

Elucidating the Design Space of Visual Generative Models

by

Maitreya Patel

A Dissertation Presented in Partial Fulfillment
of the Requirements for the Degree
Doctor of Philosophy

Approved January 2026 by the
Graduate Supervisory Committee:

Yezhou Yang, Co-Chair
Chitta Baral, Co-Chair
Jason Baldridge
Lalitha Sankar
Tejas Gokhale

ARIZONA STATE UNIVERSITY

May 2026

ABSTRACT

Visual generative models have emerged as a central paradigm in computer vision, shifting the focus from recognizing what is observed to modeling what could be observed. By learning to synthesize images and videos, these models act not only as content generators but also as powerful visual priors that support editing, simulation, reasoning, and representation learning. Despite their impressive perceptual realism, modern generative models often exhibit limitations in semantic faithfulness, compositional reasoning, efficiency, and controllability—limitations that arise from complex and brittle design choices spanning representations, objectives, architectures, and inference procedures. This dissertation studies the design space of visual generative models from both theoretical and applied perspectives, with the goal of improving their reliability, efficiency, and semantic alignment. It introduces compositional evaluation frameworks to diagnose failures in text-to-image generation, demonstrates the critical role of semantic encoders in enabling efficient and controllable generation, and establishes visual generation as a principled tool for synthetic data creation to address compositional gaps in supervision. On the theoretical side, the dissertation develops a geometric understanding of flow matching, revealing how the velocity field can enable a few-step, controllable generation and efficient solutions to downstream tasks such as image editing, inverse problems, and model alignment. Finally, it presents a scalable autoregressive framework that decouples generation cost from image resolution, enabling fast, high-quality synthesis at arbitrary scales. Together, these contributions move visual generative modeling toward a more principled and reliable approach, improving efficiency, controllability, and robustness at scale across a broad spectrum of visual generation tasks.

ACKNOWLEDGMENTS

Getting a Ph.D. is not just about the degree, but about the journey – learning new topics, building something meaningful, and gradually becoming an independent researcher. Doing a Ph.D. was my dream long before I even knew how to code. From that early aspiration to the present day, this journey has been both challenging and deeply fulfilling. This thesis is the result of a short but beautiful journey, one that would not have been possible without the direct and indirect support of many people. I would like to take this opportunity to thank everyone who helped shape my research and me as a person along the way.

First and foremost, I would like to express my deepest gratitude to my advisor, Dr. Yezhou Yang, for his constant guidance, motivation, and unwavering support. He is one of the most inspiring mentors I have ever met, and his invaluable mentorship helped me build confidence and grow as an independent researcher. He consistently gave me the freedom and encouragement to explore diverse research directions, experiment with new ideas, and collaborate with people across disciplines. Beyond research, his guidance on career choices and maintaining a healthy work – life balance has been equally impactful.

I am deeply grateful to Dr. Chitta Baral for co-advising me throughout my Ph.D. journey. His mentorship continuously pushed me to think beyond trends and focus on directions that are truly impactful. He taught me how to question assumptions, think critically, and pursue meaningful research problems. I also appreciate the collaborative environment and high-quality resources he provided, which significantly shaped my research.

I sincerely thank my thesis committee members – Dr. Tejas Gokhale, Dr. Jason Baldridge, and Dr. Lalitha Sankar – for their guidance and support. Tejas played a crucial role in shaping my early understanding of multimodal models and helped me grow my research vision. He was often the first person I reached out to, whether for small questions or major decisions, and his guidance was invaluable throughout this journey. The second half of my Ph.D. research, particularly on unified multimodal models, was strongly inspired

by Jason. His thoughtful feedback and patience in answering even my simplest questions helped me see the bigger picture. I am grateful to Lalitha for her insightful suggestions and guidance, which helped me develop a more mathematically grounded approach to research problems.

I would also like to thank my external collaborators and mentors from Adobe (Hareesh Ravi, Ajinkya Kale, Sulin Liu, Manuel Brack, and Zhonghao Wang), Sony AI (Jingtao Li, Weiming Zhuang, and Lingjuan Lv), and Intel Labs (Kyle Min). Their mentorship helped me understand what it means to do good research at scale. I am grateful for their trust in my ideas and for giving me the opportunity to pursue them.

I extend my appreciation to Changhoon Kim and Sheng Cheng for believing in me and supporting my often crazy ideas. I am also thankful to my colleagues and lab-mates – Mihir Parmar, Mirali Purohit, Paras Seth, Sangmin Jung, Chi-Yao Huang, Abhiram Kusumba, Nilay Yilmaz, Forouzan Fallah, Arpit Vaghela, Ayush Verma, Sameep Vani, Bimsara Pathiraja, Shivam Singh, Vatsal Malaviya, Song Wen, Dimitris N. Metaxas, Swaroop Mishra, Shailaja Sampat, Agneet Chatterjee, Blake Harrison, Joshua Feinglass, Neeraj Varshney, Zeel Bhatt, Yiral Luo, Man Luo, Himanshu Gupta, Ravsehaj Singh Puri, Kirby Kuznia, Mutsumi Nakamura, Aswin ARV, Md. Nayem Uddin, Divij Handa, Venkatesh Mishra, Amir Saeidi, and Shrinidhi Kumbhar – for their collaborations, discussions, and friendships throughout this journey.

Finally, I would like to express my deepest love and gratitude to my parents, sister, and family for their unconditional support and encouragement. I cannot adequately express how grateful I am to my mother, Jyoti Patel, and my father, Jitendra Patel. Their balance of toughness and warmth kept me grounded and true to my path. They taught me how to reflect, set my own goals, make decisions, and learn from my mistakes. They recognized my passion for research long before I did – without them, I might have ended up becoming a doctor (the medical kind).

I am immensely grateful to my sister, Krunali Patel, who never doubted me and constantly reminded me to enjoy life and maintain a healthy work–life balance. She also insisted that I acknowledge the unwavering support of Goofy, her pet. On a serious note, during stressful times, Goofy taught me how to relax, have fun, and live life with constant zoomies.

I am also grateful to my extended family—my grandparents, relatives, and cousins—who have played a significant role in shaping who I am today.

TABLE OF CONTENTS

	Page
LIST OF TABLES	xii
LIST OF FIGURES	xxiii
CHAPTER	
1 INTRODUCTION	1
1.1 Overview	1
1.2 Applications of Visual Generative Models	3
1.3 The Design Space of Visual Generative Models	8
1.3.1 Sources of Complexity	8
1.3.2 Why the Design Space Matters	9
1.3.3 Engineering Perspectives	10
1.3.4 Theoretical Perspectives	10
1.3.5 Toward Principled Design	11
1.4 Main Contributions	11
2 EVALUATING CONCEPT LEARNING ABILITIES OF TEXT-TO-IMAGE DIFFUSION MODELS	15
2.1 Introduction	15
2.2 Related Work	19
2.3 CONCEPTBED	20
2.3.1 CONCEPTBED: Dataset Construction	21
2.3.2 ImageNet Dataset Generation Pipeline	22
2.3.3 Composition Categorization	23
2.3.4 Question Generations	24
2.3.5 CONCEPTBED: Dataset Statistics	25
2.3.6 D: Concept Confidence Deviation	26

CHAPTER	Page
2.3.7 Task Specific Evaluation Settings	29
2.4 Experiments & Results.....	31
2.4.1 Experimental Setup	31
2.4.2 Human Annotations	33
2.4.3 Results	33
2.5 Conclusion	37
3 A RESOURCE-EFFICIENT TEXT-TO-IMAGE PRIOR FOR IMAGE GEN- ERATIONS	47
3.1 Introduction	47
3.2 Related Works	49
3.3 Methodology	51
3.3.1 Preliminaries.....	51
3.3.2 Problem Formulation	53
3.3.3 Proposed Method: <i>E LIPSE</i>	54
3.4 Experiments & Results.....	56
3.4.1 Experimental Setup	57
3.4.2 Quantitative Evaluations.....	59
3.4.3 Qualitative Evaluations	62
3.5 Analysis.....	63
3.6 Diversity With Non-Diffusion Priors.....	71
3.7 Conclusion	73
4 MULTI-CONCEPT PERSONALIZED TEXT-TO-IMAGE DIFFUSION MOD- ELS BY LEVERAGING CLIP LATENT SPACE	80
4.1 Introduction	80

CHAPTER	Page
4.2 Related Works	84
4.3 Method	87
4.3.1 Text-to-Image Prior Mapping	87
4.3.2 Image-text Interleaved Training	89
4.3.3 Additional Concept-Specific Finetuning	91
4.4 Proposed Multibench Benchmark Dataset	95
4.5 Experiments	96
4.5.1 Main Results & Analysis	98
4.5.2 λ -ECLIPSE with Finetuning	101
4.5.3 Ablations	104
4.6 Additional Qualitative Results & Failure Cases	109
4.7 Conclusion	110
5 IMPROVING COMPOSITIONAL REASONING OF CLIP VIA SYNTHETIC VISION-LANGUAGE NEGATIVES	116
5.1 Introduction	116
5.2 Related Work	118
5.3 Method	120
5.3.1 Preliminaries	121
5.3.2 TripletData: Image-text hard negative data augmentations ..	123
5.3.3 TripletCLIP	126
5.3.4 Pseudocode of TripletCLIP	127
5.4 Experiments & Results	128
5.4.1 Experiment Setup	128
5.4.2 Compositional reasoning	131

CHAPTER	Page
5.4.3	Zero-shot evaluations..... 133
5.4.4	Finetuning performance 134
5.4.5	Ablations 135
5.5	TripletData Analysis 138
5.6	Detailed Results 146
5.6.1	Compositional reasoning 146
5.6.2	Dataset-specific zero-shot classification 147
5.6.3	Detailed image-text retrieval performance 147
5.7	Conclusion 148
6	STEERING OF RECTIFIED FLOW MODELS FOR CONTROLLED GEN- ERATIONS 152
6.1	Introduction 152
6.2	Related Works 155
6.3	Preliminaries 156
6.3.1	Problem Formulation 157
6.3.2	Cost Functions 158
6.4	Proposed Method 158
6.4.1	Error Dynamics of the ODEs 159
6.4.2	FlowChef : Steering Within the Vector Field 162
6.4.3	Empirical Findings. 166
6.4.4	FlowChef Algorithm..... 168
6.5	Extended Theoretical Analysis 169
6.5.1	Numerical Accuracy for Model Steering 169
6.5.2	Error Dynamics for Diffusion Models 173

CHAPTER	Page
6.6 Experiments	175
6.6.1 Experimental Setup	176
6.6.2 Linear Inversion Problems	178
6.6.3 Image Editing	182
6.6.4 Extended Results	183
6.6.5 Hyper-parameter Study	187
6.6.6 Qualitative Results	189
6.7 Conclusion	198
7 LEARNING CONCEPT ERASURE POLICIES VIA GFLOWNET-DRIVEN ALIGNMENT	200
7.1 Introduction	200
7.2 Related Works	204
7.3 Preliminaries	205
7.3.1 Concept Erasure for Diffusion Models	206
7.3.2 Generative Flow Networks (GFlowNets)	207
7.4 Proposed Methodology: EraseFlow	207
7.4.1 EraseFlow training algorithm	213
7.5 Experimental Results	214
7.5.1 Experimental Setup	214
7.5.2 Main Results	216
7.5.3 Ablation Studies	220
7.6 Generalization to Flow Matching	225
7.7 Experimental Details	226
7.7.1 Experimental Setup	226

CHAPTER	Page
7.7.2 Flux Training Details	227
7.7.3 Detailed Evaluation Metrics	227
7.7.4 EraseFlow + AdvUnlearn Fine-Tuning Details.	229
7.7.5 Baselines Training	229
7.7.6 More Qualitative Results	230
7.8 Conclusion	240
8 SCALING 1D IMAGE TOKENIZERS AND AUTOREGRESSIVE MODELS FOR DYNAMIC RESOLUTION GENERATIONS	241
8.1 Introduction	241
8.2 Related Works	244
8.3 Preliminaries	245
8.4 Methodology	247
8.4.1 VibeToken	248
8.4.2 VibeToken-Gen	253
8.5 Experimental Results	254
8.5.1 Experimental Setup	254
8.5.2 Image Reconstruction	258
8.5.3 Image Generation	261
8.6 Ablations	265
8.6.1 VibeToken	265
8.6.2 VibeToken-Gen	266
8.6.3 Qualitative Results	267
8.7 Conclusion	274

CHAPTER	Page
9 UNDERSTANDING AND IMPROVING FEW-STEP GENERATIONS VIA FLOW MAPS: A GEOMETRIC PERSPECTIVE	275
9.1 Introduction	275
9.2 Related Works	278
9.3 Preliminaries	280
9.4 Geometric Interpretation: Curvature and the MeanFlow JVP	281
9.4.1 Curvature of the probability flow ODE	281
9.4.2 Local linearization: JVP as a curvature proxy	282
9.5 Method: Stabilizing Flow Maps with Time-Span and Curvature Control .	285
9.5.1 Piecewise MeanFlow	286
9.5.2 Curvature-Regularized MeanFlow	287
9.6 Experimental Results	288
9.6.1 Controlled Multimodal Transport	289
9.6.2 Impact of Curvature	289
9.6.3 Image Generation	290
9.7 Qualitative Results.	293
9.8 Conclusion	296
REFERENCES	297
APPENDIX	
A RELATED PUBLICATION DETAILS	324

LIST OF TABLES

Table		Page
2.1	Examples of generated existence-related questions from captions.	26
2.2	Results of Concept Alignment Evaluation. The table shows the performance of concept learners evaluated using the $\mathcal{D}$ ($\downarrow$) metric for Concepts ($\text{Domain}_{\text{PACS}}$, and $\text{Object}_{\text{ImageNet}}$), Fine-grained _{CUB} (Object-level, and Attribute-level), and Composition. The best and worst performing models are indicated by bold and underlined numbers, respectively.	27
2.3	Compositional Reasoning Evaluation Results. The table shows the performance of the prior works for Composition Alignment. CLIP ($\uparrow$) is the traditional image-text alignment metric. VQA ($\uparrow$) is the accuracy of the ViLT VQA classifier on generated boolean questions. And $\mathcal{D}$ ($\downarrow$) is the composition deviations reported from the ViLT model with respect to its performance on original images. The best-performing model is indicated by bold numbers, while the performance that is higher than the original data is reported with underline.	28
2.4	Human Evaluations. Comparison of prior quantitative metrics and $\mathcal{D}$ metric with Human evaluations. DINO based pairwise cosine similarity is the prior evaluation metric [222]. KID was used to measure the overfitting by [128]. CLIP (CLIPScore) is the traditional reference-free image-text similarity metric. $\mathcal{D}$ is our presented concept deviation-aware evaluation metric. H.S. denotes the corresponding Human Score. Here, $\text{Domain}_{\text{PACS}}$ and $\text{Object}_{\text{ImageNet}}$ evaluations are for concept alignment and composition alignment is for image-text similarity. A high negative correlation between $\mathcal{D}$ and human ratings implies strong alignment, as lower $\mathcal{D}$ and higher human ratings correspond to better performance.	29

Table	Page
2.5 Recall. Percentage of generated images highly aligned ($\mathcal{D} \leq 0.0$) with the target concept images.	29
2.6 The table summarizes the different hyper-parameter settings for all baselines considered in this work to help the reproducibility of results.	31
2.7 This table summarizes the different hyper-parameters used for Oracles. We first take the pre-trained model weights and then fine-tune them on target concepts from the CONCEPTBED dataset. Here, NLL refers to the Negative Log-Likelihood.	32
2.8 Ablation. Effect of different oracle models to measure concept alignment using $\mathcal{D}$	36
2.9 This table summarizes the results using different existing confidence quantification metrics on CONCEPTBED dataset on 80 concepts from ImageNet. Here, MSP refers to the Maximum Softmax Probability and ECE refers to Expected Calibration Error.	38
2.10 This table summarizes the $\mathcal{D}$ performance based on the different types of classifiers across the parameters range.	38
2.11 List of all concepts from CONCEPTBED library based on their data source. .	38
2.12 This table (part 1 or 2) shows the composition statistics by categories. Here, overall means the unique compositions per concept and less than or equal to the sum of all four compositions as one composite prompt can belong up to two composition categories.	41

2.13	This table (part 2 or 2) shows the composition statistics by categories. Here, overall means the unique compositions per concept and less than or equal to the sum of all four compositions as one composite prompt can belong up to two composition categories.	42
3.1	The comparison (in terms of FID and compositions) of the baselines and state-of-the-art methods with respect to the <i>ECLIPSE</i> . * indicates the official reported ZS-FID. Ψ denotes the FID performance of a model trained on MSCOCO. <i>ECLIPSE</i> consistently outperforms the SOTA big models despite being trained on a smaller subset of dataset and parameters.	55
3.2	Prior model architecture hyperparameter details.	60
3.3	Comparison of <i>ECLIPSE</i> with respect to the various baseline prior learning strategies on four categories of composition prompts in the T2I-CompBench. All prior models are of 33 million parameters and trained on the same hyperparameters.	62
3.4	Training time comparisons of various prior models in terms of resource requirements after [1].	67
3.5	This table illustrates the scaling behavior of various T2I prior learning strategies. “Small” priors are 33 million in terms of parameters. And “Large” priors have 89 million parameters. All prior models are trained on the CC12M dataset with the Karlo diffusion image decoder.	69
4.1	Our method is the first to offer multi-concept-driven generation without depending on diffusion UNet models (except for inference). We provide the extended overview table in the appendix.	82

Table	Page
4.2 The detailed overview of subject-driven text-to-image generative methodologies. * represents the backbone base models listed are subject to potential updates or modifications.	85
4.3 Example of prompt templates used for Multibench dataset. Subjects presented in Table 4.4 are placed in {}.	95
4.4 Number of occurrences of unique subject categories. The left side of the table are subjects used for two subjects prompts, and the right side of the table are subjects used for three subjects prompts.....	96
4.5 Quantitative comparisons of different methodologies on Dreambench. The bold and underline represent the metric-specific first and second-ranked methods, respectively. * represents that we re-benchmark the performance from open-source weights.	97
4.6 Quantitative comparisons of different methodologies on ConceptBed. We present results on <i>CCD</i> ($\downarrow$) across three evaluation categories. The bold and underline represent the metric-specific first and second-ranked methods, respectively. * represents our benchmarking.	99
4.7 Quantitative comparisons of different methodologies on Multibench. The bold and underline represent the metric-specific first and second-ranked methods on each metric, respectively.	99
4.8 Quantitative results of Canny edge controlled P-T2I of different methodologies on Dreambench. The bold and underline represent the metric-specific first and second-ranked methods, respectively. Red highlighted numbers indicate the relative percentage drop for concept alignment compared to Table 4.5.	100

4.9	Ablation studies w.r.t. to the key components of λ - <i>ECLIPSE</i> design. We report the concept and composition alignment for single-subject T2I using <i>CCD</i> ($\downarrow$) on the ConceptBed benchmark.	102
4.10	Ablation study on generalization ability of λ - <i>ECLIPSE</i> with respect to different pretrained UnCLIP models. We report performance on DreamBench dataset.	107
4.11	Ablation on the effect of interleaved data for training λ - <i>ECLIPSE</i> . We report performance on DreamBench dataset.	108
4.12	Ablation study on pretraining data size. We report DreamBench performance of λ - <i>ECLIPSE</i> models trained on 100k, 500k, 1M and 2M interleaved image-text pairs. * denotes the proposed λ - <i>ECLIPSE</i> model checkpoint....	108
4.13	Ablation study on pretraining model size. We report the DreamBench performance of λ - <i>ECLIPSE</i> models trained on 5M, 34M, and 70M parameters. * denotes the proposed λ - <i>ECLIPSE</i> model checkpoint.	109
5.1	Winoground-style evaluation of pretrained CLIP models on TripletData.	125
5.2	Wordnet synset analysis of captions from CC3M and TripletData.	125
5.3	Importance of image-text hard negatives. We measure the importance of various modality-specific hard negatives on SugarCrepe, image-text retrieval, and ImageNet1k. We find that TripletCLIP results into the most optimal solution. Bold number indicates the best performance.	127
5.4	Composition evaluations of the methods on SugarCrepe benchmark. Bold number indicates the best performance, and underlined number denotes the second-best performance. $\dagger$ represents the results taken from SugarCrepe benchmark.	130

5.5	Zero-shot image-text retrieval and classification results. Bold number indicates the best performance, and underlined number denotes the second-best performance.	131
5.6	Detailed pre-training hyper-parameters for CLIP training across various experiments and ablations.	132
5.7	Ablation on filtering high-quality image-text pairs from <code>TripletData</code> . We evaluate the <code>TripletCLIP</code> after applying the filters to ensure the quality similar to <code>DataComp</code> and compare the baselines on three benchmarks. We find that <code>TripletCLIP</code> results in the most optimal solution. Bold number indicates the best performance. † represents that results are borrowed from <code>DataComp</code>	132
5.8	Finetuning-based composition evaluations of the methods on <code>SugarCrepe</code> benchmark. Bold number indicates the best performance, and underlined number denotes the second-best performance.	133
5.9	Finetuning-based evaluations of the methods on Retrieval and <code>ImageNet-1k</code> benchmarks. Bold number indicates the best performance, and underlined number denotes the second-best performance.	134
5.10	LiT style fine-tuning ablations on <code>SugarCrepe</code> , image-text retrieval, and <code>ImageNet1k</code> . Bold number indicates the best performance.	135
5.11	<code>TripletData</code> as large-scale composition evaluation dataset after [262].	137
5.12	Ablation on choice pre-trained LLM. We train <code>NegCLIP++</code> (ViT-B/32) on negative captions generated from various LLMs and report <code>SugarCrepe</code> , <code>Flickr30k</code> Retrieval (R@5), and <code>ImageNet-top5</code> performances.	137

5.13 Ablation on CyCLIP with TripletLoss. To evaluate the compatibility of TripletLoss with the CyCLIP, we train the ViT-B/32 models from scratch on CC3M with 512 batch size and report SugarCrepe, Flickr30k Retrieval (R@5), and ImageNet-top5 performances.	138
5.14 Examples of LLM-generated existence-related questions from captions to evaluate the generated images.	141
5.15 Composition evaluations of the methods on various benchmarks. Bold number indicates the best performance, and underlined number denotes the second-best performance.	145
5.16 Dataset-specific zero-shot classification results. Bold number indicates the best performance, and underlined number denotes the second-best performance.	146
5.17 Composition evaluations of the methods on various benchmarks. Bold number indicates the best performance, and underlined number denotes the second-best performance.	147
6.1 Pixel-space model-based evaluations for tackling the linear inverse problems. SSIM & LPIPS results are multiplied by 100.	176
6.2 Latent-space model based evaluations for tackling the linear inverse problems. SSIM & LPIPS results are multiplied by 100.	177
6.3 Hyperparameters for solving inverse problems using pixel-space models. ...	177
6.4 Hyperparameters for solving inverse problems using latent-space models (InstaFlow).	178

6.5	Hyperparameter examples for which various editing tasks can be performed (following Algorithm 2). Notably, the FlowChef (Flux) variant can be further optimized for task-specific settings that will follow Algorithm 1 with a careful selection of hyperparameters.	178
6.6	Compute requirement comparisons on a A6000 GPU.....	180
6.7	Evaluation metrics for different methods on background preservation and CLIP similarity on PIE-Bench.....	180
6.8	Performance of Various guided sampling methods on ImageNet64x64 with 32 batch size inference on A6000 GPU.	183
6.9	Comparison of Various Classifier Guided Style Transfer.	183
7.1	Quantitative Comparison of DB vs. TB on an NSFW Prompt.	208
7.2	Adversarial Robustness Across Tasks. Bold indicates the Best Performance, and Underline indicates the Second Best. ↓ Indicates Lower Is Better; ↑ Indicates Higher Is Better.	217
7.3	NSFW Evaluation on Various Evaluation Datasets. Bold indicates the Best Performance, and Underline indicates the Second Best Performance. ↓ Indicates Lower Is Better.....	218
7.4	Fine-grained Concept Erasure Evaluation on Concept Score and Total Score. Bold indicates the Best Performance, and Underline indicates the Second Best Performance. ↓ Indicates Lower Is Better; ↑ Indicates Higher Is Better..	219
7.5	Anchor Prompt category ablation with I2P, CLIP, and FID scores. Bold indicates the Best Performance, and Underline indicates the Second Best Performance. ↓ Indicates Lower Is Better; ↑ Indicates Higher Is Better.	221

7.6	Multiconcept erasure on artistic styles benchmark. Bold indicates the Best Performance. ↓ Indicates Lower Is Better; ↑ Indicates Higher Is Better.	224
7.7	Generalization of EraseFlow to Flux. ↓: lower is better; ↑: higher is better.	225
7.14	Anchor concepts used for each concept erasure task.	229
7.8	Examples of Gecko framework-generated Nike Shoes fine-grained erasure-related questions from prompts to evaluate the generated images. The target erasure concept in this case is <i>Nike</i> logo on the shoes.	235
7.9	Examples of Gecko framework generated Coca Cola Bottle fine grained erasure-related questions from Coca Cola bottle prompts to evaluate the generated images. The target erasure concept in this case is <i>Coca Cola</i> logo on the bottle.	236
7.10	Examples of Gecko framework generated Pegasus wings fine grained erasure-themed prompts to evaluate the generated images. The target erasure concept in this case are <i>wings</i> of the Pegasus.	237
7.11	Multiconcept erasure on celebrity benchmark. Bold Indicates the Best Performance. ↓ Indicates Lower Is Better; ↑ Indicates Higher Is Better.	238
7.12	Overview of model usage for each concept-erasure task: “ X ” denotes models we trained in-house, while “ ” denotes models adopted from the official code.	238
7.13	Official GitHub repositories leveraged for model training and usage.	239
8.1	Ablation of positional embedding and adaptive patch embedding. † indicates that FLOPs for RoPE-based models will be slightly higher due to issues in the custom functions calculations. Ablations are conducted in a VAE setting on a small encoder-decoder model for 200k iterations only. VibeToken-LL represents the final proposed tokenizer, we only perturb position embeddings.	250

8.2	Tokenizer comparison across resolutions. and $\times$ indicates whether specific capability is supported or not, N/A indicates either model not available or failed to reproduce, – denotes method does not work, and $\dagger$ indicates resolution-specific (independently) trained model. Importantly, $\ddagger$ represents the compression ratios (f). Hence, token count $((H \cdot W)/f^2)$ varies <i>w.r.t.</i> target resolutions.	252
8.3	gFID across aspect ratios and resolutions (single view). Top: low resolutions (256–512); bottom: high (512–1024). The first header line shows the aspect ratio, and the second line lists the corresponding pixel resolution. * results are borrowed from low-resolution experiments due to the days of inference computing requirements. $\dagger$ FiT-v2 has high resolution trained separately for 512×512 and plus. $\ddagger$ VibeToken-Gen utilizes the native super resolution of our tokenizer.	253
8.4	VibeToken pretraining setup.....	255
8.5	Hyperparameters for <i>GPT-B</i> and <i>GPT-XXL</i> configurations of VibeToken-Gen.	256
8.6	Ablation study on ImageNet: Impact of token length.	258
8.7	Ablation study on ImageNet: Impact of resolution-agnostic training.....	258
8.8	Comparison with LlamaGen (fixed resolution) on ImageNet.	259
8.9	Comparison across models and resolutions. Columns group 256×256 and 512×512 results (Tokens/FID/IS). Among many prior works on discrete image modeling, only VibeToken-Gen scales to arbitrary resolutions effortlessly.	260
8.10	Ablation study on ImageNet: Impact of token length.	262
8.11	Ablation study on ImageNet: Impact of resolution-agnostic training.....	262

8.12	Comparison with LlamaGen (fixed resolution) on ImageNet.	263
8.13	Super-resolution ablation. Ground truth high resolution is fixed to 1024×1024 and bicubic downsampling is used to get low resolution counterparts.	264
9.1	CIF R10 Piecewise MeanFlow. We compare FID vs. NFE performance of MeanFlow variants. † represents our reproduced results using official released codebase. ‡ represents that each model is trained on NFE specific segments independently.	291
9.2	ImageNet1k 256×256 Piecewise MeanFlow. We compare FID vs. NFE performance of MeanFlow variants. † represents our reproduced results. ‡ represents that each model is trained on NFE specific segments independently.	291
9.3	DiT-XL/2 finetuning-based comparison under different NFE settings. Lower FID is better (↓) and higher IS is better (↑).	292

LIST OF FIGURES

Figure		Page
1.1	Examples of various applications of the visual generative models. All examples are generated using the work done for this dissertation.	4
1.2	Visual Generative Models resemble an iceberg; while scaling parameters may seem to be the primary driver of SoTA performance, significant complexity resides within the hidden design space. A granular understanding of these submerged components is essential for improving model robustness and efficiency in downstream tasks.	7
2.1	Visual concept learners such as textual inversion models learn to “invert” a set of images (about a concept c) into a text embedding V^* , and then use this learned textual concept in new text prompts to generate images of concept c under different contexts and by performing novel compositions with other concepts. The proposed CONCEPTBED dataset along with the evaluation metric $\mathcal{D}$ allows us to comprehensively and quantifiably evaluate concept learning abilities of text-to-image diffusion models.	16
2.2	A summary of the CONCEPTBED dataset for large-scale grounded evaluations of concept learners. The collection of concepts is categorized into three classes: (1) Domain, (2) Objects, and (3) Attributes. CONCEPTBED has 284 unique concepts and four compositional categories. Here, V^* is a learned concept.	18
2.3	Intuitive illustration of the Concept Confidence Deviation ($\mathcal{D}$) for the concept <i>rt Painting</i> . Blue and Orange are the probability distributions of the real and generated concept images.	22
2.4	Left: word-cloud of the CONCEPTBED compositions. Right: statistics of composition categories and their combinations in composite text prompts. ...	24

Figure	Page
2.5 Screenshot depicting few-shot classification using GPT3. We keep the instruction and few-shot examples constant and change the target phrase to get the corresponding categories.....	25
2.6 Qualitative examples showcasing the effectiveness of concept learners on the CONCEPTBED dataset. The leftmost column displays four instances of ground truth target concept images (V^*). Subsequent columns exhibit target concept-specific images generated by all baseline methods.	40
2.7 An example of human annotation for determining the concept alignment for object-specific concepts.	43
2.8 An example of human annotation for determining the concept alignment for style-specific concepts.	43
2.9 The example of Human Annotation for determining the image-text alignment.	43
2.10 Qualitative examples of the style-specific four concepts.	44
2.11 Qualitative examples of the object-specific four concepts.	45
2.12 Qualitative examples from Custom Diffusions at different random seeds. The leftmost four figures are the target concept images. Top-Right four images are object-specific generated images. While Bottom-Right four generated images are on different composite text prompts.	46
3.1 Comparison between SOTA text-to-image models with respect to their total number of parameters and the average performance on the three composition tasks (color, shape, and texture). <i>ECLIPSE</i> achieves better results with less number of parameters without requiring a large amount of training data. The shown <i>ECLIPSE</i> trains a T2I prior model (having only 33M parameters) using only 5M image-text pairs with Kandinsky decoder.	48

3.2	Standard T2I prior learning strategies (top) minimizes the mean squared error between the predicted vision embedding $\hat{z}_x$ w.r.t. the ground truth embedding z_x with or without time-conditioning. This methodology cannot be generalized very well to the outside training distribution (such as Orange Square). The proposed <i>ECLIPSE</i> training methodology (bottom) utilizes the semantic alignment property between z_x and z_y with the use of contrastive learning, which improves the text-to-image prior generalization.	52
3.3	Qualitative result of our text-to-image prior, <i>ECLIPSE</i> , comparing with SOTA T2I model. Our prior model reduces the model parameter requirements (from 1 Billion → 33 Million) and data requirements (from 177 Million → 5 Million → 0.6 Million). Given this restrictive setting, <i>ECLIPSE</i> performs close to its huge counterpart (i.e., Kandinsky v2.2) and even outperforms models trained on huge datasets (i.e., Würstchen, SDv1.4, and SDv2.1) in terms of compositions.	58
3.4	Qualitative evaluations by human preferences approximated by PickScore [120]. The top two figures compare <i>ECLIPSE</i> to Projection and Diffusion Baselines trained with the same amount of data and model size for both Karlo and Kandinsky decoders. In the bottom figure, we compare <i>ECLIPSE</i> with the Kandinsky v2.2 decoder trained on the LAION-HighRes dataset against SOTA models.	61
3.5	Empirical analysis of the PickScore preferences of diffusion priors with respect to the various hyper-parameters.	64

3.6	The top figure shows the qualitative examples of the biases learned by the T2I prior models. Bottom figures show the PickScore preferences of the <i>ECLIPSE</i> models trained on various datasets with respect to the other datasets (left) and Karlo (right).	66
3.7	Inference time analysis of diffusion priors having 1B and 33M parameters vs. <i>ECLIPSE</i> prior.	68
3.8	Hyperparameter (λ) ablation. This figure illustrates the PickScore preferences across the <i>ECLIPSE</i> priors trained with different values of λ w.r.t. the Projection baseline (with $\lambda = 0.0$).	68
3.9	Qualitative example for <i>ECLIPSE</i> priors (with Karlo decoder) trained with different values of hyperparameter (λ).	69
3.10	Human evaluations of the <i>ECLIPSE</i> vs. Kandinsky v2.2 generated images. It can be observed that both models are rated equally in terms of image quality and caption alignment.	70
3.11	This figure illustrates the human preferences between <i>ECLIPSE</i> prior for Kandinsky model (trained on LAION-HighRes subset) vs. Original Kandinsky v2.2 model.	71
3.12	Qualitative examples comparing (in terms of aesthetics) <i>ECLIPSE</i> with Kandinsky v2.2.	72
3.13	This figure illustrates the human preferences on the diversity of generated images between <i>ECLIPSE</i> prior with Kandinsky v2.2 diffusion image decoder vs. Kandinsky v2.2.	73

3.14	Qualitative comparisons between <i>ECLIPSE</i> and baseline priors (having 33 million parameters) trained on CC12M dataset with Karlo decoder. * prompt is: "The bold, striking contrast of the black and white photograph captured the sense of the moment, a timeless treasure memory."	74
3.15	Qualitative comparisons of <i>ECLIPSE</i> priors with Karlo decoder trained on different datasets. * prompt is: "The vibrant, swirling colors of the tie-dye shirt burst with energy and personality, a unique expression of individuality and creativity."	75
3.16	Qualitative result of our text-to-image prior, <i>ECLIPSE</i> (with Karlo decoder), along with a comparison with SOTA T2I models. Our prior model reduces the prior parameter requirements (from 1 Billion 33 Million) and data requirements (from 115 Million 12 Million 0.6 Million).	76
3.17	Qualitative comparisons between <i>ECLIPSE</i> and baseline priors (having 34 million parameters) trained on LAION-HighRes subset dataset with Kandinsky v2.2 diffusion image decoder.	77
3.18	Qualitative comparisons between <i>ECLIPSE</i> prior trained on MSCOCO and LAION datasets with Kandinsky v2.2 decoder.	77
3.19	More qualitative result of our text-to-image prior, <i>ECLIPSE</i> (with Kandinsky v2.2 decoder), along with a comparison with SOTA T2I models. Our prior model reduces the prior parameter requirements (from 1 Billion 33 Million) and data requirements (from 177 Million 5 Million 0.6 Million).	78
3.20	Instances where <i>ECLIPSE</i> encounters the challenges in following the target text prompts.	78

Figure	Page
3.21 An example of human annotation for determining the image quality and caption alignment.....	79
3.22 An example of human annotation for determining the most aesthetic image..	79
4.1 λ - <i>ECLIPSE</i> can estimate subject-specific latent image embeddings while maintaining the balance between concept and composition alignment in the CLIP latent space itself.	81
4.2 Three stages of the λ - <i>ECLIPSE</i> pipeline. 1) Create the image-text interleaved features using frozen CLIP. 2) Pre-train the λ - <i>ECLIPSE</i> (34M parameters) using Eq. 4.1, which ensures the mapping to the desired latent space given the image-text interleaved data. 3) During inference, the frozen Kandinsky v2.2 diffusion UNet model takes the output from the λ - <i>ECLIPSE</i> and generates the image.	86
4.3 CLIP(vision) features capture the semantics and fine-grained visual details. Each input is given as input to the Kandinsky v2.2 and re-generated from the decoder. (Top: Real-images, Bottom: Canny edge)	87
4.4 This figure illustrates a qualitative comparison of λ - <i>ECLIPSE</i> with contemporary approaches for single-subject T2I generations, utilizing concepts and prompts from the Dreambench dataset. For each method, concept, and prompt, we generate four images and select the one that most accurately represents the queried concept and composition.	90
4.5 Examples of image-text interleaved training data. The left column shows the input of the prior model and the right images shows the ground truth corresponding images. Note: these examples are generated from λ - <i>ECLIPSE</i> for better interpretability.	93

Figure	Page
4.6 Qualitative results categorized by generative capabilities.	94
4.7 Qualitative comparison between λ - <i>ECLIPSE</i> and other baselines.	101
4.8 Interpolation between four concepts. Here, we estimate the image embedding using λ - <i>ECLIPSE</i> corresponding to each corner and then interpolate from top to bottom and left to right. At last, we use the Kandinsky v2.2 diffusion UNet model to generate the images with fixed random seeds from these sets of image embeddings.	102
4.9 DreamBooth (Stable Diffusion v1.5) vs. λ - <i>ECLIPSE</i> (with fine-tuning) <i>w.r.t.</i> DINO and CLIP-T metrics on Dreambench.	104
4.10 Qualitative examples of λ - <i>ECLIPSE</i> without finetuning and different stages of finetuning.	105
4.11 Qualitative examples λ - <i>ECLIPSE</i> model trained with and without the interleaved pretraining strategy.	111
4.12 Qualitative examples λ - <i>ECLIPSE</i> model trained with varying pretraining data.	111
4.13 Qualitative examples λ - <i>ECLIPSE</i> model trained with varying model parameters.	112
4.14 Qualitative examples of the increasing complexity of novel visual concepts as we move from top to bottom.	113
4.15 Qualitative examples of showcasing the consistency comparisons between Kosmos-G and λ - <i>ECLIPSE</i>	113
4.16 Qualitative examples of showcasing the failure cases on human faces on Kosmos-G, IP-Adapter (SDXL), IP-Adapter (FaceID), and λ - <i>ECLIPSE</i> . ..	114

4.17	Qualitative examples of showcasing the failure cases on Multibench of Kosmos-G, Emu2, and λ -ECLIPSE.	114
5.1	Comparison of training workflows of CLIP, NegCLIP, and TripletCLIP. (x, y) represents the positive a image-text pair, and (x', y') represents the corresponding negative image-text pair.	121
5.2	Examples image-text pairs from TripletData. In each block, a positive pair from CC3M is on the left and corresponding negatives from TripletData are shown on the right.....	122
5.3	Average Results of LaCLIP and TripletCLIP for SugarCreme Compositions, Image-Text Retrieval, and ImageNet1k over increasing concept diversity.....	136
5.4	Positive vs. Negative modality-specific pair-based similarity distribution of pre-trained CLIP ViT-B/32 model <i>w.r.t.</i> the vision and text-only encoders. The left plot is the vision embedding similarities between positive and negative images. The right plot is the text embedding similarities between positive and negative captions. In the ideal scenario, the distribution should be skewed towards 0.0, which indicates that the model can correctly distinguish between the positive and negative data.....	139
5.5	Positive vs. Negative modality-specific pair-based similarity distribution of baseline LaCLIP and TripletCLIP. The left plot is the vision embedding similarities between positive and negative images. The right plot is the text embedding similarities between positive and negative captions.....	140

