

Transferability Between Understanding and Generation in Unified Multimodal Models

Jiwon Kang^{1*}, Heeji Yoon^{1*}, Jaewoo Jung¹, Jaewon Min¹, Minkyong Jeon¹,
Biyeon Hwang², Sangwon Jung^{3†}, and Seungryong Kim^{1†}

¹ KAIST AI, South Korea ² Blynx, South Korea ³ Trillionlabs, South Korea

Abstract. Unified Multimodal Models (UMMs) integrate image understanding and generation within a single architecture, yet how the two tasks interact remains understudied. We investigate **transferability** in UMMs: whether training a capability on one task improves the same capability on the other without explicit supervision. Through controlled experiments, we empirically find that transferability depends on architecture—models with fully shared transformer backbone and a unified visual encoder exhibit consistent cross-task transfer, while loosely coupled designs show little or none. Leveraging this transferability, we propose a practical training strategy. The most straightforward way to improve a target generative capability (*e.g.*, counting) is to fine-tune generation directly, but this can degrade visual quality due to distribution shift. Instead, we train the corresponding understanding task and let it *transfer* into generation, which improves capability-specific generative performance while minimizing distribution shift. We validate this across three capabilities—counting, spatial relation, and text recognition/generation—showing that cross-task transferability can be systematically exploited in UMMs.

Keywords: Unified Multimodal Models · Cross-Task Transferability · Image Understanding · Image Generation

1 Introduction

Propelled by the success of Large Language Models (LLMs) [3, 34, 64, 81], Vision-Language Models (VLMs) [5, 17, 21, 38, 56, 103] have achieved significant breakthroughs in image understanding. Concurrently, diffusion-based generative models [11, 67, 69] have driven tremendous advancements in image generation. Despite their profound progress, these two domains have mostly evolved in isolation. This historical separation has recently motivated the development of Unified Multimodal Models (UMMs) [14, 16, 19, 22, 95, 96], which integrate both tasks—image understanding and generation—within a single architecture. While recent UMMs demonstrate that such unification can achieve performance comparable to specialized experts for each task, it remains fundamentally unclear how the two tasks interact once they are trained within a single model.

* Equal contribution.

† Co-corresponding author.

Several recent studies [25, 53, 92, 97] have investigated interactions between image understanding and generation, arguing that understanding can improve generation and vice versa. However, existing analyses typically examine these interactions through aggregate performance changes on general benchmarks [31, 35, 52, 104] (*e.g.*, adding generation training improves understanding benchmark scores). While such results provide useful high-level observations, relying on broad benchmark improvements makes it difficult to determine whether the observed gains stem from genuine cross-task knowledge transfer or merely from indirect factors such as additional data or regularization effects.

We address this gap by studying a more direct and actionable signal of cross-task interaction: *transferability*. Instead of inferring interaction from aggregate benchmark gains, we examine whether training a given capability (*e.g.*, counting) in one task (understanding or generation) reliably improves the same capability in the other task, under controlled interventions. Through a series of experiments, we find that such cross-task transferability emerges in current UMMs, and this study leads to two key findings: (i) the existence and strength of transferability vary across models, empirically revealing a dependency on architectural design; (ii) when transfer is present, it can be exploited as a practical mechanism to improve targeted generative capabilities by training on the corresponding understanding task.

Specifically, we begin with a controlled analysis on architecture using *counting* as a probe task. Across several representative open-source UMMs [14, 16, 22, 95], we empirically observe that transferability depends on architectural design. In particular, models that employ a fully shared transformer backbone along with a unified visual encoder [95] exhibit consistently the strongest transfer between tasks, whereas architectures that leverage a separate backbone or an image encoder [14, 16] show little or no transfer. These results suggest that the degree of representation sharing may play an important role in enabling cross-task capability transfer.

Using the architecture with the strongest transfer, our next step is to investigate whether transferability can be utilized as a practical mechanism for enhancing generative capabilities. While directly fine-tuning with a generation objective may be the most straightforward approach to acquire a specific capability, it carries the risk of inducing a distribution shift [13, 48, 61] that may degrade overall visual quality. Instead, we find that training on the corresponding understanding task can effectively introduce the capability into generation through transfer, minimizing the distribution shift. Extensive experiments across multiple visual capabilities, including counting, spatial relation, and text recognition/generation, confirm that cross-task transfer can consistently improve generation performance on specific capabilities while preserving general multimodal understanding ability and generative quality.

Our contributions are summarized as follows:

- We introduce **transferability** as a tool for analyzing cross-task interactions in Unified Multimodal Models, and show that capabilities learned in understanding or generation can be transferred to the other task.

- Through controlled experiments, we empirically demonstrate that transferability depends on architectural design, emerging most clearly in models with fully shared transformer backbones with a unified visual encoder.
- We show that transferability can be leveraged as a practical strategy for improving targeted generative capabilities while avoiding degradation of original multimodal understanding ability or generative quality, and validate this approach across counting, spatial relation, and text recognition/generation tasks.

2 Related Work

2.1 Image Understanding and Generation in UMMs

The success of large language models [3, 34, 64, 81] has catalyzed the extension of language modeling into the multimodal domain. Vision-Language Models (VLMs) [5, 17, 21, 38, 56] enable strong image understanding by pairing a visual encoder with an LLM. More recently, Unified Multimodal Models (UMMs) have pushed this paradigm further by incorporating image generation alongside understanding within a single model [14, 22, 78, 90, 95].

Visual representations. A key design axis in UMMs is the choice of visual representation for encoding input images. One line of work adopts language-supervised encoders such as CLIP [68] and SigLIP [82, 105], whose representations are aligned with text and thus preserve high-level semantics amenable to reasoning. A complementary line leverages autoencoders [7] which are typically employed for image generation. These autoencoders, which prioritize pixel-level reconstruction fidelity, are usually trained with variational objectives [44] for continuous representation [11, 26, 46, 67, 69] or a vector quantization objective [83] for discrete representation [27, 83]. Analyses in Janus [16, 90] and related works [79] show that language-supervised representations substantially outperform compression-oriented tokenizers for understanding tasks, motivating several UMMs to employ separate encoders for each task.

Generation objectives. Image generation in UMMs follows two main paradigms. *Autoregressive* models [16, 19, 45, 78, 90, 92] employ discrete image tokenizers and generate image tokens sequentially, directly extending the next-token prediction framework of LLMs. *Diffusion*-based models instead learn a denoising objective [2, 40, 55, 58, 75, 76] and can be further divided into discrete-token diffusion models [72, 95, 102] which apply discrete diffusion [4, 60, 70] over quantized latent representations, and continuous diffusion models [14, 25, 33, 77, 94, 112] which operate on continuous latent representations.

Architectural taxonomy of unified multimodal models. Although finer-grained categorizations are possible [111], we group existing unified models into three dominant architecture families for our analysis. **(1) Single transformer** designs route all modalities through a shared transformer backbone, maximizing parameter sharing [16, 19, 72, 78, 87, 90, 95, 102]. These models either use

unified image encoders [72, 78, 87, 95, 102] or separate ones [16, 90] for understanding/generation. **(2) Mixture-of-Transformers (MoT)** designs maintain modality-specific transformer branches but couple them through joint attention layers, allowing selective parameter sharing while retaining specialized capacity [22, 50, 86, 89]. **(3) Decoupled transformer** designs employ separate transformers for language and vision, where the LLM provides conditioning signals (*e.g.*, hidden states) that are subsequently fed into a dedicated diffusion module for generation [14, 33, 65, 77, 80].

2.2 Understanding Interactions Between Understanding and Generation in UMMs

A central question in designing Unified Multimodal Models is whether multimodal understanding and generation genuinely benefit each other within a shared architecture. Several works report potential synergy, arguing that high-level semantic representations learned for understanding can improve the semantic fidelity of generation, while fine-grained visual representations acquired through generation may enhance visual grounding and compositional reasoning [45, 53, 91, 92, 95, 99, 108]. On the other hand, research such as UniFork [50] empirically demonstrates that ideal representations for these two tasks throughout the network layers are distinct, suggesting that simple parameter sharing can lead to a representational compromise. This perspective aligns with a broader trend of decoupled architectures that utilize task-specific branches in higher layers or separated visual paths to mitigate potential conflicts [16, 22, 90].

While these studies provide valuable insights into the potential benefits and trade-offs of joint training, most analyses examine interaction through aggregate performance changes on general benchmarks. For example, prior work often evaluates whether increasing generation-oriented training improves overall understanding benchmarks (*e.g.*, MMMU [104] or POPE [52]), or conversely whether additional understanding supervision improves generation benchmarks such as GenEval. Although such evaluations provide useful high-level signals, they make it difficult to isolate how specific capabilities learned in one modality influence the other. In contrast, our work investigates cross-task interaction at the capability level through the lens of transferability. Rather than inferring interaction from aggregate benchmark improvements, we directly examine whether learning a capability in one modality improves the same capability in the other modality.

3 Does knowledge transfer emerge in UMMs?

In this section, we investigate cross-task transfer in current open-source UMMs: (i) whether learning a capability in one task (understanding or generation) improves the same capability in the other, and (ii) whether this transfer depends on architectural design. To this end, we conduct an analysis using counting as a probe capability, which provides a clean and quantifiable concept that both understanding and generation should encode.

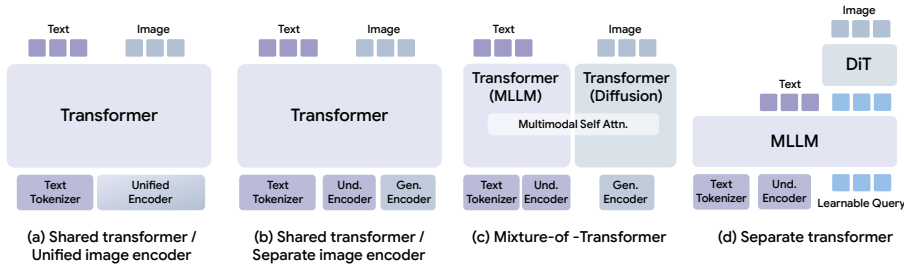


Fig. 1: Architecture of chosen unified multimodal models for cross-task transferability analysis on counting. (a) Shared Transformer with a unified image encoder. (b) Shared Transformer with separate image encoders for each task (understanding/generation). (c) Mixture-of-Transformer which uses separate transformers for understanding and generation but connects each other via joint(multimodal) self-attention. (d) Separate Transformer where the MLLM transformer provides condition signals that are fed into a diffusion transformer (DiT).

Train → Test	(a) Lumina-DiMOO		(b) Janus-Pro		(c) BAGEL		(d) BLIP3-o	
	Acc. (%) ↑	MAD ↓	Acc. (%) ↑	MAD ↓	Acc. (%) ↑	MAD ↓	Acc. (%) ↑	MAD ↓
<i>Generation Evaluation</i>								
Baseline	48.0	1.11	38.0	1.46	42.0	1.16	31.0	2.08
<i>und</i> → <i>gen</i>	57.0 (+9.0)	0.83 (-0.28)	32.0 (-6.0)	1.55 (+0.09)	47.0 (+5.0)	1.11 (-0.05)	37.0 (+6.0)	1.80 (-0.28)
<i>Understanding Evaluation</i>								
Baseline	22.0	2.48	37.0	2.70	63.0	0.65	63.0	0.54
<i>gen</i> → <i>und</i>	30.0 (+8.0)	1.88 (-0.60)	38.0 (+1.0)	3.34 (+0.64)	64.0 (+1.0)	0.62 (-0.03)	62.0 (-1.0)	0.56 (+0.02)

Table 1: Probing cross-task transferability on counting across different UMM architectures. We conduct two separate fine-tuning experiments and compare them against the baseline. (i) *und* → *gen*: training only on understanding and evaluating on generation. (ii) *gen* → *und*: training only on generation and evaluating on understanding. Acc. denotes exact-match accuracy (%) and MAD denotes mean absolute deviation from the ground-truth count. Shared transformer models with a unified visual tokenizer (Lumina-DiMOO) show the strongest transferability in both directions.

3.1 Experimental setup

Considering the diverse architectural choices introduced in Section 2.1, we select four representative open-source models: Lumina-DiMOO [95] for **(a) Shared Transformer with unified image encoder**, Janus-Pro [16] for **(b) Shared Transformer with separate image encoder**, BAGEL [22] for **(c) Mixture of Transformers**, and BLIP3-o [14] for **(d) Separate Transformer**. The architectural overview of each selected model is illustrated in Figure 1. We select 7B–8B parameter versions of each model.

For each model, we conduct two separate training runs (understanding and generation) using LoRA-based supervised fine-tuning. We construct training datasets by filtering PixMo-Points and PixMo-Count [21] to a counting range of 0–20. For understanding, the model learns counting through visual question an-

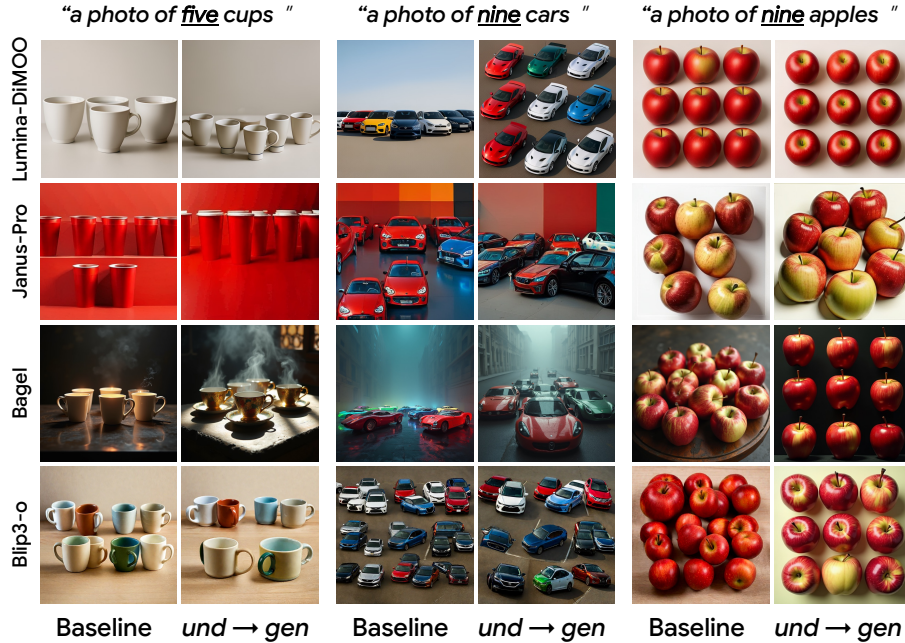


Fig. 2: Qualitative results of $und \rightarrow gen$ transferability across models. For each prompt, we compare the baseline model with the model trained only on counting understanding data ($und \rightarrow gen$) and evaluated on generation. All images are generated using the same random seed for fair comparison.

swering, where it answers how many objects appear in an image. For generation, the model learns counting through text-to-image generation, where it is supervised to produce images containing objects whose count matches the prompt. To get the generation prompts, we use Qwen3-VL [5] to caption the same images and append object count information to form the generation prompts.

