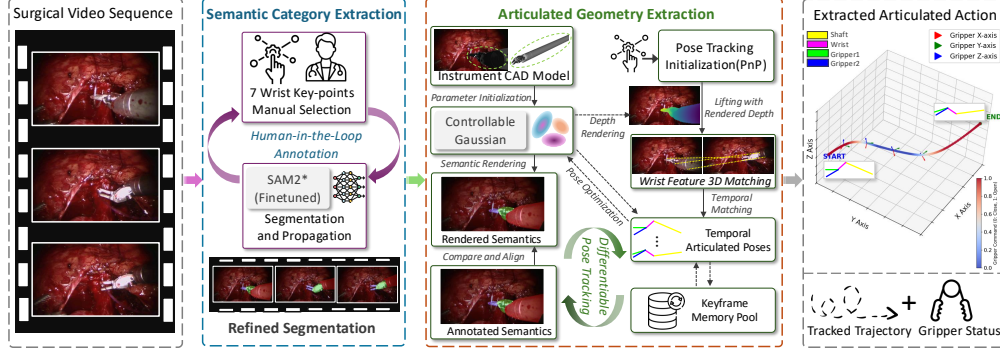


Figure 4: Data construction pipeline. Given robotic surgical videos, we obtain articulated action supervision through human-in-the-loop semantic annotation and differentiable pose tracking, producing lightweight yet visually grounded control signals for action-conditioned video generation.

**Phase 1: Action-aware distillation.** As shown in Fig. 2 top right, we first distill a few-step student generator from the full teacher under the same conditioning inputs. To preserve action-aware behavior, the student is supervised not only on the final prediction, but also on the internal control structure:

$$\mathcal{L}_{\text{distill}} = \mathcal{L}_{\text{pred}} + \lambda_r \mathcal{L}_{\text{route}} + \lambda_c \mathcal{L}_{\text{ctrl}}, \quad (7)$$

where  $\mathcal{L}_{\text{pred}}$  matches the teacher prediction,  $\mathcal{L}_{\text{route}}$  aligns routing behavior, and  $\mathcal{L}_{\text{ctrl}}$  preserves structured action-control features. More details on each loss function are provided in Sec. D.6. This transfers the teacher’s control policy into a few-step generator instead of distilling appearance alone. **Phase 2: Significance-driven execution.** We then derive a token-level action significance score from the lifted action field and routing signals:

$$s(u) = w_m e_{\text{motion}}(u) + w_t m_{\text{tool}}(u) + w_r c_{\text{route}}(u) + w_f q_{\text{fine}}(u) - w_s p_{\text{skip}}(u), \quad (8)$$

where  $u$  denotes a latent token;  $e_{\text{motion}}$ ,  $m_{\text{tool}}$ ,  $c_{\text{route}}$ ,  $q_{\text{fine}}$ , and  $p_{\text{skip}}$  denote normalized motion intensity, tool presence, routing confidence, fine-motion preference, and skip probability. The non-negative weights  $w_m$ ,  $w_t$ ,  $w_r$ ,  $w_f$ ,  $w_s$  control the relative contribution of each term. At inference,  $s(u)$  partitions latent tokens into three execution modes. As shown in Fig. 2 bottom-right, high-significance tokens perform a full control update by evaluating the KVA encoder, routers, and selected experts. Medium-significance tokens perform a light update by reusing cached routing decisions and only recomputing the final residual adapter. Low-significance tokens suppress the current control update and reuse the previous cached control residual. A frame-level refresh interval (detailed in Sec. D.6) prevents cache drift by forcing full updates at fixed temporal intervals.

**Budget-aware objective.** We optimize the efficient path with:

$$\mathcal{L}_{\text{eff}} = \mathcal{L}_{\text{distill}} + \lambda_b \mathcal{L}_{\text{budget}} + \lambda_t \mathcal{L}_{\text{temp}}, \quad (9)$$

where  $\mathcal{L}_{\text{budget}}$  enforces a target compute budget and  $\mathcal{L}_{\text{temp}}$  stabilizes reused or lightly updated regions over time. This objective reduces generation cost while preserving action-conditioned control. More details are provided in the Sec. D.6 of the appendix.

### 3 KASA Benchmark Construction

**Benchmark overview.** A key obstacle in action-conditioned surgical video generation is the lack of benchmarks pairing real surgical videos with explicit articulated controls. Existing works [8, 17, 25, 6, 3, 39] typically provide coarse language descriptions or dense visual annotations, but not lightweight action supervision aligned with visual generation. We construct the Kinematic Action-centric Surgical (KASA) benchmark for articulated tool-action-conditioned generation, associating each video with video-reconstructed annotations. These compact, visually grounded controls are compatible with robotic settings and aligned with structured control signals for kinematic-to-visual action modeling.

**Task-driven sequence selection.** We curate KASA from the SAR-RARP dataset [38], collecting 1,047 surgical scenes with 105,175 frames. The benchmark focuses on three representative sub-actions of *needle grasping*, *needle puncture*, and *knotting*, which together cover a broad range of surgical dynamics, from large transport motion to fine local manipulation. Videos are downsampled to 30 FPS, and each sequence contains 55-235 frames. KASA exhibits articulated motion, frequent self-occlusion, and complex instrument-tissue interactions, providing a challenging test platform for evaluating both action faithfulness and structured control under heterogeneous surgical dynamics.

**Human-in-the-loop semantic annotation.** To obtain reliable part-aware supervision, we adopt a Human-in-the-loop Semantic Annotation (HSA) procedure (Fig. 4 left). We first estimate part-level

Table 1: Action faithfulness on KASA. All methods are finetuned under the same split.

| Method                                      | Knotting           |                  |                  | NeedleGrasping    |                  |                   | NeedlePuncture    |                  |                  |
|---------------------------------------------|--------------------|------------------|------------------|-------------------|------------------|-------------------|-------------------|------------------|------------------|
|                                             | CD ↓               | TI ↑             | AF ↓             | CD ↓              | TI ↑             | AF ↓              | CD ↓              | TI ↑             | AF ↓             |
| <i>Text-conditioned Generation</i>          |                    |                  |                  |                   |                  |                   |                   |                  |                  |
| Wan-2.2                                     | 169.35±3.12        | 0.56±0.02        | 6.54±0.45        | 135.80±2.85       | 0.57±0.02        | 26.95±1.20        | 98.15±2.15        | 0.65±0.02        | 5.31±0.38        |
| HunyuanVideo                                | 162.81±2.95        | 0.57±0.01        | 6.07±0.50        | 131.63±2.60       | 0.59±0.02        | 26.05±1.15        | 94.11±1.95        | 0.67±0.01        | 4.82±0.35        |
| OpenSora-2.0                                | 173.42±3.40        | 0.53±0.02        | 7.02±0.62        | 140.56±3.10       | 0.55±0.02        | 28.35±1.40        | 102.67±2.45       | 0.63±0.02        | 5.72±0.45        |
| Cosmos-H-Surgical                           | 150.26±2.80        | 0.59±0.02        | 4.88±0.35        | 119.38±2.45       | 0.61±0.02        | 23.64±1.05        | 85.30±1.80        | 0.71±0.01        | 3.44±0.25        |
| <i>Raw-action-conditioned Generation</i>    |                    |                  |                  |                   |                  |                   |                   |                  |                  |
| Wan-2.2*                                    | 135.48±2.55        | 0.68±0.02        | 5.23±0.40        | 108.64±2.15       | 0.70±0.02        | 21.56±0.95        | 78.52±1.65        | 0.78±0.02        | 4.25±0.30        |
| HunyuanVideo*                               | 130.25±2.40        | 0.70±0.01        | 4.86±0.35        | 105.31±2.05       | 0.72±0.01        | 20.84±0.85        | 75.29±1.50        | 0.80±0.01        | 3.86±0.25        |
| OpenSora-2.0*                               | 138.74±2.65        | 0.65±0.02        | 5.62±0.45        | 112.45±2.25       | 0.68±0.02        | 22.68±1.10        | 82.14±1.75        | 0.76±0.02        | 4.58±0.35        |
| Cosmos-H-Surgical*                          | 120.21±2.34        | 0.74±0.01        | 3.91±0.22        | 95.51±1.86        | 0.76±0.02        | 18.91±0.74        | 68.24±1.41        | 0.84±0.01        | 2.75±0.18        |
| <i>Visual-action-conditioned Generation</i> |                    |                  |                  |                   |                  |                   |                   |                  |                  |
| SurgSora <sup>†</sup>                       | 119.50±2.10        | 0.75±0.02        | 3.85±0.20        | 94.80±1.65        | 0.77±0.01        | 18.50±0.65        | 67.50±1.50        | 0.85±0.02        | 2.65±0.15        |
| OpenSora-2.0 <sup>†</sup>                   | 113.80±1.95        | 0.84±0.02        | 1.51±0.15        | 91.10±1.80        | 0.87±0.02        | 12.20±0.55        | 61.40±1.35        | 0.87±0.01        | 1.88±0.12        |
| KVLR-fast(ours)                             | 114.32±1.91        | 0.85±0.01        | 1.54±0.14        | 91.85±1.48        | 0.88±0.01        | 12.45±0.51        | 63.45±1.12        | 0.86±0.01        | 1.95±0.12        |
| KVLR(ours)                                  | <b>110.57±1.62</b> | <b>0.88±0.01</b> | <b>1.18±0.10</b> | <b>89.22±1.35</b> | <b>0.91±0.01</b> | <b>10.67±0.43</b> | <b>60.90±0.96</b> | <b>0.89±0.01</b> | <b>1.59±0.09</b> |

Table 2: Generation Fidelity on KASA. PSNR/SSIM are averaged over generated frames, FID is computed on generated frames as images, and FVD on 17-frame clips.

| Method                                      | Knotting          |                  |                   |                    | NeedleGrasping    |                  |                   |                    | NeedlePuncture    |                  |                   |                   |
|---------------------------------------------|-------------------|------------------|-------------------|--------------------|-------------------|------------------|-------------------|--------------------|-------------------|------------------|-------------------|-------------------|
|                                             | PSNR ↑            | SSIM ↑           | FID ↓             | FVD ↓              | PSNR ↑            | SSIM ↑           | FID ↓             | FVD ↓              | PSNR ↑            | SSIM ↑           | FID ↓             | FVD ↓             |
| <i>Text-conditioned Generation</i>          |                   |                  |                   |                    |                   |                  |                   |                    |                   |                  |                   |                   |
| Wan-2.2                                     | 16.20±0.25        | 0.55±0.02        | 45.30±2.40        | 420.50±18.50       | 15.30±0.22        | 0.52±0.02        | 48.20±2.60        | 450.30±20.10       | 16.80±0.26        | 0.58±0.02        | 40.10±2.10        | 310.40±15.50      |
| HunyuanVideo                                | 16.50±0.22        | 0.57±0.01        | 42.10±2.15        | 395.20±16.20       | 15.60±0.20        | 0.54±0.01        | 44.50±2.35        | 415.80±18.40       | 17.10±0.24        | 0.60±0.01        | 37.20±1.95        | 285.20±14.20      |
| OpenSora-2.0                                | 16.10±0.28        | 0.54±0.02        | 47.50±2.55        | 445.80±19.80       | 15.10±0.25        | 0.51±0.02        | 50.40±2.80        | 470.20±22.50       | 16.60±0.25        | 0.56±0.02        | 42.30±2.25        | 335.60±16.80      |
| Cosmos-H-Surgical                           | 17.20±0.20        | 0.62±0.01        | 36.50±1.90        | 325.40±14.50       | 16.40±0.18        | 0.59±0.01        | 38.50±2.10        | 355.50±15.80       | 17.80±0.22        | 0.65±0.01        | 31.50±1.75        | 245.70±12.40      |
| <i>Raw-action-conditioned Generation</i>    |                   |                  |                   |                    |                   |                  |                   |                    |                   |                  |                   |                   |
| Wan-2.2*                                    | 17.80±0.24        | 0.66±0.02        | 31.60±1.85        | 280.20±13.50       | 16.90±0.21        | 0.63±0.02        | 33.10±2.05        | 310.80±14.80       | 18.20±0.25        | 0.68±0.02        | 26.30±1.65        | 200.50±11.50      |
| HunyuanVideo*                               | 18.10±0.22        | 0.68±0.01        | 28.40±1.75        | 255.40±12.80       | 17.30±0.19        | 0.65±0.01        | 30.20±1.95        | 285.90±13.60       | 18.60±0.23        | 0.71±0.01        | 23.40±1.55        | 175.40±10.80      |
| OpenSora-2.0*                               | 17.50±0.26        | 0.64±0.02        | 34.20±2.10        | 305.60±15.20       | 16.60±0.24        | 0.61±0.02        | 36.40±2.25        | 335.70±16.50       | 17.90±0.27        | 0.66±0.02        | 29.10±1.85        | 225.80±12.20      |
| Cosmos-H-Surgical*                          | 18.90±0.23        | 0.72±0.01        | 24.46±1.81        | 225.72±12.35       | 18.20±0.18        | 0.70±0.01        | 25.87±1.94        | 251.26±15.18       | 19.14±0.20        | 0.75±0.01        | 19.72±1.67        | 142.97±10.74      |
| <i>Visual-action-conditioned Generation</i> |                   |                  |                   |                    |                   |                  |                   |                    |                   |                  |                   |                   |
| SurgSora <sup>†</sup>                       | 18.85±0.21        | 0.73±0.02        | 24.15±1.50        | 222.40±11.50       | 18.35±0.15        | 0.71±0.01        | 25.40±1.65        | 248.50±13.20       | 19.25±0.18        | 0.76±0.01        | 19.35±1.45        | 140.60±9.80       |
| OpenSora-2.0 <sup>†</sup>                   | 20.75±0.25        | 0.79±0.01        | 16.80±1.60        | 168.90±10.80       | 19.90±0.20        | 0.75±0.02        | 18.95±1.75        | 185.20±14.50       | 20.75±0.15        | 0.83±0.01        | 12.90±1.30        | 97.50±8.90        |
| KVLR-fast(ours)                             | 20.45±0.17        | 0.78±0.01        | 18.25±1.12        | 185.32±10.46       | 19.68±0.15        | 0.76±0.01        | 20.15±1.18        | 205.40±12.33       | 20.55±0.16        | 0.81±0.01        | 14.60±0.82        | 112.50±7.64       |
| KVLR(ours)                                  | <b>21.09±0.14</b> | <b>0.82±0.01</b> | <b>14.78±0.89</b> | <b>154.56±8.72</b> | <b>20.17±0.13</b> | <b>0.80±0.01</b> | <b>16.60±0.91</b> | <b>168.94±9.85</b> | <b>21.22±0.12</b> | <b>0.85±0.01</b> | <b>11.75±0.68</b> | <b>87.10±5.96</b> |

masks for the instrument shaft, wrist, and grippers using a SAM2 model [40] finetuned on SAR-RARP, and then refine the propagated predictions through interactive human correction. This step removes part confusion, boundary leakage, and temporal inconsistency, yielding clean semantic labels that are crucial for disentangling different physical factors in our structured control representation.

**Differentiable pose tracking.** To recover continuous articulated geometry, we further introduce Differentiable Pose Tracking (DPT), which estimates temporally consistent pose trajectories directly from videos without relying on external kinematic logs or hardware tracking (Fig. 4 right). DPT uses a render-and-compare formulation initialized from instrument CAD priors and aligns rendered part-aware geometry with the observed video evidence over time. The resulting articulated trajectories provide the geometric and motion supervision, enabling downstream construction of semantic, depth, rotation, and motion-aligned control signals for action-conditioned video generation. To further validate KASA, Sec. 4.2 and Sec. E report independent-evaluator checks and annotation quality validations.

## 4 Experiment

### 4.1 Experimental Setup

**Training.** We instantiate our models based on OpenSora-2.0 [61]. The models are trained on NVIDIA H20 GPUs with a batch size of 24 using AdamW [33]. The learning rate is set to  $1 \times 10^{-5}$  with weight decay of 0.01. We split KASA into 85% training and 15% validation sets. The videos are generated at a default resolution of  $576 \times 1024$  with clip lengths aligned with the ground-truth sequences.

**Baselines.** Our full model uses 50 denoising steps, while *KVLR-fast* denotes the distilled efficient variant with a 4-step student. We compare with Wan-2.2 [47], HunyuanVideo [21], OpenSora-2.0 [61], Cosmos-H-Surgical [17], and SurgSora [6], all finetuned on KASA under the same split and input protocol. Models marked with ‘\*’ use ControlNet-style injection (same parameter scale as ours) with raw-action (articulated trajectory vectors) conditions, while models with ‘<sup>†</sup>’ leverage dense visual conditions as SurgSora [6] (*i.e.*, depth, semantic, and flow maps). The results are reported with the mean and standard deviation over three runs. Additional details are provided Sec. D of the appendix.

### 4.2 Main Results

**Action faithfulness.** We evaluate Chamfer Distance (CD), Temporal IoU (TI), and Area Flicker (AF) using an independently trained Grounded SAM [41], comparing predicted tool masks with

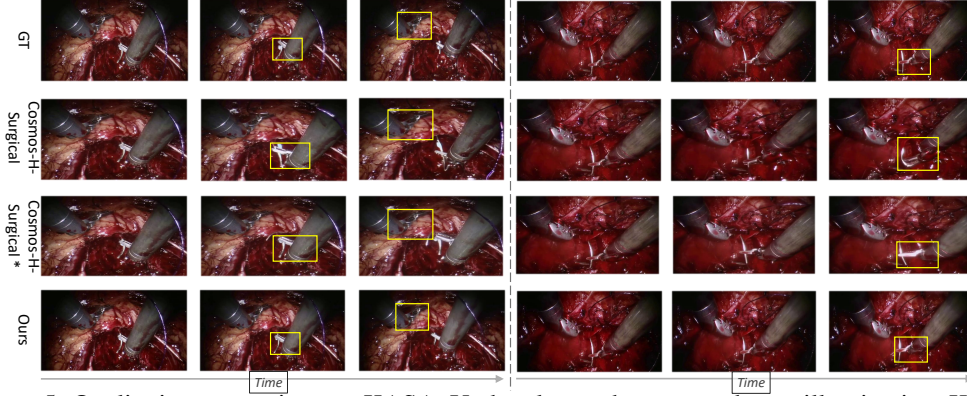


Figure 5: Qualitative comparison on KASA. Under cluttered scenes and poor illumination, KVLR better preserves tool states and action-consistent evolution than Cosmos-H-Surgical\*.

