

EraseLoRA: MLLM-Driven Foreground Exclusion and Background Subtype Aggregation for Dataset-Free Object Removal

Sanghyun Jo^{1,2,*†} , Donghwan Lee^{2,*} , Eunji Jung^{2,*} , Seong Je Oh² , and
Kyungsu Kim^{2†} 

¹ OGQ, Seoul, Korea

² Seoul National University, Seoul, Korea

Abstract. Object removal must prevent the masked target from reappearing and reconstruct the occluded background with structural and contextual fidelity, rather than merely filling a hole plausibly. Recent dataset-free approaches manipulate the diffusion model’s internal self-attention to prevent it from referencing the masked region, yet they fail in two critical ways: (i) they treat the masked region as the sole foreground, misinterpreting non-target objects as background and regenerating them, and (ii) they apply uniform attention constraints without distinguishing diverse background subtypes, leading to textural blurring and structural misalignment. Both failures stem from the absence of explicit background-aware reasoning. We propose EraseLoRA, a dataset-free framework that replaces attention surgery with background-aware reasoning and test-time adaptation. The first stage, Background-aware Foreground Exclusion (BFE), leverages a multimodal large-language model to separate target foreground, non-target foregrounds, and clean background from a single image-mask pair. The second stage, Background-aware Reconstruction with Subtype Aggregation (BRSA), performs test-time optimization that treats inferred background subtypes as complementary pieces, enforcing their consistent integration through reconstruction and alignment objectives without explicit attention intervention. As a model-agnostic plug-in applicable to diverse diffusion backbones, EraseLoRA reconstructs backgrounds at least 23% more faithful to the original scene than previous dataset-free methods while nearly halving unwanted foreground re-generation, and surpasses all dataset-driven approaches in both aspects despite requiring no training data. Code is available at <https://shjo-april.github.io/EraseLoRA>.

Keywords: object removal · multimodal large-language model · test-time adaptation · dataset-free · attention

* Equal contribution.

† Corresponding authors: shjo.april@gmail.com, kyskim@snu.ac.kr

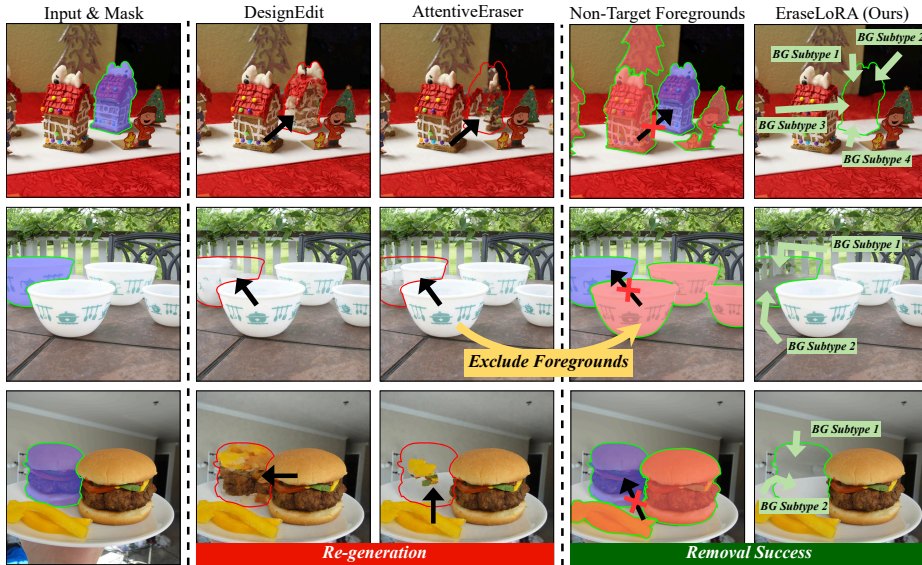


Fig. 1: Qualitative comparison with prior dataset-free methods. Previous state-of-the-art approaches [14,33] treat only the masked region as foreground, misinterpreting non-target objects as background and regenerating them. EraseLoRA identifies and excludes non-target foregrounds and reconstructs the masked region using various background subtypes, enabling faithful object removal.

1 Introduction

Image inpainting methods based on GANs [11, 21, 34] and text-to-image diffusion models [8, 12, 26, 30] can synthesize visually plausible content in missing regions, yet they primarily aim to generate realistic textures rather than restore the underlying background structure, often hallucinating new objects instead of faithfully reconstructing what lies behind the removed target.

Object removal, in contrast, requires both eliminating the target and recovering the occluded background by transferring clean background cues into the masked region with structural consistency. Recent dataset-free diffusion methods [14, 33] attempt to achieve this by redirecting or blocking self-attention within the masked region so that the model focuses on unmasked context. While effective in simple cases, these approaches share two inherent limitations. First, they treat the masked region as the only foreground and often misinterpret non-target foregrounds outside the mask as background, causing unintended re-generation of objects (see Fig. 1). This reflects a lack of background-aware reasoning. Second, constraining attention uniformly without distinguishing distinct background subtypes compromises local detail and fails to coherently integrate multiple background cues, leading to blurred textures and unnatural boundaries between disparate subtypes (see Fig. 2; further analysis in B.2).

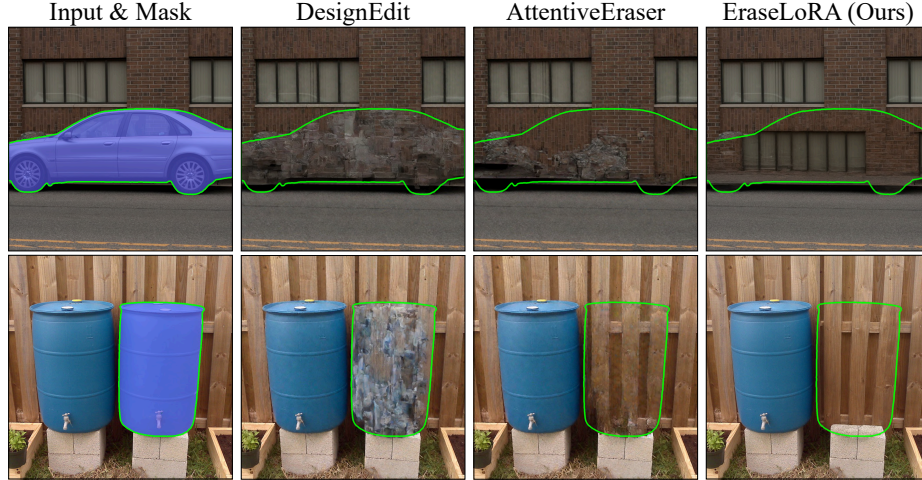


Fig. 2: Artifacts from attention manipulation. Recent dataset-free methods [14, 33] directly modify attention inside the mask without identifying background cues, leading to blurred or distorted textures, whereas EraseLoRA aggregates background subtypes without attention blocking and preserves sharp, coherent structures.

