

Prompt2Effect: Training-Free Image-to-Video Model Specialization via LoRA Generation

Xiaomeng Yang^{1,2*}, Yanyu Li¹, Gordon Guocheng Qian¹, Ivan Skorokhodov¹, Viacheslav Ivanov¹, Avalon Vinella¹, Xuan Zhang², Yanzhi Wang², Sergey Tulyakov¹, and Anil Kag¹

¹Northeastern University, ²Snap Inc.

<https://xiaomeng-yang.github.io/Prompt2Effect/>

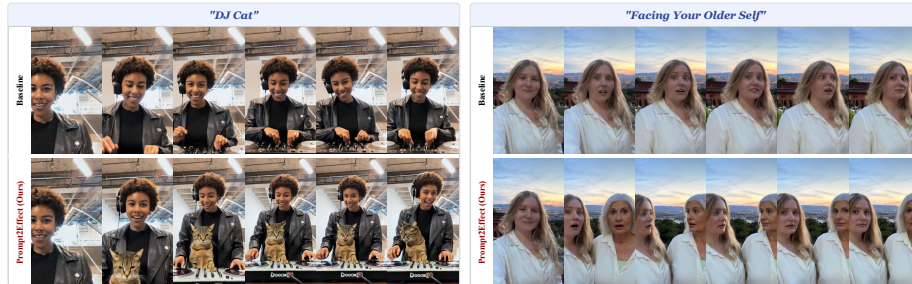


Fig. 1: Prompt2Effect enables zero-shot LoRA synthesis through a single forward pass for effect-injected image-to-video diffusion generation.

Abstract. While personalizing Image-to-Video (I2V) diffusion models with specific visual effects is increasingly demanded for high-end generation, current practice requires training a separate Low-Rank Adaptation (LoRA) module for each effect, incurring substantial data curation and iterative optimization costs that hinder interactive control. We present *Prompt2Effect*, a weight-driven hypernetwork that amortizes per-effect training by directly synthesizing effect-specific LoRA weights in a single forward pass. Unlike prior hypernetworks that regress adapter weights purely from semantics, *Prompt2Effect* is explicitly conditioned on the frozen base model weights, grounding prediction in the structural geometry of each layer. Furthermore, instead of predicting raw LoRA matrices, we introduce an SVD-canonicalized parameterization that resolves factorization ambiguity and stabilizes large-scale synthesis. Extensive experiments demonstrate that Prompt2Effect achieves on-par or superior video quality and effect alignment compared to conventional LoRA fine-tuning, while reducing the computational cost from 56 GPU training hours to 3.3 seconds of hypernetwork inference. When used as initialization for subsequent fine-tuning, our predicted weights further improve final performance and accelerate optimization by approximately 10 $\times$.

Keywords: Video Diffusion, Video Effect, LoRA, Hypernetwork, Personalization

* Work done during internship at Snap Inc.

1 Introduction

Recent Image-to-Video (I2V) diffusion models [26] have achieved remarkable success in synthesizing high-fidelity video from an input image and text guidance. Beyond general semantic control, there is a burgeoning demand for effect-level personalization: the ability to render an input image into specific visual styles or motion effects with fine-grained, immediate control. In practice, these complex dynamic effects are often distilled from dozens of unstructured reference videos, which are far too heavy to use as direct conditions during inference. Consequently, the standard approach for such personalization [2, 19] involves training a separate Low-Rank Adaptation (LoRA) [11] module on a new video dataset consisting of the target effects. However, per-effect data curation and LoRA training are inherently iterative and computationally expensive, often requiring thousands of optimization steps, which create a significant bottleneck for interactive creative workflows. This overhead is particularly acute for I2V models, where the increased model parameters and temporal complexity make effect authoring slow and difficult to scale [21].

To bypass the costs of gradient-based optimization and enable *amortized per-effect LoRA training*, weight-generation networks such as hypernetworks [6, 8] have emerged as a compelling alternative. These models directly predict model parameters in a single forward pass, conditioned on a target task. This raises a natural question: can we leverage this paradigm to synthesize effect-specific LoRA weights for I2V generation models directly from a textual effect description, without per-effect optimization? For image generation, hypernetwork approaches such as HyperDreamBooth [25] have already demonstrated the potential of training-free personalization by generating adapter weights conditioned on an input face image.

However, directly extending these techniques to I2V diffusion remains highly non-trivial. First, existing methods primarily focus on localized *subject/identity* injection in static images (typically using UNet architectures [23]), whereas our goal is to control global dynamic *visual effects and motion* in video. Second, modern I2V models (such as Diffusion Transformers [26]) involve high-dimensional temporal latent dynamics and require consistent adaptation across an exponentially larger parameter space. In practice, existing weight generation designs face a fundamental scaling challenge: they struggle to regress the large and structured LoRA parameter space of I2V models and can fail to converge to usable weight predictors under this setting (Fig. 6).

To address these challenges, we propose **Prompt2Effect**, a **weight-driven hypernetwork for SVD-canonicalized LoRA prediction** in I2V diffusion models. Unlike prior hypernetwork approaches that rely solely on input conditions such as effect prompts or identity images, our hypernetwork is explicitly conditioned on the frozen base weights of the target model. This weight-driven formulation is motivated by the nature of LoRA itself: a LoRA update ΔW is not an absolute parameter, but a structured adaptation defined with respect to the base model weight matrix W_0 . When the predictor has no access to W_0 , it must implicitly infer the geometry of the base layer purely from semantics, which

becomes increasingly ill-posed as model dimensionality grows. By conditioning on W_0 , the hypernetwork is grounded in the native input-output coordinates and correlation structure of each layer, substantially reducing the ambiguity of large-scale weight synthesis.

Furthermore, existing methods typically regress LoRA matrices in their raw factorized form. However, the decomposition $\Delta W = BA$ is inherently non-identifiable: for any invertible matrix R , the pairs $(BR, R^{-1}A)$ produce the same update. This induces a highly redundant solution space and leads to unstable optimization when predicting many layer-wise adapters simultaneously. To resolve this ambiguity, we predict the SVD-canonicalized form of the LoRA update. By decomposing $\Delta W = USV^\top$ and predicting the canonical factors $B^* = US^{1/2}$ and $A^* = S^{1/2}V^\top$, we enforce an orthogonal, energy-ordered parameterization with a unique representation up to sign. This significantly constrains the prediction space, improves convergence, and stabilizes large-scale LoRA weight synthesis. Together, these two design principles, weight-driven conditioning and SVD-canonicalized prediction, make Prompt2Effect possible to stably synthesize high-dimensional LoRA updates for modern I2V models in a single forward pass (see ablation in Fig. 6), enabling accurate training-free effect control (Fig. 1).