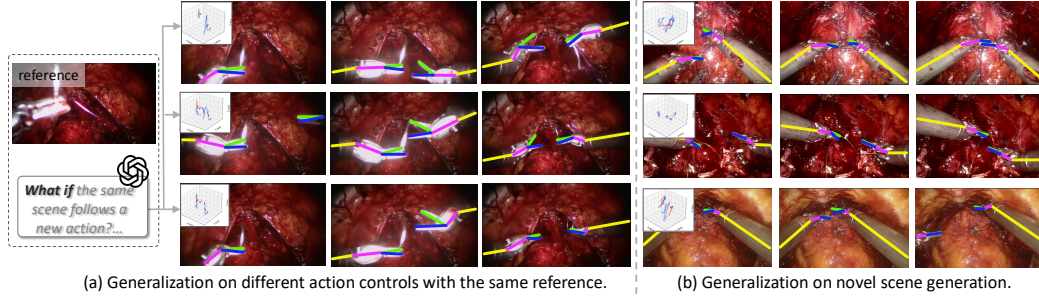


Figure 6: Generalization results. (a) Controllable synthesis under user-specified actions with the same reference. (b) Qualitative transfer to unseen scenes with different appearances and motion patterns.

ground-truth action masks; details are in Sec. D.1. As shown in Tab. 1, text-only baselines struggle to follow surgical actions, and raw-action injection improves control but remains limited. Dense-visual-condition baselines further improve action precision using semantic, depth, and flow maps, but such inputs are often unavailable in robotic settings and still underperform KVLR. *KVLR-fast* retains compelling control fidelity, showing that action alignment is largely preserved under reduced computation. **Generation fidelity.** Tab. 2 shows consistent gains in visual fidelity. Direct action injection improves over text-only generation, while structured visual-action-condition baselines further benefit from dense scene-level guidance. KVLR achieves the best overall performance, particularly on *Knotting* and *NeedleGrasping*, where realistic tool-tissue evolution requires accurate local motion control. This suggests that kinematic-to-visual lifting and routing provide benefits beyond simply introducing dense visual conditions. Fig. 5 further qualitatively confirms that KVLR better preserves action-consistent evolution and interaction outcomes under cluttered scenes and poor illumination.

**Computational efficiency.** Tab. 3 reports latency and FLOPs on H20 GPUs with 17-frame clips. Full KVLR prioritizes action faithfulness and fidelity, whereas KVLR-fast targets efficient inference. At both resolutions, KVLR-fast achieves the lowest latency and FLOPs while retaining strong action alignment, showing that routing-derived significance exposes useful sparsity for adaptive acceleration.

**Cross-dataset generalization.** We evaluate zero-shot transfer on GraSP [46, 2], which differs in tissue appearance, illumination, and viewpoint. Kinematics are extracted by our DPT pipeline without human correction, introducing realistic tracking noise. Without fine-tuning, KVLR maintains favorable faithfulness and fidelity over the evaluated baselines (Tab. 4), suggesting improved robustness to domain shift. This trend is consistent with the qualitative transfer results in Fig. 6 (b) and the “what-if” action control examples in Fig. 6 (a).

Table 3: Computational efficiency evaluation.

| Method                    | 288×512      |                     | 576×1024     |                     |
|---------------------------|--------------|---------------------|--------------|---------------------|
|                           | Latency (s)↓ | FLOPs↓<br>(PF/clip) | Latency (s)↓ | FLOPs↓<br>(PF/clip) |
| Wan-2.2*                  | 2.38         | 3.42                | 11.85        | 12.18               |
| HunyuanVideo*             | 6.12         | 9.15                | 30.54        | 32.45               |
| Cosmos-H-Surgical*        | 0.70         | 1.91                | 4.32         | 7.55                |
| SurgSora <sup>†</sup>     | 3.85         | 5.42                | 18.25        | 19.65               |
| OpenSora-2.0 <sup>†</sup> | 5.12         | 7.55                | 25.60        | 26.85               |
| KVLR (Ours)               | 4.64         | 6.94                | 23.39        | 24.77               |
| KVLR-fast (Ours)          | <b>0.37</b>  | <b>0.55</b>         | <b>1.88</b>  | <b>1.98</b>         |

Table 4: Cross-dataset generalization on GraSP.

| Method                    | CD↓              | TI↑              | AF↓              | PSNR↑           | FID↓            |
|---------------------------|------------------|------------------|------------------|-----------------|-----------------|
| Cosmos-H-Surgical*        | 134.8±3.9        | 0.68±0.03        | 8.42±0.58        | 16.9±0.3        | 28.3±1.8        |
| SurgSora <sup>†</sup>     | 127.5±3.4        | 0.71±0.02        | 7.15±0.45        | 17.6±0.2        | 24.0±1.5        |
| OpenSora-2.0 <sup>†</sup> | 115.3±3.1        | 0.77±0.01        | 4.78±0.38        | 18.9±0.2        | 17.0±1.3        |
| KVLR-fast (Ours)          | 117.6±2.8        | 0.76±0.02        | 5.05±0.31        | 18.7±0.2        | 19.1±1.1        |
| KVLR (Ours)               | <b>111.9±2.5</b> | <b>0.79±0.01</b> | <b>4.28±0.26</b> | <b>19.2±0.1</b> | <b>16.7±0.9</b> |

Table 6: Architecture ablation under parameter-matched control branches.

| Method                                 | Total Params | CD ↓         | TI ↑        | FID ↓        | PSNR ↑       | Latency (s) ↓ |
|----------------------------------------|--------------|--------------|-------------|--------------|--------------|---------------|
| Text-only baseline                     | 11.00B       | 138.88       | 0.57        | 93.31        | 12.80        | <b>4.59</b>   |
| + Raw action vector + direct condition | 12.06B       | 111.11       | 0.69        | 79.56        | 15.08        | 4.61          |
| + KVA-Field + direct condition         | 12.06B       | 96.45        | 0.78        | 45.20        | 17.65        | 4.68          |
| + KVA-Field + Tier-1 routing           | 12.05B       | 91.30        | 0.84        | 26.85        | 19.34        | 4.66          |
| + KVA-Field + Tier-1 + Tier-2 routing  | 12.06B       | <b>86.89</b> | <b>0.89</b> | <b>14.37</b> | <b>20.82</b> | 4.64          |

**Trajectory-level verification.** Rather than only computing the mask-overlap metrics, we provide an additional trajectory-level evaluation that extracts the visible tool centerline and distal endpoint from each generated video using an independently trained Grounded SAM-based tracker [41], and compares them with the projected trajectory induced by the reference articulated action. As shown in Tab. 5, KVLR reduces both centerline and endpoint errors and improves motion-direction accuracy over Cosmos-H-Surgical\*, indicating that its gains are not specific to the default mask-overlap protocol. Action-perturbation checks in appendix F.2 further confirm that the generated motion follows the specified action trajectory.

Table 5: Trajectory-level verification. Ctr. and End. are centerline and distal-endpoint errors in pixels; Dir. is motion-direction accuracy(%).

| Method             | Ctr.↓       | End.↓       | Dir.↑       |
|--------------------|-------------|-------------|-------------|
| Cosmos-H-Surgical* | 11.83       | 15.62       | 71.4        |
| KVLR-fast (Ours)   | 8.05        | 10.24       | 82.1        |
| KVLR (Ours)        | <b>7.21</b> | <b>9.13</b> | <b>84.8</b> |

### 4.3 Ablation Study

**Architecture design.** Tab. 6 ablates KVLR under parameter-matched control branches. Despite the matched capacity, direct conditioning with raw action yields limited gains, suggesting that the bottleneck is the fundamental representational mismatch between low-dimensional articulated kinematics and dense visual evolution. Replacing raw actions with the KVA-Field bridges this gap, driving a substantial performance leap (*e.g.*, CD drops to 96.45). Adding hierarchical routing dynamically allocates this fixed parameter budget across varying physical modalities and motion scales, further boosting action faithfulness without inflating latency, showing that action lifting and structured routing are complementary. Qualitative comparisons on different settings are provided in Fig. 8 of the appendix.

**Kinematic-prior losses.** Tab. 7 shows that the three priors are complementary. Replacing KP-ALB with uniform load balancing [43, 10, 60] mainly hurts spatial alignment and fidelity, while removing SRC causes the largest drop in temporal stability; removing CP also degrades all metrics, indicating that predictable routing is important for stable control. Overall, the best results require kinematic-prior-guided regularization throughout the routing process. More qualitative comparisons on the effect of different losses are provided in Fig. 9 of the appendix.

**Efficient generation.** Tab. 8 separates generic few-step acceleration from action-aware efficiency. Under the same 4-step budget, DMD2 [56] reduces latency to 0.52 s but degrades control, increasing CD from 86.89 to 98.45. By contrast, our action-aware student achieves the same latency and FLOPs while preserving much stronger control fidelity, with CD 87.50. Spatial adaptive execution further reduces latency/FLOPs to 0.42 s/0.63 PF by skipping low-significance control updates, and temporal caching reaches 0.37 s/0.55 PF with moderate degradation. These results indicate that few-step distillation provides the main acceleration, while action-aware distillation and routing-derived execution better preserve control under reduced computation. A more detailed breakdown analysis is provided in Sec. F.4.

Table 7: Kinematic-prior ablation.

| Method          | CD ↓         | TI ↑        | AF ↓        | FID ↓        |
|-----------------|--------------|-------------|-------------|--------------|
| Full w/o KP-ALB | 94.30        | 0.85        | 5.20        | 21.65        |
| Full w/o SRC    | 90.12        | 0.81        | 8.75        | 17.50        |
| Full w/o CP     | 89.45        | 0.87        | 5.12        | 16.20        |
| Full model      | <b>86.89</b> | <b>0.89</b> | <b>4.48</b> | <b>14.37</b> |

Table 8: Efficient generation ablation.

| Method                               | CD ↓         | FID ↓        | Latency ↓   | FLOPs ↓     |
|--------------------------------------|--------------|--------------|-------------|-------------|
| Teacher (full generator, 50 steps)   | <b>86.89</b> | <b>14.37</b> | 4.64        | 6.94        |
| Generic distillation (DMD2)          | 98.45        | 19.82        | 0.52        | 0.78        |
| Action-aware student (few-step only) | 87.50        | 15.60        | 0.52        | 0.78        |
| + spatial adaptive execution         | 88.65        | 16.55        | 0.42        | 0.63        |
| + temporal cache (Final efficient)   | 89.87        | 17.67        | <b>0.37</b> | <b>0.55</b> |

## 5 Conclusion

This work presents KVLR, a structured framework for action-conditioned surgical video generation from articulated inputs. KVLR lifts kinematics into a pixel-aligned KVA-Field and routes control across modalities and motion scales, improving measured action alignment and generation quality under complex surgical dynamics. We also introduce KASA, a benchmark with articulated annotations from real robotic surgical videos. Experiments show that KVLR outperforms evaluated baselines, while KVLR-fast reduces inference cost with largely preserved action control. Nevertheless, generated videos should be viewed as research aids for simulation and data augmentation rather than clinical evidence. Additional discussions on limitations and broader impact are in Sec. C.

## References

- [1] Hassan Abu Alhaija, Jose Alvarez, Maciej Bala, Tiffany Cai, Tianshi Cao, Liz Cha, Joshua Chen, Mike Chen, Francesco Ferroni, Sanja Fidler, et al. Cosmos-transfer1: Conditional world generation with adaptive multimodal control. *arXiv preprint arXiv:2503.14492*, 2025.
- [2] Nicolás Ayobi, Alejandra Pérez-Rondón, Santiago Rodríguez, and Pablo Arbeláez. Matis: Masked-attention transformers for surgical instrument segmentation. In *2023 IEEE 20th International Symposium on Biomedical Imaging (ISBI)*, pages 1–5, 2023. doi: 10.1109/ISBI53787.2023.10230819.
- [3] Diego Biagini, Nassir Navab, and Azade Farshad. Hierasurg: Hierarchy-aware diffusion model for surgical video generation. In *International Conference on Medical Image Computing and Computer-Assisted Intervention*, pages 310–319. Springer, 2025.
- [4] Liang-Chieh Chen, Yukun Zhu, George Papandreou, Florian Schroff, and Hartwig Adam. Encoder-decoder with atrous separable convolution for semantic image segmentation. In *ECCV*, 2018.
- [5] Rui Chen, Zehuan Wu, Yichen Liu, Yuxin Guo, Jingcheng Ni, Haifeng Xia, and Siyu Xia. Unimlv: Unified framework for multi-view long video generation with comprehensive control capabilities for autonomous driving, 2025. URL <https://arxiv.org/abs/2412.04842>.
- [6] Tong Chen, Shuya Yang, Junyi Wang, Long Bai, Hongliang Ren, and Luping Zhou. Surgsora: Object-aware diffusion model for controllable surgical video generation. In *International Conference on Medical Image Computing and Computer-Assisted Intervention*, pages 521–531. Springer, 2025.
- [7] Bowen Cheng, Ishan Misra, Alexander G. Schwing, Alexander Kirillov, and Rohit Girdhar. Masked-attention mask transformer for universal image segmentation. In *CVPR*, 2022.
- [8] Joseph Cho, Samuel Schmidgall, Cyril Zakka, Mrudang Mathur, Dhamanpreet Kaur, Rohan Shad, and William Hiesinger. Surgen: Text-guided diffusion model for surgical video generation. *arXiv preprint arXiv:2408.14028*, 2024.
- [9] Haoqi Fan, Yanghao Li, Bo Xiong, Wan-Yen Lo, and Christoph Feichtenhofer. Pyslowfast. <https://github.com/facebookresearch/slowfast>, 2020.
- [10] Zhengcong Fei, Mingyuan Fan, Changqian Yu, Debang Li, and Junshi Huang. Scaling diffusion transformers to 16 billion parameters. *arXiv preprint arXiv:2407.11633*, 2024.
- [11] Mingju Gao, Kaisen Yang, Huan ang Gao, Bohan Li, Ao Ding, Wenyi Li, Yangcheng Yu, Jinkun Liu, Shaocong Xu, Yike Niu, Haohan Chi, Hao Chen, Hao Tang, Yu Zhang, Li Yi, and Hao Zhao. Pam: A pose-appearance-motion engine for sim-to-real hoi video generation, 2026. URL <https://arxiv.org/abs/2603.22193>.
- [12] Ruiyuan Gao, Kai Chen, Enze Xie, Lanqing Hong, Zhenguo Li, Dit-Yan Yeung, and Qiang Xu. Magicdrive: Street view generation with diverse 3d geometry control. In *ICLR*, 2024.
- [13] Shenyuan Gao, Jiazhi Yang, Li Chen, Kashyap Chitta, Yihang Qiu, Andreas Geiger, Jun Zhang, and Hongyang Li. Vista: A generalizable driving world model with high fidelity and versatile controllability. *Advances in Neural Information Processing Systems*, 37:91560–91596, 2025.
- [14] Junhao Ge, Zuhong Liu, Longteng Fan, Yifan Jiang, Jiaqi Su, Yiming Li, Zhejun Zhang, and Siheng Chen. Unraveling the effects of synthetic data on end-to-end autonomous driving. *arXiv preprint arXiv:2503.18108*, 2025.
- [15] Jiazhe Guo, Yikang Ding, Xiwu Chen, Shuo Chen, Bohan Li, Yingshuang Zou, Xiaoyang Lyu, Feiyang Tan, Xiaojuan Qi, Zhiheng Li, et al. Dist-4d: Disentangled spatiotemporal diffusion with metric depth for 4d driving scene generation. *arXiv preprint arXiv:2503.15208*, 2025.
- [16] Yufan He, Pengfei Guo, Mengya Xu, Zhaoshuo Li, Andriy Myronenko, Dillan Imans, Bingjie Liu, Dongren Yang, Mingxue Gu, Yongnan Ji, et al. Surgworld: Learning surgical robot policies from videos via world modeling. *arXiv preprint arXiv:2512.23162*, 2025.

- [17] Yufan He, Pengfei Guo, Mengya Xu, Zhaoshuo Li, Andriy Myronenko, Dillan Imans, Bingjie Liu, Dongren Yang, Mingxue Gu, Yongnan Ji, Yueming Jin, Ren Zhao, Baiyong Shen, and Daguang Xu. Cosmos-h-surgical: Learning surgical robot policies from videos via world modeling, 2026. URL <https://arxiv.org/abs/2512.23162>.
- [18] Tianshuai Hu, Xiaolu Liu, Song Wang, Yiyao Zhu, Ao Liang, Lingdong Kong, Guoyang Zhao, Zeying Gong, Jun Cen, Zhiyu Huang, et al. Vision-language-action models for autonomous driving: Past, present, and future. *arXiv preprint arXiv:2512.16760*, 2025.
- [19] Haian Jin, Hanwen Jiang, Hao Tan, Kai Zhang, Sai Bi, Tianyuan Zhang, Fujun Luan, Noah Snively, and Zexiang Xu. Lvsm: A large view synthesis model with minimal 3d inductive bias. *arXiv preprint arXiv:2410.17242*, 2024.
- [20] Lingdong Kong, Wesley Yang, Jianbiao Mei, Youquan Liu, Ao Liang, Dekai Zhu, Dongyue Lu, Wei Yin, Xiaotao Hu, Mingkai Jia, et al. 3d and 4d world modeling: A survey. *arXiv preprint arXiv:2509.07996*, 2025.
- [21] Weijie Kong, Qi Tian, Zijian Zhang, Rox Min, Zuozhuo Dai, Jin Zhou, Jiangfeng Xiong, Xin Li, Bo Wu, Jianwei Zhang, et al. Hunyuanvideo: A systematic framework for large video generative models. *arXiv preprint arXiv:2412.03603*, 2024.
- [22] Bohan Li, Jiazhe Guo, Hongsi Liu, Yingshuang Zou, Yikang Ding, Xiwu Chen, Hu Zhu, Feiyang Tan, Chi Zhang, Tiancai Wang, et al. Uniscene: Unified occupancy-centric driving scene generation. In *Proceedings of the Computer Vision and Pattern Recognition Conference*, pages 11971–11981, 2025.
- [23] Chenxin Li, Hengyu Liu, Yifan Liu, Brandon Y Feng, Wuyang Li, Xinyu Liu, Zhen Chen, Jing Shao, and Yixuan Yuan. Endora: Video generation models as endoscopy simulators. In *International conference on medical image computing and computer-assisted intervention*, pages 230–240. Springer, 2024.
- [24] Samuel Li, Pujith Kachana, Prajwal Chidananda, Saurabh Nair, Yasutaka Furukawa, and Matthew Brown. Rig3r: Rig-aware conditioning for learned 3d reconstruction. *arXiv preprint arXiv:2506.02265*, 2025.
- [25] Wei Li, Ming Hu, Guoan Wang, Lihao Liu, Kaijing Zhou, Junzhi Ning, Xin Guo, Zongyuan Ge, Lixu Gu, and Junjun He. Ophora: a large-scale data-driven text-guided ophthalmic surgical video generation model. In *International Conference on Medical Image Computing and Computer-Assisted Intervention*, pages 425–435. Springer, 2025.
- [26] Ao Liang, Lingdong Kong, Tianyi Yan, Hongsi Liu, Wesley Yang, Ziqi Huang, Wei Yin, Jialong Zuo, Yixuan Hu, Dekai Zhu, et al. Worldlens: Full-spectrum evaluations of driving world models in real world. *arXiv preprint arXiv:2512.10958*, 2025.
- [27] Ruofan Liang, Zan Gojcic, Huan Ling, Jacob Munkberg, Jon Hasselgren, Zhi-Hao Lin, Jun Gao, Alexander Keller, Nandita Vijaykumar, Sanja Fidler, et al. Diffusionrenderer: Neural inverse and forward rendering with video diffusion models. *arXiv preprint arXiv:2501.18590*, 2025.
- [28] Chih-Hao Lin, Zian Wang, Ruofan Liang, Yuxuan Zhang, Sanja Fidler, Shenlong Wang, and Zan Gojcic. Controllable weather synthesis and removal with video diffusion models. *arXiv preprint arXiv:2505.00704*, 2025.
- [29] Han Lin, Jaemin Cho, Abhay Zala, and Mohit Bansal. Ctrl-adapter: An efficient and versatile framework for adapting diverse controls to any diffusion model, 2024.
- [30] Yaron Lipman, Ricky TQ Chen, Heli Ben-Hamu, Maximilian Nickel, and Matt Le. Flow matching for generative modeling. *arXiv preprint arXiv:2210.02747*, 2022.
- [31] Liu Liu, Xiaofeng Wang, Guosheng Zhao, Keyu Li, Wenkang Qin, Jiaxiong Qiu, Zheng Zhu, Guan Huang, and Zhizhong Su. Robotransfer: Geometry-consistent video diffusion for robotic visual policy transfer. *arXiv preprint arXiv:2505.23171*, 2025.
- [32] Ze Liu, Jia Ning, Yue Cao, Yixuan Wei, Zheng Zhang, Stephen Lin, and Han Hu. Video swin transformer. *arXiv preprint arXiv:2106.13230*, 2021.