For evaluation, we use two separate protocols. Understanding is evaluated on the test split of PixMo-Count [21] in a VQA format, where the model is asked to count target objects in an image and responds with a number. For generation, we follow the counting protocol from GenEval [35] and construct prompts in the form "A photo of [NUMBER] [OBJECT]s" with counts ranging from 2 to 10. Generated images are evaluated using an object detector [18] to verify whether the number of detected objects matches the specified count. We report both accuracy and mean absolute deviation (MAD) to capture exact correctness and proximity to the ground truth.

3.2 Transferability across architectures

Transferability varies across architectures. Table 1 summarizes the results. The most consistent finding is that knowledge transfer from understanding to generation, which we term *understanding-to-generation transfer* ($und \rightarrow gen$), is observed in all models except Janus-Pro [16]. BAGEL [22], Lumina-DiMOO [95],

and BLIP3-o [14] all show improved generation accuracy and lower MAD when trained with the understanding objective, indicating that counting capability acquired through understanding can be transferred to generation. In contrast, the inverse direction ($gen \rightarrow und$) is more selective: only BAGEL and Lumina-DiMOO show gains in counting understanding from generation training, making transfer bi-directional in these two models. Among all models, Lumina-DiMOO shows the strongest effect in both directions, underscoring that transferability is not uniform across UMM architectures. Qualitative examples illustrating this behavior are shown in Figure 2.

Which architecture results in the strongest transferability? From an architectural perspective, we empirically observe that models which employ a fully shared transformer backbone for both understanding and generation, along with a unified vision encoder (Figure 1-(a)), such as Lumina-DiMOO [95], exhibit the strongest transferability. This result is consistent with recent empirical findings that adopting a unified encoder for image understanding and generation produces synergistic effects in UMMs [91, 92]: because a shared transformer backbone with a unified encoder lets both tasks operate over a shared representation, knowledge acquired through one objective naturally propagates to the other, facilitating cross-task transfer. By contrast, architectures that maintain separate visual pathways for each task offer fewer opportunities for such propagation, which explains the weaker or absent transfer observed in models like Janus-Pro.

4 Using Transferability to Improve Generative Capabilities

In this section, we investigate how the transferability that emerged in UMMs can be practically exploited. In practice, generative models often exhibit weaknesses in specific capabilities such as spatial relation [15, 30, 37, 57, 66, 88, 93] or object counting [9, 12, 32, 51]. Generation objectives [2, 40, 55, 58, 75, 76] fundamentally learn to approximate their training data distribution, so a straightforward way to inject a specific skill (*e.g.*, counting) is fine-tuning directly on a skill-focused dataset with generation objectives. However, fine-tuning pulls the model away from its original image generation distribution, which can cause a distribution shift that may degrade generative quality [13, 48, 61].

This issue is even more pronounced for recent models whose generation distribution is carefully aligned with human preferences in a final post-training stage, using SFT over curated high-quality datasets [20, 67] or RL-based methods [10, 28, 47, 49, 85, 100]. However, direct fine-tuning readily destroys this alignment. Workarounds would be either curating a dataset that is simultaneously high-quality and skill-rich, or repeating the costly post-training alignment after fine-tuning, neither of which is practical.

We propose a practical alternative that leverages transferability. Rather than optimizing the generation objective directly, we train the target capability through



Fig. 3: Qualitative comparison of understanding-to-generation transfer ($und \rightarrow gen$) versus direct generation training ($gen \rightarrow gen$) on counting. While direct generation training ($gen \rightarrow gen$) significantly degrades image quality due to distribution shift, $und \rightarrow gen$ preserves generation quality while achieving improved counting accuracy.

Train $\rightarrow$ Test	Acc. (%) $\uparrow$	MAD $\downarrow$	IS [71] $\uparrow$	FID [39] $\downarrow$
Baseline	48.0	1.11	16.86	-
$und \rightarrow gen$	57.0 (+9.0)	0.83 (-0.28)	17.55	31.47
$gen \rightarrow gen$	57.0 (+9.0)	0.91 (-0.20)	15.29	52.51

Table 2: Quantitative comparison of understanding-to-generation transfer ($und \rightarrow gen$) versus direct generation training ($gen \rightarrow gen$) on counting. Using transferability ($und \rightarrow gen$) attains task accuracy on par with direct training ($gen \rightarrow gen$) while retaining image quality (IS [71], FID [39]). In contrast, $gen \rightarrow gen$ degrades both metrics.

the understanding and let it transfer to generation. This approach enables targeted improvements in generative capabilities while maintaining the overall generation quality of the model. To study this mechanism, we analyze understanding-to-generation transfer across three representative tasks, *i.e.*, counting, spatial relation, and text recognition/generation. We conduct the following experiments on Lumina-DiMOO [95], as its single-transformer design with a unified image encoder empirically exhibited the strongest cross-task transfer among the UMMs studied in Section 3. Result on another model [102] sharing the same architecture is reported in the Appendix.

4.1 Counting

We first revisit the counting generation from this perspective. Table 1 previously showed that counting exhibits strong transferability from understanding to generation. To examine whether this transfer can be practically beneficial, we compare the model trained with the understanding objective ($und \rightarrow gen$) with a model directly trained with the generation objective ($gen \rightarrow gen$) in Table 2.

Surprisingly, the model trained through understanding achieves a lower mean absolute deviation (MAD) in counting generation. In other words, introducing counting capability through the understanding leads to more accurate counting behavior in generation than directly optimizing the generation objective.

The qualitative results in Figure 3 further highlight this difference. When the model is trained with the generation objective ($gen \rightarrow gen$), the narrow distribution of the dataset causes the model to drift from its original training distribution, resulting in noticeable degradation of overall image quality. In contrast, the model trained via understanding ($und \rightarrow gen$) preserves the overall image quality. To quantify this, we additionally measure distribution shift using Fréchet Inception Distance (FID) [39] and image quality using Inception Score (IS) [71] in Table 2. Specifically, FID is computed between the baseline generated images and those from the trained models. Training directly with the generation objective ($gen \rightarrow gen$) yields a higher FID than $und \rightarrow gen$ and degrades IS relative to the baseline, confirming both a larger distribution shift and loss of generative quality. In contrast, the understanding-based approach ($und \rightarrow gen$) maintains generation quality with lower distribution shift, while improving counting accuracy, demonstrating that transfer through understanding provides a more stable way to introduce new capabilities into generation.

4.2 Spatial relation

We next consider a structurally different capability, spatial relation, where the model must encode the relative layout of objects. Unlike counting, this task requires capturing compositional relationships between multiple entities. This allows us to examine whether the transferability observed in counting extends to more complex relational concepts.

Experimental setting. We observe that pretrained models already handle simple spatial relations such as left, right, top, and bottom reasonably well. We therefore focus on a more challenging setting: diagonal relations between two objects (*e.g.*, top-left, top-right, bottom-right, bottom-left). Since existing spatial reasoning datasets [24, 29, 41, 74, 98, 101, 109] are difficult to use directly for our purpose—owing to their limited scale, mismatched spatial complexity, or inconsistent annotation formats—we construct a synthetic dataset of 200K samples from ImageNet [23]. For each sample, as shown in Figure 4 we select two class-image pairs and place them on a canvas according to a randomly sampled diagonal direction. Each sample is paired with a prompt of the form "[class A] is located in the [direction] of [class B]".

For evaluation, we use the same T2I prompt format as the training data and follow GenEval [35] protocol. We apply a detection model [18] to localize

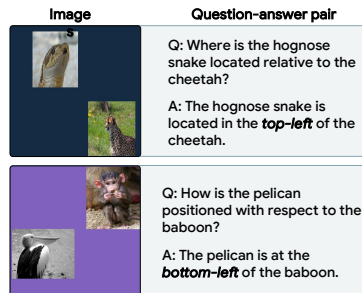


Fig. 4: Example of synthetic dataset for spatial relation.

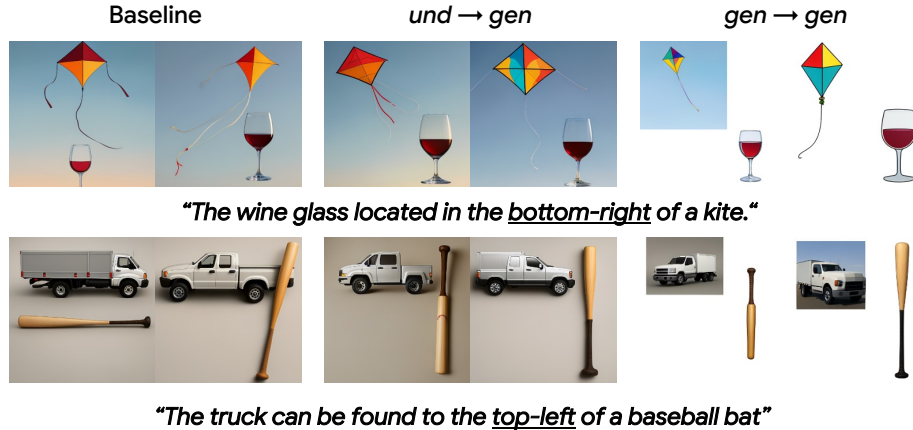


Fig. 5: Qualitative comparison between $und \rightarrow gen$ and $gen \rightarrow gen$ on spatial relation tasks. While generation-only training ($gen \rightarrow gen$) significantly degrades image quality reflecting traits of synthetic training data, $und \rightarrow gen$ preserves generation quality while achieving improved spatial relation accuracy.

Train $\rightarrow$ Test	Acc. $\uparrow$	IS [71] $\uparrow$	FID [39] $\downarrow$
Baseline	67.0	16.86	-
$und \rightarrow gen$	74.0 (+7.0)	17.50	30.95
$gen \rightarrow gen$	80.0 (+13.0)	17.41	32.28

Table 3: Quantitative comparison of understanding-to-generation transfer ($und \rightarrow gen$) versus direct generation training ($gen \rightarrow gen$) on spatial relation. Both $und \rightarrow gen$ and $gen \rightarrow gen$ improve spatial accuracy over the baseline, with $und \rightarrow gen$ attaining a substantial gain (+7.0%) from transferred understanding supervision alone, approaching direct generation training (+13.0%). $und \rightarrow gen$ also stays closer to the baseline image distribution (lower FID)

bounding boxes of the two objects and compute the center points of the detected boxes and verify whether their relative position matches the specified direction. More details of the dataset construction and evaluation pipeline can be found in the Appendix.

Results. Table 3 confirms that understanding-to-generation transferability also holds for spatial relation. The model trained through understanding ($und \rightarrow gen$) achieves higher spatial accuracy than the baseline, showing that relational knowledge transfers from understanding to generation. Training directly on generation data ($gen \rightarrow gen$) yields slightly higher accuracy but at the cost of a higher FID. Qualitative results in Figure 5 illustrate this trade-off: $gen \rightarrow gen$ drifts toward the training distribution, producing images that resemble cropped objects pasted onto a canvas rather than natural scenes, whereas $und \rightarrow gen$ improves spatial correctness while preserving realistic appearance.

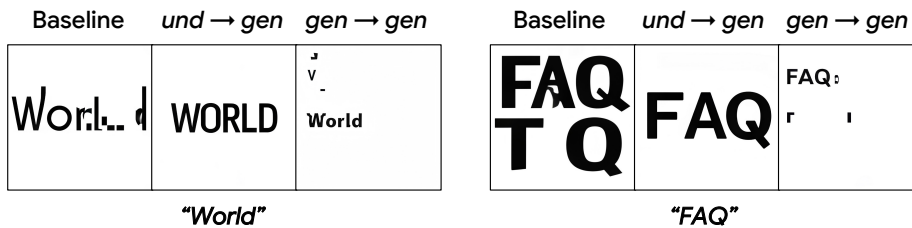


Fig. 7: Qualitative comparison between $und \rightarrow gen$ and $gen \rightarrow gen$ on text generation task. $und \rightarrow gen$ generates clear and accurate text, preserving rendering fidelity.

Train $\rightarrow$ Test	WER $\downarrow$	CER $\downarrow$	Edit Dist. $\downarrow$	BLEU $\uparrow$	METEOR $\uparrow$	F1 $\uparrow$
Baseline	0.654	0.362	36.812	0.277	0.403	0.446
$und \rightarrow gen$	0.641 (-0.013)	0.351 (-0.011)	35.636 (-1.176)	0.274 (-0.003)	0.410 (+0.007)	0.452 (+0.006)
$gen \rightarrow gen$	0.641 (-0.013)	0.367 (+0.005)	37.312 (+0.500)	0.354 (+0.077)	0.533 (+0.130)	0.525 (+0.079)

Table 4: Quantitative comparison of understanding-to-generation transfer ($und \rightarrow gen$) versus direct generation training ($gen \rightarrow gen$) on text generation. We apply an OCR model [1] to the generated images for text extraction and compute WER, CER, edit distance, BLEU, METEOR, F1 between the extracted text and the original prompt. $und \rightarrow gen$ transfer yields improvements in most metrics.

