

CountCluster: 用于文本到图像生成的对象数量指定交叉注意力聚合

Joohyeon Lee, Jin-Seop Lee, Jee-Hyong Lee*

Sungkyunkwan University, Suwon, South Korea
{leeju001, wlstjq0602, john}@skku.edu

Abstract

基于扩散的文本到图像生成模型在对象数量和多义性方面表现出色。然而，它们仍然难以生成准确反映输入提示中指定数量的对象的图像。已提出了各种方法，有些方法依赖于外部预训练模型进行迭代优化，或依赖于已定义的交叉注意力或对象在特征中得出的数量表示。然而，它们在准确反映指定对象数量方面仍然存在局限性，并且忽略了一些重要的对象特性——生成图像中的对象示例数量在去噪过程的早期步骤中很大程度上已被确定。为了正确反映图像生成的对象数量，早期步骤中对象交叉注意力的高激活区域与输入对象数量相匹配，同时每个区域得分。为了解决这个问题，我们提出了 CountCluster，一种根据输入中指定的对象数量对对象交叉注意力进行聚合的方法，无需依赖任何外部工具或额外的步骤。该方法在推理时根据注意力得分对对象交叉注意力得分进行 k 聚，定义一种理想的分布，其中每个聚在空间上良好分布，并优化在推理时以此分布为目标。与现有方法相比，我们的方法在对象数量准确性上平均提高了 18.5%。我们在各提示中展示了模型的定量控制性能。代码在：<https://github.com/JoohyeonL22/CountCluster> 中。

介绍

生成模型在文本到图像生成领域取得了显著的进展，其中生成内容完全自然语言描述中合成。尽管早期方法基于 GANs (Goodfellow et al. 2014; Radford, Metz, and Chintala 2016; Karras, Laine, and Aila 2019)，最近扩散模型 (Ho, Jain, and Abbeel 2020; Dhariwal and Nichol 2021; Rombach et al. 2022) 在对象数量和多义性方面表现出更优越的性能。随着扩散模型的进展，各种应用包括图像合成 (Gu et al. 2022; Saharia et al. 2022; Ramesh et al. 2021)、修图 (Meng et al. 2022; Lugmayr et al. 2022; Nichol et al. 2022; Couairon et al. 2023) 和图像编辑 (Kawar et al. 2023; Tumanyan et al. 2023; Cao et al. 2023) 在现实场景中越来越广泛地被采用。一种突出的生成需求是反映用户意图的图像（包括指定的对象数量）的日益增长的需求。然而，扩散模型在生成准确反映输入文本中描述的每个对象的精确数量的图像方面仍然存在困难（例如，“四只柠檬的照片”）(Paiss et al. 2023; Hu et al. 2023; Wen et al. 2023)。

最近，已进行了一些研究以解决生成图像中反映物

*Corresponding author

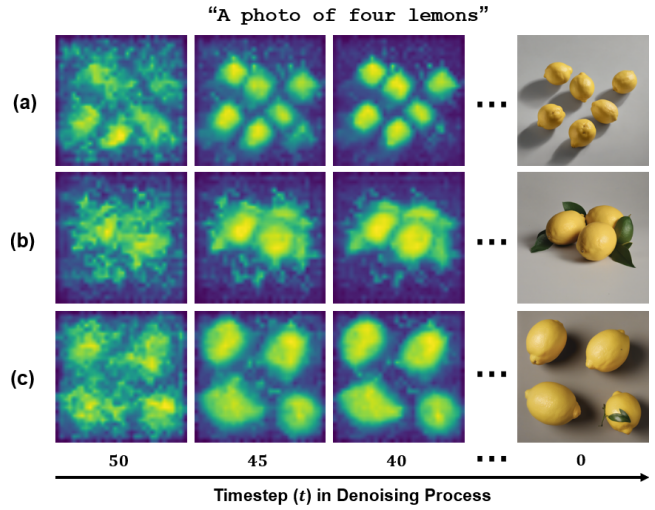


Figure 1: 使用 SDXL (Podell et al. 2024) 生成图像以提示“照片中有四只柠檬”。左列显示了去噪过程中步骤 $t = 50, 45, 40$ 的“交叉注意力”，而右列显示了步骤 $t = 0$ 的最终生成图像。(a)、(b) 和 (c) 是在相同条件下使用不同种子生成的结果。其中，(a) 和 (b) 是示例数量不正确的失败案例，而 (c) 是示例数量正确的成功案例。可以观察到，示例的位置和数量大多在去噪过程的早期步骤中确定。

法 (Binyamin et al. 2024; Kang et al. 2025) 在图像生成过程中对物体示例的数量，通过重新生成图像或迭代优化图像重新指定数量进行改进。基于数量表示引导的方法 (Zafar, Wolf, and Schwartz 2024; Zhang, Chen, and Lee 2023) 通过定义新符号或数量表示或提取控制示例数量的潜在向量中的数量信息。然后，他们利用数量表示引导其他物体生成相同数量的示例。有些方法显示出相对于一些预定义的扩散方法 (Rombach et al. 2022; Podell et al. 2024) 的改进，然而，它们仍然难以生成具有准确数量物体的图像。重要的是，它们忽略了数量控制中最基本的因素，即生成图像中示例数量主要是在扩散过程的早期步骤中确定的。

对于扩散模型而言，在对象跨注意力中高度激活的区域——某些区域表明相关对象之间的强烈关注——生成图像中对象的位置密切相关，如 (Hertz et al. 2023) 所示。如图 1 中的“lemons”跨注意力所示，

生成图像的早期步骤中,物体的位置和数量基本已确定。早期步骤中物体跨注意力图的高度激活区域的簇超入提示中指定的物数量,生成的图像往往包含更多的物体,如 fig. 1 (a)所示。相反,激活区域之的边界模糊且重,在生成过程中可能合并一些物体,致生成的物体例少,如 fig. 1 (b)所示。因此,保生成像中的物体数量反映文本提示中指定的数量,早期步骤中物体跨注意力图中高度激活区域的簇对准地期望物体相匹配,且每区域也充分,如 fig. 1 (c)所示。

此外,有的方法 (Binyamin et al. 2024; Kang et al. 2025; Zafar, Wolf, and Schwartz 2024; Zhang, Chen, and Lee 2023) 依赖于物体和布局生成等外部工具,或者需要数量表示物例信息。因此,生大量的算成本和,使量控制技的部署具有挑性。因此,有必要在不需外模和的情下数量信息。

在本文中,我提出了CountCluster,是一在生成正数量的物例的同,最大限度地少推理成本和的方法,方法不依赖于任何外部工具或外的。所提出的方法通聚物的交叉注意力指去程,以反映期的例量。

一目,我的方法旨在保在去程中,物体交叉注意力图中高度激活的区域足分,匹配入提示中指定的例量。此,我基于物体交叉注意力图中的注意力得分形成 k 聚,定一目分布,目分布鼓些聚在空上良好分。然后,我每一失函,引每聚的分布目分布,在推理程中基于目用无的在化。最,交叉注意力图中高度激活的区域根据定量分,而生成具有正例量的像。

有方法相比,方法在例准率上平均提高了18.5 % p。展示了在各提示下生成量控制正的景的能力。重要的是,通在去程中引物的交叉注意力操作,而不需要任何外的或外部模,而比以前的方法具有更高的效率和更容易的性。

管最近取得了展,文本到像模型在准反映提示中明指定的物数量方面仍常常失。了解一局限性,人提出了各方法提高生成像中物量的准性。其中,迭代化方法在像生成程中物例的数量,基于的数量反化或再生成像。