- [33] Ilya Loshchilov and Frank Hutter. Decoupled weight decay regularization. *arXiv preprint arXiv:1711.05101*, 2017.
- [34] Ning Ma, Shu Yang, Yizhao Zhou, Chaoyang Zhang, Jian Chen, and Xiaoman He. Open-world surgical video generation via dual-visual diffusion and dual-annealed generation. *Neural Networks*, page 108281, 2025.
- [35] Sabina Martyniak, Joanna Kaleta, Diego Dall’Alba, Michal Naskrket, Szymon Plotka, and Przemysław Korzeniowski. Simuscope: Realistic endoscopic synthetic dataset generation through surgical simulation and diffusion models. *2025 IEEE/CVF Winter Conference on Applications of Computer Vision (WACV)*, pages 4268–4278, 2024. URL <https://api.semanticscholar.org/CorpusID:274445994>.
- [36] Saeed Masoudnia and Reza Ebrahimpour. Mixture of experts: a literature survey. *Artificial Intelligence Review*, 42(2):275–293, 2014.
- [37] Open-H-Embodiment Consortium Nigel Nelson, Juo-Tung Chen, Jesse Haworth, Xinhao Chen, Lukas Zbinden, and etc. Dianye Huang. Open-h-embodiment: A large-scale dataset for enabling foundation models in medical robotics, 2026. URL <https://api.semanticscholar.org/CorpusID:287702178>.
- [38] Dimitrios Psychogios, Emanuele Colleoni, Beatrice Van Amsterdam, Chih-Yang Li, Shu-Yu Huang, Yuchong Li, Fucang Jia, Baosheng Zou, Guotai Wang, Yang Liu, et al. Sar-rarp50: Segmentation of surgical instrumentation and action recognition on robot-assisted radical prostatectomy challenge. *arXiv preprint arXiv:2401.00496*, 2023.
- [39] Sampath Rapuri, Lalithkumar Seenivasan, Dominik Schneider, Roger Soberanis-Mukul, Yufan He, Hao Ding, Jiru Xu, Chenhao Yu, Chenyan Jing, Pengfei Guo, Daguang Xu, and Mathias Unberath. Saw: Toward a surgical action world model via controllable and scalable video generation, 2026. URL <https://arxiv.org/abs/2603.13024>.
- [40] Nikhila Ravi, Valentin Gabeur, Yuan-Ting Hu, Ronghang Hu, Chaitanya Ryali, Tengyu Ma, Haitham Khedr, Roman Rädle, Chloe Rolland, Laura Gustafson, et al. Sam 2: Segment anything in images and videos. *arXiv preprint arXiv:2408.00714*, 2024.
- [41] Tianhe Ren, Shilong Liu, Ailing Zeng, Jing Lin, Kunchang Li, He Cao, Jiayu Chen, Xinyu Huang, Yukang Chen, Feng Yan, Zhaoyang Zeng, Hao Zhang, Feng Li, Jie Yang, Hongyang Li, Qing Jiang, and Lei Zhang. Grounded sam: Assembling open-world models for diverse visual tasks, 2024.
- [42] Olaf Ronneberger, Philipp Fischer, and Thomas Brox. U-net: Convolutional networks for biomedical image segmentation. In *MICCAI*, 2015.
- [43] Usha Ruby, Vamsidhar Yendapalli, et al. Binary cross entropy with deep learning technique for image classification. *Int. J. Adv. Trends Comput. Sci. Eng.*, 9(10), 2020.
- [44] Saumya Saksena. R3d-18 for ucf-101 action recognition. <https://huggingface.co/dronefreak/r3d-18-ucf101>, 2024.
- [45] Mehmet Kerem Turkcan, Mattia Ballo, Filippo Filicori, and Zoran Kostic. Towards suturing world models: Learning predictive models for robotic surgical tasks. *arXiv preprint arXiv:2503.12531*, 2025.
- [46] Natalia Valderrama, Paola Ruiz, Isabela Hernández, Nicolás Ayobi, Mathilde Verlyck, Jessica Santander, Juan Caicedo, Nicolás Fernández, and Pablo Arbeláez. Towards holistic surgical scene understanding. In *Medical Image Computing and Computer Assisted Intervention – MICCAI 2022*, pages 442–452, Cham, 2022. Springer Nature Switzerland.
- [47] 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

- 484 Wang, Wenmeng Zhou, Wente Wang, Wenting Shen, Wenyuan Yu, Xianzhong Shi, Xiaoming  
485 Huang, Xin Xu, Yan Kou, Yangyu Lv, Yifei Li, Yijing Liu, Yiming Wang, Yingya Zhang,  
486 Yitong Huang, Yong Li, You Wu, Yu Liu, Yulin Pan, Yun Zheng, Yuntao Hong, Yupeng Shi,  
487 Yutong Feng, Zeyinzi Jiang, Zhen Han, Zhi-Fan Wu, and Ziyu Liu. Wan: Open and advanced  
488 large-scale video generative models. *arXiv preprint arXiv:2503.20314*, 2025.
- 489 [48] Rongsheng Wang, Junying Chen, Ke Ji, Zhenyang Cai, Shunian Chen, Yunjin Yang, and Benyou  
490 Wang. Medgen: Unlocking medical video generation by scaling granularly-annotated medical  
491 videos. *arXiv preprint arXiv:2507.05675*, 2025.
- 492 [49] Enze Xie, Wenhai Wang, Zhiding Yu, Anima Anandkumar, Jose M Alvarez, and Ping Luo.  
493 Segformer: Simple and efficient design for semantic segmentation with transformers. *Advances*  
494 *in neural information processing systems*, 34:12077–12090, 2021.
- 495 [50] Jiawei Xu, Kai Deng, Zexin Fan, Shenlong Wang, Jin Xie, and Jian Yang. Ad-gs: Object-  
496 aware b-spline gaussian splatting for self-supervised autonomous driving. *arXiv preprint*  
497 *arXiv:2507.12137*, 2025.
- 498 [51] Jiazhi Yang, Kashyap Chitta, Shenyuan Gao, Long Chen, Yuqian Shao, Xiaosong Jia, Hongyang  
499 Li, Andreas Geiger, Xiangyu Yue, and Li Chen. Resim: Reliable world simulation for au-  
500 tonomous driving. *arXiv preprint arXiv:2506.09981*, 2025.
- 501 [52] Xiuyu Yang, Bohan Li, Shaocong Xu, Nan Wang, Chongjie Ye, Zhaoxi Chen, Minghan Qin,  
502 Yikang Ding, Zheng Zhu, Xin Jin, et al. Orv: 4d occupancy-centric robot video generation.  
503 *arXiv preprint arXiv:2506.03079*, 2025.
- 504 [53] Zeyu Yang, Zijie Pan, Yuankun Yang, Xiatian Zhu, and Li Zhang. Driving view synthesis on free-  
505 form trajectories with generative prior, 2025. URL <https://arxiv.org/abs/2412.01717>.
- 506 [54] Zhuoran Yang, Xi Guo, Chenjing Ding, Chiyu Wang, Wei Wu, and Yanyong Zhang. Instadrive:  
507 Instance-aware driving world models for realistic and consistent video generation. In *ICCV*,  
508 2025.
- 509 [55] Zhuoyi Yang, Jiayan Teng, Wendi Zheng, Ming Ding, Shiyu Huang, Jiazheng Xu, Yuanming  
510 Yang, Wenyi Hong, Xiaohan Zhang, Guanyu Feng, et al. Cogvideox: Text-to-video diffusion  
511 models with an expert transformer. *arXiv preprint arXiv:2408.06072*, 2024.
- 512 [56] Tianwei Yin, Michaël Gharbi, Taesung Park, Richard Zhang, Eli Shechtman, Fredo Durand,  
513 and William T Freeman. Improved distribution matching distillation for fast image synthesis.  
514 In *NeurIPS*, 2024.
- 515 [57] Lvmin Zhang, Anyi Rao, and Maneesh Agrawala. Adding conditional control to text-to-image  
516 diffusion models. In *ICCV*, 2023.
- 517 [58] Jay Zhangjie Wu, Yuxuan Zhang, Haithem Turki, Xuanchi Ren, Jun Gao, Mike Zheng Shou,  
518 Sanja Fidler, Zan Gojcic, and Huan Ling. Difx3d+: Improving 3d reconstructions with  
519 single-step diffusion models. *arXiv e-prints*, pages arXiv–2503, 2025.
- 520 [59] Shihao Zhao, Dongdong Chen, Yen-Chun Chen, Jianmin Bao, Shaozhe Hao, Lu Yuan, and  
521 Kwan-Yee K. Wong. Uni-controlnet: All-in-one control to text-to-image diffusion models.  
522 *Advances in Neural Information Processing Systems*, 2023.
- 523 [60] Youwei Zheng, Yuxi Ren, Xin Xia, Xuefeng Xiao, and Xiaohua Xie. Dense2moe: Restructuring  
524 diffusion transformer to moe for efficient text-to-image generation. In *Proceedings of the*  
525 *IEEE/CVF International Conference on Computer Vision*, pages 18661–18670, 2025.
- 526 [61] Zangwei Zheng, Xiangyu Peng, Tianji Yang, Chenhui Shen, Shenggui Li, Hongxin Liu, Yukun  
527 Zhou, Tianyi Li, and Yang You. Open-sora: Democratizing efficient video production for all.  
528 *arXiv preprint arXiv:2412.20404*, 2024.
- 529 [62] Yanqi Zhou, Tao Lei, Hanxiao Liu, Nan Du, Yanping Huang, Vincent Zhao, Andrew M Dai,  
530 Quoc V Le, James Laudon, et al. Mixture-of-experts with expert choice routing. *Advances in*  
531 *Neural Information Processing Systems*, 35:7103–7114, 2022.
- 532 [63] Fangqi Zhu, Hongtao Wu, Song Guo, Yuxiao Liu, Chilam Cheang, and Tao Kong. Irasim:  
533 Learning interactive real-robot action simulators. *arXiv preprint arXiv:2406.14540*, 2024.

# Appendix

534

|     |                                                                           |           |
|-----|---------------------------------------------------------------------------|-----------|
| 535 | <b>A Formal Motivation</b>                                                | <b>15</b> |
| 536 | A.1 Why Pixel-Aligned Lifting Helps . . . . .                             | 15        |
| 537 | A.2 When Hierarchical Routing Is Useful . . . . .                         | 15        |
| 538 | A.3 Interpretation of Kinematic-Prior Regularization . . . . .            | 16        |
| 539 | <b>B Related Work</b>                                                     | <b>17</b> |
| 540 | B.1 Surgical Video Generation . . . . .                                   | 17        |
| 541 | B.2 Controllable Generation with Structured Signals . . . . .             | 17        |
| 542 | <b>C Discussion on Limitations and Broader Impact</b>                     | <b>17</b> |
| 543 | C.1 Limitations . . . . .                                                 | 17        |
| 544 | C.2 Broader Impact . . . . .                                              | 18        |
| 545 | C.3 Asset License and Usage Note . . . . .                                | 18        |
| 546 | <b>D More Implementation Details</b>                                      | <b>18</b> |
| 547 | D.1 Detailed Experimental Setup . . . . .                                 | 18        |
| 548 | D.2 Backbone Instantiation and Training Context . . . . .                 | 19        |
| 549 | D.3 Construction of the Kinematic-to-Visual Action Field . . . . .        | 20        |
| 550 | D.4 Detailed Hierarchical Routing . . . . .                               | 21        |
| 551 | D.5 Detailed Kinematic-Prior Learning . . . . .                           | 21        |
| 552 | D.6 Detailed Action-adaptive Efficient Generation . . . . .               | 23        |
| 553 | <b>E Supplementary Details on KASA Benchmark Construction</b>             | <b>25</b> |
| 554 | E.1 Design Objective and Benchmark Scope . . . . .                        | 25        |
| 555 | E.2 Task-Driven Sequence Selection . . . . .                              | 26        |
| 556 | E.3 Human-in-the-loop Semantic Annotation . . . . .                       | 26        |
| 557 | E.4 Differentiable Pose Tracking . . . . .                                | 27        |
| 558 | E.5 Quality Validation of KASA Dataset . . . . .                          | 27        |
| 559 | E.6 From Curated Supervision to Structured Action Control . . . . .       | 28        |
| 560 | E.7 Discussion . . . . .                                                  | 28        |
| 561 | <b>F Additional Experiments</b>                                           | <b>29</b> |
| 562 | F.1 Benchmark Usability Validation of KASA . . . . .                      | 29        |
| 563 | F.2 Evaluation Independence and Bias Control . . . . .                    | 29        |
| 564 | F.3 Quantitative Analysis of Hierarchical Expert Specialization . . . . . | 30        |
| 565 | F.4 Detailed Efficient-Generation Ablation Breakdown . . . . .            | 31        |
| 566 | F.5 More Qualitative Results. . . . .                                     | 32        |

## A Formal Motivation

This section of the appendix provides a more discussions on the main design choices in KVLR.

### A.1 Why Pixel-Aligned Lifting Helps

The first design choice in KVLR is to lift sparse articulated actions into a dense image-aligned control field before injecting them into the video generator. This subsection formalizes the intuition that raw-action conditioning must learn both spatial alignment and control-feature prediction, while the lifted field externalizes part of the alignment problem.

Let  $u$  denote a spatial-temporal token and let  $z_t^*(u)$  be an ideal local control feature sufficient for action-conditioned generation at token  $u$ :

$$z_t^*(u) = \phi^*(\xi_t(u)), \quad (10)$$

where  $\xi_t(u)$  denotes the local physical state relevant to token  $u$ , such as part identity, depth, local orientation, projected motion, and motion change.

Under raw-action conditioning, the model predicts a token-wise control feature directly from the global articulated action:

$$\hat{z}_t^{\text{raw}}(u) = g_{\text{raw}}(a_t, u). \quad (11)$$

This requires the learned branch to implicitly approximate both the projection from the global articulated state to token  $u$  and the mapping from the resulting local physical state to generator features.

In contrast, KVLR first constructs a pixel-aligned action field,

$$K_t(u) = L_u(a_t), \quad (12)$$

where  $L_u(\cdot)$  denotes the deterministic lifting operation at token  $u$ . The local control feature is then predicted as:

$$\hat{z}_t^{\text{lift}}(u) = g_{\text{lift}}(K_t(u)). \quad (13)$$

The lifted representation separates the spatial-alignment step from the learned control-feature mapping.

To make this intuition explicit, suppose that the lifted field recovers the relevant local physical state up to an abstract error  $\varepsilon_{\text{lift}}$ , while a learned projection from raw actions recovers it up to  $\varepsilon_{\text{proj}}$ . Also suppose that  $\phi^*$  is  $L_\phi$ -Lipschitz and that the learned mappings have approximation errors  $\varepsilon_{\text{app}}^{\text{lift}}$  and  $\varepsilon_{\text{app}}^{\text{raw}}$ . A direct triangle-inequality decomposition gives:

$$\|\hat{z}_t^{\text{lift}}(u) - z_t^*(u)\| \lesssim L_\phi \varepsilon_{\text{lift}} + \varepsilon_{\text{app}}^{\text{lift}}, \quad (14)$$

whereas raw-action conditioning yields:

$$\|\hat{z}_t^{\text{raw}}(u) - z_t^*(u)\| \lesssim L_\phi \varepsilon_{\text{proj}} + \varepsilon_{\text{app}}^{\text{raw}}. \quad (15)$$

These quantities are the decomposition that clarifies the design assumption behind KVLR: pixel-aligned lifting is expected to help when the hand-constructed action field reduces the spatial-alignment burden more than it restricts the representation. This is consistent with the parameter-matched ablation in the main paper, where replacing raw action vectors with the KVA-Field provides a substantial gain before adding hierarchical routing.

### A.2 When Hierarchical Routing Is Useful

The second design choice in KVLR is to route lifted action features across physical modalities and motion scales. We interpret this routing mechanism as an inductive bias rather than as a universally stronger function class. A sufficiently large monolithic conditioner can, in principle, approximate heterogeneous control mappings. The motivation for routing is that, under a fixed parameter and computation budget, surgical scenes contain spatial-temporal regimes with different dominant factors: static or low-motion regions often require semantic and geometric consistency, while high-motion regions require dynamic cues and larger motion operators.

Let  $\kappa_t(u)$  denote an unobserved control regime for token  $u$ , and let  $f_{\kappa_t(u)}^*$  be the corresponding local control map. A monolithic conditioner approximates all regimes using a shared operator:

$$\hat{z}_t^{\text{mono}}(u) = f_{\text{mono}}(K_t(u)). \quad (16)$$

A routed conditioner instead selects an expert or sub-expert according to the local control regime:

$$\hat{z}_t^{\text{route}}(u) = f_{\hat{\kappa}_t(u)}(K_t(u)), \quad (17)$$

where  $\hat{\kappa}_t(u)$  denotes the selected routing branch.

