

# PRISM: A COMPOSABLE PERSON IMAGE SYNTHESIS MODEL WITH COMPOSITIONAL CONSISTENCY AND UNIFIED OPTIMIZATION

**Anonymous authors**

Paper under double-blind review

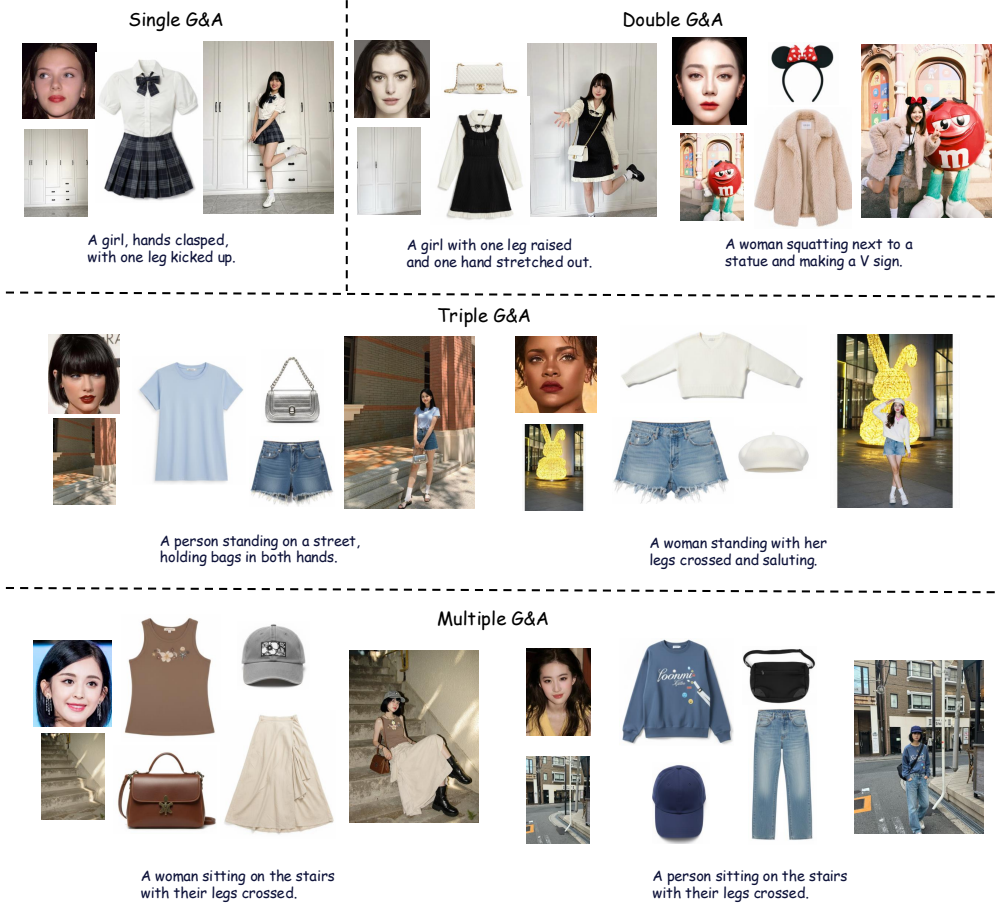


Figure 1: **Showcasing Prism’s exceptional ability** to preserve facial identity while faithfully retaining garment, accessory, and background details. The method delivers striking results across scenarios featuring one, two, three, or even more garments and accessories.

## ABSTRACT

pipeline, termed HMS-Dataset, to construct a large-scale training dataset from single images containing individuals. Building upon this, Prism first encodes facial identity, pertinent clothing elements, and background context into sequences. These sequences are subsequently fused via a novel MM-Attention mechanism. Furthermore, we propose a Compositional Consistency Losses (CCL) strategy to incorporate facial similarity, clothing feature preservation, and background consistency, which are specifically designed to boost facial fidelity, retain intricate clothing details, and enhance overall background coherence. Subsequently, guided by the Minimum Variance Distortionless Response criterion, we propose a Unified Gradient Optimization (UGO) update strategy, which enables fair perceptual optimization for multi-objective optimization problems. Ultimately, Prism demonstrates robust identity preservation and seamless human-environment interaction. Evaluated on our proposed PrismBench, Prism achieves state-of-the-art fidelity and controllability, significantly advancing practical applications in character editing and customizable scene synthesis.