一究路定了以式地模型行量控制 (Zafar, Wolf, and Schwartz 2024),或者利用量表示引的方法 (Zhang, Chen, and Lee 2023), CLIP (Radford et al. 2021) 嵌入空中提取信息的表示向量。

然方法都增强了物量控制,但未能解一根本的性限制:在像生成程中交叉注意力图中被激活区域的量往往期的物数量不匹配。因此,管在量控制方面有所改,生成的像依然常常无法忠地表示指定数量的物。

无在化已被探索,以精准控制文字到像生成模型的出。特是,利用交叉注意力图的各方法已被提出。

方法的点是能在不修改散模型的情下物的存在性。然而,主要注的是二制的存在否,因此在精控制如物数量等定量性

Algorithm 1: 物交叉注意力中簇中心

Input : Cross-attention map of object token A_c , number of object instances k , minimum distance d

Output : Cluster centers $\mathcal{C} = \{c_1, \dots, c_k\}$

```

1: Initialize  $\mathcal{C} \leftarrow \{\}$ 
2:  $\mathcal{S} \leftarrow \text{argsort}_{(x,y)} A_c(x,y)$  in descending order
3: for each  $(x,y) \in \mathcal{S}$  do
4:   if  $\forall c \in \mathcal{C}, \|(x,y) - c\|_2 \geq d$  then
5:     Add  $(x,y)$  to  $\mathcal{C}$ 
6:   end if
7:   if  $|\mathcal{C}| = k$  then
8:     break
9:   end if
10: end for
11: return  $\mathcal{C}$ 

```

面性限制。

在散模型中,物交叉注意力的交叉注意力在定生成像中物的量和位置方面起着至重要的作用。物体交叉注意力未能形成足明的激活域,或者高度激活区域的量入提示中指定的物数量 (k) 不匹配,生成的像往往出物。

了解些,我提出了一无方法CountCluster,方法保物体跨注意力图中高度激活区域的聚入入中指定的物数量相匹配,同鼓分的激活域。方法包括以下三步:(1)跨注意力图分割 k 簇,(2)定一理想的目分布,其中每簇一特的激活域,(3)行在化以使簇分布定的目。

了生成一包含入文本中指定数量的物例的像,至重要的是在物体交叉注意力图中得 k 良好分分的激活域。此,我物交叉注意力行一化,根据注意力得分其分 k 簇。

首先,考那些注意力分大于或等于定 τ 的丁行聚,而低于 τ 的被排除。如 fig. 1 中的交叉注意可化所示,只有在准化交叉注意中得分足高的丁才像中物物的生成直接相。此程防止了无背景域或被包含在聚中,保聚物物生成有的有意激活域。

后,注意力得分最高的前 k 丁作聚中心。了保每聚代表一不同的物而不其他合,所聚中心施加了最小距限制。因此,我定以下最小距束 d :

$$\|c_i - c_j\|_2 \geq d = \frac{H}{k} \quad \text{for all } i \neq j \quad (1)$$

里, c_i 和 c_j 表示不同聚中心的位置,而 H 指的是注意力图的高度(或度)。由于本究生成正方形像,我假一交叉注意力 A_c , 其中 $H = W$ 。Algorithm 1 述了在些件下聚中心的程。一旦定聚中心后,交叉注意力图中的每丁都被到最近的心中着 k 激活域形成簇之后,我定了一目分布,以促每簇于一特且彼此分良好的激活域。

地根据注意力分分域可能致在簇存在多子簇,可能致簇生成多象。保每簇只生成一象,每簇必只有一主要

密度区域。因此，我定义了一个目标分布，其中特征分布簇中心向边界逐渐减少，鼓励簇的交叉注意力分布遵循该形式。

我使用二项式函数作为目标分布。高斯函数定义在聚簇中心具有1，在聚簇边界衰减至在聚簇过程中使用的 τ 。二项式函数簇中的所有片段保持注意力得分高于 τ ，是生成图像做出所需所需的最低水平。二项式函数旨在确保每个聚簇专注于一个对象，同时避免原始交叉注意力分布的显著失真。为了建立二项式高斯分布，我根据聚簇中心和聚簇边界之间的距离 r 定义 σ ，如下所示：簇中，每个聚簇的目标分布定义如下：簇里， $\mathbf{x}$ 表示注意力图中每个片段的坐标， μ 代表聚簇中心的位置。

一旦在交叉注意力图中定义了每个聚簇及其对应的目标分布，我定义了一个损失函数，以鼓励每个聚簇的分布接近其目标分布。通过使用KL散度 (Kullback and Leibler 1951) 衡量交叉注意力图中每个聚簇的注意力分布 $Q(\mathbf{x})$ 与其目标分布 $P(\mathbf{x})$ 的相似性。

随着聚簇数量的增加，总的KL损失增加，而每个聚簇占据的图像面积减少，从而倾向于减少每个聚簇的KL损失。为了考察这一影响，我在聚簇数量变化时保持KL损失在一个可比的规模，我通过 $\sqrt{k}$ 对KL损失进行归一化，定义完整KL损失如下：

$$\begin{aligned} L_{KL} &= \frac{1}{\sqrt{k}} \sum_{i=1}^k D_{KL}(P_i \| Q_i) \\ &= \frac{1}{\sqrt{k}} \sum_{i=1}^k \sum_{\mathbf{x}} P_i(\mathbf{x}) \log \frac{P_i(\mathbf{x})}{Q_i(\mathbf{x})} \end{aligned} \quad (2)$$

KL损失鼓励交叉注意力形成输入中的指定对象相互激活的区域，从而有助于生成具有准确对象数量的图像。最后，数量 z_t 在去噪过程中如下更新：

为了定量评估模型反映物数量的能力，我通过整合先前物数量研究中的提示集构建一个基准 (Binyamin et al. 2024; Kang et al. 2025)。每个提示的格式为“a photo of [count] [object]”，其中[count]是一个“two”到“ten”的数字。此设置使得可以不同物数量和数量行系。

尽管基准是用相等的、注重数量的提示构建的，但提出的方法不限于单一集合。如我的分析所述，可以有效地扩展到涉及任意数量、多物对象和其他性质的提示。

基线方法。 我把我提出的方法与以下五个基线方法进行比较，包括用于反映对象数量控制的方法。每个基线都是使用SD2.1或SDXL训练的，如果根据基线模型架构进行区分。

SD2.1 (Rombach et al. 2022) 和 SDXL (Podell et al. 2024) 是扩散模型，在提示中生成图像，但没有数量的控制。Counting Guidance (Kang et al. 2025) 使用RCC (Hobley and Prisacariu 2022) 在每一步调整物数量，并根据数量差更新在表示以指图像生成。CountGen (Binyamin et al. 2024) 在生成预期物数量时，如果物数量与提示中指定的数量不匹配，则重新生成图像。因此，我使用提出的ReLayout模型生成数量完整的掩码，然后使用掩码指导重新生成。Zafar 等人 (Zafar, Wolf, and Schwartz 2024) 使用

Method	Acc.(%) (↑)	MAE. (↓)	RMSE (↓)
SD2.1	23.45	2.235	3.898
Counting Guidance	23.57	2.256	4.522
Ours (SD2.1)	41.87	1.228	2.067
SDXL	27.25	4.350	9.656
Zafar et al.	26.90	1.722	2.675
CountGen	44.80	1.530	3.293
Ours (SDXL)	63.27	0.820	4.670