4.3 Text recognition and generation

We next examine text-contained image generation, where the model must render legible characters with accurate strokes, spacing, and glyph structure. Unlike counting or spatial relation, this requires the model to capture highly localized and fine-grained visual details. We investigate whether such fine-grained capabilities can be introduced through understanding and transferred into generation.

Experimental setting. We find that existing OCR datasets tend to focus on either too small, sign-level text in natural scenes [73, 84, 110] or dense, full-page documents [63]. Both may be too hard for models based on discrete image tokenizers, as quantizing such images into discrete tokens discards much of the fine-grained visual details. Thus, we construct a synthetic dataset of 200K image-text pairs by rendering Markdown documents into images (Figure 6). Our rendered dataset provides diverse text content at a moderate difficulty level. For evaluation, we apply an OCR model [1] to read text from the generated image and compute five metrics against the original prompt: word error rate (WER) and character error rate (CER) measure exact correctness at the word and character level.



Fig. 6: Synthetic dataset for text recognition and generation.

Edit Distance captures the minimum number of operations needed to align the prediction with the ground truth. METEOR [8] and F1 assess partial matches by accounting for token overlap and ordering. Details of the dataset construction and evaluation pipeline can be found in the Appendix.

Results. Table 4 summarizes the results. The model trained through understanding ($und \rightarrow gen$) shows improvements over the baseline on most metrics, including WER, CER, Edit Distance, METEOR and F1, confirming that understanding-to-generation transfer also exists for text generation. The qualitative examples in Figure 7 further verify this: the understanding-trained model produces more legible text compared to the baseline. While direct training with generation objective ($gen \rightarrow gen$) yields larger gains in text accuracy, this comes at the cost of generation quality. These results together indicate that the transfer-based approach remains a safer alternative that avoids degrading generative quality when boosting specific visual capabilities.

5 Discussion

While our experiments demonstrate that understanding-to-generation transfer can effectively improve generative capabilities, several natural questions remain. In this section, we further analyze the properties and implications of transferability in UMMs. We first investigate whether the reverse direction, *i.e.*, generation-to-understanding transfer, can also provide meaningful improvements. We then examine whether jointly training understanding and generation tasks could serve as an alternative to transfer-based training. Finally, we analyze whether our understanding training negatively affects general multimodal understanding ability.

5.1 Can generation-to-understanding transfer also be useful?

While the previous section demonstrated that $und \rightarrow gen$ transfer can be exploited to improve specific generative capabilities, we observed that the strength of transferability is not uniform across capabilities. To quantify the strength, we assume direct (*e.g.*, generation) training on the same dataset as the upper bound for transfer (*e.g.*, $und \rightarrow gen$). We then measure transfer strength as the ratio of the improvement from transfer to that from direct training, with the results presented in Table 6. We found that for text generation, the transfer strength of $und \rightarrow gen$ is overall lower than that of counting and spatial relation. This raises the question of what factors determine how well a capability transfers.

We hypothesize that the strength of $und \rightarrow gen$ transfer depends on the nature of the visual knowledge each capability requires. Counting and spatial relation rely on relatively high-level, structural concepts, such as numerical quantity and relative object layout. Consequently, transferring from understanding to generation for these tasks does not strictly require pixel-level precision. The text capability, by contrast, intrinsically demands highly fine-grained visual details, such as stroke geometry, character spacing, and glyph composition. Such

Train → Test	Counting		Spatial Rel.	Text Recognition/Generation						
	Acc. (%) ↑	MAD ↓		WER ↓	CER ↓	Edit Dist. ↓	BLEU ↑	METEOR ↑	F1 ↑	
Baseline	22.0	2.48	61.0	0.129	0.073	7.464	0.835	0.899	0.885	
<i>gen</i> → <i>und</i>	30.0 (+8.0)	1.88 (-0.60)	67.0 (+6.0)	0.118 (-0.011)	0.061 (-0.012)	6.232 (-1.232)	0.845 (+0.010)	0.907 (+0.008)	0.892 (+0.007)	
<i>und</i> → <i>und</i>	57.0 (+35.0)	0.89 (-1.59)	89.0 (+28.0)	0.091 (-0.038)	0.042 (-0.031)	4.276 (-3.188)	0.875 (+0.040)	0.928 (+0.029)	0.912 (+0.027)	

Table 5: Generation-to-understanding transfer (*gen* → *und*) versus direct understanding training (*und* → *und*) on counting, spatial relation, and text recognition for Lumina-DiMOO. Across all three capabilities, *gen* → *und* transfer consistently improves over the baseline.

Transfer Direction	Counting		Spatial Rel.	Text Recog./Gen.		
	Acc.	MAD		WER	METEOR	F1
<i>und</i> → <i>gen</i>	100%	140%	53.8%	100%	5.4%	7.6%
<i>gen</i> → <i>und</i>	22.9%	37.7%	21.4%	28.9%	27.6%	25.9%

Table 6: Transfer strength comparison across tasks and directions. Assuming direct training as an upper bound that transfer can achieve, we quantify transfer strength as the ratio of the improvement obtained by transfer to that obtained by direct training. For counting and spatial-relation, *und* → *gen* transfer is stronger than *gen* → *und*, whereas for text recognition/generation the opposite holds, indicating that the dominant transfer direction may be capability-dependent.

fine-grained precision is difficult to acquire purely through the understanding objective [36, 42, 43, 106, 107], which often relies on partial or contextual cues to perform recognition without fully encoding the exact visual details needed for rendering. This gap between the coarse-grained representations learned through the understanding objective [62, 92] and the fine-grained precision inherently required by the text capability explains the weaker *und* → *gen* transfer.

This reasoning leads to a natural follow-up question: if text generation requires pixel-level precision, and text recognition fundamentally demands the same level of detailed visual discrimination, does the *reverse transfer direction* (*gen* → *und*) *work more effectively* for such capabilities? Since the generation objective inherently provides more detailed supervision than the understanding objective, we expect that learning to generate text forces the model to capture the fine-grained representations highly beneficial for text recognition.

To investigate this, we extend our evaluation to the remaining capabilities—text and spatial relation—by reporting the *gen* → *und* transfer performance on Lumina-DiMOO alongside direct training (*und* → *und*) for comparison (Table 5), and calculate their respective transfer strengths (Table 6). Consistent with our hypothesis, Table 6 demonstrates that while *und* → *gen* transfer remains stronger for counting and spatial relation, the reverse direction (*gen* → *und*) more robustly improves performance for text capabilities compared to *und* → *gen*. These findings confirm that the dominant direction and strength of transferability are capability-dependent and governed by the granularity of supervision required.

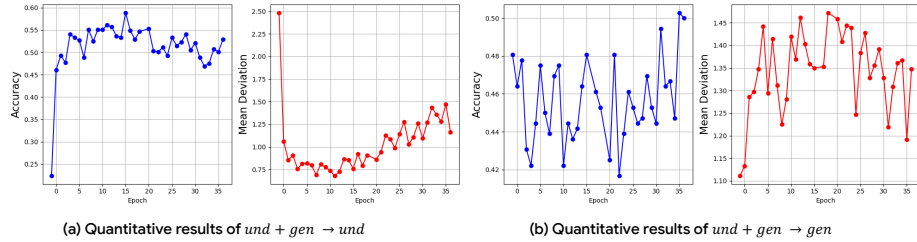


Fig. 8: Trade-off in joint *und+gen* training. (a) Counting understanding accuracy and mean absolute deviation (MAD) over training epochs. (b) Counting generation accuracy and MAD over training epochs. Blue lines denote accuracy and red lines denote MAD.

5.2 Does joint training lead to stronger transferability?

Throughout this work, we examined transferability in UMMs and showed that training a capability through one task can improve it in the other task. This naturally raises the question of whether training both tasks simultaneously leads to improved performance for both.

To investigate this, we trained Lumina-DiMOO on understanding and generation jointly (*und + gen*) for counting. Results shown in Figure 8 indicate that the two objectives do not improve simultaneously during training. Instead, their performance exhibits an oscillatory pattern. Understanding accuracy increases in the first half of training while generation accuracy declines, and the trend reverses in the latter half. This behavior suggests that the presence of transferability does not necessarily translate into consistent gains when the two are optimized simultaneously.

One possible explanation is that understanding and generation require different alignment patterns within the network, which may lead to conflicting optimization signals during joint training, as suggested by [50]. At the same time, our experimental setting may limit how clearly synergy can be observed. We conduct our experiments as post-training on a pretrained model whose capabilities are already largely established, which may lead to different training dynamics from those that arise during large-scale pre-training. In addition, the joint training uses the same concept-focused dataset constructed for the individual experiments in Section 3, where each image-text pair is tightly aligned around *counting*. Such paired data distributions are uncommon in typical multimodal pre-training corpora and may introduce an unfamiliar training configuration when the two objectives are optimized together. Distinguishing between these possibilities is difficult within the scope of post-training experiments and therefore remains an open question.

5.3 Does understanding training harm general multimodal ability?

A potential concern of capability injection through understanding training is that it may negatively affect the model’s overall multimodal ability. Since our

Train (<i>und</i>)	POPE [52]	MMBench [59]	MMM U [104]	MME-P [31]
Baseline	86.7	90.8	57.7	1534
Counting	86.4	90.9	62.0	1492
Spatial Relation	87.3	90.8	61.3	1554
Text recognition	86.7	90.5	61.3	1538

Table 7: Performance on general multimodal benchmarks after understanding-based training. Despite introducing task-specific capabilities through understanding training for *und* $\rightarrow$ *gen* transfer, the model’s general multimodal understanding remains largely preserved across all benchmarks.

approach introduces additional supervision targeting specific capabilities, one might expect such training to overfit to the target task or degrade performance on general multimodal benchmarks.

To examine this, we evaluate the models on several widely used multimodal understanding benchmarks, including POPE [52], MMBench [59], MMMU [104], and MME [31]. The results are summarized in Table 7. Across all benchmarks, we observe that understanding-based training does not degrade general multimodal performance. The scores on POPE and MMBench remain nearly unchanged compared to the baseline model, while performance on MMMU even improves slightly after capability-focused training. These results suggest that injecting capabilities through understanding tasks introduces minimal changes to the model parameters while preserving the model’s general multimodal knowledge.

6 Conclusion

In this work, we studied cross-task transferability between understanding and generation in UMMs. Through controlled experiments, we empirically showed that capabilities transfer between the two tasks and that the strength of this transfer depends on architectural design. We further exploited this as a practical strategy: introducing a target capability through understanding, rather than fine-tuning generation directly, improves capability-specific performance while preserving overall generative quality. We hope our findings provide new insights into the interaction between understanding and generation and open promising directions for post-training in unified multimodal models.

Acknowledgements

This research was supported by Institute of Information & Communications Technology Planning & Evaluation (IITP) grant funded by the Korea government (MSIT) (RS-2019-II190075, RS-2024-00509279, RS-2025-II212068, RS-2023-00227592, RS-2025-02214479, RS-2024-00457882, RS-2025-25441838, RS-2025-02217259, RS-2026-25519202) and the Culture, Sports, and Tourism R&D Program through the Korea Creative Content Agency grant funded by the Ministry of Culture, Sports and Tourism (RS-2024-00345025, RS-2024-00333068), and National Research Foundation of Korea (RS-2024-00346597).

References

1. AI, Z.: Glm-ocr. <https://github.com/zai-org/GLM-OCR> (2026), accessed: 2026-07-01
2. Albergo, M.S., Vanden-Eijnden, E.: Building normalizing flows with stochastic interpolants. arXiv preprint arXiv:2209.15571 (2022)
3. Anthropic: Introducing the next generation of claude (2024), <https://www.anthropic.com/news/claude-3-family>, accessed: 2026-07-01
4. Austin, J., Johnson, D.D., Ho, J., Tarlow, D., Van Den Berg, R.: Structured denoising diffusion models in discrete state-spaces. *Advances in neural information processing systems* **34**, 17981–17993 (2021)
5. Bai, S., Cai, Y., Chen, R., Chen, K., Chen, X., Cheng, Z., Deng, L., Ding, W., Gao, C., Ge, C., et al.: Qwen3-v1 technical report. arXiv preprint arXiv:2511.21631 (2025)
6. Bai, S., Chen, K., Liu, X., Wang, J., Ge, W., Song, S., Dang, K., Wang, P., Wang, S., Tang, J., et al.: Qwen2. 5-v1 technical report. arXiv e-prints pp. arXiv–2502 (2025)
7. Ballard, D.H.: Modular learning in neural networks. In: *Proceedings of the sixth National conference on Artificial intelligence-Volume 1*. pp. 279–284 (1987)
8. Banerjee, S., Lavie, A.: Meteor: An automatic metric for mt evaluation with improved correlation with human judgments. In: *Proceedings of the acl workshop on intrinsic and extrinsic evaluation measures for machine translation and/or summarization*. pp. 65–72 (2005)
9. Binyamin, L., Tewel, Y., Segev, H., Hirsch, E., Rassin, R., Chechik, G.: Make it count: Text-to-image generation with an accurate number of objects. In: *Proceedings of the Computer Vision and Pattern Recognition Conference*. pp. 13242–13251 (2025)
10. Black, K., Janner, M., Du, Y., Kostrikov, I., Levine, S.: Training diffusion models with reinforcement learning. In: *International Conference on Learning Representations*. vol. 2024, pp. 4965–4987 (2024)
11. Black Forest Labs: FLUX.2: Analyzing and enhancing the latent space of FLUX – representation comparison (2025), <https://bfl.ai/research/representation-comparison>, accessed: 2026-07-01
12. Boo, H., Kim, H., Lee, M., Lee, S., Lee, J., Choi, J.H., Cho, H.: Countsteer: Steering attention for object counting in diffusion models. arXiv preprint arXiv:2511.11253 (2025)
13. Chae, J., Kim, J., Choi, J., Kim, K., Hwang, S.: Apt: Adaptive personalized training for diffusion models with limited data. In: *Proceedings of the Computer Vision and Pattern Recognition Conference*. pp. 28619–28628 (2025)
14. Chen, J., Xu, Z., Pan, X., Hu, Y., Qin, C., Goldstein, T., Huang, L., Zhou, T., Xie, S., Savarese, S., et al.: Blip3-o: A family of fully open unified multimodal models-architecture, training and dataset, 2025. URL <https://arxiv.org/abs/2505.09568>
15. Chen, M., Laina, I., Vedaldi, A.: Training-free layout control with cross-attention guidance. In: *Proceedings of the IEEE/CVF winter conference on applications of computer vision*. pp. 5343–5353 (2024)
16. Chen, X., Wu, Z., Liu, X., Pan, Z., Liu, W., Xie, Z., Yu, X., Ruan, C.: Janus-pro: Unified multimodal understanding and generation with data and model scaling. arXiv preprint arXiv:2501.17811 (2025)
17. Chen, Z., Wu, J., Wang, W., Su, W., Chen, G., Xing, S., Zhong, M., Zhang, Q., Zhu, X., Lu, L., et al.: Internvl: Scaling up vision foundation models and aligning