5.6	Qualitative examples of positive and hard negative image-text pairs from TripletData. In each block, left image-text pairs are positive images from CC3M, and right pairs are corresponding negatives from TripletData.	150
5.7	Examples of T2I failures.	151
6.1	FlowChef steers the trajectory of Rectified Flow Models during inference to tackle linear inverse problems, image editing, and classifier guidance. We extend FlowChef to SOTA models like Flux and InstaFlow, enabling gradient- and inversion-free control for efficient, controlled image generation.	152
6.2	Motivation behind FlowChef based on rectified flow models' trajectory space. Let $p_1 \sim N(0, I)$ and p_0 be noise and posterior distributions. (a) Stochasticity and nonlinear trajectories can complicate gradient estimation at each denoising step t . (b) Our method FlowChef enables efficient trajectory steering to guide x_t along the trajectory towards x_0^{ref} . (c) As $t \rightarrow 0$, the cosine similarity of RFMs gradients, with or without ODESolver backpropagation, increases, in contrast to the erratic gradients of DMs. (d) FlowChef improves the latency by 7x and performance by 2x compared to the prior diffusion SOTA baselines on inverse problems using pixel models.	154
6.3	Illustration of impact of number guided control steps on baseline MPGD (DMs) and proposed FlowChef (RFMs) with equal guidance values: $\nabla_{\hat{x}_0} \mathcal{L} = (\hat{x}_0 - x_0^{ref})/N$. FlowChef deterministically steers RFMs, achieving high-quality results with only five guided steps. In contrast, MPGD, applied to stochastic DMs, exhibits reduced performance due to inherent randomness.	154

6.4	Empirical analysis of gradient similarity (a, b, and c) and convergence rate. (a) and (b) analyzes the gradients without model steering. (c) contains the gradient similarity during the active model steering. And (d) shows the trajectory similarity at each timestep t <i>w.r.t.</i> the inversion based trajectory. . .	166
6.5	Qualitative results on image editing. As illustrated, our method attains the SOTA performance on comparison inversion-free methods. Meanwhile, the FlowChef (Flux) variant achieves better quality and edits. * denotes the concurrent works.	179
6.6	Qualitative results on linear inverse problems. All baselines are implemented on stable diffusion v1.5, except FlowChef Flux variant. Results are reported for VRAM and time on an A100 GPU at 512 x 512 resolution, with Flux experiments at 1024 x 1024.	181
6.7	Extending FlowChef to 3D multiview synthesis.	184
6.8	FlowChef (Flux) multi object editing examples.	185
6.9	Human preference analysis for image editing.	186
6.10	Hyperparameter sensitivity analysis.	187
6.11	Effect of FlowChef learning rate with fixed 20 max steps and one optimization step on InstaFlow.	188
6.12	Effect of FlowChef optimization steps with fixed 20 max steps and 0.02 learning rate on InstaFlow.	188
6.13	Effect of various FlowChef’s steering parameters with increasing maximum optimization steps on InstaFlow.	189
6.14	FlowChef (Flux) model failures on inverse problems and image editing.	189

Figure	Page
6.15 Qualitative results on image editing. Additional qualitative comparisons of FlowChef with the baselines.	191
6.16 Qualitative examples of various methods for easy box inpainting task on RF++.	192
6.17 Qualitative examples of various methods for hard box inpainting task on RF++.	193
6.18 Qualitative examples of various methods for an easy deblurring task on RF++.	194
6.19 Qualitative examples of various methods for the hard deblurring task on RF++.	195
6.20 Qualitative examples of various methods for an easy super-resolution task on RF++.	196
6.21 Qualitative examples of various methods for the hard super-resolution task on RF++.	197
7.1 EraseFlow (ours) achieves effective concept erasure when compared with various concept erasure methods across diverse tasks—removing NSFW content (top), artistic styles like “Van Gogh” (middle), and fine-grained elements such as the “Nike” logo from shoes (bottom)—while preserving image quality and fidelity.	201
7.2 Comparison of EraseFlow variants and baselines on erasing the nudity concept in Stable Diffusion v1.4. Robustness is measured by adversarial attack success rate () and utility by FID (). Lower is better on both axes. Circle size reflects training cost.	203
7.3 Qualitative Comparison of DB vs. TB on an NSFW Prompt.	209

7.4	Illustration of probability redistribution during <code>EraseFlow</code> optimization. (a) In the pretrained model, probability mass is concentrated around target regions, increasing the likelihood of generating target concepts. (b) During training, the trajectory-balance objective redistributes probability mass from target regions toward anchor regions. (c) After optimization, the learned distribution aligns with anchor regions, effectively suppressing target concepts while maintaining visual and semantic fidelity.	211
7.5	Image generations by SDv1.4 and concept-erasure methods on different prompts. (top) A nudity prompt attacked by UDAtk; <code>EraseFlow</code> effectively suppresses NSFW content while most methods fail. (middle) A Van Gogh-style prompt attacked by UDAtk; <code>EraseFlow</code> removes the artistic style successfully. (bottom) A prompt with fine-grained elements; <code>EraseFlow</code> removes “wings” from “Pegasus” while preserving all other details like the horse and the mountain range.	216
7.6	Ablation results. (Left) Effect of $\log \beta$ on erasure and fidelity scores. (Right) Effect of early stopping (<code>STOP_SAMPLING</code>) on performance stability.	220
7.7	Ablation on $\log \beta$ and $\log z_\phi$. Left: I2P performance across two anchor prompts—“ <i>Fully dressed</i> ” and “ <i>a person wearing complete clothing</i> ”—showing a sharp drop once $\log \beta \approx 1.6$. Right: Varying $\log z_\phi$ with fixed $\log \beta=1.6$ demonstrates that larger $\log z_\phi$ values accelerate convergence and improve erasure.	222
7.8	Timestep selection ablation. Effect of different sampling strategies on erasure and fidelity. <code>EraseFlow</code> performs best with a balanced mix of timesteps.	223

Figure	Page
7.9 Nudity erasure qualitative results on Flux showing generalization of EraseFlow beyond diffusion-based models. Compared to baselines, EraseFlow more effectively suppresses target concepts while preserving overall fidelity and text-image alignment.	226
7.10 More qualitative examples of UDAtk on NSFW erasure	231
7.11 More qualitative examples of UDAtk on Van Gogh erasure	231
7.12 Qualitative examples of UDAtk on Caravaggio style erasure. To hide inappropriate content, few images are blurred for publication purposes.	232
7.13 Qualitative examples of erasing "Nike" logo from shoes.	232
7.14 Qualitative examples of erasing "Coca Cola" brand from glass bottle.	233
7.15 Qualitative examples of erasing "wings" from Pegasus.	233
7.16 Examples of EraseFlow Concept Erasure Failures.	234
8.1 VibeToken-Gen image generation examples. A single resolution-generalist visual autoregressive model trained on ImageNet1k that can synthesize arbitrary resolution and aspect ratio images. The examples shown are from various resolutions between 256×256 and 1024×1024. VibeToken-Gen leverages our novel resolution-agnostic 1D visual tokenizer, VibeToken, that can efficiently encode any resolution into as few as 64 tokens, making our AR model 63× efficient than the baseline (LlamaGen) on a 1024×1024 resolution.	242
8.2 Compute comparisons. Across resolution and metrics, VibeToken achieves strong efficiency (fewer tokens/FLOPs) compared to 2D baselines. This, in turn, leads to VibeToken-Gen being significantly more efficient with a constant 179 GFLOPs.	243

8.3	Overview of VibeToken tokenizer. VibeToken introduces four key components: 1) Dynamic grid position embedding, 2) Dynamic patch embedding, 3) Adaptive decoder resolution, and 4) Variable latent token-based resolution-agnostic encoding. These components provide full flexibility to 1D tokenizers, supporting arbitrary resolutions and compute controls.	248
8.4	Dynamic token length. rFID vs. latent length L at 256^2 . VibeToken-Gen is the only model that pairs competitive quality with resolution generalization.	261
8.5	rFID performance of VibeToken-LL with different token lengths across the resolutions. Notably, reference images are different for each resolution, hence, rFID scale and performance are not comparable between each resolution.	265
8.6	gFID performance of VibeToken-Gen and baseline LlamaGen models as training progresses. The results are shown on 256×256 without classifier-free guidance.	266
8.7	gFID performance of VibeToken-Gen (B/XXL) with respect to classifier-free guidance scale on 256×256	267
8.8	Class 985. Qualitative examples of generated images using VibeToken-Gen (XXL/64) on arbitrary resolutions.	268
8.9	Class 980. Qualitative examples of generated images using VibeToken-Gen (XXL/64) on arbitrary resolutions.	269
8.10	Class 387. Qualitative examples of generated images using VibeToken-Gen (XXL/64) on arbitrary resolutions.	270
8.11	Class 250. Qualitative examples of generated images using VibeToken-Gen (XXL/64) on arbitrary resolutions.	271

Figure	Page
8.12 Class 88. Qualitative examples of generated images using VibeToken-Gen (XXL/64) on arbitrary resolutions.	272
8.13 Class 33. Qualitative examples of generated images using VibeToken-Gen (XXL/64) on arbitrary resolutions.	273
8.14 Qualitative examples of native 4× super-resolution ability of VibeToken with respect to simple bilinear resize operation.	273
9.1 Geometric Perspective Intuition. This figure first illustrates MeanFlow identity (Eq. 9.4) over the oracle flow trajectory, where the time derivative is represented as JVP correction (a). Overall velocity field of flow matching models have high time derivative, which is a problem. There are two ways to reduce the impact of time derivative: (b) reducing the time delta and (c) reducing the curvature. By doing either of them instantaneous velocity (v_t) becomes very close to the average velocity (u_θ) and time derivative becomes smaller.	276
9.2 Multimodal toy distribution. Qualitative results on 2D-dataset for few-step generations. MeanFlow performs very poorly and slightly improves as inference steps increases. Stable MeanFlow significantly improves over the baseline and keeps improving.	284
9.3 Impact of curvature in Flow Maps learning. We illustrate the impact of the curvature regularizer on distilling Flow Maps. Specifically, we perform up to 40k optimization steps for the curvature regularizer followed by 10k steps of Flow Map distillation on top of a pretrained REPA model.	285

Figure	Page
9.4 Qualitative Examples at NFE=1. Qualitative comparison between randomly selected images from MeanFlow and Stable MeanFlow. Stable MeanFlow significantly improves over the baseline.....	293
9.5 Qualitative Examples at NFE=4. Qualitative comparison between randomly selected images from MeanFlow and Stable MeanFlow. Stable MeanFlow significantly improves over the baseline.....	294
9.6 Qualitative figure comparing MeanFlow and Stable MeanFlow together across the NFE.	294
9.7 Qualitative figure comparing MeanFlow and Stable MeanFlow together across the NFE.	295
9.8 Qualitative figure comparing MeanFlow and Stable MeanFlow together across the NFE.	295

Chapter 1

INTRODUCTION

1.1 Overview

Perception has long stood at the center of machine learning research. Much of the field’s progress has been driven by the goal of interpreting visual data—recognizing objects, assigning semantic labels, and extracting structure from images and videos. These advances have enabled machines to answer increasingly complex questions about visual inputs. Yet interpretation alone provides only a partial account of visual intelligence. A system that genuinely understands the visual world should not only be able to analyze what it observes, but also to construct visual instances that are consistent with that understanding.

Visual generative models address this complementary objective. Rather than mapping visual inputs to predefined outputs, they aim to model the process by which images and videos are formed. In doing so, generative models attempt to capture the regularities, constraints, and degrees of freedom that define visual data. This shift—from recognition to synthesis—changes the nature of the learning problem. Generation requires a model to internalize how objects, scenes, and events could plausibly appear, not merely how they can be categorized.

The motivation for visual generation is therefore not limited to producing realistic imagery. To generate a coherent image, a model must reason—implicitly or explicitly—about geometry, appearance, composition, and context. For video generation, these requirements extend further to motion, temporal consistency, and causal structure. When generative models fail, they often reveal precisely which aspects of visual structure have not been captured, making generation a diagnostic tool for visual understanding rather than a peripheral

application.

Historically, visual generative models have evolved alongside broader trends in computer vision. Early approaches emphasized probabilistic formulations and density estimation, treating images as high-dimensional signals drawn from unknown distributions. Later methods introduced latent representations that enabled limited control over visual attributes, reflecting early attempts at disentangling factors of variation. Recent large-scale generative models have dramatically expanded the scope of what can be synthesized, producing images and videos with high perceptual fidelity and responding flexibly to diverse forms of conditioning, including natural language.

As these models have matured, their role has shifted. Visual generative models are no longer viewed solely as tools for sampling or reconstruction, but increasingly as general-purpose visual priors. They are used to edit images, complete missing content, simulate alternative scenarios, and support multimodal reasoning. In many settings, generation is not the final objective; rather, it enables interaction with visual representations in a way that is flexible, compositional, and open-ended. This evolution reflects an implicit goal of generative modeling: to learn representations that can be manipulated and queried, not just observed.

At the same time, the capabilities of modern generative models raise important questions about what they truly capture. High-quality synthesis does not guarantee semantic correctness, structural consistency, or faithful adherence to conditioning signals. Models may generate visually plausible outputs that violate spatial relationships, physical constraints, or temporal logic. Such behaviors highlight the tension between perceptual realism and structural understanding, and underscore the need to reason carefully about the objectives and inductive biases embedded in generative systems.

These challenges are further compounded by the diversity of design choices that underlie visual generative models. Differences in architectures, representations, training objectives,

conditioning mechanisms, and evaluation protocols can lead to qualitatively different behaviors, even when models appear similar in output quality. As generative modeling extends from static images to videos, and from unconditional sampling to increasingly structured forms of generation, understanding these design dimensions becomes central to interpreting model behavior.

Taken together, visual generative models occupy a distinctive position within computer vision. They blur the boundary between perception and synthesis, serving both as powerful tools for creating visual content and as probes into the nature of visual representation itself. By forcing models to construct, rather than merely classify, generative modeling reframes the question of visual understanding—shifting it from recognizing what is seen to modeling what could be seen. This perspective provides a foundation for thinking about the future of visual intelligence, where generation, interaction, and understanding are tightly intertwined.

1.2 Applications of Visual Generative Models

Visual generative models support a wide range of applications that extend beyond the act of image or video synthesis itself. These applications reflect different ways in which generative capabilities interact with visual representations, from creation and manipulation to reasoning and simulation. While the boundaries between these categories are often fluid, they provide a useful lens for understanding how generative models are used in practice and what demands these uses place on model design.

Content Synthesis and Creation. One of the most direct applications of visual generative models is the creation of images and videos. In this setting, generation serves as the primary objective: models are tasked with producing novel visual content that is coherent, diverse, and perceptually plausible. This includes unconditional generation, where samples are drawn without explicit constraints, as well as conditional generation based on class labels,



Dynamic Resolution Image Synthesis



Multi-Subject Driven T2I



Super Resolution

Deblurring



Editing

Inpainting

Classifier Guidance

Figure 1.1: Examples of various applications of the visual generative models. All examples are generated using the work done for this dissertation.

text descriptions, sketches, or other inputs. Content synthesis places strong emphasis on visual fidelity and diversity. Models must strike a balance between generating high-quality samples and adequately covering the underlying data distribution. For video generation, additional challenges arise in maintaining temporal coherence and consistent identity across frames. These applications often act as stress tests for a model’s capacity to represent the full complexity of visual data.

Image and Video Editing. Beyond creating content from scratch, generative models are widely used for editing existing images and videos. Editing tasks include attribute manipulation, object insertion or removal, style transfer, inpainting, and temporal editing in videos. In these scenarios, generation is constrained: the model must modify certain aspects of the input while preserving others. Editing applications highlight the importance of controllability and locality in generative models. Successful edits require representations that separate semantic factors and allow targeted changes without unintended side effects. For video, edits must also remain consistent over time, avoiding flicker or drift. These requirements expose limitations in models that rely primarily on global generation without explicit structural constraints.

Data Augmentation and Simulation. Visual generative models are frequently used to augment datasets or simulate visual environments. In supervised learning settings, synthetic images and videos can be generated to increase data diversity, address class imbalance, or expose models to rare or extreme conditions. In simulation, generative models can create environments or scenarios that are difficult, expensive, or unsafe to capture in the real world. In these applications, realism alone is insufficient. Generated data must be semantically meaningful and aligned with the intended use case. For example, augmentations that alter superficial appearance without respecting semantic structure may fail to improve downstream

performance. As a result, simulation-oriented applications emphasize semantic fidelity, coverage of relevant variations, and alignment with task-specific constraints.

Representation Learning and Priors. Generative models also play an important role as tools for learning visual representations. By training models to reconstruct or generate visual data, latent representations are encouraged to capture salient factors of variation. These representations can then be reused for downstream tasks such as classification, retrieval, or reasoning. In this context, generation serves as a means rather than an end. The quality of the learned representation is often more important than the visual quality of the generated outputs. Applications that rely on generative priors benefit from models that encode structure in a compact and interpretable manner, enabling generalization beyond the training distribution.

Multimodal and Interactive Systems. With the integration of language, vision, and other modalities, visual generative models increasingly operate within multimodal systems. Text-to-image and text-to-video generation enable users to interact with visual models through natural language, while iterative prompting and refinement support interactive content creation. These applications place new demands on generative models. Outputs must be faithful to conditioning inputs, robust to linguistic variation, and capable of compositional reasoning. Failures in spatial relationships, object counts, or temporal logic are particularly salient in multimodal settings, as they reveal misalignment between visual and symbolic representations.

Reasoning, Planning, and World Modeling. A growing class of applications treats visual generative models as components of reasoning or planning systems. In such settings, generation is used to imagine future states, explore counterfactuals, or simulate the outcomes of actions. For video generation, this perspective aligns closely with the idea of modeling

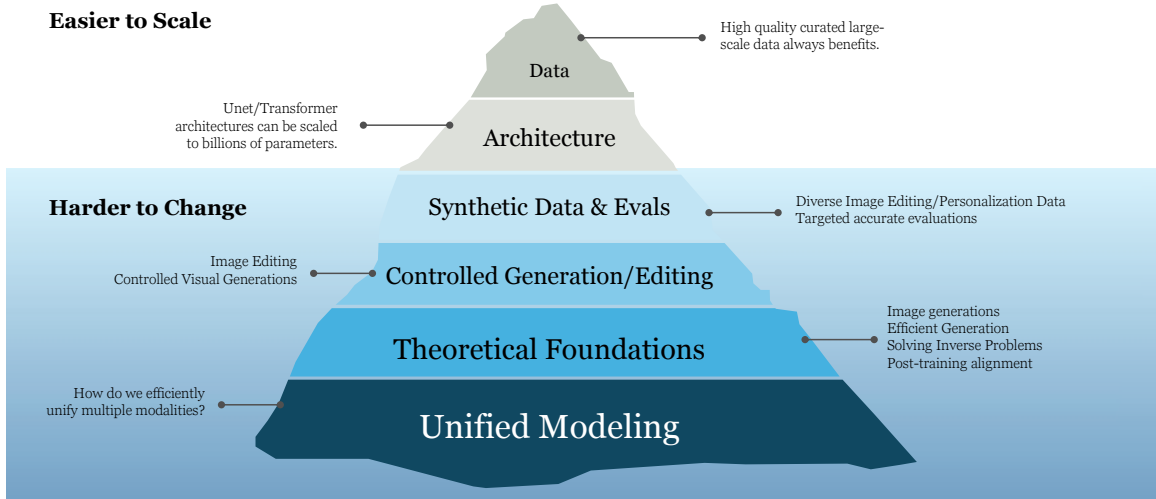


Figure 1.2: Visual Generative Models resemble an iceberg; while scaling parameters may seem to be the primary driver of SoTA performance, significant complexity resides within the hidden design space. A granular understanding of these submerged components is essential for improving model robustness and efficiency in downstream tasks.

dynamics and causality in visual environments. These applications shift the emphasis from perceptual realism to consistency, predictability, and structural correctness. Generative models must support coherent rollouts and respect constraints imposed by the environment. This use case underscores the connection between generative modeling and broader notions of world models, where visual generation serves as a substrate for reasoning rather than content creation alone.

Evaluation and Diagnostics. Finally, visual generative models are increasingly used as diagnostic tools. By analyzing what a model can and cannot generate, researchers gain insight into its learned representations and biases. Generation-based evaluation can reveal weaknesses that may not appear in discriminative benchmarks, such as failures in compositionality or spatial reasoning. In this role, generative models are not judged solely by output quality, but by what their outputs reveal about underlying structure. This perspective reinforces the view of generation as a probe for understanding visual intelligence.

1.3 The Design Space of Visual Generative Models

Visual generative models are defined not by a single architectural paradigm, but by a collection of interdependent design choices. These choices span representations, network architectures, training objectives, conditioning mechanisms, inference procedures, and evaluation criteria. As models have grown in scale and capability, this design space has expanded rapidly, resulting in systems that are increasingly powerful, but also increasingly complex and opaque.

This complexity is not incidental. Visual generation is a fundamentally challenging problem: images and videos are high-dimensional, structured, and multi-scale, with dependencies that span space, time, and semantics. Capturing these properties requires models to balance expressivity with tractability. Design decisions that improve performance in one dimension—such as perceptual fidelity—may introduce trade-offs in others, such as efficiency, controllability, or reliability. Understanding these trade-offs is a central motivation for studying the design space of generative models.

1.3.1 Sources of Complexity

The complexity of visual generative models arises from several sources. First, representations differ in how they encode visual information: some emphasize pixel-level fidelity, while others operate in latent spaces that abstract away low-level detail. Second, architectures vary in their inductive biases, with different mechanisms for modeling spatial structure, long-range dependencies, or temporal dynamics. Third, training objectives introduce additional layers of complexity, combining reconstruction, likelihood-based, adversarial, or consistency-driven losses, often with implicit assumptions about what constitutes a “good” generation.

Inference procedures further complicate the picture. Some models rely on iterative

sampling or refinement processes, trading computation for quality, while others favor direct generation for speed. Conditioning mechanisms—such as text, images, or control signals—add another axis of variation, influencing how faithfully models respond to external constraints. Taken together, these choices define a rich and multidimensional design space, in which no single configuration is universally optimal.

1.3.2 Why the Design Space Matters

Careful reasoning about the design space is essential for both practical and conceptual reasons. From an engineering perspective, design choices directly affect efficiency. Large generative models often incur substantial computational and memory costs during training and inference. Understanding which components are essential, and which can be simplified or replaced, enables more efficient models without sacrificing core capabilities. This is particularly important as generative models are deployed in resource-constrained or real-time settings.

Performance is another key consideration. Improvements in generation quality are often achieved through architectural refinements or objective redesigns, but such gains may not generalize across tasks or modalities. A design-space view helps distinguish improvements that reflect fundamental modeling advances from those that exploit narrow evaluation criteria. It also clarifies how performance metrics interact with design decisions, revealing why models that excel under one metric may fail under another. Controllability and interpretability further motivate attention to design. Applications such as editing, simulation, and reasoning require models to respond predictably to inputs and constraints. Design choices that entangle factors of variation or rely on implicit correlations can undermine this goal, leading to brittle or unintuitive behavior.

1.3.3 *Engineering Perspectives*

From an engineering standpoint, the design space frames generative modeling as a system-level problem. Choices about architecture, training, and inference must be aligned with deployment requirements, such as latency, scalability, and robustness. Engineering considerations often expose practical trade-offs that are not apparent from isolated benchmarks. For example, autoregressive visual generative models may outperform other alternatives on 256×256 image generation but fail to generalize to arbitrary resolutions due to quadratic increases in compute requirements.

Systematic exploration of the design space enables informed compromises. It allows practitioners to tailor generative models to specific applications, rather than relying on monolithic solutions. In this sense, understanding the design space is a prerequisite for turning generative modeling from a research artifact into a reliable system that works at scale.

1.3.4 *Theoretical Perspectives*

Beyond engineering concerns, the design space of visual generative models raises fundamental theoretical questions. Different modeling choices reflect different assumptions about the structure of visual data—what should be represented explicitly, what can be learned implicitly, and how uncertainty should be handled. Studying these choices helps clarify the relationship between model architecture, learning dynamics, and generalization.

Theoretical perspectives also seek to explain why certain designs work. For instance, questions about sample efficiency, expressivity, and inductive bias are central to understanding how generative models scale with data and compute. A principled view of the design space can reveal connections between seemingly disparate models, highlighting shared mechanisms and underlying principles.

1.3.5 Toward Principled Design

As visual generative models continue to evolve, the need for principled design becomes increasingly pressing. Without a clear understanding of the design space, progress risks becoming driven primarily by scale and empirical trial-and-error. By contrast, a structured view of design choices enables more systematic exploration, clearer evaluation, and deeper insight into the nature of visual generation itself.

In this context, the design space is not merely a catalog of architectures or techniques. It is a framework for reasoning about how visual generative models are built, why they behave as they do, and how they can be improved—both as engineering systems and as objects of scientific study.

1.4 Main Contributions

As noted above, visual generative models are made of various integrated components in a very brittle way which makes them hard to scale and resource consuming. In this dissertation, I will address these shortcomings by focusing on different design aspects of visual generative models from theory to applied, from diffusion to flow matching to flow maps and autoregressive modeling, and from semantic encoders to reconstruction encoders to improve visual generative models’ performance, efficiency, and controllability for various tasks such as text-to-image generation, image editing, solving inverse problems, few-shot image generation, etc. In this section, I will detail key contributions and provide a high-level summary of various chapters. This involves:

- (evaluation) **Compositional evaluations** for text-to-image generative models to understand the failure modes on visual and text based concept following (such as concept alignment and caption alignment).
- (text-to-image mapping) **Understanding the role of semantic encoders** for signifi-

cantly boosting performance, training efficiency and controllability for text-to-image and personalized image generative models.

- (synthetic data) The role of visual generative models for **synthetic data creation** for complex image editing (such as referring expression) and compositional semantic encoders (such as “mug in the grass” vs. “grass in the mug”).
- (theoretical foundation) **Building the theory** of flow matching velocity field dynamics for few-step generative models and efficient solving downstream tasks (such as image editing, solving inverse problems and removing unwanted content).
- (autoregressive generation) **Efficiently scaling autoregressive models** to support fast and accurate image generations on arbitrary resolutions to bring autoregressive models closer to diffusion models at scale.

Chapter 2 presents ConceptBed, a compositional evaluation framework for text-to-image generative models that probes how well learned visual concepts are followed under complex textual compositions. It introduces two complementary criteria—concept alignment and composition (caption) alignment—and a large-scale benchmark of 284 concepts with $\sim 33\text{K}$ composite prompts spanning objects, attributes, styles, and relations. To quantify failures beyond photorealism, the chapter proposes Concept Confidence Deviation (CCD), a confidence-based metric using oracle classifiers and VQA that correlates strongly with human judgments. Extensive evaluations reveal a consistent trade-off: methods that strongly memorize concepts often fail at compositional reasoning, while composition-preserving methods weaken concept fidelity. Together, these findings motivate compositional evaluation as essential for understanding and improving personalized text-to-image models.

The chapters 3 and 4 study the role of semantic encoders—particularly CLIP-style vision–language representations—in improving text-to-image generation, showing that

strong semantic alignment in latent space is a key driver of performance, efficiency, and controllability. It introduces ECLIPSE, a contrastive learning framework for training compact text-to-image priors that distill semantic structure from pre-trained vision–language models, achieving state-of-the-art compositional text following with orders-of-magnitude fewer parameters and training data than diffusion-based priors. Building on this insight, the chapter 4 presents λ -ECLIPSE, which operates directly in the CLIP latent space to enable fine-tuning-free, multi-concept personalized image generation, balancing concept fidelity and compositional alignment while dramatically reducing compute. Together, these works demonstrate that carefully designed semantic encoders and priors can replace brute-force scaling, offering a principled path toward efficient, controllable, and compositional text-to-image and personalized generative models.

The chapter 5 investigates how visual generative models enable high-quality synthetic data creation to address compositional challenges that are difficult to capture with web-scale supervision alone. It presents TripletCLIP, which leverages text-to-image generative models to synthesize hard negative image–text pairs that explicitly encode compositional distinctions (e.g., “mug in the grass” vs. “grass in the mug”), and shows that such synthetic negatives substantially improve the compositional reasoning of semantic encoders. This work establishes synthetic data as a principled and scalable tool for learning compositional semantics, highlighting how generative models can actively shape the training distributions needed for robust visual understanding models as well.

The chapters 6 and 7 develop a theoretical foundation for flow matching that explains why few-step generative models can be both efficient and controllable, and how their vector fields can be systematically steered to solve downstream tasks. We first analyze the error dynamics of ODE-based generative processes, showing that diffusion-style stochastic trajectories suffer from curvature-induced errors and trajectory crossovers, while rectified flows exhibit near-straight paths that enable stable, deterministic control. Building on this

insight, we introduce FlowChef, a theory-grounded framework that steers rectified flow trajectories directly in the velocity field, yielding gradient- and inversion-free solutions for image editing, inverse problems, and classifier guidance, and scaling efficiently to large latent models. We then extend this trajectory-centric view to model alignment and safety via EraseFlow, which formulates concept erasure as trajectory redistribution and proves that a reward-free, trajectory-balance objective can provably align denoising dynamics to safe anchors while preserving the prior. Finally, we study few-step generation through a geometric lens, deriving curvature-aware identities for MeanFlow and proposing piecewise and curvature-regularized flow maps that stabilize training and improve the FID–vs.–NFE trade-off. Together, these results unify efficiency, controllability, and safety under a single theoretical perspective on **flow velocity fields and their geometry**, providing principled tools for fast generation and robust downstream control.

At last, the chapter 8 presents a scalable approach to autoregressive image generation at arbitrary resolutions. It introduces VibeToken, a resolution-agnostic 1D tokenizer that compresses images of any size or aspect ratio into a short, dynamic token sequence, and VibeToken-Gen, an AR generator whose inference cost is independent of image resolution. By decoupling token length from spatial size, the model enables fast, high-quality generation up to 1024×1024 with orders-of-magnitude lower compute than prior AR methods, bringing autoregressive models closer to diffusion-level performance at scale.

Chapter 2

EVALUATING CONCEPT LEARNING ABILITIES OF TEXT-TO-IMAGE DIFFUSION MODELS

2.1 Introduction

Humans reason about the visual world by aggregating entities that they see into “visual concepts”: both *cats* and *elephants* are *animals*, and both *palms* and *pinos* are *trees*. We use natural language to describe images and things that we see. Although this type of visual concept learning is well-defined in human psychology [180], it remains elusive in the context of data-driven techniques capable of learning and reasoning from images and their natural language descriptions.

Text-to-Image (T2I) generative models are trained to translate natural language phrases into images that correspond to that input. High-quality T2I models, therefore, serve as a link between human-level concepts (expressed in natural language) and their visual representations and are one way to reproduce visual concepts. On the other hand, this has also sparked interest in visual concept learning (*a.k.a.* personalized T2I) through the procedure of “*image inversion*” – to translate one or many images corresponding to a visual concept into a latent representation of that visual concept. While earlier methods primarily explored image inversion using generative adversarial networks [287], methods such as Textual Inversion [72] and Dreambooth [222] combine image inversion with T2I – this has led to an effective way to quickly learn concepts from a few images and reproduce them in novel combinations and compositions with other concepts, attributes, styles, etc. These methods aim to learn concepts with minimal reference images by fine-tuning pre-trained text-conditioned diffusion models (Figure 2.1). Therefore this paradigm of T2I and image

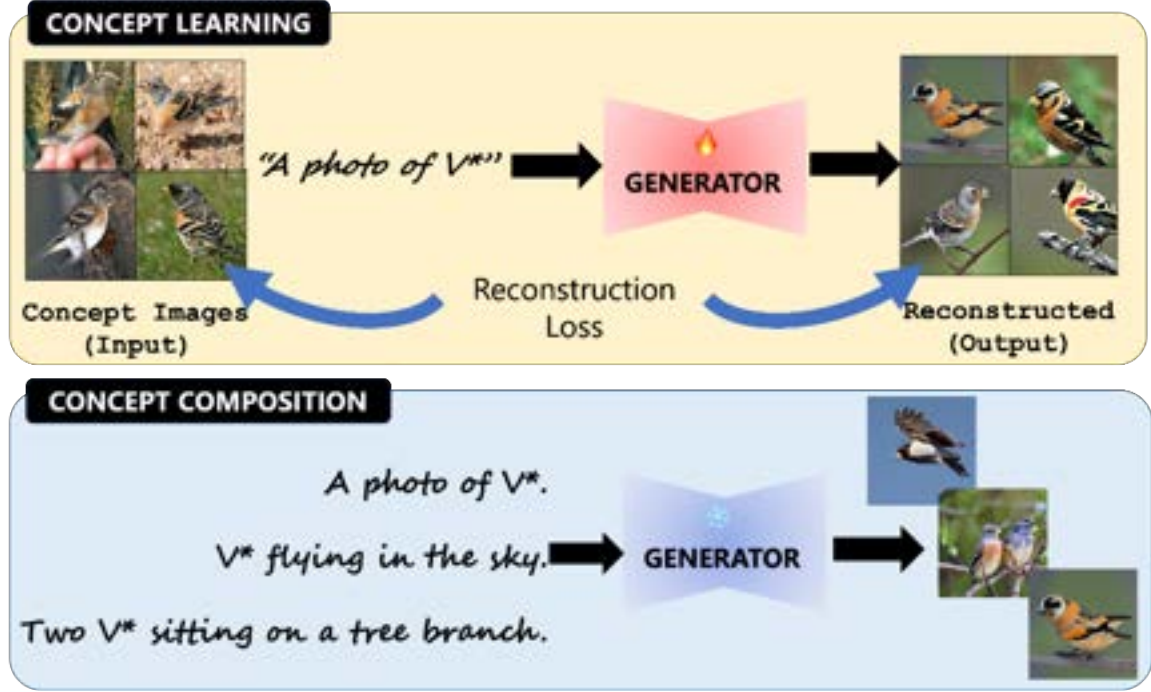


Figure 2.1: Visual concept learners such as textual inversion models learn to “invert” a set of images (about a concept c) into a text embedding V^* , and then use this learned textual concept in new text prompts to generate images of concept c under different contexts and by performing novel compositions with other concepts. The proposed CONCEPTBED dataset along with the evaluation metric $\mathcal{D}$ allows us to comprehensively and quantifiably evaluate concept learning abilities of text-to-image diffusion models.

inversion is a powerful new way of learning and reproducing concepts.

Within this paradigm of novel visual concept learning via image inversion, two primary evaluation criteria have emerged: (1) concept alignment, which assesses the correspondence between the generated images and the target concept images, and (2) composition alignment, which evaluates whether the generated images maintain compositionality. Previous studies have been small scale, evaluating only a small number of hand-picked concepts and compositions; as such making generic claims via such findings is difficult. Furthermore, the established evaluation metrics such as DINO-based cosine similarity [222] (for measuring concept alignment), KID [128] (for measuring the amount of concept overfitting), and CLIP-Score [93] (for evaluating compositionality), have encountered challenges in accurately

capturing human preferences. Consequently, there is a growing need for better automated evaluations.

Therefore, we introduce CONCEPTBED, a comprehensive dataset and evaluation framework that is aligned with human preferences. The CONCEPTBED dataset comprises 284 distinct concepts and approximately 33,000 composite text prompts, which can be further extended using the provided automatic realistic dataset creation pipeline. The dataset focuses on four diverse concept learning evaluation tasks: learning styles, learning objects, learning attributes, and compositional reasoning. To gain a deeper understanding of previous methodologies, we incorporate four composition categories – action, attribution, counting, and relations.

We use our large-scale dataset to evaluate concept learners, by developing a novel evaluation metric called Concept Confidence Deviation (CCD). We conduct a human study and find that relative evaluations of models in terms of CCD are well aligned with human preferences. Therefore, CCD combined with the CONCEPTBED dataset, offers an alternative to existing evaluation strategies, facilitating more effective large-scale evaluations. For each evaluation criterion, we train supervised classifiers (oracles) to detect whether generated concept images are accurate. Subsequently, the confidence scores from these oracles are utilized to calculate the instance-level concept deviations of the generated concept images in relation to the reference target ground truth images using the proposed CCD metric. This approach enables us to assess concept and composition alignment more effectively. We further show that CCD calculated using a pre-trained few-shot classifier also maintains a high correlation with human preferences. This allows CCD to measure concept alignment on unseen concepts.

We conduct extensive experiments on four recently proposed concept learning methodologies. In total, we fine-tune approximately 1100 models (one model per concept) and generate over 500,000 concept-specific images. Our results reveal a surprising trade-off be-

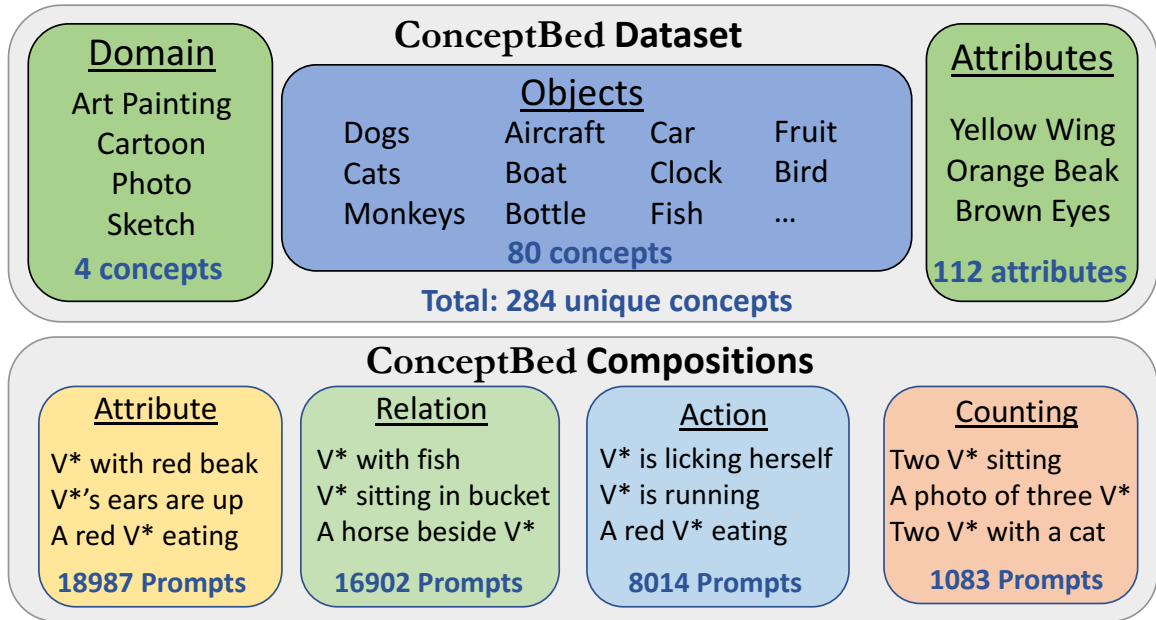


Figure 2.2: A summary of the CONCEPTBED dataset for large-scale grounded evaluations of concept learners. The collection of concepts is categorized into three classes: (1) Domain, (2) Objects, and (3) Attributes. CONCEPTBED has 284 unique concepts and four compositional categories. Here, V^* is a learned concept.

tween concept alignment and composition alignment, wherein methods excelling at concept alignment tend to fall short in preserving compositions and vice versa. This suggests that previous concept learning approaches are either highly overfitted or severely underfitted. Furthermore, our experiments demonstrate that utilizing a pre-trained CLIP [210] textual encoder aids in maintaining compositionality, but it lacks the flexibility required to learn complex concepts, such as *sketch*.

In summary, we make the following **key contributions**:

- We introduce CONCEPTBED, a comprehensive benchmark for grounded quantitative evaluations of text-conditioned concept learners.
- The Concept Confidence Deviation ($\mathcal{D}$) evaluation metric measures the learners' ability to preserve concepts and compositions. We demonstrate a strong correlation between $\mathcal{D}$ and human preferences.

- Through extensive experiments with 1,100+ models, we identify shortcomings in prior works and suggest future research directions. CONCEPTBED sets a standard for evaluating personalized text-to-image generative models.