If the average within-regime approximation error is denoted by  $\bar{\varepsilon}$ , and if the probability of selecting an unsuitable branch is denoted by  $\eta$ , then the routed error can be informally decomposed as

$$\mathbb{E}_{t,u} [\|\hat{z}_t^{\text{route}}(u) - z_t^*(u)\|] \lesssim \bar{\varepsilon} + C\eta, \quad (18)$$

where  $C$  is a bounded mismatch constant measuring the cost of choosing an unsuitable expert. This expression separates two intuitive error sources: expert approximation error and routing error. It is not meant to provide an estimable guarantee, because  $\bar{\varepsilon}$ ,  $\eta$ , and  $C$  are not measured in the actual model.

This decomposition motivates two aspects of KVLR. First, hierarchical routing can be useful when different physical regimes are easier to approximate with specialized operators than with a single shared pathway. Second, the benefit of specialization depends on routing quality, which motivates the kinematic-prior objectives introduced in the main method. In practice, the value of this inductive bias is evaluated empirically through the parameter-matched ablation comparing raw-action injection, KVA-Field injection, Tier-1 routing, and Tier-1 plus Tier-2 routing.

### A.3 Interpretation of Kinematic-Prior Regularization

The kinematic-prior objectives are introduced to bias the router toward physically plausible and temporally stable assignments. They constrain observable aspects of routing behavior that are relevant to surgical motion.

**Kinematic-prior adaptive load balancing.** The KP-ALB objective aligns modality usage with the physical signal strength in the KVA-Field:

$$\mathcal{L}_{\text{KP-ALB}} = \sum_{i \in \mathcal{M}} (\ell_i - \pi_i(K_t))^2, \quad (19)$$

where  $\pi_i(K_t)$  denotes the normalized physical prior of modality  $i$ , and  $\ell_i$  denotes its routing load. Unlike uniform load balancing, this term does not force all experts to be used equally. Instead, it discourages routing patterns that ignore the relative strength of semantic, geometric, and dynamic action cues.

**Spatiotemporal routing consistency.** The SRC objective penalizes frame-to-frame routing variation around tool regions:

$$\mathcal{L}_{\text{SRC}} = \frac{1}{T-1} \sum_{t=1}^{T-1} \sum_{h,w} M_{\text{tool}}^{(t)}(h,w) \|R_t(h,w) - R_{t-1}(h,w)\|_2^2. \quad (20)$$

This term reflects the assumption that routing decisions around articulated tools should change smoothly unless the underlying motion changes. It reduces routing flicker but does not enforce temporal smoothness in static background regions.

**Capacity prediction.** The capacity-prediction objective trains a lightweight predictor to approximate routing decisions from detached action features:

$$\mathcal{L}_{\text{CP}} = \text{BCE}(\sigma(\hat{R}), \text{sg}(R^*)). \quad (21)$$

This makes routing decisions more predictable from action-derived signals, which later supports the efficient execution policy. Its role is practical rather than theoretical: it helps expose a stable approximation of where full control computation is likely to matter.

Overall, these regularizers reduce the degrees of freedom of the router and make the learned routing easier to interpret.

## B Related Work

### B.1 Surgical Video Generation

Surgical video generation has recently attracted growing attention for simulation, training, and predictive modeling in robotic surgery [48, 25, 8, 34, 45, 39]. Early progress mainly focuses on scaling surgical video corpora and improving visual realism. MedGen [48] and Ophora [25] build large-scale captioned datasets for generalist and ophthalmic video synthesis, while Endora [23] studies endoscopic video generation as a simulator for minimally invasive scenes. Simulation-assisted efforts such as SimuScope [35] combine surgical simulation with diffusion-based image translation to improve synthetic endoscopic realism. Beyond visual synthesis, recent works explore controllable surgical video generation. SurGen [8] and Open-world [34] rely on language guidance for flexible scene generation, while SurgSora [6] introduces structured RGBD-flow decomposition for controllable surgical video synthesis. HieraSurg [3] further exploits hierarchy-aware segmentation to improve surgical video generation. More recent surgical world-modeling efforts move toward action-conditioned synthesis: SAW [39] conditions video diffusion on lightweight signals such as language, reference scenes, tissue affordance masks, and 2D tool-tip trajectories, while Surg-World [16] uses generative world modeling and pseudo-kinematics to support surgical policy learning. Open-H-Embodiment [37] also highlights the importance of synchronized medical robotic videos and kinematics for embodied surgical AI. Despite these advances, existing methods still primarily rely on coarse text prompts, 2D trajectory cues, or dense visual priors such as masks, depth, or optical flow. These controls are either insufficient for fine-grained articulated manipulation or costly to obtain at inference time. In contrast, our work targets *action-conditioned* surgical video generation from sparse articulated controls. By lifting kinematics into a pixel-aligned action field and hierarchically routing conditioning across physical modalities and motion scales, KVLR aims to improve control precision, visual fidelity, and efficient inference under accessible robotic action inputs.

### B.2 Controllable Generation with Structured Signals

Beyond surgery, controllable video generation has been actively studied in robotics, autonomous driving, and general scene modeling [61, 55, 11, 52, 19, 24, 26, 51, 22, 5, 54, 50, 14, 53, 15, 20, 18, 26]. In robotics, IRASim [63] studies action-to-video prediction for robotic manipulation, while Cosmos-Transfer [1] and RoboTransfer [31] condition synthesis on scene-level geometric cues such as depth or surface normals. In driving and embodied scene generation, UniScene [22] uses hierarchical occupancy priors, and MagicDrive [12] and Vista [13] exploit structured multi-view or latent control for controllable synthesis. Related efforts in controllable rendering and 3D-aware generation also rely on dense scene representations or explicit geometric guidance [28, 58, 27]. These approaches demonstrate the value of structured conditioning, but they typically depend on dense maps, occupancy fields, or other heavy intermediate representations that are difficult to access in robotic surgery, where sensing is limited and occlusion is frequent [52, 58, 22, 27]. Differently, our method focuses on sparse articulated kinematic actions as the primary control interface, and in transforming them into pixel-aligned visual control through hierarchical kinematic-to-visual routing rather than relying on dense geometric priors.

## C Discussion on Limitations and Broader Impact

### C.1 Limitations

Our framework is developed for action-conditioned surgical video generation under articulated robotic control, and is evaluated on representative suturing-related sub-actions with curated supervision. While the proposed structured kinematic-to-visual control improves faithfulness, fidelity, and efficiency in this setting, extending the framework to broader surgical procedures and more diverse instruments may require additional adaptation in data curation and model scaling. In addition, as with other generative models, quantitative improvements on benchmark metrics do not fully capture all aspects of clinical realism or downstream utility, and human-centered evaluation remains important for future study.

## 691 C.2 Broader Impact

692 This work has the potential to benefit surgical simulation, training, and data augmentation by  
693 enabling more controllable video generation from lightweight robotic action signals. At the same  
694 time, synthetic surgical videos should not be interpreted as a substitute for real clinical evidence or  
695 expert judgment. We do not advocate direct deployment of generated videos for autonomous clinical  
696 decision-making. More broadly, responsible use of such models requires appropriate attention to  
697 data governance, benchmarking transparency, and the intended educational or research context.

## 698 C.3 Asset License and Usage Note

699 This work builds on several existing datasets, models, and software components, including SAR-  
700 RARP for source video data, SAM2 for segmentation-assisted annotation, OpenSora and related  
701 pretrained video-generation components for model initialization, and publicly available baseline  
702 methods used for comparison. We credit the original creators of these assets in the main paper and  
703 references, and use them in accordance with their respective licenses, terms of use, and research-only  
704 restrictions where applicable. For all external assets, we follow the usage conditions specified by  
705 their official releases or project pages. Any future release of our code, models, or derived benchmark  
706 assets will also be conditioned on compatibility with the licenses and usage terms of these underlying  
707 resources. We refer to the official releases/pages of these assets for the exact license terms and usage  
708 conditions, and will include explicit license identifiers in any public release package.

## 709 D More Implementation Details

### 710 D.1 Detailed Experimental Setup

711 **Training details.** Unless otherwise specified, all models are trained on NVIDIA H20 GPUs with  
712 a batch size of 24 using AdamW [33], with a learning rate of  $1 \times 10^{-5}$  and weight decay of 0.01. We  
713 split the KASA benchmark into 85% training data and 15% validation data. Our model is trained with  
714 a progressive two-stage protocol for stable convergence. In the first stage, we finetune the generator  
715 using text-only conditioning to preserve the pretrained video prior and adapt the backbone to the  
716 surgical domain. In the second stage, we continue training with joint text-and-action conditioning,  
717 enabling the model to learn structured action control on top of the adapted generator. The overall  
718 training is implemented on 24 GPUs for approximately 7 days. Unless otherwise stated, generated  
719 videos use a spatial resolution of  $576 \times 1024$ , and the generated sequence length is aligned with the  
720 corresponding ground-truth clip. For inference, the full model uses 50 denoising steps. The efficient  
721 variant, denoted as *Ours-fast*, uses the distilled 4-step student described in Sec. 2.5.

722 **Baseline setup.** We compare with Wan-2.2 [47], HunyuanVideo [21], OpenSora-2.0 [61], and  
723 Cosmos-H-Surgical [17]. All evaluated baselines are finetuned on KASA under the same train and  
724 validation split and the same input protocol for fair comparison. For text-conditioned baselines, we  
725 use the original conditioning interface of each model. For action-conditioned comparison, models  
726 marked with ‘\*’ are augmented with ControlNet-style direct action injection, so that articulated  
727 actions are provided as explicit additional controls. Models with ‘†’ leverage dense visual conditions  
728 as SurgSora [6], including depth, semantic, and flow maps. During inference, each baseline uses  
729 its default sampling setting when applicable. Our method uses the default settings described above  
730 unless otherwise specified in the corresponding experiment.

731 **Evaluation metrics.** We evaluate three aspects: *action faithfulness*, *generation fidelity*, and *computa-*  
732 *tional efficiency*. For action faithfulness, each generated video is matched to its corresponding KASA  
733 annotation by video ID. Let  $\Omega \subset \mathbb{Z}^2$  denote the image pixel domain. For an instrument instance  $m$   
734 at frame  $t$ , we render the articulated annotation into a ground-truth action tube mask  $T_t^m \subset \Omega$  by  
735 projecting the recovered 3D instrument skeleton into the image plane and drawing a fixed-width tube  
736 around the projected skeleton. We also render a thin centerline mask  $C_t^m$  and projected articulated  
737 keypoints for auxiliary validation. To extract the generated instrument region, SAM2 or Grounded  
738 SAM is initialized with the first valid ground-truth action tube mask and propagated through the  
739 generated video, producing a predicted binary mask  $P_t^m \subset \Omega$  for each instrument instance and frame.

740 The reported Chamfer Distance (CD) measures the symmetric pixel distance between the generated  
 741 instrument mask and the target action tube:

$$\text{CD}_t^m = \frac{1}{2} \left( \frac{1}{|P_t^m|} \sum_{p \in P_t^m} d(p, T_t^m) + \frac{1}{|T_t^m|} \sum_{g \in T_t^m} d(g, P_t^m) \right), \quad (22)$$

742 where  $|\cdot|$  denotes the number of foreground pixels and  $d(x, A) = \min_{a \in A} \|x - a\|_2$  is the Euclidean  
 743 distance from pixel  $x$  to the nearest foreground pixel in mask  $A$ , computed using distance transforms.  
 744 Lower CD indicates that the generated instrument region is spatially closer to the projected target  
 745 action.

746 Temporal IoU (TI) evaluates frame-to-frame consistency of the generated tool region. Let  $\bar{P}_t =$   
 747  $\bigcup_m P_t^m$  be the union of predicted masks over all instrument instances at frame  $t$ . We compute:

$$\text{TI}_t = \frac{|\bar{P}_t \cap \bar{P}_{t-1}|}{|\bar{P}_t \cup \bar{P}_{t-1}|}. \quad (23)$$

748 Higher TI indicates more temporally stable generated tool masks. Area Flicker (AF) measures relative  
 749 frame-to-frame area variation:

$$\text{AF}_t = \frac{||\bar{P}_t| - |\bar{P}_{t-1}||}{\max(|\bar{P}_{t-1}|, 1)}. \quad (24)$$

750 Lower AF indicates less abrupt temporal fluctuation in the generated action region. All action-  
 751 faithfulness metrics are averaged over valid frames and instrument instances with non-empty projected  
 752 action targets.

753 For generation fidelity, we report PSNR, SSIM, FID, and FVD. PSNR and SSIM are averaged over all  
 754 generated frames. FID is computed by treating generated video frames as images. FVD is computed  
 755 on 17-frame clips to evaluate temporal realism and video-level consistency.

756 For efficiency, we report inference latency and FLOPs, and additionally report peak memory in  
 757 the efficient-generation ablations. Unless otherwise stated, all validation metrics are averaged over  
 758 sequences. The computational efficiency results in the main paper are measured on NVIDIA H20  
 759 GPUs using 17-frame clips.

760 Main benchmark tables report averages over the three action subsets, while uncertainty estimates are  
 761 provided for the representative quantitative analyses in Sec. F.

## 762 D.2 Backbone Instantiation and Training Context

763 Our generator  $G$  is instantiated on top of OpenSora, which consists of a pretrained text encoder, a 3D  
 764 VAE, and a trainable DiT backbone. In the main paper, we keep Sec. 2.1 intentionally compact and  
 765 mainly describe the two-stage control decomposition. Here, we provide the omitted implementation  
 766 context.

767 Given a reference image  $I^{\text{ref}}$ , a text prompt  $y$ , and an articulated action sequence  $\mathbf{a}_{1:T}$ , we first encode  
 768 the visual input into the latent space of the pretrained 3D VAE. The text condition is processed by the  
 769 pretrained text encoder, while the articulated action sequence is transformed into the Kinematic-to-  
 770 Visual Action Field described in Sec. D.3. The resulting structured action features are then processed  
 771 by the hierarchical routing module before being injected into the DiT backbone through a lightweight  
 772 conditional branch.

773 Unless otherwise specified, we preserve the original latent dimensionality and temporal tokenization  
 774 of the OpenSora backbone. The action branch is designed to be additive rather than replacing the  
 775 original visual-text conditioning pathway. This choice stabilizes optimization in the early stage  
 776 of finetuning and allows the pretrained generator to progressively absorb structured action control  
 777 without disrupting its original generative prior.

778 The full model is trained with the primary video generation objective together with the kinematic-  
 779 prior routing objectives introduced in the main paper. Articulated supervision is obtained from  
 780 our curated KASA pipeline, which provides lightweight action signals consisting of a temporally  
 781 consistent articulated pose and gripper status. Additional details on the routing objectives and efficient  
 782 generation strategy are provided in Sec. 2.3 and the later supplementary sections.

### 783 D.3 Construction of the Kinematic-to-Visual Action Field

784 **Articulated action representation.** For each frame  $t$ , we represent the articulated action as:

$$a_t = [p_t, r_t, q_t^{\text{sw}}, q_t^{\text{lg}}, q_t^{\text{rg}}] \in \mathbb{R}^9, \quad (25)$$

785 where the wrist is connected to the shaft by a shaft-to-wrist joint, and the left and right grippers  
786 are connected to the wrist by separate wrist-to-gripper joints.  $p_t \in \mathbb{R}^3$  and  $r_t \in \mathbb{R}^3$  denote the  
787 wrist translation and wrist orientation, respectively, with  $r_t$  expressed in axis-angle form. The scalar  
788 variables  $q_t^{\text{sw}}$ ,  $q_t^{\text{lg}}$ , and  $q_t^{\text{rg}}$  denote the shaft-to-wrist, wrist-to-left-gripper, and wrist-to-right-gripper  
789 joint states.

790 **Dense field layout.** We lift each sparse articulated state to a pixel-aligned tensor:

$$K_t \in \mathbb{R}^{H \times W \times 9}, \quad (26)$$

791 whose channels encode part semantics, depth, rotation, velocity, and acceleration:

$$K_t(h, w) = [s_t(h, w), d_t(h, w), \rho_t(h, w), v_t(h, w), \alpha_t(h, w)]. \quad (27)$$

792 Here,  $s_t \in \mathbb{R}^3$  denotes part-aware semantic channels,  $d_t$  is rendered depth,  $\rho_t$  is a local rotation  
793 descriptor,  $v_t \in \mathbb{R}^3$  is projected motion, and  $\alpha_t$  is acceleration magnitude. We use  $\alpha_t$  for acceleration  
794 to avoid overloading the symbol  $a_t$  for articulated action.

795 **Semantic channels.** The semantic channels are derived from the Human-in-the-loop Semantic  
796 Annotation (HSA) stage of KASA. Specifically, for each clip, we initialize part-aware annotation  
797 from a reference frame and obtain temporally propagated masks using a SAM2 model finetuned on  
798 SAR-RARP. Annotators then refine the propagated masks to correct leakage, part confusion, and  
799 ambiguous boundaries. In our default implementation, the three semantic channels correspond to the  
800 articulated instrument parts used throughout the paper: shaft, wrist, and gripper. We rasterize these  
801 refined part masks into three binary maps and stack them as:

$$s_t(h, w) = [s_t^{\text{shaft}}(h, w), s_t^{\text{wrist}}(h, w), s_t^{\text{gripper}}(h, w)]. \quad (28)$$

802 When multiple rendered parts overlap at a pixel, we resolve the assignment using the front-most  
803 visible part under the rendered depth order.

804 **Depth and rotation channels.** The geometric channels are obtained from the Differentiable Pose  
805 Tracking (DPT) stage. DPT estimates temporally consistent articulated poses directly from surgical  
806 videos by aligning rendered part-aware silhouettes from a controllable Gaussian representation,  
807 initialized from the instrument CAD prior, to the observed segmentation masks. Once the articulated  
808 pose  $M_t$  is estimated, we render a depth map in the camera view to obtain  $d_t$ .

809 For the rotation channel, we use a local descriptor  $\rho_t(h, w)$  derived from the projected articulated  
810 orientation of the visible instrument part at pixel  $(h, w)$ . In practice, this descriptor can be instantiated  
811 using the wrist/tool orientation most relevant to the visible part, *e.g.*, projected roll/yaw information  
812 or another normalized scalar orientation cue consistent across frames. In all cases, the role of  $\rho_t$  is to  
813 encode orientation-related geometry that is not captured by depth alone.

814 **Temporal dynamic channels.** We compute motion channels from consecutive articulated states:

$$v_t = \frac{\Pi(M_t) - \Pi(M_{t-1})}{\Delta t}, \quad \alpha_t = \frac{\|v_t - v_{t-1}\|_2}{\Delta t}, \quad (29)$$

815 where  $\Pi(\cdot)$  denotes camera projection and  $M_t$  is the articulated instrument state at frame  $t$ . Intuitively,  
816  $v_t$  measures projected motion in image space, while  $\alpha_t$  captures the magnitude of temporal motion  
817 change. This yields a field that encodes not only where the instrument is, but also how it moves.