- for generic visual-linguistic tasks. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 24185–24198 (2024)
18. Cheng, B., Misra, I., Schwing, A.G., Kirillov, A., Girdhar, R.: Masked-attention mask transformer for universal image segmentation. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 1290–1299 (2022)
 19. Cui, Y., Chen, H., Deng, H., Huang, X., Li, X., Liu, J., Liu, Y., Luo, Z., Wang, J., Wang, W., et al.: Emu3. 5: Native multimodal models are world learners. arXiv preprint arXiv:2510.26583 (2025)
 20. Dai, X., Hou, J., Ma, C.Y., Tsai, S., Wang, J., Wang, R., Zhang, P., Vandenhende, S., Wang, X., Dubey, A., et al.: Emu: Enhancing image generation models using photogenic needles in a haystack. arXiv preprint arXiv:2309.15807 (2023)
 21. Deitke, M., Clark, C., Lee, S., Tripathi, R., Yang, Y., Park, J.S., Salehi, M., Muennighoff, N., Lo, K., Soldaini, L., et al.: Molmo and pixmo: Open weights and open data for state-of-the-art vision-language models. In: Proceedings of the Computer Vision and Pattern Recognition Conference. pp. 91–104 (2025)
 22. Deng, C., Zhu, D., Li, K., Gou, C., Li, F., Wang, Z., Zhong, S., Yu, W., Nie, X., Song, Z., et al.: Emerging properties in unified multimodal pretraining. arXiv preprint arXiv:2505.14683 (2025)
 23. Deng, J., Dong, W., Socher, R., Li, L.J., Li, K., Fei-Fei, L.: Imagenet: A large-scale hierarchical image database. In: 2009 IEEE conference on computer vision and pattern recognition. pp. 248–255. Ieee (2009)
 24. Deng, N., Gu, L., Ye, S., He, Y., Chen, Z., Li, S., Wang, H., Wei, X., Yang, T., Dou, M., et al.: Internspatial: A comprehensive dataset for spatial reasoning in vision-language models. arXiv preprint arXiv:2506.18385 (2025)
 25. Dong, R., Peng, Y., Qi, Z., Ge, Z., Yang, J., Zhao, L., Sun, J., Zhou, H., Wei, H., Kong, X., et al.: Dreamllm: Synergistic multimodal comprehension and creation. In: International Conference on Learning Representations. vol. 2024, pp. 6666–6702 (2024)
 26. Esser, P., Kulal, S., Blattmann, A., Entezari, R., Müller, J., Saini, H., Levi, Y., Lorenz, D., Sauer, A., Boesel, F., et al.: Scaling rectified flow transformers for high-resolution image synthesis. In: Forty-first international conference on machine learning (2024)
 27. Esser, P., Rombach, R., Ommer, B.: Taming transformers for high-resolution image synthesis. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 12873–12883 (2021)
 28. Fan, Y., Watkins, O., Du, Y., Liu, H., Ryu, M., Boutilier, C., Abbeel, P., Ghavamzadeh, M., Lee, K., Lee, K.: Dpok: Reinforcement learning for fine-tuning text-to-image diffusion models. *Advances in Neural Information Processing Systems* **36**, 79858–79885 (2023)
 29. Fan, Y., He, X., Yang, D., Zheng, K., Kuo, C.C., Zheng, Y., Narayanaraju, S.J., Guan, X., Wang, X.E.: Grit: Teaching mllms to think with images (2025), <https://arxiv.org/abs/2505.15879>, accessed: 2026-07-01
 30. Feng, W., He, X., Fu, T.J., Jampani, V., Akula, A., Narayana, P., Basu, S., Wang, X.E., Wang, W.Y.: Training-free structured diffusion guidance for compositional text-to-image synthesis. arXiv preprint arXiv:2212.05032 (2022)
 31. Fu, C., Chen, P., Shen, Y., Qin, Y., Zhang, M., Lin, X., Yang, J., Zheng, X., Li, K., Sun, X., et al.: Mme: A comprehensive evaluation benchmark for multimodal large language models. arXiv preprint arXiv:2306.13394 (2023)

32. Fu, S., Zhou, J., Chen, Q., Jing, H., Nguyen, H.A., Liu, X., Zeng, Z., Ma, L., Zhang, Q., Wu, Q.: Counting hallucinations in diffusion models. arXiv preprint arXiv:2510.13080 (2025)
33. Ge, Y., Zhao, S., Zhu, J., Ge, Y., Yi, K., Song, L., Li, C., Ding, X., Shan, Y.: Seed-x: Multimodal models with unified multi-granularity comprehension and generation. arXiv preprint arXiv:2404.14396 (2024)
34. Gemini Team Google: Gemini: A family of highly capable multimodal models. arXiv preprint arXiv:2312.11805 (2023)
35. Ghosh, D., Hajishirzi, H., Schmidt, L.: Geneval: An object-focused framework for evaluating text-to-image alignment. *Advances in Neural Information Processing Systems* **36**, 52132–52152 (2023)
36. Ghosh, D., Zhang, Y., Schmidt, L.: Understanding the fine-grained knowledge capabilities of vision-language models. arXiv preprint arXiv:2602.17871 (2026)
37. Gokhale, T., Palangi, H., Nushi, B., Vineet, V., Horvitz, E., Kamar, E., Baral, C., Yang, Y.: Benchmarking spatial relationships in text-to-image generation. arXiv preprint arXiv:2212.10015 (2022)
38. Guo, D., Yang, D., Zhang, H., Song, J., Wang, P., Zhu, Q., Xu, R., Zhang, R., Ma, S., Bi, X., et al.: Deepseek-r1: Incentivizing reasoning capability in llms via reinforcement learning. arXiv preprint arXiv:2501.12948 (2025)
39. Heusel, M., Ramsauer, H., Unterthiner, T., Nessler, B., Hochreiter, S.: Gans trained by a two time-scale update rule converge to a local nash equilibrium. *Advances in neural information processing systems* **30** (2017)
40. Ho, J., Jain, A., Abbeel, P.: Denoising diffusion probabilistic models. *Advances in neural information processing systems* **33**, 6840–6851 (2020)
41. Hudson, D.A., Manning, C.D.: Gqa: A new dataset for real-world visual reasoning and compositional question answering. In: *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*. pp. 6700–6709 (2019)
42. Jung, W., Go, J., Jeon, M., Yoon, S., Kim, J.: Visual funnel: Resolving contextual blindness in multimodal large language models. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 8962–8971 (2026)
43. Khayatkhoei, M., Chhikara, P., Ilievski, F., et al.: Mllms know where to look: Training-free perception of small visual details with multimodal llms. In: *International Conference on Learning Representations*. vol. 2025, pp. 68194–68213 (2025)
44. Kingma, D.P., Welling, M.: Auto-encoding variational bayes. arXiv preprint arXiv:1312.6114 (2013)
45. Kou, S., Jin, J., Liu, Z., Liu, C., Ma, Y., Jia, J., Chen, Q., Jiang, P., Deng, Z.: Orthus: Autoregressive interleaved image-text generation with modality-specific heads. In: *International Conference on Machine Learning*. pp. 31619–31633. PMLR (2025)
46. Labs, B.F., Batifol, S., Blattmann, A., Boesel, F., Consul, S., Diagne, C., Dockhorn, T., English, J., English, Z., Esser, P., et al.: Flux. 1 konteks: Flow matching for in-context image generation and editing in latent space. arXiv preprint arXiv:2506.15742 (2025)
47. Lee, K., Liu, H., Ryu, M., Watkins, O., Du, Y., Boutilier, C., Abbeel, P., Ghavamzadeh, M., Gu, S.S.: Aligning text-to-image models using human feedback. arXiv preprint arXiv:2302.12192 (2023)
48. Lee, K., Kwak, S., Sohn, K., Shin, J.: Direct consistency optimization for robust customization of text-to-image diffusion models. *Advances in neural information processing systems* **37**, 103269–103304 (2024)

49. Li, S., Kallidromitis, K., Gokul, A., Kato, Y., Kozuka, K.: Aligning diffusion models by optimizing human utility. *Advances in Neural Information Processing Systems* **37**, 24897–24925 (2024)
50. Li, T., Lu, Q., Zhao, L., Li, H., Zhu, X., Qiao, Y., Zhang, J., Shao, W.: Unifork: Exploring modality alignment for unified multimodal understanding and generation. *arXiv preprint arXiv:2506.17202* (2025)
51. Li, Y., Wan, P., Han, L., Wang, Y., Nie, L., Zhang, M.: Countdiffusion: Text-to-image synthesis with training-free counting-guidance diffusion. *arXiv preprint arXiv:2505.04347* (2025)
52. Li, Y., Du, Y., Zhou, K., Wang, J., Zhao, W.X., Wen, J.R.: Evaluating object hallucination in large vision-language models. In: *Proceedings of the 2023 conference on empirical methods in natural language processing*. pp. 292–305 (2023)
53. Li, Z., Li, H., Shi, Y., Farimani, A.B., Kluger, Y., Yang, L., Wang, P.: Dual diffusion for unified image generation and understanding. In: *Proceedings of the Computer Vision and Pattern Recognition Conference*. pp. 2779–2790 (2025)
54. Lin, T.Y., Maire, M., Belongie, S., Hays, J., Perona, P., Ramanan, D., Dollár, P., Zitnick, C.L.: Microsoft coco: Common objects in context. In: *European conference on computer vision*. pp. 740–755. Springer (2014)
55. Lipman, Y., Chen, R.T., Ben-Hamu, H., Nickel, M., Le, M.: Flow matching for generative modeling. *arXiv preprint arXiv:2210.02747* (2022)
56. Liu, H., Li, C., Wu, Q., Lee, Y.J.: Visual instruction tuning. *Advances in neural information processing systems* **36**, 34892–34916 (2023)
57. Liu, N., Li, S., Du, Y., Torralba, A., Tenenbaum, J.B.: Compositional visual generation with composable diffusion models. In: *European conference on computer vision*. pp. 423–439. Springer (2022)
58. Liu, X., Gong, C., Liu, Q.: Flow straight and fast: Learning to generate and transfer data with rectified flow. *arXiv preprint arXiv:2209.03003* (2022)
59. Liu, Y., Duan, H., Zhang, Y., Li, B., Zhang, S., Zhao, W., Yuan, Y., Wang, J., He, C., Liu, Z., et al.: Mmbench: Is your multi-modal model an all-around player? In: *European conference on computer vision*. pp. 216–233. Springer (2024)
60. Lou, A., Meng, C., Ermon, S.: Discrete diffusion modeling by estimating the ratios of the data distribution. *arXiv preprint arXiv:2310.16834* (2023)
61. Lv, H., Xiao, J., Li, L.: Pick-and-draw: Training-free semantic guidance for text-to-image personalization. In: *Proceedings of the 32nd ACM International Conference on Multimedia*. pp. 10535–10543 (2024)
62. Marsili, D., Mehta, A., Lin, R.Y., Gkioxari, G.: Same or not? enhancing visual perception in vision-language models. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 17303–17315 (2026)
63. NVIDIA: Llama-nemotron-vlm-dataset-v1. <https://huggingface.co/datasets/nvidia/Llama-Nemotron-VLM-Dataset-v1> (2025), accessed: 2026-07-01
64. OpenAI: Gpt-4 technical report. *arXiv preprint arXiv:2303.08774* (2023)
65. Pan, X., Shukla, S.N., Singh, A., Zhao, Z., Mishra, S.K., Wang, J., Xu, Z., Chen, J., Li, K., Juefei-Xu, F., et al.: Transfer between modalities with metaqueries. *arXiv preprint arXiv:2504.06256* (2025)
66. Phung, Q., Ge, S., Huang, J.B.: Grounded text-to-image synthesis with attention refocusing. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 7932–7942 (2024)
67. 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: *International Conference on Learning Representations*. vol. 2024, pp. 1862–1874 (2024)

