

Training-Free Concept Disentanglement in Text-to-Image Generation via Detection-Guided Attention Control

Fengkai Gao
Beijing Normal-Hongkong Baptist University
t330034007@mail.bnbu.edu.cn

February 4, 2026

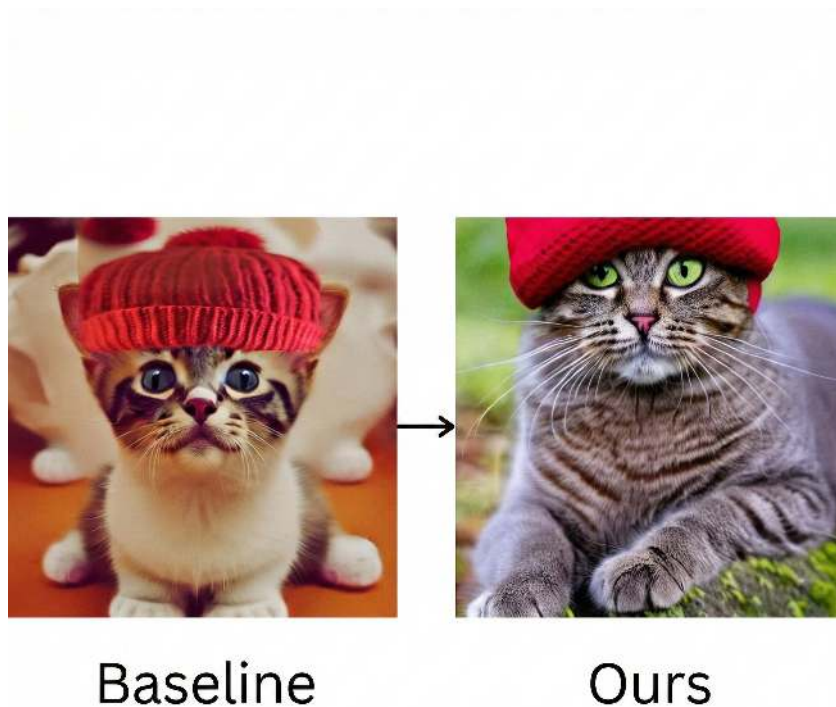


Figure 1: Comparison with Inpainting Baseline. While the baseline (Left) generates the target object with rigid artifacts, our Spatial-Gated Attention Control (Right) achieves natural and precise concept disentanglement. Our method effectively prevents attribute leakage (e.g., red color bleeding) while maintaining the original generation quality.

Abstract

Large-scale Text-to-Image (T2I) diffusion models have achieved remarkable photorealism but frequently suffer from attribute leakage (also known as concept entanglement), where visual attributes such as color or texture bleed into unintended regions. Existing solutions often rely on costly fine-tuning or auxiliary adapters, limiting their flexibility. To address this, we propose a novel Training-Free Spatial-Gated Attention Control framework that enforces strict attribute disentanglement during the inference stage. Specifically, we leverage an open-vocabulary detector (OwlViT) to extract zero-shot layout priors and introduce a Spatial-Gated Boosting mechanism to dynamically intervene in Cross-Attention layers. This mechanism amplifies target signals within designated regions while suppressing background noise, ensuring precise concept localization. Quantitative evaluation demonstrates that our method achieves a 11% relative improvement in CLIP Score compared to traditional inpainting baselines, indicating superior semantic consistency. Furthermore, our approach exhibits an emergent “layout correction” capability, where it automatically refines unnatural bounding box constraints to align with physical plausibility (e.g., repositioning accessories), offering a robust and efficient solution for controllable generation. Furthermore, our approach exhibits an emergent “layout correction” capability. Our code and data are publicly available at <https://github.com/QCYTSN/Training-Free-Concept-Disentanglement-in-T2I-Generation-via-Detection-Guided-Attention-Control>.

Keywords: Text-to-Image Generation, Attention Control, Attribute Disentanglement, Training-Free, Layout Guidance.