We introduce EraseLoRA, a dataset-free framework that addresses these issues through background-aware reasoning and test-time adaptation. The first stage, Background-aware Foreground Exclusion (BFE), leverages a multimodal large-language model (MLLM) [2, 48] to produce clean background cues by separating target foreground, non-target foregrounds, and background from a single image-mask pair. The second stage, Background-aware Reconstruction with Subtype Aggregation (BRSA), performs test-time optimization with Low-Rank Adaptation (LoRA) [13] to inject these cues into the masked area and aggregate multiple inferred background subtypes into a coherent reconstruction without a dataset-level removal prior or explicit attention blocking. EraseLoRA is validated as a plug-in across diverse pretrained diffusion backbones [3, 8] and standard object-removal benchmarks [19, 31], demonstrating consistent gains over dataset-free baselines and competitive performance with dataset-driven methods. By combining the reasoning capability of MLLMs with the generative fidelity of diffusion models, EraseLoRA establishes an extensible formulation of dataset-free object removal that requires no additional data or retraining.

Our key contributions are as follows:

- We identify a fundamental failure mode in object removal: non-target foregrounds are frequently misinterpreted as background, causing their unintended regeneration across recent dataset-free methods.
- We propose EraseLoRA, a background-aware, dataset-free object-removal framework that combines MLLM-guided separation of target and non-target foregrounds from background with a multi-background-aware test-time adap-

Table 1: Conceptual comparison of EraseLoRA (Ours) with previous approaches for object removal.

Properties	[ECCV'24] PowerPaint [51]	[CVPR'25] EntityErasure [50]	[CVPR'25] SmartEraser [15]	[CVPR'26] ObjectClear [44]	[AAAI'25] DesignEdit [14]	[AAAI'25] AttentiveEraser [33]	EraseLoRA
Dataset-free object removal	✗	✗	✗	✗	✓	✓	✓
Identifies non-target foregrounds with backgrounds	✗	✗	✗	✗	✗	✗	✓
Leverages multiple background subtypes	✗	✗	✗	✗	✗	✗	✓
Model-agnostic applicability	✗	✗	✗	✗	✓	✓	✓

tation scheme, preventing foreground regeneration while maintaining contextual coherence.

- We provide three-label ground-truth annotations that distinguish target foreground, non-target foreground, and background, along with two evaluation metrics designed for unpaired object-removal settings.
- EraseLoRA improves background fidelity by at least 23% over previous dataset-free methods while nearly halving unwanted foreground re-generation, and retains these gains when diffusion backbones and MLLMs are each replaced by alternatives.

2 Related Work

2.1 Image Inpainting with Generative Models

Image inpainting aims to complete missing regions using the visible context. Early approaches [32, 46, 52] based on GANs have been surpassed by diffusion-based methods [40, 42], which produce stable, detailed, and high-fidelity completions. Building on text-to-image diffusion backbones [26, 30], existing methods [25, 40] fine-tune these backbones on paired inpainting datasets so that they can exploit text prompts while learning to fill masked regions. However, these approaches are still optimized for context-consistent completion and tend to hallucinate plausible new objects rather than preserving the original scene content.

2.2 Diffusion Models for Object Removal

Object removal is a specialized form of inpainting that must not only erase the masked target but also restore the occluded background with structural and contextual fidelity while preserving target-unrelated regions. Existing methods fall into two categories: dataset-driven approaches [7, 15, 24, 51], which learn removal priors from additional, often paired before/after data, and dataset-free approaches [6, 14, 33], which operate directly on pretrained text-to-image diffusion models [8, 26, 30] without additional data. However, dataset-driven methods inherit a static training distribution and do not explicitly distinguish non-target foregrounds from background, which can leave object traces or perturb target-unrelated regions in complex scenes. Moreover, constructing paired examples

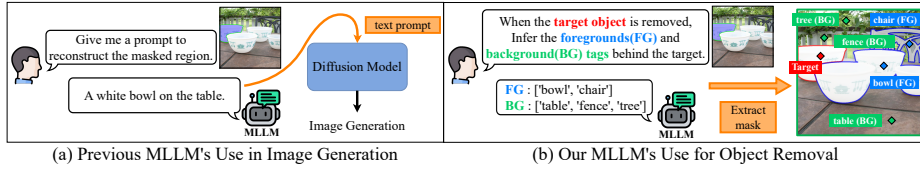


Fig. 3: Background-aware reasoning power of MLLM. Unlike prior works [16, 27, 37, 47] employ MLLMs for visual reasoning over the visible scene, we first leverage MLLMs to infer background cues behind the masked target.

is expensive and often unrealistic, since most pairs must be synthesized or extracted from video frames. These limitations motivate dataset-free methods that control the diffusion process at inference time. Recent extensions address object effects such as shadows and reflections [10, 39, 44]. While related, they pursue a broader removal target and can distort target-unrelated regions, which object removal should preserve. We focus on removal within the given mask, prioritizing faithful restoration and preservation of target-unrelated regions.

Recent state-of-the-art dataset-free approaches [14, 33] redirect or suppress self-attention within the masked region to guide the model toward unmasked context. While this attention manipulation reduces unwanted regeneration to some extent, it introduces two fundamental limitations. First, these methods remain background-agnostic: by treating only the masked area as foreground, they often misinterpret non-target foregrounds as background and regenerate them, as shown in Fig. 1. Second, blocking or altering attention disrupts fine textures and prevents consistent integration of multiple background cues, leading to blurry or structurally inconsistent results, as illustrated in Fig. 2.

To address these limitations, we introduce a background-aware, dataset-free framework that leverages visual reasoning from an external model to identify and exclude non-target foregrounds and to produce clean background cues for reconstruction. We further aggregate multiple inferred background subtypes coherently through test-time adaptation, enabling structurally consistent background restoration without a dataset-level removal prior or explicit attention blocking. Our method is plug-and-play and model-agnostic, applicable across diverse diffusion backbones. Tab. 1 summarizes this design, highlighting how EraseLoRA differs from previous state-of-the-art methods [14, 15, 33, 50].

2.3 MLLMs for Visual Reasoning in Image Editing

Multimodal large-language models (MLLMs) [1, 2, 22, 48] have gained traction in vision-language tasks due to their ability to interpret visual scenes and reason about object relations. Recent editing methods [16, 37] leverage this capability to extract semantic descriptions, generate editing instructions, or guide global scene manipulation, while inpainting works [9, 41, 47] use MLLMs to analyze the visible context and propose content to fill masked regions. However, these