2.2 Related Work

Concept Learning. Concept learning encompasses various problem statements and approaches, depending on the perspective adopted. Concept Bottleneck Models (CBMs) [122] and Concept Embedding Models (CEMs) [308] treat object attributes as concepts and propose classification strategies to identify these concepts. Neuro Symbolic Concept Learner (NS-CL) [172] aims to learn visual concepts by associating them with language semantics, enabling the model to perform visual question answering. Image Inversion Style Concept Learning [287], takes a different approach. Its objective is to invert a given concept image back into the latent space of a pre-trained model. However, text-based concept composition is not possible for such models.

Text-to-Image Generative Models. With advances in vector quantization [266] and diffusion modeling [216], text-to-image generation has improved its performance. Notable works such as DALL-E [212] train transformer models. While current state-of-the-art, diffusion-based text-to-image models such as GLIDE [182], LDM [216], and Imagen [226], have surpassed prior approaches (such as StackGAN [313], StackGAN++ [314], TReCS [121], and DALL-E [212]) and achieved superior performance. Pixart- [35] and ECLIPSE [194] further enhances T2I methods without depending on heavy compute. Additionally, as shown by [233], these T2I models also have multilingual concept understanding to a certain extent.

Text-to-Image Concept Learning. Text-conditioned diffusion models, such as LDM, have demonstrated their potential for learning novel visual concepts with only a few reference images. Textual Inversion [72] proposes learning the embedding corresponding to the placeholder (V^*) through optimization. DreamBooth [222] suggests optimizing the UNet

parameters instead of optimizing the placeholder embedding. Custom Diffusion [128] combines both approaches by optimizing the placeholder and key/value weights from the cross-attention layers for faster concept learning. These concept learners are essentially text-conditioned diffusion models and inherit the same limitations of diffusion models. One limitation is the overfitting of concepts and language drift. By optimizing model parameters on a handful of reference images, it is highly likely that the model might overfit the given concept and cannot maintain compositionality. Therefore, in this paper, we propose CONCEPTBED for systematic evaluations.

Text-to-Image Generative Model Evaluations. Evaluating generative models is not widely studied. The FID [94] score is commonly used to measure generated image quality. CLIPScore [93] is another popular evaluation metric for reference-free image-text alignment. Another study focuses on compositional evaluations of text-to-image models on small subsets (CU-Birds and Oxford-Flowers) [188]. DALL-Eval [42] evaluates reasoning skills on synthetic datasets and social biases of text-to-image generative models. DALL-Eval, VISOR [82], and LAYOUTBENCH [41] evaluate spatial reasoning abilities. Parallel work T2I CompBench [100] also adopts the idea of VQA for accurate composition evaluations. Although text-to-image model evaluations are well-explored, they lack concept-specific assessments and cannot be used for evaluating concept learning. Therefore, CONCEPTBED attempts to overcome this gap in evaluations of novel visual concept learning abilities.

2.3 CONCEPTBED

In this section, we introduce CONCEPTBED, a comprehensive collection of concepts, designed to accurately estimate concept and composition alignment by quantifying deviations in the generated images. Later, we introduce the novel evaluation framework associated with CONCEPTBED. Please refer to the Appendix for additional insights on the proposed dataset and evaluation framework.

2.3.1 CONCEPTBED: Dataset Construction

CONCEPTBED incorporates existing datasets such as ImageNet [49], PACS [138], CUB [269], and Visual Genome [124], enabling the creation of a labeled dataset. Figure 7.2 provides an overview of the CONCEPTBED dataset.

Learning Styles. We use styles from the PACS dataset: *rt Painting*, *Cartoon*, *Photo*, and *Sketch*. Each style contains images corresponding to seven categories. The concept learner aims to use examples from one style as a reference and generate style-specific images for all seven entities.

Learning Objects. Extracting object-level concepts is accomplished through the utilization of the ImageNet dataset. It comprises 1000 low-level concepts from the WordNet [67] hierarchy. However, due to the presence of noise in ImageNet images and the lack of relevance to daily life for many concepts, we employ an automated filtering pipeline to ensure the usefulness and quality of the reference concept images. The pipeline involves extracting a list of low-level concepts and their parent concepts from ImageNet, followed by extracting text phrases from Visual Genome containing the concept as a subject in the caption. If an insufficient number of such captions exists (less than 10 in Visual Genome) or they cannot be found, the concepts are discarded. This filtering process results in 80 concepts such as (*brambling*, *squirrel monkey*, etc.). We select the top 100 high-quality images for each concept that will be used to train the concept learning methodologies.

Learning Attributes. Since ImageNet dataset images are not labeled based on the attributes present in the image, it is necessary to rely on datasets that provide attribute-level grounded labels. Therefore, we additionally employ the CUB dataset, which offers attribute-level labels (such as *orange wing*, *blue forehead*, etc.), enabling the CONCEPTBED to perform evaluations and measure the attribute-level performance of concept learners.

Compositional Reasoning. In addition to learning new concepts, it is crucial to maintain

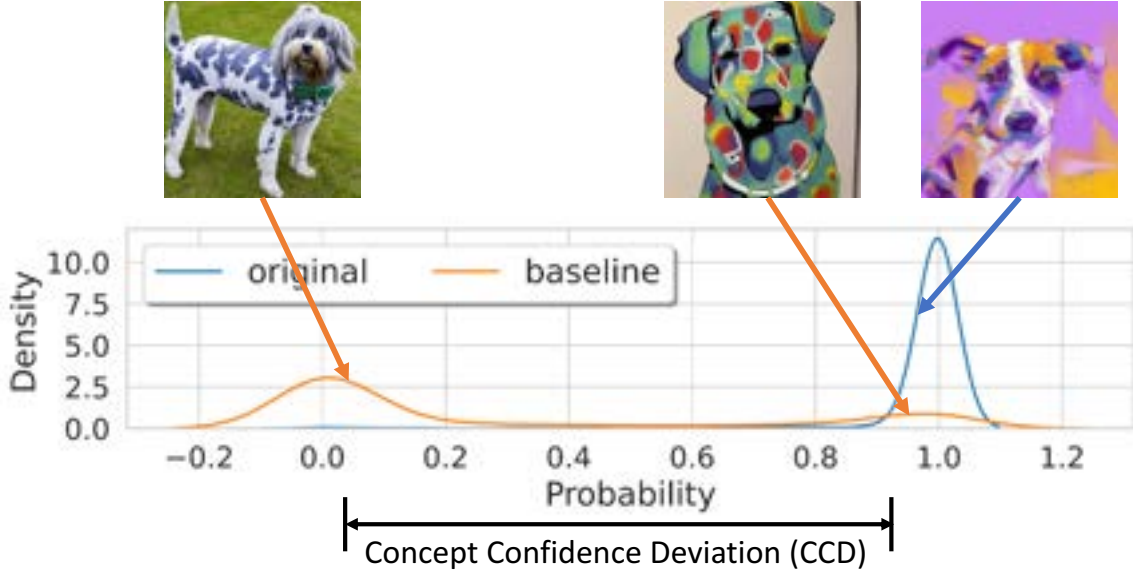


Figure 2.3: Intuitive illustration of the Concept Confidence Deviation (CCD) for the concept *Painting*. Blue and Orange are the probability distributions of the real and generated concept images.

prior knowledge and associate the acquired concepts with it. To conduct these evaluations holistically, we use Visual Genome to extract captions in which the concept appears as the subject of the sentence. These captions are categorized into four composition categories (actions, attributes, counting, and relation) through few-shot classification using GPT3 [25]. This categorization allows us to measure the performance of the baselines on each category, and provides an in-depth understanding of the varying difficulty levels of different compositions. The concepts and their respective categories are presented in Table 2.11. We use the CUB dataset for attribute-level analysis, while the ImageNet concepts are included to ensure a diversity of concepts.

2.3.2 ImageNet Dataset Generation Pipeline

As mentioned in Section 2.3.1, ImageNet contains 1000 classes but not all of them are used in day-to-day interactions. Moreover, performing experiments on each of these 1000 classes is computationally very extensive, as one needs to train 4000 models and

Algorithm 1: CONCEPTBED Object-Concepts Collection Pipeline

```

1 [0] Input:  $Y_{VG}$  = Visual Genome Captions;
2    $C_{ImageNet} = \{(c, X_c)\}_1^N$ ;
3 Output: Estimated  $C_{CONCEPTBED} = \{(\hat{X}_c, c)\}_1^M$  [1] Initialize:  $C_{ConceptBed} = []$ ;  $M = 0$ 
    $(c, X_c) \in C_{ImageNet}$   $count = 0$   $y \in Y_{VG}$   $subject(y) == c$   $count = count + 1$   $count \geq 10$ 
    $\hat{X}_c = []$   $x \in X_c$   $area = \frac{\# \text{ of pixels}(c)}{\# \text{ of total pixels}}$   $area \geq 0.4$   $\hat{X}_c \leftarrow x$   $\hat{X}_c = sorted(\hat{X}_c)$ 
    $C_{CONCEPTBED} \leftarrow (\hat{X}_c[:100], c)$   $M = M + 1$ 

```

generate 400,000 images. Therefore, it is important to filter out concepts that are commonly encountered in daily life. To measure real-life importance, we check if any concept (such as *dog*) is the subject of a caption prompt in the whole Visual Genome dataset. If there exist at least 10 captions having the concept as subject, then we add the concept to the CONCEPTBED library. Additionally, the concept learning methodologies can learn new concepts using as little as 4 images. Using all ImageNet images as training data can potentially add more noise, as these images are not high resolution. Hence, we further filter out the top 100 images based on the percentage of the object pixels (with a ratio of at least 0.4) within the image. We provide Algorithm 1 for readers' understanding of the data generation pipeline. It is worth noting that this pipeline can be used to extend CONCEPTBED and even to train future concept learning methodologies.

2.3.3 Composition Categorization

We leverage the Visual Genome dataset to create composite prompts for each of the 80 ImageNet concepts. This process yields over 33,000 compositions, resulting in a rich variety of prompts. Table 2.13 provides detailed statistics on the compositions for each concept. Furthermore, Figure 2.4 illustrates the distribution of composition categories within the CONCEPTBED dataset. For the sake of simplicity, CONCEPTBED contains composite

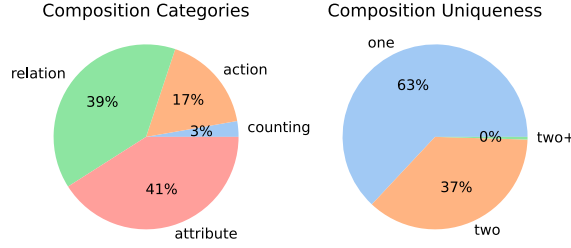
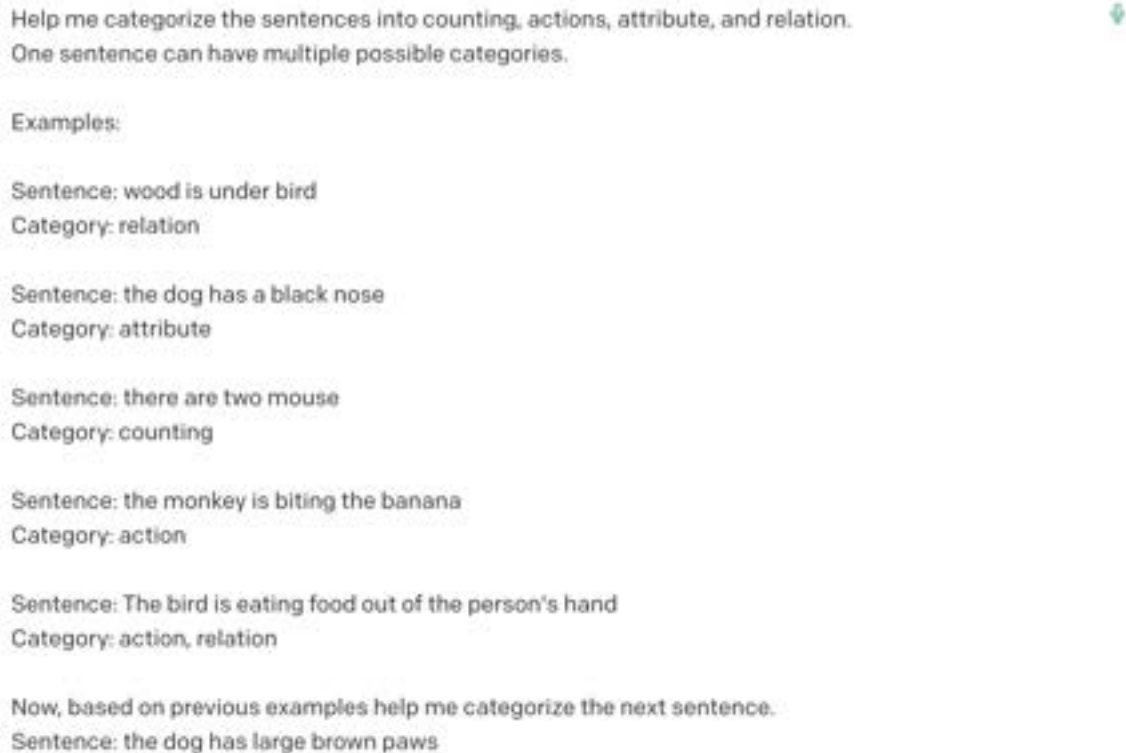


Figure 2.4: Left: word-cloud of the CONCEPTBED compositions. Right: statistics of composition categories and their combinations in composite text prompts.

prompts that combine up to two different compositions. To determine the composition type, we employ the GPT3 (text-davinci-003) model for few-shot classification. Figure 2.5 showcases the instruction and in-context examples used to categorize each text phrase.

2.3.4 Question Generations

Rather than relying solely on image-text similarity, we evaluate compositions through VQA performance using synthetically generated boolean questions (i.e., *yes* or *no*) based on the composite text phrases. To create these questions, we manually filter out salient words such as nouns, attributes, and verbs, and formulate questions corresponding to each of these words. This process enables the creation of existence-related questions where the ground truth answer is always *yes*. Table 7.8 provides examples of the questions generated for different composite text prompts.



Help me categorize the sentences into counting, actions, attribute, and relation.
One sentence can have multiple possible categories.

Examples:

Sentence: wood is under bird
Category: relation

Sentence: the dog has a black nose
Category: attribute

Sentence: there are two mouse
Category: counting

Sentence: the monkey is biting the banana
Category: action

Sentence: The bird is eating food out of the person's hand
Category: action, relation

Now, based on previous examples help me categorize the next sentence.
Sentence: the dog has large brown paws

Figure 2.5: Screenshot depicting few-shot classification using GPT3. We keep the instruction and few-shot examples constant and change the target phrase to get the corresponding categories.

2.3.5 CONCEPTBED: *Dataset Statistics*

The CONCEPTBED dataset consists of 284 unique concepts, comprising 80 concepts from ImageNet, 200 concepts from CUB, and 4 concepts from PACS. In total, the dataset contains approximately 33,000 composite prompts for the evaluation of all 80 processed concepts from ImageNet, with each composite prompt having up to two composition categories. Out of these composite prompts, 18987, 16902, 8014, and 1083 prompts contribute to the attribute, relation, action, and counting categories, respectively.

Our dataset curation pipeline is flexible to be extended to larger datasets such as OpenImages-v7 [130] and LAION-5B [237]. However, it is important to note that this extension would significantly increase the resource requirements. With the introduction of this dataset, our primary objective is to provide a standardized and benchmarked evaluation

Caption	Generated Questions
two <i>birds</i> standing on branches	are there two birds ? are there branches ?
tongue hanging out a <i>dog</i> 's mouth	is there tongue ? is there a dog's mouth ?
<i>cat</i> is licking herself	is there cat ? is the cat licking herself ?
the <i>monkey</i> is holding onto a red bar	is there the monkey ? is there a red bar ?
Propeller blade on an <i>aircraft</i>	is there propeller blade ? is there an aircraft ?
the vine is around <i>clock</i>	is there the vine ? is there clock ?
man driving the <i>boat</i>	is there man ? is there the boat ?

Table 2.1: Examples of generated existence-related questions from captions.

framework for concept learners, enhancing research in the field.

2.3.6 D : Concept Confidence Deviation

Problem Statement. Consider a pre-trained text-conditioned diffusion model $g(\cdot)$, which can be further fine-tuned on a specific concept c such that $c \in \mathcal{C}_{\text{CONCEPTBED}}$. We assume the availability of concept-specific target images from the CONCEPTBED dataset, denoted as $\mathcal{D}_c^{\text{real}} \in \mathcal{D}_{\text{CONCEPTBED}}^{\text{test}}$. Denote the concept learner $g(\cdot)$ fine-tuned on concept c using $\mathcal{D}_c^{\text{real}}$ as $g_c(\cdot)$. First, we generate a collection of N images using the learned concept c , and denote

Algorithm 2: Concept Confidence Deviation

```

1 [0] Input: Concept fine-tuned models  $G \in \{g_c\}$ ,  $c \in C_{\text{CONCEPTBED}}$ ;
2   Oracles  $F_t \in \{F_{PAC}, F_{\text{Imagenet}}, F_{\text{CUBS}}, F_{\text{VQA}}\}$ ;
3   Reference set of concept images  $X^{ref} \in \{x_c\}$ ;
4   Target set of prompts  $Y \in \{y_c\}$ ;
5 Output: Estimated  $\mathcal{D}$  [1] Initialize:  $score = []$ ;  $p^{real} = []$   $c \in C_{\text{CONCEPTBED}}$   $p^{real} = []$ 
    $t = VQA$   $c = 3$ 
6  $x = 1 \dots M$   $p^{real} \leftarrow F_t(x_i, c)$   $p^{real} = \frac{1}{M} \sum_{i=1}^M p_i^{real}$   $n = 1 \dots N$   $x_{gen} = g_c(y_c)$ 
    $score \leftarrow -1 * (F(x_{gen}, c) - p^{real})$  // Eq. 2.1  $\mathcal{D} = \frac{1}{NC} \sum_{i=1}^{NC} score_i$ 

```

Model	Concepts		Fine-grained _{CUB}		Composition
	Domain _{PACS}	Object _{ImageNet}	Object-level	Attribute-level	
TI (LDM)	0.0478	0.0955	0.2289	0.1174	0.1906
TI (SD)	0.2456	0.0472	0.0859	0.0332	0.1090
DB	<u>0.6825</u>	0.0678	0.0963	0.0469	0.3527
CD	0.6206	<u>0.2085</u>	<u>0.3934</u>	<u>0.1743</u>	<u>0.4916</u>
Original	0.0000	0.0000	0.0000	0.000	0.0000

Table 2.2: Results of Concept Alignment Evaluation. The table shows the performance of concept learners evaluated using the $\mathcal{D}$ ($\downarrow$) metric for Concepts (Domain_{PACS}, and Object_{ImageNet}), Fine-grained_{CUB} (Object-level, and Attribute-level), and Composition. The best and worst performing models are indicated by bold and underlined numbers, respectively.

this set of images as $\mathcal{D}_c^{gen} = \{x_i^{gen} = g_c(p_c^i, s^i); \forall i \in [0, N]\}$, where p_c^i is the concept-specific prompt and s is the random seed.

The alignment between two distributions (i.e., D_c^{real} and D_c^{gen}) is typically computed by first extracting features from the model m (i.e., $f_{real} = m(D^{real})$; $f_{gen} = m(D^{gen})$) and then employing a distance metric d (i.e., $score = d(f_{real}, f_{gen})$). Several combinations of models (m) and distance measures (d) have been used in prior work. For concept alignment, Ruiz et al. [222] use m =DINO with d =Cos and Kumari et al. [128] use m =Inception with d =KID. For composition alignment, all prior work utilizes m =CLIP with d =Cos. However,

Models	Relation			Action			Attribute			Counting		
	CLIP	VQA	D	CLIP	VQA	D	CLIP	VQA	D	CLIP	VQA	D
TI (LDM)	0.6589	66.60%	0.2074	0.6523	68.69%	0.2098	0.6599	72.22%	0.1331	0.6515	65.78%	0.1231
TI (SD)	0.6294	70.09%	0.1735	0.6274	70.81%	0.1884	<u>0.6360</u>	74.75%	0.1091	0.6301	68.38%	0.1020
DB	<u>0.7051</u>	82.20%	0.0542	<u>0.6995</u>	84.61%	0.0496	<u>0.6862</u>	82.24%	0.0355	0.6924	<u>78.90%</u>	<u>-0.0016</u>
CD	0.7065	82.94%	0.0471	0.7053	86.35%	0.0347	0.6940	84.20%	0.0163	<u>0.6921</u>	79.36%	-0.0054
SD	<u>0.7222</u>	83.42%	0.0403	<u>0.7178</u>	87.39%	0.0256	<u>0.7053</u>	83.85%	0.0184	<u>0.7085</u>	<u>81.07%</u>	<u>-0.0206</u>
Original	0.6626	87.45%	0.0000	0.6831	89.78%	0.0000	0.6306	85.79%	0.0000	0.6553	78.32%	0.0000

Table 2.3: Compositional Reasoning Evaluation Results. The table shows the performance of the prior works for Composition Alignment. CLIP ($\uparrow$) is the traditional image-text alignment metric. VQA ($\uparrow$) is the accuracy of the ViLT VQA classifier on generated boolean questions. And D ($\downarrow$) is the composition deviations reported from the ViLT model with respect to its performance on original images. The best-performing model is indicated by bold numbers, while the performance that is higher than the original data is reported with underline.

these methods fail to accurately capture the concept deviations within the generated images; rendering them ineffective in comparing performance across the methodologies (as shown in Section 7.5.2).

Concept Confidence Deviation (D). To address the above limitations, we propose training the oracle classifier F , specifically for the concept detection task using the CONCEPTBED training dataset, $\mathcal{D}_{\text{CONCEPTBED}}^{\text{train}}$. Then one can simply use $m = F$ and $d = \text{Accuracy}$ to verify whether x_{gen} is aligned with x_{real} . However, measuring accuracy does not allow instance-level evaluations. By leveraging the output probabilities of the oracle (concerning the concept label y_c), we can estimate the deviations associated with each generated image x_{gen} w.r.t. the output probabilities of real target images x_{real} . Concept Confidence Deviation is defined as:

$$D = - \left[\mathbb{E}_{x_{\text{gen}}} \left[F(y_c | x_{\text{gen}}) - \mathbb{E}_{x_{\text{real}}} F(y_c | x_{\text{real}}) \right] \right]. \quad (2.1)$$

D first calculates the mean target probability on the test ground truth images and then

Models	Domain _{PACS}				Object _{ImageNet}				Compositional Reasoning		
	DINO ()	KID (↓)	D (↓)	H.S. ()	DINO ()	KID (↓)	D (↓)	H.S. ()	CLIP ()	D (↓)	H.S. ()
TI (LDM)	0.5073	0.0117	0.0478	4.028	0.4708	0.0552	0.0955	4.069	0.6611	0.1684	2.851
TI (SD)	0.4104	0.0422	0.2456	4.084	0.4457	0.0294	0.0472	4.159	0.6309	0.1432	3.694
DB	0.3925	0.1101	0.6825	3.083	0.4525	0.0290	0.0678	4.075	0.6919	0.0344	3.556
CD	0.3956	0.0593	0.6206	3.164	0.4450	0.0492	0.2085	3.803	0.6968	0.0232	4.178
Correlation	0.6557	<u>-0.8252</u>	-0.9515	1.000	0.2787	<u>-0.5347</u>	-0.9892	1.000	<u>0.3486</u>	-0.7342	1.000

Table 2.4: Human Evaluations. Comparison of prior quantitative metrics and $\mathcal{D}$ metric with Human evaluations. DINO based pairwise cosine similarity is the prior evaluation metric [222]. KID was used to measure the overfitting by [128]. CLIP (CLIPScore) is the traditional reference-free image-text similarity metric. $\mathcal{D}$ is our presented concept deviation-aware evaluation metric. H.S. denotes the corresponding Human Score. Here, Domain_{PACS} and Object_{ImageNet} evaluations are for concept alignment and composition alignment is for image-text similarity. A high negative correlation between $\mathcal{D}$ and human ratings implies strong alignment, as lower $\mathcal{D}$ and higher human ratings correspond to better performance.

Model	PACS		ImageNet	
	Domain	Object	Composition	
TI (LDM)	72.84	64.53	58.28	
TI (SD)	52.25	70.79	65.42	
DB	24.71	67.45	39.42	
CD	20.12	52.06	26.31	

Table 2.5: Recall. Percentage of generated images highly aligned ($\mathcal{D} \leq 0.0$) with the target concept images.

measures the difference in probability of the generated images. $\mathcal{D}$ with negative or close to 0.0 values indicates that the generated images closely follow the distribution of the ground truth concept images. A positive $\mathcal{D}$ value suggests that the generated images deviate from the original distribution. Figure 2.3 shows an intuitive example of $\mathcal{D}$ by calculating the distance between two probability densities corresponding to the real and generated target concept.

2.3.7 Task Specific Evaluation Settings

To efficiently leverage the CONCEPTBED evaluation pipeline, we trained separate oracles on the corresponding CONCEPTBED datasets. Two different types of evaluations

are conducted, each with its respective set of oracles: 1) concept alignment, measured by concept classifiers, and 2) compositional reasoning, measured by a VQA model.

Concept Alignment. Concept alignment evaluation was performed on all tasks, including the generated concept images with different composite text prompts. To evaluate the style, a ResNet18 [87] model is trained to distinguish the images between four style concepts. To evaluate the objects, a ConvNeXt [158] model is fine-tuned on 80 classes from the CONCEPTBED using the ImageNet training subset. The Concept Embedding Model (CEM) [308] was trained on CUB to detect the concepts and attributes. Images corresponding to the concepts were generated for each task by following the prompts: “A photo of V^* ” for objects and “A photo of a *<entity-name>* in the style of V^* ” for styles. Here, *<entity-name>* belongs to the seven classes from PACS. The remaining task, composition, utilizes the same pre-trained ConvNeXt model for concept alignment, as CONCEPTBED compositions are specifically for 80 ImageNet concepts.

Compositional Reasoning. To measure the image-text alignment with respect to the input prompts, the concept-specific token (V^*) was removed and replaced with the corresponding ground truth label (i.e., *dogs*, *cats*, etc.). The image-text similarity was then measured. Unlike previous works, CLIP was not used due to its inability to capture compositions [262]. Instead, taking after [42], we propose to use a pre-trained ViLT [118] as a VQA model for composition evaluations. Specifically, from each composite prompt, the boolean questions with positive answers are generated [12]. As ViLT is essentially a classifier, the $\mathcal{D}$ can be calculated with respect to the confidence of the model associated with a “yes” answer.

2.4 Experiments & Results

In this section, we benchmark four state-of-the-art concept learning methodologies. We first explain the experimental setup and report the evaluation results using the CONCEPTBED framework along with human preferences. Additional details about the experimental setup, results, and human evaluations are in the appendix.

2.4.1 Experimental Setup

In our experiments, we study four text-conditioned diffusion modeling-based concept learning strategies: Textual Inversion (TI) on LDM and SD, DreamBooth (DB) [222], and Custom Diffusion (CD) [128]. We generate $N = 100$ images for all concepts to measure the concept alignment and $N = 3$ images for 33K composite text prompts. For a total of 284 concepts, we train all four baselines. This leads to 1100+ concept-specific fine-tuned models and we generate a total of 500,000 images for evaluations. To show the stability of D, we report the mean performance across the three seeds of oracle training.

Hyper-Parameter	Textual Inversion (LDM)	Textual Inversion (SD)	DreamBooth	Custom Diffusion
Base Model	LDM	Stable Diffusion - v1.5	Stable Diffusion - v1.5	Stable Diffusion - v1.5
Optimized	V*	V*	UNet	V* + CrossAtten(k,v)
Optimization Steps	3000	3000	400	250
Learning Rate	5e-4	5e-4	5e-6	1e-5
Place-Holder Token	*	<object>	sks	<new1>
Regularizer	-	-	✓	✓
Regularization Images	-	-	200	200
# if inference steps	50	50	50	50
Guidance Scale	7.5	7.5	7.5	7.5
Noise Scheduler	-	PNDMScheduler	PNDMScheduler	PNDMScheduler

Table 2.6: The table summarizes the different hyper-parameter settings for all baselines considered in this work to help the reproducibility of results.

Table 2.6 presents the hyperparameter details for the baseline methodologies employed in

Hyper-parameters	Domain _P CS	Object _{ImageNet}	Fine-Grained _{CUB}	Compositions
Model Architecture	ResNet18	ConceNeXt-base	-	ViLT
Pre-training Dataset	ImageNet	ImageNet	-	MSCOCO, GCC, SBU, VG
# of target concepts	4	80	200/112	1
Objective Function	NLL	NLL (w outlier exposure)	NLL	NLL

Table 2.7: This table summarizes the different hyper-parameters used for Oracles. We first take the pre-trained model weights and then fine-tune them on target concepts from the CONCEPTBED dataset. Here, NLL refers to the Negative Log-Likelihood.

our benchmarking process. To ensure fair comparisons, we generate images for all concepts using the same number of inference steps and guidance scale. Specifically, Textual Inversion (SD), DreamBooth, and Custom Diffusion utilize the Stable Diffusion V1.5 pre-trained model ¹. Regarding Textual Inversion, we explore two variants: 1) Latent Diffusion Model-based, and 2) Stable Diffusion-based. By incorporating different pre-trained models, we aim to investigate their impact on learning novel concepts, and our findings reveal significant differences. For instance, Textual Inversion (LDM) outperforms Textual Inversion (SD) when learning style as a concept, while the SD version excels in adapting object-level concepts.

Table 2.7 outlines the hyperparameter settings for each oracle. It is important to note that we can employ any type of classifier as an oracle, as elaborated in the subsequent section. In our approach, we initially take pre-trained models and further, fine-tune them on concepts using the negative log-likelihood objective. However, it is well-established that classifiers may exhibit misclassification tendencies with high confidence. To address this, for Objects_{ImageNet}, we additionally incorporate an outlier-exposure objective function [90].

¹<https://huggingface.co/runwayml/stable-diffusion-v1-5>

2.4.2 Human annotations

To evaluate the effectiveness of our CONCEPTBED evaluation framework, we conducted human evaluations consisting of three distinct studies: object-based concept similarity, style-based concept similarity, and traditional image-text similarity for evaluating compositions. For the object and style-based concept similarity, we asked participants to rate the likelihood of the target image being the same as three reference images on a scale of 1 to 5. A rating of 1 indicated the least similarity, while a rating of 5 indicated an exact match in terms of concept. We ensured that human annotators did not compare generated images from different concept learning strategies; instead, they rated each image independently. Regarding the composition evaluation, we simply asked annotators to rate the image-text similarities on the same 1-5 scale, with 1 representing the least similarity and 5 representing an exact match. Figures 2.7, 2.8, and 2.9 present screenshots of the MTurk interface used for each type of human evaluation.

To ensure comprehensive coverage, we randomly selected 100 generated images and obtained evaluations from three unique workers for each image. This resulted in a total of 900 evaluations from human annotators. To assess the relationship between human evaluations and various baseline evaluation metrics, as well as our $\mathcal{D}$ method, we computed Pearson’s correlation. Our findings indicate **a strong correlation between the human evaluations and our $\mathcal{D}$ evaluation metric.**

2.4.3 Results

Concept alignment. Table 3.1 shows the overall performance of the baselines in terms of $\mathcal{D}$, where lower score indicates better performance. First, we can observe that $\mathcal{D}$ for *concept alignment* is low for the original images; suggesting that the oracle is certain about its predictions. Second, it can be inferred that Custom Diffusion performs poorly,

while Textual Inversion (SD) outperforms the other methodologies except for the case of the learning styles. We attribute this behavior to differences in textual encoders. LDM trains the BERT-style textual encoder from scratch while SD uses pre-trained CLIP to condition the diffusion model. CLIP contains vast image-text knowledge leading to better performance on learning objects but less flexibility to learn different styles as a concept. Surprisingly, if we compare the *concept alignment* performance with and without composite prompts, we observe that the performance further drops significantly for all baseline methodologies when composite prompts are used. This shows that existing concept learning methodologies find it difficult to maintain the concepts whenever the prompt contains the composition.

Compositional Reasoning. Previously, we discussed concept alignment on composite prompts. Table 2.3 summarizes the evaluations on composition tasks. Here, we observe the complete opposite trend in results. Custom Diffusion outperforms the other approaches across the composition categories. This result shows the trade-off between learning concepts and at the same time maintaining compositionality in recent concept learning methodologies. Moreover, CLIPScore estimates the better performance of the baselines compared to the original image-text pairs which are inaccurate.

Qualitative Results. Figure 2.6 provides the qualitative examples of the concept learning. It can be inferred that Textual Inversion (LDM) learns the *sketch* concept very well (the first row), while DreamBooth and Custom Diffusion struggle to learn it. All baselines perform comparatively well in reproducing the learned concept (the second row). Interestingly, in the case of compositions, DreamBooth and Custom Diffusion perform well with the cost of losing the concept alignment (the last two rows). At the same time, textual inversion approaches cannot reproduce the compositions (like, “*Two V**”) but they maintain concept alignment. Overall, these qualitative examples align with our quantitative results

and strengthen our evaluation framework.

Human Evaluations. We perform Human Evaluations using Amazon Mechanical Turk for both types of evaluations: 1) concept alignment – to measure the alignment between generated images and ground truth reference images on $\text{Domain}_{\text{PACS}}$ and $\text{Object}_{\text{ImageNet}}$, and 2) compositional reasoning – to measure the image-text alignment. For concept alignment, we ask human annotators to rate the likelihood of the target image the same as three reference images. While for compositional reasoning we simply ask the annotators to rate the likelihood alignment of the image and the corresponding caption. Table 2.4 summarizes the performance of prior and proposed ($\mathcal{D}$) quantitative metrics *w.r.t.* the Human Score. KID performs better for domains than objects as image dynamics varies a lot in domains. [128] proposed to use KID with LAION-retrieved concept images as a reference instead of ground truth due to the scarcity of reference images. However, CONCEPTBED alleviates this limitation. Therefore, we use actual ground truth images to report KID which is more accurate. It can be inferred that the $\mathcal{D}$ is strongly correlated with human preferences and outperforms the prior evaluation metrics by a large amount.

Figure 2.10 presents qualitative examples showcasing the performance of various baseline methods across different style concepts. Notably, the textual inversion methods demonstrate limitations in preserving object-specific features and accurately learning the desired style. Furthermore, both DreamBooth and Custom Diffusion exhibit challenges in effectively capturing and reproducing the intended styles. In Figure 2.11, we delve into the object-specific learned concepts obtained through the baseline methodologies. Notably, Custom Diffusion struggles in acquiring and comprehending new concepts, thus explaining its relatively lower performance in terms of concept alignment. To gain further insights, Figure 2.12 offers a comparison of the generated images using Custom Diffusion at different random seeds. The results indicate that Custom Diffusion successfully generates the learned

Models	ConvNeXt	Inception	ViT	Few-Shot
TI (LDM)	0.0955	0.0773	0.1165	0.0823
TI (SD)	0.0472	0.0201	0.0599	0.0489
DB	0.0678	0.0485	0.0786	0.0596
CD	0.2085	0.1845	0.2286	0.1384
Correlation	-0.9892	-0.9888	-0.9816	-0.9763

Table 2.8: Ablation. Effect of different oracle models to measure concept alignment using $\mathcal{D}$.

concepts in three out of four instances. However, when tasked with generating concept-specific images based on composite text prompts, Custom Diffusion struggles to maintain fidelity to the learned concept.

Percentage of highly aligned instances. Using $\mathcal{D}$, we can further measure the recall of the concept learning models. DINO and KID metrics do not allow us to measure the recall. Hence, it becomes hard to investigate the actual quality of the generated images. Table 2.5 shows the recall ($\frac{\text{sample with } \mathcal{D} \leq 0.0}{\text{total samples}} * 100$) for the concept alignment shown in Table 3.1. It can be inferred that Custom-Diffusion can work once in every four generation attempts. While Textual Inversion will work at least once in every two attempts. At the same time, when composition prompts are provided, Textual Inversion consistently maintains the concept alignment at the cost of achieving the composition alignment.

Generalization. Fine-tuned oracles cannot be generalized to unseen concepts, making $\mathcal{D}$ unreliable on OOD concepts. Hence, we propose to utilize a few-shot classifier (5-way 5-shot) instead, which can allow the generalization to unseen concepts while maintaining a high correlation (shown in Table 4.9). This shows the effectiveness of using confidence and $\mathcal{D}$ as the alternative to the DINO, KID, and CLIP.

Different Confidence Measures Table 2.9 presents a comprehensive comparison of various confidence quantification metrics employed in Out-Of-Distribution (OOD) detection.

Notably, all these metrics outperform the baseline metrics DINO and KID, as evidenced by their consistently high correlation scores, reaching an absolute high correlation of at least 0.9. This implies that our evaluation framework supports multiple metrics to measure the alignment as we are performing supervised learning to train the oracles. Importantly, Accuracy and ECE measure the performance of a large collection of generated images. While MSP and $\mathcal{D}$ measure the performance at the instance level, which is more useful in practical scenarios where we don't have access to a lot of generated images to estimate the performance. Although MSP also achieves a high correlation, in some cases, there might be a chance that an oracle can predict the wrong class with high confidence (as it is class-label independent). For instance, MSP on domain alignment leads to only a 0.14 correlation with human preferences. Hence, conditional probability is important to measure the instance-level alignment. It is worth noting that the negative sign in the correlation coefficients stems from the inherent differences between the nature of these metrics. Specifically, lower values of $\mathcal{D}$, and ECE indicate better performance, while higher scores in human evaluations indicate superior performance.

Choice of Classifiers for Oracles. In Table 2.10, we explore the impact of utilizing different types of classifiers as oracles. Our analysis encompasses four distinct classifiers, each characterized by an increasing number of parameters. Intriguingly, the choice of classifier appears to have a negligible effect, as $\mathcal{D}$ consistently demonstrates strong correlations with human scores, surpassing a Pearson's correlation of at least -0.98 .

2.5 Conclusion

In this paper, we introduce a novel benchmark called CONCEPTBED designed to assess the efficacy of text-conditioned diffusion models in learning new concepts (*a.k.a.* personalized T2I). The CONCEPTBED benchmark encompasses an end-to-end evaluation

Models	Accuracy ()	MSP ()	ECE (↓)	D(↓)
Textual Inversion (LDM)	80.07%	0.8734	0.0755	0.0955
Textual Inversion (SD)	84.30%	0.9022	0.0623	0.0472
DreamBooth	83.17%	0.8923	0.0647	0.0678
Custom Diffusion	69.73%	0.8311	0.1382	0.2085
Original	89.31%	0.9276	0.0436	-0.0000

Table 2.9: This table summarizes the results using different existing confidence quantification metrics on CONCEPTBED dataset on 80 concepts from ImageNet. Here, MSP refers to the Maximum Softmax Probability and ECE refers to Expected Calibration Error.

Models	ResNet18	Inception-V4	ViT-Large	ConvNeXt
Textual Inversion (LDM)	0.0107	0.0773	0.1165	0.0955
Textual Inversion (SD)	-0.0100	0.0201	0.0599	0.0472
DreamBooth	0.0214	0.0485	0.0786	0.0678
Custom Diffusion	0.1538	0.1845	0.2286	0.2085
original	0.0000	0.0000	0.0000	0.0000

Table 2.10: This table summarizes the D performance based on the different types of classifiers across the parameters range.

pipeline, a comprehensive concept library, and a novel Concept Confidence Deviation (D) evaluation metric. We conduct evaluations based on two key criteria: concept alignment and composition alignment. Through extensive experiments, we demonstrate that existing text-conditioned diffusion model-based concept learners exhibit significant limitations in their performance. We perform human evaluations to validate the effectiveness of our proposed evaluation metric (D), which showcases a strong correlation with human preferences. This finding positions D as a viable alternative to human judgments, enabling

Table 2.11: List of all concepts from CONCEPTBED library based on their data source.

large-scale and comprehensive evaluations. CONCEPTBED represents the first large-scale concept-learning dataset that facilitates precise and accurate evaluations of personalized text-to-image generative models.