Our key **contributions** are summarized as follows:

- **Weight-Driven LoRA Synthesis for I2V Diffusion.** We introduce a weight-conditioned formulation for LoRA prediction that explicitly leverages frozen base weights to ground adaptation in layer geometry, enabling stable scaling to high-dimensional I2V diffusion models.
- **SVD-Canonicalized LoRA Prediction.** We identify the non-identifiability of weight regression as a fundamental obstacle to stable training. To address this, we propose predicting SVD-canonicalized LoRA factors, which impose an orthogonal and energy-ordered parameterization, reduce redundancy in the solution space, and significantly improve convergence stability.
- **Scalable Training-Free Effect Control.** Our framework amortizes the training process, enabling instantaneous effect generation without per-effect optimization by synthesizing large collections of layer-wise LoRA updates in a single forward pass. Prompt2Effect achieves on-par or superior video quality and effect alignment compared to conventional LoRA fine-tuning, while reducing the computational cost from 56 GPU hours of training to 3.3 seconds of inference time per effect.
- **Fast Adaptation and Composable Control.** When maximum fidelity is required, our predicted weights serve as strong initialization, substantially accelerating conventional LoRA fine-tuning by $10\times$. Moreover, Prompt2Effect supports interpolation and composition directly in weight space, enabling flexible and continuous effect manipulation.

2 Related Work

Controllable Video Generation. While controllable video generation includes *structure-controlled* methods leveraging pixel-aligned conditions, our work fo-

cuses on *semantic/effect* control to enforce non-aligned target behaviors like concepts, styles, or VFX without pixel-wise correspondence.

Per-effect / per-video adaptation. A dominant approach to semantic control is *condition-specific overfitting*, adapting a pretrained T2V model to each new effect via lightweight finetuning like DreamBooth [24] or LoRA [11]. For instance, Set-and-Sequence [2] isolates identity and motion through two-stage LoRA finetuning, but still requires per-concept optimization. Similarly, models like VFX Creator [16], Omni-Effects [19], and EasyVFX [18] train specialized adapters or effect experts, yet they rely on predefined categories or effect-specific training that cannot generalize to unseen effects without further adaptation.

Task-specific design and inference-time control. Another line of work improves controllability through task-specific modules or inference-time mechanisms. Tune-A-Video [27] preserves structure while injecting new semantics via one-shot video adaptation with inversion-based initialization. For reference-driven stylization, StyleMaster [30] employs adapter-like cross-attention to condition generation on style images, similar to IP-Adapter [29]. Recent works also steer generation without explicit parameter updates by augmenting inputs with additional modalities or latents [1, 4, 13]. While ControlNet-style branches and image-conditioned adapters excel with compact conditions like masks or static images, they struggle with our target application. Global dynamic effects are unstructured and typically require dozens of reference videos, making them too computationally heavy for inference-time conditioning. In contrast, our goal is *zero-shot* effect control at inference time by *synthesizing* plug-and-play LoRA parameters from an effect prompt in a single forward pass.

Weight Prediction Networks. Hypernetworks [8] that predict adaptation weights have been explored as an efficient alternative to per-instance finetuning. HyperDreamBooth [25] predicts personalized diffusion-model weights for fast identity adaptation, demonstrating that weight generation can substantially reduce test-time cost. Subsequent works [6, 17] improved parameter efficiency by predicting LoRA [11] adapters rather than full weights. Notably, LoFA [10] serves as recent prior art in weight-driven LoRA prediction. However, LoFA targets compact conditions (e.g., style, pose reference images) and directly regresses ambiguous raw LoRA factors, which can lead to unstable optimization. In contrast, our framework synthesizes LoRAs for unstructured, dynamic effects distilled from dozens of videos, and we propose an SVD-canonicalized prediction to resolve factorization ambiguity and stabilize large-scale weight synthesis.

More recently, generative approaches have been proposed to synthesize LoRA parameters, including diffusion-based LoRA generation for LLMs and image diffusion models [14, 15, 28]. Despite this progress, existing methods primarily focus on LLMs or localized *subject/identity* injection in static images (typically UNet architectures). Extending *prompt-conditioned LoRA synthesis* to large-scale *video* diffusion backbones (such as Diffusion Transformers [26]) remains underexplored. This transition is highly non-trivial due to the explosion in parameter dimensionality and the need to adapt temporally structured representations for global visual effects.

Our Prompt2Effect addresses this challenge by introducing a hypernetwork that predicts SVD-canonicalized LoRA weights from effect prompts. While recent works like PiSSA [20] and Make-a-LoRA [7] have shown that SVD-based LoRA initialization improves gradient-based training convergence, we uniquely leverage SVD to canonicalize the regression target space for hypernetwork training. This ensures a structurally aligned and stable optimization landscape for the hypernetwork, allowing it to reliably predict complex video effect parameters.

3 Method

In this section, we introduce Prompt2Effect, a framework for training-free LoRA synthesis at inference time for controllable video effects. Our overall pipeline consists of two main phases: (1) a one-time pretraining stage, where a hypernetwork is trained to regress LoRA weights from effect prompts, and (2) a zero-shot inference stage, where the hypernetwork directly predicts effect-specific LoRA weights from a given effect prompt through a single forward pass. We also introduce an optional lightweight testing-time optimization that uses the predicted weights as an initialization to improve the zero-shot quality while being 10× faster than standard LoRA training from scratch.

3.1 Preliminaries

Video Diffusion Models. Video diffusion models generate video clips by iteratively denoising a sequence of latent variables. Given a clean video latent z_0 , a forward diffusion process adds Gaussian noise to produce a noisy latent z_t at timestep t . The model, typically parameterized by a 3D-UNet [23] or spatial-temporal Transformer [21], is trained to reverse the process and generate data z_0 with conditions such as textual prompts (c).

Low-Rank Adaptation (LoRA). LoRA [11, 24] is a parameter-efficient fine-tuning technique that adapts a frozen pretrained model to new concepts or domains. Instead of updating the original weight matrix $W \in \mathbb{R}^{d \times k}$, LoRA optimizes a low-rank residual ΔW :

$$W' = W + \Delta W = W + BA$$

where $B \in \mathbb{R}^{d \times r}$ and $A \in \mathbb{R}^{r \times k}$ are trainable low-rank matrices, and the rank $r \ll \min(d, k)$.