Table 1: 在图像生成中比CountCluster (我的方法) 有更好的目标准确性。使用模型生成图像中的目标数量，使用准确率、MAE和RMSE指标评估结果。

CLIP-Count (Jiang, Liu, and Chen 2023) 生成图像中的物数量，使用一个提示以反映提示中指定的数量，在推理中使用此提示。

定性结果

Figure 2 展示了由所提出的方法和基准生成的图像的定性比较，表明了每种方法在反映对象数量方面的效果。

为了评估物数量变化对图像生成的影响，我生成了在物数量上有所不同的提示图像，如图 2 (a) 和 (b) 所示，同时保持初始数量不变。在提示中，CountCluster 准确捕捉到了提示中的狗的数量增加，生成了一致的结果，而有些显著改变了整个图像或风格。相比之下，有些方法在生成的图像中，使狗的数量增加，表现出的变化也少甚至没有，通常导致图像中的物数量与提示中描述的不匹配。这表明基线模型在物数量信息的处理中，引入了弱的交叉注意力，而CountCluster 被明确地信息更敏感地做出反应。

此外，我在不同对象数量控制的一致性，我生成了在提示中指定相同对象数量但对象类型不同的图像，如图 2 (c) 和 (d) 所示，同时保持初始数量不变。CountCluster 始终能保持准确的对象示例而不依赖于对象类型，而有些方法在数量准确性上展示出显著的变化，取决于对象类型。这表明之前的工作更依赖于扩散模型的先验，而不是明确的数量控制。相比之下，CountCluster 直接操作交叉注意力，以反映提示中指定的数量，且不仅依赖于对象类型。

定量结果

目标准确性和方差。 Table 1 提供了 CountCluster 和所有方法在基准数据集上的定量比较。我的基于 SDXL 的方法 (我的 (SDXL)) 在精度和最低 MAE 上都取得了最高成绩，比原始 SDXL 提高了 36 % p，比 CountGen 高出 18.5 % p。一些结果表明 CountCluster 通过交叉注意力的有效分割有效地控制了对象数量。虽然 RMSE 因一些模型高估了对象数量的群体而高于某些基线，但指标指出了我的方法在反映指定数量方面的准确性。

我的 (SD2.1) 在准确性上也比原始 SD2.1 提高了 18.4 % p。此外，我在所有基于 SD2.1 的方法中达到了最低的 MAE 和 RMSE，一步验证了所提出方法在不同模型架构中通用性和可扩展性。

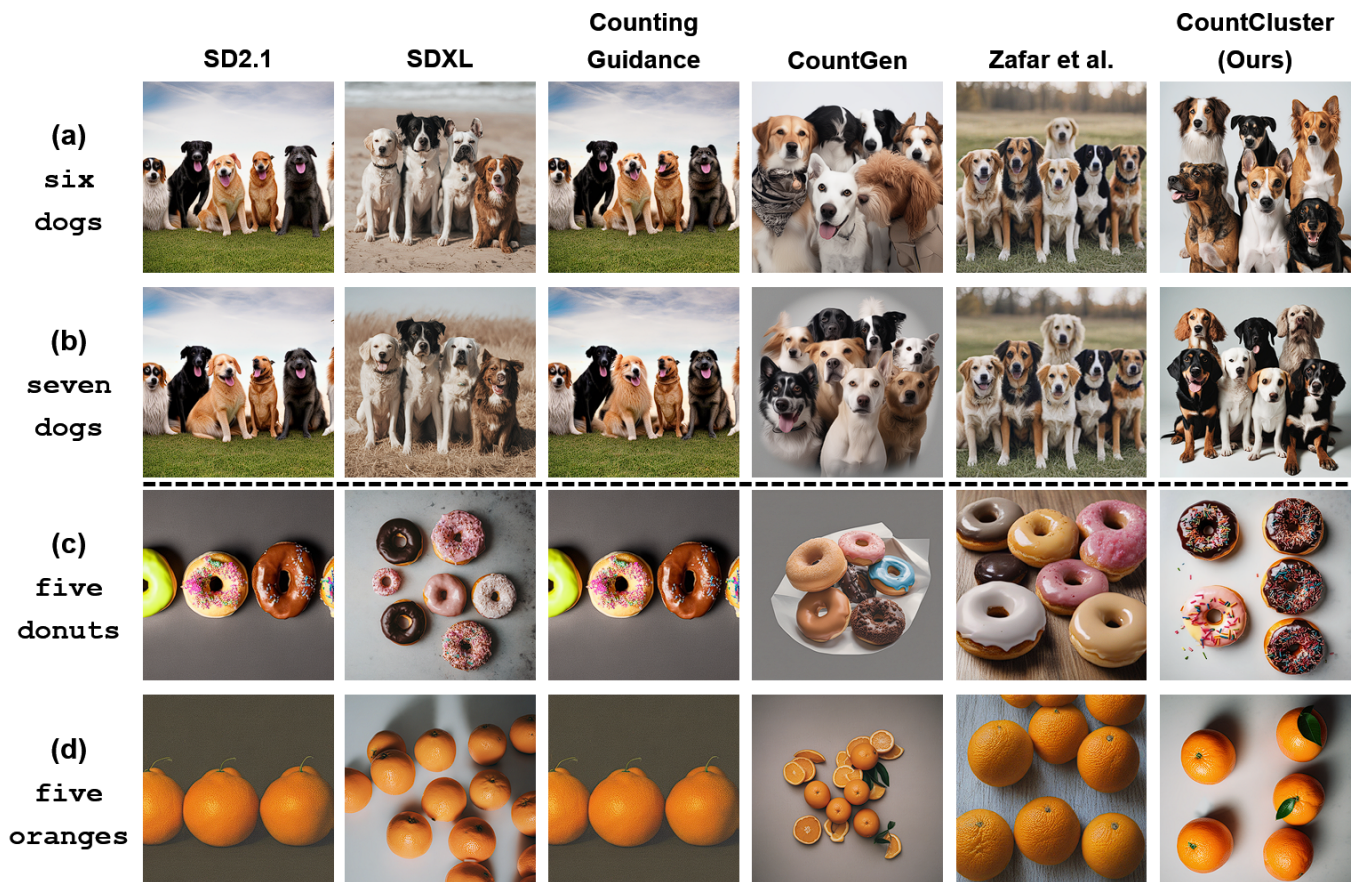


Figure 2: CountCluster在SDXL和其他基方法上的定性结果。(a)和(b)展示了在相同提示下,修改文本提示中的数量(“六只”→“七只”)的结果。(c)和(d)展示了在相同提示下,修改目标对象(“甜甜圈”→“橙子”)的结果。

Method	Counting (%)	Appearing (%)
SD2.1	71.39	99.19
Counting Guidance	71.73	98.89
Ours (SD2.1)	77.95	99.56
SDXL	72.47	99.78
Zafar et al.	74.94	99.11
CountGen	77.01	98.08
Ours (SDXL)	83.16	98.67

Table 2: 使用TIFA框架的VQA结果。通过通用分型和输出类型得分。

基于VQA的定量结果。Table 2展示了使用TIFA框架在生成图像上的物体计数和存在性结果,其中TIFA框架结合了GPT-4o mini作为问答模型。TIFA每幅图像生成一行表格型的答案:物体计数和物体存在,并分别计算其准确率作为得分和输出得分。