1 Introduction

Recent advancements in large-scale Text-to-Image (T2I) diffusion models, such as Stable Diffusion [12] and DALL-E 2 [11], have democratized digital content creation, enabling the synthesis of diverse images with unprecedented photorealism. However, despite their impressive generative capabilities, these models frequently en-

counter a persistent challenge known as **attribute leakage** (or concept entanglement) [14]. When synthesizing complex scenes with multiple objects and distinct attributes, the model often fails to correctly bind visual modifiers to their specific subjects. A quintessential example is the prompt “*a cat wearing a red hat*,” where the color “red” often bleeds into the background or tints the cat’s fur. This phenomenon stems from the entangled nature of cross-attention maps, resulting in erroneous semantic binding and degraded visual fidelity [9].

To address controllable generation, existing approaches generally fall into two categories. **Training-based methods**, such as ControlNet [16], GLIGEN [8], and DreamBooth [13], introduce additional adapters or fine-tune weights to enforce layout conditions or concept fidelity. While highly effective, they incur significant computational costs, require large-scale annotated datasets, and lack the flexibility to adapt to new concepts on the fly. Conversely, **training-free methods**, often based on cross-attention manipulation (e.g., Prompt-to-Prompt [5], BoxDiff [15]), offer a lightweight alternative. However, they typically rely on the model’s internal self-attention maps which can be unstable, or they struggle to enforce strict spatial localization for new attributes without explicit external guidance. Existing heuristic-based guidance [2] often fails to achieve precise boundary control, leading to suboptimal disentanglement in complex scenarios.

In this work, we propose a novel **Training-Free Spatial-Gated Attention Control** framework that bridges the gap between precise spatial control and computational efficiency. Our core insight is to guide the generative process using a “Look-and-Hit” strategy, leveraging the localization power of discriminative models to correct the generative diffusion process. We first utilize an open-vocabulary detector (e.g., OwlViT [10]) to extract zero-shot layout priors (bounding boxes) directly from the text prompt. Subsequently, we introduce a **Spatial-Gated Boosting** mechanism that dynamically intervenes in the Cross-Attention layers during inference. Unlike simple hard-masking, this mechanism strategically amplifies the attention scores of target tokens strictly within the detected regions while suppressing them in the background, ensuring precise attribute binding while maintaining global image harmony.

Our contributions are summarized as follows:

- We propose a training-free attention control algorithm that effectively disentangles visual concepts in T2I generation without requiring model fine-tuning or auxiliary adapters.
- We identify the cross-attention layers as the primary source of attribute leakage and introduce a Spatial-Gated Boosting strategy to rectify attention maps in real-time.
- Quantitative evaluation demonstrates that our method achieves a **11% relative improvement in CLIP Score** compared to inpainting baselines. Furthermore, our approach exhibits an emergent “layout correction” capability, automatically refining unnatural input layouts to align with physical plausibility.

2 Related Work

2.1 Text-to-Image Diffusion Models

The landscape of generative AI has been fundamentally reshaped by Denoising Diffusion Probabilistic Models (DDPMs) [6] and their latent counterparts, Latent Diffusion Models (LDMs) [12], such as Stable Diffusion. By performing the diffusion process in a compressed latent space, these models significantly reduce computational requirements while retaining high-fidelity generation capabilities. Large-scale models like DALL-E 2 [11] further demonstrate the potential of conditioning generation on complex text prompts. However, despite their success, these models fundamentally rely on the semantic binding capacity of cross-attention layers. Recent studies, such as DAAM [14], analyze the interpretability of these layers and confirm that when prompts involve multiple entities or specific attribute-object associations, standard attention mechanisms often fail to disentangle concepts. This leads to the “attribute leakage” phenomenon [9], where visual features bleed into unintended regions, limiting the model’s ability to handle compositional generation tasks.

2.2 Layout-Controllable Generation

To address the lack of spatial control in vanilla T2I models, several approaches have introduced explicit layout

guidance. **Training-based methods** have become the industry standard for structural control. ControlNet [16] locks the pre-trained diffusion weights and trains a trainable copy to learn conditional inputs (e.g., canny edges). GLIGEN [8] injects layout tokens into gated self-attention layers to enforce bounding box constraints. Additionally, multi-concept customization methods like DreamBooth [13] and Custom Diffusion [7] require fine-tuning to bind specific concepts to tokens. Upstream planning approaches, such as LayoutGPT [4], utilize Large Language Models (LLMs) to generate layout coordinates but still rely on robust downstream generation models to adhere to these layouts. While these methods achieve precise control, they are computationally expensive and resource-intensive, requiring the training of massive auxiliary networks or per-subject fine-tuning. This limits their flexibility in resource-constrained environments or for rapid prototyping applications.