Hypernetworks for LoRA Prediction. A hypernetwork H_ϕ parameterized by ϕ is a meta-model designed to generate the weights of another neural network. In our context, we aim to learn a mapping function that takes an effect prompt c_{eff} as input and directly predicts a collection of LoRA weight updates for a pretrained video diffusion model:

$$\Delta\theta = H_\phi(c_{\text{eff}}), \quad \Delta\theta = \{\Delta W_i\}_{i=1}^L \quad (1)$$

where $\Delta\theta$ represents the set of low-rank adaptations applied to the L selected layers of the base model with pretrained parameters θ_0 .

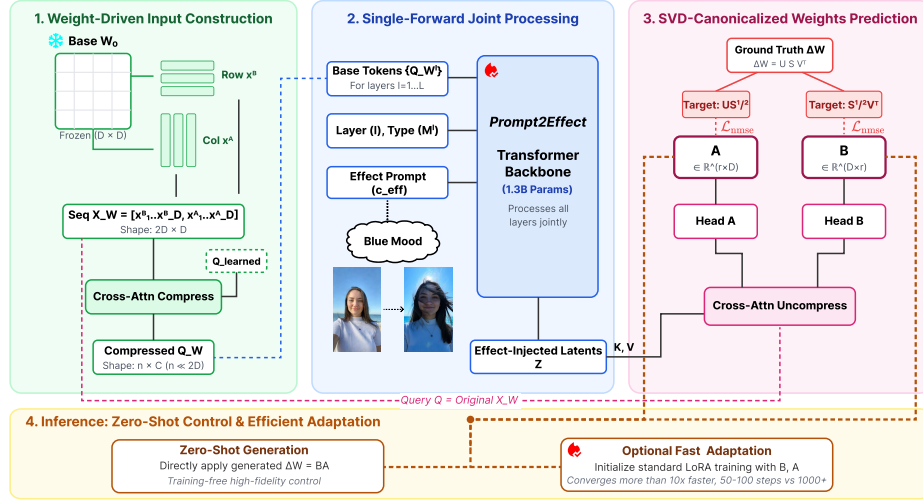


Fig. 2: Overview of the Prompt2Effect framework. The pipeline consists of three stages: (1) **Weight-Driven Input Construction**, where the frozen base weights W_0 are decomposed and compressed into compact weight tokens; (2) **Hypernetwork Backbone**, which jointly processes the compressed weight tokens, effect prompt, and layer metadata to produce effect-injected latents; and (3) **SVD-Canonicalized Weight Prediction**, which projects the latents back to the original parameter space to predict SVD-canonicalized LoRA matrices. The predicted weights can be directly applied for zero-shot generation, or used as initialization for fast adaptation on more complex effects to further improve performance with more than 10 $\times$ training reduction compared to LoRA training from scratch. Here we use D for simplicity.

The adapted model parameters are obtained via

$$\Theta = \Theta_0 + \Delta\Theta. \quad (2)$$

Under the LoRA parameterization, each layer update is decomposed as

$$\Delta W^l = B^l A^l, \quad A^l \in \mathbb{R}^{r \times d_{in}^l}, \quad B^l \in \mathbb{R}^{d_{out}^l \times r}, \quad l = 1, \dots, L, \quad (3)$$

3.2 Weight-Driven SVD LoRA Prediction

Existing hypernetworks [25,28] typically generate target network weights directly from abstract semantic embeddings, such as noise vectors or text features. These approaches are *purely semantic-driven*: the predictor has no explicit access to the base model weights it is meant to adapt. As a result, weight generation lacks structural grounding in the underlying network, which empirically leads to suboptimal performance when predicting high-dimensional updates (Sec. 4.4).

Moreover, prior methods generally regress weights in their raw parameterization. Directly predicting unconstrained weight matrices introduces two additional challenges: (1) the solution space is highly non-identifiable due to equivalent factorizations of $\Delta W = BA$; thus, for any invertible matrix R , the pairs

($BR, R^{-1}A$) produce the same update, and (2) the absence of canonical structure makes optimization less stable, especially for large networks.

To address these limitations, we propose a *weight-driven SVD LoRA prediction* framework. First, instead of relying solely on semantic embeddings, we condition the hypernetwork explicitly on the frozen base weights W_0 , providing strong structural priors for predicting ΔW . Second, rather than regressing raw LoRA matrices, we predict their SVD-canonicalized form, which removes factorization ambiguity and stabilizes training.

Weight-Driven Input Formulation. We formulate a novel *weight-driven* prediction mechanism, motivated by the fact that a LoRA update is an adaptation defined with respect to the original weight matrix. Therefore, we propose to use the pre-trained network weights directly as an extra input along with the textual prompt to the hypernetwork to predict the desired LoRA update.

Specifically, for a target layer l with frozen baseline weights $W_0^l \in \mathbb{R}^{d_{\text{out}}^l \times d_{\text{in}}^l}$, we construct a structured input sequence by slicing W_0^l into *row tokens* and *column tokens*. The rows form B tokens $x_i^{B,l} = W_0^l[i, :] \in \mathbb{R}^{d_{\text{in}}^l}$, and the columns form A tokens $x_j^{A,l} = W_0^l[:, j] \in \mathbb{R}^{d_{\text{out}}^l}$. Since row and column tokens have different raw dimensionalities, we project them into a shared token space of width d :

$$\tilde{x}_i^{B,l} = x_i^{B,l} P_B^l, \quad P_B^l \in \mathbb{R}^{d_{\text{in}}^l \times d}, \quad \tilde{x}_j^{A,l} = x_j^{A,l} P_A^l, \quad P_A^l \in \mathbb{R}^{d_{\text{out}}^l \times d}. \quad (4)$$

We then concatenate them into a row/column token sequence:

$$X_W^l = [\tilde{x}_1^{B,l}, \dots, \tilde{x}_{d_{\text{out}}^l}^{B,l}, \tilde{x}_1^{A,l}, \dots, \tilde{x}_{d_{\text{in}}^l}^{A,l}] \in \mathbb{R}^{(d_{\text{out}}^l + d_{\text{in}}^l) \times d}. \quad (5)$$

This construction is inspired by classical CUR decompositions, which approximate a matrix using its actual columns and rows [5, 9]. We use the rows and columns of W_0^l directly as tokens to expose the hypernetwork to the layer’s native input/output feature coordinates and cross-feature correlations, instead of computing a CUR factorization here.