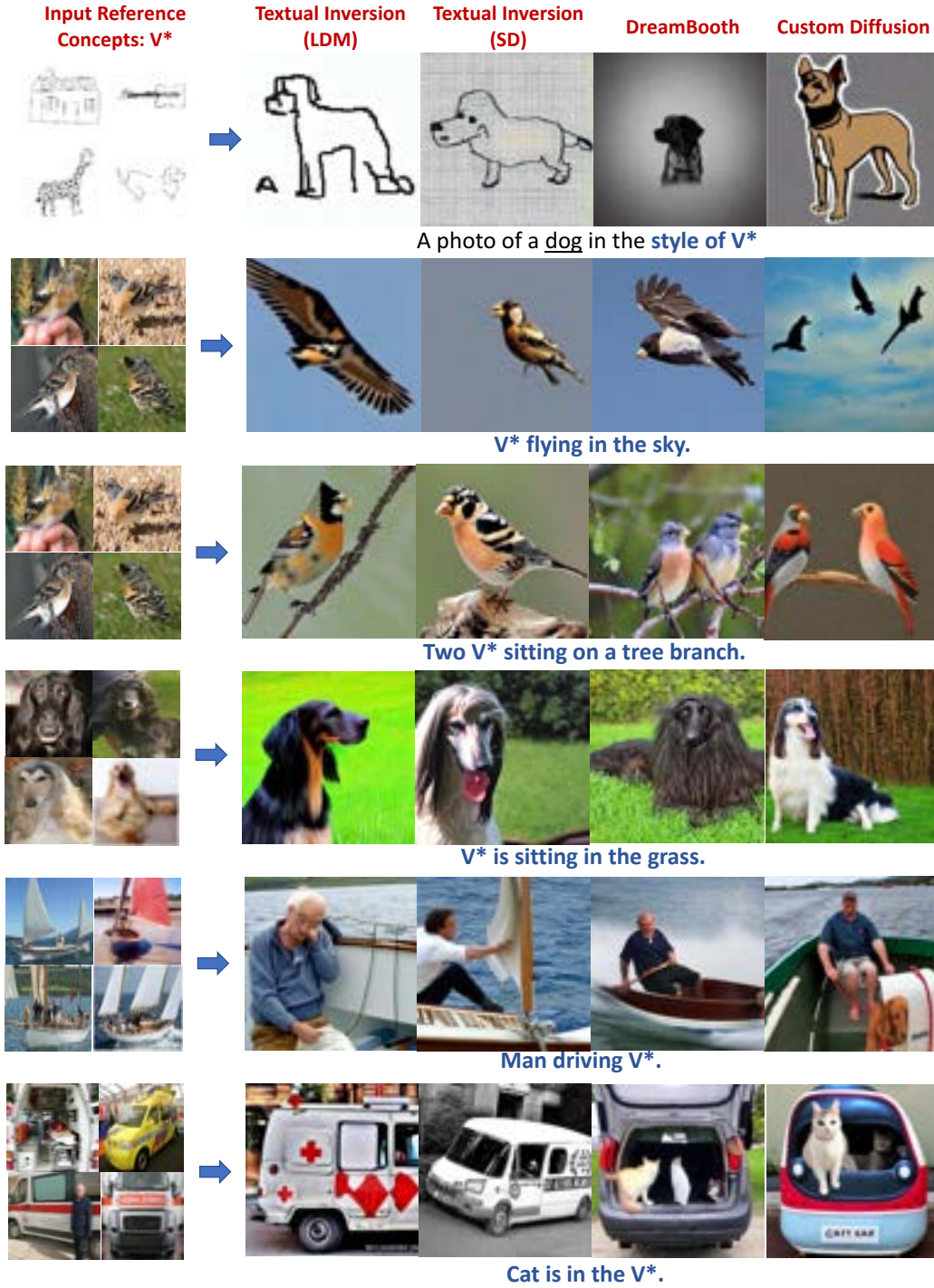


Figure 2.6: Qualitative examples showcasing the effectiveness of concept learners on the CONCEPTBED dataset. The leftmost column displays four instances of ground truth target concept images (V^*). Subsequent columns exhibit target concept-specific images generated by all baseline methods.

Concept	action	attribute	Counting	Relation	Overall
<i>laptop</i>	17	18	2	40	52
<i>tow_truck</i>	97	348	35	409	645
<i>hand-held_computer</i>	17	18	2	40	52
<i>gorilla</i>	9	11	0	11	19
<i>chimpanzee</i>	9	11	0	11	19
<i>pickup</i>	97	348	35	409	645
<i>yawl</i>	116	178	36	466	567
<i>beagle</i>	380	807	24	496	1222
<i>bulbul</i>	62	270	11	142	374
<i>spider_monkey</i>	9	11	0	11	19
<i>borzoi</i>	380	807	24	496	1222
<i>analog_clock</i>	1	61	10	71	109
<i>letter_opener</i>	11	10	0	21	31
<i>water_ouzel</i>	62	270	11	142	374
<i>web_site</i>	17	18	2	40	52
<i>garbage_truck</i>	97	348	35	409	645
<i>bloodhound</i>	380	807	24	496	1222
<i>basset</i>	380	807	24	496	1222
<i>proboscis_monkey</i>	9	11	0	11	19
<i>Dutch_oven</i>	58	85	13	111	194
<i>fireboat</i>	116	178	36	466	567
<i>black-and-tan_coonhound</i>	380	807	24	496	1222
<i>speedboat</i>	116	178	36	466	567
<i>beach_wagon</i>	98	213	17	363	497
<i>airliner</i>	4	8	2	11	20
<i>titi</i>	9	11	0	11	19
<i>marmoset</i>	9	11	0	11	19
<i>beer_bottle</i>	1	9	0	14	20
<i>magpie</i>	62	270	11	142	374
<i>Irish_wolfhound</i>	380	807	24	496	1222
<i>lifeboat</i>	116	178	36	466	567
<i>brambling</i>	62	270	11	142	374
<i>rotisserie</i>	58	85	13	111	194
<i>junco</i>	62	270	11	142	374
<i>ambulance</i>	98	213	17	363	497
<i>gondola</i>	116	178	36	466	567
<i>tabby</i>	424	992	42	727	1592
<i>cleaver</i>	11	10	0	21	31
<i>limousine</i>	98	213	17	363	497
<i>desktop_computer</i>	17	18	2	40	52
<i>colobus</i>	9	11	0	11	19
<i>house_finch</i>	62	270	11	142	374
<i>chickadee</i>	62	270	11	142	374
<i>cab</i>	98	213	17	363	497
<i>notebook</i>	17	18	2	40	52
<i>squirrel_monkey</i>	9	11	0	11	19
<i>digital_clock</i>	1	61	10	71	109
<i>canoe</i>	116	178	36	466	567

Table 2.12: This table (part 1 or 2) shows the composition statistics by categories. Here, overall means the unique compositions per concept and less than or equal to the sum of all four compositions as one composite prompt can belong up to two composition categories.

Concept	ction	ttribute	Counting	Relation	Overall
<i>indri</i>	9	11	0	11	19
<i>English_foxhound</i>	380	807	24	496	1222
<i>airship</i>	4	8	2	11	20
<i>capuchin</i>	9	11	0	11	19
<i>tiger_cat</i>	424	992	42	727	1592
<i>bluetick</i>	380	807	24	496	1222
<i>fgghan_hound</i>	380	807	24	496	1222
<i>moving_van</i>	97	348	35	409	645
<i>jay</i>	62	270	11	142	374
<i>police_van</i>	97	348	35	409	645
<i>howler_monkey</i>	9	11	0	11	19
<i>langur</i>	9	11	0	11	19
<i>gibbon</i>	9	11	0	11	19
<i>redbone</i>	380	807	24	496	1222
<i>organ</i>	3	24	12	41	68
<i>slide_rule</i>	17	18	2	40	52
<i>goldfinch</i>	62	270	11	142	374
<i>pill_bottle</i>	1	9	0	14	20
<i>siamang</i>	9	11	0	11	19
<i>convertible</i>	98	213	17	363	497
<i>baboon</i>	9	11	0	11	19
<i>Walker_hound</i>	380	807	24	496	1222
<i>guenon</i>	9	11	0	11	19
<i>indigo_bunting</i>	62	270	11	142	374
<i>grand_piano</i>	3	24	12	41	68
<i>fire_engine</i>	97	348	35	409	645
<i>robin</i>	62	270	11	142	374
<i>macaque</i>	9	11	0	11	19
<i>orangutan</i>	9	11	0	11	19
<i>jeep</i>	98	213	17	363	497
<i>patas</i>	9	11	0	11	19
<i>Madagascar_cat</i>	9	11	0	11	19

Table 2.13: This table (part 2 or 2) shows the composition statistics by categories. Here, overall means the unique compositions per concept and less than or equal to the sum of all four compositions as one composite prompt can belong up to two composition categories.

(0) The target image is a/an robin (a type of bird)

(lowest) ☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 (highest)

Target Image:



Reference Images:



Figure 2.7: An example of human annotation for determining the concept alignment for object-specific concepts.

(1) Does the target image belong to style: art_painting ?

☐ 1 (lowest) ☐ 2 ☐ 3 ☐ 4 ☐ 5 (highest)

Target Image:



Reference Images for art_painting:



Figure 2.8: An example of human annotation for determining the concept alignment for style-specific concepts.

Caption: a dog has a red leash



(1) Rate the quality of the image.
(*1' being artificial (noise, blurry) and '5' being natural (a real photograph))

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5

(2) What is the similarity between the caption and image?
(Rate from "1" (least similar) to "5" (most similar))

☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5

(3) Is there a dog ?

☐ True
☐ False

(4) Is there a red leash ?

☐ True
☐ False

Figure 2.9: The example of Human Annotation for determining the image-text alignment.

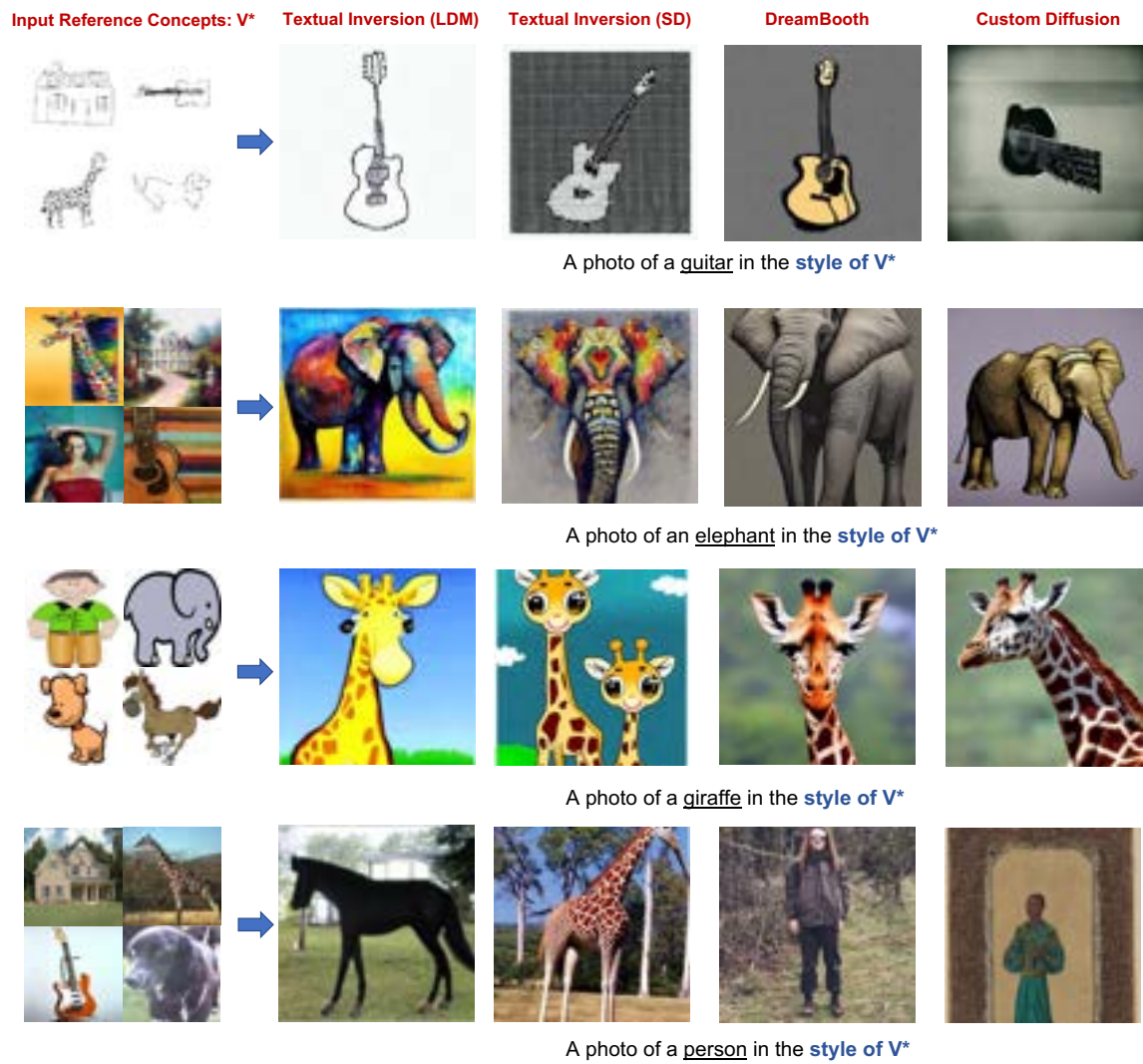


Figure 2.10: Qualitative examples of the style-specific four concepts.



Figure 2.11: Qualitative examples of the object-specific four concepts.

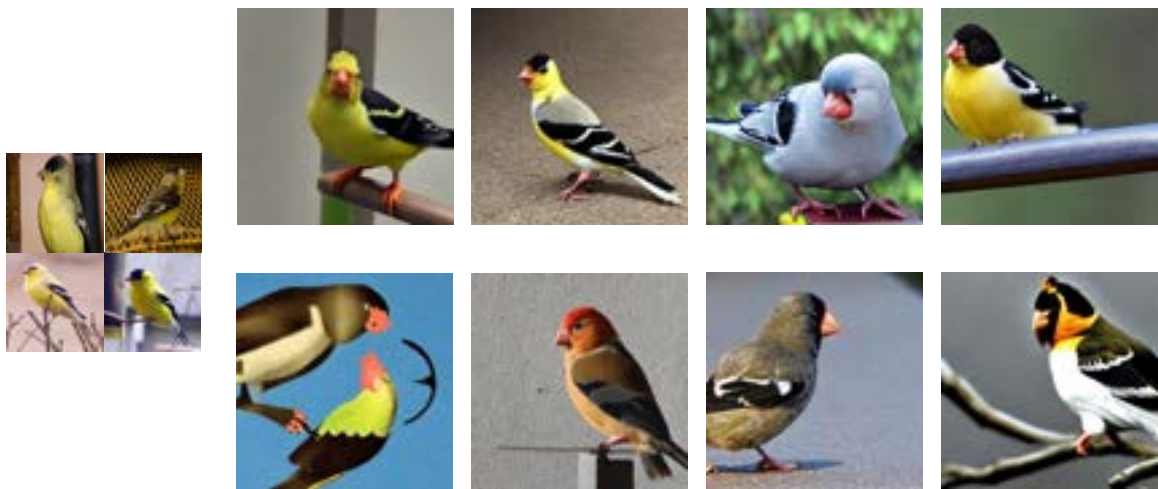


Figure 2.12: Qualitative examples from Custom Diffusions at different random seeds. The leftmost four figures are the target concept images. Top-Right four images are object-specific generated images. While Bottom-Right four generated images are on different composite text prompts.

A RESOURCE-EFFICIENT TEXT-TO-IMAGE PRIOR FOR IMAGE GENERATIONS

3.1 Introduction

Diffusion models [247, 96, 211, 216] have demonstrated remarkable success in generating high-quality images conditioned on text prompts. This Text-to-Image (T2I) generation paradigm has been effectively applied to various downstream tasks such as subject / segmentation / depth-driven image generation [32, 36, 191, 72, 142]. Central to these advancements are two predominant text-conditioned diffusion models: Latent Diffusion Models (LDM) [216], also known as Stable Diffusion, and unCLIP models [211]. The LDM, notable for its open-source availability, has gained widespread popularity within the research community. On the other hand, unCLIP models have remained under-studied. Both model types fundamentally focus on training the diffusion models conditioned on text prompts. The LDM contains a singular text-to-image diffusion model, while unCLIP models have a text-to-image prior, and a diffusion image decoder. Both model families work within the vector quantized latent space of the image [266]. In this paper, we focus on unCLIP models because they consistently outperform other SOTA models in various composition benchmarks such as T2I-CompBench [100] and HRS-Benchmark [10].

These T2I models, typically large in parameter count, require massive amounts of high-quality image-text pairs for training. unCLIP models like DALL-E-2 [211], Karlo [55], and Kandinsky [214], feature prior module containing approximately 1 billion parameters, resulting in a significant increase in overall model size ($\geq 2B$) compared to LDMs. These unCLIP models are trained on 250M, 115M, and 177M image-text pairs, respectively. Therefore, two critical questions remain: 1) *Does the incorporation of a text-to-image*

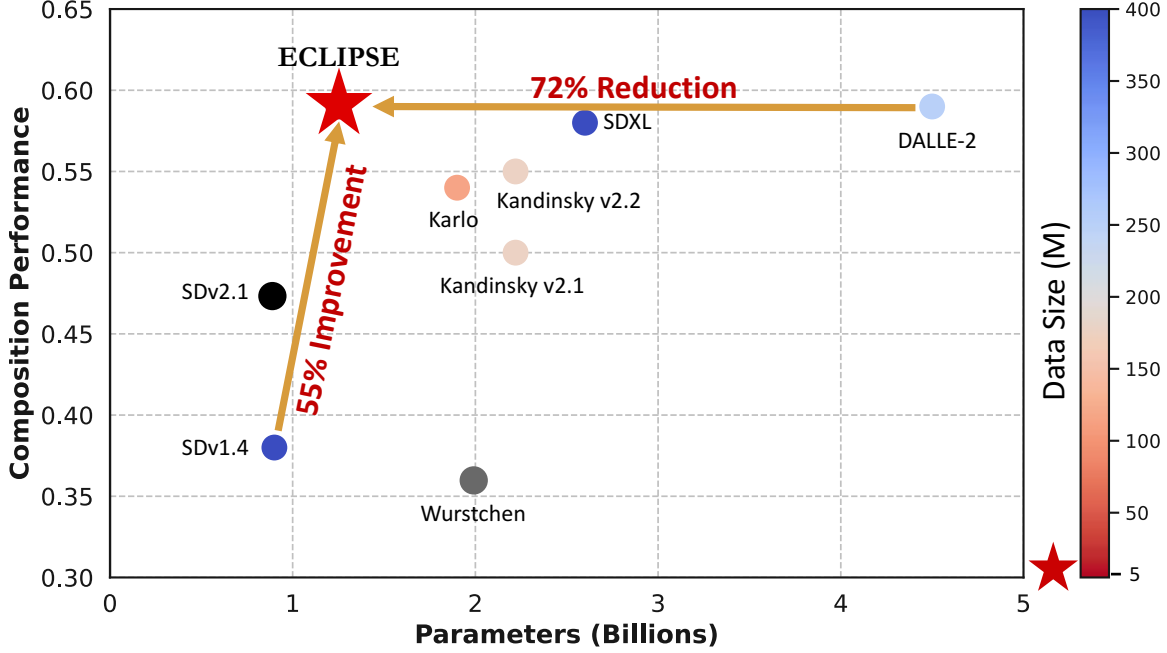


Figure 3.1: Comparison between SOTA text-to-image models with respect to their total number of parameters and the average performance on the three composition tasks (color, shape, and texture). *ECLIPSE* achieves better results with less number of parameters without requiring a large amount of training data. The shown *ECLIPSE* trains a T2I prior model (having only 33M parameters) using only 5M image-text pairs with Kandinsky decoder.

prior contribute to SOT performance on text compositions? 2) Or is scaling up model size the key factor? In this study, we aim to deepen the understanding of T2I priors and propose substantial enhancements to existing formulations by improving parameter and data efficiency.

As proposed by Ramesh et al. [211], T2I priors are also diffusion models, which are designed to directly estimate the noiseless image embedding at any timestep of the diffusion process. We perform an empirical study to analyze this prior diffusion process. We find that this diffusion process has a negligible impact on generating accurate images and having the diffusion process slightly hurts the performance. Moreover, diffusion models require substantial GPU hours for training due to the slower convergence. Therefore, in this work, we use the non-diffusion model as an alternative. While this approach may reduce the compositional capabilities due to the absence of classifier-free guidance [97], it significantly

enhances parameter efficiency and decreases the dependencies on the data.

To overcome the above limitations, in this work, we introduce *ECLIPSE*, a novel contrastive learning strategy to improve the T2I non-diffusion prior. We improve upon the traditional method of maximizing the Evidence Lower Bound (ELBO) for generating the image embedding from the given text embedding. We propose to utilize the semantic alignment (between the text and image) property of the pre-trained vision-language models to supervise the prior training. Utilizing *ECLIPSE*, we train compact (97% smaller) non-diffusion prior models (having 33 million parameters) using a very small portion of the image-text pairs (0.34% - 8.69%). We train *ECLIPSE* priors for two unCLIP diffusion image decoder variants (Karlo and Kandinsky). The *ECLIPSE*-trained priors significantly surpass baseline prior learning strategies and rival the performance of 1 billion parameter counterparts. Our results indicate a promising direction for T2I generative models, achieving better compositionality without relying on extensive parameters or data. As illustrated in Fig. 3.1, by simply improving the T2I prior for unCLIP families, their overall parameter and data requirements drastically reduce and achieve the SOTA performance against similar parameter models.

Contributions. 1) We introduce *ECLIPSE*, the first attempt to employ contrastive learning for text-to-image priors in the unCLIP framework. 2) Through extensive experimentation, we demonstrate *ECLIPSE*'s superiority over baseline priors in resource-constrained environments. 3) Remarkably, *ECLIPSE* priors achieve comparable performance to larger models using only 2.8% of the training data and 3.3% of the model parameters. 4) We analyze and offer empirical insights on the shortcomings of T2I diffusion priors.

3.2 Related Works

Text-to-Image Generative Models. Advancements in vector quantization and diffusion modeling have notably enhanced text-to-image generation capabilities. Notable works

like DALL-E [212] have leveraged transformer models trained on quantized latent spaces. Contemporary state-of-the-art models, including GLIDE [181], Latent Diffusion Model (LDM) [216], DALL-E-2 [211], and Imagen [226], have significantly improved over earlier approaches like StackGAN [313] and TReCS [121]. As these models achieve remarkable photorealism, several works focus on making T2I models more secure [111, 68, 183, 115]. LDM models primarily focus on unified text-to-image diffusion models that incorporate the cross-attention layers [216]. Additionally, several studies aim at refining Stable Diffusion models during inference through targeted post-processing strategies [32, 36, 203]. In contrast, unCLIP models, exemplified by DALL-E-2 [110], Karlo [55], and Kandinsky [214], incorporate a two-step process of text-to-image diffusion transformer prior model and diffusion image decoder having the same model architecture as LDMs. Recent benchmarks have highlighted the superior compositional capabilities of DALL-E-2 over LDM methods [100, 10]. Our work examines and enhances existing prior learning strategies in open-source pre-trained unCLIP models, Karlo and Kandinsky.

Efficient Text-to-Image Models. The current generation of T2I models is characterized by extensive parameter sizes and demanding training requirements, often necessitating thousands of GPU days. Research efforts have primarily centered on model refinement through knowledge distillation, step distillation, and architectural optimization [143, 228, 164]. Würstchen [200] presents an efficient unCLIP stack requiring less training time. Concurrently, Pixart- [35] leverages pre-trained Diffusion-Transformers (DiT) [197] as base diffusion models, further reducing training time. Distinctively, *ECLIPSE* focuses on refining text-to-image priors within the unCLIP framework using a mere 3.3% of the original model parameters, thereby significantly reducing the training duration to approximately 50 GPU hours. Our work falls orthogonal to the existing efficient T2I methodologies that mainly focus on knowledge and step distillation, and/or architectural compression. When integrated with these model compression strategies, *ECLIPSE* can position the unCLIP

family models as a compact yet highly accurate and efficient methodology.

Contrastive Learning in Generative Models. Contrastive learning, traditionally applied in visual discriminative tasks, has seen utilization in image-text alignment models like CLIP [210], LiT [310], and SigLIP [309]. However, its application in generative models, particularly in Generative Adversarial Networks (GANs), remains limited [328, 144, 45]. For instance, Lafite [328] employs a contrastive approach for image-to-text prior training in language-free T2I GANs. StyleT2I [144] attempts to learn the latent edit direction for StyleGAN [109], which is supervised via spatial masks on the images making the method not scalable. ACTIG [45] introduces an attribute-centric contrastive loss to enhance discriminator performance. These methods are constrained by their domain-specific knowledge requirements and inability to be directly applied to diffusion models [144, 45]. In contrast, *ECLIPSE* applies CLIP-based contrastive learning to train more effective T2I prior models in diffusion-based T2I systems. This strategy is not only resource-efficient but significantly enhances the traditional text-to-image diffusion priors by exploiting the semantic latent space of pre-trained vision-language models.

3.3 Methodology

This section elaborates on the Text-to-Image (T2I) methodologies, beginning with an overview of unCLIP, followed by the formal problem statement. We then delve into our proposed training strategy, *ECLIPSE*, for T2I prior in detail. Figure 3.2 provides the overview of baselines and *ECLIPSE* training strategies.

3.3.1 Preliminaries

Without the loss of generality, let’s assume that $y \in Y$ denotes the raw text and $x \in X$ denotes the raw image. z_x and z_y denote the image and text latent embeddings extracted using the pre-trained vision and text encoders ($z_x = C_{vision}(x)$; $z_y = C_{text}(y)$). Ideally, these

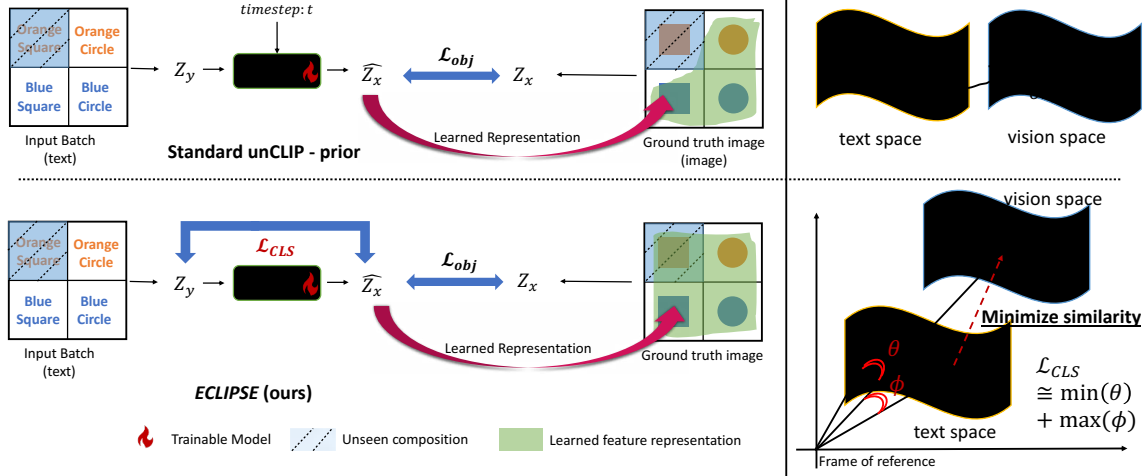


Figure 3.2: Standard T2I prior learning strategies (top) minimizes the mean squared error between the predicted vision embedding $\hat{z}_x$ w.r.t. the ground truth embedding z_x with or without time-conditioning. This methodology cannot be generalized very well to the outside training distribution (such as Orange Square). The proposed *ECLIPSE* training methodology (bottom) utilizes the semantic alignment property between z_x and z_y with the use of contrastive learning, which improves the text-to-image prior generalization.

C_{text} and C_{vision} can be any model (e.g., T5-XXL, ViT, and CLIP). Both model families (LDM and unCLIP) fundamentally focus on learning a mapping function $f_\theta : Y \rightarrow X$. The LDMs contain a singular text-to-image decoder model (f_θ), while unCLIP framework ($f_\theta = h_\theta \circ g_\phi$) contains two primary modules:

- **Text-to-Image Prior** ($g_\phi : z_y \rightarrow z_x$): This module maps the text embeddings to the corresponding vision embeddings. Ramesh et al. [211] showed that the diffusion model as T2I prior leads to slightly better performance than the autoregressive models. For each timestep t and a noised image embedding $z_x^{(t)} \sim q(t, z_x)$ (here, q is a forward diffusion process), the diffusion prior directly estimates noiseless z_x rather than estimating Gaussian noise distribution $\epsilon \sim \mathcal{N}(0, \mathcal{I})$ as:

$$\mathcal{L}_{prior} = \mathop{\mathbb{E}}_{\substack{t \sim [0, T], \\ z_x^{(t)} \sim q(t, z_x)}} \left[\|z_x - g_\phi(z_x^{(t)}, t, z_y)\|_2^2 \right]. \quad (3.1)$$

- **Diffusion Image Decoder** ($h_\theta : (z_x, z_y) \rightarrow x$): This module generates the final image conditioned on the z_x and the input text features z_y . This diffusion decoder follows

the standard diffusion training procedure by estimating $\epsilon \sim \mathcal{N}(0, I)$ after [96]:

$$\mathcal{L}_{\text{decoder}} = \mathop{\mathbb{E}}_{\substack{\epsilon \sim N(0, I) \\ t \sim [0, T], \\ (z_x, z_y)}} \left[\|\epsilon - h_\theta(x^{(t)}, t, z_x, z_y)\|_2^2 \right]. \quad (3.2)$$

Different versions of the unCLIP decoder (i.e., Kandinsky and Karlo) vary in whether they include text conditioning (z_y) in the diffusion image decoder. Both approaches yield comparable results, provided that image conditioning (z_x) is accurate. The training objectives, $\mathcal{L}_{\text{prior}}$ and $\mathcal{L}_{\text{decoder}}$, integrate Classifier-Free Guidance (CFG) [97], enhancing the model’s generative capabilities.

3.3.2 Problem Formulation

Given the pivotal role of the T2I prior module in image generation from text, in this paper, our focus is on enhancing g_ϕ , while keeping the pre-trained h_θ frozen. Let’s consider a training distribution P_{XY} , comprising input pairs of image and text (x, y) . Maximizing the Evidence Lower Bound (ELBO) on the training distribution P_{XY} facilitates this mapping of $z_y \rightarrow z_x$. However, such a strategy does not inherently assure generalization, especially when the input text prompt (y) deviates from the assumed independently and identically distributed (i.i.d.) pattern of P_{XY} [267]. Therefore, attaining a more diverse and representative P_{XY} becomes crucial for improving the performance. While a diffusion prior combined with CFG has been shown to bolster generalization, especially with diverse training data and extensive training iterations [184], it is computationally expensive and is not always reliable (especially, in low resource constraint settings) as shown in Section 3.4.2. Given these constraints, our goal is to develop an alternative prior learning methodology that improves parameter efficiency (97% reduction) and mitigates the need for large-scale high-quality data ($\leq 5\%$) while maintaining the performance.

3.3.3 Proposed Method: *ECLIPSE*

This section elaborates on *ECLIPSE*, our model training strategy to learn text-to-image prior (g_ϕ). We focus on enhancing non-diffusion prior models through the effective distillation of pre-trained vision-language models, such as CLIP, while preserving the semantic alignment between the input text embedding z_y and corresponding estimated vision embeddings $\hat{z}_x$ by using the contrastive loss.

Base Prior Model. T2I diffusion prior deviates from the standard diffusion objective (such as Eq. 3.2). Unlike the standard ϵ prediction diffusion objective, the T2I diffusion prior objective instead estimates the z_x which is noiseless. Despite its convertibility, the empirical analysis (Section 3.5) shows that having more diffusion prior steps does not benefit the overall text-to-image generation abilities. During the inference, for diffusion priors, we still adhere to the conventional denoising process, introducing additional noise ($\sigma_t\epsilon$) at each step, except for the final step according to Ho et al. [96]. This degradation in performance is likely due to the compressed latent space of the CLIP models. Moreover, if we repeat this for T timesteps, it can lead to the accumulation of errors, which is undesirable.

Therefore, to mitigate this unnecessary computing, we use non-diffusion T2I prior, making the prior model both parameter-efficient and less demanding in terms of computational resources. This non-diffusion architecture forms our base model, and we introduce the training objective that leverages pre-trained vision-language models trained on extensive datasets to improve generalization outside the P_{XY} .

Projection Objective. Despite vision-language models aligning the semantic distributions across modalities, each modality may exhibit unique distributions. Therefore, our approach involves projecting the text embedding onto the vision embedding. This is achieved using a mean squared error objective between the predicted vision embedding ($\hat{z}_x$) and the ground

truth vision embedding (z_x):

$$\mathcal{L}_{proj} = \mathbb{E}_{\substack{\epsilon \sim \mathcal{N}(0, I) \\ z_y, z_x}} \left[\|z_x - g_\phi(\epsilon, z_y)\|_2^2 \right], \quad (3.3)$$

where ϵ is the Gaussian input noise. Notably, as discussed previously, this is an approximation of the diffusion prior objective (Eq. 3.1) with $t = T$ and without CFG. $\mathcal{L}_{proj}$ learns latent posterior distribution with the *i.i.d.* data assumption. However, this model, fine-tuned on P_{XY} , may not generalize well beyond its distribution. The optimal solution would be to train on a dataset that encapsulates all potential distributions to cover all possible scenarios, which is an impractical and resource-consuming task.

Table 3.1: The comparison (in terms of FID and compositions) of the baselines and state-of-the-art methods with respect to the *ECLIPSE*. * indicates the official reported ZS-FID. Ψ denotes the FID performance of a model trained on MSCOCO. *ECLIPSE* consistently outperforms the SOTA big models despite being trained on a smaller subset of dataset and parameters.

Methods	Model	Training	Total	Data	ZS-	T2I-CompBench				
	Type	Params [M]*	Params [B]	Size [M]	FID (↓)	Color (↑)	Shape (↑)	Texture (↑)	Spatial (↑)	Non-Spatial (↑)
Stable Diffusion v1.4	LDM	900	0.9	400	16.31*	0.3765	0.3576	0.4156	0.1246	0.3076
Stable Diffusion v2.1	LDM	900	0.9	2000	14.51*	0.5065	0.4221	0.4922	0.1342	0.3096
Würstchen	unCLIP	1000	2.0	1420	23.60*	0.3216	0.3821	0.3889	0.0696	0.2949
Kandinsky v2.1	unCLIP	1000	2.22	177	18.09	0.4647	0.4725	0.5613	0.1219	0.3117
DALL-E-2	unCLIP	1000	4.5	250	10.65*	0.5750	0.5464	0.6374	0.1283	0.3043
Karlo	unCLIP	1000	1.9	115	20.64	0.5127	0.5277	0.5887	0.1337	0.3112
		33	0.93	0.6_{MSCOCO}	23.67	0.5965	0.5063	0.6136	0.1574	0.3235
		33	0.93	2.5_{CC3M}	26.73	0.5421	0.5090	0.5881	0.1478	0.3213
<i>E LIPSE (ours)</i>	Karlo	33	0.93	10.0_{CC12M}	26.98	0.5660	0.5234	0.5941	0.1625	0.3196
Kandinsky v2.2	unCLIP	1000	2.22	177	20.48	0.5768	0.4999	0.5760	0.1912	0.3132
<i>E LIPSE (ours)</i>	Kandinsky v2.2	34	1.26	0.6_{MSCOCO}	16.53	0.5785	0.4951	0.6173	0.1794	0.3204
		34	1.26	5.0_{HighRes}	19.16	0.6119	0.5429	0.6165	0.1903	0.3139

CLIP Contrastive Learning. To address these limitations, we propose utilizing the CLIP more effectively, which contains the semantic alignment between image and language. Specifically, we apply the CLIP Contrastive Loss after [210] to train the T2I priors. For a given input batch $\{(z_x^i, z_y^i)\}_{i=1}^N$ from the P_{XY} distribution, we calculate the text-conditioned

image contrastive loss for the i^{th} image embedding prediction relative to the all input ground truth text embeddings as:

$$\mathcal{L}_{CLS; y \rightarrow x} = -\frac{1}{N} \sum_{i=0}^N \log \frac{\exp(\langle \hat{z}_x^i, z_y^i \rangle / \tau)}{\sum_{j \in [N]} \exp(\langle \hat{z}_x^i, z_y^j \rangle / \tau)}, \quad (3.4)$$

where τ is the temperature parameter, $\langle \cdot, \cdot \rangle$ denotes the cosine similarity, and N is the batch size. This loss encourages the model to understand and follow the input text better, effectively reducing overfitting to the P_{XY} , as illustrated in Figure 3.2. Consequently, the final objective function is:

$$\mathcal{L}_{ECLIPSE} = \mathcal{L}_{proj} + \lambda * \mathcal{L}_{CLS; y \rightarrow x}, \quad (3.5)$$

where λ is the hyperparameter balancing the regularizer’s effect. Overall, the final objective function aims to map the text latent distribution to the image latent distribution via $\mathcal{L}_{proj}$ and such that it preserves the image-text alignment using $\mathcal{L}_{CLS; y \rightarrow x}$. This makes the prior model generalize beyond the given training distribution P_{XY} such that it can follow the semantic alignment constraint. Importantly, we cannot use $\mathcal{L}_{CLS; y \rightarrow x}$ alone or with a high value of λ as the prior model will converge outside the vision latent distribution that optimizes the contrastive loss (such input text latent space itself). And keeping λ to a very low value cannot do knowledge distillation well enough. Empirical studies suggest setting $\lambda = 0.2$ for optimal performance, balancing knowledge distillation, and maintaining alignment within the vision latent distribution.

3.4 Experiments & Results

This section introduces the datasets, training specifications, comparative baselines, and evaluation metrics utilized in our experiments. We conduct an extensive assessment of our proposed *ECLIPSE* methodology and its variants.

3.4.1 Experimental Setup

Dataset. Our experiments span four datasets of varying sizes: MSCOCO [146], CC3M [239], CC12M [29], and LAION-HighResolution ¹ [236]. MSCOCO comprises approximately 0.6 million image-text pairs, while CC3M and CC12M contain around 2.5 and 10 million pairs, respectively ². We select a very small subset of 5 million (2.8%) image-text pairs from the LAION-HighRes dataset (175M). We perform Karlo diffusion image decoder-related experiments on MSCOCO, CC3M, and CC12M as these datasets are subsets of the data used to train the Karlo diffusion image decoder. Similarly, we use MSCOCO and LAION-HighRes for the Kandinsky decoder.

Baselines. *ECLIPSE* variants are compared against leading T2I models, including Stable Diffusion, Würstchen, Karlo, Kandinsky, and DALL-E-2. Additionally, we introduce two more baselines along with *ECLIPSE* to evaluate the impact of our training strategy in a resource-constrained environment: 1) Projection: A non-diffusion prior model trained with $\mathcal{L}_{proj}$ (Eq. 3.3). 2) Diffusion-Baseline: A diffusion prior model trained with $\mathcal{L}_{prior}$ (Eq. 3.1) – the traditional T2I prior, and 3) *ECLIPSE*: A non-diffusion prior model trained with our proposed methodology $\mathcal{L}_{ECLIPSE}$ (Eq. 3.5).