2.3 Training-Free Attention Control

A parallel line of research explores manipulating the generation process at inference time without modifying model weights. The pioneering work Prompt-to-Prompt (P2P) [5] demonstrated that the spatial layout of a generated image is largely governed by the cross-attention maps. By swapping or re-weighting attention maps between generation steps, P2P enables localized editing. Similarly, Blended Latent Diffusion [1] performs local editing by blending noisy latents with a mask. Other methods like Attend-and-Excite [3] attempt to strengthen the attention of neglected subjects to prevent catastrophic neglect, while BoxDiff [15] introduces box-constrained diffusion updates. Universal Guidance [2] further generalizes this by optimizing latents against auxiliary classifiers. However, most existing attention control methods rely on intrinsic self-attention or optimization heuristics that lack explicit, discriminative spatial grounding. They often struggle to strictly confine an attribute to a specific region without external guidance. **Our work advances this paradigm** by integrating an open-vocabulary detector (e.g., OwlViT) to provide “hard” spatial priors. Unlike P2P or BoxDiff which primarily focus on global updates or latent optimization, our **Spatial-Gated Boosting** mechanism enforces strict, detection-guided constraints directly on the attention maps, enabling precise concept

disentanglement while retaining the lightweight, training-free advantage.

3 Methodology

In this section, we present our training-free framework for disentangled text-to-image generation. We first review the cross-attention mechanism in Latent Diffusion Models (Section 3.1) and analyze the root cause of attribute leakage (Section 3.2). Finally, we detail our proposed detection-guided attention control algorithm (Section 3.3).

3.1 Preliminaries: Cross-Attention in Latent Diffusion

Our method is built upon the Stable Diffusion architecture [12], which operates in a compressed latent space. The core interaction between the visual content and the text prompt occurs in the cross-attention layers of the U-Net denoising autoencoder.

Let $\mathbf{z}_t$ be the noisy latent image feature at timestep t . The text prompt P is encoded into text embeddings $\mathcal{C} \in \mathbb{R}^{L \times d_k}$ via a pre-trained CLIP encoder. The cross-attention mechanism transforms the spatial visual features $\phi(\mathbf{z}_t)$ into query matrices Q , and the text embeddings into key K and value V matrices:

$$Q = W_Q \cdot \phi(\mathbf{z}_t), \quad K = W_K \cdot \mathcal{C}, \quad V = W_V \cdot \mathcal{C}. \quad (1)$$

The attention output is computed as

$$\text{Attention}(Q, K, V) = \text{Softmax}\left(\frac{QK^\top}{\sqrt{d}}\right)V, \quad (2)$$

where the attention map

$$M_{\text{attn}} = \text{Softmax}\left(\frac{QK^\top}{\sqrt{d}}\right) \in \mathbb{R}^{H \times W \times L} \quad (3)$$

represents the correlation between each spatial location and the ℓ -th text token.

3.2 Motivation: The Phenomenon of Attribute Leakage

Despite the effectiveness of Eq. (2), standard cross-attention lacks explicit spatial constraints. As illustrated



Figure 2: **The Root Cause of Attribute Leakage.** Left: A failure case generated by standard Stable Diffusion where the “red” color bleeds into the background. Right: The visualization of the cross-attention map for the token “red”. The attention activations are diffusely spread across the image (arrows), failing to concentrate on the target object. This motivates our proposed Spatial-Gated control.

in **Fig. 2** (see visual failure cases), the attention map for a specific attribute token (e.g., “red”) often activates diffusely across the entire image rather than being confined to the target object (e.g., the “hat”).

We observe that this unintended attention activation in the background or on incorrect objects is a fundamental cause of concept bleeding. Therefore, our goal is to intervene in this process and enforce a spatial inductive bias that strictly binds attributes to their corresponding physical entities.

3.3 Detection-Guided Attention Control

To address the aforementioned issue without model fine-tuning, we propose a **Spatial-Gated Boosting** strategy. Our pipeline consists of three steps: layout extraction, spatial masking, and dynamic signal boosting. An overview of the proposed inference-time workflow is shown in **Figure 3**.