818 **Rendering and normalization.** All channels are constructed in the image plane and resized to the  
819 spatial resolution used by the conditional branch. Binary semantic masks are kept in  $[0, 1]$ . The depth,  
820 rotation, velocity, and acceleration channels are normalized per channel using statistics computed on  
821 the training split, which stabilizes optimization across heterogeneous motion magnitudes. Unless  
822 otherwise specified, the same normalization parameters are reused during inference.

823 **Why this representation is sufficient.** The Kinematic-to-Visual Action Field serves as the minimal  
824 structured control basis used by the subsequent routing module. It preserves the accessibility of sparse  
825 articulated control, establishes explicit spatial correspondence with the image plane, and separates  
826 semantic, geometric, and dynamic factors before routing. This separation is important because the  
827 downstream hierarchical router allocates conditional computation according to physical modalities  
828 and motion scales rather than learning such disentanglement implicitly from raw control tokens.

## 829 D.4 Detailed Hierarchical Routing

830 **Tier-1 cross-modality routing.** Here we provide the detailed formulation of the outer router. We use  
 831 five modality-specific experts corresponding to semantics, depth, rotation, velocity, and acceleration.  
 832 Each expert processes only its associated channels from the Kinematic-to-Visual Action Field. The  
 833 outer gating network predicts token-wise routing weights conditioned on the action representation  
 834 and diffusion timestep. We adopt sparse top- $k$  routing during training, with a capacity schedule that  
 835 gradually transitions from dense soft routing to sparse expert selection.

836 **Tier-2 intra-modal motion-scale routing.** Within each modality branch, we further introduce three  
 837 sub-experts: a fine-manipulation branch for localized high-frequency motion, a transport-motion  
 838 branch for larger displacements, and a skip branch for static or weakly changing regions. The inner  
 839 router operates locally on modality features and selects the appropriate motion-scale operator for  
 840 each token.

841 **Sparse fusion and stable injection.** The final routed action representation is obtained by sparse  
 842 fusion over the selected Tier-1 experts using re-normalized routing weights. We inject the fused  
 843 control features into the DiT backbone through zero-initialized additive layers, which preserve the  
 844 pretrained generative prior at initialization and enable stable action-conditioned finetuning.

845 **Additional architecture details.** In practice, we apply the routing branch at selected DiT blocks  
 846 and match its spatial resolution to the latent token grid used by the generator. We provide the exact  
 847 layer placement, hidden dimensions, normalization layers, and activation functions in the provided  
 848 supplementary code.

## 849 D.5 Detailed Kinematic-Prior Learning

850 This section expands the kinematic-prior learning objective introduced in Sec. 2.4. The purpose of  
 851 these losses is not to force the router to use all experts uniformly, but to make expert usage agree  
 852 with the physical content of the lifted action field and remain stable across time. In implementation,  
 853 the auxiliary objective is attached to the action encoder and added to the flow-matching training loss  
 854 whenever the structured MoE action pathway is enabled.

855 **Notation and routing statistics.** Let  $K_t \in \mathbb{R}^{B \times T \times H \times W \times 9}$  be the Kinematic-to-Visual Action  
 856 Field. The nine channels are ordered as semantic part indicators  $K_t^{\text{sem}} \in \mathbb{R}^3$ , depth  $K_t^{\text{dep}} \in \mathbb{R}$ ,  
 857 local rotation  $K_t^{\text{rot}} \in \mathbb{R}$ , velocity  $K_t^{\text{vel}} \in \mathbb{R}^3$ , and acceleration  $K_t^{\text{acc}} \in \mathbb{R}$ . The outer router predicts  
 858 token-wise logits:

$$G_{\text{outer}} \in \mathbb{R}^{BT \times |\mathcal{M}| \times H' \times W'}, \quad \mathcal{M} = \{\text{sem}, \text{dep}, \text{rot}, \text{vel}, \text{acc}\}, \quad (30)$$

859 where  $(H', W')$  is the spatial resolution of the routing grid. We denote the soft routing probability by  
 860  $P = \text{softmax}(G)$  and the binary top- $k$  routing mask by  $A \in \{0, 1\}^{BT \times |\mathcal{M}| \times H' \times W'}$ . In our default  
 861 setting, the outer router uses top- $k = 2$  experts per token. For each modality expert  $i$ , the realized  
 862 sparse routing load is measured by:

$$f_i = \mathbb{E}_{b,t,h,w}[A_{b,t,i,h,w}], \quad p_i = \mathbb{E}_{b,t,h,w}[P_{b,t,i,h,w}], \quad \ell_i = f_i p_i. \quad (31)$$

863 Here  $f_i$  captures the hard fraction of tokens assigned to expert  $i$ , while  $p_i$  captures the average soft  
 864 confidence assigned to that expert. Their product follows sparse MoE load estimation and prevents a  
 865 high-probability but rarely selected expert, or a frequently selected but low-confidence expert, from  
 866 being treated as fully utilized.

867 **Kinematic-prior adaptive load balancing.** Uniform load balancing is a poor target for surgical  
 868 manipulation, because the physical causes of appearance change are highly non-uniform across  
 869 frames and actions. We therefore compute a sample-dependent target distribution directly from  $K_t$ .  
 870 First,  $K_t$  is adaptively average-pooled to the router resolution  $(H', W')$ . At each routed token, we  
 871 compute modality energies:

$$\begin{aligned} e_{\text{sem}} &= \|K_t^{\text{sem}}\|_2, & e_{\text{dep}} &= |K_t^{\text{dep}}|, & e_{\text{rot}} &= |K_t^{\text{rot}}|, \\ e_{\text{vel}} &= \|K_t^{\text{vel}}\|_2, & e_{\text{acc}} &= |K_t^{\text{acc}}|. \end{aligned} \quad (32)$$

872 These terms respectively measure tool-part presence, rendered camera-depth magnitude, articulated  
 873 rotation magnitude, projected motion magnitude, and acceleration magnitude. The local physical

874 prior is normalized over modalities:

$$\pi_i(K_t^{b,h,w}) = \frac{e_i(K_t^{b,h,w})}{\sum_{j \in \mathcal{M}} e_j(K_t^{b,h,w})}. \quad (33)$$

875 The batch-level target used by the loss is the mean physical prior:

$$\bar{\pi}_i(K_t) = \mathbb{E}_t^{b,h,w} [\pi_i(K_t^{b,h,w})]. \quad (34)$$

876 The KP-ALB loss then aligns the realized routing load with this physical target:

$$\mathcal{L}_{\text{KP-ALB}} = \frac{1}{|\mathcal{M}|} \sum_{i \in \mathcal{M}} (\ell_i - \text{sg}(\bar{\pi}_i(K_t)))^2, \quad (35)$$

877 where  $\text{sg}(\cdot)$  denotes stop-gradient. This loss differs from standard MoE balancing in two ways. First,  
878 the target is not uniform; it changes with the current action field. Second, the target is grounded in  
879 interpretable kinematic quantities, so semantic/depth experts are favored in static visible tool regions,  
880 rotation experts are favored during wrist articulation, and velocity/acceleration experts are favored  
881 during fast tool motion.

882 **Spatiotemporal routing consistency.** The SRC term reduces frame-to-frame routing flicker around  
883 articulated instruments. Instead of applying temporal smoothness to the entire image, we derive a  
884 binary tool mask from the semantic channels of  $K$ . Specifically, the semantic channels are max-pooled  
885 to  $(H', W')$ , and a token is marked as tool-present when any semantic channel is active. Let:

$$M^{\text{tool}} \in \{0, 1\}^{B \times T \times 1 \times H' \times W'} \quad (36)$$

886 denote this mask. The continuous routing probability used by SRC is the capacity-predictor  
887 probability:

$$R = \sigma(\hat{A}) \in [0, 1]^{B \times T \times |\mathcal{M}| \times H' \times W'}. \quad (37)$$

888 We penalize temporal differences only inside tool regions:

$$\mathcal{L}_{\text{SRC}} = \frac{\sum_{b,t=2}^T \sum_{i,h,w} M_{b,t,1,h,w}^{\text{tool}} (R_{b,t,i,h,w} - R_{b,t-1,i,h,w})^2}{|\mathcal{M}| \sum_{b,t=2}^T \sum_{h,w} M_{b,t,1,h,w}^{\text{tool}}}. \quad (38)$$

889 This formulation stabilizes expert selection on moving tools and tool tips, while avoiding unnecessary  
890 constraints on static background regions. If a clip contains only one frame after temporal sampling,  
891 the SRC term is set to zero.

892 **Tier-2 routing stabilizer.** Within each modality expert, KVLR further routes features to three  
893 sub-experts: fine manipulation, transport motion, and skip. The inner router uses top-1 routing.  
894 Although the main paper groups this under hierarchical routing, the implementation includes a small  
895 auxiliary stabilizer to prevent the collapse of the Tier-2 router:

$$\mathcal{L}_{\text{sub}} = \frac{1}{|\mathcal{M}|} \sum_{m \in \mathcal{M}} |\mathcal{S}| \sum_{s \in \mathcal{S}} f_{m,s} p_{m,s}, \quad \mathcal{S} = \{\text{fine, transport, skip}\}. \quad (39)$$

896 Here  $f_{m,s}$  and  $p_{m,s}$  are the hard assignment fraction and mean soft probability of sub-expert  $s$  inside  
897 modality branch  $m$ . This term is lightweight and acts as an implementation-level stabilization term  
898 for the nested router.

899 **Capacity schedule and final objective.** The routing path is trained with a three-stage capacity  
900 schedule. Stage I uses dense softmax fusion so that all modality experts receive the gradient signal.  
901 Stage II linearly interpolates between dense fusion and sparse top- $k$  fusion. Stage III uses fully  
902 sparse top- $k$  fusion. In the default configuration, Stage I occupies the first 40% of training and Stage  
903 II ends at 75% of training. The auxiliary losses above are active during training whenever a routing  
904 mask is available.

905 The primary generation objective is the flow-matching loss. With latent endpoints  $x_0$  and  $x_1$ , the  
906 training target is:

$$v_t = (1 - \sigma_{\min})x_1 - x_0, \quad (40)$$

907 and the model prediction is trained with mean-squared error:

$$\mathcal{L}_{\text{Flow}} = \mathbb{E} [\|\hat{v}_t - v_t\|_2^2]. \quad (41)$$

908 For image-to-video conditioning, the implementation can exclude conditioned reference frames from  
909 this MSE so that the loss focuses on generated frames.

910 The full training objective for the structured action model is:

$$\mathcal{L}_{\text{total}} = \mathcal{L}_{\text{Flow}} + \lambda_{\text{kp}} \mathcal{L}_{\text{KP-ALB}} + \lambda_{\text{src}} \mathcal{L}_{\text{SRC}} + \lambda_{\text{cp}} \mathcal{L}_{\text{CP}} + \lambda_{\text{sub}} \mathcal{L}_{\text{sub}}. \quad (42)$$

911 Unless otherwise specified, we use:

$$\lambda_{\text{kp}} = 0.01, \quad \lambda_{\text{src}} = 0.005, \quad \lambda_{\text{cp}} = 0.01, \quad \lambda_{\text{sub}} = 0.005. \quad (43)$$

912 In the compact objective in Sec. 2.4,  $\mathcal{L}_{\text{sub}}$  can be viewed as a small nested-router stabilization term  
913 associated with hierarchical routing, while the three primary kinematic-prior losses are KP-ALB,  
914 SRC, and CP.

915 **Capacity-predictor supervision.** The capacity predictor is a lightweight  $1 \times 1$  convolutional  
916 network that predicts which outer experts will be active without recomputing full top- $k$  routing at  
917 inference time. It receives the detached shared action feature  $c_{\text{act}}$  and outputs raw logits:

$$\hat{A} = Q_{\psi}(\text{sg}(c_{\text{act}})) \in \mathbb{R}^{BT \times |\mathcal{M}| \times H' \times W'}. \quad (44)$$

918 The target is the router-induced binary routing mask  $A$ . We train the predictor with binary cross-  
919 entropy with logits:

$$\mathcal{L}_{\text{CP}} = \text{BCE}(\hat{A}, \text{sg}(A)). \quad (45)$$

920 Detaching  $c_{\text{act}}$  and  $A$  makes this term supervise the predictor rather than changing the main router  
921 to become easier to predict. During dense warm-up, all experts contribute through softmax fusion,  
922 but the same top- $k$  mask induced by the router probabilities is used as a synthetic target for this loss.  
923 During sparse training,  $A$  is the actual top- $k$  routing decision.

924 At inference, the predictor probability  $\sigma(\hat{A})$  is compared with an expert-wise threshold. These  
925 thresholds are updated during training by an exponential moving average of expert-wise quantiles:

$$\tau_i \leftarrow \beta \tau_i + (1 - \beta) \text{Quantile}_{1 - \bar{a}_i}(\sigma(\hat{A}_i)), \quad (46)$$

926 where  $\bar{a}_i = \mathbb{E}[A_i]$  is the observed fraction of routed tokens for expert  $i$  and  $\beta = 0.95$  in our implemen-  
927 tation. The quantile level is clamped for numerical stability, but the threshold is updated even when  
928 an expert receives nearly zero or nearly all tokens. This avoids positive-feedback expert starvation:  
929 an unused expert still receives an updated threshold instead of being permanently suppressed.

## 930 D.6 Detailed Action-adaptive Efficient Generation

931 This section provides additional details for the efficient generation extension in Sec. 2.5. The goal is to  
932 reduce inference cost while preserving the action-conditioned behavior of the full structured-control  
933 model. The key principle is to reuse the same signals that drive action-conditioned generation of  
934 structured action lifting and hierarchical routing to determine where dense action-control computation  
935 remains necessary.

936 **Overview.** The efficient model combines two complementary mechanisms. First, a few-step student  
937 reduces the number of denoising evaluations. Second, an action-adaptive execution policy modulates  
938 the cost of action-control updates within each step. Thus, acceleration is achieved at both the  
939 sampling-depth level and the per-step control-computation level.

940 Let  $G_t$  denote the full teacher and  $G_s$  the efficient student. Both receive the same reference image  
941  $I^{\text{ref}}$ , text prompt  $y$ , and articulated action sequence  $a_{1:T}$ . The teacher provides full structured-control  
942 supervision, while the student operates with a reduced denoising schedule and an execution policy  
943 induced by the control pathway.

944 **Action-aware distillation.** The student preserves the teacher’s conditioning interface, including  
945 structured action lifting, hierarchical routing, and action-control injection. We distill the teacher at  
946 three levels.

947 *Prediction-level distillation.* We first match the teacher’s prediction target:

$$\mathcal{L}_{\text{pred}} = \mathbb{E} \left[ \left\| \hat{z}^s - \text{sg}(\hat{z}^t) \right\|_2^2 \right], \quad (47)$$

948 where  $\hat{z}^s$  and  $\hat{z}^t$  denote the student and teacher velocity/noise predictions at the same training state,  
949 and  $\text{sg}(\cdot)$  denotes stop-gradient.

950 *Routing distillation.* Since the routing pathway encodes which physical modalities and motion  
951 operators are active, we further align the outer routing distribution, the inner sub-expert routing  
952 distribution, and the skip probability:

$$\mathcal{L}_{\text{route}} = \text{KL}(\mathbf{R}_{\text{out}}^t \parallel \mathbf{R}_{\text{out}}^s) + \text{KL}(\mathbf{R}_{\text{in}}^t \parallel \mathbf{R}_{\text{in}}^s) + \text{BCE}(p_{\text{skip}}^s, \text{sg}(p_{\text{skip}}^t)). \quad (48)$$

953 *Control-path distillation.* We additionally preserve selected control-path signals:

$$\mathcal{L}_{\text{ctrl}} = \sum_{h \in \mathcal{H}} \left\| h^s - \text{sg}(h^t) \right\|_2^2, \quad (49)$$

954 where  $\mathcal{H}$  includes intermediate action-control features and structured action statistics, such as motion  
955 energy and tool masks. This term keeps the student aligned with the teacher’s action-conditioning  
956 behavior even under a reduced denoising schedule.

957 The overall distillation objective is:

$$\mathcal{L}_{\text{distill}} = \mathcal{L}_{\text{pred}} + \lambda_r \mathcal{L}_{\text{route}} + \lambda_c \mathcal{L}_{\text{ctrl}}. \quad (50)$$

958 **Action significance estimation.** Not all spatio-temporal locations require the same amount of control  
959 computation. Surgical videos contain highly dynamic tool tips and contact regions, but also large  
960 static or slowly changing areas. We therefore estimate a token-level action significance score using  
961 the structured control pathway:

$$s(u) = \alpha e_{\text{motion}}(u) + \beta m_{\text{tool}}(u) + \gamma c_{\text{route}}(u) + \eta q_{\text{fine}}(u) - \delta p_{\text{skip}}(u), \quad (51)$$

962 where  $e_{\text{motion}}$  is computed from the velocity and acceleration channels of the lifted action field,  $m_{\text{tool}}$   
963 is derived from the semantic tool map,  $c_{\text{route}}$  is the confidence of the outer routing distribution,  $q_{\text{fine}}$   
964 is the probability of selecting the fine-manipulation sub-expert, and  $p_{\text{skip}}$  is the skip probability. The  
965 score is normalized per sample to obtain a token-level significance map. In our default implementation,  
966 we use the explicit weighted formulation above for interpretability and stable optimization.

967 **Selective action-control execution.** Given  $s(u)$ , tokens are assigned to three execution modes:

$$\pi(u) \in \{\text{full}, \text{light}, \text{reuse}\}. \quad (52)$$

968 We use budgeted top-ratio selection: the top fraction of tokens receives full updates, the next fraction  
969 receives light updates, and the remaining tokens use reuse. In the default setting, the full and light  
970 ratios are 0.2 and 0.3, respectively.

971 In the current implementation, the execution policy modulates the action-control residual injected  
972 into the student. Full tokens receive the complete action-control update, light tokens receive a scaled  
973 update, and reuse tokens use cached or suppressed residuals. This design is minimally invasive to  
974 the backbone architecture: the student generator remains unchanged, while the structured control  
975 branch induces non-uniform conditioning strength and computation. Unlike generic token pruning,  
976 the execution policy is not based solely on visual redundancy, but is induced by articulated motion,  
977 tool presence, routing confidence, and learned skip behavior.

978 **Temporal refresh and feature caching.** We further derive a frame-level refresh policy from the  
979 average significance of each frame. Let  $\bar{s}(t)$  denote the mean significance over frame  $t$ . We define  
980 the refresh interval as:

$$r(t) = \begin{cases} 1, & \bar{s}(t) \geq \tau_h, \\ 2, & \tau_m \leq \bar{s}(t) < \tau_h, \\ K, & \bar{s}(t) < \tau_m, \end{cases} \quad (53)$$

981 where high-significance frames are refreshed every step and low-significance frames are refreshed  
982 less frequently. During inference, cached action-control residuals can therefore be reused for low-  
983 significance regions, while high-significance regions continue to receive dense updates. This temporal

reuse further reduces the cost of structured control on stable segments without changing the main sampling path.