Training and inference details. We evaluate *ECLIPSE* using two pre-trained image decoders: Karlo-v1-alpha and Kandinsky v2.2, trained on distinct CLIP vision encoders. Our prior architecture is based on the standard PriorTransformer model [211], modified to be time-independent. The detailed architecture is outlined in the appendix. We configure prior models with 33 and 34 million parameters for Karlo and Kandinsky, respectively. This contrasts with larger models in the field, which often use up to 1 billion parameters (as summarized in Table 3.1). The Projection, Diffusion-Baseline, and *ECLIPSE* priors

¹<https://huggingface.co/datasets/laion/laion-high-resolution>

²According to the download date: 08/26/2023

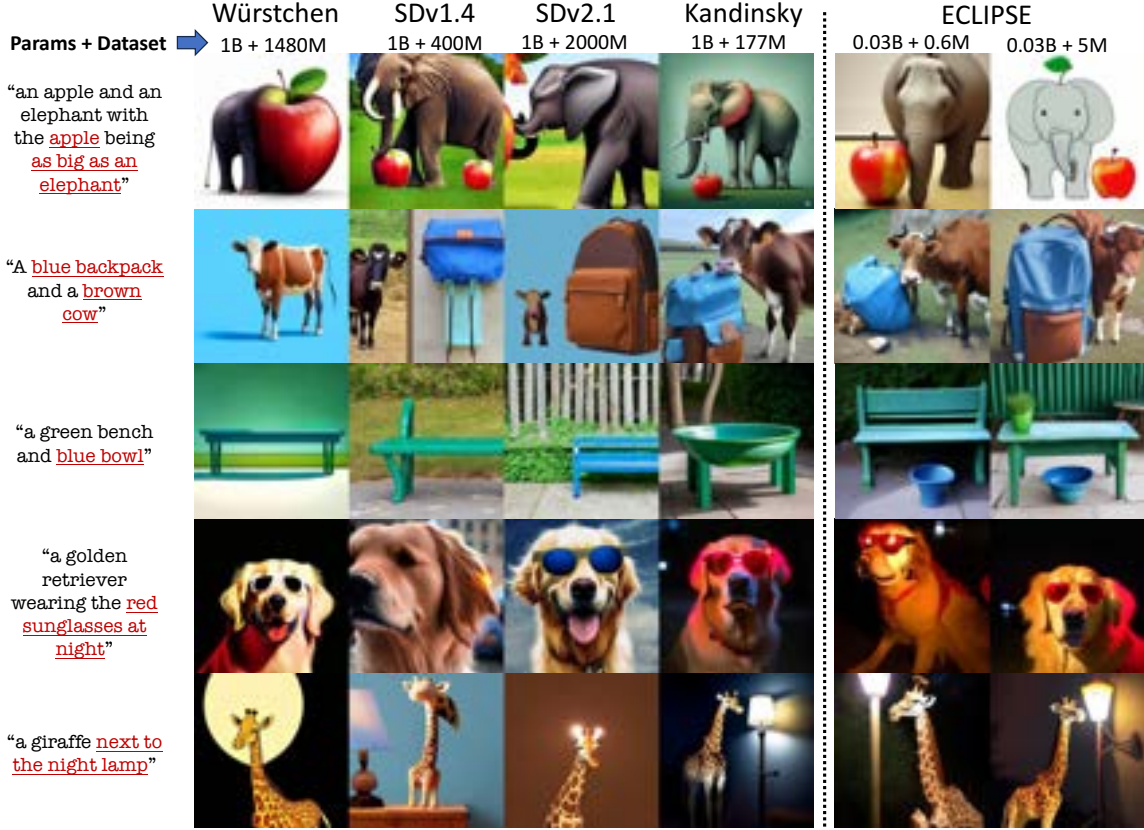


Figure 3.3: Qualitative result of our text-to-image prior, *ECLIPSE*, comparing with SOTA T2I model. Our prior model reduces the model parameter requirements (from 1 Billion 33 Million) and data requirements (from 177 Million 5 Million 0.6 Million). Given this restrictive setting, *ECLIPSE* performs close to its huge counterpart (i.e., Kandinsky v2.2) and even outperforms models trained on huge datasets (i.e., Würstchen, SDv1.4, and SDv2.1) in terms of compositions.

are trained for both diffusion image decoders, maintaining consistent hyperparameters (including total number of parameters) across all models. Training on CC12M, CC3M, and LAION-HighRes is performed on 4 x RTX A6000 GPUs with a 256 per-GPU batch size, a learning rate of 0.00005, and the CosineAnnealingWarmRestarts scheduler [161]. Each model undergoes approximately 60,000 iterations, totaling around 200 GPU hours. For MSCOCO, training takes about 100 GPU hours. This can be further reduced to ≤ 50 GPU hours if image-text pairs are preprocessed beforehand. The diffusion prior is trained with a linear scheduler and 1000 DDPM timesteps. Inferences utilize 25 DDPM steps

with 4.0 classifier-free guidance, while Projection and *ECLIPSE* models do not require diffusion sampling. Image diffusion decoders are set to 50 DDIM steps and 7.5 classifier-free guidance.

Evaluation setup. Our evaluation framework encompasses various metrics. We employ MS-COCO 30k to assess FID [94] and T2I-CompBench [100] for evaluating composition abilities in color, shape, texture, spatial, and non-spatial compositions. Given the impracticality of large-scale human studies, we approximate human preferences using PickScore [120], reporting results on the T2I-CompBench validation set comprising about 1500 unique prompts.

Implementation Details. Table 3.2 shows the comparison between *ECLIPSE*, Karlo, and Kandinsky priors. Notably, *ECLIPSE* prior uses very compressed architecture across the possible avenues (i.e., number of layers, number of attention heads, attention head dimension, etc.). Karlo uses CLIP-Vit-L/14 with 768 projection dimensions. While Kandinsky v2.2 uses the ViT-bigG-14-laion2B-39B-b160k with 1280 projection dimensions. Overall, the total number of parameters in *ECLIPSE* priors is about 33 million compared to 1 billion parameters of Karlo/Kandinsky priors. Additionally, Projection and Diffusion-Baseline use the same architecture as *ECLIPSE* prior for better comparisons. Except the Diffusion-Prior contains the additional time embeddings for diffusion modeling.

3.4.2 Quantitative Evaluations

In Table 3.1, we present a performance comparison between *ECLIPSE* variants and leading T2I models. Our evaluation metrics include zero-shot Fréchet Inception Distance (FID) on MS-COCO 30k for image quality assessment and T2I-CompBench [100] for evaluating compositionality. *ECLIPSE* priors, trained with both types of diffusion image decoders, demonstrate notable improvements. *ECLIPSE* consistently surpasses various

	<i>E LIPSE</i>	Karlo / Kandinsky Priors
Num Attention Heads	16	32
Attention Head Dim	32	64
Num Layers	10	20
Embedding Dim	768/1280	768/1280
Additional Embeddings	3	4
Dropout	0.0	0.0
Time Embed	No	Yes
Total Parameters	33/34 M	1 B

Table 3.2: Prior model architecture hyperparameter details.

baselines in terms of compositionality, irrespective of the dataset size. Its performance is comparable to that of DALL-E-2 and other SOTA models, a significant improvement considering *ECLIPSE*’s parameter efficiency. Standard T2I priors usually incorporate 1 billion parameters, while *ECLIPSE* operates with only 3.3% of these parameters, maintaining competitive performance levels. When combined with corresponding diffusion image decoders, the total parameter count of *ECLIPSE* is close to that of Stable Diffusion models, yet it outperforms them, especially considering that the latter are trained on a massive set of image-text pairs. A noticeable decline in zero-shot FID (ZS-FID) is observed in comparison to the original Karlo. We attribute this variation to the image quality differences in the training dataset, suggesting a potential area for further investigation and improvement. At the same time, if we utilize the smaller subset of high-resolution datasets then we can still maintain better FID and improve the compositions, as shown in the last row of Table 3.1. *ECLIPSE* prior with Kandinsky v2.2 decoder trained on LAION-HighRes subset achieves similar FID to other original Kandinsky v2.2 unCLIP model and at the same time

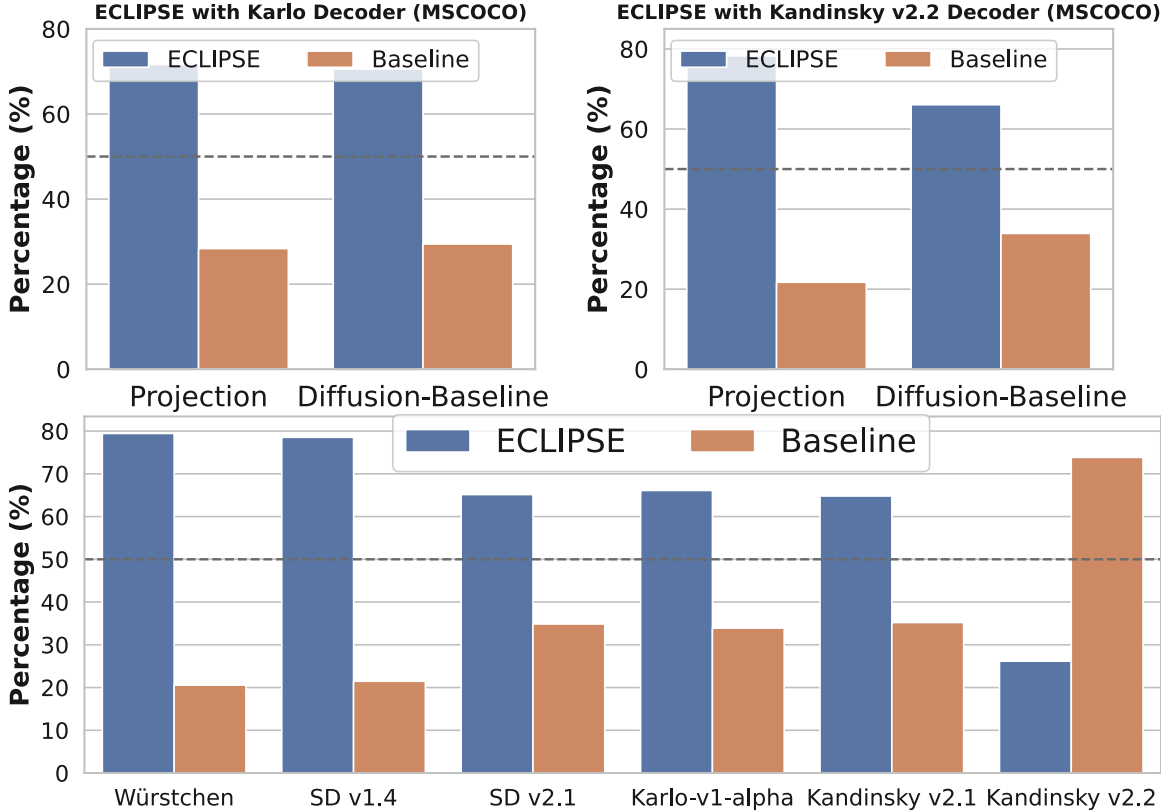


Figure 3.4: Qualitative evaluations by human preferences approximated by PickScore [120]. The top two figures compare *ECLIPSE* to Projection and Diffusion Baselines trained with the same amount of data and model size for both Karlo and Kandinsky decoders. In the bottom figure, we compare *ECLIPSE* with the Kandinsky v2.2 decoder trained on the LAION-HighRes dataset against SOTA models.

outperforming in terms of compositions.

Table 3.3 provides a comparison of various baseline training strategies for small prior models, using identical datasets and hyperparameters. *ECLIPSE* exhibits superior performance across all datasets. We also note that diffusion priors benefit from larger datasets, supporting our premise that such priors necessitate extensive training data for optimal results, which is also attributed to the CFG. In contrast, *ECLIPSE* demonstrates the consistent performance on compositions irrespective of the amount of image-text pairs.

Table 3.3: Comparison of *ECLIPSE* with respect to the various baseline prior learning strategies on four categories of composition prompts in the T2I-CompBench. All prior models are of 33 million parameters and trained on the same hyperparameters.

Methods	T2I-CompBench			
	Color (↑)	Shape (↑)	Texture (↑)	Spatial (↑)
MSCOCO with Karlo				
Projection	0.4667	0.4421	0.5051	0.1478
Diffusion-Baseline	0.4678	0.4797	0.4956	0.1240
<i>E LIPSE</i>	0.5965	0.5063	0.6136	0.1574
CC3M with Karlo				
Projection	0.4362	0.4501	0.4948	0.1126
Diffusion-Baseline	0.5493	0.4809	0.5462	0.1132
<i>E LIPSE</i>	0.5421	0.5091	0.5881	0.1477
CC12M with Karlo				
Projection	0.4659	0.4632	0.4995	0.1318
Diffusion-Baseline	0.5390	0.4919	0.5276	0.1426
<i>E LIPSE</i>	0.5660	0.5234	0.5941	0.1625
MSCOCO with Kandinsky v2.2				
Projection	0.4678	0.3736	0.4634	0.1268
Diffusion-Baseline	0.4646	0.4403	0.4834	0.1566
<i>E LIPSE</i>	0.5785	0.4951	0.6173	0.1794
HighRes with Kandinsky v2.2				
Projection	0.5379	0.4983	0.5217	0.1573
Diffusion-Baseline	0.5706	0.5182	0.5067	0.1687
<i>E LIPSE</i>	0.6119	0.5429	0.6165	0.1903

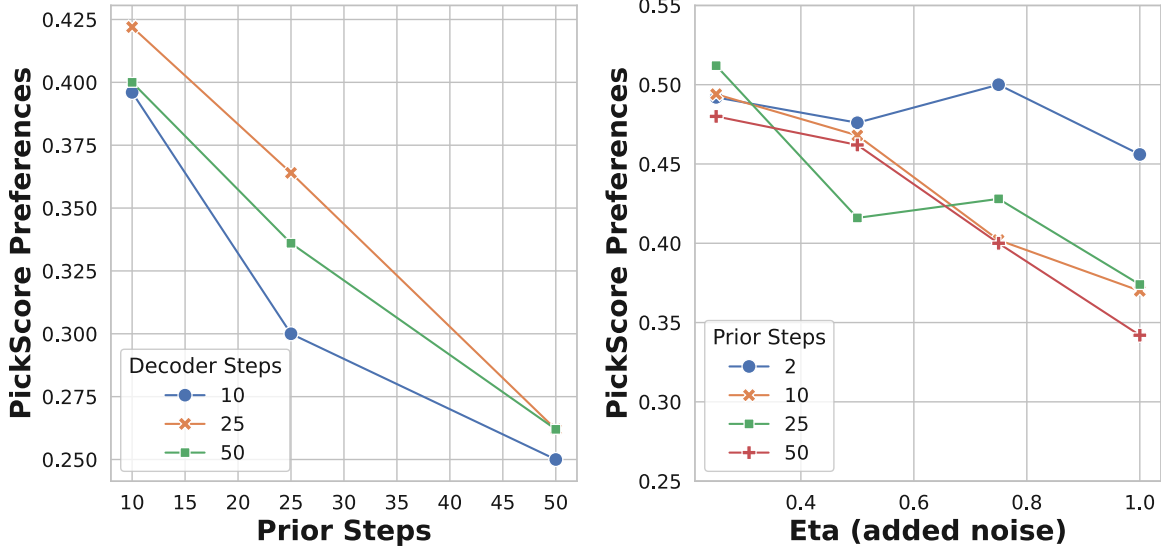
3.4.3 Qualitative Evaluations

In Figure 3.3, we display qualitative examples from various methods responding to complex prompts. *ECLIPSE* demonstrates superior performance in comparison to Stable Diffusion v1.4, Stable Diffusion v2.1, and Würstchen, while closely matching the quality of its big counterpart, Kandinsky v2.2. Interestingly, *ECLIPSE* trained on only 0.6 million images maintains the compositions with minor degradation in image quality. These observations align with our previously established quantitative results. Beyond numerical metrics,

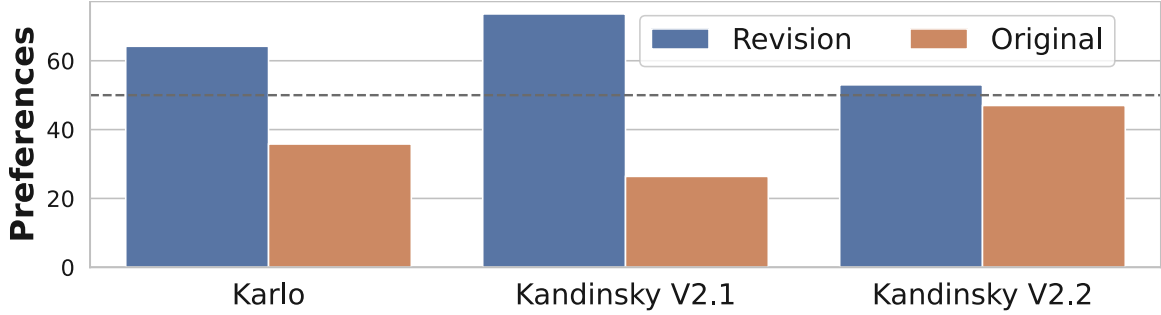
understanding human preferences is crucial. To this end, we selected 1500 unique validation prompts from T2I-CompBench and assessed PickScore preferences. The results, illustrated in Figure 3.4, reveal that *ECLIPSE* notably surpasses its baselines in respective restrictive settings with an average score of 71.6%. We can also observe that the best *ECLIPSE* variant (with Kandinsky decoder and trained on LAION-HighRes) consistently outperforms the other big SOTA models achieving an average performance of 63.36%. We observe that in terms of preferences, the original Kandinsky v2.2 diffusion prior (with a 1 billion parameter) trained on LAION-HighRes (175M) performs better than the *ECLIPSE* prior (having 33 million parameters). We hypothesize that this might be due to its use of a large-scale dataset that contains more aesthetically pleasing images. We provide a set of qualitative results in the appendix to show that *ECLIPSE* performs similarly well, if not better, *w.r.t.* semantic understanding of the text.

3.5 Analysis

analyzing the traditional diffusion priors. To further support our choice of using non-diffusion prior models, we analyze the existing diffusion prior formulation. We conducted two key empirical studies: 1) Evaluating the Impact of Prior Steps: We examined how the number of prior steps influences model performance. 2) Assessing the Influence of Added Noise ($\sigma_t\epsilon$): We focused on understanding how the introduction of noise affects human preferences. For these studies, we utilized PickScore preferences, and the outcomes, depicted in Figure 3.5, corroborate our hypothesis: both the prior steps and the addition of ($\sigma_t\epsilon$) detrimentally affect performance. Furthermore, as indicated in Table 3.3, diffusion prior surpasses the projection baseline if provided with more high-quality data. We attribute this enhanced performance to the incorporation of classifier-free guidance, which bolsters the model’s generalization capabilities to a certain extent. However, it is worth noting that both baselines are still outperformed by *ECLIPSE*. This observation underscores the



(a) Left: Performance comparison by varying the prior steps and decoder steps *w.r.t.* the fixed prior steps ($t = 2$). Right: Performance comparison by varying the mean η of the added scheduler noise ($\sigma_t \epsilon$) *w.r.t.* the noiseless predictions ($\eta = 0$). Both experiments are on the Kandinsky v2.1.



(b) Overall performance comparisons on various pre-trained unCLIP models before and after reducing the prior steps to two and η to 0.0.

Figure 3.5: Empirical analysis of the PickScore preferences of diffusion priors with respect to the various hyper-parameters.

effectiveness of our proposed methodology in comparison to traditional approaches in the realm of T2I.

Importance of data selection. In our previous analysis (Table 3.1 and 3.3), we demonstrated that *ECLIPSE* attains competitive performance on composition benchmarks regardless of dataset size. This achievement is largely due to the integration of the contrastive loss $\mathcal{L}_{CLS}$ (Eq.5.1). However, the final objective function also incorporates the $\mathcal{L}_{proj}$ (Eq.3.3), which is pivotal in estimating the vision latent distribution. This estimation is fundamentally

dependent on the training distribution (P_{XY}), leading the model to learn spurious correlations within P_{XY} . Consequently, the model’s image quality could directly correlate with the overall quality of images in the training set. To further substantiate this, we evaluated the preferences for *ECLIPSE* models trained on MSCOCO, CC3M, and CC12M, in comparison to among themselves and Karlo-v1-alpha. The outcomes, presented in Figure 3.6, reveal that the *ECLIPSE* model trained on CC12M outperforms those trained on other datasets, exhibiting performance on par with its big counterpart. *ECLIPSE* prior (w Karlo decoder) trained on the CC12M dataset performs comparably to Karlo-v1-alpha while *ECLIPSE* priors trained on other datasets struggle to do so. Furthermore, as illustrated in Figure 3.6, the *ECLIPSE* model trained on MSCOCO demonstrates a tendency to learn spurious correlations, such as associating the term “young tiger” with the person.

Training and Inference Efficiency. In this section, we assess the efficiency of various text-to-image (T2I) prior models, examining their resource utilization during training and inference. This includes an analysis of GPU hours, data requirements, and model parameters. A comparative analysis, as shown in Table 3.4, highlights the efficiency of diverse T2I priors, including stable diffusion. However, specific training details for several T2I priors like LAION, Kandinsky, and Karlo remain undisclosed, prompting us to draw comparisons with domain-specific priors known for their relatively streamlined training processes. These comparisons reveal that even specialized domain priors necessitate substantial resources, entailing millions of parameters and extensive GPU processing time. In contrast, *ECLIPSE* emerges as an efficient model, requiring merely 50 GPU hours to achieve state-of-the-art (SOTA) results. Moreover, Figure 3.7 compares the inference times of traditional diffusion priors against *ECLIPSE*. Whereas conventional models demand approximately 0.8 seconds for inference, *ECLIPSE* significantly reduces this to just 0.005 seconds, attributing to its lesser parameters and single-step estimations. This efficiency underscores a pivotal

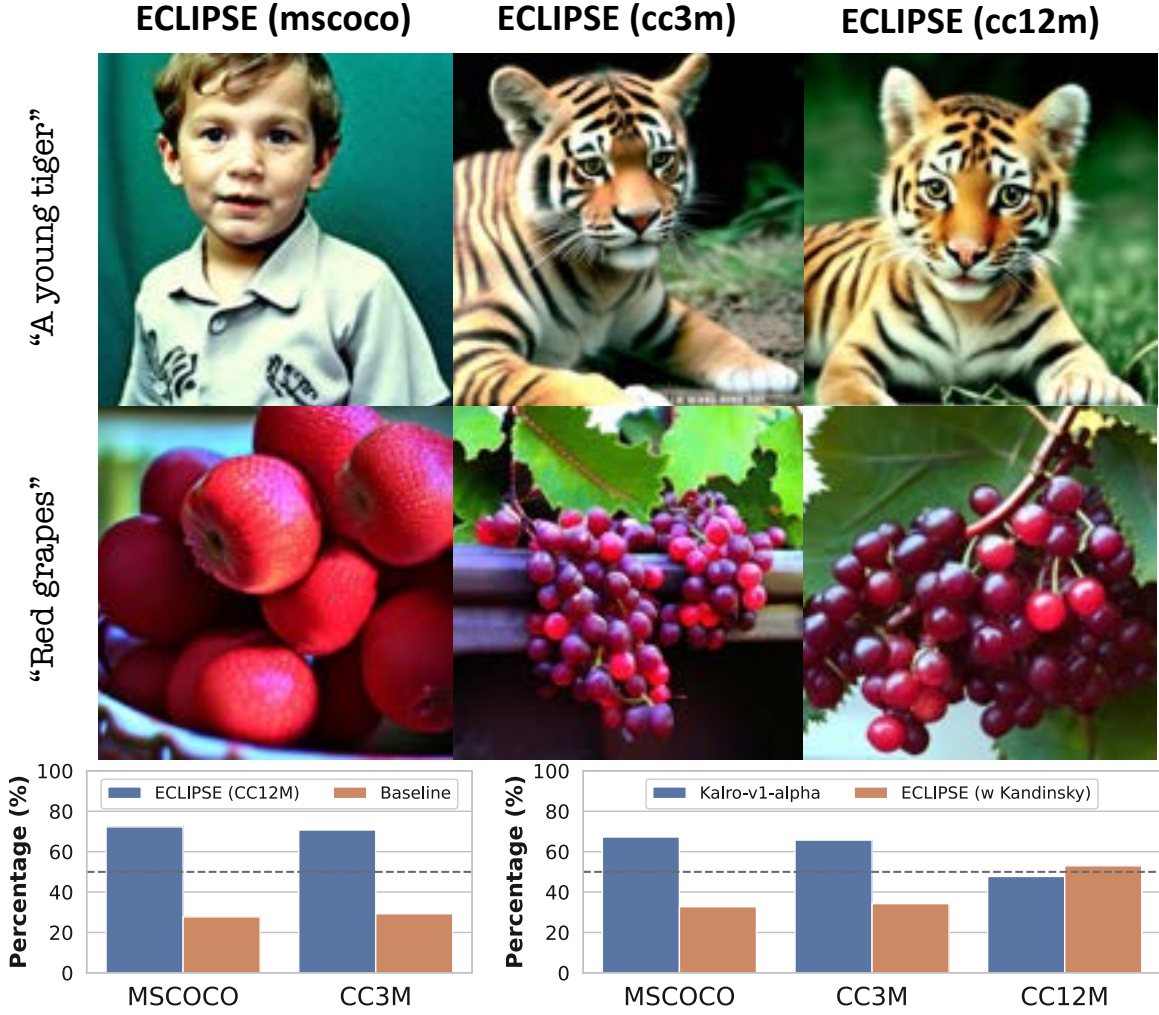


Figure 3.6: The top figure shows the qualitative examples of the biases learned by the T2I prior models. Bottom figures show the PickScore preferences of the *ECLIPSE* models trained on various datasets with respect to the other datasets (left) and Karlo (right).

insight: the process of text-to-image mapping does not necessitate the use of expansive models like Stable Diffusion. Instead, we demonstrate that T2I conversion can be executed more proficiently within the latent space, marking a significant stride towards enhancing model efficiency without compromising performance.

Hyper-parameter analysis. *ECLIPSE* only contains one important hyperparameter (λ) that controls the contrastive learning. As discussed in Section 3.3.3, a higher value of λ can make the prior model learn the different distributions that are highly aligned with text

Table 3.4: Training time comparisons of various prior models in terms of resource requirements after [1].

Methods	Compute	Parameters	Data Size
	100 GPU Hours (↓)	Millions (↓)	Millions (↓)
Stable Diffusion	150000	859.92	2000
Isolated Prior	1344	249.22	20
Vector Prior	1680	101.76	26
Texture Prior	576	249.22	10
Color Prior	3072	249.98	61
LAION Prior (T2I)	N/	1000	2000
Karlo Prior (T2I)	N/	1000	115
Kandinsky Prior (T2I)	N/	1000	117
<i>E LIPSE</i>	50	33~34	< 10

distributions. A lower value of λ may not benefit in terms of generalization to unseen prompts. Hence, we conducted a small study on the MSCOCO dataset. We train the *ECLIPSE* priors for Karlo decoder on 20,000 iterations with the OneCycle learning rate. Figure 3.8 illustrates the pickscore preferences on T2I-CompBench of various values of λ . It can be observed that higher values of λ lead to the same performance as the baseline. While lower values of λ outperform the baseline by significant margins. Additionally, Figure 3.9 shows one qualitative example across the range of λ . It can be seen that the generated image quality drops as λ increases. **Hence, the optimal range is: $\lambda \in [0.2, 0.4]$.**

***E LIPSE* Prior Model Scaling Behaviour.** To analyze the scaling behavior of different prior learning strategies to a certain extent, we increase the prior model size from 33M to 89M. Table 3.5 shows the results when small and large priors are trained on the same dataset (CC12M) with the Karlo image diffusion decoder. We train both versions of the prior models on 60,000 iterations (about 350 GPU hours) with exactly the same hyperparameters. First, we observe that *E LIPSE* prior improves the performance slightly. Second, the Projection baseline gets the same performance, which suggests that **data is the bottleneck**

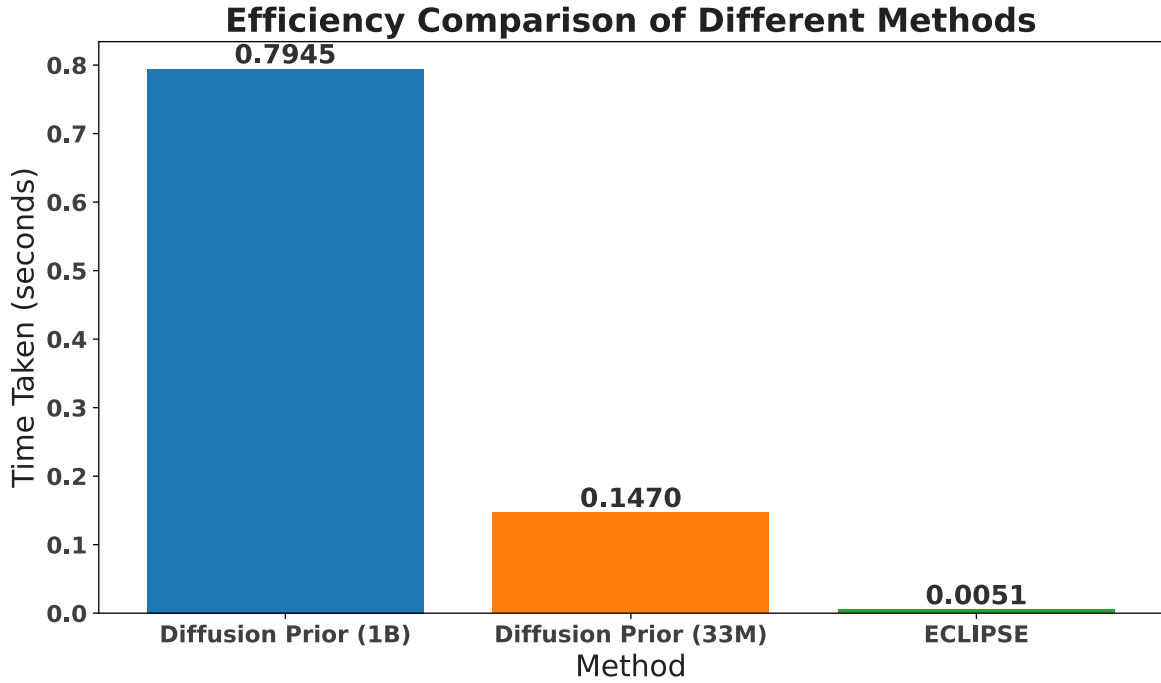


Figure 3.7: Inference time analysis of diffusion priors having 1B and 33M parameters vs. *ECLIPSE* prior.

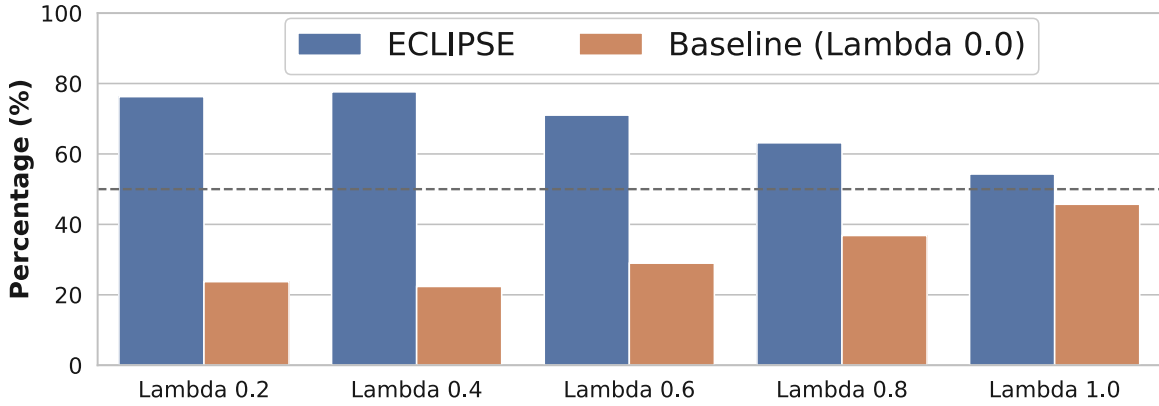


Figure 3.8: Hyperparameter (λ) ablation. This figure illustrates the PickScore preferences across the *ECLIPSE* priors trained with different values of λ w.r.t. the Projection baseline (with $\lambda = 0.0$).

Table 3.5: This table illustrates the scaling behavior of various T2I prior learning strategies. “Small” priors are 33 million in terms of parameters. And “Large” priors have 89 million parameters. All prior models are trained on the CC12M dataset with the Karlo diffusion image decoder.

Methods	ZS	T2I-CompBench			
	FID	Color ($\uparrow$)	Shape ($\uparrow$)	Texture ($\uparrow$)	Spatial ($\uparrow$)
33M Priors					
Projection	28.84	0.4659	0.4632	0.4995	0.1318
Diffusion-Baseline	26.13	0.5390	0.4919	0.5276	0.1426
<i>E LIPSE</i>	26.98	0.5660	0.5234	0.5941	0.1625
89M Priors					
Projection	28.81	0.4579	0.4625	0.4761	0.1343
Diffusion-Baseline	29.78	0.4988	0.4790	0.4604	0.1247
<i>E LIPSE</i>	25.77	0.5712	0.5358	0.6194	0.16665



Figure 3.9: Qualitative example for *ECLIPSE* priors (with Karlo decoder) trained with different values of hyperparameter (λ).

for the Projection prior. Third, interestingly Diffusion prior degrades the performance. Upon further inspection, we found that 60,000 iterations are insufficient for the Diffusion model to converge. Therefore, this verifies that **Diffusion-priors are resource-hungry**. Importantly, *ECLIPSE* priors easily converge irrespective of the data and number of parameters; suggesting that *ECLIPSE* do not depend upon the huge resource constraints.

aesthetics: Kandinsky v2.2 vs. *ECLIPSE*. As was observed in Figure 3.4, the Kandinsky v2.2 model outperforms the *ECLIPSE* prior when evaluated in terms of human preferences measured by Pickscore. We attribute this behavior to the differences in the aesthetic quality of the generated images. Therefore, we conduct additional actual human studies to analyze this behavior further. In total, we randomly selected 200 prompts from the MSCOCO validation set (instead of T2I-CompBench as reported in Figure 3.4) and asked the human evaluators to perform two studies:

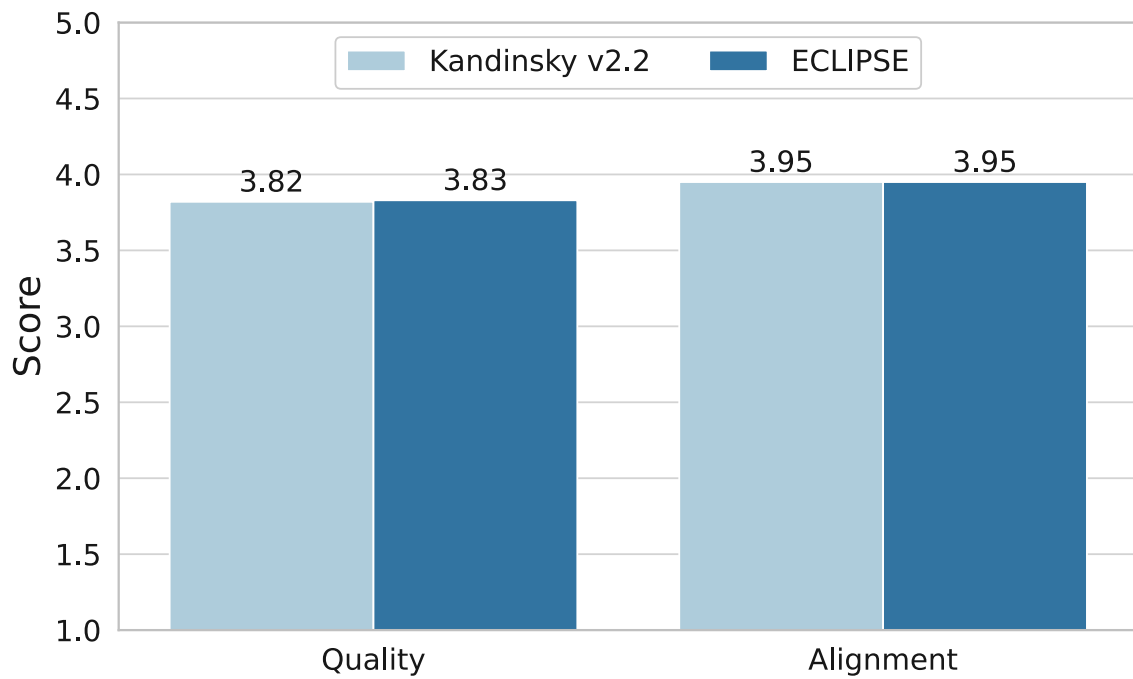


Figure 3.10: Human evaluations of the *ECLIPSE* vs. Kandinsky v2.2 generated images. It can be observed that both models are rated equally in terms of image quality and caption alignment.

- Rate each image in terms of quality and caption alignment between 1-5. Where 1 is the artificial-looking image and caption alignment is poor. While 5 represents a very high-quality image and is perfectly aligned with the captions.
- Image preferences in terms of aesthetics. We show images from both models and ask the evaluators to choose one which looks more aesthetically pleasing.

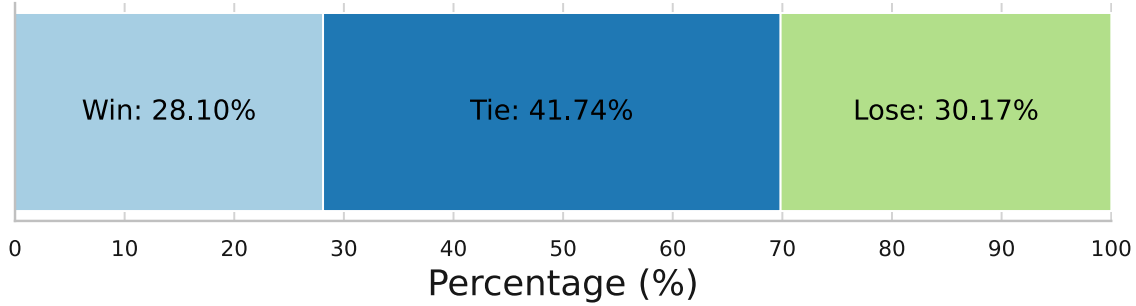


Figure 3.11: This figure illustrates the human preferences between *ECLIPSE* prior for Kandinsky model (trained on LAION-HighRes subset) vs. Original Kandinsky v2.2 model.

Interestingly, as shown in Figure 3.10, both models are rated equally when evaluated independently. Additionally, according to Figure 3.11, Kandinsky v2.2 is preferred slightly more than the *ECLIPSE* in terms of aesthetic quality. This finding suggests that smaller prior trained with *ECLIPSE* can perform equally (if not better) to those big prior models. Figure 3.12 shares three examples from the MSCOCO. Both models perform equally well but Kandinsky is more aesthetically pleasing. Figure 3.21 and 3.22 show the MTurk human evaluation instructions.

3.6 Diversity With Non-Diffusion Priors

One important aspect of the diffusion models is the diversity of the generated images. Therefore, diversity and caption alignment go hand-in-hand. We further analyze whether having the non-diffusion prior hurts diversity or not. We perform additional qualitative evaluations and given a prompt – we ask the human evaluators to select which of the two grids of six images are more diverse. This experiment is performed between *ECLIPSE* and Kandinsky v2.2. As shown in Figure 3.13, even if we use the non-diffusion prior model it does not hurt the diversity. Diffusion image decoder is the main reason that contributes to the diversity and having diffusion or non-diffusion prior does not contribute that significantly.



Figure 3.12: Qualitative examples comparing (in terms of aesthetics) *ECLIPSE* with Kandinsky v2.2.