## 1 INTRODUCTION

Recent advances in diffusion-based models have reshaped image synthesis, delivering major gains in both fidelity and controllability. Early approaches (Rombach et al., 2022; Ramesh et al., 2022; Song et al., 2020; Sohl-Dickstein et al., 2015; Gal et al., 2022) relied solely on text prompts, and later work (Ye et al., 2023; Wang & Shi, 2023; Ruiz et al., 2023; Zhang et al., 2023; Xu et al., 2023) broadened conditioning to include visual inputs. To meet the demand for higher precision, many methods now provide specialized controls, such as identity preservation (ID) (Li et al., 2024; Wang et al., 2024b; Ye et al., 2025; Yuan et al., 2025), pose transfer (Zhang et al., 2023), depth map guidance (Zhang et al., 2023), and stylization (Hertz et al., 2024). As models have grown in scale and capability, subject-driven generation using single or multiple reference objects has become a central focus. At the same time, practical applications such as virtual try-on (VTON) are emerging, bringing generative AI closer to real-world deployment.

- We build HMS-Dataset and propose **Prism**, featuring MM-Attention, Compositional Consistency Losses, and Unified Gradient Optimization.
- The performance of our proposed PrismBench demonstrates state-of-the-art performance in fidelity, controllability, and human-environment alignment.

## 2 RELATED WORK

**Diffusion Model.** Diffusion has become the dominant text-to-image paradigm, evolving from text-conditioned systems with strong language encoders and classifier-free guidance (Ho & Salimans, 2021) to efficient latent-space models with high fidelity and broad controllability (Esser et al., 2024; Ramesh et al., 2022; Podell et al., 2023; Feng et al., 2023; Balaji et al., 2022). Controllability, specialization, and personalization are advanced by external conditioning branches and parameter-efficient tuning (e.g., LoRA), enabling structural guidance, rapid domain adaptation, and image-conditioned generation via adapters or attention controls (Zhang et al., 2023; Mou et al., 2024; Hu et al., 2022; Gal et al., 2022; Ruiz et al., 2023; Ye et al., 2023; Chen et al., 2024; Cao et al., 2023), motivating unified, instruction-driven pipelines (Xiao et al., 2025; Wu et al., 2025a; Chen et al., 2025b). A concurrent shift replaces convolutional U-Nets with transformer backbones operating



Figure 2: **Data Curation Pipeline.** Structured conditions—face, garments, and background—are extracted using InsightFace (Deng et al., 2019), Qwen-VL-Max (Wang et al., 2024a), and inpainting-based background reconstruction.

with target velocity  $v^* = z_0 - \epsilon$ . The network regresses  $v_\theta(z_t, t, C_T)$  using mean-squared error:

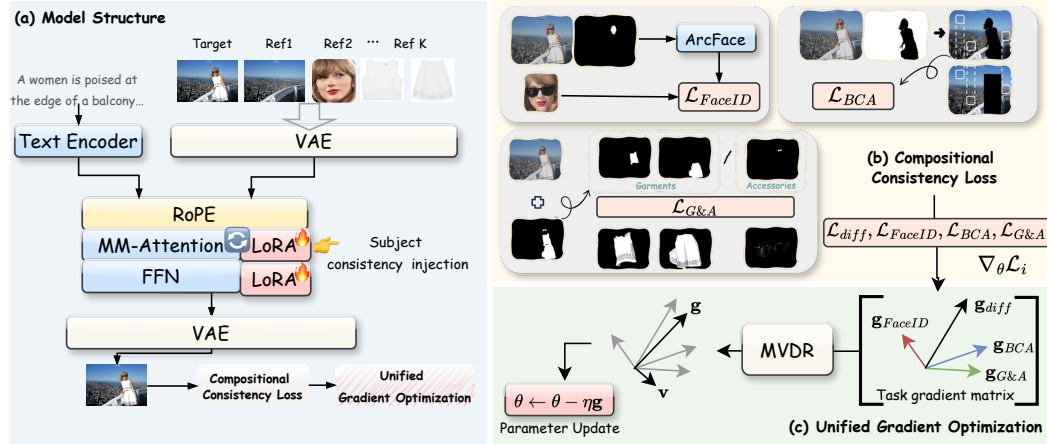


Figure 3: An overview of the **Prism** framework. (a) Model Structure: Multi-condition references—including face, garments, and background- are encoded via a VAE, with subject consistency injected through MM-Attention and LoRA modules. (b) Compositional Consistency Loss: Identity preservation, garment/accessory fidelity, and background alignment are supervised by  $\mathcal{L}_{\text{FaceID}}$ ,  $\mathcal{L}_{\text{G\&A}}$ , and  $\mathcal{L}_{\text{BCA}}$ , respectively. (c) Unified Gradient Optimization: Multi-objective gradients are harmonized under

where  $A_{\text{garment} \rightarrow \text{tgt}}$  is the attention matrix and  $\mathbf{G}_k$  are normalized reference coordinates. Smoothness is enforced via total variation within the garment mask  $\mathbf{M}_{\text{garment}}$ :

$$\mathcal{L}_{\text{G\&AL}} = \|\nabla(\mathbf{F}_{\text{garment}} \odot \mathbf{M}_{\text{garment}})\|_1. \quad (8)$$

This encourages locally coherent attention, transferring textures and patterns accurately.

**BCAL: Background Correspondence Attention Loss.** Maintaining background consistency is important for seamless composition. Instead of dense supervision, we use sparse semantic correspondences  $\mathcal{C}_{\text{bg}} = \{(u_j, v_j)\}$  between the reference background and target latent:

$$\mathcal{L}_{\text{BCAL}} = -$$

Table 1: Quantitative comparison of composable person image generation on **PrismBench**. Gray rows indicate *closed-source models* evaluated via official interfaces.

| Method                      | FaceSim $\uparrow$ | G&ASim $\uparrow$ | BackSim $\uparrow$ | LPIPS $\downarrow$ | Sem. Cons. $\uparrow$ |
|-----------------------------|--------------------|-------------------|--------------------|--------------------|-----------------------|
| UNO (Wu et al., 2025b)      | 40.52              | 79.78             | 59.71              | 70.16              | 33.32                 |
| DreamO (Mou et al., 2025)   | 52.42              | 79.67             | 57.44              | 72.37              | 34.29                 |
| OmniGen (Xiao et al., 2025) | 51.30              | 84.27             | 65.86              | 67.20              | 33.07                 |