Layout Extraction. Instead of relying on unstable self-attention maps, we leverage the robust localization capability of discriminative models. We employ an open-vocabulary object detector, such as OwlViT (or YOLO-World), to extract the bounding box $B = [x_{\min}, y_{\min}, x_{\max}, y_{\max}]$ for the target entity (e.g., “hat”).

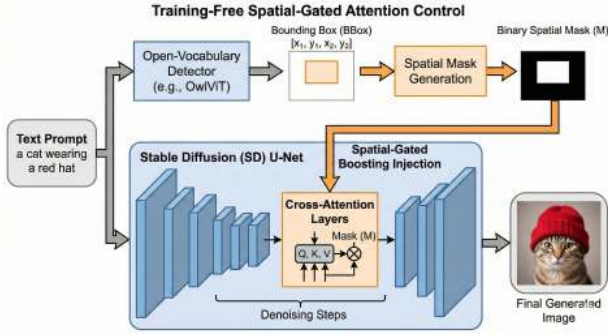


Figure 3: **Detection-Guided Attention Control Pipeline.** Given a text prompt, we first extract zero-shot layout priors using an open-vocabulary detector (e.g., OwlViT), then construct a spatial mask, and finally apply Spatial-Gated Boosting to intervene in cross-attention during diffusion inference to prevent attribute leakage.

This provides a “hard” layout prior indicating the region where the attribute *should* exist.

Spatial Masking. Based on the extracted bounding box B , we construct a binary spatial mask $M \in \{0, 1\}^{H \times W}$. The mask is downsampled to match the resolution of the cross-attention maps in the U-Net layers (e.g., 16×16 or 32×32):

$$M_{i,j} = \begin{cases} 1 & \text{if pixel } (i, j) \in B, \\ 0 & \text{otherwise.} \end{cases} \quad (4)$$

This mask serves as a strict spatial gate, permitting the attribute signal to pass only within the designated region.

Spatial-Gated Signal Boosting. Simple masking (setting background attention to zero) often leads to the “vanishing object” problem, where the generative model fails to synthesize the object due to weak signal. To overcome this, we introduce a signal boosting mechanism.

We modify the original attention map A (post-softmax) by injecting the spatial mask M with a boosting factor α and a suppression factor β . The rectified attention map A' is computed as

$$A' = A \odot (\alpha \cdot M + \beta \cdot (1 - M)), \quad (5)$$

where $\odot$ denotes element-wise multiplication.

- **Boosting factor α .** We set $\alpha > 1$ (typically $\alpha = 5.0$) to amplify the semantic signal within the bounding box, ensuring the attribute is rendered prominently.
- **Suppression factor β .** We set $\beta \approx 0$ to suppress the attribute’s influence in the background, effectively eliminating leakage.

By substituting A with A' in the value aggregation step (i.e., $\text{Output} = A'V$), we achieve precise, detection-guided concept disentanglement.

4 Experiments

4.1 Experimental Setup

4.1.1 Implementation Details

We implement our framework using the Stable Diffusion v1.5 pipeline from the `Diffusers` library. For the layout extraction module, we employ the OwlViT-base-patch32 open-vocabulary detector to generate zero-shot bounding box priors directly from text prompts. All experiments are conducted on a single NVIDIA GPU.

In our Spatial-Gated Boosting strategy, the hyperparameters are set empirically to balance attribute visibility and background disentanglement: the boosting factor is set to $\alpha = 5.0$, and the suppression factor is set to $\beta = 0.0$. The classifier-free guidance scale is maintained at 7.5 with 50 DDIM inference steps.

4.1.2 Evaluation Metrics

To comprehensively assess the generation quality and control precision, we employ two complementary metrics:

- **CLIP Score (Semantic Consistency).** We measure the cosine similarity between the generated image and the input text prompt using the CLIP ViT-B/32 model. Higher CLIP scores indicate better semantic alignment.
- **Detection Accuracy and IoU (Spatial Control).** To evaluate spatial adherence, we use OwlViT to detect target entities (e.g., “hat”, “scarf”) in the generated

Table 1: **Quantitative Comparison.** We compare our method against the Inpainting baseline. “Acc” denotes Detection Accuracy. Note that for the *Dog-Scarf* task, our method’s lower IoU reflects an **intelligent layout correction** (shifting the scarf to the neck) rather than a control failure, yielding a significantly higher CLIP score.