Since LoRA parameterizes adaptation as a low-rank update $\Delta W = BA$ defined with respect to the frozen base weight W_0 , conditioning the predictor on W_0 provides a strong structural prior that is absent when mapping solely from text or noise embeddings. Empirically, we observe a clear compressibility gap: W_0 concentrates most of its spectral energy in the top singular components, while ΔW distributes its energy across a substantially

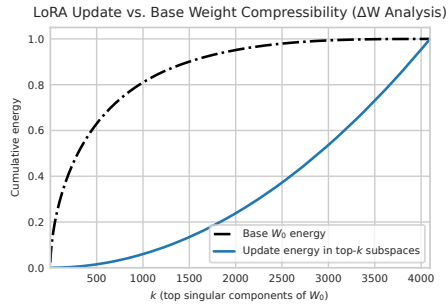


Fig. 3: Compressibility (measured by cumulative SVD energy) gap between the frozen base weights (W_0 in black) and LoRA updates (ΔW in blue). The plotted curves aggregate all LoRA layer pairs across the training set. While W_0 is highly compressible, ΔW spreads its energy across a much broader set of base singular directions, motivating full-rank base weight tokenization for accurate LoRA prediction.

distributes its energy across a substantially

broader set of singular directions (Fig. 3). This observation motivates the use of *full-rank* weight tokenization of the base model as input to the hypernetwork. Aggressive compression of W_0 would discard singular directions that carry critical information for accurately predicting ΔW . Refer to Sec. 4.4 for the ablation.

Weight Token Compression. Although full-rank weight tokenization preserves critical structural directions, directly feeding the entire sequence X_W^l into a hypernetwork is computationally prohibitive due to its large token length ($d_{\text{out}}^l + d_{\text{in}}^l$). Importantly, our goal is not to perform spectral truncation or low-rank approximation of W_0 , which would discard singular directions that are informative for predicting ΔW . Instead, we introduce a learnable token aggregation mechanism that reduces sequence length while still allowing the hypernetwork to attend to all original weight directions. Specifically, we apply a cross-attention operation using learnable queries to aggregate the high-dimensional weight tokens:

$$Q_W^l = \text{CrossAttn}(Q_{\text{learned}}, X_W^l, X_W^l), \quad (6)$$

where $Q_{\text{learned}} \in \mathbb{R}^{n \times d}$ is a learnable query embedding, and the compressed weight token sequence $Q_W^l \in \mathbb{R}^{n \times d}$ has a much shorter token length $n \ll d_{\text{out}}^l + d_{\text{in}}^l$.

SVD-Canonicalized LoRA Prediction. Instead of directly predicting the full unconstrained LoRA matrices A^l and B^l that are non-identifiable, we propose predicting the *SVD-decomposed* LoRA weights. We pre-extract ground-truth LoRAs across diverse dynamic concepts and apply Singular Value Decomposition (SVD) to canonicalize the matrices A^l and B^l :

$$\Delta W^l = U^l S^l (V^l)^T \Rightarrow B^{*,l} = U^l (S^l)^{1/2}, A^{*,l} = (S^l)^{1/2} (V^l)^T \quad (7)$$

The hypernetwork is then trained to minimize the reconstruction loss against these SVD-canonicalized matrices. Because SVD enforces an orthogonal, energy-compacted basis, *we empirically find that learning to predict these canonicalized matrices converges significantly faster and produces more stable generation than predicting raw, non-canonicalized weights.* See Sec. 4.4 for the ablation.

3.3 Prompt2Effect Architecture

Designing a hypernetwork that generates high-dimensional adaptation weights for a large-scale video diffusion model is challenging due to the scale and layer-wise heterogeneity of the backbone. Prompt2Effect employs a Transformer-based hypernetwork H_ϕ that predicts all LoRA weights in a single forward pass, conditioned on the effect prompt and layer-specific metadata. H_ϕ consists of two parts: (i) a Transformer backbone $H_\mathcal{E}$ that fuses weighttext and weight tokens, we inject semantics through cross-attention so that structurally-aware weight tokens attend to the effect semantics, and (ii) an attention-based weight prediction head $H_\mathcal{D}$ that uncompresses layer tokens and outputs LoRA factors.

Transformer backbone $H_\mathcal{E}$. $H_\mathcal{E}$ models the interaction between the structural priors of the base model and the semantic guidance of the effect prompt. For each adapted layer $l \in \{1, \dots, L\}$, the compressed weight tokens Q_W^l serve as

the primary input sequence. To provide spatial and structural awareness, we add learned layer-index embeddings e_l and module-type embeddings e_M^l directly to the base tokens:

$$Z^l = Q_W^l + e_l + e_M^l \quad (8)$$

where $e_l = E_{\text{layer}}(l) \in \mathbb{R}^d$ and $e_M^l = E_{\text{mod}}(M^l) \in \mathbb{R}^d$. The semantic effect prompt embedding c_{eff} is encoded via a frozen text encoder to obtain a sequence of text tokens $T_{\text{eff}} = E_{\text{text}}(c_{\text{eff}}) \in \mathbb{R}^{n_t \times d}$. Rather than concatenating text and weight tokens, we inject semantics through cross-attention so that structurally-aware weight tokens attend to the prompt.

$H_{\mathcal{E}}$ takes the layer-wise conditions as input and predicts the effect-injected latents used to generate LoRA weights:

$$X_{\mathcal{E}} = H_{\mathcal{E}}(Z, T_{\text{eff}}) \in \mathbb{R}^{L \times n \times d}. \quad (9)$$

Here, Z is the stacking of $\{Z^l\}_{l=1}^L$, L denotes the number of adapted layers, n is the number of latent tokens generated per layer, and d is the hidden width of the hypernetwork. Each slice $X_{\mathcal{E}}^l \in \mathbb{R}^{n \times d}$ corresponds to the effect-injected latent associated with layer l , which is subsequently used to predict its LoRA weights.

LoRA Weight Prediction Head $H_{\mathcal{D}}$. The goal of $H_{\mathcal{D}}$ to predict *full uncompressed weights of SVD-Canonicalized LoRA matrix* A^* and B^* from the compressed latent $X_{\mathcal{E}}$. $H_{\mathcal{D}}$ consists a cross-attention uncompression operation:

$$X_{\mathcal{D}}^l = \text{CrossAttn}(Q = X_W^l, K = X_{\mathcal{E}}^l, V = X_{\mathcal{E}}^l) \in \mathbb{R}^{(d_{\text{out}}^l + d_{\text{in}}^l) \times d} \quad (10)$$

This operation projects the learned effect dynamics directly back into the native structural dimensionality of the base weights by using the uncompressed original weights X_W^l as queries. We then split tokens into row-part and column-part:

$$X_{\mathcal{D}}^{l,(B)} = X_{\mathcal{D}}^l[:, d_{\text{out}}^l], \quad X_{\mathcal{D}}^{l,(A)} = X_{\mathcal{D}}^l[d_{\text{out}}^l :]. \quad (11)$$

Finally, $H_{\mathcal{D}}$ employs two linear projection heads to predict the canonical LoRA matrices $\hat{B}$ and $\hat{A}$:

$$\hat{B}^l = \text{Head}_B(X_{\mathcal{D}}^{l,(B)}) \in \mathbb{R}^{d_{\text{out}}^l \times r}, \quad \hat{A}^l = \left(\text{Head}_A(X_{\mathcal{D}}^{l,(A)}) \right)^\top \in \mathbb{R}^{r \times d_{\text{in}}^l}, \quad (12)$$

By processing all layer tokens jointly, $H_{\mathcal{D}}$ produces $\{(\hat{A}^l, \hat{B}^l)\}_{l=1}^L$ in a single forward pass, enabling efficient LoRA synthesis without per-layer optimization.