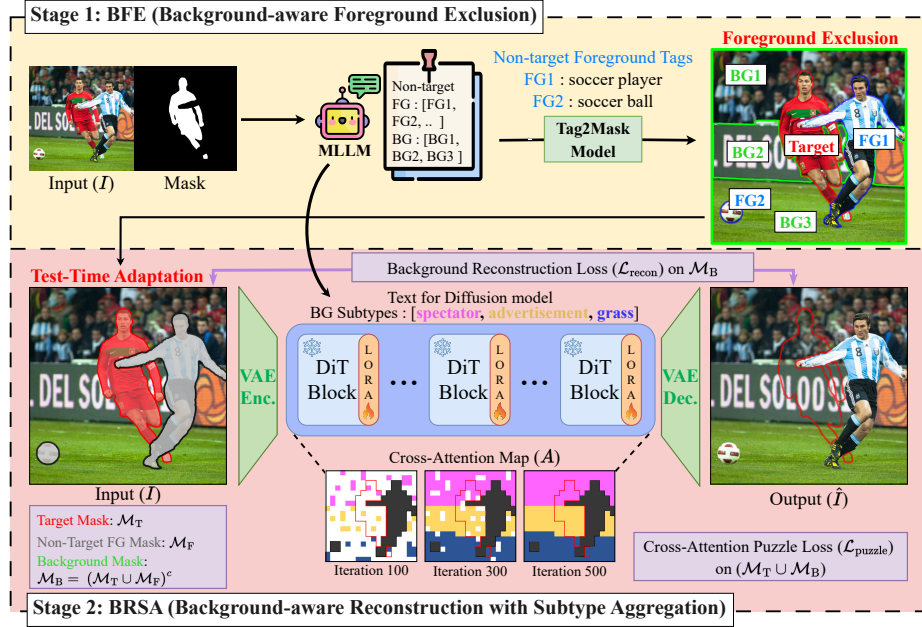


Fig. 4: Overview of EraseLoRA. BFE (Sec. 3.1) separates target foreground, non-target foregrounds, and background from a single image-mask pair using an MLLM [48] and Tag2Mask models [23, 28]. After producing clean background cues, BRSA (Sec. 3.2) performs test-time adaptation [36] with reconstruction and alignment objectives, coherently integrating background subtypes into the masked region.

approaches rely on visible context and aim to generate new objects, rather than inferring what lies behind a removed target.

We leverage MLLMs in a fundamentally different role: as background-aware reasoners for object removal. Rather than generating new foreground content, we use MLLMs to infer the occluded background behind the target and to identify non-target foregrounds that cause unintended regeneration, producing clean background cues that are not directly visible in the input image (see Fig. 3).

3 Method

Our proposed EraseLoRA is a dataset-free object removal framework that leverages MLLM-guided background reasoning and test-time adaptation to achieve coherent background reconstruction. It consists of two stages: Background-aware Foreground Exclusion (BFE; Sec. 3.1) and Background-aware Reconstruction with Subtype Aggregation (BRSA; Sec. 3.2). The overall pipeline is illustrated in Fig. 4, and we provide diffusion and attention preliminaries in Appendix A.



Fig. 5: Identification of non-target foregrounds. Prior methods [6, 33] treat the entire unmasked region as background, which causes regeneration of non-target foregrounds. In contrast, EraseLoRA explicitly identifies non-target foregrounds within the unmasked region and excludes them, producing clean background.

3.1 Background-aware Foreground Exclusion

The first stage, BFE, prevents unintended object regeneration by explicitly excluding non-target foregrounds from reference regions and extracting clean background cues for contextually coherent reconstruction. Given an input image I and a target mask M_T , we leverage the background-aware reasoning of MLLMs [2, 48] to partition target foreground, non-target foregrounds, and background. The MLLM first identifies all semantic tags in the image and classifies the masked object as the target foreground, visible objects that may cause regeneration as non-target foreground tags $\mathcal{F}$, and occluded objects or scene components behind the target as background subtype tags $\mathcal{B}$. For each non-target foreground tag in $\mathcal{F}$, we use Tag2Mask models (*e.g.*, Grounding DINO [23] and SAM2 [28]) to localize its corresponding region. The union of these localized regions defines the non-target foreground mask; the remaining pixels outside both the target and non-target foreground regions are treated as clean background.

We define three binary masks M_T , M_F , and M_B over the latent spatial domain $\Omega = \{1, \dots, h \times w\}$. Following the VAE architecture of the employed diffusion backbone (*e.g.*, [8]), we set $h=H/8$ and $w=W/8$ for an input of size $H \times W$:

$$\Omega = M_T \cup M_F \cup M_B, \quad (1)$$

where M_T , M_F , and M_B denote the target, non-target foreground, and clean background masks in the latent space, respectively. This partition explicitly distinguishes non-target foreground objects from the unmasked background region, allowing us to isolate distractors that previous dataset-free methods [14, 33] mistakenly treat as background (see Fig. 5), which in turn yields cleaner background supervision for subsequent reconstruction.

3.2 Background-aware Reconstruction with Subtype Aggregation

Based on the clean background cues obtained in BFE, BRSA performs test-time optimization with Low-Rank Adaptation (LoRA) [13] to effectively aggregate

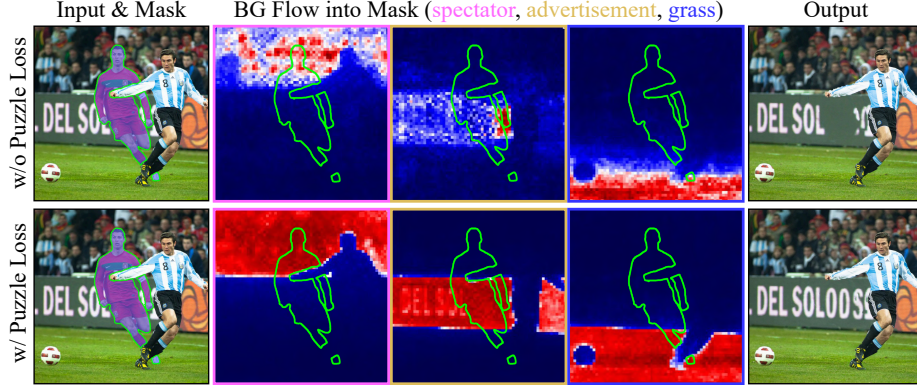


Fig. 6: Effect of the background puzzle loss. We visualize how each background subtype (spectator, advertisement, grass) is represented inside the mask. The background puzzle loss ensures structurally coherent integration of background subtypes within the mask, unlike the weak integration without it.

multiple background subtypes and reconstruct the masked region with structural and contextual consistency. This adaptation provides image-specific background fidelity in dataset-free setting without artifacts from attention manipulation. To achieve this, BRSA jointly optimizes two complementary objectives: the Background Reconstruction Loss ($\mathcal{L}_{\text{recon}}$) and the Background Puzzle Loss ($\mathcal{L}_{\text{puzzle}}$). The overall objective is formulated as $\mathcal{L}_{\text{total}} = \mathcal{L}_{\text{recon}} + \lambda \mathcal{L}_{\text{puzzle}}$, where $\lambda=0.2$ balances background reconstruction fidelity and subtype aggregation. Together, these losses guide how background information is integrated into the masked region, softly regulating attention flow without the hard attention-blocking used in prior methods [14, 33].

Background Reconstruction Loss. To preserve regions that are confidently identified as clean background by BFE (Sec. 3.1), we impose a reconstruction loss only on M_B . Let $z = \text{Enc}(I)$ be the latent representation of the input image I , and $\hat{z}$ be the reconstructed latent after denoising. The background reconstruction loss is defined as

$$\mathcal{L}_{\text{recon}} = \frac{1}{|M_B|} \sum_{p \in M_B} \|\hat{z}[p] - z[p]\|_2^2, \quad (2)$$

where p denotes spatial indices in the latent feature map. By anchoring $\hat{z}$ to z on these background locations, $\mathcal{L}_{\text{recon}}$ enforces fidelity to the original background and promotes globally coherent reconstruction.