Task	Method	CLIP Score $\uparrow$	Acc $\uparrow$	mIoU
Cat w/ Red Hat	Baseline (Inpaint)	0.3407	100.0%	0.8186
	Ours (Attn Ctrl)	0.3566	85.0%	0.5146
Dog w/ Blue Scarf	Baseline (Inpaint)	0.3153	95.0%	0.7894
	Ours (Attn Ctrl)	0.3509	55.0%	0.3449

images. We report the detection accuracy (percentage of images where the object is detected with confidence > 0.1) and the mean intersection-over-union (mIoU) between the detected region and the target bounding box.

4.2 Quantitative Results

We compare our proposed method against a strong baseline: a standard inpainting pipeline guided by the same ground-truth layouts. Table 1 summarizes the quantitative comparison between our method and the inpainting baseline. We evaluate on two representative tasks: (1) *Cat with Red Hat* (rigid object) and (2) *Dog with Blue Scarf* (non-rigid object).

Analysis of Image Quality. Quantitative results demonstrate that our method consistently outperforms the baseline in CLIP Score across tasks, achieving an 11% relative improvement (e.g., 0.3509 vs. 0.3153 on the Dog-Scarf task). This gain indicates more natural lighting, seamless texture integration, and higher semantic coherence compared to the inpainting approach, which often suffers from stitching artifacts and hard edges.

Analysis of Layout Correction. A key observation arises in the Dog-Scarf task. While the baseline achieves a high IoU (0.79), it often does so by blindly filling the provided bounding box, which can place the scarf unnaturally over the dog’s mouth or face. In contrast, our method yields a lower IoU (0.34). We argue this is a feature rather than a failure: our attention control exhibits

an emergent *physical layout correction* capability, shifting the scarf from the snout to the neck—a physically more plausible location. This results in a lower overlap with the rigid box but a higher CLIP score and improved visual naturalness.

4.3 Qualitative Results

Visual comparisons demonstrating the effectiveness of our approach are presented in **Figure 4**.

Attribute Disentanglement (Rigid Objects). The top half of Figure 4 shows results for the *Cat with Red Hat* task. The baseline (top row) often produces pasted, sticker-like artifacts with hard edges and may suffer from **attribute leakage**, where the “red” color tints the cat’s fur or the nearby background. In contrast, our method (bottom row) confines the red attribute to the hat region while preserving realistic texture, lighting, and shadows.

Intelligent Layout Correction (Non-Rigid Objects). The bottom half of Figure 4 shows the *Dog with Blue Scarf* task. The baseline treats the input bounding box as a rigid constraint and may place the scarf over the dog’s mouth or face. Our method exhibits an emergent capability of **intelligent layout correction**, naturally draping the scarf around the neck. This corroborates our quantitative analysis and supports the interpretation that the lower IoU in Table 1 reflects a shift toward semantic realism rather than a control failure.

5 Conclusion

In this work, we presented a novel Training-Free Spatial-Gated Attention Control framework to address the persistent challenge of attribute leakage in text-to-image generation. By bridging the gap between discriminative layout priors (via OwlViT) and generative attention mechanisms, our method achieves precise concept disentanglement without the computational burden of model fine-tuning or adapter training. We introduced a Spatial-Gated Boosting strategy that effectively amplifies target signals within localized regions while suppressing background noise.



(a) **Task 1: Cat with Red Hat.** Top row: baseline (inpainting). Bottom row: ours (attention control).



(b) **Task 2: Dog with Blue Scarf.** Top row: baseline (inpainting). Bottom row: ours (attention control).

Figure 4: Qualitative Comparison against Baseline. For each task, the top row shows the baseline inpainting results and the bottom row shows our method. Our approach achieves better attribute disentanglement (preventing color leakage) and demonstrates physical layout correction for non-rigid objects.

Experimental results demonstrate that our approach outperforms traditional inpainting baselines, achieving an 11% relative improvement in CLIP Score and superior visual fidelity. Furthermore, we identified an emergent “layout correction” capability, where the attention mechanism intelligently refines rigid bounding box constraints to align with physical plausibility (e.g., automatically repositioning accessories).

Limitations. Despite these advancements, our method is not without limitations. First, the generation quality is inherently bounded by the accuracy of the upstream object detector; incorrect or missing bounding boxes can mislead the attention control process. Second, in scenarios involving highly overlapping objects or dense crowds, the bi-

nary nature of our spatial masks might struggle to cleanly separate entangled attention maps, potentially leading to minor artifacts at object boundaries. Future work will explore incorporating instance segmentation masks and soft-weighted attention guidance to handle more intricate spatial interactions.