Failure Cases. Figure 3.20 shows some examples where *ECLIPSE* model fails to follow the prompt precisely. It is still difficult for the prior to learn something very unconventional. The model fails at generating some composition prompts (first four images). It has been shown that vision-language models also suffer from such composition understanding (e.g.,

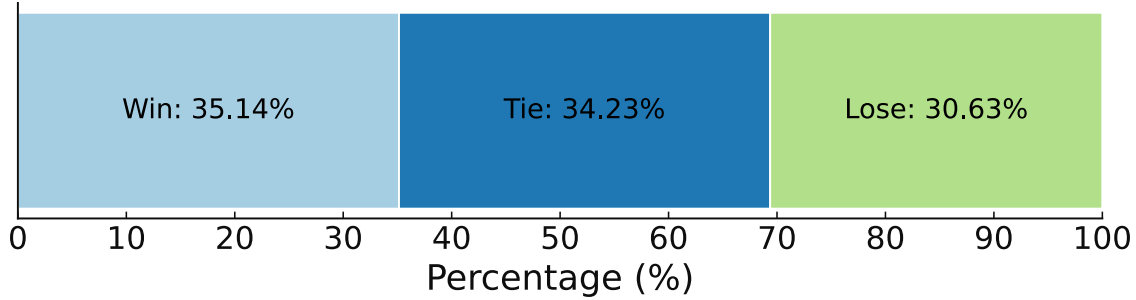


Figure 3.13: This figure illustrates the human preferences on the diversity of generated images between *ECLIPSE* prior with Kandinsky v2.2 diffusion image decoder vs. Kandinsky v2.2.

“grass in the mug” vs. “mug in the grass”). Therefore, improving the Vision-Language model can further improve the capabilities of unCLIP priors. Notably, *ECLIPSE* finds it difficult to generate artistic imaginary images (such as “nebula explosion that looks like corgi”). However, such corner cases can be only solved with more diverse high-quality datasets.

3.7 Conclusion

In this paper, we introduce a novel text-to-image prior learning strategy, named *ECLIPSE*, which leverages pre-trained vision-language models to provide additional supervision for training the prior model through contrastive learning. This approach significantly enhances the training efficiency of prior models in a parameter-efficient way. Through comprehensive quantitative and qualitative evaluations, we assessed *ECLIPSE* priors alongside various diffusion image decoders. The results indicate that *ECLIPSE* surpasses both the baseline projection models and traditional diffusion-prior models. Remarkably, *ECLIPSE* achieves competitive performance alongside larger, state-of-the-art T2I models. It demonstrates that priors can be trained with merely 3.3% of the parameters and 2.8% of image-text pairs typically required, without compromising performance. This advancement directly leads to at least 43% overall compression of the unCLIP models. Our findings show that pre-trained vision-language models can be utilized more effectively, suggesting a promising research



Figure 3.14: Qualitative comparisons between *ECLIPSE* and baseline priors (having 33 million parameters) trained on CC12M dataset with Karlo decoder. * prompt is: "The bold, striking contrast of the black and white photograph captured the sense of the moment, a timeless treasure memory."

direction where improving vision-language models may directly benefit T2I generation.

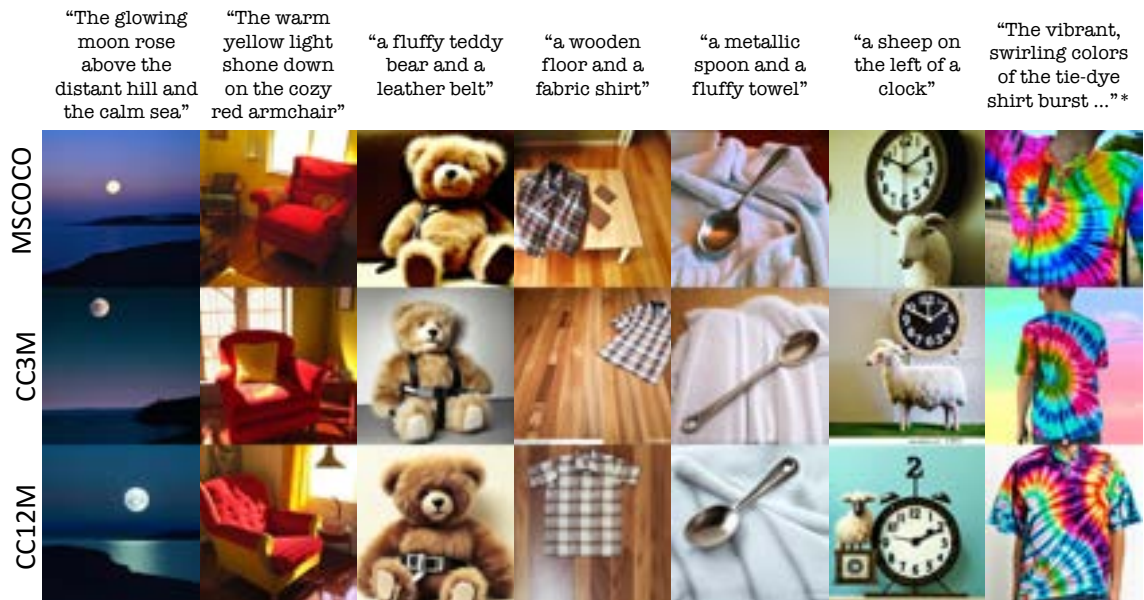


Figure 3.15: Qualitative comparisons of *ECLIPSE* priors with Karlo decoder trained on different datasets. * prompt is: "The vibrant, swirling colors of the tie-dye shirt burst with energy and personality, a unique expression of individuality and creativity."



Figure 3.16: Qualitative result of our text-to-image prior, *ECLIPSE* (with Karlo decoder), along with a comparison with SOTA T2I models. Our prior model reduces the prior parameter requirements (from 1 Billion → 33 Million) and data requirements (from 115 Million → 12 Million → 0.6 Million).



Figure 3.17: Qualitative comparisons between *ECLIPSE* and baseline priors (having 34 million parameters) trained on LAION-HighRes subset dataset with Kandinsky v2.2 diffusion image decoder.



Figure 3.18: Qualitative comparisons between *ECLIPSE* prior trained on MSCOCO and LAION datasets with Kandinsky v2.2 decoder.



Figure 3.19: More qualitative result of our text-to-image prior, *ECLIPSE* (with Kandinsky v2.2 decoder), along with a comparison with SOTA T2I models. Our prior model reduces the prior parameter requirements (from 1 Billion → 33 Million) and data requirements (from 177 Million → 5 Million → 0.6 Million).



Figure 3.20: Instances where *ECLIPSE* encounters the challenges in following the target text prompts.

Visual Question Answering

Consent Checked: yes

Context: The purpose of this research project is to study how accurately Artificial Intelligence (Machine Learning) models generate images from textual descriptions. We study a special type of models called text-to-image generators – users can enter a sentence to the model, and the model generates an image for this this sentence. The goal is to understand if the generated image is aligned with the input text.

Your task: In this HIT, we will show you an image. Your task is to answer several questions about this image about the objects present in the image and the quality of the image to measure the image-text alignment.

Solved Examples: In order to help you build an understanding of the task, here are a few examples of solved HITs: [\[Example 1\]](#) [\[Example 2\]](#) [\[Example 3\]](#)

Caption: a man in a chefs hat chopping food



(1) Rate the *quality* of the image.
(*1" being artificial (noisy, blurry) and "5" being natural (a real photograph))

☐ 1 ☐ 2 ☒ 3 ☐ 4 ☐ 5

(2) What is the *similarity* between the caption and image?
(Rate from "1" (least similar) to "5" (most similar))

☐ 1 ☐ 2 ☒ 3 ☐ 4 ☐ 5

Figure 3.21: An example of human annotation for determining the image quality and caption alignment.

AI generated image preferences

Consent Checked: yes

Context: The purpose of this research project is to study the quality of the AI generated images with respect to each other, given the input caption. In this study, we will provide a caption and two different images. The goal is to select one of these two as the preferred choice.

Your task: In this HIT, we will give you four tasks. For each task, you will get two images and one target prompt/caption. Your goal is to select the best aesthetically pleasing image.

Quick Guide:

1) If both images are equally good then you can select the **EQUAL** option.

Solved Examples: In order to help you build an understanding of the task, here are a few examples of solved HITs: [\[Example 1\]](#) [\[Example 2\]](#) [\[Example 3\]](#)

(1) Select the image with the *best aesthetics* that follow the caption: A bike parked on the side walk and a car on the street



Figure 3.22: An example of human annotation for determining the most aesthetic image.

Chapter 4

MULTI-CONCEPT PERSONALIZED TEXT-TO-IMAGE DIFFUSION MODELS BY LEVERAGING CLIP LATENT SPACE

4.1 Introduction

The field of text-to-image (T2I) diffusion models has recently witnessed remarkable advancements, achieving greater photorealism and enhanced adherence to textual prompts. This has catalyzed the emergence of diverse applications, notably subject-driven personalized T2I (P-T2I) models. In particular, this encompasses the intricate task of learning and reproducing novel visual concepts or subjects in varied contexts requiring high concept and compositional alignment. The complexity escalates further when multi-subject personalization is desired.

Early works employed concept-specific optimization strategies involving fine-tuning certain parameters within T2I diffusion models [72, 223, 129, 265, 73]. Although these methods achieve state-of-the-art (SOTA) performance, they struggle with generalization and are time-intensive. Contemporary research is pivoting towards fast personalization techniques. Within this paradigm, there are two types of popular approaches: 1) Methods that involve training hypernetworks and integrating new layers or parameters within pre-trained diffusion UNet models [282, 296, 261, 238, 224], and 2) MLLM-based learning of prior models that focuses on leveraging text-latent space of frozen diffusion UNet model [186, 257].

The hypernetwork-based strategy achieves single-concept customization but has not been extended to multi-concepts. Moreover, when combined with additional control (i.e. Canny edge map), they struggle to maintain the concept alignment ($\sim 30\%$ drop in performance;

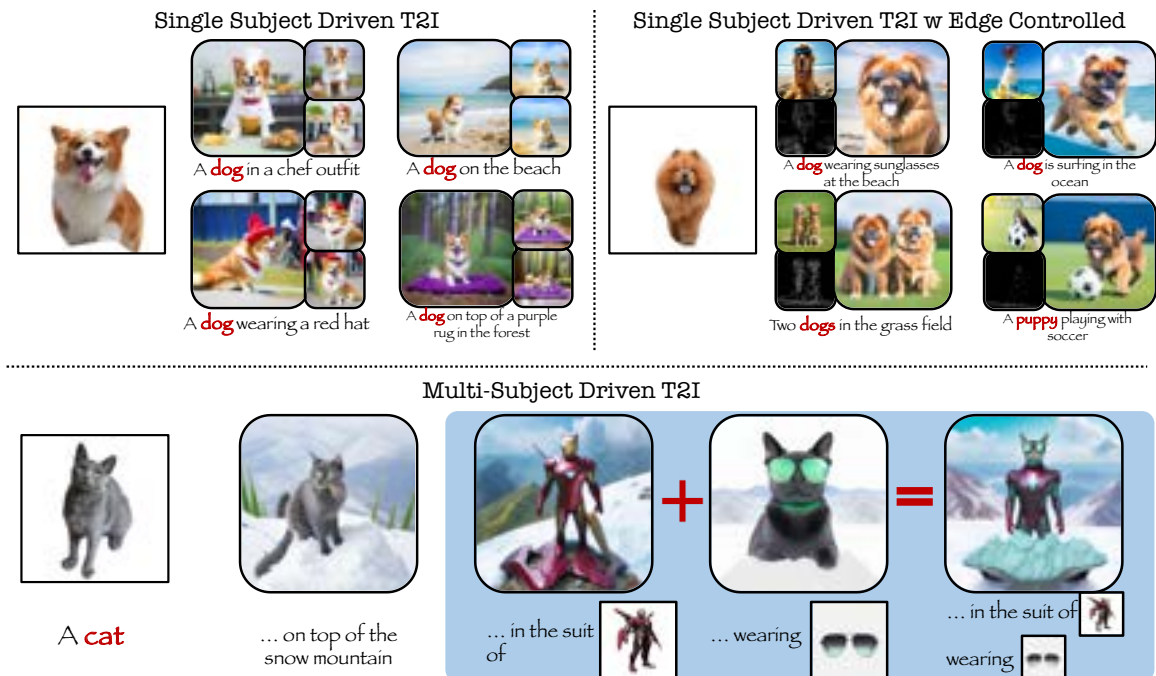


Figure 4.1: λ -ECLIPSE can estimate subject-specific latent image embeddings while maintaining the balance between concept and composition alignment in the CLIP latent space itself.

Section 7.5) and strongly favor the edge map. At the same time, MLLM-based approaches can perform fast multi-concept customization but require heavy computing resources. In Table 4.1, we provide the overview of various single and multi-concept customization methodologies in terms of total parameters, iterations, and GPU hours required to train the models. It can be observed that multi-concept customization methodologies further increase the resource requirements. For instance, Kosmos-G [186] consumes 18x more resources than IP-Adapter [296]. And Emu2 [257] requires training of 19x more parameters compared to Kosmos-G. Hence, despite MLLMs’ seemingly useful scenarios, it is not viable to blindly train them.

Upon further investigation, we find that most subject-driven T2I approaches build upon variants of the Latent Diffusion Model (LDM) [216], specifically Stable Diffusion models. These LDMs employ cross-attention layers to condition diffusion models with text

Table 4.1: Our method is the first to offer multi-concept-driven generation without depending on diffusion UNet models (except for inference). We provide the extended overview table in the appendix.

Method	Multi Concepts	Finetuning Free	Diffusion Free	Total opt. params	Training Steps	Dataset Size	GPU Hours
Textual Inversion [72]	✗	✗	✗	768	5000	-	1
DreamBooth [223]	✗	✗	✗	0.9B	800	-	0.2
ELITE [282]	✗		✗	77M	135K	125K	336
BLIP-Diffusion [137]	✗		✗	1.5B	500K	129M	2304
IP-Adapter [296]	✗		✗	22M	1M	-	672
Custom Diffusion [129]		✗	✗	57M	500	-	0.1
Subject-Diffusion [168]			✗	252M	300K	76M	-
Kosmos-G [186]			✗	1.9B	800K	200M	12300
Emu-2 [257]			✗	37B	70K	162M	-
λ - <i>ECLIPSE</i> (ours)				34M	100K	2M	74

embeddings, necessitating a mapping of target subject images to latent spaces compatible with the diffusion models at the prior training stage. This is also known as score distillation instruction tuning for MLLMs [186]. As there is no choice but to learn this text-to-image diffusion latent space, it involves backpropagation through the entire diffusion model often comprising over a billion parameters, contributing to the inefficiency of existing P-T2I methods. Therefore, in this work, we focus on answering one question: Do we really need diffusion models to train the customization models?

To answer this question and improve the resource efficiency for multi-concept image generation, we present λ -*ECLIPSE*¹, which leverages the properties of UnCLIP T2I models (e.g. DALL-E 2 [211] and Kandinsky v2.2 [214]) and performs P-T2I in the compressed

¹The designation λ -*ECLIPSE* is inspired by its conceptual alignment with the λ -calculus. In this context, the λ -*ECLIPSE* model functions similarly to a functional abstraction within λ -calculus, where it effectively binds variables. These variables, in our case, represent novel visual concepts that are integrated through composition prompts. Here, *ECLIPSE* indicates our architecture design choice.

latent space of frozen CLIP model. Specifically, unlike previous MLLM-based methodologies, λ -*ECLIPSE* aligns the output space of priors with CLIP vision space instead of the CLIP text space. λ -*ECLIPSE* takes multiple images and text instructions as input and estimates the respective vision embeddings, which can be used by the frozen diffusion UNet model from the UnCLIP stack to generate the resulting image. This elevates the training time dependencies on diffusion models for P-T2I; significantly contributing to the resource efficiency. Additionally, as diffusion or MLLM-based priors are still compute heavy due to a huge number of parameters and slower convergence, we build upon *ECLIPSE* [194] and SEED [78], which shows that text-to-image mapping can be optimized through contrastive pre-training. Here, we select *ECLIPSE* as preferred choice of prior architecture for best efficiency. At last, we propose a subject-driven instruction tuning task based on the image-text interleaved data as a pre-training strategy. This involves creating 2 million high-quality image-text pairs, where text embeddings linked to subjects are substituted with the respective image embeddings, which in return are considered as input to the λ -*ECLIPSE*. While λ -*ECLIPSE* can be plugged with these pre-trained methods, we explore the possibility of λ -*ECLIPSE* to incorporate Canny edge as an additional coarse-level control to synergetically work with subject-driven image generation tasks. Figure 4.1 provides the overview of λ -*ECLIPSE* capabilities.

Overall, we propose λ -*ECLIPSE* as an initial attempt to motivate future works on designing resource-efficient solutions for MLLM-based approaches. We summarize our main contributions as follows:

- We show that P-T2I mapping can be learned in CLIP latent space without training-time dependency on the diffusion models; enabling efficient training and fast multi-subject customization.
- Extensive experiments on Dreambench, Multibench, and ConceptBed reveal that

λ -*ECLIPSE* (34M parameter model) trained on a mere 74 GPU hours can achieve competitive performance to that of big counterparts (having 2B-37B parameters) and improve text-composition alignment.

- At last, λ -*ECLIPSE* inherits the smooth CLIP latent space. This allows us to perform the seamless transition between multi-concept generated images.

4.2 Related Works

We further expand our comparative analysis of P-T2I methods encompassing a total of 33 approaches including ours and parallel works. Table 4.2 summarizes them into four crucial aspects: 1) multi-subject support, 2) fine-tuning free, 3) base model types, and 4) the required number of input images. To summarize, λ -*ECLIPSE* is the only methodology built on top of the UnCLIP models while supporting multi-subject driven image generation with fine-tuning free, and only requires a single reference image input for the training. We detail the comparison below:

Multi-Subject Generation. Multi-subject generation enables users to integrate multiple personal subjects to generate an image that follows the text prompts and aligns with all the concept visuals. In total, 15 of the 33 methods offer this capability, while 6 methods support fast multi-subject personalization, others demand separate training for each subject to be learned and then an additional fusing step for combining the learned subjects is required (i.e. Zip-LoRA, Mix-of-Show). Among these methods, only a few can learn auxiliary guided information such as canny edge, depth maps, or open-pose and adapt style variation (i.e. Kosmos-G).

Fine-tuning Free (Fast Personalization). Many methods require test-time fine-tuning. Each varies on which part alteration occurs, as early models tend to modify the whole UNet.

Table 4.2: The detailed overview of subject-driven text-to-image generative methodologies.
 * represents the backbone base models listed are subject to potential updates or modifications.

Method	Multi-Subject	Finetuning-Free	Base-Model	# of Input Images
Re-Imagen [38]	✗		Imagen	Single
Textual Inversion [72]	✗	✗	SDv1.4	Multiple
DreamBooth [223]	✗	✗	SDv1.4	Multiple
Custom Diffusion [129]		✗	SDv1.4	Multiple
ELITE [282]	✗		SDv1.4	Single
E4T [73]	✗	✗	SD	Single
Cones [156]		✗	SDv1.4	Single
SVDiff [86]		✗	SD	Multiple
UMM-Diffusion [169]	✗		SDv1.5	Single
XTI [268]	✗	✗	SDv1.4	Multiple
Continual Diffusion [246]		✗	-	Multiple
InstantBooth [242]	✗		SDv1.4	Multiple
SuTi [37]	✗		Imagen	Multiple
Taming [104]	✗		Imagen	Single
BLIP-Diffusion [137]	✗		SDv1.5	Single
Cones 2 [160]		✗	SDv2.1	Single
DisenBooth [34]	✗	✗	SDv2.1	Single
FastComposer [288]			SDv1.5	Single
Perfusion [261]		✗	SDv1.5	Multiple
Mix-of-Show [84]		✗	Chilloutmix	Multiple
NeTI [2]	✗	✗	SDv1.4	Multiple
Break-A-Scene [8]		✗	SDv2.1	Single*
ViCo [265]	✗	✗	SDv1.4	Multiple
Domain-Agnostic [7]	✗	✗	-	Single
Subject-Diffusion [168]			SDv2	Single
HyperDreamBooth [224]	✗	✗	SDv1.5	Single
IP-Adapter [296]	✗		SDv1.5	Single
Kosmos-G [186]			SDv1.5	Single
Zip-LoRA [238]		✗	SDXL	Multiple
CatVersion [323]	✗	✗	SDv1.5	Multiple
SSR-Encoder [322]			SDv1.5	Single
Emu2 [257]			SDXL	Single
λ -E LIPSE (ours)			Kv2.2	Single

In contrast, recent models tune a small portion of the cross-attention layers or introduce additional layers performing as adapters. In our analysis of P-T2I methodologies, 14 out of

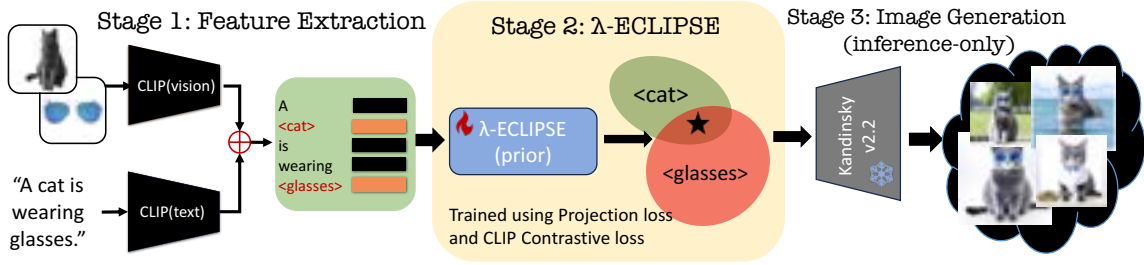


Figure 4.2: Three stages of the λ -ECLIPSE pipeline. 1) Create the image-text interleaved features using frozen CLIP. 2) Pre-train the λ -ECLIPSE (34M parameters) using Eq. 4.1, which ensures the mapping to the desired latent space given the image-text interleaved data. 3) During inference, the frozen Kandinsky v2.2 diffusion UNet model takes the output from the λ -ECLIPSE and generates the image.

33 methods employ a finetuning-free approach which enables fast personalization.

Diffusion Independent. A majority of the reviewed models utilize diffusion models, with Stable Diffusion being the predominant choice, spanning versions 1.4, 1.5, 2.1, and XL. Few adapt Imagen (SuTi, Taming) and Mix-of-show employs ChillOutMix as their pre-trained model, known for its adeptness at preserving realistic concepts like human faces. A unique outlier in this landscape is our λ -ECLIPSE, the only one that eschews the use of any diffusion prior model.

Easiness of Use. A more user-friendly model typically requires a single reference image per subject, as opposed to multiple images of the same subject. In our study, 19 methods offer P-T2I capabilities with just one input image. In contrast, others often require 4 to 5 images of the subject. Additionally, some methods necessitate storage space for learned concepts, ranging from a few hundred kilobytes (e.g., Perfusion, HyperDreamBooth) to several megabytes (e.g., Zip-LoRA). Our method stands out by eliminating the need for individual concept pre-learning or storing any artifacts for P-T2I utilization, offering a streamlined, efficient user experience.

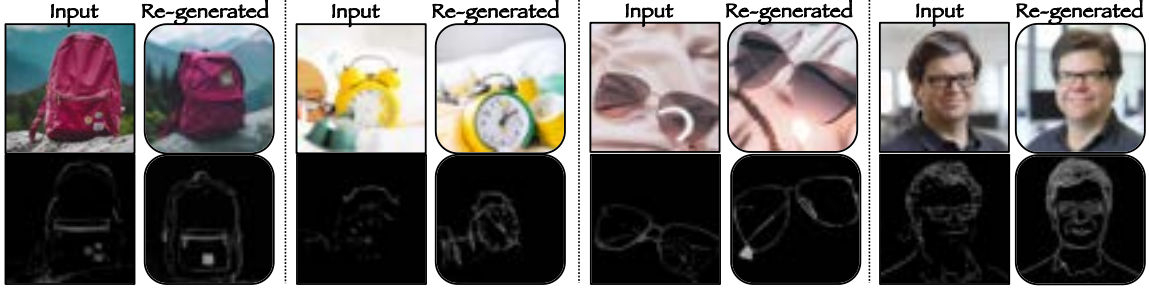


Figure 4.3: CLIP(vision) features capture the semantics and fine-grained visual details. Each input is given as input to the Kandinsky v2.2 and re-generated from the decoder. (Top: Real-images, Bottom: Canny edge)

4.3 Method

In this section, we introduce λ -*ECLIPSE*, our approach to multi-subject personalized text-to-image generation. Our method combines the contrastive text-to-image strategy from *ECLIPSE* with the novel image-text interleaved pretraining strategy, notably omitting the need for explicit diffusion modeling. Our approach mainly capitalizes on the efficient utilization of the CLIP latent space. Figure 4.2 outlines the end-to-end framework.

The primary objective of λ -*ECLIPSE* is to facilitate single and multi-subject P-T2I generation processes, accommodating edge maps as conditional guidance. Initially, we detail the problem formulation and elaborate on the UnCLIP stack design of the λ -*ECLIPSE*. Subsequently, we delve into the image-text interleaved training methodology. This fine-tuning process enables the λ -*ECLIPSE* to harness semantic correlations between CLIP image and text latent spaces while preserving subject-specific visual features.

4.3.1 Text-to-Image Prior Mapping

In the UnCLIP T2I models, the objective of the text-to-image prior model (f_θ) is to establish a proficient text-to-image embedding mapping. This model is designed to adeptly map textual representations to their corresponding visual embeddings, denoted as $(f_\theta : z_y \rightarrow z_x)$, where $z_{x/y}$ represent the embeddings for images and text, respectively. The visual

embedding predictions ($\hat{z}_x = f_\theta(z_y)$) are then effectively utilized by the diffusion image generators (h_ϕ), which are inherently conditioned on these vision embeddings ($h_\phi : z_x \rightarrow x$). In our experiments, we utilize the Kandinsky v2.2 diffusion UNet model as h_ϕ .

As shown in Figure 4.3, the CLIP vision encoder is very expressive and preserves the fine-grained details in z_x that are required to reconstruct the input image. CLIP image embedding itself achieves the high concept alignment score (DINO: 0.66) similar to the finetuning-based DreamBooth method [223].

Our goal is to accurately estimate the image embedding $\hat{z}_x$, incorporating the subject representations, thereby eliminating reliance on h_ϕ during training. Existing LDM-based P-T2I methods are limited by the LDM’s singular module approach ($h_\phi : z_y \rightarrow x$). Consequently, mastering the latent space of h_ϕ becomes essential for effective P-T2I for the baseline methodologies, which limits the previous methodologies.

We propose a new mapping function, f_θ , which processes text representations (z_y) alongside subject (x_k) specific visual representations (z_{x_k}), to derive an image embedding that encapsulates both text prompts and subject visuals ($\hat{z}_x$). The challenge lies in harmonizing z_{x_k} and z_y within $f_\theta : (z_y, z_{x_k}) \rightarrow \hat{z}_x$, ensuring alignment while preventing overemphasis on either aspect, as this could compromise composition alignment. To address this, we employ the contrastive pre-training strategy after [194]:

$$\mathcal{L}_{prior} = \mathop{\mathbb{E}}_{\epsilon \sim \mathcal{N}(0, I)} \left[\|z_x - f_\theta(\epsilon, z_y, z_{x_k})\|_2^2 \right] - \frac{\lambda}{N} \sum_{i=0}^N \log \frac{\exp(\langle \hat{z}_x^i, z_y^i \rangle / \tau)}{\sum_{j \in [N]} \exp(\langle \hat{z}_x^i, z_y^j \rangle / \tau)}. \quad (4.1)$$

Here, λ serves as the hyperparameter. i and j represent the index of the given input batch with the size N . $\langle \cdot \rangle$ represents the inner-product and τ is the temperature parameter. The first loss term (projection loss) measures the mean-squared error between the estimated and actual image embeddings, primarily ensuring concept alignment. However, our preliminary studies reveal that exclusive reliance on this term diminishes composition alignment. Therefore, we

stick with the contrastive loss component (the second loss term) to bolster compositional generalization, with λ balancing concept and composition alignment.

Additional Coarse-level Control-based T2I Prior Mapping. Acknowledging the limitations in existing methods, which necessitate learning the diffusion latent space even for additional control inputs, we endeavor to achieve a more nuanced balance between subject, text, and supplementary conditions. Consequently, we have augmented λ -*ECLIPSE* to accommodate an additional modality, a Canny edge map, providing more refined control over subject-driven image generation. This entails modifying the prior model to accept additional conditions ($f'_\theta : (z_y, z_{x_i}, z_c) \rightarrow \hat{z}_x$, here z_c is the additional modality embedding).

Additionally, during training, we drop z_c for 1% to improve the image generations without relying on the edge map. This enhances the stability and broadens the generalization capabilities of λ -*ECLIPSE*, yielding benefits even in the absence of these controls during inference. Our results demonstrate that λ -*ECLIPSE* learns a unified mapping function, accurately estimating target image representations through the effective integration of text, image, and edge maps – leading to learning coarse-level controls instead of hard constraints.

4.3.2 Image-text Interleaved Training

Our approach targets developing a versatile prior model capable of processing diverse inputs to estimate target visual outputs. Drawing from earlier methodologies, a straightforward solution involves concatenating different inputs, like combining text (“a dog wearing sunglasses”) with respective concept-specific images. Preliminary experiments indicated that this method does not effectively capture the intricate relationships between target text tokens (e.g. “dog”) and the corresponding concept images, especially when multiple concepts are present.

To address this, we adopt the interleaved pre-training strategy used in Kosmos-G, but



Figure 4.4: This figure illustrates a qualitative comparison of λ -ECLIPSE with contemporary approaches for single-subject T2I generations, utilizing concepts and prompts from the Dreambench dataset. For each method, concept, and prompt, we generate four images and select the one that most accurately represents the queried concept and composition.

with a notable modification to enhance resource efficiency. We incorporate pretrained frozen CLIP text and vision encoders for extracting modality-specific embeddings—separating text-only from subject-specific images. The key refinement in our process is the substitution of subject token-specific text embeddings with corresponding vision embeddings instead of introducing additional trainable tokens to handle the image embeddings via resampler [3]. First, we extract reference concept visual features ($z_{x_k} \in \mathcal{R}^{1 \times 1280}$) from the CLIP vision encoder. Similarly, we also extract the text prompt features ($z_y \in \mathcal{R}^{77 \times 1280}$) from the last

layer of the CLIP text encoder. Here, 1280 is the CLIP-specific feature dimension. At last, we replace the concept noun corresponding latent features from z_y with z_{x_k} ; resulting in image-text interleaved features while preserving the contextual information of the text features. This alteration allows us to bypass the need to train the big priors models (e.g. MLLMs), significantly improving the model’s proficiency in handling interleaved data.

For the generation of high-quality training datasets, we carefully selected 2 million high-quality images from the LAION dataset [236], each with a resolution of 1024x1024. Utilizing BLIP2, we generate captions for these images and employ SAM [119] for extracting noun or subject-specific segmentation masks. Given the CLIP model’s requirement for 224x224 resolution images, we avoid resizing the masks within their original resolutions. Instead, we opt for cropping the area of interest using Grounding DINO [152], followed by resizing the masked object while preserving its aspect ratio. This technique is crucial in retaining maximum visual information for each subject during the training phase. We provide more details about the filters used in the appendix.

4.3.3 Additional Concept-Specific Finetuning

Due to the nature of UnCLIP models, even if λ -ECLIPSE is very accurate, the diffusion UNet model (h_ϕ) may not be effective in generating very unique visual representations. In Figure 4.3, we can observe that generated images do not always precisely follow the reference image (e.g. hair style of the person). However, such behavior is common across the fast P-T2I methods and they lack in terms of maintaining performance compared to the finetuning-based methods (as outlined in Table 4.5). Therefore, we extend the λ -ECLIPSE and perform concept-specific finetuning.

Compared to the traditional finetuning methodologies (e.g. DreamBooth), λ -ECLIPSE provides very unique advantages. As λ -ECLIPSE prior model (f_θ) is pre-trained for personalization, there is no need for further finetuning the f_θ and we need to only finetune diffusion

UNet model. Importantly, the fine-tuning of the h_ϕ does not depend on the text embeddings (z_y). Hence, this leads to stable fine-tuning of the h_ϕ ; unlike DreamBooth on stable diffusion that observes catastrophic forgetting. The new fine-tuning objective is:

$$\mathcal{L}_{decoder} = \mathop{\mathbb{E}}_{\substack{\epsilon \sim N(0, I) \\ t \sim [0, T], (z_x)}} \left[\|\epsilon - h_\phi(x^{(t)}, t, z_x)\|_2^2 \right]. \quad (4.2)$$

Here, z_x is the visual feature of the reference concept image x . Notably, we do not need to use regularization from the DreamBooth as text alignment is already ensured during the pretraining stage of λ -*ECLIPSE*. Moreover, this finetuning can be performed across the set of given visual concepts altogether in a single model without degrading performance.

Dataset Creation. In constructing the dataset, we adhered to the data processing pipeline of Subject Diffusion [168]. We utilized the LAION-5B High-Res dataset, requiring a minimum image size of 1024x1024 resolution. Original captions were replaced with new captions generated by BLIP-2 (flan-t5-xl)². Subjects were extracted using Spacy³. For each subject, we identified bounding boxes employing Grounding DINO [152], setting both box-threshold and text-threshold values to 0.2. We retained images with 1 to 8 detected bounding boxes, discarding the rest. Additionally, captions with multiple instances of identical objects were filtered, allowing a maximum of 6 identical objects. Following bounding box detection, individual subject masks were isolated using Segment-Anything (SAM) [119]. To enhance the efficiency of the training process, we pre-processed the dataset by pre-extracting features from CLIP vision and text encoders. During this phase, images predominantly featuring a background (white portion) exceeding 50% of the total area were excluded. We preserved bounding boxes with an area ranging from 0.08 to 0.7 of the total image area and logit scores of at least 0.3. Masks constituting less than 40% of the bounding

²<https://huggingface.co/Salesforce/blip2-flan-t5-xl>

³<https://spacy.io>

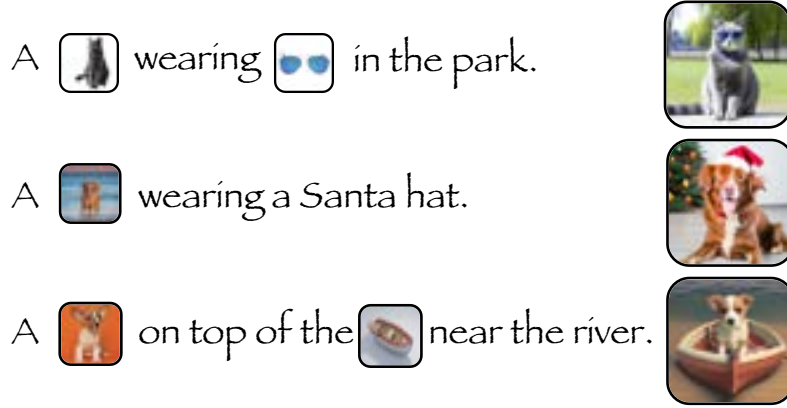


Figure 4.5: Examples of image-text interleaved training data. The left column shows the input of the prior model and the right images shows the ground truth corresponding images. Note: these examples are generated from λ -ECLIPSE for better interpretability.

box area were discarded. For the selection of subjects in images, we constrained the range to 1-4 subjects per image, excluding those with more than 4 subjects. At last, the interleaved image-text examples with respective ground truth images are shown in Figure 4.5.

Dataset Statistics. In the final analysis, our dataset comprised a total of 1,990,123 images. The distribution of subjects per image exhibited a range from 1 to 4, with the following breakdown: 1,479,785 images featuring one subject, 432,831 images with two subjects, 65,597 images containing three subjects, and 11,910 images showcasing four subjects. The overall count of unique subjects acquired from this dataset amounted to 30,358. We partitioned our dataset into an 80:20 split between training and validation, reserving the remaining 1.6 million images for training and the rest for validation.

In summary, the prior model, trained with our image-text interleaved data and supplementary condition, presents an efficient pathway for resource-efficient multi-subject-driven image generations.



Figure 4.6: Qualitative results categorized by generative capabilities.

Table 4.3: Example of prompt templates used for Multibench dataset. Subjects presented in Table 4.4 are placed in {}.

Two subjects	Three subjects
{} in the {}	{} with a {} and {}
{} wearing a {}	{} is playing with {} in {}
{} chasing a {}	{} with {} in front of {}
{} looking at a {}	{} with a {} and a view of the {}
{} is sitting on a {}	{} with a {} and {} in the background
{} standing on a {}	
{} and {} playing in the garden	
{} and {} on top of the mountain	
{} and {} in the jungle	
{} and {} in the snow	
{} and {} on the beach	
{} and {} on a cobblestone street	
{} and {} standing next to each other	

4.4 Proposed Multibench Benchmark Dataset

We provide additional qualitative results in Figure 4.6. For the multi-subject image benchmark, our dataset comprises 2,308 unique prompts, segmented into 904 two-subject and 1,476 three-subject prompts. This dataset integrates subjects from the original Dream-Bench dataset, featuring 30 distinct concepts. We expanded the dataset by incorporating additional concepts vital for two and three subject-specific prompts, such as various parks, hats, glasses, and more. Prompt templates and the count of unique subject categories featured in prompts are detailed in Tables 4.3 and 4.4, respectively. Overall, the dataset includes 217 two-subject compositions and 405 three-subject compositions, enriching the

Table 4.4: Number of occurrences of unique subject categories. The left side of the table are subjects used for two subjects prompts, and the right side of the table are subjects used for three subjects prompts.

Two subjects				Three subjects			
dog	76	boat	5	dog	81	rainbow	35
cat	76	park	4	stuffed animal	105	ruins	35
bird	76	ruins	9	toy	105	tower	35
horse	73	castle	5	cat	81	horse	81
guinea pig	73	desert	4	desert	60	bird	81
glasses	5	rainbow	5	hill	60	guinea pig	81
hat	5	candle	5	castle	45	guitar	25
tower	10	backpack	3	backpack	65	french horn	25
				can	130	vase	25
				candle	65	robot	25
				church	35		

benchmark’s diversity and comprehensiveness.

4.5 Experiments

In this section, we first introduce the experimental and evaluation setups. Later, we delve into the qualitative and quantitative results.

Training and inference details. We initialize our model, λ -*ECLIPSE*, equipped with 34M parameters. We train our model on an image-text interleaved dataset of 2M instances, partitioned into 1.6M for training and 0.4M for validation. The model is specifically tuned for the Kandinsky v2.2 diffusion image decoder. Therefore, we use pre-trained OpenCLIP-

Table 4.5: Quantitative comparisons of different methodologies on Dreambench. The bold and underline represent the metric-specific first and second-ranked methods, respectively. * represents that we re-benchmark the performance from open-source weights.