Background Puzzle Loss. While the reconstruction loss $\mathcal{L}_{\text{recon}}$ preserves high-fidelity background appearance, it does not explicitly control how different background subtypes are filled and integrated into the masked area, often leading to structurally inconsistent or partially missing patterns (*e.g.*, misaligned background context in Fig. 6). To address this, we introduce a background puzzle loss that treats each background subtype as a distinct puzzle piece that must con-

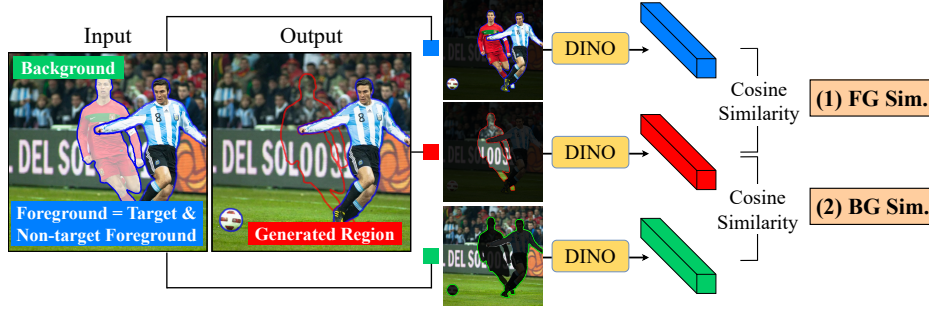


Fig. 7: Illustration of evaluation metrics (*i.e.*, BG Sim. and FG Sim.) for unpaired object removal.

tribute coherently to reconstructing the masked region. This loss enforces that background attention is concentrated only on valid regions (target foreground or clean background), preventing attention distraction toward non-target foregrounds and improving spatial consistency:

$$\mathcal{L}_{\text{puzzle}} = 1 - \text{Dice}(A^{\text{dom}}, \mathbf{1}_{M_T \cup M_B}), \quad (3)$$

where $\text{Dice}(\cdot, \cdot)$ computes the soft spatial overlap between the continuous attention map and the binary valid-region indicator (see Appendix B for details). $A^{\text{dom}}[p] = \max_{b \in \mathcal{B}} A_b[p]$ selects the strongest attention response at each location p across background subtype tags $b \in \mathcal{B}$ inferred by BFE (Sec. 3.1), ensuring that every position is accounted for by its most relevant subtype.

4 Experiments

4.1 Experimental Setup

Implementation details. We implement EraseLoRA on three text-to-image diffusion backbones [3, 8, 26]. For a fair comparison, we strictly follow the official inference configurations (scheduler, guidance scale, resolution, and the number of sampling steps) and apply EraseLoRA in a purely dataset-free, test-time manner. During test-time adaptation (TTA), we freeze all backbone parameters and optimize only the inserted LoRA adapters [13] (rank $r=32$) for 500 iterations with respect to the final loss defined in Sec. 3.2. This configuration is kept consistent across all backbones and benchmarks. For BFE (Sec. 3.1), we use InternVL3-78B [48] as the default MLLM and Grounded SAM2 [28, 29] for tag-to-mask conversion. Additional details are provided in Appendix B.

Benchmarks. We evaluate EraseLoRA on two benchmarks: 200 samples from OpenImages V7 [19] and 343 frames from RORD [31]. In both datasets, the original annotations do not distinguish non-target foregrounds from background, so we annotate them with three-label (target / non-target foreground /

Table 2: Quantitative comparison with previous state-of-the-art methods on test datasets [19, 31]. The best results are in **bold**.

	OpenImages V7			RORD		
	BG Sim.($\uparrow$)	FG Sim.($\downarrow$)	BG Pres.($\uparrow$)	BG Sim.($\uparrow$)	FG Sim.($\downarrow$)	BG Pres.($\uparrow$)
<i>Dataset-Free Approaches:</i>						
SD3.5-M [8]	0.605	0.286	0.934	0.582	0.319	0.907
+ AttentiveEraser [33]	0.559	0.276	0.931	0.541	0.302	0.901
+ DesignEdit [14]	0.600	0.255	0.933	0.597	0.273	0.908
+ EraseLoRA	0.743	0.151	0.924	0.779	0.138	0.901
<i>Dataset-Driven Approaches:</i>						
SDXL-Inpainting [26]	0.677	0.212	0.742	0.645	0.234	0.720
PowerPaint [51]	0.669	0.217	0.719	0.729	0.176	0.687
CLIPAway [7]	0.656	0.223	0.713	0.744	0.156	0.705
SmartEraser [15]	0.709	0.185	0.727	0.768	0.148	0.672
EntityErasure [50]	0.679	0.204	0.728	0.766	0.175	0.716

background) ground-truth masks to capture distractor regions. These refined annotations will be released and form the basis of the evaluation metrics described below. Additional results on RemovalBench [39] are provided in Appendix C.

Evaluation metrics. For unpaired object removal, we use DINO similarity [5], following recent image generation and editing works [20, 43]. Foreground Similarity (FG Sim.) and Background Similarity (BG Sim.) measure, respectively, how much the reconstructed region stays similar to the foreground (lower is better) and how well it aligns with the background (higher is better; see Fig. 7). We additionally report Background Preservation (BG Pres.), computed via SSIM [38] on unmasked regions, following the evaluation protocol of recent editing methods [18, 49], to assess fidelity outside the masked area.

4.2 Comparison with State-of-the-art Approaches

Table 2 shows the quantitative results on two benchmarks [19, 31]. Compared to the baseline [8], EraseLoRA substantially improves BG Sim., from 0.605 to 0.743 on OpenImages V7 [19] and from 0.582 to 0.779 on RORD [31] (absolute gains of +0.14 and +0.20, corresponding to roughly 23% and 34% relative improvements). At the same time, FG Sim. is almost halved (0.286→0.151 on OpenImages V7, 0.319→0.138 on RORD), indicating that EraseLoRA suppresses foreground re-generation while filling the mask with background-consistent content. EraseLoRA also outperforms dataset-driven methods [7, 15, 26, 50, 51]: it attains the highest BG Sim. and the lowest FG Sim. on both benchmarks [19, 31], while maintaining background preservation around 0.90, which is about 0.18 higher than all five dataset-driven methods. This indicates that EraseLoRA predominantly modifies only the masked region and leaves the unmasked background almost unchanged, unlike dataset-driven models that often perturb surrounding content. Qualitative comparisons in Fig. 8 reflect the same trend: EraseLoRA removes the target object cleanly while preserving sharp background details and