68. Radford, A., Kim, J.W., Hallacy, C., Ramesh, A., Goh, G., Agarwal, S., Sastry, G., Askell, A., Mishkin, P., Clark, J., et al.: Learning transferable visual models from natural language supervision. In: International conference on machine learning. pp. 8748–8763. PmLR (2021)
69. Rombach, R., Blattmann, A., Lorenz, D., Esser, P., Ommer, B.: High-resolution image synthesis with latent diffusion models. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 10684–10695 (2022)
70. Sahoo, S.S., Arriola, M., Schiff, Y., Gokaslan, A., Marroquin, E., Chiu, J.T., Rush, A., Kuleshov, V.: Simple and effective masked diffusion language models. *Advances in Neural Information Processing Systems* **37**, 130136–130184 (2024)
71. Salimans, T., Goodfellow, I., Zaremba, W., Cheung, V., Radford, A., Chen, X.: Improved techniques for training gans. *Advances in neural information processing systems* **29** (2016)
72. Shi, Q., Bai, J., Zhao, Z., Chai, W., Yu, K., Wu, J., Tong, Y., Li, X., Li, X., Yan, S.: Muddit: Liberating generation beyond text-to-image with a unified discrete diffusion model. *arXiv preprint arXiv:2505.23606* (2025)
73. Singh, A., Pang, G., Toh, M., Huang, J., Galuba, W., Hassner, T.: Textocr: Towards large-scale end-to-end reasoning for arbitrary-shaped scene text. In: Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. pp. 8802–8812 (2021)
74. Song, C.H., Blukis, V., Tremblay, J., Tyree, S., Su, Y., Birchfield, S.: Robospatial: Teaching spatial understanding to 2d and 3d vision-language models for robotics. In: Proceedings of the Computer Vision and Pattern Recognition Conference. pp. 15768–15780 (2025)
75. Song, J., Meng, C., Ermon, S.: Denoising diffusion implicit models. *arXiv preprint arXiv:2010.02502* (2020)
76. Song, Y., Ermon, S.: Generative modeling by estimating gradients of the data distribution. *Advances in neural information processing systems* **32** (2019)
77. Sun, Q., Yu, Q., Cui, Y., Zhang, F., Zhang, X., Wang, Y., Gao, H., Liu, J., Huang, T., Wang, X.: Emu: Generative pretraining in multimodality. In: International conference on learning representations. vol. 2024, pp. 12352–12380 (2024)
78. Team, C.: Chameleon: Mixed-modal early-fusion foundation models. *arXiv preprint arXiv:2405.09818* (2024)
79. Tong, P., Brown, E., Wu, P., Woo, S., Iyer, A.J.V., Akula, S.C., Yang, S., Yang, J., Middepogu, M., Wang, Z., et al.: Cambrian-1: A fully open, vision-centric exploration of multimodal llms. *Advances in Neural Information Processing Systems* **37**, 87310–87356 (2024)
80. Tong, S., Fan, D., Li, J., Xiong, Y., Chen, X., Sinha, K., Rabbat, M., LeCun, Y., Xie, S., Liu, Z.: Metamorph: Multimodal understanding and generation via instruction tuning. In: Proceedings of the IEEE/CVF International Conference on Computer Vision. pp. 17001–17012 (2025)
81. Touvron, H., Lavril, T., Izacard, G., Martinet, X., Lachaux, M.A., Lacroix, T., Rozière, B., Goyal, N., Hambro, E., Azhar, F., et al.: Llama: Open and efficient foundation language models. *arXiv preprint arXiv:2302.13971* (2023)
82. Tschannen, M., Gritsenko, A., Wang, X., Naeem, M.F., Alabdulmohsin, I., Parthasarathy, N., Evans, T., Beyer, L., Xia, Y., Mustafa, B., et al.: Siglip 2: Multilingual vision-language encoders with improved semantic understanding, localization, and dense features. *arXiv preprint arXiv:2502.14786* (2025)
83. Van Den Oord, A., Vinyals, O., et al.: Neural discrete representation learning. *Advances in neural information processing systems* **30** (2017)

84. Veit, A., Matera, T., Neumann, L., Matas, J., Belongie, S.: Coco-text: Dataset and benchmark for text detection and recognition in natural images. arXiv preprint arXiv:1601.07140 (2016)
85. Wallace, B., Dang, M., Rafailov, R., Zhou, L., Lou, A., Purushwalkam, S., Ermon, S., Xiong, C., Joty, S., Naik, N.: Diffusion model alignment using direct preference optimization. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 8228–8238 (2024)
86. Wang, X., Zhang, Z., Zhang, H., Lin, Z., Zhou, Y., Liu, Q., Zhang, S., Li, Y., Liu, S., Zheng, H., et al.: Hbridge: H-shape bridging of heterogeneous experts for unified multimodal understanding and generation. arXiv preprint arXiv:2511.20520 (2025)
87. Wang, X., Zhang, X., Luo, Z., Sun, Q., Cui, Y., Wang, J., Zhang, F., Wang, Y., Li, Z., Yu, Q., et al.: Emu3: Next-token prediction is all you need. arXiv preprint arXiv:2409.18869 (2024)
88. Wang, Z., Hu, X., Wang, Y., Xiong, F., Zhang, M., Chu, X.: Everything in its place: Benchmarking spatial intelligence of text-to-image models. In: The Fourteenth International Conference on Learning Representations (2026), <https://openreview.net/forum?id=ddFN3lWpIr>
89. Wang, Z., Chen, Z., Gou, C., Li, F., Deng, C., Zhu, D., Li, K., Yu, W., Tu, H., Fan, H., et al.: Lightfusion: A light-weighted, double fusion framework for unified multimodal understanding and generation. arXiv preprint arXiv:2510.22946 (2025)
90. Wu, C., Chen, X., Wu, Z., Ma, Y., Liu, X., Pan, Z., Liu, W., Xie, Z., Yu, X., Ruan, C., et al.: Janus: Decoupling visual encoding for unified multimodal understanding and generation. In: Proceedings of the Computer Vision and Pattern Recognition Conference. pp. 12966–12977 (2025)
91. Wu, J., Jiang, Y., Ma, C., Liu, Y., Zhao, H., Yuan, Z., Bai, S., Bai, X.: Liquid: Language models are scalable and unified multi-modal generators. International Journal of Computer Vision **134**(1), 39 (2026)
92. Wu, S., Zhang, W., Xu, L., Jin, S., Wu, Z., Tao, Q., Liu, W., Li, W., Loy, C.C.: Harmonizing visual representations for unified multimodal understanding and generation. In: Proceedings of the IEEE/CVF International Conference on Computer Vision. pp. 17739–17750 (2025)
93. Xiao, J., Lv, H., Li, L., Wang, S., Huang, Q.: R&b: Region and boundary aware zero-shot grounded text-to-image generation. In: International Conference on Learning Representations. vol. 2024, pp. 26088–26116 (2024)
94. Xie, J., Mao, W., Bai, Z., Zhang, D.J., Wang, W., Lin, K.Q., Gu, Y., Chen, Z., Yang, Z., Shou, M.Z.: Show-o: One single transformer to unify multimodal understanding and generation. In: International Conference on Learning Representations. vol. 2025, pp. 28240–28264 (2025)
95. Xin, Y., Qin, Q., Luo, S., Zhu, K., Yan, J., Tai, Y., Lei, J., Cao, Y., Wang, K., Wang, Y., et al.: Lumina-dimoo: An omni diffusion large language model for multi-modal generation and understanding. arXiv preprint arXiv:2510.06308 (2025)
96. Xu, J., Guo, Z., Hu, H., Chu, Y., Wang, X., He, J., Wang, Y., Shi, X., He, T., Zhu, X., et al.: Qwen3-omni technical report. arXiv preprint arXiv:2509.17765 (2025)
97. Xu, Z., Chen, J., Lin, Z., Pan, X., Huang, L., Zhou, T., Khabsa, M., Wang, Q., Jin, D., Yasunaga, M., et al.: Pisces: An auto-regressive foundation model for image understanding and generation. arXiv preprint arXiv:2506.10395 (2025)

98. Xue, L., Shu, M., Awadalla, A., Wang, J., Yan, A., Purushwalkam, S., Zhou, H., Prabhu, V., Dai, Y., Ryoo, M.S., Kendre, S., Zhang, J., Qin, C., Zhang, S., Chen, C.C., Yu, N., Tan, J., Awalganekar, T.M., Heinecke, S., Wang, H., Choi, Y., Schmidt, L., Chen, Z., Savarese, S., Niebles, J.C., Xiong, C., Xu, R.: xgen-mm(blip-3): A family of open large multimodal models. *arXiv preprint* (August 2024)
99. Yang, J., Yin, D., Zhou, Y., Rao, F., Zhai, W., Cao, Y., Zha, Z.J.: Mmar: Towards lossless multi-modal auto-regressive probabilistic modeling. In: *Proceedings of the Computer Vision and Pattern Recognition Conference*. pp. 7974–7985 (2025)
100. Yang, K., Tao, J., Lyu, J., Ge, C., Chen, J., Shen, W., Zhu, X., Li, X.: Using human feedback to fine-tune diffusion models without any reward model. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. pp. 8941–8951 (2024)
101. Yang, K., Russakovsky, O., Deng, J.: Spatialsense: An adversarially crowdsourced benchmark for spatial relation recognition. In: *International Conference on Computer Vision (ICCV)* (2019)
102. Yang, L., Tian, Y., Li, B., Zhang, X., Shen, K., Tong, Y., Wang, M.: Mmada: Multimodal large diffusion language models. *arXiv preprint arXiv:2505.15809* (2025)
103. Yoon, H., Jung, J., Kim, J., Choi, H., Shin, H., Lim, S., An, H., Kim, C., Han, J., Kim, D., et al.: Visual representation alignment for multimodal large language models. *arXiv preprint arXiv:2509.07979* (2025)
104. Yue, X., Ni, Y., Zhang, K., Zheng, T., Liu, R., Zhang, G., Stevens, S., Jiang, D., Ren, W., Sun, Y., et al.: Mmmu: A massive multi-discipline multimodal understanding and reasoning benchmark for expert agi. In: *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*. pp. 9556–9567 (2024)
105. Zhai, X., Mustafa, B., Kolesnikov, A., Beyer, L.: Sigmoid loss for language image pre-training. In: *Proceedings of the IEEE/CVF international conference on computer vision*. pp. 11975–11986 (2023)
106. Zhang, J., Hu, J., Khayatkhoei, M., Ilievski, F., Sun, M.: Exploring perceptual limitation of multimodal large language models. *arXiv preprint arXiv:2402.07384* (2024)
107. Zhang, J., Khayatkhoei, M., Chhikara, P., Ilievski, F.: Visual cropping improves zero-shot question answering of multimodal large language models. In: *R0-FoMo: Robustness of Few-shot and Zero-shot Learning in Large Foundation Models* (2023)
108. Zhang, J., Li, T., Li, L., Yang, Z., Cheng, Y.: Are unified vision-language models necessary: Generalization across understanding and generation. *arXiv preprint arXiv:2505.23043* (2025)
109. Zhang, P., Li, C., Qiao, L., Cheng, Z., Pu, S., Niu, Y., Wu, F.: Vsr: a unified framework for document layout analysis combining vision, semantics and relations. In: *International conference on document analysis and recognition*. pp. 115–130. Springer (2021)
110. Zhang, Y., Gueguen, L., Zharkov, I., Zhang, P., Seifert, K., Kadlec, B.: Uber-text: A large-scale dataset for optical character recognition from street-level imagery. In: *SUNw: Scene Understanding Workshop - CVPR 2017, Hawaii, U.S.A.* (2017)
111. Zhao, S., Zhang, X., Guo, J., Hu, J., Duan, L., Fu, M., Chng, Y.X., Wang, G.H., Chen, Q.G., Xu, Z., et al.: Unified multimodal understanding and generation models: Advances, challenges, and opportunities. *arXiv preprint arXiv:2505.02567* (2025)

112. Zhou, C., Yu, L., Babu, A., Tirumala, K., Yasunaga, M., Shamis, L., Kahn, J., Ma, X., Zettlemoyer, L., Levy, O.: Transfusion: Predict the next token and diffuse images with one multi-modal model. arXiv preprint arXiv:2408.11039 (2024)

Appendix

Table of Contents

A	Training Details	2
B	Task Details	3
B.1	Counting	3
B.2	Spatial Relation	4
B.3	Text Recognition and Generation	7
C	Additional Results and Analysis	10
C.1	Additional results on shared transformer with unified image encoder using MMaDA [102]	10
C.2	Attention-level Analysis of Transferability	11
C.3	Quantitative Analysis of Generation Distribution Shift After Fine-Tuning	13
C.4	Ablation on how LoRA rank affects transfer strength	14
C.5	Can transferred capabilities be naturally combined at generation time to enable more complex generation?	15
D	Additional Qualitative Results	16
E	Limitations	19

A Training Details

In the experiments of Section 3, we evaluate four representative UMMs spanning different architectural families (Figure 1): BAGEL [22] (7B), Janus-Pro [16] (7B), Lumina-DiMOO [95] (8B), and BLIP3-o [14] (8B). All models fall within the 7B–8B parameter range to enable a fair comparison across architectures. To investigate whether transferability emerges in the shared LLM backbone, we perform LoRA-based supervised fine-tuning on top of publicly available pretrained weights. We apply LoRA with rank 128 to the query, key, and value projection layers as well as the MLP layers of the transformer blocks. All other components—including the image encoder, projector/aligner, language modeling head, and diffusion head—are kept frozen. We use 8-bit AdamW with bfloat16 training across all runs. We train for 15 epochs with early stopping to get the best checkpoint and use one capability-specific dataset per training run without additional data balancing. Unless otherwise specified, we follow the official codebases and recommended training configurations of each model. The model-specific settings are described below.

BAGEL [22]. We use a learning rate of 1×10^{-4} and retain the model’s original dynamic image resizing strategy.

Janus-Pro [16]. We use a learning rate of 4×10^{-5} . For both understanding and generation, we preserve the aspect ratio, resize the longer side to 384 pixels, and pad the shorter side, resulting in a 384×384 input.

Lumina-DiMOO [95]. We use a learning rate of 3×10^{-6} . For both understanding and generation, we preserve the aspect ratio, resize the longer side to 512 pixels, and pad the shorter side, resulting in a 512×512 input. We use the same training configuration for the experiments in Sections 4 and 5.1.

BLIP3-o [14]. We use a learning rate of 2×10^{-5} . Since BLIP3-o employs separate transformer branches for understanding and generation, we fine-tune only the branch corresponding to each task. We retain the original image processing strategy of Qwen2.5-VL-7B-Instruct [6], which serves as the LLM backbone of BLIP3-o.