Method	Base Model	Params	GPU Hours	DINO (↑)	CLIP-I (↑)	CLIP-T (↑)	
Finetuning	Textual Inversion	SDv1.5	768	1	0.569	0.780	0.255
	DreamBooth	SDv1.5	0.9B	0.2	0.668	<u>0.803</u>	0.305
	Custom Diffusion	SDv1.5	57M	0.2	0.643	0.790	0.305
	BLIP-Diffusion	SDv1.5	0.9B	0.1	<u>0.670</u>	0.805	0.302
	λ - E LIPSE*	Kv2.2	0.9B	0.2	0.682	0.796	<u>0.304</u>
Finetuning-free	Re-Imagen	Imagen	-	-	0.600	0.740	0.270
	ELITE	SDv1.4	77M	336	0.621	0.771	0.293
	Subject-Diffusion	SDv1.5	252M	-	0.711	0.787	0.293
	BLIP-Diffusion*	SDv1.5	1.5B	2304	0.603	0.793	0.291
	IP- dapter*	SDv1.5	22M	672	<u>0.629</u>	0.827	0.264
	IP- dapter*	SDXL	22M	672	0.613	<u>0.810</u>	0.292
	Kosmos-G*	SDv1.5	1.9B	12300	0.618	0.822	0.250
	Emu2*	SDXL	37B	-	0.563	0.765	<u>0.273</u>
	λ - E LIPSE*	Kv2.2	34M	74	<u>0.613</u>	<u>0.783</u>	0.307

ViT-G/14⁴ as the text and vision encoders, ensuring alignment with Kandinsky v2.2 image embeddings. Training is executed on 2 x A100 GPUs, leveraging a per-GPU batch size of 512 and a peak learning rate of 0.00005, across approximately 100,000 iterations, summing up to 74 GPU hours. During inference, the model employs 50 DDIM steps and 7.5 classifier-free guidance for the Kandinsky v2.2 diffusion image generator. Adhering to baseline methodologies, we perform the P-T2I following the baseline papers’ protocols. For λ -ECLIPSE, target subject pixel regions in reference images are segmented before embedding extraction via the CLIP(vision) encoder. We drop the Canny edge map during inference unless specified explicitly. Unless specified all results (quantitative and qualitative)

⁴<https://huggingface.co/laion/CLIP-ViT-g-14-laion2B-s12B-b42K>

are without concept-specific additional fine-tuning.

The λ -*ECLIPSE* transformer prior architecture is significantly more compact compared to other Text-to-Image (T2I) methodologies. Our model employs a configuration of 16 Attention Heads, each with a dimension size of 32, alongside a total of 10 layers. Additionally, the embedding dimension size for our model is set at 1280, supplemented by 4 auxiliary embeddings (including, one for canny edge map). As λ -*ECLIPSE* is not a diffusion prior model, we do not keep time embedding layers. Overall, the λ -*ECLIPSE* model comprises approximately 34 million parameters, establishing it as a streamlined yet effective solution for Personalized-T2I. Notably, the standard UnCLIP T2I priors contain 1 billion parameters.

Evaluation setup. We primarily utilize Dreambench (encompassing 30 unique concepts with 25 prompts per concept) for qualitative and quantitative evaluations using DINO and CLIP-based metrics [223]. Due to their limitations, we extend our evaluations on the ConceptBed [191] benchmark (covering 80 diverse imagenet concepts and a total of 33K composite prompts), where we report performance on concept replication, concept, and composition alignment using the Concept Confidence Deviation (*CCD*) metric [191]. We extend Dreambench for multi-subject customization and present the Multibench dataset. Multibench contains about 24 unique concepts and 15 diverse prompts that result in 904 two-subject specific prompts and 1476 three-subject specific prompts. We provide further details about the Multibench in the appendix.

4.5.1 Main Results & analysis

Quantitative comparison. The quantitative assessments detailed in Table 4.5 and Table 4.6 focus on the single-concept T2I task, while Table 4.7 shows the results on multi-concept-driven image generation. For Dreambench and Multibench, we generate and evaluate four images per prompt, reporting average performance on three metrics (DINO, CLIP-I,

Table 4.6: Quantitative comparisons of different methodologies on ConceptBed. We present results on *CCD* ($\downarrow$) across three evaluation categories. The bold and underline represent the metric-specific first and second-ranked methods, respectively. * represents our benchmarking.

Method	Base Model	Concept Replication ($\downarrow$)	Concept Alignment ($\downarrow$)	Composition Alignment ($\downarrow$)	Average ($\downarrow$)
Textual Inversion	SDv1.4	0.0662	0.1163	0.1436	0.1087
Dreambooth	SDv1.4	<u>0.0880</u>	<u>0.3551</u>	<u>0.0360</u>	<u>0.1597</u>
Custom Diffusion	SDv1.4	0.2309	0.4882	0.0204	0.2465
ELITE*	SDv1.4	<u>0.3195</u>	0.4666	0.1832	0.3231
BLIP-Diffusion*	SDv1.5	0.3510	0.3245	0.1589	0.2781
IP- adapter*	SDXL	0.3665	<u>0.3571</u>	<u>0.0641</u>	<u>0.2626</u>
λ -ECLIPSE*	Kv2.2	0.2853	0.3619	-0.0200	0.2091

Table 4.7: Quantitative comparisons of different methodologies on Multibench. The bold and underline represent the metric-specific first and second-ranked methods on each metric, respectively.

	Two Subjects			Three Subjects	
	Kosmos-G	Emu2	λ -ECLIPSE	Emu2	λ -ECLIPSE
DINO ($\uparrow$)	0.4549	0.4451	<u>0.4478</u>	0.3168	0.3420
CLIP-I ($\uparrow$)	0.7759	0.7397	<u>0.7409</u>	0.6231	0.6463
CLIP-T ($\uparrow$)	0.2493	<u>0.2673</u>	0.3327	0.2819	0.3469

and CLIP-T). In the case of ConceptBed, we process each of the 33K prompts to generate a single concept image. The results, as depicted in these tables, highlight λ -ECLIPSE’s superior performance in composition alignment while maintaining competitive concept alignment. Analysis on ConceptBed (Table 4.6) indicates that λ -ECLIPSE exhibits a notable proficiency in concept replication, albeit with a marginal trade-off in concept alignment for enhanced composition fidelity. Comparatively, all baselines prioritize concept alignment,

Table 4.8: Quantitative results of Canny edge controlled P-T2I of different methodologies on Dreambench. The bold and underline represent the metric-specific first and second-ranked methods, respectively. Red highlighted numbers indicate the relative percentage drop for concept alignment compared to Table 4.5.

Method	DINO ($\uparrow$)	CLIP-I ($\uparrow$)	CLIP-T ($\uparrow$)
BLIP-Diffusion*	0.4234 _{29.7%}	0.7119	<u>0.3152</u>
IP- dapter*	<u>0.4281</u> _{31.9%}	<u>0.7315</u>	0.3034
λ-ECLIPSE*	0.5173 _{14.3%}	0.7437	0.3158

often at the expense of composition alignment. While λ -ECLIPSE improves the CLIP-T while preserving the DINO; achieved with significantly fewer resources. Notably, Multi-bench results (Table 4.7) indicate that λ -ECLIPSE significantly outperforms the Kosmos-G (2B params) and Emu2 (37B params) in terms of CLIP-T while maintaining the DINO performance. Therefore, we can conclude that λ -ECLIPSE is the most resource-efficient compared, especially when compared to the MLLM-based methods.

Qualitative comparisons. In Figure 4.4, we present a range of single subject-specific images generated by various methodologies including BLIP-Diffusion, IP-Adapter, Kosmos-G, Emu2, and λ -ECLIPSE. λ -ECLIPSE demonstrates exemplary proficiency in composition while ensuring concept alignment. In contrast, the baselines often overemphasize reference images or exhibit concept dilution, leading to higher concept alignment but compromised composition. Interestingly, we find that Emu2 can capture the single-subjects but it fails to reproduce them with complex text compositions (as shown in Figure 4.4). Similarly, Figure 4.7a exhibits λ -ECLIPSE’s multi-concept generation prowess, in comparison to ZipLoRA (fine-tuning-based approach) along with Kosmos-G and Emu2 (Multimodal LLM-based approaches), underscoring its capability to rival compute-intensive methods. We discuss additional examples and limitations in the appendix. That said, even though λ -ECLIPSE improves the performance over the baselines, this is still not enough and it

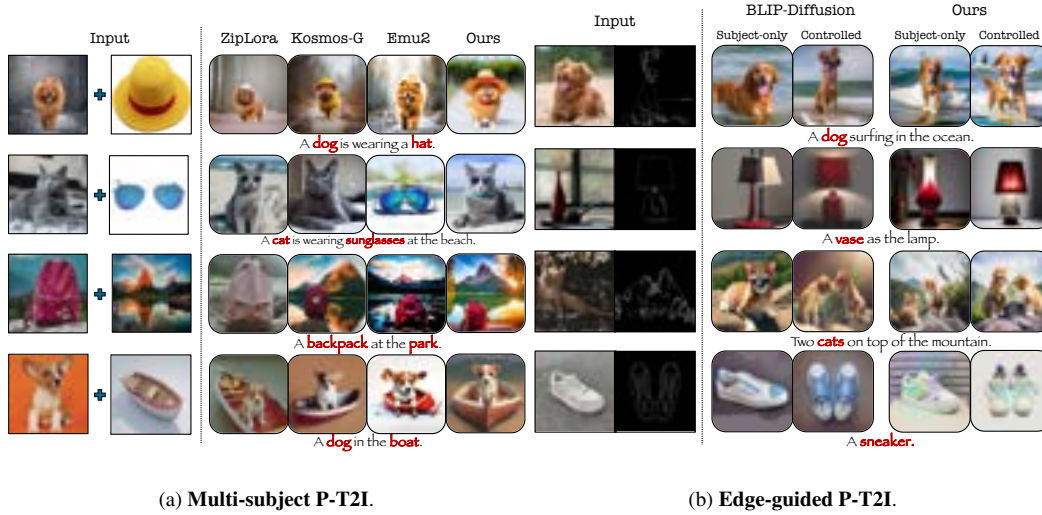


Figure 4.7: Qualitative comparison between λ -ECLIPSE and other baselines.

signifies the challenges associated with fast multi-concept personalization.

Canny edge controlled image generation. As shown in Figure 4.7b, the baseline (BLIP-Diffusion) adheres strictly to the imposed edge maps, often at the cost of concept retention (rows 1, 3, and 4). This leads to a large number of unwanted artifacts in the generated images. To further ground this behavior, we first generated images using Stable Diffusion v1.5 for Dreambench prompts without customization then we extracted the Canny edge map and used this edge map to control the subject-driven image generations. At last, we report the performance in Table 4.8. It can be observed that both baselines IP-Adapter and BLIP-Diffusion drop the DINO score by 30%, which follows the qualitative results. While λ -ECLIPSE do not follow the Canny edge precisely but preserves the concept alignment and improves the performance relatively by 21%.

4.5.2 λ -ECLIPSE with Finetuning

uxiliary finetuning. We further perform concept-specific fine-tuning (as described in Section 4.3.3). After finetuning, as shown in Table 4.5, λ -ECLIPSE outperforms the

Table 4.9: Ablation studies w.r.t. to the key components of λ -*ECLIPSE* design. We report the concept and composition alignment for single-subject T2I using *CCD* ($\downarrow$) on the ConceptBed benchmark.

Method	Concept	Composition
	lignment ($\downarrow$)	lignment ($\downarrow$)
Projection loss (i.e. $\lambda=0.0$)	0.394	0.008
w/ contrastive loss ($\lambda=0.5$)	0.435	-0.043
w/ contrastive loss ($\lambda=0.2$)	0.402	-0.026
w/ edge conditions ($\lambda=0.2$)	0.362	-0.020

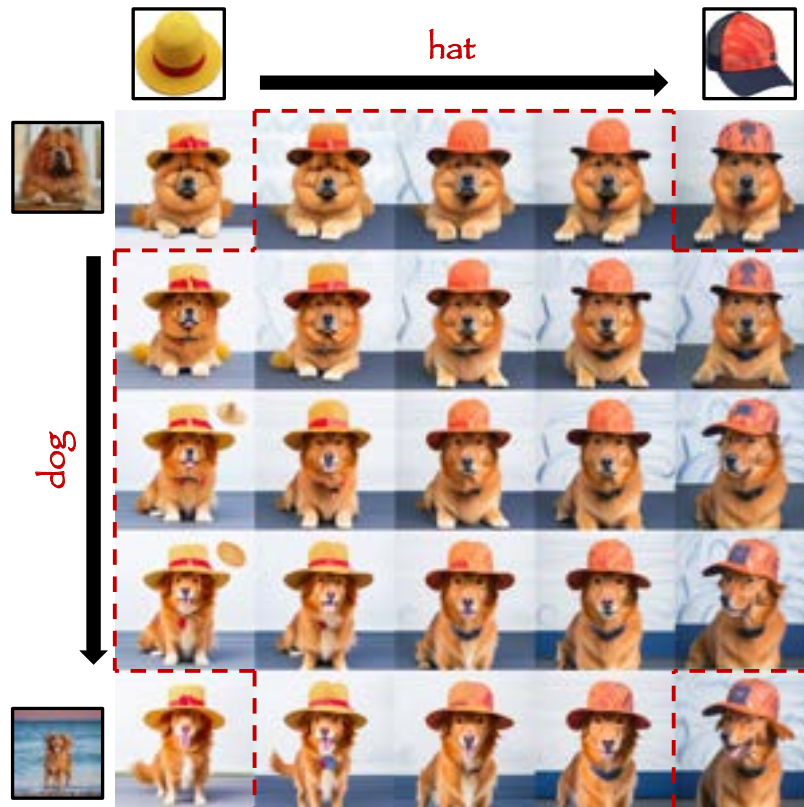


Figure 4.8: Interpolation between four concepts. Here, we estimate the image embedding using λ -*ECLIPSE* corresponding to each corner and then interpolate from top to bottom and left to right. At last, we use the Kandinsky v2.2 diffusion UNet model to generate the images with fixed random seeds from these sets of image embeddings.

DreamBooth and BLIP-Diffusion in terms of concept alignment (DINO) while maintaining the performance on composition alignment (CLIP-T). Our findings, illustrated in Figure 4.9, reveal that λ -*ECLIPSE* and DreamBooth exhibit improved performance with incremental fine-tuning steps. Notably, the DINO score improved from 0.61 to 0.68 with few optimization steps and outperforms the baselines (see Table 4.5). A detailed analysis indicates that while DreamBooth’s DINO score improves, its CLIP-T performance diminishes, hinting at concept overfitting. Conversely, λ -*ECLIPSE* consistently improves in DINO scoring without adversely impacting the CLIP-T performance, underscoring the efficacy of our image-text interleaved training approach at the prior stage. Qualitative comparisons, as shown in Figure 4.10, further highlight the benefits of fine-tuning λ -*ECLIPSE* with minimal steps.

Experimental Setup. Given that λ -*ECLIPSE* effectively trains the T2I prior, capturing concept-specific features to a significant degree, we opted not to further optimize this component. Our focus shifted to exclusively fine-tuning the diffusion UNet model (h_ϕ), employing the AdamW optimizer at a learning rate of 1e-5, without warm-up steps. For the DreamBooth application within the Stable Diffusion v1.5 model, we selected a learning rate of 5e-6, maintaining consistency in other hyperparameters. To simplify, we excluded the use of a prior preservation regularizer and conducted training on the Dreambench platform using a single RTX A6000 GPU.

Advantages of fine-tuning λ -*ECLIPSE*. The fine-tuning of λ -*ECLIPSE*, in comparison to the baselines, reveals two key benefits: 1) Achieving state-of-the-art (SOTA) performance within a few finetuning steps. 2) Unlike the Stable Diffusion model, which exhibits catastrophic forgetting of nearby concepts post-DreamBooth fine-tuning, λ -*ECLIPSE* maintains previous knowledge. This suggests that a single model is sufficient to effectively fine-tune across multiple concepts together.

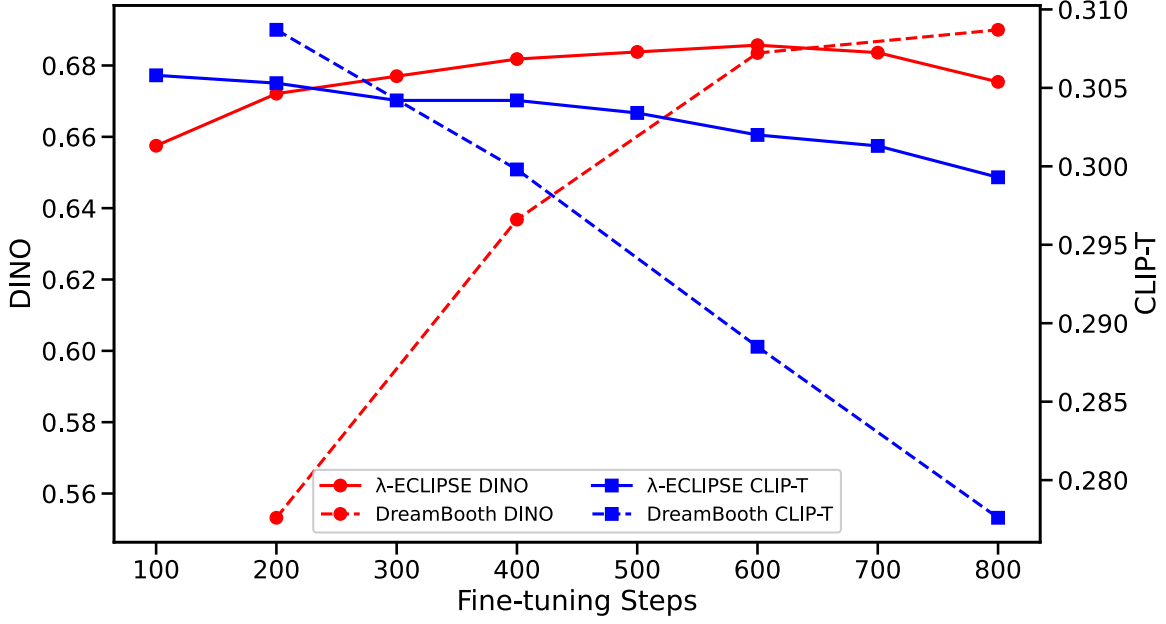


Figure 4.9: DreamBooth (Stable Diffusion v1.5) vs. λ -ECLIPSE (with fine-tuning) w.r.t. DINO and CLIP-T metrics on Dreambench.

This analysis underscores the strategic advantages and enhanced efficiency of fine-tuning λ -ECLIPSE for personalized applications in complex visual data processing.

4.5.3 Ablations

Loss weight study. We extend our study to evaluate the individual contributions of different components in λ -ECLIPSE. Initially, the model’s performance with solely the projection loss (referenced in Eq.4.1) is assessed. Subsequent experiments involve training λ -ECLIPSE variants with varying hyperparameters for the contrastive loss, specifically λ values of 0.2 and 0.5. A comparative analysis of these baselines is conducted against the fully equipped λ -ECLIPSE model, which incorporates $\mathcal{L}_{prior}$ (Eq.4.1) with $\lambda = 0.2$ and utilizes Canny edge maps during training. Relying solely on projection loss results in high concept alignment but compromises compositions (Table 4.9). The contrastive loss variant with $\lambda = 0.5$ enhances composition alignment at the expense of concept alignment, whereas $\lambda = 0.2$ achieves a more balanced performance. Crucially, the integration of Canny

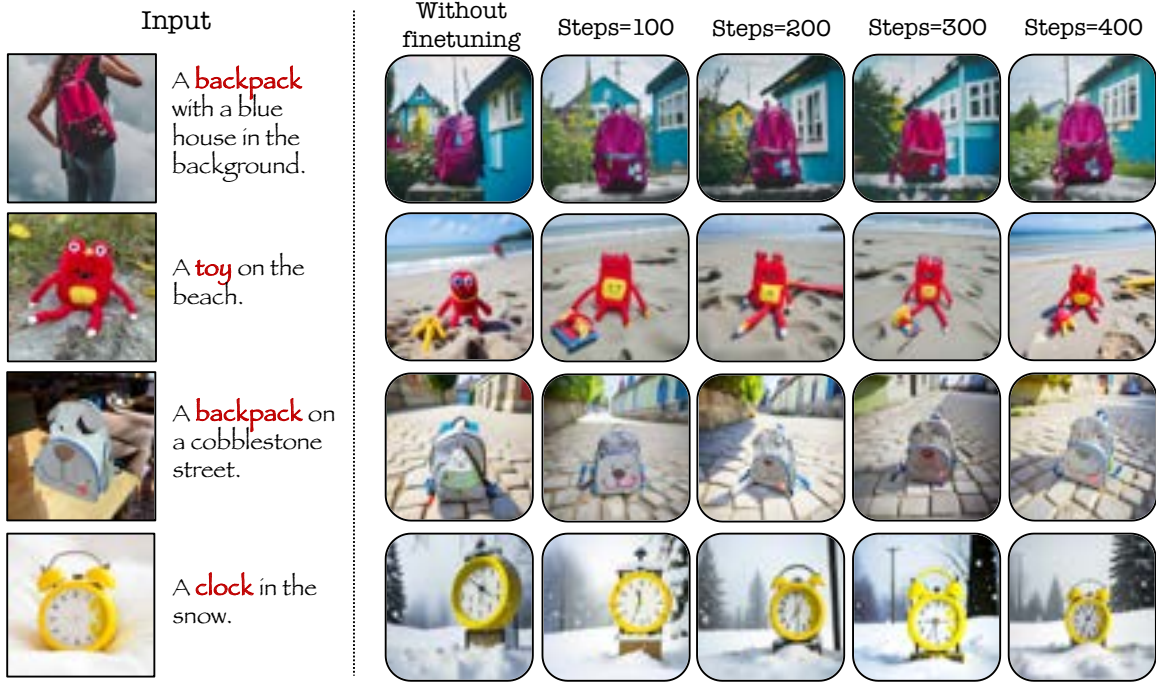


Figure 4.10: Qualitative examples of λ -ECLIPSE without finetuning and different stages of finetuning.

edge maps during training optimally balances both alignments and, specifically, improves the concept alignment. The negative values indicate that the *CCD* oracle model is highly confident in the generated images.

Multi-subject interpolation. A key attribute of the CLIP latent space is the ability to perform smooth interpolation between two sets of embeddings. We conducted experiments to demonstrate λ -ECLIPSE’s ability to learn and replicate this smooth latent space transition. We selected two distinct dog breeds ($\langle \text{dog1} \rangle$, $\langle \text{dog2} \rangle$) and two types of hats ($\langle \text{hat1} \rangle$, $\langle \text{hat2} \rangle$) as the concepts. λ -ECLIPSE was then used to estimate image embeddings for all four possible combinations, each corresponding to the input phrase “a $\langle \text{dog} \rangle$ wearing a $\langle \text{hat} \rangle$.” Fig. 4.8 showcases a gradual and seamless transition in the synthesized images from the top left to the bottom right. Unlike current diffusion models, which often exhibit sensitivity to input variations requiring numerous iterations of user interactions for desired

outcomes, λ -*ECLIPSE* inherits CLIP’s smooth latent space. This not only facilitates progressive changes in the conceptual domain but also extends the model’s utility by enabling personalized **multi-subject interpolations**.

Generalization to Pretrained UnCLIP Diffusion Decoders. Our approach is designed to generalize across any pretrained UnCLIP diffusion models, including Stable-UnCLIP/SDv2.1, Karlo, and Kandinsky v2.2. We conducted additional pretraining experiments to substantiate this claim and demonstrated the generalization ability of λ -*ECLIPSE*. We would like to reiterate the core pipeline of λ -*ECLIPSE* or UnCLIP models. The UnCLIP model comprises two key modules: (1) the Prior model, which maps the text embedding to the image embedding, and (2) the Diffusion rendering model, which synthesizes the image based on the estimated image embeddings from the prior model. For our generalization studies, we adhered to two essential criteria:

- Selection of an open-source UnCLIP (2 stage) model, where we chose Stable-UnCLIP, a modified version of Stable Diffusion v2.1, to accept image embeddings as input.
- Choice of pretrained CLIP model, with Stable-UnCLIP utilizing the OpenAI-CLIP-ViT-L/14 model. Accordingly, we trained λ -*ECLIPSE* with this model as a preliminary feature extractor.

Following these guidelines, we trained λ -*ECLIPSE* with the same parameters and evaluated its performance on DreamBench (see Table 4.10). Our results confirm that λ -*ECLIPSE* effectively generalizes to different diffusion models and CLIP variants. Notably, the Stable-UnCLIP variant of λ -*ECLIPSE* achieves performance similar to Emu2, reinforcing the superiority of our initial design choices.

Impact of Interleaved Pretraining. Section 4.3.2 mentions that training λ -*ECLIPSE* without interleaved pretraining would yield similar results for “single-concept” P-T2I tasks.

Table 4.10: Ablation study on generalization ability of λ -*ECLIPSE* with respect to different pretrained UnCLIP models. We report performance on DreamBench dataset.

Method	DINO ($\uparrow$)	CLIP-I ($\uparrow$)	CLIP-T ($\uparrow$)
Stable UnCLIP (Stable Diffusion v2.1)	0.564	0.778	0.276
Kandinsky v2.2	0.613	0.783	0.307

However, for “multi-concept” personalization, our experiments reveal that models trained without interleaved data sometimes struggle to synthesize the desired images. We conducted additional experiments by training the model without interleaved data. Specifically, we concatenated the prompt embedding from the CLIP text encoder with the concept-specific image embedding and trained the model with identical hyperparameters. The performance comparisons on DreamBench (single concept) and Multibench (multiple concepts) are shown below. Our findings indicate that without interleaved data, the model’s ability to align concepts decreases as the number of concepts increases, resulting in a sharp 6% drop in the DINO score (see Table 4.11). This performance degradation is primarily due to attribute leakage. As shown in Figure 4.11, when the model is tasked with generating “A backpack at the ruins” + $\langle$ backpack $\rangle$ + $\langle$ ruins $\rangle$, the non-interleaved pretrained models tend to generate the backpack with the color of the ruins.

Effect of Data and Model Sizes. We also conduct ablation on data sizes (100k, 500k, 1M, and 2M) and model parameter sizes (5M, 35M, and 70M). Tables 4.12 and 4.13 report the performance of these newly trained models on DreamBench. All models were trained with identical hyperparameters as the proposed λ -*ECLIPSE*, though this may not fully optimize the larger models (e.g., 70M parameters). The results show that data size influences model performance, with larger datasets improving concept understanding and prompt compositions. This also follows the qualitative results in Figure 4.12. Notably, λ -*ECLIPSE* with only 5M parameters deliver performance close to that of larger models,

Table 4.11: Ablation on the effect of interleaved data for training λ -*ECLIPSE*. We report performance on DreamBench dataset.

# of Subjects	Metric	w Interleaved	w/o Interleaved	Difference
1	DINO	0.6130	0.6006	-2.02%
	CLIP-I	0.7830	0.7882	0.66%
	CLIP-T	0.3070	0.3048	-0.72%
2	DINO	0.4478	0.4332	-3.26%
	CLIP-I	0.7409	0.7384	-0.34%
	CLIP-T	0.3327	0.3271	-1.68%
3	DINO	0.3420	0.3202	-6.37%
	CLIP-I	0.6463	0.6480	0.26%
	CLIP-T	0.3469	0.3469	0.00%

Table 4.12: Ablation study on pretraining data size. We report DreamBench performance of λ -*ECLIPSE* models trained on 100k, 500k, 1M and 2M interleaved image-text pairs. * denotes the proposed λ -*ECLIPSE* model checkpoint.

Data Size	DINO ($\uparrow$)	CLIP-I ($\uparrow$)	CLIP-T ($\uparrow$)
100K	0.593	0.786	0.301
500K	0.595	0.777	0.305
1M	0.596	0.778	0.306
2M*	0.613	0.783	0.307

and the 34M parameter model even surpasses the 70M parameter model in terms of DINO score. However, qualitative results (see Figure 4.13) show that increasing model parameters enhances qualitative performance and concept alignment. Specifically, the 70M parameter model excels in generating finer details of the reference concept while adhering closely to text prompts. The 34M model offers a more balanced trade-off between performance and

Table 4.13: Ablation study on pretraining model size. We report the DreamBench performance of λ -*ECLIPSE* models trained on 5M, 34M, and 70M parameters. * denotes the proposed λ -*ECLIPSE* model checkpoint.

Model Size	DINO ($\uparrow$)	CLIP-I ($\uparrow$)	CLIP-T ($\uparrow$)
5M	0.586	0.779	0.305
34M*	0.613	0.783	0.307
70M	0.593	0.775	0.309

resource efficiency.

4.6 Additional Qualitative Results & Failure Cases

In this section, we showcase a collection of detailed qualitative examples from the P-T2I generation process, highlighting the challenges of crafting complex compositions within λ -*ECLIPSE* and comparative models. As depicted in Figure 4.14, the complexity of the showcased examples progressively increases, illustrating a noticeable escalation in the intricacy of visual concepts from the top to the bottom of the figure. With the rising complexity, we note a universal decline in the ability of all methodologies, including λ -*ECLIPSE*, to preserve subject fidelity accurately. Interestingly, despite these challenges, λ -*ECLIPSE* demonstrates a better grasp of compositional integrity, unlike the baseline models which falter across all complexity levels.

Moreover, we present instances demonstrating the variability in outcomes produced by P-T2I methods across different trials. As illustrated in Figure 4.15, while there is a semblance of consistency in generating single and multiple concepts between models, Kosmos-G specifically shows variability in rendering multiple concepts—occasionally misplacing elements of the Ironman suit on a dog or failing to include it altogether. This phenomenon suggests that λ -*ECLIPSE* minimizes image diversity to enhance result consistency, a trait observed across the UnCLIP model family.

Figure 4.10 offers qualitative insights into the performance of λ -*ECLIPSE* without and with minimal fine-tuning. It is evident that in certain edge cases, where λ -*ECLIPSE* initially struggles to fully grasp novel visual concepts without finetuning, a modest application of few optimization iterations significantly enhances concept capture. Further optimization not only preserves text composition but also enriches minor, subject-specific details, underscoring the adaptability and finesse of λ -*ECLIPSE* in nuanced image generation.

Moreover, in our evaluations using the Multibench dataset, we noticed that both the baseline models (Kosmos-G and Emu2) and λ -*ECLIPSE* encounter difficulties in precisely maintaining all subject-specific details, as depicted in Figure 4.17. **This underscores that zero-shot multi-subject P-T2I generation remains a significant challenge in the field.** Further, we explored how well each model preserves genuine human facial characteristics in various scenarios, particularly when combined with differing captions. The qualitative examples in Figure 4.16 shed light on this aspect. Although each model strives to maintain the original facial features, none succeeds in replicating the specific personal facial details accurately. These instances typically fall short of precisely conveying the intended compositions, with the exception of one scenario in IP-Adapter FaceID, indicating a notable area for future improvements in model performance.

4.7 Conclusion

In this paper, we have introduced a novel training-time diffusion-agnostic methodology for personalized text-to-image (P-T2I) applications, utilizing the latent space of the pre-trained CLIP model with high efficiency. Our λ -*ECLIPSE* model, trained on an image-text interleaved dataset, achieves the capability to execute single-concept, multi-concept, and edge-guided controlled P-T2I tasks using a singular model framework, while simultaneously minimizing resource utilization. Notably, λ -*ECLIPSE* sets a new benchmark in achieving competitive performance in terms of concept and composition alignment. Furthermore,



Figure 4.11: Qualitative examples λ -ECLIPSE model trained with and without the interleaved pretraining strategy.

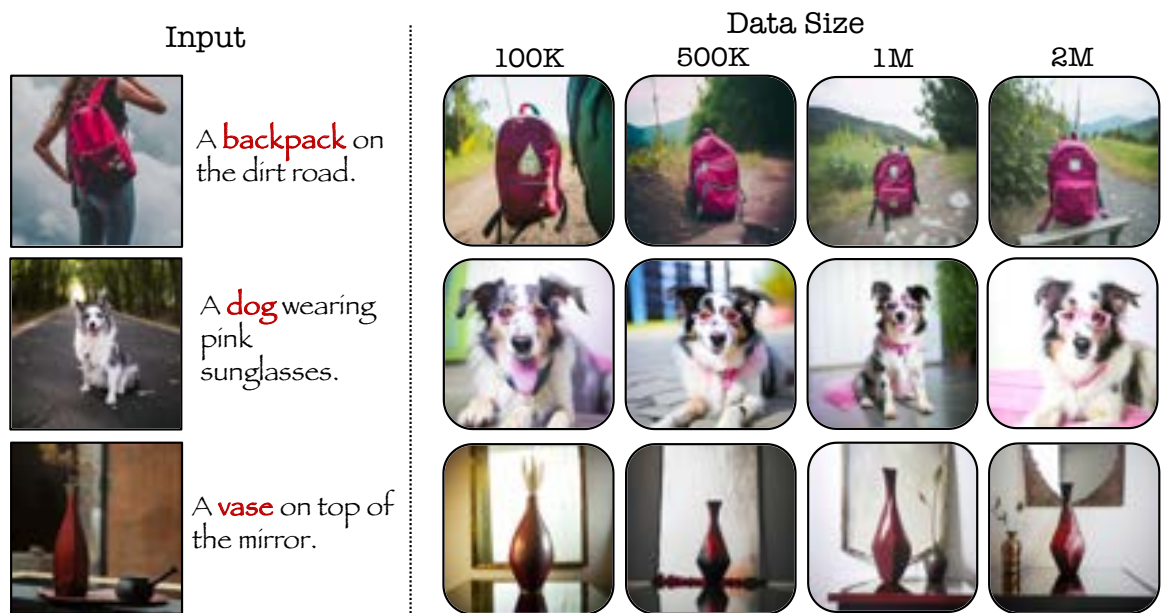


Figure 4.12: Qualitative examples λ -ECLIPSE model trained with varying pretraining data.

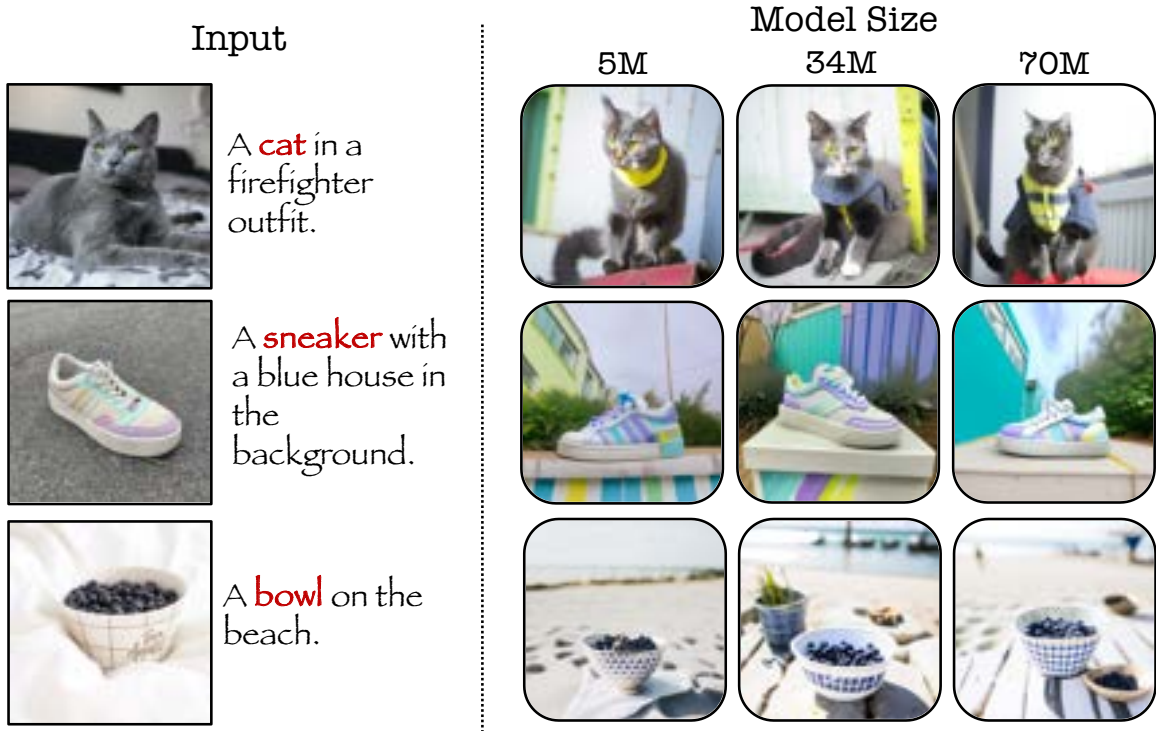


Figure 4.13: Qualitative examples λ -*ECLIPSE* model trained with varying model parameters.

our research illuminates the potential of λ -*ECLIPSE* in exploring and leveraging the smooth latent space. This capability unlocks new avenues for interpolating between multiple concepts, thereby generating entirely novel concepts. Our findings underscore a promising pathway to improve MLLMs to effectively control the pre-trained diffusion image models without necessitating extra supervision.

Limitations Primarily, despite its strengths, CLIP’s inability to perfectly capture hierarchical representations adds the upper bound on performance. Hence, enhancing CLIP’s representations could significantly boost our framework’s efficacy in P-T2I mapping. Even though λ -*ECLIPSE* model, trained on 34 million parameters and 1.6 million images, presents a substantial foundation, yet, there’s potential for further refinement, as increasing the quality of data and the number of parameters could yield even better outcomes. However,



Figure 4.14: Qualitative examples of the increasing complexity of novel visual concepts as we move from top to bottom.

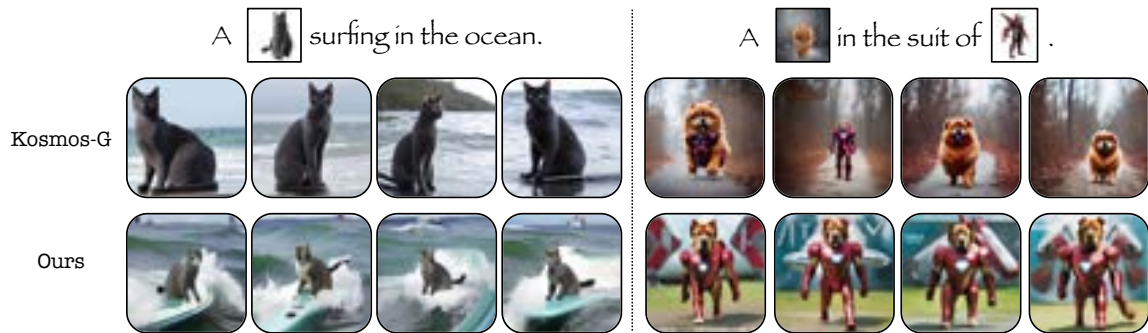


Figure 4.15: Qualitative examples of showcasing the consistency comparisons between Kosmos-G and λ -ECLIPSE.

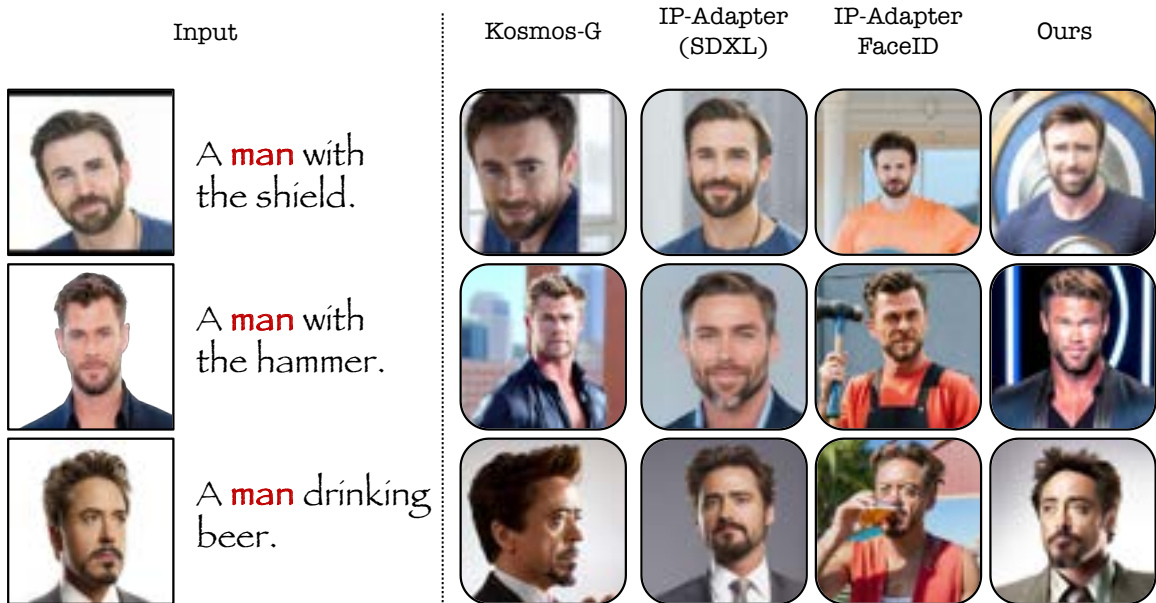


Figure 4.16: Qualitative examples of showcasing the failure cases on human faces on Kosmos-G, IP-Adapter (SDXL), IP-Adapter (FaceID), and λ -ECLIPSE.