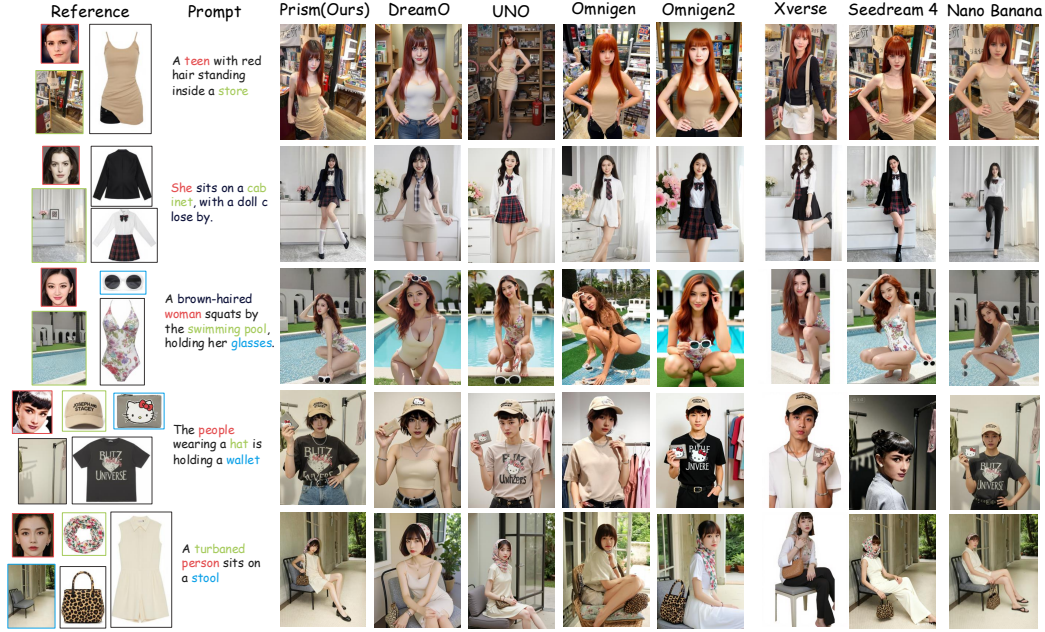


Figure 4: Qualitative Comparison. We compare with Nano Banana and SeedDream 4.0 using their official interfaces, enabling a fair evaluation against closed-source models.

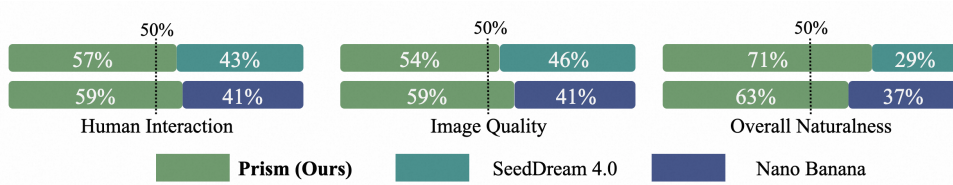


Table 2: Impact of each loss terms on the performance of **Prism**, e.g., FIDL, G&AL, and BCAL.

| Method   | FaceSim $\uparrow$ | G&ASim $\uparrow$ | BackSim $\uparrow$ | LPIPS $\downarrow$ | Sem. Cons. $\uparrow$ |
|----------|--------------------|-------------------|--------------------|--------------------|-----------------------|
| w/o FIDL | 49.34              | 88.24             | 84.21              | 52.77              | <b>34.57</b>          |
| w/o G&AL | 55.12              | 77.96             | 83.71              | 57.81              | 31.75                 |
| w/o BCAL | 54.91              | 88.16             | 78.51              | 68.83              | 29.56                 |

## REFERENCES

- Yogesh Balaji, Seungjun Nah, Xun Huang, Arash Vahdat, Jiaming Song, Qinsheng Zhang, Karsten Kreis, Miika Aittala, Timo Aila, Samuli Laine, et al. ediff-i: Text-to-image diffusion models with an ensemble of expert denoisers. *arXiv preprint arXiv:2211.01324*, 2022.
- ByteDance. Seed dream4.0. <https://jimeng.jianying.com/>, 2025.
- Mingdeng Cao, Xintao Wang, Zhongang Qi, Ying Shan, Xiaohu Qie, and Yinqiang Zheng. Masactrl: Tuning-free mutual self-attention control for consistent image synthesis and editing. In *ICCV*, pp. 22560–22570, 2023.
- Bowen Chen, Mengyi Zhao, Haomiao Sun, Li Chen, Xu Wang, Kang Du, and Xinglong Wu. Xverse: Consistent multi-subject control of identity and semantic attributes via dit