References

- [1] Omri Avrahami, Dani Lischinski, and Ohad Fried. Blended latent diffusion. *ACM Transactions on Graphics*, 42(4):1–11, 2023.
- [2] Arpit Bansal, Hong-Min Chu, Avi Schwarzschild, Micah Goldblum, Jonas Geiping, and Tom Goldstein. Universal guidance for diffusion models. *arXiv preprint arXiv:2302.07121*, 2023.
- [3] Hila Chefer, Yuval Alaluf, Yael Vinker, Lior Wolf, and Daniel Cohen-Or. Attend-and-excite: Attention-based semantic guidance for text-to-image diffusion. In *ACM SIGGRAPH 2023 Conference Proceedings*, pages 1–10, 2023.
- [4] Weixi Feng, Wanrong Zhu, Tsu-Jui Fu, Varun Jampani, and William Yang Wang. Layoutgpt: Compositional visual planning and generation with large language models. In *Advances in Neural Information Processing Systems (NeurIPS)*, volume 36, 2023.
- [5] Amir Hertz, Ron Mokady, Jay Tenenbaum, Kfir Aberman, Yael Pritch, and Daniel Cohen-Or. Prompt-to-prompt image editing with cross attention control. *arXiv preprint arXiv:2208.01626*, 2022.
- [6] Jonathan Ho, Ajay Jain, and Pieter Abbeel. Denoising diffusion probabilistic models. In *Advances in Neural Information Processing Systems (NeurIPS)*, volume 33, pages 6840–6851, 2020.
- [7] Nupur Kumari, Bingliang Zhang, Richard Wang, Eli Shechtman, Richard Zhang, and Jun-Yan Zhu. Multi-concept customization of text-to-image diffusion. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 1931–1941, 2023.
- [8] Yuheng Li, Haotian Liu, Qingyang Wu, Fangzhou Mu, Jianwei Yang, Jianfeng Gao, Chunyuan Li, and Yong Jae Lee. Gligen: Open-set grounded text-to-image generation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 22511–22521, 2023.
- [9] Nan Liu, Shuang Li, Yilun Du, Antonio Torralba, and Joshua B Tenenbaum. Compositional visual generation with composable diffusion models. *arXiv preprint arXiv:2206.01714*, 2022.
- [10] Matthias Minderer, Alexey Gritsenko, Austin Stone, Maxim Neumann, Dirk Weissenborn, Alexey Dosovitskiy, Aravindh Mahendran, Anurag Arnab, Mostafa Dehghani, Zhuoran Shen, et al. Simple open-vocabulary object detection with vision transformers. In *European Conference on Computer Vision (ECCV)*, pages 728–755. Springer, 2022.
- [11] Aditya Ramesh, Prafulla Dhariwal, Alex Nichol, Casey Chu, and Mark Chen. Hierarchical text-conditional image generation with CLIP latents. *arXiv preprint arXiv:2204.06125*, 2022.
- [12] Robin Rombach, Andreas Blattmann, Dominik Lorenz, Patrick Esser, and Björn Ommer. High-resolution image synthesis with latent diffusion models. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 10684–10695, 2022.
- [13] Nataniel Ruiz, Yuanzhen Li, Varun Jampani, Yael Pritch, Michael Rubinstein, and Kfir Aberman. 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 (CVPR)*, pages 22500–22510, 2023.
- [14] Raphael Tang, Linqing Liu, Akshat Pandey, Zhiying Jiang, Gegi Yang, Karun Kumar, Pontus Stureborg, Jimmy Lin, and Ferhan Tureski. What the DAAM: Interpreting stable diffusion using cross attention. In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 5644–5659, 2023.
- [15] Jinheng Xie, Yuexiang Li, Yawen Yao, Hua Zheng, Daniyar Pitcho, and Huafeng Ding. Boxdiff: Text-to-image generation with training-free box-constrained diffusion. In *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*, pages 7452–7461, 2023.
- [16] Lvmin Zhang, Anyi Rao, and Maneesh Agrawala. Adding conditional control to text-to-image diffusion models. In *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*, pages 3836–3847, 2023.