Training Objective. Given our predictions $\{(\hat{A}^l, \hat{B}^l)\}_{l=1}^L$ and ground-truth SVD-canonicalized LoRA weights $\{(A^{*,l}, B^{*,l})\}_{l=1}^L$, we train H_{ϕ} using a *normalized MSE (NMSE)* loss that measures the *relative* Frobenius error for both A and B :

$$\mathcal{L}_{\text{LoRA}} = \frac{1}{2L} \sum_{l=1}^L \left(\frac{1}{N} \sum_{b=1}^N \frac{\|\hat{A}_b^l - A_b^{*,l}\|_F^2}{\|A_b^{*,l}\|_F^2 + \epsilon} + \frac{1}{N} \sum_{b=1}^N \frac{\|\hat{B}_b^l - B_b^{*,l}\|_F^2}{\|B_b^{*,l}\|_F^2 + \epsilon} \right), \quad (13)$$

where b indexes the batch (N samples) and ϵ is a small constant for numerical stability. This relative formulation reduces sensitivity to scale variations across layers and effects, leading to more stable hypernetwork training.

At inference time, we directly apply the predicted LoRA parameters to the frozen video diffusion backbone to enable zero-shot effect generation without per-effect optimization.

4 Experiments

4.1 Implementation Details

Hypernetwork. Unless otherwise specified, we use a 1.3B-parameter Transformer hypernetwork (14 layers, hidden width $d = 2048$) and set the target LoRA rank to $r = 128$. For each adapted layer, we compress weight tokens to $n = 256$ latent tokens. It takes 3.3 seconds for our hypernetwork to generate the entire ΔW for the base model on A100. We primarily conduct experiments on our proprietary image-to-video (I2V) diffusion model. To validate that Prompt2Effect is model-agnostic, we additionally reproduce our core findings on the public WAN video model (see Appendix for details).

Baselines. We compare against (i) the baseline I2V controlled using text prompt, (ii) a HyperDreamBooth-style hypernetwork baseline [25] adapted to video LoRA prediction, and (iii) a fully optimized LoRA trained per effect from scratch, which serves as an upper bound. To ensure fairness, all hypernetworks predict LoRA updates for the same set of target layers and use the same training/evaluation data and inference settings.

Data. We compiled 75 effects prompts, each accompanied by approximately 50 manually selected, captioned videos. Roughly 40% of this data was sourced from public video models, while the remainder was generated and collected internally. Five of the 75 prompts were reserved as a held-out evaluation set. Further details are provided in the Appendix.

Training. We train the hypernetwork using the AdamW optimizer across 4 nodes $\times$ 8 NVIDIA A100 GPUs for 4k epochs, with a constant learning rate of $5e^{-5}$ and a 0.01% warmup ratio. For LoRA training, we utilize a single node for 1k steps, employing a cosine learning rate with an initial learning rate of $1e^{-4}$ with a 0.01 warmup ratio and decaying to a minimum learning rate of $1e^{-5}$.

Evaluation. We use 70 in-distribution (IID) effects and 32 out-of-distribution (OOD) effects to evaluate our hypernetwork. For each effect, we generate 76 videos at 512×288 resolution using a fixed prompt suite spanning diverse subject categories (single/multi-person scenes, dynamic backgrounds, objects, and animals). We use identical prompts and random seeds across methods for generating the evaluation videos to ensure fairness.

We measure text-video alignment using CLIP Score [22] computed between predicted and effect prompt, and overall video fidelity by Aesthetic Quality metric from VBench [12]. To directly assess whether the intended effect is executed (e.g., “polar bear spa”), we employ Qwen3-VL [3] as an expert VLM judge. The

Table 1: Quantitative results across evaluation metrics for IID effects. **Bold** denotes the best and underline denotes the second best. The LoRA baseline that requires per-effect optimization is deemphasized.

Method	CLIP	Aesthetic	Motion	VLM Evaluation(↑)			Overall
	Score(↑)	Quality(↑)	Smoothness(↑)	Effect	Context	Temporal	
Baseline (Base I2V model)	24.23	53.36	98.33	15.99	14.21	12.37	42.57
HyperDreambooth [25]	22.34	35.30	-	0.00	0.00	0.00	0.00
Prompt2Effect (Ours)	25.10	56.30	98.54	27.89	15.34	13.70	56.93
LoRA	<u>25.06</u>	<u>57.60</u>	<u>98.44</u>	<u>27.68</u>	<u>15.82</u>	<u>14.07</u>	<u>57.57</u>

Table 2: Quantitative results of zero-shot and fast adaptation on OOD effects.

Method	CLIP	Aesthetic	Motion	VLM Evaluation(↑)			Overall
	Score(↑)	Quality(↑)	Smoothness(↑)	Effect	Context	Temporal	
Baseline (Base I2V model)	23.17	52.26	98.42	15.24	14.78	12.44	42.46
HyperDreambooth [25]	22.36	36.57	-	0.00	0.00	0.00	0.00
Prompt2Effect (Ours)	23.52	53.44	98.58	20.50	15.42	12.95	48.87
Init-50	25.57	56.06	98.47	26.14	15.51	13.23	54.89
Init-100	25.82	56.85	<u>98.53</u>	<u>26.89</u>	<u>15.65</u>	<u>13.79</u>	<u>56.33</u>
LoRA-100	25.11	56.50	98.20	24.66	14.70	12.91	52.28
LoRA (1000 step)	<u>25.73</u>	<u>56.66</u>	98.14	26.94	15.85	13.91	56.70

VLM scores three criteria: (i) *Effect Execution*, (ii) *Context Preservation*, and (iii) *Temporal Quality*, and we report their sum as *Overall*.

4.2 Prompt2Effect Results

We evaluate Prompt2Effect as a *zero-shot* prompt-to-effect system that predicts LoRA weights in a single forward pass, and compare against the base model, HyperDreamBooth [25], and fully optimized LoRA (upper bound).