B Task Details

B.1 Counting

Dataset details. Our *visual instruction tuning dataset* for counting consists of 65.6K images and approximately 200K question-answer pairs. We derive the VQA pairs from the training sets of PixMo-Count [21] and PixMo-Points [21]. For counting understanding, we directly adopt the original question-answer pairs from PixMo-Count, as it was originally designed for counting tasks. To scale up the training data, we additionally incorporate samples from PixMo-Points. Since PixMo-Points is a point prediction dataset rather than a counting dataset, we reformulate it for counting by treating the number of points to be predicted as the object count. Note that PixMo-Points originally contains annotations of more than 1K points per image. We find that such extreme cases introduce excessive difficulty during training. To maintain a reasonable difficulty level, we filter out samples whose object count exceeds 20. For the *counting-aware generation dataset*, we construct the training set from the same 65.6K images used in the understanding split to ensure a controlled experimental setting. Specifically, we leverage the object count annotations from the understanding training set and use Qwen3-VL [5] to generate count-aware generation descriptions conditioned on these annotations. Table 8 presents representative samples from both the counting understanding and generation training sets.

Understanding evaluation. We evaluate counting understanding using 540 QA pairs constructed from the test split of PixMo-Count. Each prompt includes a **Response Example** field that instructs the model to place the predicted count inside markdown markers **** ****. Each QA pair follows the template:

Question: "How many {OBJECT}s are there in the image? Response
Example: There are ****<COUNT>**** of {OBJECT}s in the image."
Answer: "There are ****{COUNT}**** of {OBJECT}s in the image."

During evaluation, we parse the substring enclosed by the markers and compare it with the ground-truth count. If parsing fails, we fall back to checking whether the ground-truth count appears anywhere in the model output, either as an Arabic numeral or as an English number word.

Generation evaluation. Following the prompt generation protocol of GenEval [35], we construct a total of 90 prompts by combining 10 object classes and 9 count words. Each prompt follows the template:

"a photo of {COUNT} {OBJECT}s"

where $\text{OBJECT} \in \{\text{person, dog, cat, car, cup, chair, book, bottle, apple, tie}\}$ and $\text{COUNT} \in \{\text{two, three, four, five, six, seven, eight, nine, ten}\}$.

Image	PixMo-Points/Counts	Understanding QA	Description
	<pre><points x1="9.388" y1="75.817" x2="26.667" y2="78.214" x3="47.483" y3="80.392" x4="65.578" y4="77.124" x5="93.469" y5="73.638" alt="thumbs">, <points x1="89.932" y1="57.081" x2="70.612" y2="56.209" x3="51.429" y3="56.427" x4="34.558" y4="60.784" x5="14.014" y5="59.477" alt="index fingers"> ...</pre>	Q: How many thumbs are there in the image? Response Example : There are <number> thumbs in the image. A: There are ** 5 ** thumbs in the image.	5 thumbs, 5 index fingers, 5 middle fingers, 5 pinky fingers, 5 number of right thumbs, 5 right pointer fingers are arranged in a horizontal row against a solid black background. Each hand displays exactly five fingers, including the thumb, with the right thumb and right pointer finger clearly visible on each. The skin tones vary slightly, and the lighting highlights the natural texture and creases of the palms and fingers. A total of five hands are shown, each positioned to emphasize the count of fingers and thumbs, with no additional objects present.
	<pre><points x1="81.15" y1="35.982" x2="67.25" y2="46.327" x3="68.85" y3="21.589" x4="62.65" y4="11.094" x5="42.75" y5="34.933" x6="16.35" y6="38.081" alt="people">, <points x1="68.35" y1="20.69" x2="63.95" y2="25.937" x3="61.45" y3="6.147" x4="39.15" y4="11.394" x5="14.75" y5="14.543" alt="heads"> ...</pre>	Q: How many people are there in the image? Response Example : There are <number> people in the image. A: There are ** 6 ** people in the image.	Six people, five heads, six people on court, eleven shoes, ten arms, nine socks, six humans, nine letters, seven hands, five numbers, eight eyes, six white lines, nine elbows, five shorts. The scene shows exactly six people on court, each wearing blue uniforms with white lettering and numbers, standing on a wooden floor marked with six white lines. A total of eleven shoes are visible, paired with nine socks, and the players' arms and elbows are clearly defined. The uniforms feature five numbers and nine letters, with eight eyes and seven hands visible. Five shorts are worn by the players, and the composition is lit evenly under bright overhead lighting.
	<pre><points x1="15.918" y1="63.424" x2="28.613" y2="65.377" x3="42.099" y3="66.796" x4="51.367" y4="64.591" x5="63.477" y5="66.537" x6="75.977" y6="65.24" x7="87.695" y7="65.377" alt="people sitting">, <points x1="22.363" y1="23.476" x2="89.746" y2="39.559" ... alt="female"> ...</pre>	Q: How many people sitting are there in the image? Response Example : There are <number> people sitting in the image. A: There are ** 7 ** people sitting in the image.	7 people sitting, 13 female, 15 males, 6 people with glasses, 14 ties, 8 scarves, 5 people wearing glasses, 7 wooden chairs, 13 Females, 7 Sitting people. The group is arranged in three rows, with exactly 7 people sitting on 7 wooden chairs. A total of 13 females and 15 males are visible, wearing white shirts and blue patterned ties or scarves. Exactly 6 people have glasses, and 5 of them are wearing glasses, with 14 ties and 8 scarves distributed among them. The background is a textured blue wall, and the polished wooden floor reflects the studio lighting.
	<pre><points x1="57.325" y1="47.714" x2="58.008" y2="50.741" x3="85.059" y3="65.194" x4="82.813" y4="70.468" alt="pillows">, <points x1="42.188" y1="67.782" x2="42.286" y2="76.962" x3="42.872" y3="83.602" x4="49.415" y4="81.356" x5="65.919" y5="67.489" alt="boxes"> ...</pre>	Q: How many pillows are there in the image? Response Example : There are <number> pillows in the image. A: There are ** 4 ** pillows in the image.	4 pillows, 5 boxes, 8 shoes, 11 hangers, 6 leaves, 13 folded items, 15 folded textile items, 10 shirts hanging up, 11 hanger body, 11 metal hooks, 6 leads, 9 grey buttons on white shirt on end, 10 clothes on clothes hangers. The metal-framed shelving unit holds exactly 10 shirts hanging up and a total of 15 folded textile items. Four pillows rest on lower shelves, and exactly 8 shoes are arranged neatly. A total of 5 boxes, including one red and one black, sit on the bottom shelf. Six leaves are visible near the window, and the 11 hangers support 10 clothes on clothes hangers.

Table 8: Examples of counting dataset.

For each prompt, we generate four images with different random seeds and report the average accuracy across all generated samples. To obtain object counts from generated images, we follow the same detection-based approach as GenEval [35].

B.2 Spatial Relation

Dataset details. Our *visual instruction tuning dataset* for spatial relation understanding consists of 200K image-question-answer pairs. We construct composite images using object crops drawn from the ImageNet-1K [23] training set. Specifically, for each sample, we randomly select two images from different ImageNet classes and paste them onto a 512×512 canvas with a random color background. We resize each selected image to a random size (between 150 and 300 pixels per side, with an aspect ratio constrained near 1:1) and place them such that the two images do not overlap and exhibit a clearly unambiguous relative spatial relationship.

We consider four relative positions—**top-left**, **top-right**, **bottom-left**, and **bottom-right**—defined by the difference between the centroids of the two images, and enforce an even distribution of 50K samples per position. For each sample, a question-answer pair is generated by randomly selecting one of 10 diverse QA templates that ask about the relative position of object *A* with respect to object *B*. Each question is appended with a **Response format** hint that specifies the expected answer structure. For the *spatial-relation-aware generation dataset*, we construct the training set from the same images used in the understanding split. We leverage the spatial relation annotations from the understanding training set and use Qwen3-VL [5] to generate spatially-aware generation prompts conditioned on these annotations. Table 9 presents representative samples from both the spatial relation understanding and generation training sets.

Each QA pair follows the template:

Question: "{QUESTION}\nResponse format: {RESPONSE_FORMAT}"
Answer: "{ANSWER}"

For example, given a template sampled at random:

Question: "Where is the {OBJECT_A} located relative to the {OBJECT_B}?"
Response format: The {OBJECT_A} is located in the top-left/top-right/bottom-left/bottom-right of the {OBJECT_B}."
Answer: "The {OBJECT_A} is located in the {POSITION} of the {OBJECT_B}."

where $POSITION \in \{\text{top-left}, \text{top-right}, \text{bottom-left}, \text{bottom-right}\}$, and $OBJECT_A, OBJECT_B$ are ImageNet class names.

Understanding evaluation. We evaluate spatial relation understanding following the multiple-choice setting with 250 validation samples. Each original open-ended QA pair is converted into a four-way multiple-choice question by removing the **Response format** line and appending the following fixed choices:

- a. top-left
- b. top-right
- c. bottom-left
- d. bottom-right

The ground-truth answer is mapped to the corresponding option letter according to the target spatial relation.

An example evaluation pair after this transformation is:

Question:
 "Where is the goose located relative to the water buffalo?"
 a. top-left
 b. top-right

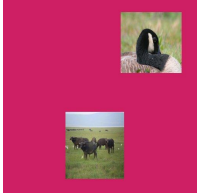
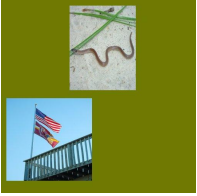
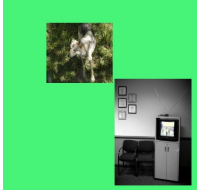
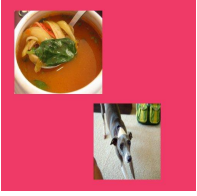
Image	Prompt	Image	Prompt
	Q: Where is the goose located relative to the water buffalo? Response format: The goose is located in the top-left/top-right/bottom-left/bottom-right of the water buffalo. A: The goose is located in the top-right of the water buffalo.		Q: If you look at the flagpole, where is the thunder snake? Response format: The thunder snake is to the top-left/top-right/bottom-left/bottom-right of the flagpole. A: The thunder snake is to the top-right of the flagpole.
	Q: Looking at the image, where does the timber wolf appear relative to the television? Response format: The timber wolf appears to the top-left/top-right/bottom-left/bottom-right of the television. A: The timber wolf appears to the top-left of the television.		Q: How is the paper towel positioned with respect to the school bus? Response format: The paper towel is at the top-left/top-right/bottom-left/bottom-right of the school bus. A: The paper towel is at the top-left of the school bus.
	Q: Describe the spatial relationship between the tailed frog and the hermit crab. Response format: The tailed frog is in the top-left/top-right/bottom-left/bottom-right relative to the hermit crab. A: The tailed frog is in the bottom-right relative to the hermit crab.		Q: How is the Italian greyhound positioned with respect to the soup bowl? Response format: The Italian greyhound is at the top-left/top-right/bottom-left/bottom-right of the soup bowl. A: The Italian greyhound is at the bottom-right of the soup bowl.
	Q: Looking at the image, where does the trailer truck appear relative to the bighorn? Response format: The trailer truck appears to the top-left/top-right/bottom-left/bottom-right of the bighorn. A: The trailer truck appears to the bottom-left of the bighorn.		Q: What is the position of the flagpole relative to the necklace? Response format: The flagpole is positioned to the top-left/top-right/bottom-left/bottom-right of the necklace. A: The flagpole is positioned to the bottom-left of the necklace.

Table 9: Examples of spatial relation dataset.

c. bottom-left
d. bottom-right"
Answer: "b"

Generation evaluation. We generate 100 prompts describing spatial relationships between two objects, following the same answer-style declarative format used during training. Each prompt is a sentence that states the diagonal relation between two object classes, *e.g.*, “The truck can be found to the top-left of the baseball.” The full list of evaluation prompts is provided in Fig. 9. For each prompt, we generate four images with different random seeds and report the average accuracy across all generated samples.

For evaluation, we follow the same detection-based approach as GenEval [35], obtaining bounding boxes for both objects. The relative position between objects is determined by the difference between bounding box centroids: if object *A* is

centered at (x_A, y_A) and object B is centered at (x_B, y_B) , following standard image coordinate conventions where the y -axis increases downward, then:

- **top-left:** $x_A < x_B$ and $y_A < y_B$
- **top-right:** $x_A > x_B$ and $y_A < y_B$
- **bottom-left:** $x_A < x_B$ and $y_A > y_B$
- **bottom-right:** $x_A > x_B$ and $y_A > y_B$

A sample is counted as correct if the detected relation matches the one specified in the prompt. If the detector fails to localize either object, the sample is excluded from evaluation.

B.3 Text Recognition and Generation

Dataset details. We construct the synthetic OCR training set from a public English sentence corpus on Hugging Face. Specifically, we use the training split of `agentlans/high-quality-english-sentences` and keep the first 200K samples. Each sample is normalized by replacing line breaks with spaces, trimming whitespace, and truncating the target string to at most 120 characters. The normalized text is then formatted as `# <sentence>`, *i.e.*, a single level-1 Markdown heading rather than a full free-form document. We convert this markdown string to HTML using the Python `markdown` library with standard extensions for fenced code blocks, tables, automatic line breaks, and list handling, and then rasterize the HTML into a PNG image with `wkhtmltoimage`. We define four hand-designed CSS templates, `clean_light`, `dark`, `sepia`, and `blue_gray`, which vary the background, text color, and font family. During understanding training, one of these templates is sampled at random for each example, whereas generation training uses the `clean_light` template throughout. Figure 6 shows representative examples.

We use the same underlying text-image pair for both modalities. For understanding, each rendered image is paired with the fixed instruction `Extract all text from the image.`, and the normalized sentence is used as the target. For generation, the same sentence is embedded into a fixed caption template describing a Mathpix-style rendering with sharp, legible text in natural reading order. This construction lets us compare *und* and *gen* transfer on identical textual content while changing only the direction of supervision. The explicit understanding and generation formats are shown below.

Understanding question: "Extract all text from the image."