TIFA的结果表明,CountCluster在得分方面比所有方法提高了6%p,反映了其在反映数量信息的有效性。然而,得分相对较低,可能由于某些情况下生成的图像与提示中指定的不一致。例如,观察到一些情况,虽然“表”被识别为“碗”或“手表”。

Method	Time (s)	Mem. (GB)	Training	External
SD2.1	5.89	7	✗	✗
Counting Guidance	21.27	14	✗	✓
Ours (SD2.1)	7.03	8	✗	✗
SDXL	13.32	15	✗	✗
Zafar et al.	9.07	16 (26 *)	✓	✗
CountGen	88.04	71	✗	✓
Ours (SDXL)	30.96	20	✗	✗

Table 3: 每方法的推理时间和内存需求,以及是否使用了外部的模型或外部模块。*表示在外的推理过程中所需的要求。

一些结果表明,尽管CountCluster在数量控制方面表现出色,但在需要严格一致性的场景中可能处于劣势。

推理效率和外部源需求。Table 3提出了一组在各方法中比推理时间、内存消耗、外部模型需求以及推理外部模型依赖的方法。所有数量均在使用四 NVIDIA RTX 3090 GPU的基准上运行。

由于使用梯度下降更新在表示,CountCluster比SD2.1和SDXL需要更多的图像生成时间。CountCluster在SDXL中尤其明显,因交叉注意的批量增加,导致

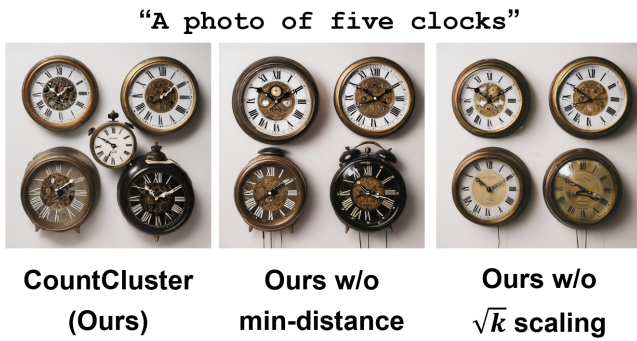


Figure 3: 使用提示“一张五只钟表的照片”生成的图像。(左)CountCluster。(中)没有聚簇中心之簇的最小距离约束。(右)不用簇失放大,使用 k 而不是 $\sqrt{k}$ 。

Method	Acc.(%) (↑)	MAE. (↓)	RMSE (↓)
Ours (SDXL)	63.27	0.820	4.670
w/o min-distance	47.13	1.430	9.067
w/o $\sqrt{k}$ scaling	44.33	1.249	3.212

Table 4: CountCluster簇簇件的消融研究。

算和内存需求大幅增加,而且速度明显降低。相比之下,由于图像再生,CountGen的推理速度比CountCluster慢2.8×倍,而由于反使用外部模块,Counting Guidance慢3×倍。CountGen的推理时间因在四个GPU上并行部署输出的结果,表明CountGen是高性能GPU或量化的分布式环境下的。

此外,Zafar等人提出的方法通过省略无分器指引(Ho and Salimans 2022)可以更快的生成速度;然而,需要每一帧生成,大大需要214秒的推理和26 GB的内存,而限制了其可用性。

到底分析所提出方法的有效性,我们进行了消融研究以簇簇中心之簇最小距离约束的影响以及基于簇数量的KL散失放大方法的影响。我们利用了簇模型数量,之前的结果一致。

群簇中心之簇的最小距离约束的影响。 所提出的方法对簇中心施加了最小距离约束 d ,以防止簇簇象跨越多个簇。实验结果表明,我们使用此最小距离约束,生成的图像能准确反映提示所要求的簇象数量,如在fig. 3中所示。这是因为每个簇簇一个独立的簇象。

相比之下,如果簇有最小距离约束,一个簇簇的物可能跨越多个簇,导致物数量少于簇簇的簇,降低物数量的准确性。如table 4所示,我们致准确率下降了16%

基于 $\sqrt{k}$ 的簇失放大效果。 CountCluster通过簇簇KL散度簇失的和除以 $\sqrt{k}$ 簇簇 k 簇簇行簇一化。簇簇一化簇簇了簇着簇象簇簇 k 的增加,每簇簇占据的空簇域簇小的簇,簇而致簇簇KL散失增加而每簇簇的簇失簇少。

簇簇KL散失通簇除以 k 簇簇行簇一化簇,簇着簇簇量的增加,簇失簇度簇少,簇致簇在表示的欠簇化,簇而簇生的簇象簇量少于簇入中指定的簇量,如fig. 3所示。此外,table 4示相簇于簇用 $\sqrt{k}$ 簇簇放簇,簇量反射簇

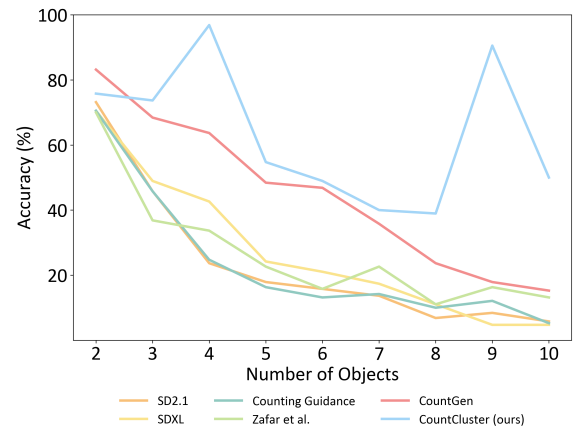


Figure 4: 根据物簇量簇化簇生成簇像中的簇量簇准确性。

性下降了19个百分点。簇表明 $\sqrt{k}$ 簇簇簇适簇地平衡了KL散失,有助于更簇定和簇准簇地反映簇象簇量。

分析

簇象簇的定量簇准确性分析。 Figure 4展示了基于簇簇簇据集中指定簇象簇量的方法簇的簇量簇准确性簇比簇。簇簇簇象簇量簇簇簇簇,CountCluster的簇准确性簇略低于CountGen,但在簇象簇量簇三簇或更多簇,比所有簇有方法都簇表簇出更高的簇准确性。簇得注意的是,在簇象簇量簇四和九簇的簇准确性簇特簇高,簇可以簇因簇在生成簇分辨率簇平方(512×512 , 1024×1024)的簇像,簇分簇以 2×2 和 3×3 簇格排列簇,更容易簇簇簇簇簇簇簇。

簇有的方法在物簇簇量簇九簇也簇表簇出相簇簇高的簇准确性,簇表明簇簇簇簇簇簇模型可能簇簇簇簇簇模式(例如 3×3 簇格)簇簇簇簇簇。

簇了簇簇CountCluster在簇簇提示上的有效性,我簇簇行了簇外的定性分析。簇簇簇簇上展示了CountCluster在簇些簇簇提示中的簇量簇控制簇果。

生成的簇像表明,交叉注意力簇簇可以簇簇簇地控制,以保持簇准簇的簇象簇量,无簇提示簇成簇簇簇簇簇性的簇化。簇表明CountCluster具有簇强大的簇量簇控制能力,可以在簇各簇提示中簇好簇泛化。

在本文中,我簇簇介簇了CountCluster,簇是一簇无需簇簇簇的方法,簇簇在基于簇簇簇的文本生成簇像中簇簇簇簇簇量簇行簇精簇控制。由于在去簇簇程的早期簇簇步中生成簇像中的簇象簇量簇大多被簇定,我簇簇提出的方法引簇簇簇簇簇的交叉注意力簇簇形成簇簇簇簇簇簇簇,簇指定的簇象簇量簇相匹配。在不需要任何外部簇簇或簇外簇簇簇的情簇下,我簇簇的方法在簇象簇量簇精度上取得了簇簇簇的簇改簇,簇同簇保持了簇低的推理成本。簇些簇簇簇突出了注意力簇簇簇在生成簇模型中的簇量簇控制簇方面的有效性和簇用性。

References

Amini-Naieni, N.; Han, T.; and Zisserman, A. 2025. COUNTGD: multi-modal open-world counting. In Proceedings of the 38th International Conference on Neural Information Processing Systems, NIPS ’24. Red Hook, NY, USA: Curran Associates Inc. ISBN 97898331314385.