IID effects. Table 1 reports results on IID effects. Prompt2Effect substantially improves over the base model across all evaluation metrics and performs closely matching the LoRA upper bound. In contrast, the HyperDreamBooth [25] baseline exhibits unstable training when scaled to video generation, leading to noisy outputs, near-zero VLM scores, and the lowest CLIP and aesthetic scores. These results highlight the effectiveness of our weight-driven, SVD-canonicalized hypernetwork design in enabling stable and high-quality weight prediction.

OOD Effects. While Prompt2Effect demonstrates strong zero-shot performance on IID concepts, highly out-of-distribution (OOD) effects may lack certain fine-grained, idiosyncratic details, although ours still outperforms both the base model and HyperDreamBooth, as shown in Table 2. To further enhance OOD performance, we use the predicted weights as an initialization for subsequent LoRA optimization. Initializing from our zero-shot prediction (Init-50 and Init-100) significantly accelerates convergence compared to training from scratch (LoRA-100). Notably, with only 100 optimization steps (Init-100), our method approaches the performance of a fully converged standard LoRA model, which typically requires 1,000 steps. Specifically, Init-100 achieves an overall VLM score of 56.33, comparable to 56.70 for standard LoRA, while yielding a higher CLIP

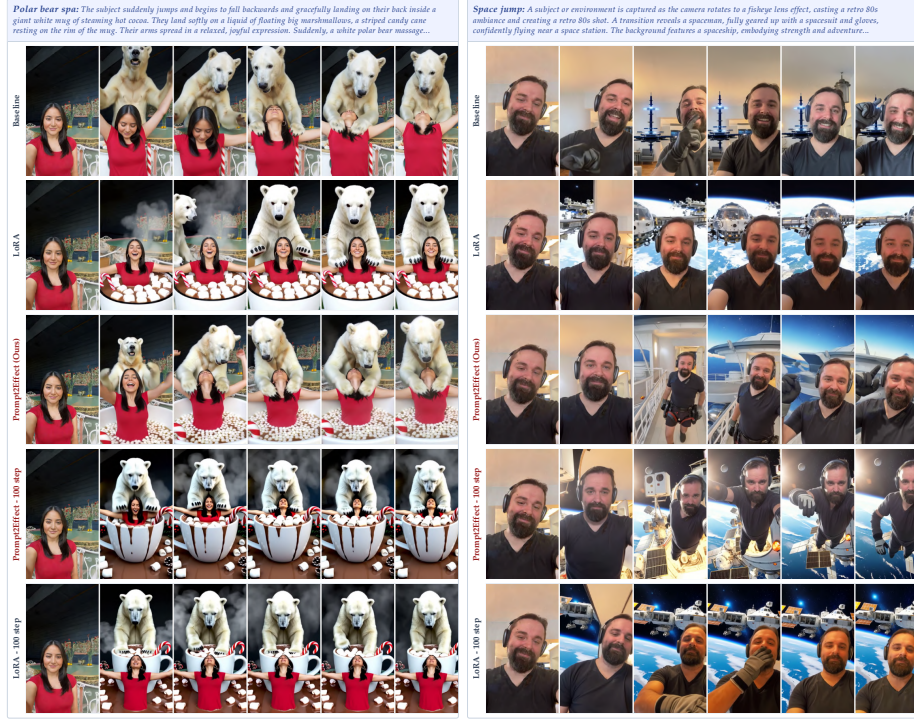


Fig. 4: Out-of-distribution (OOD) qualitative results and fast adaptation. We compare the base model, a fully optimized LoRA, our **Prompt2Effect**’s one-shot predicted weights (zero-shot), and 100-step adaptation initialized from our prediction (Init-100) versus 100-step LoRA training from scratch (LoRA-100). Prompt2Effect’s initialization substantially improves OOD effect within a small optimization budget, reducing $10\times$ training time compared to training LoRA from scratch (100 steps vs. 1000 steps).

score (25.82 vs. 25.73). Qualitative results further show that this 100-step fast adaptation recovers fine-grained OOD textures that may be slightly smoothed in the zero-shot prediction, as illustrated in Fig. 4. This indicates a $10\times$ training speed reduction using our Prompt2Effect as initialization compared to training LoRA from scratch.

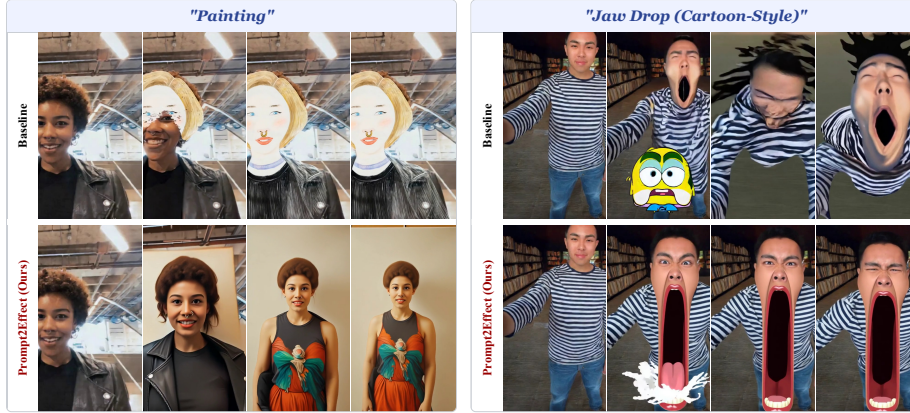
4.3 Public Model Experiments

To validate that the proposed Prompt2Effect is model-agnostic, we additionally reproduce our core findings on the public Wan2.1-I2V-14B-480P¹ model (see Appendix for details). As shown in Tab. 3, Prompt2Effect achieves highly competitive quantitative performance on the Wan backbone, closely tracking the fully optimized, per-effect LoRA baselines in both text-video alignment (CLIP Score)

¹ <https://huggingface.co/Wan-AI/Wan2.1-I2V-14B-480P>

Table 3: Quantitative results on the Wan2.1-I2V-14B-480P model.

Method	CLIP Score(↑)	Aesthetic Quality(↑)	Motion Smoothness(↑)	VLM Evaluation(↑)			
				Effect	Context	Temporal	Overall
Baseline (Wan2.1-I2V)	23.62	58.24	97.63	17.05	14.69	11.35	43.09
Prompt2Effect (Ours)	24.79	62.31	98.02	34.51	15.90	14.08	64.49
LoRA	24.68	62.90	97.88	34.16	16.43	14.57	65.16


Fig. 5: Qualitative effect generation applying Prompt2Effect on the Wan2.1-I2V-14B backbone.

and effect execution (VLM Score), while actively improving motion smoothness over the baseline. Qualitative examples of these synthesized effects demonstrating the robust translation to the Wan architecture are provided in Fig. 5.