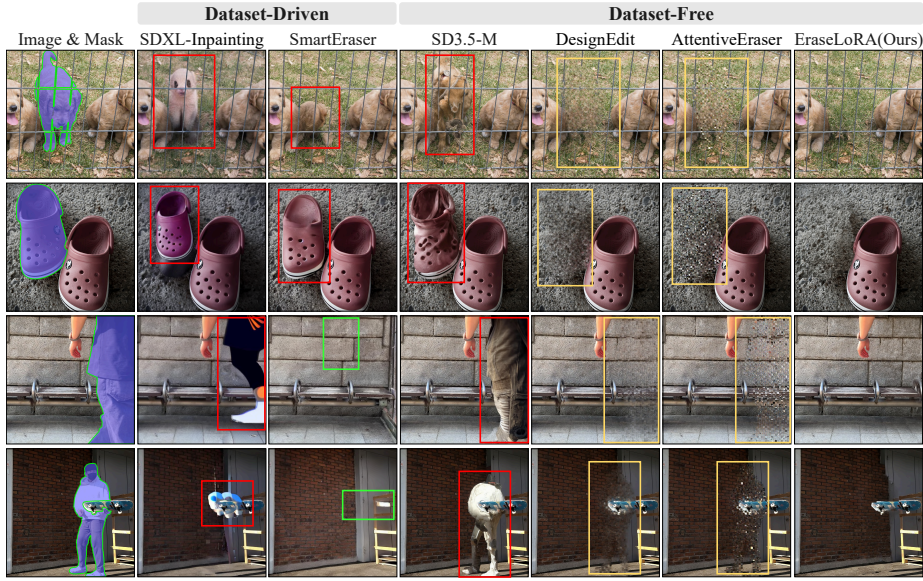


Fig. 8: Qualitative comparison on OpenImages V7 [19] and RORD [31]. Without any paired data, EraseLoRA avoids unwanted background changes (green), copying nearby objects (red), and residual foreground artifacts (yellow), achieving cleaner object removal and more faithful backgrounds than both dataset-driven and dataset-free methods [8, 14, 15, 26, 33].

avoiding unwanted edits in unmasked regions. Additional quantitative and qualitative results are provided in Appendix C and F.

4.3 Discussion

To understand EraseLoRA’s performance gains, we conduct component-wise ablations on OpenImages V7 [19], measuring the effect of each stage and loss term.

Effect of BFE with prior dataset-free methods. To test whether our foreground exclusion module (BFE; Sec. 3.1) generalizes beyond EraseLoRA, we plug it into prior dataset-free object-removal methods [14, 33] without modifying their inference pipelines. By explicitly excluding non-target foregrounds from their reference regions, background similarity improves by up to 6.6%, while foreground similarity decreases by 8.6% (Tab. 3, left). As shown in the left panel of Fig. 9, non-target foregrounds that were previously regenerated are successfully suppressed once BFE provides clean background references, confirming that BFE alleviates the re-generation limitation of existing dataset-free approaches [14, 33] in a fully model-agnostic manner.

Effect of losses in BRSA. To determine the most effective objective for BRSA (Sec. 3.2), we examine the roles of the two complementary losses through a loss-combination study (see Tab. 3, right). The background reconstruction loss

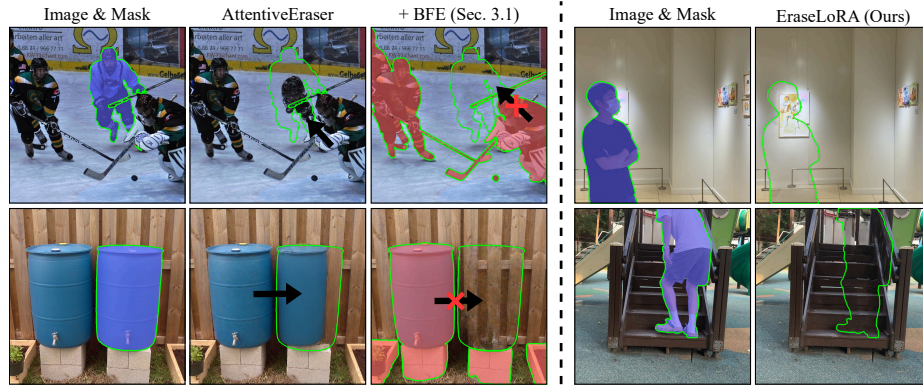


Fig. 9: (Left) Effect of BFE (Sec. 3.1) on AttentiveEraser [33]. Red masks denote excluded non-target foregrounds. (Right) Examples of EraseLoRA in occlusion cases. EraseLoRA reconstructs occluded content (*e.g.*, a painting, stairs) as background from surrounding context.

Table 3: (Left) Effect of non-target foreground exclusion (BFE; Sec. 3.1) in previous state-of-the-art dataset-free methods [14, 33] and (Right) loss component in BRSA (Sec. 3.2).

Method	Metrics		Method	Loss Components		Metrics	
	BG Sim.(↑)	FG Sim.(↓)		$\mathcal{L}_{\text{recon}}$	$\mathcal{L}_{\text{puzzle}}$	BG Sim.(↑)	FG Sim.(↓)
AttentiveEraser [33]	0.559	0.276	SD3.5-M [8]	✗	✗	0.605	0.286
+ BFE (Ours; Sec. 3.1)	0.596	0.252	+ EraseLoRA	✓	✗	0.736	0.158
DesignEdit [14]	0.600	0.255	+ EraseLoRA	✗	✓	0.561	0.278
+ BFE (Ours; Sec. 3.1)	0.603	0.251	+ EraseLoRA	✓	✓	0.743	0.151

($\mathcal{L}_{\text{recon}}$; Eq. (2)) preserves structural background consistency, improving BG Sim. from 0.605 to 0.736 (+21.7%), but can still leave faint foreground traces. In contrast, the background puzzle loss ($\mathcal{L}_{\text{puzzle}}$; Eq. (3)) suppresses foreground artifacts via background subtype aggregation, but alone lacks explicit background anchoring and fails to capture fine background structure and global patterns, often leading to visually inconsistent completions (see Fig. 10). Together (*i.e.*, EraseLoRA), the two losses achieve the best results, yielding coherent and detail-preserving background restoration while overcoming the limitations observed when either loss is used alone.

Flexibility. To evaluate the general applicability of EraseLoRA, we apply our method to different text-to-image diffusion backbones [3, 8, 26, 30], MLLMs of varying scales [2, 22, 45, 48], and Tag2Mask models [4, 17, 29, 35]. EraseLoRA demonstrates generalization ability by yielding reliable gains in both BG Sim. and FG Sim. regardless of the underlying diffusion backbones (see Tab. 7 and Fig. 15). Detailed quantitative and qualitative analyses are provided in Appendix D.1.

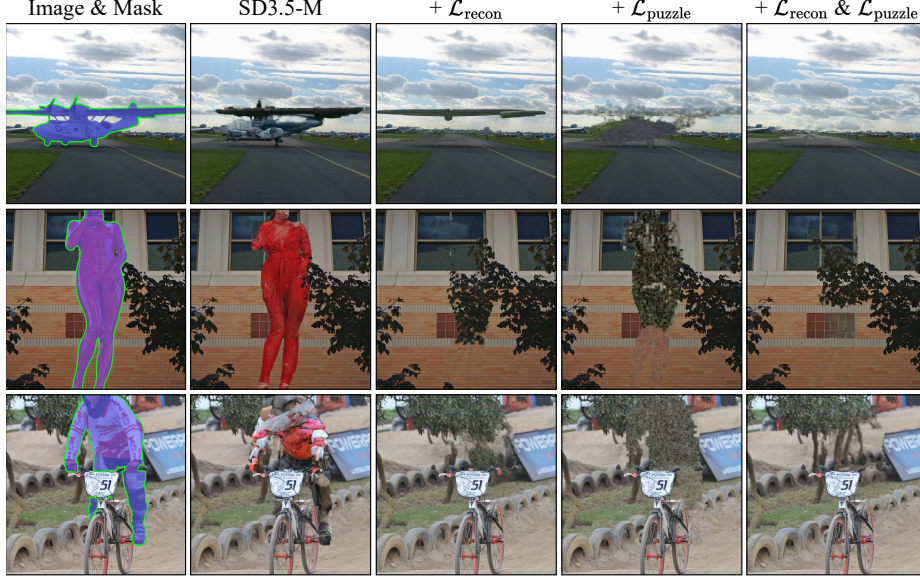


Fig. 10: Visualization of loss components in BRSA (Sec. 3.2). $\mathcal{L}_{\text{recon}}$ preserves background structure and $\mathcal{L}_{\text{puzzle}}$ completes the masked region; using both (*i.e.*, EraseLoRA) produces the most coherent and detailed background.