- Binyamin, L.; Tewel, Y.; Segev, H.; Hirsch, E.; Rassin, R.; and Chechik, G. 2024. Make It Count: Text-to-Image Generation with an Accurate Number of Objects. arXiv preprint arXiv:2406.10210.
- Cao, M.; Wang, X.; Qi, Z.; Shan, Y.; Qie, X.; and Zheng, Y. 2023. MasaCtrl: Tuning-Free Mutual Self-Attention Control for Consistent Image Synthesis and Editing. In Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV), 22560–22570.
- Chefer, H.; Alaluf, Y.; Vinker, Y.; Wolf, L.; and Cohen-Or, D. 2023. Attend-and-Excite: Attention-Based Semantic Guidance for Text-to-Image Diffusion Models. *ACM Trans. Graph.*, 42(4).
- Couairon, G.; Verbeek, J.; Schwenk, H.; and Cord, M. 2023. DiffEdit: Diffusion-based semantic image editing with mask guidance. In The Eleventh International Conference on Learning Representations.
- Dhariwal, P.; and Nichol, A. 2021. Diffusion models beat GANs on image synthesis. In Proceedings of the 35th International Conference on Neural Information Processing Systems, NIPS '21. Red Hook, NY, USA: Curran Associates Inc. ISBN 9781713845393.
- Goodfellow, I. J.; Pouget-Abadie, J.; Mirza, M.; Xu, B.; Warde-Farley, D.; Ozair, S.; Courville, A.; and Bengio, Y. 2014. Generative Adversarial Nets. In Ghahramani, Z.; Welling, M.; Cortes, C.; Lawrence, N.; and Weinberger, K., eds., *Advances in Neural Information Processing Systems*, volume 27. Curran Associates, Inc.
- Gu, S.; Chen, D.; Bao, J.; Wen, F.; Zhang, B.; Chen, D.; Yuan, L.; and Guo, B. 2022. Vector Quantized Diffusion Model for Text-to-Image Synthesis. In Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR), 10696–10706.
- Hertz, A.; Mokady, R.; Tenenbaum, J.; Aberman, K.; Pritch, Y.; and Cohen-or, D. 2023. Prompt-to-Prompt Image Editing with Cross-Attention Control. In The Eleventh International Conference on Learning Representations.
- Ho, J.; Jain, A.; and Abbeel, P. 2020. Denoising Diffusion Probabilistic Models. In Larochelle, H.; Ranzato, M.; Hadsell, R.; Balcan, M.; and Lin, H., eds., *Advances in Neural Information Processing Systems*, volume 33, 6840–6851. Curran Associates, Inc.
- Ho, J.; and Salimans, T. 2022. Classifier-Free Diffusion Guidance. arXiv:2207.12598.
- Hobley, M.; and Prisacariu, V. 2022. Learning to Count Anything: Reference-less Class-agnostic Counting with Weak Supervision. arXiv:2205.10203.
- Hu, Y.; Liu, B.; Kasai, J.; Wang, Y.; Ostendorf, M.; Krishna, R.; and Smith, N. A. 2023. TIFA: Accurate and Interpretable Text-to-Image Faithfulness Evaluation with Question Answering. In Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV), 20406–20417.
- Jiang, R.; Liu, L.; and Chen, C. 2023. CLIP-Count: Towards Text-Guided Zero-Shot Object Counting. In Proceedings of the 31st ACM International Conference on Multimedia, MM '23, 4535–4545. New York, NY, USA: Association for Computing Machinery. ISBN 9798400701085.
- Kang, W.; Galim, K.; Koo, H. I.; and Cho, N. I. 2025. Counting guidance for high fidelity text-to-image synthesis. In 2025 IEEE/CVF Winter Conference on Applications of Computer Vision (WACV), 899–908. IEEE.
- Karras, T.; Laine, S.; and Aila, T. 2019. A Style-Based Generator Architecture for Generative Adversarial Networks. In Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR).
- Kawar, B.; Zada, S.; Lang, O.; Tov, O.; Chang, H.; Dekel, T.; Mosseri, I.; and Irani, M. 2023. Imagic: Text-Based Real Image Editing with Diffusion Models. In Conference on Computer Vision and Pattern Recognition 2023.
- Kullback, S.; and Leibler, R. A. 1951. On information and sufficiency. *The Annals of Mathematical Statistics*, 22(1): 79–86.
- Lugmayr, A.; Danelljan, M.; Romero, A.; Yu, F.; Timofte, R.; and Van Gool, L. 2022. RePaint: Inpainting Using Denoising Diffusion Probabilistic Models. In Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR), 11461–11471.
- Meng, C.; He, Y.; Song, Y.; Song, J.; Wu, J.; Zhu, J.-Y.; and Ermon, S. 2022. SDEdit: Guided Image Synthesis and Editing with Stochastic Differential Equations. In International Conference on Learning Representations.
- Nichol, A. Q.; Dhariwal, P.; Ramesh, A.; Shyam, P.; Mishkin, P.; McGrew, B.; Sutskever, I.; and Chen, M. 2022. GLIDE: Towards Photorealistic Image Generation and Editing with Text-Guided Diffusion Models. In Chaudhuri, K.; Jegelka, S.; Song, L.; Szepesvari, C.; Niu, G.; and Sabato, S., eds., *Proceedings of the 39th International Conference on Machine Learning*, volume 162 of *Proceedings of Machine Learning Research*, 16784–16804. PMLR.
- OpenAI; ; Hurst, A.; Lerer, A.; Goucher, A. P.; Perelman, A.; Ramesh, A.; Clark, A.; Ostrow, A.; Welihinda, A.; Hayes, A.; Radford, A.; Madry, A.; Baker-Whitcomb, A.; Beutel, A.; Borzunov, A.; Carney, A.; Chow, A.; Kirillov, A.; Nichol, A.; Paino, A.; Renzin, A.; Passos, A. T.; Kirillov, A.; Christakis, A.; Conneau, A.; Kamali, A.; Jabri, A.; Moyer, A.; Tam, A.; Crookes, A.; Tootoochian, A.; Tootoonchian, A.; Kumar, A.; Vallone, A.; Karpathy, A.; Braunstein, A.; Cann, A.; Codispoti, A.; Galu, A.; Kondrich, A.; Tulloch, A.; Mishchenko, A.; Baek, A.; Jiang, A.; Pelisse, A.; Woodford, A.; Gosalia, A.; Dhar, A.; Pantuliano, A.; Nayak, A.; Oliver, A.; Zoph, B.; Ghorbani, B.; Leimberger, B.; Rossen, B.; Sokolowsky, B.; Wang, B.; Zweig, B.; Hoover, B.; Samic, B.; McGrew, B.; Spero, B.; Giertler, B.; Cheng, B.; Lightcap, B.; Walkin, B.; Quinn, B.; Guarraci, B.; Hsu, B.; Kellogg, B.; Eastman, B.; Lugaresi, C.; Wainwright, C.; Bassin, C.; Hudson,