**Budget-aware training.** To make the trade-off between speed and quality explicit, we introduce a budget regularizer. Let  $\rho_{\text{full}}$ ,  $\rho_{\text{light}}$ , and  $\rho_{\text{refresh}}$  denote the observed full-update, light-update, and refresh ratios. We estimate the active compute ratio as:

$$\rho_{\text{compute}} = w_f \rho_{\text{full}} + w_l \rho_{\text{light}}, \quad (54)$$

and optimize the target-budget objective:

$$\mathcal{L}_{\text{budget}} = |\rho_{\text{compute}} - \rho_{\text{target}}| + w_r |\rho_{\text{refresh}} - \rho_{\text{refresh}}^*|. \quad (55)$$

For learnable significance weights, we additionally use the mean normalized significance as a differentiable proxy for active computation.

To suppress flickering in lightly updated or reused regions, we further impose temporal feature consistency:

$$\mathcal{L}_{\text{temp}} = \sum_{(u,t) \in \Omega_{\text{light/reuse}}} \|f^s(u, t) - \text{sg}(f^s(u, t-1))\|_2^2, \quad (56)$$

where  $\Omega_{\text{light/reuse}}$  denotes the set of tokens assigned to light or reuse modes. The full efficient objective is:

$$\mathcal{L}_{\text{eff}} = \mathcal{L}_{\text{distill}} + \lambda_b \mathcal{L}_{\text{budget}} + \lambda_t \mathcal{L}_{\text{temp}}, \quad (57)$$

where the coefficients are set as  $\lambda_b = 0.5$ ,  $\lambda_t = 0.35$ .

**Training protocol and implementation notes.** The teacher remains unchanged and provides prediction, routing, skip, and control-path supervision. The student uses a four-step generation setting by default. In the first stage, we optimize only the distillation objective in Eq. (50) to stabilize the few-step student. In the second stage, we enable  $\mathcal{L}_{\text{budget}}$  and  $\mathcal{L}_{\text{temp}}$  to train the adaptive execution policy. The resulting efficient extension stays close to the original structured-control model while exposing an interpretable trade-off between generation speed, visual quality, and action fidelity.

## E Supplementary Details on KASA Benchmark Construction

This section provides additional details on the construction of the **Kinematic Action-centric Surgical (KASA)** benchmark introduced in the main paper. Our goal is not to build a generic surgical video corpus, but to curate a benchmark that explicitly supports *action-conditioned* video generation under sparse articulated control. Accordingly, the benchmark is designed to provide lightweight yet visually grounded supervision that can be transformed into the structured control representation used by our model. In particular, each sequence is associated with a temporally consistent action signal, consisting of articulated pose and gripper status, which later supports the construction of semantic, depth, rotation, velocity, and acceleration-aware control cues.

### E.1 Design Objective and Benchmark Scope

A central motivation of this work is that existing controllable surgical video generation resources do not provide a suitable middle ground between two extremes. On one hand, text descriptions are lightweight but too coarse to specify the fine-grained tool motions and contact transitions that drive surgical scene evolution. On the other hand, dense visual priors such as optical flow, full segmentation maps, or manually designed intermediate annotations are expensive to obtain and often unavailable at inference time in realistic robotic settings. The KASA benchmark is constructed to address this gap.

The design principle of KASA is therefore to pair *realistic surgical video sequences* with *compact articulated action supervision* that remains physically meaningful and visually projectable. Rather than storing only categorical labels or global task descriptions, KASA provides an articulated representation that preserves the tool pose, orientation, and gripper status over time. This makes the benchmark particularly suitable for studying controllable generation under sparse action inputs, where the challenge is not only to synthesize plausible appearance, but also to maintain visual evolution consistent with the underlying surgical manipulation.

Another design consideration is alignment with the model structure introduced in the main paper. Our method decomposes action conditioning into multiple physical factors, including part semantics,

1028 geometry, orientation, and motion. For this reason, the benchmark construction pipeline is organized  
1029 around two complementary supervision sources:

- 1030 • **Human-in-the-loop Semantic Annotation (HSA)**, which provides reliable part-aware  
1031 semantic masks for articulated instrument components;
- 1032 • **Differentiable Pose Tracking (DPT)**, which recovers temporally consistent articulated  
1033 geometry directly from videos.

1034 Together, these two branches provide the minimal supervision needed to bridge sparse articulated  
1035 control and dense video synthesis.

## 1036 E.2 Task-Driven Sequence Selection

1037 We curate KASA from the SAR-RARP dataset, focusing on robotic surgical sequences performed  
1038 with da Vinci Large Needle Drivers. The final benchmark contains 1,047 scenes and 105,175  
1039 frames. Rather than uniformly sampling all available surgical actions, we deliberately focus on three  
1040 representative sub-actions: *needle grasping*, *needle puncture*, and *knotting*.

1041 This task selection is driven by the type of action variability we want the benchmark to capture.  
1042 These three sub-actions jointly span a broad range of articulated surgical dynamics. Needle grasping  
1043 often contains localized pose adjustment and contact preparation. Needle puncture requires coordi-  
1044 nated tool reorientation and short-horizon motion transitions around tissue interaction. Knotting  
1045 further introduces more complex trajectories, including larger transport motion, interaction-induced  
1046 deformation, and temporally extended manipulation. As a result, the benchmark contains both  
1047 high-velocity transport motion and low-velocity fine manipulation, which is important for evaluating  
1048 the motion-scale specialization behavior of our hierarchical routing framework.

1049 To preserve temporal continuity for pose tracking while keeping the data manageable, all videos are  
1050 downsampled to 30 FPS. After preprocessing, each sequence contains between 55 and 235 frames.  
1051 This sequence length is sufficient to preserve meaningful action evolution while avoiding excessive  
1052 redundancy. Importantly, the resulting clips remain challenging due to articulated motion, frequent  
1053 self-occlusion, viewpoint variation, and complex instrument interaction of tissue. These properties  
1054 make KASA a suitable benchmark not only for generation fidelity, but also for evaluating whether a  
1055 model can remain faithful to sparse action control under heterogeneous surgical dynamics.

## 1056 E.3 Human-in-the-loop Semantic Annotation

1057 The first branch of the KASA pipeline is devoted to obtaining reliable part-aware semantic supervision.  
1058 Because our action representation later separates physical modalities, it is important that the semantic  
1059 component be clean and unambiguous. In robotic surgery, however, instrument appearance is  
1060 challenging to segment consistently due to specular highlights, motion blur, partial occlusion, and  
1061 visual similarity across articulated parts. Therefore, naively relying on automatic segmentation tends  
1062 to produce part confusion, boundary leakage, or inconsistent labels over time.

1063 To address this issue, we adopt a Human-in-the-loop Semantic Annotation (HSA) strategy. For each  
1064 video, we begin from a reference frame and manually specify a small set of wrist keypoints used to  
1065 initialize the geometric tracking stage described later. In parallel, we estimate part-level segmentation  
1066 masks for the main articulated instrument components of the shaft, wrist, and gripper, using a SAM2  
1067 model finetuned on SAR-RARP. These predictions are then temporally propagated across frames to  
1068 provide an initial dense annotation.

1069 Automatic propagation alone is not sufficient for the level of semantic precision required by structured  
1070 action conditioning. We therefore refine the propagated masks through interactive human correction.  
1071 The purpose of this stage is not simply to improve visual mask quality in a generic sense, but  
1072 specifically to eliminate errors that would contaminate the semantic branch of the action representation.  
1073 Typical corrections include:

- 1074 • resolving part confusion between adjacent articulated components;
- 1075 • correcting mask leakage into background tissue or other instruments;
- 1076 • refining ambiguous boundaries under self-occlusion or motion blur;

1077 • enforcing temporally consistent labeling across consecutive frames.

1078 The final output of HSA is a set of high-quality part-aware semantic masks for each sequence. These  
1079 masks form the basis of the semantic channels in the Kinematic-to-Visual Action Field. In practice,  
1080 the three semantic channels correspond to the shaft, wrist, and gripper regions, respectively. This  
1081 explicit part decomposition is important because different articulated regions often play different  
1082 functional roles during manipulation, and collapsing them into a single foreground mask would  
1083 weaken the structured control signal.

#### 1084 **E.4 Differentiable Pose Tracking**

1085 While HSA provides semantic supervision, action-conditioned generation also requires continuous  
1086 geometric and motion information. For this reason, the second branch of the KASA pipeline recovers  
1087 temporally consistent articulated pose trajectories directly from videos. Our goal is to avoid reliance  
1088 on external robot logs or dedicated hardware tracking systems, since such signals are often unavailable  
1089 or difficult to synchronize in retrospective video collections. We therefore design a video-based  
1090 **Differentiable Pose Tracking (DPT)** procedure.

1091 At a high level, DPT formulates articulated pose estimation as a render-and-compare problem. We  
1092 start from a controllable geometric representation initialized from instrument CAD priors. Given a  
1093 candidate articulated pose, we render part-aware observations in the camera view and compare them  
1094 against the observed video evidence. This allows the optimization to remain tied to the image plane  
1095 while still preserving a physically meaningful articulated parameterization.

1096 The optimization objective is designed to capture both global alignment and local manipulation  
1097 details. A standard part-aware silhouette alignment term encourages the rendered geometry to match  
1098 the observed instrument masks. To better capture fine-grained distal motion, especially around the  
1099 tool tip and gripper region, we additionally use a targeted matching constraint for the articulated  
1100 end region. This is important because small pose errors near the distal components can lead to  
1101 disproportionately large errors in the inferred manipulation dynamics.

1102 Temporal coherence is introduced in two ways. First, the initial frame is anchored using Perspective-  
1103 n-Point initialization from manually specified wrist keypoints. Second, for subsequent frames, we  
1104 propagate pose estimates using temporally local correspondences. Concretely, wrist-related visual  
1105 features are tracked across time and lifted to 3D using rendered depth, producing robust 2D-3D  
1106 correspondences for coarse initialization before iterative optimization. This design stabilizes the  
1107 trajectory and reduces the chance of drifting into implausible local minima.

1108 In addition, long sequences may still contain severe occlusion or abrupt appearance changes that  
1109 temporarily degrade tracking quality. To improve robustness, DPT maintains a keyframe memory  
1110 pool of high-confidence articulated states. When the current tracking quality drops, the system can  
1111 recover by matching the observation to this memory pool and reinitializing pose estimation through a  
1112 new PnP stage. This reinitialization mechanism is particularly useful in clips with self-occlusion or  
1113 rapid tool reorientation.

1114 The output of DPT is a temporally consistent articulated pose sequence for each video. This output  
1115 provides the geometric foundation in KASA. Once the articulated states are estimated, they can be  
1116 rendered or projected into multiple visual control channels, including depth, local rotation, projected  
1117 velocity, and acceleration-related cues.

#### 1118 **E.5 Quality Validation of KASA Dataset**

1119 We validate the quality of KASA annotations through a render-and-compare consistency check. For  
1120 each recovered articulated pose, we re-pose the controllable instrument Gaussian representation  
1121 and render its part-aware silhouette into the imaging space. The rendered silhouette is then compared  
1122 with the corresponding segmentation mask using the Dice score. This provides a direct geometric  
1123 measure of whether the recovered pose explains the observed instrument region. We use this score as  
1124 a quality indicator for DPT and remove low-confidence cases with poor silhouette-mask agreement  
1125 before constructing the final action-supervised dataset.

1126 Tab. 9 reports the Dice scores of KASA annotations across three surgical sub-actions: *NeedleGrasp-*  
1127 *ing*, *Knotting*, and *NeedlePuncture*. We evaluate the overall instrument region and three articulated  
1128 parts, including the wrist, shaft, and gripper. We compare DPT with and without matching-based

Table 9: Quality validation of KASA annotations. Dice scores are computed between rendered silhouettes from recovered poses and segmentation masks.

| Setting               | Surgical Task  | Dice Score $\uparrow$ |                   |                   |                   |
|-----------------------|----------------|-----------------------|-------------------|-------------------|-------------------|
|                       |                | Overall Instrument    | Wrist Part        | Shaft Part        | Gripper Part      |
| w/ Mem-pool Recovery  | NeedleGrasping | $0.967 \pm 0.056$     | $0.892 \pm 0.031$ | $0.978 \pm 0.016$ | $0.775 \pm 0.067$ |
|                       | Knotting       | $0.966 \pm 0.027$     | $0.885 \pm 0.047$ | $0.979 \pm 0.020$ | $0.757 \pm 0.082$ |
|                       | NeedlePuncture | $0.947 \pm 0.081$     | $0.887 \pm 0.038$ | $0.982 \pm 0.015$ | $0.787 \pm 0.070$ |
| w/o Mem-pool Recovery | NeedleGrasping | $0.936 \pm 0.063$     | $0.864 \pm 0.051$ | $0.946 \pm 0.049$ | $0.758 \pm 0.070$ |
|                       | Knotting       | $0.937 \pm 0.029$     | $0.852 \pm 0.047$ | $0.947 \pm 0.022$ | $0.726 \pm 0.082$ |
|                       | NeedlePuncture | $0.919 \pm 0.085$     | $0.841 \pm 0.067$ | $0.953 \pm 0.035$ | $0.749 \pm 0.073$ |

recovery using memory pool. Without memory-pool recovery, the tracker continues optimization from the current pose estimate when tracking quality drops, rather than retrieving reliable historical pose candidates for re-initialization. As a result, pose drift can accumulate and lead to complete tracking loss, producing missing or severely misaligned predictions with lower Dice scores. With memory-pool recovery, reliable historical pose candidates are used to re-initialize pose optimization once low-quality tracking is detected. As shown in Tab. 9, memory-pool recovery improves silhouette-mask alignment across most tasks and parts, suggesting more stable pose recovery under challenging surgical dynamics. Remaining low-Dice cases are mainly caused by unstable segmentation under complex illumination, instruments moving partially outside the camera field of view, or severe tissue occlusion that hides the instrument and disrupts pose tracking. In this way, the dice-based validation helps identify unreliable annotations and improves the consistency of the final KASA dataset.

## E.6 From Curated Supervision to Structured Action Control

The KASA benchmark is constructed not merely as an annotation resource, but as the data foundation for the structured action representation used by our generation framework. The HSA and DPT branches serve different but complementary roles in this process.

HSA provides reliable *part-aware semantics*. These labels determine where each articulated component appears in image space and prevent semantic ambiguity between the shaft, wrist, and gripper regions. DPT provides temporally consistent *articulated geometry and motion*. These trajectories determine how the instrument is positioned, oriented, and displaced over time. By combining the outputs of both branches, we can derive the set of physically interpretable channels used in the Kinematic-to-Visual Action Field.

More specifically, the semantic masks obtained from HSA are rasterized into part-aware semantic maps. The articulated poses recovered by DPT are used to render depth and local orientation descriptors, and to compute projected velocity and acceleration cues over time. These quantities are then aligned in image space to form the structured control tensor introduced in the main paper. Therefore, the purpose of KASA is not only to provide a benchmark for evaluation, but also to make the sparse articulated control signal learnable by transforming it into a visually grounded supervision source.

## E.7 Discussion

We emphasize that the contribution of KASA is not to propose a standalone tracking or annotation system optimized for surgical reconstruction as an end goal. Instead, the benchmark is constructed to support a specific research question: *can realistic surgical videos be generated faithfully from sparse articulated action control* ? From this perspective, the most important property of KASA is that it exposes a compact action interface that remains consistent with dense visual scene evolution.

This benchmark design offers two practical advantages. First, the action signal is lightweight and closer to the type of supervision or control interface that can plausibly be used in robotic settings. Second, because the supervision remains visually grounded, it can be directly transformed into structured intermediate control for generation. We believe this makes KASA a useful benchmark for studying controllable surgical video generation beyond text-only conditioning and beyond heavy dense priors that are difficult to access during inference.

Table 10: KASA usability validation through instrument segmentation. Standard perception baselines are trained on the KASA training split and evaluated on the held-out validation split.

| Method            | Part mIoU $\uparrow$ | Tool IoU $\uparrow$ | Boundary F-score $\uparrow$ |
|-------------------|----------------------|---------------------|-----------------------------|
| U-Net [42]        | 68.4                 | 86.7                | 71.5                        |
| DeepLabV3+ [4]    | 71.8                 | 88.9                | 74.2                        |
| SegFormer-B1 [49] | 75.6                 | 91.3                | 77.8                        |
| Mask2Former-S [7] | <b>77.1</b>          | <b>92.0</b>         | <b>79.4</b>                 |

Table 11: Visual-to-action consistency on KASA. Standard video encoders predict action trajectories from held-out validation clips.

| Method                      | Trans. MAE (mm) $\downarrow$ | Rot. MAE (deg) $\downarrow$ | Gripper Acc. (%) $\uparrow$ | DTW $\downarrow$ |
|-----------------------------|------------------------------|-----------------------------|-----------------------------|------------------|
| Mean trajectory             | 9.84                         | 13.72                       | 61.5                        | 1.00             |
| R3D-18 regressor [44]       | 5.62                         | 8.41                        | 82.3                        | 0.58             |
| Video Swin-T regressor [32] | 4.91                         | 7.36                        | 86.7                        | 0.51             |
| TimeSformer-S regressor [9] | <b>4.58</b>                  | <b>6.92</b>                 | <b>88.1</b>                 | <b>0.47</b>      |

## F Additional Experiments

### F.1 Benchmark Usability Validation of KASA

The previous section focuses on evaluation independence. We further validate KASA as a standalone benchmark artifact using only the final KASA annotations, without relying on intermediate annotations or tracking outputs. Since intermediate correction states are not retained for all videos, we do not claim a full audit of every annotation step. Instead, we provide end-to-end usability checks showing that the final annotations are learnable, visually grounded, and useful for downstream evaluation.

**Instrument part segmentation.** We first train standard segmentation baselines on the KASA training split and evaluate them on the held-out validation split. The models predict four classes: background, shaft, wrist, and gripper. We report mean IoU over instrument parts, foreground IoU for the merged tool region, and boundary F-score. As shown in Tab. 10, standard perception models obtain consistent validation performance, suggesting that the final semantic annotations are sufficiently coherent to support downstream perception tasks.

**Visual-to-action consistency.** To assess whether the final action supervision is visually grounded, we train video-to-action regressors that predict articulated trajectories from the corresponding video clips. The models are trained only on the KASA training split and evaluated on held-out validation clips. We report translation MAE, rotation MAE, gripper-state accuracy, and trajectory DTW distance. As shown in Tab. 11, learned video encoders predict the action trajectories substantially better than a mean-trajectory baseline, suggesting that the final action supervision is correlated with observable tool motion.