- Minchul Kim, Anil K Jain, and Xiaoming Liu. Adaface: Quality adaptive margin for face recognition. In *CVPR*, pp. 18750–18759, 2022.
- Diederik P Kingma. Adam: A method for stochastic optimization. In *ICLR*, 2015.
- Nupur Kumari, Bingliang Zhang, Richard Zhang, Eli Shechtman, and Jun-Yan Zhu. Multi-concept customization of text-to-image diffusion. In *CVPR*, pp. 1931–1941, 2023.
- Black Forest Labs, Stephen Batifol, Andreas Blattmann, Frederic Boesel, Saksham Consul, Cyril Diagne, Tim Dockhorn, Jack English, Zion English, Patrick Esser, et al. Flux. 1 kontekst: Flow matching for in-context image generation and editing in latent space. *arXiv preprint arXiv:2506.15742*, 2025.
- Dongxu Li, Junnan Li, and Steven Hoi. Blip-diffusion: Pre-trained subject representation for controllable text-to-image generation and editing

- Chenyuan Wu, Pengfei Zheng, Ruiran Yan, Shitao Xiao, Xin Luo, Yueze Wang, Wanli Li, Xiyan Jiang, Yexin Liu, Junjie Zhou, et al. Omnigen2: Exploration to advanced multimodal generation. *arXiv preprint arXiv:2506.18871*, 2025a.
- Shaojin Wu, Mengqi Huang, Wenxu Wu, Yufeng Cheng, Fei Ding, and Qian He. Less-to-more generalization: Unlocking more controllability by in-context generation. *arXiv preprint arXiv:2504.02160*, 2025b.
- Shitao Xiao, Yueze Wang, Junjie Zhou, Huaying Yuan, Xingrun Xing, Ruiran Yan, Chaofan Li, Shuting Wang, Tiejun Huang, and Zheng Liu. Omnigen: Unified image generation. In *CVPR*, pp. 13294–13304, 2025.
- Yiming Xie, Varun Jampani, Lei Zhong, Deqing Sun, and Huaizu Jiang. Omnicontrol: Control any joint at any time for human motion generation. *arXiv preprint arXiv:2310.08580*

**Algorithm 1** Unified Gradient Optimization (Pseudocode, PyTorch-like)

---

```

def UGO_step(model, losses, optimizer, damping=1e-4, weights=None):
    params = [p for p in model.parameters() if p.requires_grad]
    Gs = []
    for L in losses:
        g = torch.autograd.grad(L, params, retain_graph=True, allow_unused=True)
        g = [torch.zeros_like(p) if gi is None else gi for gi, p in zip(g, params)]
        Gs.append(torch.cat([gi.reshape(-1) for gi in g]).detach())

    G = torch.stack(Gs, dim=1) # [D, m]
    m, dev, dt = G.shape[1], G.device, G.dtype

    w = torch.ones(m, device=dev, dtype=dt) / m if weights is None else torch.as_tensor(
        weights, device=dev, dtype=dt)
    w = w / (w.sum()
```

## B.1 GARMENTS &amp; ACCESSORIES EXTRACTION

## Prompt of Garment &amp; Accessory Catalogue Generation

You are an expert fashion analyst and a precise object detection AI. Your primary function is to meticulously analyze the provided image to identify every single garment and accessory worn by the person.

Based on your analysis, you must generate a JSON object that strictly adheres to the following rules:

- **Scope:** Your analysis must be strictly limited to garments (clothing) and accessories. Ignore the person, handheld items (like phones or cups), and any background elements.
- **item\_name:** For each item, create a descriptive but concise name. Include its primary color, pattern (if any), and type.

## B.2 BACKGROUND INPAINTING

### Prompt of Background Reconstruction Guide

Seamlessly fill in the background, analyzing the surrounding environment with high precision. Isolate and erase the human figure(s) and all associated objects. Reconstruct the missing background area by meticulously continuing all lines, patterns, and textures from the surrounding context. Pay critical attention to maintaining correct perspective, shadow continuity, and realistic depth of field. The inpainted section must be photorealistically integrated, leaving no visible seams or artifacts.