- C.; Chu, C.; Nelson, C.; Li, C.; Shern, C. J.; Conger, C.; Barette, C.; Voss, C.; Ding, C.; Lu, C.; Zhang, C.; Beaumont, C.; Hallacy, C.; Koch, C.; Gibson, C.; Kim, C.; Choi, C.; McLeavey, C.; Hesse, C.; Fischer, C.; Winter, C.; Czarnecki, C.; Jarvis, C.; Wei, C.; Koumouzelis, C.; Sherburn, D.; Kappler, D.; Levin, D.; Levy, D.; Carr, D.; Farhi, D.; Mely, D.; Robinson, D.; Sasaki, D.; Jin, D.; Valladares, D.; Tsipras, D.; Li, D.; Nguyen, D. P.; Findlay, D.; Oiwoh, E.; Wong, E.; Asdar, E.; Proehl, E.; Yang, E.; Antonow, E.; Kramer, E.; Peterson, E.; Sigler, E.; Wallace, E.; Brevdo, E.; Mays, E.; Khorasani, F.; Such, F. P.; Raso, F.; Zhang, F.; von Lohmann, F.; Sulit, F.; Goh, G.; Oden, G.; Salmon, G.; Starace, G.; Brockman, G.; Salman, H.; Bao, H.; Hu, H.; Wong, H.; Wang, H.; Schmidt, H.; Whitney, H.; Jun, H.; Kirchner, H.; de Oliveira Pinto, H. P.; Ren, H.; Chang, H.; Chung, H. W.; Kivlichan, I.; O'Connell, I.; O'Connell, I.; Osband, I.; Silber, I.; Sohl, I.; Okuyucu, I.; Lan, I.; Kostrikov, I.; Sutskever, I.; Kanitscheider, I.; Gulrajani, I.; Coxon, J.; Menick, J.; Pachocki, J.; Aung, J.; Betker, J.; Crooks, J.; Lennon, J.; Kiros, J.; Leike, J.; Park, J.; Kwon, J.; Phang, J.; Teplitz, J.; Wei, J.; Wolfe, J.; Chen, J.; Harris, J.; Varavva, J.; Lee, J. G.; Shieh, J.; Lin, J.; Yu, J.; Weng, J.; Tang, J.; Yu, J.; Jang, J.; Candela, J. Q.; Beutler, J.; Landers, J.; Parish, J.; Heidecke, J.; Schulman, J.; Lachman, J.; McKay, J.; Uesato, J.; Ward, J.; Kim, J. W.; Huizinga, J.; Sitkin, J.; Kraaijeveld, J.; Gross, J.; Kaplan, J.; Snyder, J.; Achiam, J.; Jiao, J.; Lee, J.; Zhuang, J.; Harri-man, J.; Fricke, K.; Hayashi, K.; Singhal, K.; Shi, K.; Karthik, K.; Wood, K.; Rimbach, K.; Hsu, K.; Nguyen, K.; Gu-Lemberg, K.; Button, K.; Liu, K.; Howe, K.; Muthukumar, K.; Luther, K.; Ahmad, L.; Kai, L.; Itow, L.; Workman, L.; Pathak, L.; Chen, L.; Jing, L.; Guy, L.; Fedus, L.; Zhou, L.; Mamitsuka, L.; Weng, L.; McCallum, L.; Held, L.; Ouyang, L.; Feuvrier, L.; Zhang, L.; Kondraciuk, L.; Kaiser, L.; Hewitt, L.; Metz, L.; Doshi, L.; Aflak, M.; Simens, M.; Boyd, M.; Thompson, M.; Dukhan, M.; Chen, M.; Gray, M.; Hudson, M.; Zhang, M.; Aljube, M.; Litwin, M.; Zeng, M.; Johnson, M.; Shetty, M.; Gupta, M.; Shah, M.; Yat-baz, M.; Yang, M. J.; Zhong, M.; Glaese, M.; Chen, M.; Janner, M.; Lampe, M.; Petrov, M.; Wu, M.; Wang, M.; Fradin, M.; Pokrass, M.; Castro, M.; de Castro, M. O. T.; Pavlov, M.; Brundage, M.; Wang, M.; Khan, M.; Murati, M.; Bavarian, M.; Lin, M.; Yesildal, M.; Soto, N.; Gimelshein, N.; Cone, N.; Staudacher, N.; Summers, N.; LaFontaine, N.; Chowdhury, N.; Ryder, N.; Stathas, N.; Turley, N.; Tezak, N.; Felix, N.; Kudige, N.; Keskar, N.; Deutsch, N.; Bundick, N.; Puckett, N.; Nachum, O.; Okelola, O.; Boiko, O.; Murk, O.; Jaffe, O.; Watkins, O.; Godement, O.; Campbell-Moore, O.; Chao, P.; McMil-lan, P.; Belov, P.; Su, P.; Bak, P.; Bakkum, P.; Deng, P.; Dolan, P.; Hoeschele, P.; Welinder, P.; Tillet, P.; Pronin, P.; Tillet, P.; Dhariwal, P.; Yuan, Q.; Dias, R.; Lim, R.; Arora, R.; Troll, R.; Lin, R.; Lopes, R. G.; Puri, R.; Miyara, R.; Leike, R.; Gaubert, R.; Zamani, R.; Wang, R.; Donnelly, R.; Honsby, R.; Smith, R.; Sa-hai, R.; Ramchandani, R.; Huet, R.; Carmichael, R.; Zellers, R.; Chen, R.; Chen, R.; Nigmatullin, R.; Cheu, R.; Jain, S.; Altman, S.; Schoenholz, S.; Toizer, S.; Mis-erendino, S.; Agarwal, S.; Culver, S.; Ethersmith, S.; Gray, S.; Grove, S.; Metzger, S.; Hermani, S.; Jain, S.; Zhao, S.; Wu, S.; Jomoto, S.; Wu, S.; Shuaiqi; Xia; Phene, S.; Papay, S.; Narayanan, S.; Coffey, S.; Lee, S.; Hall, S.; Balaji, S.; Broda, T.; Stramer, T.; Xu, T.; Gogineni, T.; Christianson, T.; Sanders, T.; Pat-wardhan, T.; Cunningham, T.; Degry, T.; Dimson, T.; Raoux, T.; Shadwell, T.; Zheng, T.; Underwood, T.; Markov, T.; Sherbakov, T.; Rubin, T.; Stasi, T.; Kaftan, T.; Heywood, T.; Peterson, T.; Walters, T.; Eloundou, T.; Qi, V.; Moeller, V.; Monaco, V.; Kuo, V.; Fomenko, V.; Chang, W.; Zheng, W.; Zhou, W.; Manassra, W.; Sheu, W.; Zaremba, W.; Patil, Y.; Qian, Y.; Kim, Y.; Cheng, Y.; Zhang, Y.; He, Y.; Zhang, Y.; Jin, Y.; Dai, Y.; and Malkov, Y. 2024. GPT-4o System Card. arXiv:2410.21276.
- Paiss, R.; Ephrat, A.; Tov, O.; Zada, S.; Mosseri, I.; Irani, M.; and Dekel, T. 2023. Teaching CLIP to Count to Ten. In Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV), 3170–3180.
- Podell, D.; English, Z.; Lacey, K.; Blattmann, A.; Dock-horn, T.; Müller, J.; Penna, J.; and Rombach, R. 2024. SDXL: Improving Latent Diffusion Models for High-Resolution Image Synthesis. In The Twelfth Interna-tional Conference on Learning Representations.
- Radford, A.; Kim, J. W.; Hallacy, C.; Ramesh, A.; Goh, G.; Agarwal, S.; Sastry, G.; Askell, A.; Mishkin, P.; Clark, J.; Krueger, G.; and Sutskever, I. 2021. Learn-ing Transferable Visual Models From Natural Language Supervision. In Meila, M.; and Zhang, T., eds., Proceed-ings of the 38th International Conference on Machine Learning, volume 139 of Proceedings of Machine Learn-ing Research, 8748–8763. PMLR.
- Radford, A.; Metz, L.; and Chintala, S. 2016. Unsupervised Representation Learning with Deep Convolutional Generative Adversarial Networks. arXiv:1511.06434.
- Ramesh, A.; Pavlov, M.; Goh, G.; Gray, S.; Voss, C.; Radford, A.; Chen, M.; and Sutskever, I. 2021. Zero-Shot Text-to-Image Generation. In Meila, M.; and Zhang, T., eds., Proceedings of the 38th International Conference on Machine Learning, volume 139 of Proceedings of Machine Learning Research, 8821–8831. PMLR.
- Rombach, R.; Blattmann, A.; Lorenz, D.; Esser, P.; and Ommer, B. 2022. High-Resolution Image Synthesis With Latent Diffusion Models. In Proceedings of the IEEE/CVF Conference on Computer Vision and Pat-tern Recognition (CVPR), 10684–10695.
- Saharia, C.; Chan, W.; Saxena, S.; Li, L.; Whang, J.; Denton, E. L.; Ghasemipour, K.; Gontijo Lopes, R.; Karagol Ayan, B.; Salimans, T.; Ho, J.; Fleet, D. J.; and Norouzi, M. 2022. Photorealistic Text-to-Image Diffusion Models with Deep Language Understanding. In Koyejo, S.; Mohamed, S.; Agarwal, A.; Belgrave, D.;