Understanding answer: "{SENTENCE}"

Generation prompt: "A Mathpix Markdown format with sharp, legible black text. High-resolution typography, top-down view. The text is rendered in natural left-to-right, top-to-bottom reading order. The text reads: {SENTENCE}"

Understanding evaluation. We evaluate text recognition on the first 250 samples from the test split of the same corpus. For each example, we normalize the reference sentence in the same way as during training and re-render it as a `clean_light` markdown image. At evaluation time, we use the explicit instruction `Extract all text from the image in reading order.` and otherwise follow the default understanding prompt format and decoding procedure of each model. After generation, both prediction and reference are lowercased, line breaks are replaced with spaces, and leading and trailing whitespace is stripped. We report WER, CER, BLEU, METEOR, average edit distance, and token-level precision, recall, and F1.

Generation evaluation. We evaluate text generation on the first 250 samples from the test split. For each target sentence, we construct the same Mathpix-style caption template shown above and generate an image following the default text-to-image inference setting of each model. We then apply GLM-OCR [1] to the generated image and extract text. The extracted text and the ground-truth sentence are normalized with the same lowercasing and whitespace processing used in understanding evaluation, and we report the same set of metrics: WER, CER, BLEU, METEOR, average edit distance, and token-level precision, recall, and F1.

1. The dog located in the top-left of a teddy bear.
2. The wine glass located in the bottom-right of a kite.
3. The couch is located in bottom-left of a cup.
4. The laptop can be found to the top-right of a cow.
5. The fork is located in top-right of a hair drier.
6. The tie can be found to the bottom-left of a baseball bat.
7. The stop sign is located in top-right of a fork.
8. The bird is located in bottom-left of a skateboard.
9. The an apple located in the bottom-right of a tv.
10. The train located in the bottom-right of a potted plant.
11. The truck can be found to the top-right of a refrigerator.
12. The tv remote is located in bottom-right of a cow.
13. The bottle can be found to the bottom-left of a train.
14. The dog is located in top-right of a cow.
15. The skateboard located in the bottom-left of a person.
16. The baseball glove is located in bottom-left of an umbrella.
17. The dining table located in the top-right of an oven.
18. The hot dog can be found to the top-left of a suitcase.
19. The bus is located in top-left of a toothbrush.
20. The backpack located in the top-right of a sandwich.
21. The cake is located in bottom-right of a baseball bat.
22. The dog located in the top-left of a tie.
23. The suitcase can be found to the bottom-left of a boat.
24. The bear located in the bottom-right of a clock.
25. The tv remote can be found to the top-right of an umbrella.
26. The sports ball can be found to the top-left of an umbrella.
27. The train located in the top-right of a dining table.
28. The hair drier is located in bottom-right of an elephant.
29. The tennis racket can be found to the bottom-left of a spoon.
30. The wine glass can be found to the bottom-left of a hot dog.
31. The computer mouse can be found to the top-right of a bench.
32. The carrot can be found to the top-right of an orange.
33. The kite located in the bottom-right of a toothbrush.
34. The toaster is located in top-left of a traffic light.
35. The cat is located in top-left of a baseball glove.
36. The skis located in the top-left of a zebra.
37. The stop sign located in the bottom-left of a chair.
38. The stop sign located in the bottom-right of a parking meter.
39. The hot dog located in the top-left of a skateboard.
40. The pizza is located in bottom-left of a computer keyboard.
41. The hair drier can be found to the top-left of a toilet.
42. The cow can be found to the bottom-right of a stop sign.
43. The suitcase is located in top-right of a skis.
44. The book located in the bottom-left of a laptop.
45. The toothbrush is located in top-left of a pizza.
46. The toilet can be found to the top-right of a kite.
47. The tie is located in top-right of a sink.
48. The bird can be found to the bottom-right of a couch.
49. The bed can be found to the bottom-left of a sports ball.
50. The an elephant is located in top-left of a surfboard.
51. The frisbee can be found to the bottom-left of a motorcycle.
52. The vase located in the bottom-right of a fire hydrant.
53. The zebra can be found to the top-left of an elephant.
54. The bench can be found to the bottom-right of a bear.
55. The donut located in the top-left of a bench.
56. The frisbee is located in bottom-left of a horse.
57. The computer keyboard located in the bottom-right of a snowboard.
58. The tv is located in bottom-right of a cow.
59. The an elephant is located in top-left of a horse.
60. The suitcase can be found to the top-right of a banana.
61. The train is located in top-left of an airplane.
62. The cat is located in bottom-right of a backpack.
63. The backpack is located in bottom-left of a cake.
64. The sandwich is located in bottom-left of a knife.
65. The bicycle is located in top-right of a parking meter.
66. The knife located in the top-left of a suitcase.
67. The hot dog located in the bottom-left of a knife.
68. The zebra located in the top-left of a parking meter.
69. The chair can be found to the bottom-right of a zebra.
70. The cow is located in top-left of an airplane.
71. The cup can be found to the bottom-right of an umbrella.
72. The zebra is located in bottom-left of a computer keyboard.
73. The zebra is located in bottom-left of a broccoli.
74. The laptop is located in bottom-right of a sports ball.
75. The truck can be found to the top-left of a baseball bat.
76. The refrigerator is located in top-right of a baseball bat.
77. The tv located in the bottom-right of a baseball bat.
78. The baseball glove located in the top-left of a bear.
79. The refrigerator is located in bottom-right of a scissors.
80. The dining table located in the bottom-right of a suitcase.
81. The parking meter located in the bottom-right of a broccoli.
82. The frisbee located in the bottom-right of a truck.
83. The pizza located in the top-left of a banana.
84. The bus located in the bottom-left of a boat.
85. The cell phone can be found to the bottom-right of a tennis racket.
86. The horse located in the top-right of a broccoli.
87. The broccoli located in the bottom-left of a bottle.
88. The vase can be found to the bottom-left of a horse.
89. The bear is located in top-right of a spoon.
90. The zebra located in the top-left of a bed.
91. The cow can be found to the bottom-left of a laptop.
92. The bed located in the top-left of a frisbee.
93. The tie located in the top-right of a motorcycle.
94. The laptop can be found to the bottom-left of a tv.
95. The cell phone can be found to the bottom-left of a chair.
96. The couch is located in top-left of a potted plant.
97. The clock is located in bottom-left of a tv.
98. The couch is located in bottom-right of a vase.
99. The donut is located in bottom-left of a cat.
100. The couch can be found to the top-left of a toaster.

Fig. 9: Full list of spatial relation evaluation prompts. All 100 prompts used for spatial relation evaluation. Each prompt describes the relative position of two COCO objects using one of the predefined position templates.

C Additional Results and Analysis

C.1 Additional results on shared transformer with unified image encoder using MMaDA [102]

Train $\rightarrow$ Test	Counting		Spatial Rel.
	Acc. (%) $\uparrow$	MAD $\downarrow$	Acc. (%) $\uparrow$
<i>Understanding Evaluation</i>			
Baseline	38.7	1.16	40.8
<i>gen</i> $\rightarrow$ <i>und</i>	43.9 (+5.2)	1.07 (-0.09)	48.8 (+8.0)
<i>Generation Evaluation</i>			
Baseline	15.8	3.44	33.0
<i>und</i> $\rightarrow$ <i>gen</i>	20.0 (+4.2)	3.23 (-0.21)	46.3 (+13.3)

Table 10: Bidirectional transferability of MMaDA [102]. MMaDA shares the same architecture and generation objective as Lumina-DiMOO (Fig. 1-(a), a shared transformer with a unified image encoder). We report counting (accuracy and MAD) and spatial-relation (accuracy) results under both *und* $\rightarrow$ *gen* transfer and *gen* $\rightarrow$ *und* transfer. Consistent with Lumina-DiMOO, MMaDA improves on both tasks in both directions, indicating that transferability is shared across models of this architectural family.

Due to computational cost and limited public code availability, our main experiments select one representative model from each architectural family in Fig. 1. Among them, Lumina-DiMOO, which belongs to the shared transformer with a unified image encoder family (Fig. 1-(a)), exhibits the strongest cross-task transferability, so we center our analysis on it in Secs. 4 and 5. To verify that this behavior is not specific to a single model, we additionally evaluate MMaDA [102], which shares the same architecture (Fig. 1-(a)) and generation objective but differs in training data and scale.

As shown in Table 10, MMaDA exhibits transferability in both directions, mirroring the trend observed for Lumina-DiMOO. On the understanding side, *gen* $\rightarrow$ *und* transfer improves counting accuracy from 38.7% to 43.9% (+5.2) and reduces MAD from 1.16 to 1.07 (−0.09), while spatial-relation accuracy rises from 40.8% to 48.8% (+8.0). On the generation side, *und* $\rightarrow$ *gen* transfer raises counting accuracy from 15.8% to 20.0% (+4.2) with MAD dropping from 3.44 to 3.23 (−0.21), and spatial-relation accuracy improves substantially from 33.0% to 46.3% (+13.3). The gains are consistent across both tasks and both transfer directions, which is in line with cross-task knowledge sharing rather than task-specific artifacts.

Although empirical, these results offer additional support for our main observation. A second, independently trained model from the shared transformer with a unified image encoder family (Fig. 1-(a)) again exhibits bidirectional transferability, consistent with the trend in Table 1. We read this as further evidence that transferability is associated with this architectural design rather than with any individual model. We note, however, that a stronger claim would require

broader coverage across more models and families, which remains constrained by computational cost and limited public code availability.

C.2 Attention-level Analysis of Transferability

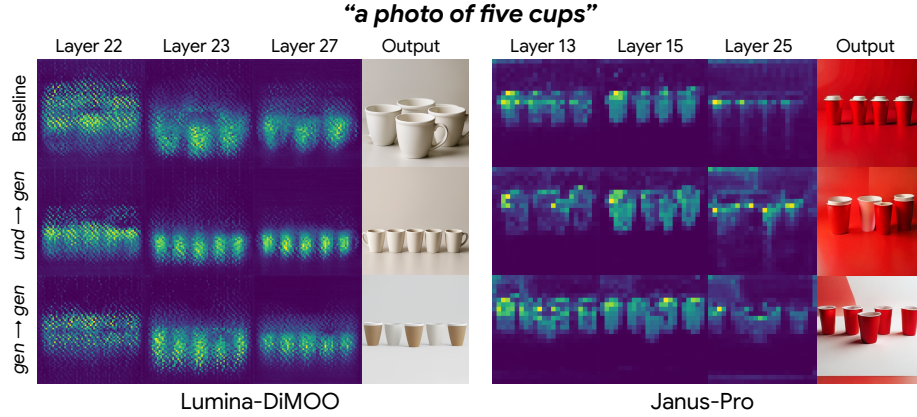


Fig. 10: Attention analysis on Lumina-DiMOO and Janus-Pro on counting. We visualize the attention maps associated with the text token “cups” in the prompt “a photo of five cups”. The attention scores are averaged over heads and reshaped into the spatial layout of image tokens. In Lumina-DiMOO, which exhibits strong transferability, *und* $\rightarrow$ *gen* (middle row) produces sharper and more prompt-aligned attention over cup regions than the baseline (top row), to a degree comparable to direct *gen* $\rightarrow$ *gen* (bottom row). In Janus-Pro, where such transfer is absent, the attention maps remain largely unchanged across training settings.

Beyond the output-level evaluation, we further inspect the intermediate behavior of Lumina-DiMOO [95] and Janus-Pro [16] by qualitatively analyzing their multimodal attention maps on the counting task (Fig. 10). Specifically, we visualize the text-to-image attention scores in the multimodal attention layers, where the query corresponds to a text token and the keys correspond to image tokens involved in generation. Since the two models follow different generation paradigms, we extract attention maps at model-specific generation stages. For Lumina-DiMOO, a discrete diffusion model, we cache the multimodal attention maps at each denoising step and visualize the attention scores at an early denoising stage (after two denoising steps out of total 64 denoising steps) when the coarse image structure begins to emerge. For Janus-Pro, an autoregressive image generation model, we visualize the attention scores after the image generation process is completed.

The resulting attention patterns reveal a qualitative difference that aligns with the transferability trends observed in Table 1. In Lumina-DiMOO, where understanding-to-generation transfer is strong, the *und* $\rightarrow$ *gen* model produces

attention that is not only more localized over the generated cup regions compared with the baseline, but also more consistent with the count-object relationship specified in the prompt. This pattern appears across several layers and is visually comparable to that of direct generation training ($gen \rightarrow gen$). In contrast, Janus-Pro, which does not show clear understanding-to-generation transfer, exhibits much weaker changes in text-to-image attention across the baseline, $und \rightarrow gen$ and $gen \rightarrow gen$ models. These observations suggest that, when transferability emerges, understanding-side training may also affect how textual concepts are routed to image tokens during generation.

C.3 Quantitative Analysis of Generation Distribution Shift After Fine-Tuning

Task-Specific Prompts						COCO Validation Prompts							
Model	Baseline		<i>und</i> $\rightarrow$ <i>gen</i>		<i>gen</i> $\rightarrow$ <i>gen</i>		Model	Baseline		<i>und</i> $\rightarrow$ <i>gen</i>		<i>gen</i> $\rightarrow$ <i>gen</i>	
	IS $\uparrow$	FID $\downarrow$	IS $\uparrow$	FID $\downarrow$	IS $\uparrow$	FID $\downarrow$		IS $\uparrow$	FID $\downarrow$	IS $\uparrow$	FID $\downarrow$	IS $\uparrow$	FID $\downarrow$
<i>Counting</i>						<i>Counting</i>							
Bagel	10.37	-	10.39	9.03	10.80	12.85	Bagel	19.32	-	18.94	15.91	18.97	21.89
Janus-Pro	8.45	-	8.61	13.69	9.23	43.00	Janus-Pro	17.31	-	18.74	37.10	17.72	43.89
Lumina-DiMOO	7.93	-	7.88	14.74	9.20	80.45	Lumina-DiMOO	16.86	-	17.55	31.47	15.29	52.51
BLIP3-o	7.69	-	7.90	15.66	8.06	85.76	BLIP3-o	17.84	-	18.61	29.17	15.69	59.40
<i>Spatial Relation</i>						<i>Spatial Relation</i>							
Lumina-DiMOO	16.25	-	16.28	13.26	15.30	24.90	Lumina-DiMOO	16.86	-	17.50	30.96	17.41	32.28
<i>Text Recognition / Generation</i>						<i>Text Recognition / Generation</i>							
Lumina-DiMOO	1.14	-	1.14	4.29	1.10	21.45	Lumina-DiMOO	16.86	-	18.03	30.75	17.51	31.19

Table 11: Generation quality across unified multimodal models. Left: task-specific prompts. Right: COCO validation prompts. We report Inception Score (IS) and Fréchet Inception Distance (FID) for the baseline, $und \rightarrow gen$, and $gen \rightarrow gen$ checkpoints. Since FID is computed with images generated by baseline, baseline FID is undefined and is therefore shown as ‘-’.