Table 4: Applicability across different (Left) MLLMs and (Right) Tag2Mask models.

Method	MLLM	Metrics		Method	Tag2Mask Model	Metrics	
		BG Sim.(↑)	FG Sim.(↓)			BG Sim.(↑)	FG Sim.(↓)
SD3.5-M [8]	N/A	0.605	0.286	SD3.5-M [8]	N/A	0.605	0.286
+ EraseLoRA	LLaVA-7B [22]	0.728	0.164	+ EraseLoRA	Seg4Diff [17]	0.666	0.205
+ EraseLoRA	LLaVA-7B+MARINE [45]	0.727	0.161	+ EraseLoRA	YOLOE [35]	0.709	0.177
+ EraseLoRA	Qwen2.5-VL-72B [2]	0.726	0.165	+ EraseLoRA	SAM3 [4]	0.708	0.177
+ EraseLoRA	InternVL3-78B [48]	0.743	0.151	+ EraseLoRA	G.SAM2 [29]	0.743	0.151

We also evaluate EraseLoRA across MLLMs of various scales [2, 22, 45, 48] from 7B to 78B parameters, including a hallucination-mitigated variant [45]. Across all these models, EraseLoRA provides consistent gains (Tab. 4, left), confirming that our framework reliably leverages the background-aware reasoning ability of MLLMs rather than depending on a specific architecture. Notably, even lightweight 7B MLLMs yield substantial improvements comparable to much larger models [2], achieving gains of up to 20.3% in BG Sim. and 42.7% in FG Sim., which indicates that EraseLoRA remains highly effective in resource-constrained settings without requiring a large MLLM. Despite their effectiveness, we adopt a large-scale MLLM [48] as our default configuration in the main experiments, as it provides the strongest background-aware reasoning and achieves the best overall performance.

We further validate EraseLoRA’s effectiveness across different Tag2Mask models [4, 17, 29, 35] in BFE (Sec. 3.1). Across all Tag2Mask variants, EraseLoRA

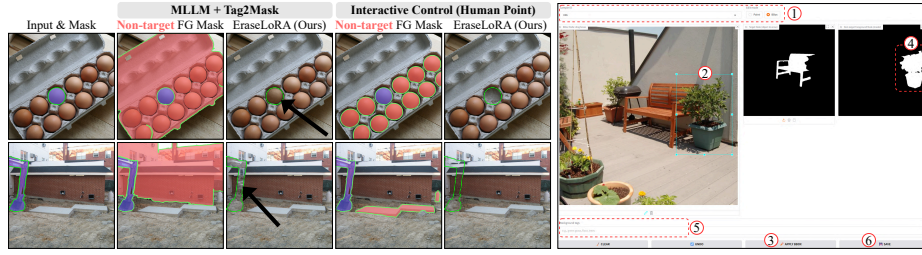


Fig. 11: Interactive Control. (Right) Our interactive interface allows users to generate customized non-target foreground masks for BFE (Sec. 3.1) based on manual masks for foreground exclusion and background tag selection effectively correct challenging failure cases of EraseLoRA.

consistently improves background reconstruction and foreground suppression, yielding at least 10.0% gains in BG Sim. and 28.3% reductions in FG Sim. over the SD3.5-M [8] baseline (see Tab. 4, right). Grounded SAM2 achieves the best performance, improving BG Sim. by up to 22.8% and reducing FG Sim. by up to 47.2%, resulting in the cleanest and most faithful background reconstruction. Additional qualitative results and detailed analysis regarding the impact of MLLM hallucinations are provided in Appendix D.1 and D.3.

MLLM-guided removal under occlusion. Occlusion poses a unique challenge for object removal: content normally regarded as foreground may be hidden behind the target, yet must be recovered as background during inpainting. By leveraging the MLLM’s scene-level understanding, EraseLoRA correctly identifies such occluded elements and reconstructs them with semantically coherent completions (see Fig. 9, right).

Interactive Control. While EraseLoRA effectively removes the target automatically via MLLM’s background-aware reasoning [48], we offer an interactive variant where users provide minimal guidance via an interactive interface to enhance user control and computational efficiency. This mode allows users to bypass MLLM overhead or correct potential imperfect predictions of its reasoning. To support these needs, our interactive variant follows a streamlined workflow (Fig. 11, right): (i) selecting a specific sample and the edit mode (*e.g.*, Point or BBox prompts), (ii) providing visual guidance via points or bounding boxes on the input image, (iii) clicking the apply button to (iv) update and refine the precise non-target foreground mask by observing the extracted results, (v) generating descriptive background tags based on the input image, and (vi) clicking the save button to store the manual results for BRSA (Sec. 3.2).

Limitations. While EraseLoRA achieves strong performance in object removal, several limitations remain. First, the use of large-scale MLLMs such as InternVL3-78B [48] and test-time optimization for fixed 500 iterations with LoRA adapters introduces additional computational overhead. Specifically, the full optimization process takes approximately three minutes per test image on

the default backbone, SD3.5-M [8]. Although no paired data are required, iterative optimization and MLLM queries increase latency and memory usage at inference time. However, this cost can be significantly mitigated by (i) employing lightweight MLLMs (*e.g.*, 7B scale), which show comparable quality as shown in Tab. 4, and (ii) applying an early stopping strategy. We observe that employing an early stopping strategy can reduce the average optimization cost to approximately 141 iterations, achieving a more than $3.5\times$ speedup compared to the fixed 500-iterations baseline while preserving over 96.5% of the background similarity (see Appendix E.1 for more details). Second, the MLLM-guided background definition can be imperfect in complex scenes with subtle semantic boundaries or heavy occlusion, which may lead to incomplete or inaccurate removal. A more detailed analysis is provided in Appendix D.3.

5 Conclusion

In this paper, we introduce EraseLoRA, a dataset-free object removal framework that leverages MLLM-guided background-aware reasoning and test-time adaptation to enable faithful background restoration. Building on the failure modes identified in existing dataset-free methods, EraseLoRA explicitly separates target, non-target foreground, and background, filters out distractors to prevent their regeneration, and integrates multiple background subtypes. As a plug-and-play and model-agnostic module, EraseLoRA consistently produces cleaner and more coherent background reconstructions across diffusion backbones and benchmarks without requiring any paired data. An interesting future direction is video object removal, where shared background context across frames can amortize the per-image optimization cost.