4.4 Ablation Study

HyperNetwork Design. We investigate two architectural variants:

- *Per-Block Network:* A modular design where separate, smaller hypernetworks (or heads) predict the weights for specific layers or blocks of the base model. This approach (denoted as “Ensemble” in Table 4) achieves suboptimal performance overall, suggesting that global coordination across layers is important for consistent effect synthesis.
- *Single Big Network:* Our unified, large-scale hypernetwork (1.3B parameters) models the joint distribution of all layer weights simultaneously. By routing features jointly across layers, this architecture achieves superior semantic consistency and higher generation quality.

Weight-Driven Input Formulation. To verify the necessity of providing structural priors, we replace our weight-driven input formulation with standard abstract noise embeddings (denoted as “Noise” in Table 4). The noise-driven model struggles to map semantic text directly to functional weights, yielding

Table 4: Ablation study comparing Ensemble and Noise methods against our approach. Metrics are formatted as IID/OOD.

Method	CLIP Score	Aesthetic Quality	Effect	VLM Evaluation		Overall
				Context	Temporal	
Ensemble	<u>24.93/23.46</u>	55.93/ 53.63	<u>23.21/17.93</u>	<u>14.9/14.74</u>	<u>13.19/12.43</u>	<u>51.29/45.11</u>
Noise	24.41/23.21	53.83/52.21	19.35/15.99	14.33/ <u>14.92</u>	12.45/ <u>12.67</u>	46.14/43.58
Ours	25.10/23.52	56.30/53.44	27.89/20.50	15.34/15.42	13.70/12.95	56.93/48.87

Table 5: Ablation study comparing different compression ratio (n in Eq. (6)). Metrics are formatted as IID/OOD.

Method	CLIP Score	Aesthetic Quality	Effect	VLM Evaluation		Overall
				Context	Temporal	
n=128	24.23/23.17	53.21/52.76	21.10/17.54	15.29/14.84	13.36/12.63	49.75/45.01
n=256	25.10/23.52	<u>56.30/53.44</u>	27.89/20.50	<u>15.34/15.42</u>	<u>13.70/12.95</u>	56.93/48.87
n=512	<u>24.83/23.45</u>	56.98/54.02	<u>27.25/19.46</u>	15.36/15.54	<u>13.61/13.05</u>	<u>56.22/48.03</u>

a steep performance drop in VLM metrics (Overall: 46.14/43.58) and aesthetic quality. Training dynamics in Fig. 6 further verify that, with the weight-driven input formulation, convergence is substantially accelerated. These confirm that exposing the hypernetwork to the principal subspaces of the base weights is crucial for predicting aligned low-rank updates. See Appendix for ablations of other designs other than weight-driven.

SVD-Canonicalized Weight Prediction.

We compare training the hypernetwork to predict raw LoRA weights against predicting SVD-canonicalized targets (A^*, B^*). As shown in Fig. 6, SVD-canonicalization significantly accelerates convergence and improves optimization stability. The NMSE of the reconstructed update ΔW decreases smoothly throughout training, whereas the non-canonicalized variant exhibits pronounced oscillations and converges to a substantially higher final error.

Compression Ratio. Tab. 5 ablates the number of learned queries used to compress weight tokens (n). We find that $n = 256$ offers the best trade-off between compute and fidelity, outperforming $n = 128$ and slightly improving over $n = 512$ on average.

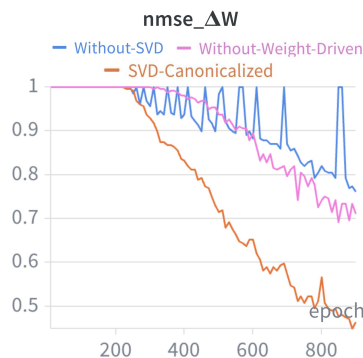
**Fig. 6:** Training curve to predict raw LoRA weights versus SVD-canonicalized prediction versus without weight-driven.



Fig. 7: Zero-shot compositional control by blending predicted LoRA updates, showing that Prompt2Effect learns a LoRA-like space supporting semantic composition.

4.5 Zero-Shot Controllability

A key advantage of Prompt2Effect is that the synthesized adaptations behave like LoRA weights trained conventionally, enabling training-free zero-shot model control. We experiment with semantic composition by interpolating the predicted weights of two distinct effects ($\Delta W = 0.5\Delta W_A + 0.5\Delta W_B$). Visual results in Fig. 7 confirm that our synthesized weights can be seamlessly blended, producing a combined video effect that exhibits characteristics of both concepts.

5 Conclusion

In this work, we introduced Prompt2Effect, a training-free hypernetwork framework that synthesizes LoRA weights for controllable video generation in a single forward pass. By using a weight-driven input representation and predicting SVD-canonicalized adaptations, our method exploits the structural priors of large-scale video diffusion models. Extensive experiments show that Prompt2Effect achieves high-quality zero-shot video effects comparable to optimization-based LoRA, while reducing the computational cost from 56 GPU training hours to 3.3 seconds. Moreover, it provides a strong initialization that accelerates standard LoRA fine-tuning by up to $10\times$. Despite these gains, out-of-distribution performance remains constrained by the diversity and scale of the training data for our hypernetwork. Future work will explore scaling data and training, and extending Prompt2Effect to stronger video backbones to further improve robustness and generalization.