We measure FID [39] and IS [71] under two prompt settings to assess how generation-side fine-tuning affects the output distribution of the baseline image generation. First, we evaluate with task-specific prompts that match the format used during generation training. We use 1.8K samples for counting, 2K for spatial relation, and 1.8K for text recognition and generation. Second, to examine the effect on more general image synthesis, we repeat the evaluation using 1K prompts from the COCO [54] validation set. The results are reported in Table 11.

We first consider FID to quantify the distribution shift, which is computed with images generated by the baseline. Across both prompt settings and all tasks, $und \rightarrow gen$ consistently yields lower FID than $gen \rightarrow gen$. This result indicates that understanding-based tuning preserves the baseline’s original output image distribution more faithfully than direct generation tuning.

By contrast, IS is considerably more task-dependent. Under the task-specific prompt setting, the IS of $und \rightarrow gen$ is generally close to that of the baseline, whereas $gen \rightarrow gen$ often attains higher IS on counting. One possible explanation is that direct generation tuning on the counting dataset, which contains multiple instances of the same object category, makes object appearance more classifier-friendly within the narrow counting prompt distribution. However, this tendency does not extend consistently to spatial relation or text recognition and generation, and under the COCO [54] validation prompt set the apparent advantage of $gen \rightarrow gen$ weakens or disappears. We therefore interpret changes in IS not as direct evidence of improved general image quality, but rather as reflecting adaptation to a narrow task-specific prompt distribution.

Overall, these results suggest that $und \rightarrow gen$ perturbs the baseline image-generation distribution substantially less than $gen \rightarrow gen$.

C.4 Ablation on how LoRA rank affects transfer strength

LoRA Rank	Counting	
	Acc.(%) $\uparrow$	MAD $\downarrow$
<i>Generation Evaluation</i>		
Baseline	48.0	1.11
8	51.9 (+3.9)	1.01 (-0.10)
32	53.0 (+5.0)	0.95 (-0.16)
128	53.6 (+5.6)	0.88 (-0.23)

Table 12: Ablation on how LoRA rank affects *und* $\rightarrow$ *gen* transfer strength. We evaluate counting *und* $\rightarrow$ *gen* transfer on Lumina-DiMOO across LoRA ranks 8, 32, and 128. Both counting accuracy and MAD improve monotonically with the rank, indicating that larger ranks yield stronger transfer.

A natural question follows from the observation that transfer emerges through training: does stronger training induce stronger transfer? To examine whether the extent of parameter adaptation correlates with the strength of transfer, we conduct an ablation over the LoRA rank, which controls how much the model parameters are allowed to change during fine-tuning. We compare counting *und* $\rightarrow$ *gen* transfer on Lumina-DiMOO across ranks 8, 32, and 128 at equal 3K training steps. As shown in Table 12, the transfer strength grows monotonically with the rank: counting accuracy rises from 48.0% at the baseline to 51.9%, 53.0%, and 53.6% at ranks 8, 32, and 128, with MAD dropping from 1.11 to 0.88.

We also conducted full fine-tuning experiments. However, when we trained the full model on the counting dataset alone, training was much more unstable than LoRA (rank 128), and the resulting generation quality collapsed entirely.

Overall, these results indicate that although stronger training generally leads to stronger transfer, properly exploiting transferability also requires choosing an appropriate tuning strength. We leave a systematic exploration of this trade-off to future work.

C.5 Can transferred capabilities be naturally combined at generation time to enable more complex generation?

Train $\rightarrow$ Test	Joint Acc.(%) $\uparrow$
<i>Generation Evaluation</i>	
Baseline	24.7
<i>und</i> $\rightarrow$ <i>gen</i>	28.5 (+3.8)

Table 13: Composing two transferred capabilities at generation time. After fine-tuning Lumina-DiMOO on simple mixture of counting and spatial-relation *understanding* data, we evaluate on 400 complex generation prompts requiring **both** counting and spatial reasoning, e.g., “three cats are on the top-right of a sofa”, which is an unseen scenario during training. Joint accuracy measures the fraction of target objects that satisfy both (1) the required object count and (2) the specified spatial relation. Trained model *naturally composes* the two transferred capabilities enabling *more complex generation behaviors*, despite never seeing such compositional training examples during training.

Our main experiments transfer a single capability at a time. A natural follow-up question is what happens when several capabilities are transferred together: do they remain isolated, each improving only its own attribute independently, or can they be *combined* at generation time? To probe this, we train a single model on a mixture of counting and spatial-relation *understanding* data, without any example that requires both capabilities jointly. We then evaluate it on 400 compositional generation prompts that demand the two capabilities simultaneously, e.g., “three cats are on the top-right of a sofa”, which is an unseen scenario during training. Following our main protocol, we report joint accuracy, defined as the fraction of target objects that satisfy *both* (1) the required object count and (2) the specified spatial relation.

As shown in Table 13, the model improves joint accuracy from 24.7% to 28.5% (+3.8%) over the baseline, even though it was never trained on compositional examples. This indicates that the capabilities transferred through understanding are not merely injected in isolation but can be composed during generation. We view this as preliminary evidence that more complex generation behaviors may be assembled from simpler component capabilities via *und* $\rightarrow$ *gen* transfer.

D Additional Qualitative Results

We provide additional qualitative examples for Lumina-DiMOO in Figures 11, 12, and 13. These figures extend the Lumina-DiMOO results shown in Figures 3, 5, and 7 in the main paper.

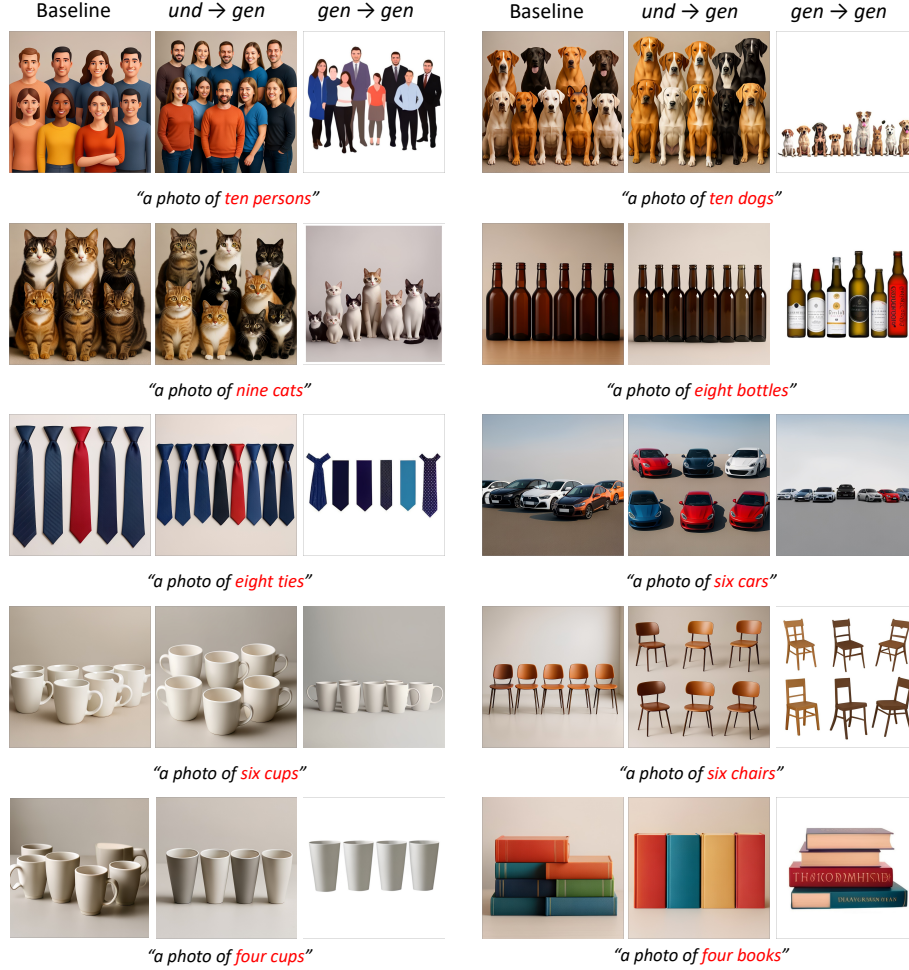


Fig. 11: Additional qualitative comparisons for Lumina-DiMOO between $und \rightarrow gen$ and $gen \rightarrow gen$ counting tasks.

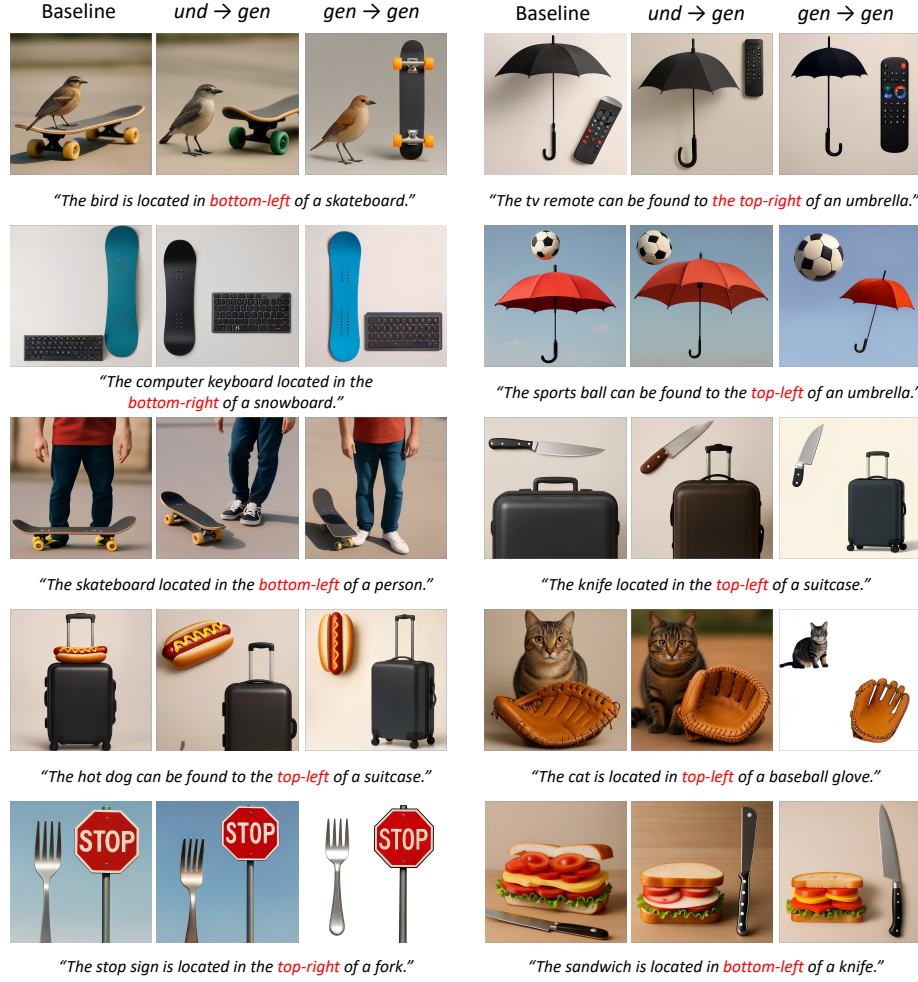


Fig. 12: Additional qualitative comparisons for Lumina-DiMOO between $und \rightarrow gen$ and $gen \rightarrow gen$ spatial relation tasks.

Baseline	<i>und</i> → <i>gen</i>	<i>gen</i> → <i>gen</i>	Baseline	<i>und</i> → <i>gen</i>	<i>gen</i> → <i>gen</i>
Acicerate	ACIERATE	gacicerate.	censu:red	censured	censured
<i>"acierate"</i>			<i>"censured"</i>		
barosinusltis /	barosinusitis	barosiinusitis	PROLETAIRATE	proletariate	.proletariate.
<i>"barosinusitis"</i>			<i>"proletariate"</i>		
TRANSHIP	Transship	transship	proannion	PROAMNION	- proamnion.
<i>"transship"</i>			<i>"proamnion"</i>		
sulpho seénium	SULPHOSELENIUM	sulphoselenium	CL,APERED	clappered	clappered
<i>"sulphoselenium"</i>			<i>"clappered"</i>		
recomputation	RECOMPUTATION	recomputation	newFoundander	newFoundlander	newFoundalader
<i>"recomputation"</i>			<i>"newfoundlander"</i>		

Fig.13: Additional qualitative comparisons for Lumina-DiMOO between *und* → *gen* and *gen* → *gen* text generation tasks.

E Limitations

In this study, we examine four representative UMMs drawn from different architectural families. While this selection allows us to explore how transferability varies with architectural design, it does not cover the full spectrum of possible UMM configurations. Evaluating a wider range of models would help assess how broadly our observations hold.

We also find that transferability cannot be attributed to architectural factors alone. Both its direction and magnitude depend on the type of visual knowledge required by each task. Although we present empirical results and qualitative insights highlighting this task dependence, determining the optimal transfer direction in advance remains challenging.

Future work could expand the analysis to more diverse model families and, importantly, investigate principled ways to anticipate transferability across tasks.