**Downstream perception with synthetic augmentation.** We finally examine whether KASA can support downstream perception evaluation and whether action-conditioned synthetic videos provide useful training data. We train a SegFormer-B1 [49] segmentation model under different augmentation settings and evaluate all models on the same real KASA validation split. Synthetic videos are generated from the training split only, and validation videos are never used for generation or training. As shown in Tab. 12, adding synthetic data from KVLRL improves real validation performance more consistently than adding synthetic data from the strongest baseline. This suggests that KASA provides a meaningful downstream benchmark and that action-conditioned generation yields useful synthetic supervision beyond improving generation metrics.

### F.2 Evaluation Independence and Bias Control

A potential concern is that the KASA benchmark construction, the proposed Kinematic-to-Visual Action Field, and the action-faithfulness metrics all involve articulated surgical actions. We therefore provide additional analyses to clarify the separation between supervision, conditioning, and evaluation, and to verify that the reported gains are not merely due to optimizing a representation-specific metric.

Table 12: Downstream perception evaluation on KASA with synthetic data augmentation. SegFormer-B1 [49] is trained under different augmentation settings and evaluated on the real KASA validation split.

| Training Data                  | Part mIoU $\uparrow$ | Tool IoU $\uparrow$ | Boundary F-score $\uparrow$ | Temp. Consistency $\uparrow$ |
|--------------------------------|----------------------|---------------------|-----------------------------|------------------------------|
| Real only                      | 75.6                 | 91.3                | 77.8                        | 84.5                         |
| Real + Cosmos-H-Surgical* [17] | 76.7                 | 91.9                | 78.6                        | 85.2                         |
| Real + KVLR-fast               | 78.4                 | 92.8                | 80.1                        | 86.8                         |
| Real + KVLR                    | <b>79.6</b>          | <b>93.4</b>         | <b>81.0</b>                 | <b>87.5</b>                  |

Table 13: Cross-evaluator validation of action faithfulness. Grounded SAM [41] is independently trained and is not used for KASA construction or KVLR training. Results are averaged over the three KASA action subsets.

| Evaluator         | Method                  | CD $\downarrow$ | TI $\uparrow$ | AF $\downarrow$ |
|-------------------|-------------------------|-----------------|---------------|-----------------|
| Grounded SAM [41] | Cosmos-H-Surgical* [17] | 94.65           | 0.78          | 8.52            |
| Grounded SAM [41] | KVLR                    | 86.90           | 0.89          | 4.48            |
| SAM2 [40]         | Cosmos-H-Surgical* [17] | 101.34          | 0.75          | 9.17            |
| SAM2 [40]         | KVLR                    | 90.82           | 0.86          | 5.03            |

**Separation between action supervision, model conditioning, and evaluation.** KASA provides temporally consistent articulated supervision obtained from human-in-the-loop semantic annotation and differentiable pose tracking. During training, KVLR uses this supervision to construct the Kinematic-to-Visual Action Field as a conditioning input. During evaluation, however, the generated videos are first processed by an external visual extractor to obtain instrument masks or trajectories, and these extracted visual measurements are then compared with the projected action references. Therefore, the evaluation does not directly access the internal routed action features or the model’s conditioning branch. Nevertheless, because the default evaluator uses the same semantic categories as KASA, we further introduce representation-independent checks below.

**Cross-evaluator action-faithfulness evaluation.** To reduce dependence on the annotation model used during KASA construction, we evaluate generated videos with an independent instrument extractor that is not used for constructing KASA or training KVLR. Specifically, we use a separately trained surgical instrument segmentation model (Grounded SAM [41]), and recompute Chamfer Distance (CD), Temporal IoU (TI), and Area Flicker (AF) on the validation set. As shown in Tab. 13, KVLR remains consistently better than the strongest action-conditioned baseline under both evaluators. This suggests that the improvement is not specific to the evaluation pipeline.

**Trajectory-level verification independent of mask overlap.** Mask-overlap metrics may still favor methods that better match the projected action masks. We therefore evaluate action consistency using trajectory-level measurements extracted from generated videos. For each generated clip, we estimate the visible tool centerline and tip trajectory using Grounded SAM [41], and compare them with the projected reference trajectory. We report mean centerline error, endpoint error, and motion-direction accuracy. As shown in Tab. 14, KVLR also improves trajectory-level alignment, indicating that the generated motion better follows the specified articulated action beyond mask-overlap agreement.

**Action perturbation sanity check.** Finally, we verify that KVLR uses the provided action trajectory rather than only relying on the reference image or dataset prior. We evaluate three perturbed-control settings: *shuffled action*, where the action trajectory is randomly sampled from another clip of the same category; *wrong-category action*, where the action is sampled from a different sub-action category; and *reversed action*, where the temporal order of the action sequence is reversed. Tab. 15 shows that perturbing the action substantially degrades action-faithfulness metrics while only moderately affecting image-level fidelity. This confirms that the action input has a causal influence on the generated motion, and that the default improvements are not caused by simply replaying appearance priors from the reference image.

### F.3 Quantitative Analysis of Hierarchical Expert Specialization

Fig. 7 further quantifies how the learned routing behavior evolves with motion intensity. We compute clip-level routing statistics over motion-magnitude quantiles and report the corresponding means and 95% confidence intervals. At Tier 1, the activation mass gradually shifts away from the semantic

Table 14: Trajectory-level action verification with an independent visual tracker. Lower centerline and endpoint errors are better; higher direction accuracy is better.

| Method                  | Centerline Error (px) ↓ | Endpoint Error (px) ↓ | Direction Acc. (%) ↑ |
|-------------------------|-------------------------|-----------------------|----------------------|
| Cosmos-H-Surgical* [17] | 11.83                   | 15.62                 | 71.4                 |
| KVLR-fast               | 8.05                    | 10.24                 | 82.1                 |
| KVLR                    | <b>7.21</b>             | <b>9.13</b>           | <b>84.8</b>          |

Table 15: Action perturbation sanity check for KVLR. Perturbing the input action reduces action faithfulness, showing that the model depends on the specified action trajectory. Results are averaged over the validation subset.

| Control Input         | CD ↓         | TI ↑        | AF ↓        | FID ↓        |
|-----------------------|--------------|-------------|-------------|--------------|
| Matched action        | <b>86.90</b> | <b>0.89</b> | <b>4.48</b> | <b>14.38</b> |
| Shuffled action       | 126.75       | 0.64        | 10.82       | 18.94        |
| Wrong-category action | 138.42       | 0.58        | 12.36       | 21.15        |
| Reversed action       | 119.63       | 0.67        | 9.74        | 17.88        |

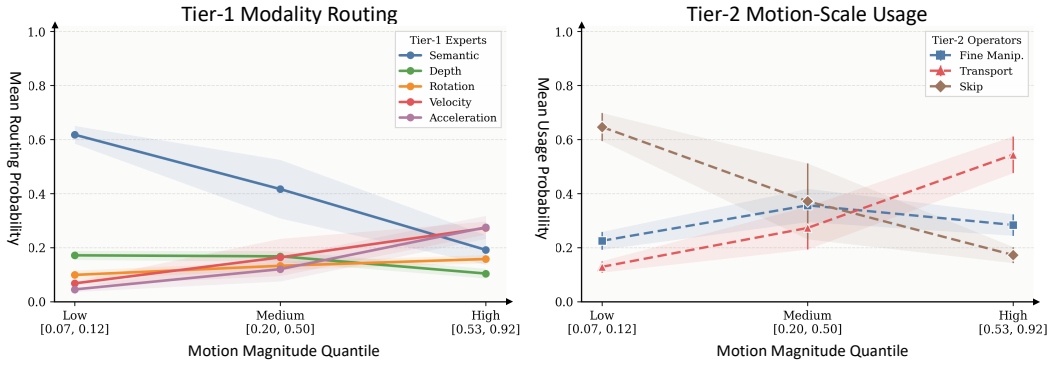


Figure 7: Motion-specialized routing statistics. We group tokens by motion-magnitude quantiles and report the mean routing mass with 95% confidence intervals. At Tier 1, expert usage shifts from semantic-dominant routing in low-motion regions to velocity- and acceleration-dominant routing as motion increases, indicating that the model allocates more computation to dynamic cues when articulated changes become larger. At Tier 2, the routing pattern transitions from skip-dominant behavior in low-motion bins to transport-dominant behavior in high-motion bins, while fine-manipulation operators are most active in the intermediate regime. These results quantitatively confirm that the proposed hierarchical routing learns a motion-dependent decomposition of surgical control rather than a static or uniform expert allocation.

branch and toward the velocity and acceleration branches as motion magnitude increases, showing that the router does not amplify all experts uniformly, but selectively emphasizes dynamics-sensitive pathways when stronger articulated motion is present. At Tier 2, the operator distribution changes from skip-dominant behavior in low-motion regions to transport-dominant behavior in high-motion regions, while fine-manipulation experts remain most competitive in the intermediate range. This intermediate regime also exhibits wider confidence intervals, suggesting that it contains the most heterogeneous micro-dynamics across clips, where subtle local adjustments and larger motion transitions co-exist. Overall, these statistics provide quantitative evidence that hierarchical routing captures a meaningful motion-specialized control decomposition, aligning low-motion states with semantic/static processing and high-motion states with transport-oriented dynamic computation.

#### F.4 Detailed Efficient-Generation Ablation Breakdown

We expand the main efficiency ablation to isolate what each component contributes. For every row, we report the latency and FLOP speedup relative to the 50-step teacher, the marginal compute saving relative to the immediately preceding row, the CD-based control retention, and the FID increase relative to the teacher. The control-retention ratio is computed as  $CD_{\text{teacher}}/CD_{\text{method}}$ , since lower CD indicates better action alignment.

Table 16: Expanded Efficient-Generation Ablation. All metrics are evaluated on the KASA validation split. Latency and FLOPs are measured per 17-frame clip.

| Group                                            | Setting                             | Latency<br>(s) ↓ | FLOPs<br>(PF/clip) ↓ | CD<br>↓ | FID<br>↓ |
|--------------------------------------------------|-------------------------------------|------------------|----------------------|---------|----------|
| <b>Teacher</b>                                   | Full 50-step structured generator   | 4.64             | 6.94                 | 86.89   | 14.37    |
| <b>Distillation Signal</b><br>(Base: 4-step)     | Prediction-only (Generic)           | 0.52             | 0.78                 | 98.45   | 19.82    |
|                                                  | Prediction + Route                  | 0.52             | 0.78                 | 92.10   | 17.50    |
|                                                  | Prediction + Control-path           | 0.52             | 0.78                 | 90.50   | 16.80    |
|                                                  | Pred + Route + Ctrl (Action-aware)  | 0.52             | 0.78                 | 87.50   | 15.60    |
| <b>Spatial Execution</b><br>(Base: Action-aware) | Full update for all tokens          | 0.52             | 0.78                 | 87.50   | 15.60    |
|                                                  | Full/Reuse only (0.35, 0.00, 0.65)  | 0.38             | 0.57                 | 91.20   | 17.90    |
|                                                  | Full/Light only (0.20, 0.80, 0.00)  | 0.46             | 0.69                 | 88.00   | 16.10    |
|                                                  | Full/Light/Reuse (0.20, 0.30, 0.50) | 0.42             | 0.63                 | 88.65   | 16.55    |
| <b>Temporal Cache</b><br>(Base: Spatial)         | Cache disabled (1, 1, 1)            | 0.42             | 0.63                 | 88.65   | 16.55    |
|                                                  | Conservative (1, 2, 2)              | 0.39             | 0.59                 | 89.10   | 17.05    |
|                                                  | Default intervals (1, 2, 4)         | 0.37             | 0.55                 | 89.87   | 17.67    |
|                                                  | Aggressive reuse (1, 4, 8)          | 0.33             | 0.50                 | 92.40   | 19.30    |

**Overall Interpretation.** The baseline table separates three effects that are conflated in the final KVLr-fast number: few-step acceleration, action-aware preservation of the teacher’s control policy, and routing-derived adaptive execution. Generic few-step distillation (DMD2) [56] is computationally effective but action-destructive: it gives an  $8.92\times$  latency speedup, yet retains only 88.3% of the teacher’s CD-based control accuracy and increases FID by 5.45. The action-aware student keeps exactly the same measured cost as DMD2 but recovers the control retention to 99.3% and reduces the FID increase to +1.23. Spatial adaptive execution and temporal caching then provide additional compute reductions, yielding the final  $12.54\times$  latency speedup while retaining 96.7% of teacher control accuracy.

**Expanded Component Validations.** To further isolate the internal dynamics of the efficient student and the adaptive execution policy, we introduce an expanded ablation matrix. Tab. 16 summarizes these diagnostic settings, evaluating the independent contributions of the distillation signals, spatial execution policies, and temporal refresh rates.

**Analysis of Distillation Signals.** The distillation ablation isolates the value of the structured control pathway. While prediction-only distillation (DMD2) [56] severely degrades action faithfulness (CD 98.45), adding routing alignment explicitly informs the student which physical modalities are active, dropping CD to 92.10. However, preserving the intermediate action-control features (Control-path) provides an even stronger independent signal (CD 90.50). The combination of both yields the full action-aware student (CD 87.50), proving that route and control-path supervision are synergistic and transfer the teacher’s action policy at zero additional inference cost.

**Analysis of Spatial Execution.** The spatial execution ablation tests whether the scaled “light” branch is more effective than hard token suppression. A strict Full/Reuse policy (allocating 35% of tokens to full updates and aggressively reusing the rest to match the target FLOPs) drops latency to 0.38 s but penalizes quality heavily (CD 91.20), as critical edge cases are suppressed. Introducing the “light” branch (Full/Light/Reuse at 20/30/50) allows the model to maintain broad contextual awareness with cheaper cached-routing updates, yielding a strictly better Pareto frontier (Latency 0.42 s, CD 88.65).

**Analysis of Temporal Cache.** The temporal cache ablation exposes the deployment trade-offs of feature caching. Disabling the cache entirely relies strictly on spatial token sparsification (0.42 s). Enabling conservative caching (1, 2, 2) yields “free” latency improvements (0.39 s) with negligible CD degradation (+0.45). Our default (1, 2, 4) setting pushes this efficiency further. However, aggressive cache reuse (1, 4, 8) introduces noticeable temporal drift; as the instrument moves out of the cached bounding regions, the student fails to update the active kinematics, resulting in a sharp penalty to both CD (92.40) and visual fidelity (FID 19.30).

## F.5 More Qualitative Results.

**Qualitative results on architecture design.** Fig. 8 provides additional qualitative comparisons for the architecture ablation. The text-only baseline fails to maintain a clear instrument appearance and exhibits inaccurate tool motion, especially in the later frames. Directly injecting raw action vectors yields only limited improvement: although the generated motion is partially guided by the input

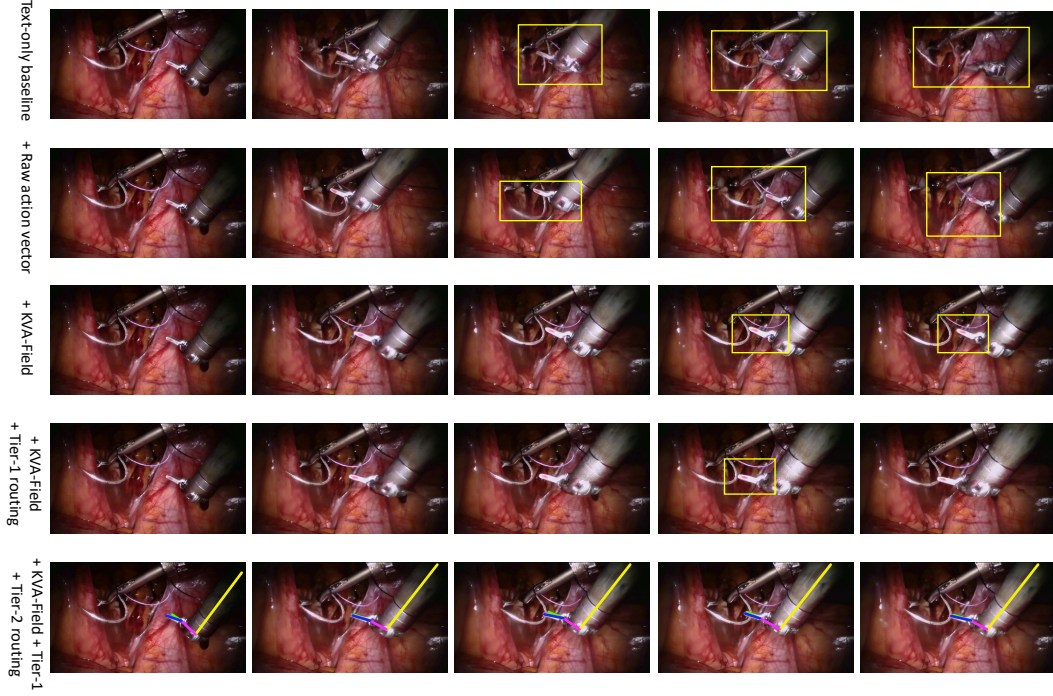


Figure 8: Additional qualitative comparisons on architecture design. We compare representative variants from the architecture ablation, including text-only generation, direct raw-action conditioning, KVA-Field conditioning, Tier-1 routing, and the full hierarchical routing model. The full model produces clearer tool appearance and more action-consistent temporal evolution.

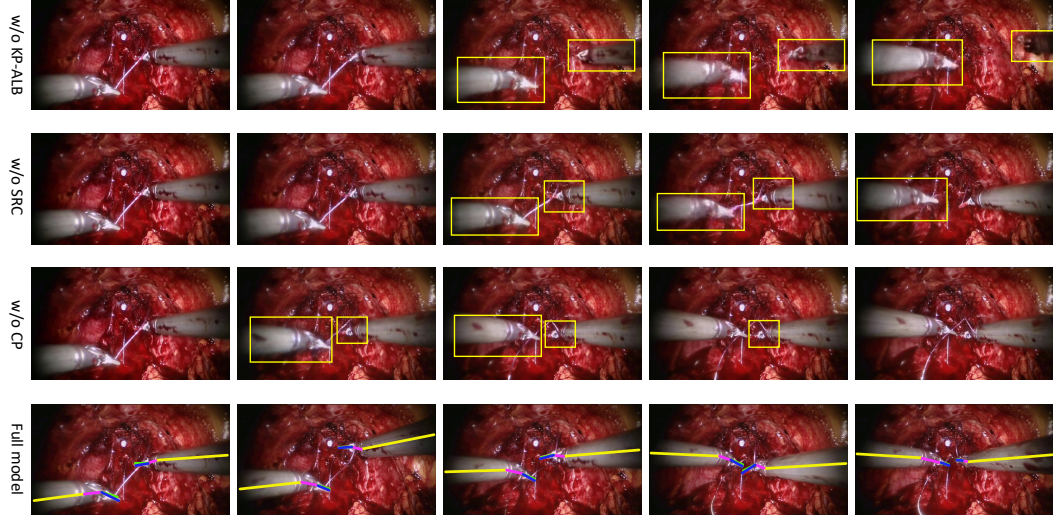


Figure 9: Additional qualitative comparisons on kinematic-prior losses. We compare the full model with variants that remove KP-ALB, SRC, or CP. The proposed kinematic-prior losses improve spatial alignment and temporal consistency, leading to more stable action-conditioned generation.