Acknowledgements

This work was partly supported by the KHIDI grant funded by the Korean government (MOHW) [No.RS-2025-02307233], the NRF or IITP grants funded by the Korean government (MSIT) [No.05-26-04-0094, No.RS-2026-25472075, No.RS-2026-25483206, No.RS-2025-02305581, No.RS-2025-25442338, and No.RS-2021-II211343], the ITIP grant funded by the Korean government (MOTIR) [No.RS-2026-25549946], the Research grant from SNU, and the Strategic Hub grant for International Research Collaboration of SNU.

Kyungsu Kim is affiliated with the School of Transdisciplinary Innovations, Department of Biomedical Science, Interdisciplinary Program in Artificial Intelligence (IPAI), Medical Research Center, and AI Institute at SNU.

References

1. 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)
2. Bai, S., Chen, K., Liu, X., Wang, J., Ge, W., Song, S., Dang, K., Wang, P., Wang, S., Tang, J., et al.: Qwen2.5-VL technical report. arXiv preprint arXiv:2502.13923 (2025)
3. Black Forest Labs: Flux. <https://github.com/black-forest-labs/flux> (2024), accessed 2026-06-26
4. Carion, N., Gustafson, L., Hu, Y.T., Debnath, S., Hu, R., Coll-Vinent, D.S., Ryali, C., Alwala, K.V., Khedr, H., Huang, A., Lei, J., Ma, T., Guo, B., Kalla, A., Marks, M., Greer, J., Wang, M., Sun, P., Rädle, R., Afouras, T., Mavroudi, E., Xu, K., Wu, T.H., Zhou, Y., Momeni, L., Hazra, R., Ding, S., Vaze, S., Porcher, F., Li, F., Li, S., Kamath, A., Cheng, H.K., Dollar, P., Ravi, N., Saenko, K., Zhang, P., Feichtenhofer, C.: SAM 3: Segment anything with concepts. In: ICLR (2026)
5. Caron, M., Touvron, H., Misra, I., Jégou, H., Mairal, J., Bojanowski, P., Joulin, A.: Emerging properties in self-supervised vision transformers. In: ICCV. pp. 9650–9660 (2021)
6. Chen, Z., Wang, W., Yang, Z., Yuan, Z., Chen, H., Shen, C.: FreeCompose: Generic zero-shot image composition with diffusion prior. In: ECCV. pp. 70–87 (2024). https://doi.org/10.1007/978-3-031-72643-9_5
7. Ekin, Y., Yildirim, A.B., Caglar, E.E., Erdem, A., Erdem, E., Dundar, A.: CLIPAway: Harmonizing focused embeddings for removing objects via diffusion models. In: NeurIPS. vol. 37, pp. 17572–17601 (2024). <https://doi.org/10.52202/079017-0559>
8. Esser, P., Kulal, S., Blattmann, A., Entezari, R., Müller, J., Saini, H., Levi, Y., Lorenz, D., Sauer, A., Boesel, F., Podell, D., Dockhorn, T., English, Z., Rombach, R.: Scaling rectified flow transformers for high-resolution image synthesis. In: ICML. vol. 235, pp. 12606–12633 (2024)
9. Fanelli, N., Vessio, G., Castellano, G.: I Dream My Painting: Connecting mlms and diffusion models via prompt generation for text-guided multi-mask inpainting. In: IEEE WACV. pp. 6073–6082 (2025)
10. Fu, Y., Zheng, Y., Dai, Z., Ding, H.: EffectErase: Joint video object removal and insertion for high-quality effect erasing. In: CVPR. pp. 2005–2014 (2026)
11. Goodfellow, I.J., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A., Bengio, Y.: Generative adversarial nets. In: NeurIPS. vol. 27 (2014)
12. Ho, J., Jain, A., Abbeel, P.: Denoising diffusion probabilistic models. In: NeurIPS. vol. 33, pp. 6840–6851 (2020)
13. Hu, E.J., Shen, Y., Wallis, P., Allen-Zhu, Z., Li, Y., Wang, S., Wang, L., Chen, W.: LoRA: Low-rank adaptation of large language models. In: ICLR (2022)
14. Jia, Y., Cheng, A., Yuan, Y., Wang, C., Li, J., Jia, H., Zhang, S.: DesignEdit: Unify spatial-aware image editing via training-free inpainting with a multi-layered latent diffusion framework. In: AAAI. vol. 39, pp. 3958–3966 (2025). <https://doi.org/10.1609/aaai.v39i4.32414>
15. Jiang, L., Wang, Z., Bao, J., Zhou, W., Chen, D., Shi, L., Chen, D., Li, H.: SmartEraser: Remove anything from images using masked-region guidance. In: CVPR. pp. 24452–24462 (2025)
16. Kim, B.S., Kim, J., Ye, J.C.: Chain-of-Zoom: Extreme super-resolution via scale autoregression and preference alignment. In: NeurIPS. vol. 38 (2025)