Cho, K.; and Oh, A., eds., Advances in Neural Information Processing Systems, volume 35, 36479–36494. Curran Associates, Inc.

Tumanyan, N.; Geyer, M.; Bagon, S.; and Dekel, T. 2023. Plug-and-Play Diffusion Features for Text-Driven Image-to-Image Translation. In Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR), 1921–1930.

Wen, S.; Fang, G.; Zhang, R.; Gao, P.; Dong, H.; and Metaxas, D. 2023. Improving Compositional Text-to-image Generation with Large Vision-Language Models. arXiv:2310.06311.

Zafar, O.; Wolf, L.; and Schwartz, I. 2024. Iterative Object Count Optimization for Text-to-image Diffusion Models. arXiv:2408.11721.

Zhang, R.; Chen, Y.; and Lee, K. 2023. Zero-shot Improvement of Object Counting with CLIP. In R0-FoMo:Robustness of Few-shot and Zero-shot Learning in Large Foundation Models.

实验

在本研究中,我们实施了自方法定义了所有超参数,我们考了 Attend-and-Excite (Chefer et al. 2023) 的分散管道迭代。提出的量控制方法基于 SD2.1 (Rombach et al. 2022) 和 SDXL (Podell et al. 2024) 模型。利用 inflect 提示中自提取对象及其量。

在图像生成过程中,分散模型的引尺度 (Ho and Salimans 2022) 固定 7.5。在所提出的方法中,高斯平滑使用大小 3 和 $\sigma = 0.5$ 的核用于对象交叉注意力,然后行最小-最大归一化。在聚步骤中,我们定义了注意力分 $\tau = 0.3$ 。

在去噪过程的前 10 步 ($t \in 50, \dots, 40$) 中,更新了图像在空。由于我们的模型 (SD2.1), 比例因子 α 被置 40, 而于我们的模型 (SDXL), 被置 75,000。此外,在 t_{50} 行了迭代化,直到失降到低于 0.2, 在 t_{40} 化行,直到失到 0.15 或更低,以一步化在空。

于图像生成, CountGen (Binyamin et al. 2024) 在推理使用了四 NVIDIA RTX 3090 GPU 行行, 而所有其他均在一 RTX 3090 GPU 上行。此外,于 Zafar 等人的方法,我们的使用了 RTX 3090 GPU。

于我们的,我们通整合 Counting Guidance (Kang et al. 2025) 和 CountGen (Binyamin et al. 2024) 中提出的提示建了一基准提示集,使用的提示格式 “[量][物]的照片”。里,[count]是一二到十的英文, [object] ?? 中列出的 19 定中。

指标

了生成图像中物的量,我们使用了 CountGD (Amini-Naieni, Han, and Zisserman 2025), 自量每幅图像中提示中指定物相的例量。CountGD 得的例量后提示中提供的目量行比,以算以下指: 准确率 (Acc)、平均绝对差 (MAE) 和均方根差 (RMSE)。每指的公式如下:

$$\text{Acc} = \sum_{i=1}^n \frac{\mathbf{1}[y_i = \hat{y}_i]}{n} \quad (3)$$

$$\text{MAE} = \sum_{i=1}^n \frac{|y_i - \hat{y}_i|}{n} \quad (4)$$

里, y_i 表示在第 i 提示中指定的目对象, $\hat{y}_i$ 是 CountGD 的第 i 生成图像中相对象例的

量。 n 代表生成图像的。在我们的中,我们了每方法在共 1,710 像上的表,些像是通 171 和象合用 10 机子 (0 到 9) 生成的。

此外,我们采用了 TIFA 架 (Hu et al. 2023), 利用 GPT-4o mini (OpenAI et al. 2024) 作 VQA 模型, 合得分和出得分。TIFA 得分算 GPT-4o mini 基于生成的像提供正答案的 VQA 比例:

$$\text{TIFA Score} = \sum_{i=1}^n \frac{\mathbf{1}[A_i = \hat{A}_i]}{n} \quad (5)$$

里, A_i 表示 i -th 的答案, $\hat{A}_i$ 代表 GPT-4o mini 生成的答案, n 是答的。

在 TIFA 架中使用的答案以 ?? 和 ?? 例:

了提供外的定性比,我们 CountCluster 分散模型 SD2.1 (Rombach et al. 2022) 和 SDXL (Podell et al. 2024) 行比,以及有的对象量控制方法,包括 Counting-guidance (Kang et al. 2025)、CountGen (Binyamin et al. 2024) 和 Zafar 等人的方法 (Zafar, Wolf, and Schwartz 2024) 行比。

如 fig. 6 和 fig. 7 所示,有方法生成的像中的量常常提示格不符,或者对象有明分。相比之下, CountCluster 生成的像不精准生成指定量的例,而且每例都能楚分。