1292 actions, the instrument trajectory remains imprecise, and the visual quality degrades over time. This  
 1293 supports our observation that low-dimensional articulated kinematics are difficult to use directly for  
 1294 dense image-space video evolution. Replacing raw actions with the proposed KVA-Field substantially  
 1295 improves the spatial alignment between action inputs and visual generation, leading to clearer tool  
 1296 structures and more consistent motion. Introducing Tier-1 routing further improves the use of different  
 1297 physical control modalities, while the full model with both Tier-1 and Tier-2 routing achieves the best  
 1298 qualitative results by adapting control to both modality types and motion scales. These qualitative

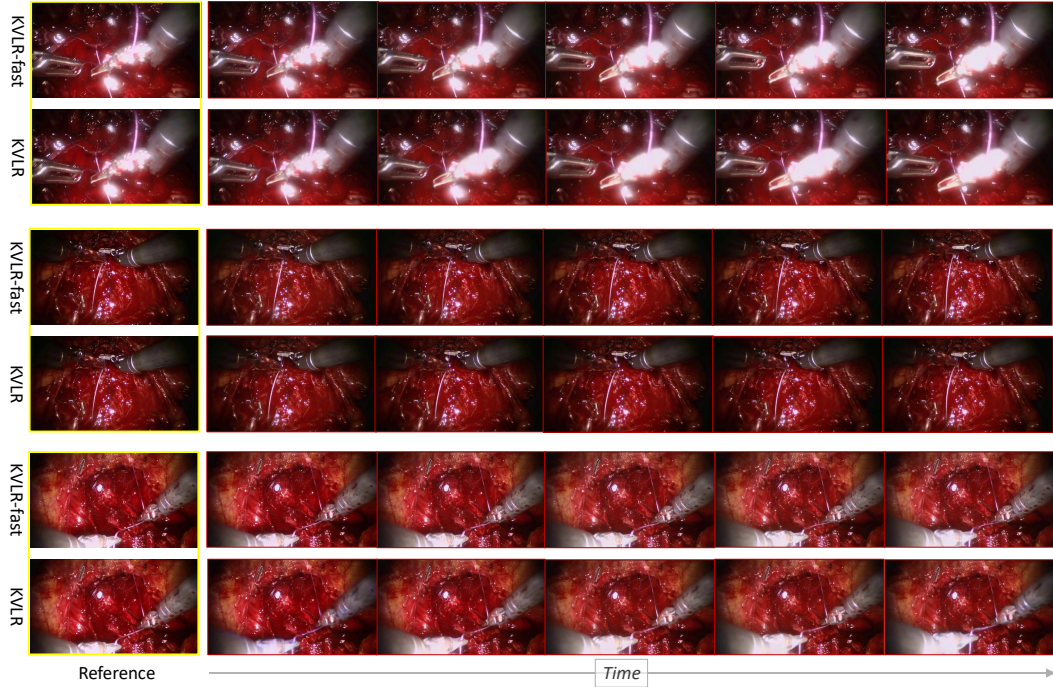


Figure 10: Additional generation results of KVLR and KVLR-fast. We show more examples across diverse surgical scenes and actions. Both the full model and the efficient variant preserve action-consistent visual evolution, while KVLR generally yields slightly sharper appearance details and stronger temporal stability under complex motion. These examples further show that the action-adaptive efficient model retains most of the control fidelity of the full generator despite using substantially reduced computation.

trends are consistent with the quantitative architecture ablation in Tab. 6, showing that action lifting and hierarchical routing are jointly important for action-controlled surgical video generation.

**Qualitative results on kinematic-prior losses.** Fig. 9 shows additional qualitative comparisons for the kinematic-prior loss ablation. Replacing KP-ALB with uniform load balancing [43, 10, 60] weakens the correspondence between routing decisions and physical action cues, resulting in degraded spatial alignment and blurrier tool appearance. Removing SRC leads to visibly less stable temporal evolution, with tool structures and interaction regions becoming inconsistent in the later frames. This suggests that temporal regularization of routing decisions is important for maintaining coherent action-conditioned generation across video frames. Removing CP also degrades the results, indicating that predictable routing helps stabilize expert selection and supports reliable control during generation. Overall, the full model produces the most stable and visually faithful results, confirming that the three kinematic-prior losses are complementary and collectively improve the quality of routed visual control.

**Additional qualitative results.** We provide additional qualitative results in Figs. 10 and 11. These visualizations cover a wider range of scenes, action trajectories, and procedural stages across *Knotting*, *NeedleGrasping*, and *NeedlePuncture*, and further illustrate two key properties of the proposed framework. As shown in Fig. 10, both KVLR and KVLR-fast preserve action-consistent scene evolution under diverse articulated controls, including large tool transport, fine local adjustment, and contact-heavy interaction. Moreover, as illustrated in Figs. 11, the generated videos remain visually coherent across challenging conditions such as appearance variation, illumination change, cluttered backgrounds, and different instrument poses. Together with the main qualitative comparisons, these results further support that the proposed structured kinematic-to-visual control improves both controllability and visual stability across heterogeneous surgical dynamics.

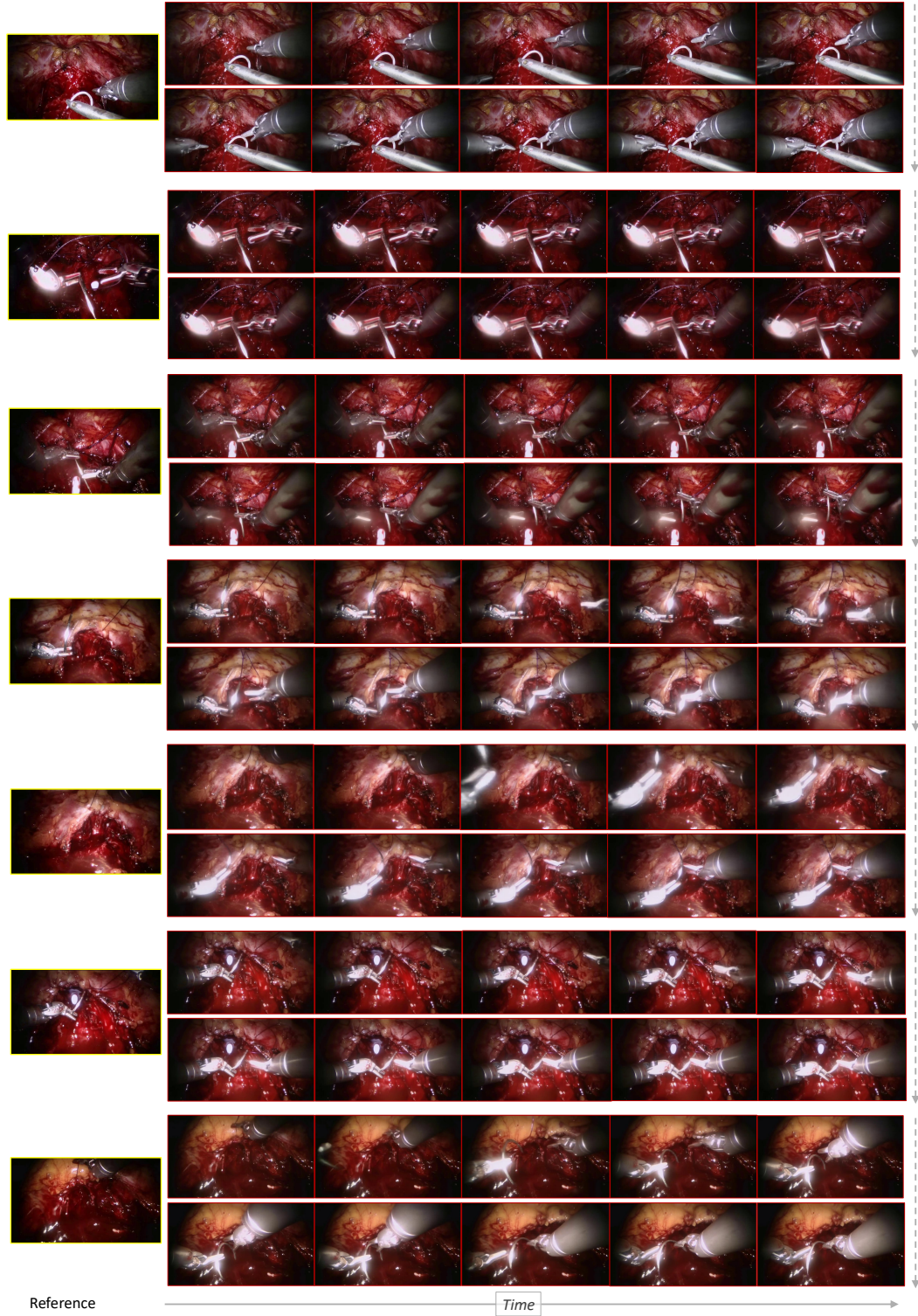


Figure 11: Additional diverse generation results of KVLR on different surgical actions. The generated videos remain consistent with the prescribed articulated controls across different motion patterns, including fine manipulation, larger transport motion, and varying tool–tissue interaction stages. These results highlight the robustness of the proposed structured control interface across multiple action categories and scene configurations.

## 1322 NeurIPS Paper Checklist

### 1323 1. Claims

1324 Question: Do the main claims made in the abstract and introduction accurately reflect the  
1325 paper’s contributions and scope?

1326 Answer: [Yes]

1327 Justification: The abstract and introduction state the paper’s main claims—structured  
1328 kinematic-to-visual control, hierarchical routing, efficient generation, and the KASA  
1329 benchmark.

1330 Guidelines:

- 1331 • The answer [N/A] means that the abstract and introduction do not include the claims  
1332 made in the paper.
- 1333 • The abstract and/or introduction should clearly state the claims made, including the  
1334 contributions made in the paper and important assumptions and limitations. A [No]  
1335 or [N/A] answer to this question will not be perceived well by the reviewers.
- 1336 • The claims made should match theoretical and experimental results, and reflect how  
1337 much the results can be expected to generalize to other settings.
- 1338 • It is fine to include aspirational goals as motivation as long as it is clear that these goals  
1339 are not attained by the paper.

### 1340 2. Limitations

1341 Question: Does the paper discuss the limitations of the work performed by the authors?

1342 Answer: [Yes]

1343 Justification: [Yes]

1344 Guidelines: The Appendix C includes a dedicated discussion of limitations and broader  
1345 impact. It explicitly scopes the method to articulated robotic surgical video generation with  
1346 curated 7-DoF supervision and notes that broader procedures and instruments may require  
1347 additional adaptation.

- 1348 • The answer [N/A] means that the paper has no limitation while the answer [No] means  
1349 that the paper has limitations, but those are not discussed in the paper.
- 1350 • The authors are encouraged to create a separate “Limitations” section in their paper.
- 1351 • The paper should point out any strong assumptions and how robust the results are to  
1352 violations of these assumptions (e.g., independence assumptions, noiseless settings,  
1353 model well-specification, asymptotic approximations only holding locally). The  
1354 authors should reflect on how these assumptions might be violated in practice and what  
1355 the implications would be.
- 1356 • The authors should reflect on the scope of the claims made, e.g., if the approach was  
1357 only tested on a few datasets or with a few runs. In general, empirical results often  
1358 depend on implicit assumptions, which should be articulated.
- 1359 • The authors should reflect on the factors that influence the performance of the approach.  
1360 For example, a facial recognition algorithm may perform poorly when image resolution  
1361 is low or images are taken in low lighting. Or a speech-to-text system might not be  
1362 used reliably to provide closed captions for online lectures because it fails to handle  
1363 technical jargon.
- 1364 • The authors should discuss the computational efficiency of the proposed algorithms  
1365 and how they scale with dataset size.
- 1366 • If applicable, the authors should discuss possible limitations of their approach to  
1367 address problems of privacy and fairness.
- 1368 • While the authors might fear that complete honesty about limitations might be used  
1369 by reviewers as grounds for rejection, a worse outcome might be that reviewers  
1370 discover limitations that aren’t acknowledged in the paper. The authors should use  
1371 their best judgment and recognize that individual actions in favor of transparency play  
1372 an important role in developing norms that preserve the integrity of the community.  
1373 Reviewers will be specifically instructed to not penalize honesty concerning limitations.

### 1374 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: [Yes]

Justification: The paper's theoretical results are provided in Appendix A, and for these results, the assumptions and proofs are explicitly stated.

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.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

#### 4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [Yes]

Justification: The paper specifies the problem setup, model components, training/evaluation protocol, metrics, dataset construction pipeline, and ablations in Secs. 2–4 and the appendix, which together provide the information needed to reproduce the reported main results.

Guidelines:

- The answer [N/A] means that the paper does not include experiments.
- If the paper includes experiments, a [No] answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
  - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
  - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
  - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
  - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in

1429 some way (e.g., to registered users), but it should be possible for other researchers  
1430 to have some path to reproducing or verifying the results.

## 1431 5. Open access to data and code

1432 Question: Does the paper provide open access to the data and code, with sufficient  
1433 instructions to faithfully reproduce the main experimental results, as described in  
1434 supplemental material?

1435 Answer: [Yes]

1436 Justification: We provide anonymized access to the code and supporting artifacts used in this  
1437 work for peer review. The core codebase is included in the supplementary material, while  
1438 larger checkpoints and data access instructions are shared through anonymized external  
1439 links that preserve double-blind reviewing. Any public de-anonymized release will follow  
1440 the licenses and usage restrictions of the underlying assets.

1441 Guidelines:

- 1442 • The answer [N/A] means that paper does not include experiments requiring code.
- 1443 • Please see the NeurIPS code and data submission guidelines (<https://neurips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- 1444 • While we encourage the release of code and data, we understand that this might not  
1445 be possible, so [No] is an acceptable answer. Papers cannot be rejected simply for not  
1446 including code, unless this is central to the contribution (e.g., for a new open-source  
1447 benchmark).
- 1448 • The instructions should contain the exact command and environment needed to run to  
1449 reproduce the results. See the NeurIPS code and data submission guidelines (<https://neurips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- 1450 • The authors should provide instructions on data access and preparation, including how  
1451 to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- 1452 • The authors should provide scripts to reproduce all experimental results for the new  
1453 proposed method and baselines. If only a subset of experiments are reproducible, they  
1454 should state which ones are omitted from the script and why.
- 1455 • At submission time, to preserve anonymity, the authors should release anonymized  
1456 versions (if applicable).
- 1457 • Providing as much information as possible in supplemental material (appended to the  
1458 paper) is recommended, but including URLs to data and code is permitted.

## 1461 6. Experimental setting/details

1462 Question: Does the paper specify all the training and test details (e.g., data splits, hyper-  
1463 parameters, how they were chosen, type of optimizer) necessary to understand the results?

1464 Answer: [Yes]

1465 Justification: The paper specifies the main training and evaluation details needed to  
1466 understand the results, including the data split, optimizer, learning rate, weight decay, batch  
1467 size, default resolution, number of denoising steps, baseline setup, and evaluation metrics  
1468 in Sec. 4.1, with additional implementation details deferred to the appendix.

1469 Guidelines:

- 1470 • The answer [N/A] means that the paper does not include experiments.
- 1471 • The experimental setting should be presented in the core of the paper to a level of  
1472 detail that is necessary to appreciate the results and make sense of them.
- 1473 • The full details can be provided either with the code, in appendix, or as supplemental  
1474 material.

## 1475 7. Experiment statistical significance

1476 Question: Does the paper report error bars suitably and correctly defined or other appropriate  
1477 information about the statistical significance of the experiments?

1478 Answer: [Yes]

Justification: The paper reports additional statistical information in the experimental section and Appendix F, including uncertainty estimates, definitions of the averaging protocol, and clarification of how the reported confidence intervals or summary statistics are computed for the supporting analyses.

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, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g., negative error rates).
- If error bars are reported in tables or plots, the authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

## 8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: The paper describes the compute resources in the experimental section and Appendix D, including the hardware used for training and evaluation, as well as additional resource details such as the relevant worker configuration and runtime information needed for reproduction.

Guidelines:

- The answer [N/A] means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

## 9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: The research described in this submission conforms with the NeurIPS Code of Ethics. The paper preserves anonymity and uses standard research and evaluation practices for academic study.

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 that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

## 10. Broader impacts

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

Answer: [Yes]

Justification: The Appendix C includes a broader-impact discussion covering both potential positive impacts.

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 their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate Deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

## 11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pre-trained language models, image generators, or scraped datasets)?

Answer: [No]

Justification: The paper does not describe concrete safeguards for releasing data or models, such as gated access, usage restrictions, or release-time filtering.

Guidelines:

- The answer [N/A] means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

## 12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: The Appendix C includes an asset-license note that credits the external datasets, models, and software components used in this work and explicitly states their license or usage conditions, together with the commitment to respect the corresponding terms of use and release constraints where applicable.

Guidelines:

- The answer [N/A] means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, `paperswithcode.com/datasets` has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

### 13. New assets

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

Answer: [No]

Justification: Although this work introduces new derived benchmark annotations and model artifacts for evaluation, we do not claim a formal public release of new assets in the current submission. Materials shared for peer review are anonymized and provided only for evaluation purposes. Any future release will be subject to the licenses and usage constraints of the underlying resources.

Guidelines:

- The answer [N/A] means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

### 14. Crowdsourcing and research with human subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [N/A]

Justification: The paper does not report crowdsourcing experiments or human-subject studies involving participants who were instructed, compensated, or evaluated. The human-in-the-loop annotation used for dataset curation is not presented as a crowdsourcing study in the submission.

Guidelines:

- The answer [N/A] means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

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

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [N/A]

Justification: The paper does not report crowdsourcing studies or human-subject experiments aimed at collecting responses from participants. The human-in-the-loop annotation used for dataset curation is part of internal data annotation rather than a crowdsourcing or behavioral study.

Guidelines:

- The answer [N/A] means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

#### 16. Declaration of LLM usage

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does *not* impact the core methodology, scientific rigor, or originality of the research, declaration is not required.

Answer: [N/A]

Justification: The core methodology in this research does not use LLMs as an important, original, or non-standard component, so declaration is not required under the checklist policy.

Guidelines:

- The answer [N/A] means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy in the NeurIPS handbook for what should or should not be described.