17. Kim, C., Shin, H., Hong, E., Yoon, H., Arnab, A., Seo, P.H., Hong, S., Kim, S.: Seg4Diff: Unveiling open-vocabulary segmentation in text-to-image diffusion transformers. In: NeurIPS. vol. 38 (2025)
18. Kim, J., Lee, Z., Cho, D., Jo, S., Jung, Y., Kim, K., Yang, E.: Early timestep zero-shot candidate selection for instruction-guided image editing. In: ICCV. pp. 18844–18854 (2025)
19. Kuznetsova, A., Rom, H., Alldrin, N., Uijlings, J., Krasin, I., Pont-Tuset, J., Kamali, S., Popov, S., Mallocci, M., Kolesnikov, A., Duerig, T., Ferrari, V.: The Open Images Dataset V4: Unified image classification, object detection, and visual relationship detection at scale. IJCV **128**, 1956–1981 (2020). <https://doi.org/10.1007/s11263-020-01316-z>
20. Li, P., Nie, Q., Chen, Y., Jiang, X., Wu, K., Lin, Y., Liu, Y., Peng, J., Wang, C., Zheng, F.: Tuning-free image customization with image and text guidance. In: ECCV. pp. 233–250 (2024). https://doi.org/10.1007/978-3-031-73116-7_14
21. Li, W., Lin, Z., Zhou, K., Qi, L., Wang, Y., Jia, J.: MAT: Mask-aware transformer for large hole image inpainting. In: CVPR. pp. 10758–10768 (2022)
22. Liu, H., Li, C., Wu, Q., Lee, Y.J.: Visual instruction tuning. In: NeurIPS. vol. 36, pp. 34892–34916 (2023)
23. Liu, S., Zeng, Z., Ren, T., Li, F., Zhang, H., Yang, J., Jiang, Q., Li, C., Yang, J., Su, H., Zhu, J., Zhang, L.: Grounding DINO: Marrying DINO with grounded pre-training for open-set object detection. In: ECCV. pp. 38–55 (2024). https://doi.org/10.1007/978-3-031-72970-6_3
24. Liu, Y., Zhou, H., Cui, B., Shang, W., Lin, R.: Erase Diffusion: Empowering object removal through calibrating diffusion pathways. In: CVPR. pp. 2418–2427 (2025)
25. Manukyan, H., Sargsyan, A., Atanyan, B., Wang, Z., Navasardyan, S., Shi, H.: HD-Painter: High-resolution and prompt-faithful text-guided image inpainting with diffusion models. In: ICLR. pp. 96301–96330 (2025)
26. Podell, D., English, Z., Lacey, K., Blattmann, A., Dockhorn, T., Müller, J., Penna, J., Rombach, R.: SDXL: Improving latent diffusion models for high-resolution image synthesis. In: ICLR. pp. 1862–1874 (2024)
27. Qu, L., Li, H., Wang, W., Liu, X., Li, J., Nie, L., Chua, T.S.: SILMM: Self-improving large multimodal models for compositional text-to-image generation. In: CVPR. pp. 18497–18508 (2025)
28. Ravi, N., Gabeur, V., Hu, Y.T., Hu, R., Ryali, C., Ma, T., Khedr, H., Rädle, R., Rolland, C., Gustafson, L., Mintun, E., Pan, J., Alwala, K.V., Carion, N., Wu, C.Y., Girshick, R., Dollár, P., Feichtenhofer, C.: SAM 2: Segment anything in images and videos. In: ICLR (2025)
29. Ren, T., Liu, S., Zeng, A., Lin, J., Li, K., Cao, H., Chen, J., Huang, X., Chen, Y., Yan, F., Zeng, Z., Zhang, H., Li, F., Yang, J., Li, H., Jiang, Q., Zhang, L.: Grounded SAM: Assembling open-world models for diverse visual tasks. arXiv preprint arXiv:2401.14159 (2024)
30. Rombach, R., Blattmann, A., Lorenz, D., Esser, P., Ommer, B.: High-resolution image synthesis with latent diffusion models. In: CVPR. pp. 10684–10695 (2022)
31. Sagong, M.C., Yeo, Y.J., Jung, S.W., Ko, S.J.: RORD: A real-world object removal dataset. In: BMVC. p. 542 (2022)
32. Sargsyan, A., Navasardyan, S., Xu, X., Shi, H.: MI-GAN: A simple baseline for image inpainting on mobile devices. In: ICCV. pp. 7335–7345 (2023)
33. Sun, W., Dong, X.M., Cui, B., Tang, J.: Attentive Eraser: Unleashing diffusion model’s object removal potential via self-attention redirection guidance. In: AAAI. vol. 39, pp. 20734–20742 (2025). <https://doi.org/10.1609/aaai.v39i19.34285>

34. Suvorov, R., Logacheva, E., Mashikhin, A., Remizova, A., Ashukha, A., Silvestrov, A., Kong, N., Goka, H., Park, K., Lempitsky, V.: Resolution-robust large mask inpainting with fourier convolutions. In: IEEE WACV. pp. 2149–2159 (2022)
35. Wang, A., Liu, L., Chen, H., Lin, Z., Han, J., Ding, G.: YOLOE: Real-time seeing anything. In: ICCV. pp. 24591–24602 (2025)
36. Wang, D., Shelhamer, E., Liu, S., Olshausen, B., Darrell, T.: Tent: Fully test-time adaptation by entropy minimization. In: ICLR (2021)
37. Wang, Z., Li, A., Li, Z., Liu, X.: GenArtist: Multimodal llm as an agent for unified image generation and editing. In: NeurIPS. vol. 37, pp. 128374–128395 (2024)
38. Wang, Z., Bovik, A.C., Sheikh, H.R., Simoncelli, E.P.: Image quality assessment: from error visibility to structural similarity. IEEE TIP **13**(4), 600–612 (2004). <https://doi.org/10.1109/TIP.2003.819861>
39. Wei, R., Yin, Z., Zhang, S., Zhou, L., Wang, X., Ban, C., Cao, T., Sun, H., He, Z., Liang, K., Ma, Z.: OmniEraser: Remove objects and their effects in images with paired video-frame data. arXiv preprint arXiv:2501.07397 (2025)
40. Xie, S., Zhang, Z., Lin, Z., Hinz, T., Zhang, K.: SmartBrush: Text and shape guided object inpainting with diffusion model. In: CVPR. pp. 22428–22437 (2023)
41. Xie, T., Ma, R., Wang, Q., Ye, X., Liu, F., Tai, Y., Zhang, Z., Wang, L., Yi, Z.: Anywhere: A multi-agent framework for user-guided, reliable, and diverse foreground-conditioned image generation. In: AAAI. vol. 39, pp. 7410–7418 (2025). <https://doi.org/10.1609/aaai.v39i7.32797>
42. Yang, B., Gu, S., Zhang, B., Zhang, T., Chen, X., Sun, X., Chen, D., Wen, F.: Paint by Example: Exemplar-based image editing with diffusion models. In: CVPR. pp. 18381–18391 (2023)
43. Zhang, K., Mo, L., Chen, W., Sun, H., Su, Y.: MagicBrush: a manually annotated dataset for instruction-guided image editing. In: NeurIPS. vol. 36, pp. 31428–31449 (2023)
44. Zhao, J., Wang, Z., Yang, P., Zhou, S.: Precise object and effect removal with adaptive target-aware attention. In: CVPR. pp. 19370–19379 (2026)
45. Zhao, L., Deng, Y., Zhang, W., Gu, Q.: Mitigating object hallucination in large vision-language models via image-grounded guidance. In: ICML. vol. 267, pp. 77461–77486 (2025)
46. Zhao, S., Cui, J., Sheng, Y., Dong, Y., Liang, X., Chang, E.I., Xu, Y.: Large scale image completion via co-modulated generative adversarial networks. In: ICLR (2021)
47. Zhou, J., Li, J., Xu, Z., Li, H., Cheng, Y., Hong, F.T., Lin, Q., Lu, Q., Liang, X.: FireEdit: Fine-grained instruction-based image editing via region-aware vision language model. In: CVPR. pp. 13093–13103 (2025)
48. Zhu, J., Wang, W., Chen, Z., Liu, Z., Ye, S., Gu, L., Tian, H., Duan, Y., Su, W., Shao, J., et al.: InternVL3: Exploring advanced training and test-time recipes for open-source multimodal models. arXiv preprint arXiv:2504.10479 (2025)
49. Zhu, T., Zhang, S., Shao, J., Tang, Y.: KV-Edit: Training-free image editing for precise background preservation. In: ICCV. pp. 16607–16617 (2025)
50. Zhu, Y., Zhang, Q., Wang, Y., Nie, Y., Zheng, W.S.: EntityErasure: Erasing entity cleanly via amodal entity segmentation and completion. In: CVPR. pp. 28274–28283 (2025)
51. Zhuang, J., Zeng, Y., Liu, W., Yuan, C., Chen, K.: A task is worth one word: Learning with task prompts for high-quality versatile image inpainting. In: ECCV. pp. 195–211 (2024)

52. Zuo, Z., Zhao, L., Li, A., Wang, Z., Zhang, Z., Chen, J., Xing, W., Lu, D.: Generative image inpainting with segmentation confusion adversarial training and contrastive learning. In: AAAI. vol. 37, pp. 3888–3896 (2023). <https://doi.org/10.1609/aaai.v37i3.25502>