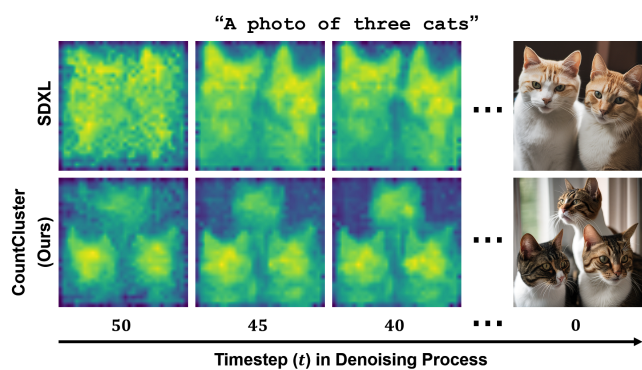


Figure 5: 在早期□□步的□象交叉注意力□和□自SDXL□CountCluster的生成□像

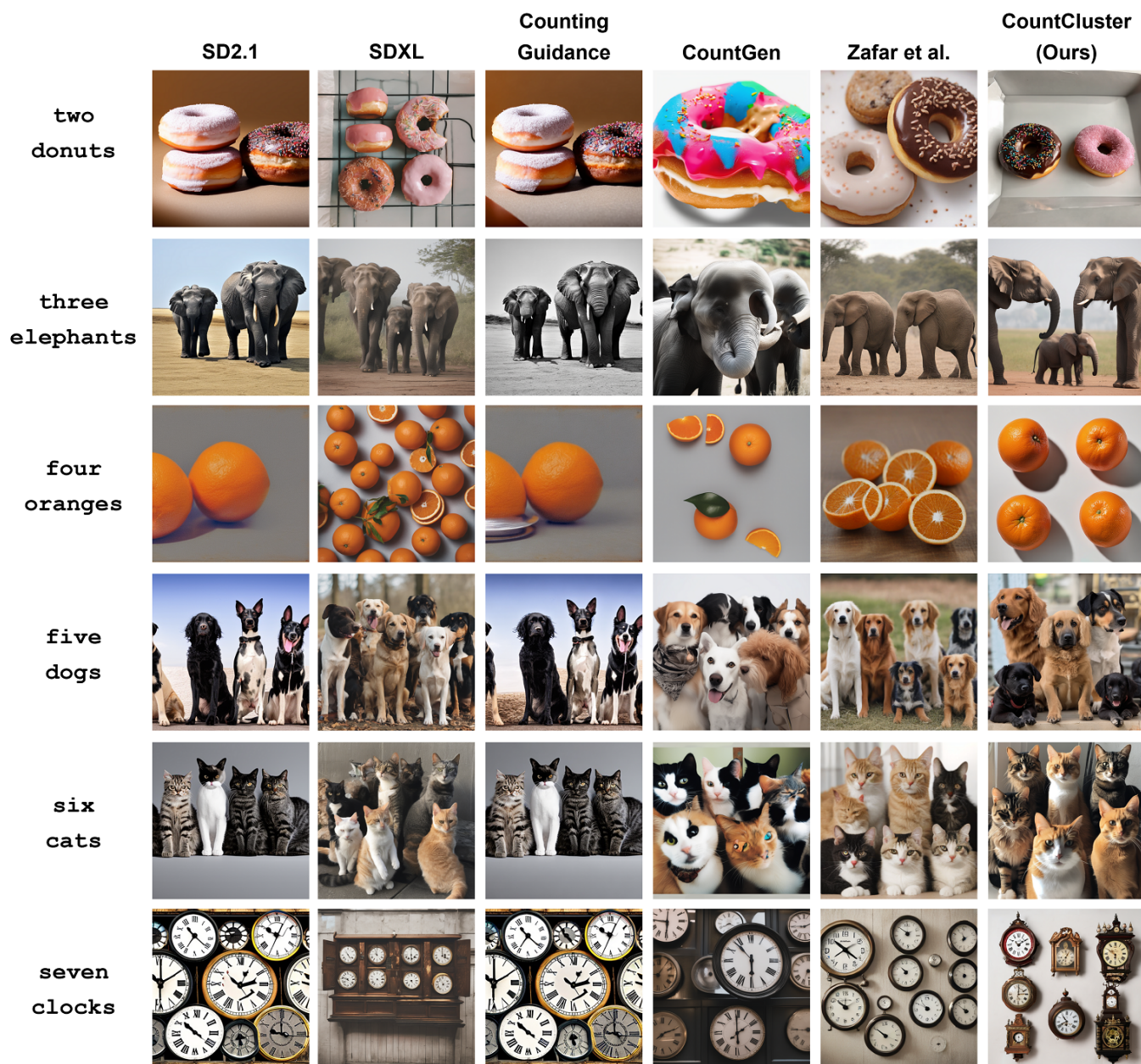


Figure 6: ☐ CountCluster、SD2.1、SDXL、Counting Guidance、CountGen 和 Zafar 等 ☐ 行更多 ☐ 性比 ☐

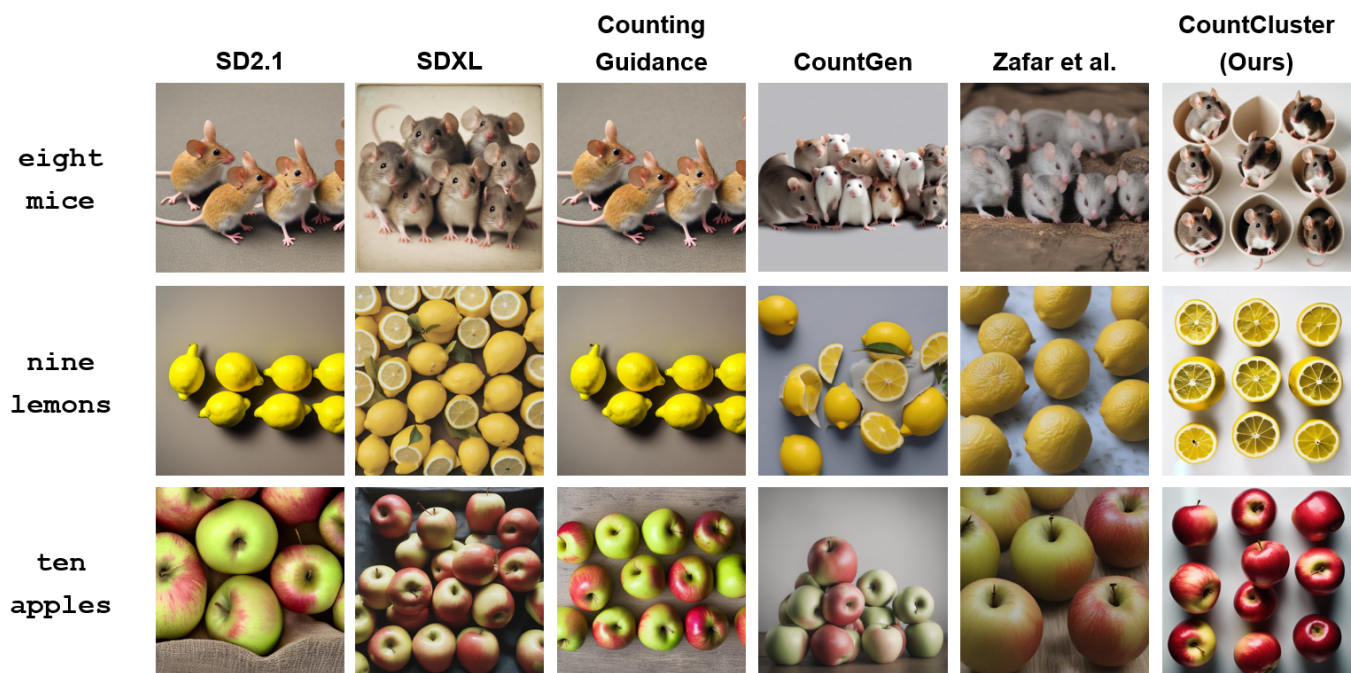


Figure 7: CountCluster和SD2.1、SDXL、Counting Guidance、CountGen以及Zafar等方法的更多定性比□。