References

1. Abdal, R., Patashnik, O., Deyneka, E., Chen, H., Siarohin, A., Tulyakov, S., Cohen-Or, D., Aberman, K.: Zero-shot dynamic concept personalization with grid-based lora. In: *Proceedings of the SIGGRAPH Asia 2025 Conference Papers*. pp. 1–10 (2025)
2. Abdal, R., Patashnik, O., Skorokhodov, I., Menapace, W., Siarohin, A., Tulyakov, S., Cohen-Or, D., Aberman, K.: Dynamic concepts personalization from single videos. In: *Proceedings of the Special Interest Group on Computer Graphics and Interactive Techniques Conference Papers*. pp. 1–9 (2025)
3. Bai, S., Cai, Y., Chen, R., Chen, K., Chen, X., Cheng, Z., Deng, L., Ding, W., Gao, C., Ge, C., et al.: Qwen3-vl technical report. *arXiv preprint arXiv:2511.21631* (2025)
4. Bian, Y., Chen, X., Li, Z., Zhi, T., Sang, S., Luo, L., Xu, Q.: Video-as-prompt: Unified semantic control for video generation. *arXiv preprint arXiv:2510.20888* (2025)
5. Boutsidis, C., Woodruff, D.P.: Optimal cur matrix decompositions (2014), <https://arxiv.org/abs/1405.7910>
6. Charakorn, R., Cetin, E., Tang, Y., Lange, R.T.: Text-to-lora: Instant transformer adaption. *arXiv preprint arXiv:2506.06105* (2025)
7. Fan, C., Lu, Z., Liu, S., Gu, C., Qu, X., Wei, W., Cheng, Y.: Make lora great again: Boosting lora with adaptive singular values and mixture-of-experts optimization alignment. *arXiv preprint arXiv:2502.16894* (2025)
8. Ha, D., Dai, A., Le, Q.V.: Hypernetworks. *arXiv preprint arXiv:1609.09106* (2016)
9. Hamm, K., Huang, L.: Perspectives on cur decompositions (2019), <https://arxiv.org/abs/1907.12668>
10. Hao, Y., Xu, M., Ye, C., Qin, J., Lu, S., Qin, Y., Han, X.: Lofa: Learning to predict personalized priors for fast adaptation of visual generative models. *arXiv preprint arXiv:2512.08785* (2025)
11. Hu, E.J., Shen, Y., Wallis, P., Allen-Zhu, Z., Li, Y., Wang, S., Wang, L., Chen, W., et al.: Lora: Low-rank adaptation of large language models. *Iclr* 1(2), 3 (2022)
12. Huang, Z., He, Y., Yu, J., Zhang, F., Si, C., Jiang, Y., Zhang, Y., Wu, T., Jin, Q., Chanpaisit, N., et al.: Vbench: Comprehensive benchmark suite for video generative models. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 21807–21818 (2024)
13. Junhao Zhang, D., Li, D., Le, H., Shou, M.Z., Xiong, C., Sahoo, D.: Moonshot: Towards controllable video generation and editing with multimodal conditions. *arXiv e-prints* pp. arXiv–2401 (2024)
14. Khan, R.M.S., Tang, D., Li, P., Wang, K., Chen, T.: Oral: Prompting your large-scale loras via conditional recurrent diffusion. *arXiv preprint arXiv:2503.24354* (2025)
15. Liang, Z., Tang, D., Zhou, Y., Zhao, X., Shi, M., Zhao, W., Li, Z., Wang, P., Schürholt, K., Borth, D., Bronstein, M.M., You, Y., Wang, Z., Wang, K.: Drag-and-drop LLMs: Zero-shot prompt-to-weights. In: *The Thirty-ninth Annual Conference on Neural Information Processing Systems* (2025), <https://openreview.net/forum?id=fTkBZLxBzV>
16. Liu, X., Zeng, A., Xue, W., Yang, H., Luo, W., Liu, Q., Guo, Y.: Vfx creator: Animated visual effect generation with controllable diffusion transformer. *arXiv preprint arXiv:2502.05979* (2025)

17. Lv, C., Li, L., Zhang, S., Chen, G., Qi, F., Zhang, N., Zheng, H.T.: Hyperloras: Efficient cross-task generalization via constrained low-rank adapters generation. In: Findings of the Association for Computational Linguistics: EMNLP 2024. pp. 16376–16393 (2024)
18. Ma, Y., Ye, X., Wang, Q., Wang, Y., Liu, H., Zhang, Y., Wang, X., Che, Y., Mo, S., Liang, P., Zhan, F., Chen, Q.: Easyvfx: Frequency-driven decoupling for resource-efficient vfx generation. In: ACM SIGGRAPH (2026)
19. Mao, F., Hao, A., Chen, J., Liu, D., Feng, X., Zhu, J., Wu, M., Chen, C., Wu, J., Chu, X.: Omni-effects: Unified and spatially-controllable visual effects generation. arXiv preprint arXiv:2508.07981 (2025)
20. Meng, F., Wang, Z., Zhang, M.: PiSSA: Principal singular values and singular vectors adaptation of large language models. In: The Thirty-eighth Annual Conference on Neural Information Processing Systems (2024), <https://openreview.net/forum?id=6ZBHIEtdP4>
21. Peebles, W., Xie, S.: Scalable diffusion models with transformers. In: Proceedings of the IEEE/CVF international conference on computer vision. pp. 4195–4205 (2023)
22. Radford, A., Kim, J.W., Hallacy, C., Ramesh, A., Goh, G., Agarwal, S., Sastry, G., Askell, A., Mishkin, P., Clark, J., et al.: Learning transferable visual models from natural language supervision. In: International conference on machine learning. pp. 8748–8763. PmLR (2021)
23. Rombach, R., Blattmann, A., Lorenz, D., Esser, P., Ommer, B.: High-resolution image synthesis with latent diffusion models. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 10684–10695 (2022)
24. Ruiz, N., Li, Y., Jampani, V., Pritch, Y., Rubinstein, M., Aberman, K.: Dreambooth: Fine tuning text-to-image diffusion models for subject-driven generation. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 22500–22510 (2023)
25. Ruiz, N., Li, Y., Jampani, V., Wei, W., Hou, T., Pritch, Y., Wadhwa, N., Rubinstein, M., Aberman, K.: Hyperdreambooth: Hypernetworks for fast personalization of text-to-image models. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 6527–6536 (2024)
26. Wan, T., Wang, A., Ai, B., Wen, B., Mao, C., Xie, C.W., Chen, D., Yu, F., Zhao, H., Yang, J., et al.: Wan: Open and advanced large-scale video generative models. arXiv preprint arXiv:2503.20314 (2025)
27. Wu, J.Z., Ge, Y., Wang, X., Lei, S.W., Gu, Y., Shi, Y., Hsu, W., Shan, Y., Qie, X., Shou, M.Z.: Tune-a-video: One-shot tuning of image diffusion models for text-to-video generation. In: Proceedings of the IEEE/CVF international conference on computer vision. pp. 7623–7633 (2023)
28. Wu, Y., Shi, Y., Wei, J., Sun, C., Yang, Y., Shen, H.T.: DiffLora: Generating personalized low-rank adaptation weights with diffusion. arXiv preprint arXiv:2408.06740 (2024)
29. Ye, H., Zhang, J., Liu, S., Han, X., Yang, W.: Ip-adapter: Text compatible image prompt adapter for text-to-image diffusion models. arXiv preprint arXiv:2308.06721 (2023)
30. Ye, Z., Huang, H., Wang, X., Wan, P., Zhang, D., Luo, W.: Stylemaster: Stylize your video with artistic generation and translation. In: Proceedings of the Computer Vision and Pattern Recognition Conference. pp. 2630–2640 (2025)