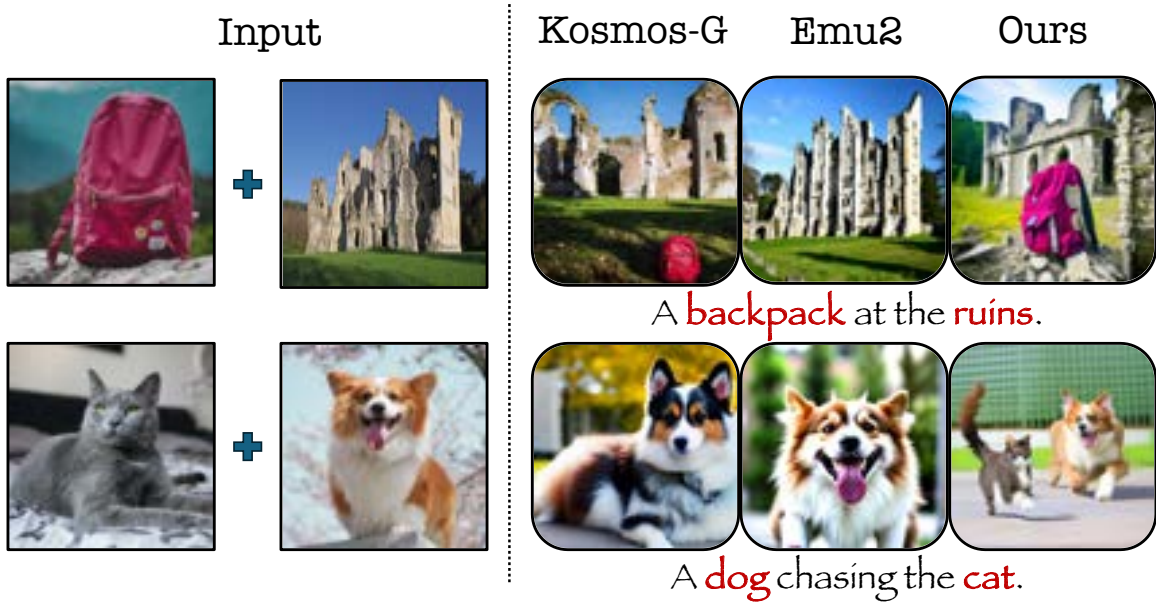


Figure 4.17: Qualitative examples of showcasing the failure cases on Multibench of Kosmos-G, Emu2, and λ -ECLIPSE.

this is outside the scope of this paper. Additionally, we validate the λ -*ECLIPSE* on only upto three concepts (due to the real-life usecases) and adding more concepts could be explored in future works.

Broader Impact The current landscape of text-to-image (T2I) generative models is dominated by approaches that rely on extensive data and large-scale models to achieve state-of-the-art (SOTA) performance, which demands significant computational resources. In contrast, our work with λ -*ECLIPSE* demonstrates that it is feasible to attain competitive performance relative to SOTA large models while achieving a tenfold reduction in resource consumption. This advancement not only makes T2I generative models more accessible and cost-effective but also promotes sustainable AI practices by significantly lowering the environmental impact associated with large-scale model training.

IMPROVING COMPOSITIONAL REASONING OF CLIP VIA SYNTHETIC VISION-LANGUAGE NEGATIVES

5.1 Introduction

Large-scale vision-language models, such as CLIP [210], have significantly advanced multi-modal learning by employing contrastive learning to acquire shared semantic representations from paired datasets. This approach has resulted in improved performance in vision-language tasks as well as zero-shot image classification [250] and segmentation [105, 329]. Beyond vision-language tasks, the individual components of these models, such as the vision encoder and the language encoder, are integral to several multimodal architectures and generative models such as multimodal large language models (MLLMs) [149, 133] and text-to-image (T2I) diffusion models [229, 194, 195]. Yet, compositional reasoning remains challenging and multimodal models continue to exhibit naïve “bag of words” behavior, frequently failing to distinguish between expressions like “bulb in the grass” and “grass in the bulb” [307, 262]. Addressing this challenge remains critical for enhancing vision-language models and their downstream applications.

Contrastive learning of representations benefits from “hard negative samples” (i.e., points that are difficult to distinguish from an anchor point) [215]. However, at each optimization step for training CLIP, image-text pairs are *randomly* sampled from the training dataset – this random sampling seldom exposes the model to highly similar negative pairs. We hypothesize that the limited compositional understanding of CLIP may stem from such issues in the optimization objective and sampling from training datasets. A straightforward solution could involve iteratively identifying hard negative pairs for each training iteration. However, due

to the noisy captions and the scarcity of such pairs in existing datasets, prior work generates hard negative captions as a form of augmentation using rule-based strategies [307, 316]. For instance, given an image-text pair labeled “a brown horse”, an additional negative caption “a blue horse” might be introduced. However in prior work, image data is not subjected to similar hard negative semantic augmentation during training; this is mainly because of the difficulty of making semantic perturbations at the pixel levels compared to sentence perturbation. While the text-only augmentation strategies have improved the models’ compositional understanding to a certain extent, it raises an intriguing question: *could incorporating hard negative augmentation for both text and image modality further enhance the compositional reasoning capabilities of vision-language models?*

Motivated by this, in this paper, we introduce a novel, simple, and yet highly effective strategy for integrating hard negative images as well as hard negative text to enhance the compositional understanding of vision-language models. Recent developments in text-to-image diffusion models have opened up possibilities for performing semantic perturbations within images [101]. Existing works have evaluated the impact of creating synthetic data for text-to-image generative models [19, 31]. However, it remains underexplored how these generative models can benefit CLIP-like models. To tackle this challenge, our approach leverages the in-context learning capabilities of LLMs to produce realistic, linguistically accurate negative captions [276]. We then employ a pre-trained text-to-image diffusion model to create images corresponding to these captions, thereby enriching any given image-text dataset with valuable hard negatives that foster improved reasoning. This resulting `TripletData` comprises 13M image-text pairs to complement the CC3M and CC12M datasets [30].

We developed `TripletCLIP`, which incorporates hard negative image-text pairs effectively by using them to optimize a novel triplet contrastive loss function. Extensive experiments on the CC3M and CC12M datasets and various downstream tasks with an equal

compute budget demonstrate that `TripletCLIP` significantly enhances compositional reasoning. Notably, `TripletCLIP` results in more than 9% and 6% absolute improvement on the SugarCrepe benchmark compared to LaCLIP [132] and NegCLIP [307], respectively. `TripletCLIP` also improves zero-shot classification and image-text retrieval performance with similar training-time concept diversity. An investigation into the effects of increasing training-time concept diversity revealed that baseline models consistently under-performed in compositional tasks despite an increase in integrated knowledge, while `TripletCLIP` demonstrated significant improvements. In summary, our **key contributions** are as follows:

- We introduce a novel CLIP pre-training strategy that employs hard negative text and images in conjunction with triplet contrastive learning to enhance compositionality.
- `TripletCLIP` consistently improves across downstream tasks, demonstrating the effectiveness of synthesizing hard negative image-text pairs.
- Our extensive ablations on the choice of the loss function, modality-specific pre-training, the increase in concept diversity, and filtering high-quality `TripletData` provide deeper insights into the utility of hard negative image-text pairs for CLIP pre-training.
- Ultimately, we present a promising avenue where synthetic contrastive datasets significantly improve reasoning capabilities, leading to the creation and release of the `TripletData` — a 13M contrastive image-text dataset.

5.2 Related Work

Vision-Language Models. Recent advancements, including ALIGN [103] and CLIP [210], have gained significant interest due to their capability to learn transferable semantic representations across multiple modalities through contrastive learning. These models facilitate

downstream tasks such as zero-shot classification [250], image-text retrieval [315, 11], visual grounding/reasoning [150], text-to-image generation [259, 194, 195, 114], semantic segmentation [329, 105], and various evaluations [93, 232]. Subsequent research has sought to enhance various aspects of these models, including data efficiency [66], hierarchical representation learning [50], and the quantization of latent spaces for more stable pre-training [39]. LiT [310] employs a pre-trained frozen CLIP vision encoder to fine-tune a BERT-like text encoder [51], achieving notable improvements in zero-shot transfer performance. Similarly, BLIP-2 [137] combines contrastive pre-training with the next-token prediction for image captioning during training. However, these approaches generally presume the availability of high-quality data. In contrast, `TripletCLIP` focuses on leveraging the proposed hard negative contrastive dataset and incorporating triplet contrastive pre-training for compositional data. This approach is orthogonal to prior works.

Data for Contrastive Pre-training. The effectiveness of maximizing mutual information between modalities heavily relies on the quality of extensive, web-scraped datasets that ideally encompass all possible concepts and knowledge. For instance, despite its noise, the LAION dataset [236, 71], which includes more than 5 billion internet images paired with alt-text captions, is a primary resource. Studies show that over 1 billion data points are necessary to match the performance of the original CLIP model [71, 289]. Recent works like DataComp [71] and MetaCLIP [289] have focused on creating smaller, high-quality datasets by applying stringent filters and ensuring wordnet [176] synset-level concept diversity. Nevertheless, the inherently noisy nature of internet-scraped datasets can degrade model performance. Studies such as SynthCLIP [85] demonstrate that tripling the volume of fully synthetic data is required to equal the efficacy of real data. Other efforts like VeCLIP [132] and LaCLIP [64] enhance dataset quality by using generative language models to re-caption existing images, significantly boosting performance.

Compositionality for vision-language. Despite the increased emphasis on data quality

and modeling techniques, mastering compositionality remains a significant challenge for vision-language models. Benchmarks like ARO [307], VALSE [187], and CREPE [170] have been developed to assess models’ abilities to handle compositional data. SugarCrepe [98], in particular, offers a large-scale, systematic framework for such evaluations. Previous methods primarily focused on identifying hard negatives within existing datasets or generating synthetic negative captions [307, 316, 58, 57, 292, 244]. However, these rule-based generated captions are often unrealistic and linguistically flawed, leading to sub-optimal model performance on complex datasets like SugarCrepe. A handful of works focus on finding negative images. [275] propose utilizing the video data. [227] focuses on object-centric image-editing to synthesize the negative images. [198] utilizes the simulation-based data negative data.

Contrary to prior approaches that predominantly add unrealistic negative captions or very constrained negative images that are either very synthetic or object-focused, this work introduces `TripletCLIP`, which centers on generating naturally occurring *hard negative image-text pairs*. We propose a novel triplet contrastive learning strategy that effectively utilizes these challenging data pairs. Additionally, while our method is distinct, integrating advancements that refine contrastive learning could potentially boost `TripletCLIP`’s efficacy further.

5.3 Method

This section begins with an overview of the contrastive learning algorithm used by CLIP and NegCLIP. We then describe the synthetic data generation pipeline for generating hard negatives using LLMs and T2I models and introduce triplet contrastive learning which forms the basis of `TripletCLIP`. A high-level comparison between prior work and `TripletCLIP` can be found in Figure 5.1.

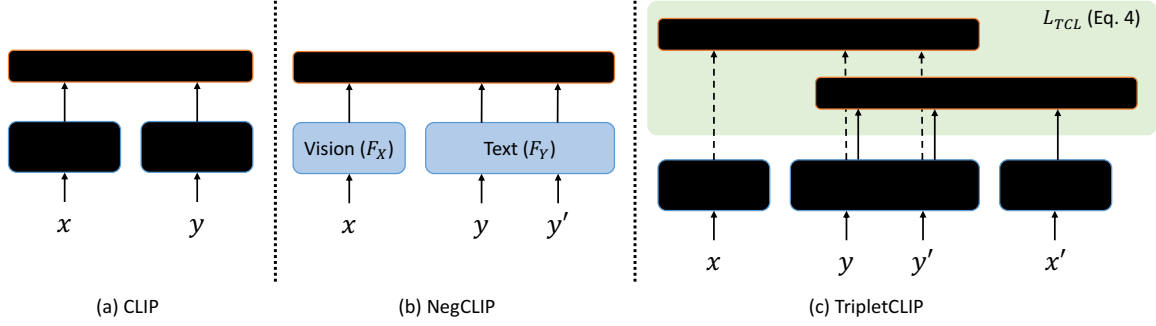


Figure 5.1: Comparison of training workflows of CLIP, NegCLIP, and TripletCLIP. (x, y) represents the positive a image-text pair, and (x', y') represents the corresponding negative image-text pair.

5.3.1 Preliminaries

The goal for self-supervised contrastive learning [60], when dealing with inputs from a single modality, is to use a feature extractor (F) to encode inputs and their augmentations and minimize the InfoNCE loss [185] between the two encodings. CLIP is designed for multimodal settings (for example, vision and language inputs) – this entails using two encoders (one for each modality).

Let $\mathcal{X}$ and $\mathcal{Y}$ represent two modalities and $\mathcal{D} = \{(x_i, y_i)\}_{i=1, \dots, M}$, be the training dataset, where $x_i \in \mathcal{X}$ and $y_i \in \mathcal{Y}$. The goal is to train two modality-specific encoders, $F_{\mathcal{X}}$ and $F_{\mathcal{Y}}$, by minimizing the InfoNCE loss between the normalized features extracted from the encoders.

$$\mathcal{L}_{\mathcal{Y}}^{CL} = \frac{-1}{N} \sum_{i=1}^N \log \frac{\exp(\langle F_{\mathcal{X}}(x_i), F_{\mathcal{Y}}(y_i) \rangle / \tau)}{\sum_{k=1}^N \exp(\langle F_{\mathcal{X}}(x_i), F_{\mathcal{Y}}(y_k) \rangle / \tau)}, \quad (5.1)$$

where $\langle \cdot \rangle$ represents cosine similarity and τ is the trainable temperature parameter. For simplicity, we do not show feature normalization in the InfoNCE loss. Similarly, we can define the $\mathcal{L}_{\mathcal{X}}^{CL}$ training loss. The combined CLIP total training objective is given as, $\mathcal{L}_{CLIP} = \mathcal{L}_{\mathcal{X}}^{CL} + \mathcal{L}_{\mathcal{Y}}^{CL}$. By minimizing this training loss, both encoders learn representations that maximize the mutual information between two modalities.

NegCLIP introduces synthetic augmentations to generate “hard” negative captions ($y'_i \in \mathcal{Y}'$) by performing semantic inverting perturbations to the reference captions ($y_i \in \mathcal{Y}$).

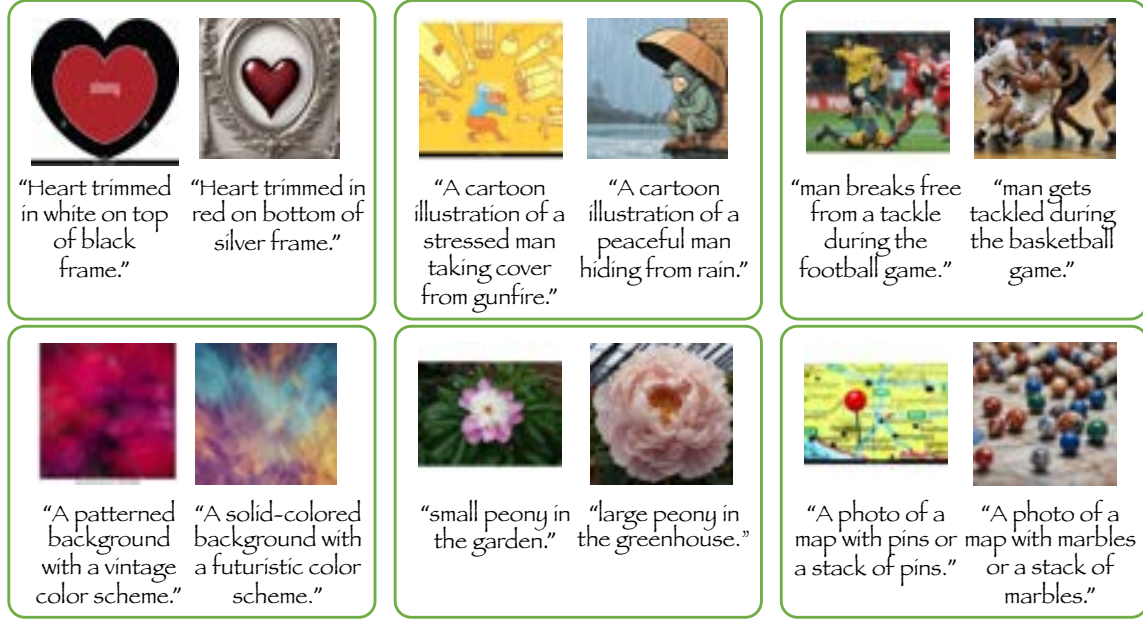


Figure 5.2: Examples image-text pairs from `TripletData`. In each block, a positive pair from CC3M is on the left and corresponding negatives from `TripletData` are shown on the right.

Therefore, the single modality-specific hard negative augmentation-based training loss can be formulated as:

$$\mathcal{L}_{\mathcal{X} \mathcal{Y}; \mathcal{Y}'}^{NegCL} = \frac{-1}{N} \sum_{i=1}^N \log \frac{\exp(\langle F_{\mathcal{X}}(x_i), F_{\mathcal{Y}}(y_i) \rangle / \tau)}{\sum_{k=1}^N \exp(\langle F_{\mathcal{X}}(x_i), F_{\mathcal{Y}}(y_k) \rangle / \tau) + \sum_{m=1}^N \exp(\langle F_{\mathcal{X}}(x_i), F_{\mathcal{Y}}(y'_m) \rangle / \tau)}. \quad (5.2)$$

The total loss for NegCLIP for image modality ($\mathcal{X}$) and text modality ($\mathcal{Y}$) is given by:

$$\mathcal{L}_{NegCLIP}(\mathcal{X}, \mathcal{Y}, \mathcal{Y}') = \mathcal{L}_{\mathcal{Y} \mathcal{X}}^{CL} + \mathcal{L}_{\mathcal{X} \mathcal{Y}; \mathcal{Y}'}^{NegCL}. \quad (5.3)$$

In Eq. 5.2, the negative samples are generated only for language modality as it is easy to make semantic-level perturbations. Existing methods have not explored performing semantic perturbations in the image modality to create hard negatives. In this work, we demonstrate how hard negatives can be created in the image modality by leveraging the semantic language grounding and photorealism of text-to-image diffusion models. Our novel hard negative generation pipeline and refined training objective seeks to bridge the significant gap identified in literature.

5.3.2 *TripletData*: Image-text hard negative data augmentations

To generate high-quality hard negative image-text pairs, we follow a two-step procedure. The first stage is to generate hard negative captions from the ground truth positive caption. Second, to generate images corresponding to the hard negative captions as negative images. The AltText captions from the existing web-scraped datasets are very noisy, leading to the noisy and unreliable generation of hard negatives. Therefore, we build upon the existing work LaCLIP, which first rewrites the captions using LLM from the existing data that are linguistically accurate. Figure 5.2 illustrates several examples of positive and corresponding negative image-text pairs.

Generating hard negative captions. Existing works perform random swapping, replacing, and adding actions between the nouns, attributes, and relations of the positive caption [307, 316]. This method results in nonsensical and grammatically incorrect artifacts, such as “a person riding on four slope,” which impedes the generation of negative images, ultimately leading to diminishing performance on harder benchmarks [98]. Therefore, we utilize the in-context learning ability of LLMs to generate negative captions. The choice of LLM is a trivial task as long as they provide hard negative captions. We find that Mistral-7B-Instruct-v0.2¹ performs reasonably better on our goal, and the output is easy to parse. We generate the negative captions in batches to speed up the generation process. Generating the 13M negative captions takes only 3 days on 8xRTX A6000. Instead of generating multiple hard negative captions, we find that a single high-quality hard negative caption is enough to improve the performance compared to the traditional NegCLIP style caption generation (see NegCLIP++ results in Table 5.4). We provide examples of various types of negative captions in the appendix. Specifically, we provide the following prompt to LLM:

¹<https://huggingface.co/mistralai/Mistral-7B-Instruct-v0.2>

You are given the description of an image. You should provide a mostly similar description, changing the original one slightly but introducing enough significant differences such that the two descriptions could not possibly be for the same image. Keep the description length the same. Finally, only a few things (such as counting, objects, attributes, and relationships) can be modified to change the image structure significantly. Provide just the updated description. Examples: Input: A dog to the left of the cat. Output: A dog to the right of the cat. Input: A person wearing a red helmet drives a motorbike on a dirt road. Output: A person in a blue helmet rides a motorbike on a gravel path. Now, do the same for the following captions: Input: { } Output:

Generating hard negative images. Typically, semantic perturbations within images require tools like image editing, which are resource-intensive and cannot be scaled. Remember that we want to provide additional ground truth references for the negative caption. Therefore, we propose to utilize the negative captions from the previous stage to generate the respective reference images directly for pre-training. As the previous stage generates negative captions that are linguistically correct, it becomes easier for image-generative models to synthesize the respective images precisely. We utilize pre-trained text-to-image diffusion models to generate the corresponding images. Specifically, we select SDXL-turbo [230] due to its relatively faster generation speed. After applying various inference time optimizations, we can generate 13M negative images within 2 days using 30 v100 GPUs. We provide various examples of the hard negative image-text pairs in the appendix.

analyzing difficulty of the hard TripletData. Let’s assume we have positive and negative image-text pairs from the `TripletData`, (x_i, y_i) and (x'_i, y'_i) , respectively. If the data is truly hard negative, existing pre-trained models should struggle to find the correct image-text pairs (i.e., $\cos(x_i, y_i) > \cos(x_i, y'_i)$). Following winoground, we measure the text-score, image-score, and group-score to evaluate the popular pretrained CLIP

Table 5.1: Winoground-style evaluation of pretrained CLIP models on `TripletData`.

	Img Score	Text Score	Grp Score
ViT-B/32	40.29	68.17	36.53
ViT-L/14	44.84	69.21	40.91
ViT-bigG	42.94	77.61	40.98
Siglip-so400m	44.24	71.27	26.10
Humans (on Winoground)	88.50	89.50	85.50

Table 5.2: Wordnet synset analysis of captions from CC3M and `TripletData`.

	CC3M	TripletData (Negative Only)	TripletData	Intersection
# unique	59094	59616	62741	55969
# total synsets	231M	215M	446M	-

models. Table 5.1 shows that even CLIP models trained on billions of data struggle to get near human performance on `TripletData`, which is less difficult than winoground. Importantly, the goal of generating hard negative samples isn't to add more diversity in terms of unique concepts during the training but to add diversity in semantic meanings. Therefore, we measure the unique wordnet synsets in CC3M *vs.* `TripletData`. From Table 5.2, it can be observed that `TripletData` does not add any new concepts but uses existing concepts to provide negative samples that are semantically different. To summarize, `TripletData` contains the relatively hard negative image-text pairs that current models find difficult to differentiate.

5.3.3 TripletCLIP

Prior works have demonstrated the value of hard negative captions for enhancing the compositionality of CLIP models via $\mathcal{L}_{NegCLIP}$ as the key training objective (Eq. 5.2) [307, 316]. However, it remains elusive if negative images alone can benefit or not. We conduct modality-specific ablations, reporting the average performance across the diverse set of benchmarks in Table 5.3 (we provide more details about experiments in Section 5.4). Our findings indicate that both “hard” negative captions and images individually boost performance when compared to LaCLIP. However, this initial empirical experiments to train the CLIP model on hard negative images (i.e., NegImage) by minimizing $\mathcal{L}_{NegCLIP}(\mathcal{Y}, \mathcal{X}, \mathcal{X}')$ reveal that negative images alone cannot improve the compositionality significantly (see Table 5.3). We hypothesize that images contain low-level information, making it difficult to train the model using images as negative examples. Aligning with our initial motivation and building upon this crucial insight, we propose to utilize the negative images to regularize the effect of negative captions and to stabilize the pre-training. Therefore, to utilize these hard negative image-text pairs from the previous stage more effectively, we propose to focus on two triplets $(\mathcal{X}, \mathcal{Y}, \mathcal{Y}')$ and $(\mathcal{X}', \mathcal{Y}', \mathcal{Y})$, hence, the final triplet contrastive learning training objective is defined as:

$$\mathcal{L}_{TCL} = \mathcal{L}_{NegCLIP}(\mathcal{X}, \mathcal{Y}, \mathcal{Y}') + \mathcal{L}_{NegCLIP}(\mathcal{X}', \mathcal{Y}', \mathcal{Y}). \quad (5.4)$$

Intuitively, the second term introduces the additional form of supervision that hard negative images are closer to the corresponding negative captions than positive captions. This allows the system to understand that if the positive image does not represent the negative caption “blue horse,” then what does this caption entail? Through this strategic alternation of hard negative image-text pairs for the TripletCLIP, we improve compositionality and image-text understanding of the vision-language model (see Table 5.3). We provide the pseudo-code in the appendix and the code in supplementary materials. This simple yet

Table 5.3: Importance of image-text hard negatives. We measure the importance of various modality-specific hard negatives on SugarCrepe, image-text retrieval, and ImageNet1k. We find that `TripletCLIP` results into the most optimal solution. Bold number indicates the best performance.

Models	Negative Captions	Negative Images	SugarCrepe	Retrieval	ImageNet1k
LaCLIP	×	×	54.09	8.19	3.79
NegImage	×	✓	56.28	9.20	4.48
NegCLIP++	✓	×	61.69	8.36	3.84
TripletCLIP	✓	✓	63.49	16.42	7.31

effective strategy elevates the training of the CLIP, offering a scalable framework to improve overall performance.

5.3.4 Pseudocode of `TripletCLIP`

```
def forward(batch_positive, batch_negative):
    # get positive and negative image-text pairs
    img_pos, txt_pos = batch_positive
    img_neg, txt_neg = batch_negative

    # compute positive image and text representations
    image_features_pos = clip.encode_image(img_pos)
    text_features_pos = clip.encode_text(txt_pos)

    # compute negative image and text representations
    image_features_neg = clip.encode_image(img_neg)
    text_features_neg = clip.encode_text(txt_neg)

    positive_txt = torch.cat([text_features_pos,
```

```

text_features_neg]))
    negative_txt = torch.cat([text_features_neg,
text_features_pos]))

    # compute NegCLIP losses
    loss_1 = negclip_loss(image_features_pos, positive_txt)
    loss_2 = negclip_loss(image_features_neg, negative_txt)
    loss = loss_1 + loss_2

    return loss

```

5.4 Experiments & Results

5.4.1 Experiment Setup

Pretraining Datasets. We utilize the CC3M and CC12M datasets, which comprise 2.6M and 8.6M image-text pairs, respectively. Following the approach demonstrated by LaCLIP, we use LLM-rewritten captions to replace noisy original captions. For NegCLIP, we introduce four negative captions per positive image-text pair, focusing on semantic inverting perturbations across four categories: attribute, relation, object, and action [316]. This generates approximately 10.4M and 34.4M text-only augmentations for CC3M and CC12M, respectively. To train the TripletCLIP, we create augmentations (TripletData) for both datasets to integrate hard negatives effectively. We produce one augmentation per image-text pair, adding 2.6M and 8.6M image-text augmented pairs for CC3M and CC12M, respectively. Finally, we perform all the ablations on the CC3M dataset.

Baselines. We train LaCLIP, LaCLIP with real hard negatives (LaCLIP+HN), and NegCLIP from scratch to ensure consistency and fairness in our comparisons. As NegCLIP’s rule-based augmentations closely resemble some compositional benchmarks, so we introduce NegCLIP++ as an improved baseline. NegCLIP++ incorporates hard negative captions generated using LLM from `TripletData`, enhancing the language comprehension compared to standard NegCLIP.

Implementation Details. Our experiments employ the ViT-B/32 [56] model architecture. To guarantee fair comparisons, we retrain all baseline models using identical hyperparameters. Since the overall training data for NegCLIP and `TripletData` is more than the baseline datasets, we align the number of iterations across all models to equalize the number of image-text pairs seen during training, similar to the strategy used in DataComp. The batch size is fixed to 1024 with the AdamW optimizer at a maximum learning rate of 0.0005, employing cosine decay. Training durations are set at approximately 100k iterations for CC3M and 200k iterations for CC12M. All models are trained on a single A100 (80GB) GPU using bf16 precision. The final training-related experiments and ablations will cost about 1200 A100 GPU hours. We leave the experiments on increasing the data and model size as future works for the community, as scaling further is not viable in the academic budget.

Hyperparameters. Table 5.6 provides a comprehensive overview of the pre-training hyperparameters employed across all baseline models and `TripletCLIP`. To ensure fair comparisons, we standardized the hyperparameters across all methodologies. Although larger batch sizes are typically associated with improved performance in contrastive learning, computational constraints necessitated fixing the batch size at 1024 for all experiments. To accommodate this batch size on a single A100 GPU, we employed bf16 precision. In terms

Table 5.4: Composition evaluations of the methods on SugarCrepe benchmark. Bold number indicates the best performance, and underlined number denotes the second-best performance. † represents the results taken from SugarCrepe benchmark.

Methods	Replace			Swap		dd		Overall
	Object	ttribute	Relation	Object	ttribute	Object	ttribute	vg.
CC3M	LaCLIP	59.44	53.17	51.42	54.69	49.25	55.29	54.09
	LaCLIP + HN	63.44	55.96	50.71	50.60	48.57	56.98	53.92
	NegCLIP	62.71	58.12	54.48	56.33	51.20	56.26	57.18
	NegCLIP++ (<i>ours</i>)	<u>64.77</u>	<u>66.12</u>	65.93	<u>55.51</u>	<u>55.41</u>	<u>59.65</u>	<u>61.69</u>
	TripletCLIP (<i>ours</i>)	69.92	69.03	<u>64.72</u>	56.33	57.96	<u>62.61</u>	63.49
Performance Gain w.r.t. LaCLIP		10.48%	18.56%	13.30%	1.64%	8.71%	7.32%	9.40%
CC12M	LaCLIP	75.06	65.48	58.68	53.47	57.66	67.65	63.54
	NegCLIP	77.84	69.29	63.23	66.53	62.31	67.17	68.00
	NegCLIP++ (<i>ours</i>)	<u>82.99</u>	<u>78.68</u>	<u>75.75</u>	61.63	65.47	70.08	72.94
	TripletCLIP (<i>ours</i>)	83.66	81.22	79.02	<u>64.49</u>	<u>63.66</u>	<u>73.67</u>	74.45
Performance Gain w.r.t. LaCLIP		8.60%	15.75%	20.34%	11.02%	6.00%	8.67%	10.91%
DataComp	small: ViT-B/32† (13M)	56.90	56.85	51.99	50.81	50.00	53.93	54.43
	medium: ViT-B/32† (128M)	77.00	69.54	57.68	57.72	57.06	66.73	64.37
	large: ViT-B/16† (1B)	92.68	79.82	63.94	56.10	57.66	84.34	73.31
	xlarge: ViT-L/14† (13B)	95.52	84.52	69.99	65.04	66.82	91.03	79.70

of computational resources, experiments using the CC3M dataset required approximately 16 GPU hours, while those involving the CC12M dataset utilized up to 56 GPU hours per experiment.

Downstream Datasets. The primary objective of this study is to enhance the compositional capabilities of CLIP models. We mainly evaluate TripletCLIP and the baseline models using the challenging SugarCrepe composition benchmark, with additional performance assessments provided in the appendix for older benchmarks. Models are also tested on image-text retrieval tasks for broader evaluation using the Flickr30k [204] and MSCOCO [146] datasets. Zero-shot classification performance is assessed across approximately 18 different datasets. Evaluations adhere to the methodologies outlined in the

Table 5.5: Zero-shot image-text retrieval and classification results. Bold number indicates the best performance, and underlined number denotes the second-best performance.

Methods		Retrieval (R@5)				Zero-shot Classification			
		Image-to-Text		Text-to-Image		VT B		ImageNet1k	
		MSCOCO	Flickr30k	MSCOCO	Flickr30k	top-1	top-5	top-1	top-5
CC3M	LaCLIP	5.06	10.90	5.97	10.84	11.56	34.72	3.79	10.49
	LaCLIP + HN	<u>8.08</u>	<u>16.10</u>	<u>8.64</u>	<u>16.64</u>	12.31	<u>37.14</u>	<u>5.75</u>	<u>15.22</u>
	NegCLIP	6.32	13.80	6.61	12.96	<u>12.25</u>	36.38	4.67	12.69
	NegCLIP++ (<i>ours</i>)	5.8	11.20	6.19	10.24	11.65	35.47	3.84	10.52
	TripletCLIP (<i>ours</i>)	10.38	22.00	11.28	22.00	12.31	41.45	7.32	18.34
Performance Gain		5.32%	11.1%	5.31%	11.16%	0.75%	6.73%	3.53%	7.85%
CC12M	LaCLIP	25.86	42.70	19.78	36.30	19.08	49.06	19.72	41.39
	NegCLIP	<u>30.16</u>	<u>46.60</u>	<u>23.11</u>	41.70	<u>19.12</u>	<u>50.56</u>	<u>20.22</u>	<u>42.63</u>
	NegCLIP++ (<i>ours</i>)	26.96	43.90	22.69	<u>42.86</u>	18.48	50.38	19.06	40.91
	TripletCLIP (<i>ours</i>)	33.00	55.90	28.50	52.38	20.81	53.40	23.31	47.33
Performance Gain		7.14%	13.2%	8.72%	16.08%	1.73%	4.34%	3.59%	5.94%

CLIP-Benchmark ² or the official benchmark implementations.

5.4.2 Compositional reasoning

We comprehensively analyze the compositional understanding of models on the Sugar-Crepe benchmark, as detailed in Table 5.4. Notably, TripletCLIP consistently outperforms all baseline models across all sub-categories of SugarCrepe on both the CC12M/CC3M training datasets. Specifically, TripletCLIP surpasses LaCLIP and NegCLIP by **10.91%/9.4%** and **6.45%/6.31%** on the CC12M and CC3M datasets, respectively. Our enhanced baseline, NegCLIP++, also shows improvement over standard NegCLIP, highlighting the benefits of LLM-generated negatives. Nevertheless, TripletCLIP further advances performance, underscoring the critical role of hard negative image-text pairs, not just text. Additional com-

²https://github.com/LION- I/CLIP_benchmark

Table 5.6: Detailed pre-training hyper-parameters for CLIP training across various experiments and ablations.

Hyperparameters	CC3M	CC12M	LiT	Concept Coverage	ablations
Batch size	1024	1024	1024	1024	
Optimizer	AdamW	AdamW	AdamW	AdamW	
Learning rate	5×10^{-4}	5×10^{-4}	5×10^{-4}	5×10^{-4}	
Weight decay	0.5	0.5	0.5	0.5	
dam β	(0.9, 0.999)	(0.9, 0.999)	(0.9, 0.999)	(0.9, 0.999)	
dam ϵ	1×10^{-8}	1×10^{-8}	1×10^{-8}	1×10^{-8}	
Total steps	90,000	230,000	90,000	200,000	
Learning rate schedule	cosine decay	cosine decay	cosine decay	cosine decay	

Table 5.7: Ablation on filtering high-quality image-text pairs from `TripletData`. We evaluate the `TripletCLIP` after applying the filters to ensure the quality similar to DataComp and compare the baselines on three benchmarks. We find that `TripletCLIP` results in the most optimal solution. Bold number indicates the best performance. $\dagger$ represents that results are borrowed from DataComp.

Models	Filtering Strategy	Data Size	augmentations	SugarCrepe	Retrieval	ImageNet1k
CLIP$\dagger$	No filtering	12.8	-	55.61	6.49	2.7
	CLIP Score	3.8	-	57.31	9.08	5.1
	Image-based $\cap$ CLIP Score	1.4	-	54.75	5.63	3.9
LaCLIP	No filtering (CC3M)	2.6	-	54.09	8.19	3.79
TripletCLIP	No filtering (CC3M)	2.6	2.6	<u>63.49</u>	<u>16.42</u>	<u>7.31</u>
TripletCLIP++	CLIP Score (from CC12M)	1.4	1.4	66.09	19.85	8.85

parisons on older composition benchmarks (Valse [187], Cola [213], and Winoground [262]) in the appendix reveal `TripletCLIP`’s consistent performance. Table 5.4 also contrasts `TripletCLIP` with models trained using the DataComp approach, which involves more parameters and training data, demonstrating that `TripletCLIP` achieves comparable performance to a ViT-B/16 model trained on 1 billion image-text pairs.

Table 5.8: Finetuning-based composition evaluations of the methods on SugarCreme benchmark. Bold number indicates the best performance, and underlined number denotes the second-best performance.

Methods	Replace			Swap		dd		Overall
	Object	ttribute	Relation	Object	ttribute	Object	ttribute	vg.
CLIP	90.92	80.08	69.13	61.22	64.26	77.16	68.64	73.06
CLIP (finetuned)	90.92	79.69	64.01	60.82	64.26	84.67	78.76	74.73
NegCLIP	91.53	83.25	73.97	72.24	67.72	86.95	<u>88.44</u>	80.59
Baseline [227]	93.22	84.39	67.35	62.04	70.12	88.31	79.48	77.84
CoN-CLIP [245]	<u>93.58</u>	80.96	63.3	87.29	79.62	59.18	65.16	75.58
TSVLC (RB) [58]	91.34	81.34	64.15	68.16	69.07	79.49	91.33	77.84
TSVLC (LLM+RB) [58]	88.13	76.78	62.73	64.08	66.67	75.80	81.07	73.61
D C [57]	94.43	89.48	84.35	<u>75.10</u>	<u>74.17</u>	<u>89.67</u>	97.69	86.41
TripletCLIP (ours)	94.43	<u>85.53</u>	<u>80.94</u>	69.80	69.82	90.40	86.27	<u>82.46</u>

5.4.3 Zero-shot evaluations

Image-Text Retrieval. In Table 5.5, we summarize the performance of models on text-to-image (T2I) and image-to-text (I2T) retrieval tasks on MSCOCO and Flickr30k datasets, where we report R@5 scores. Remarkably, TripletCLIP significantly outperforms baseline models by an average of **8%/10%** and **8%/12.5%** on I2T and T2I tasks, respectively, on the CC3M and CC12M datasets. Intriguingly, while LaCLIP+HN performs better than NegCLIP, TripletCLIP outstrips both.

Zero-shot Classification. Table 5.5 also presents the average zero-shot classification performance on 18 standard datasets, including ImageNet1k. TripletCLIP consistently enhances top-1 accuracy by an average of **3%** and top-5 accuracy by **5-7%** compared to LaCLIP. Like the retrieval performance, LaCLIP+HN exceeds NegCLIP, yet TripletCLIP maintains the highest performance. Dataset-specific results are in the appendix.

Table 5.9: Finetuning-based evaluations of the methods on Retrieval and ImageNet-1k benchmarks. Bold number indicates the best performance, and underlined number denotes the second-best performance.

Methods	Retrieval				ImageNet 1k	
	Text Retrieval (R@5)		Image Retrieval (R@5)		top-1	top-5
	MSCOCO	Flickr30k	MSCOCO	Flickr30k		
CLIP	74.9	<u>94.60</u>	55.92	83.38	63.31	88.22
CLIP (finetuned)	68.9	88.40	53.50	81.10	49.95	79.16
NegCLIP	66.00	88.60	53.41	81.12	48.85	78.34
Baseline [227]	81.4	96.0	67.49	89.84	<u>61.40</u>	<u>88.10</u>
TSVLC (RB) [58]	71.70	93.00	62.01	87.12	58.81	85.97
TSVLC (LLM+RB) [58]	<u>71.82</u>	92.50	62.24	87.46	59.77	87.02
D C [57]	54.5	79.60	<u>63.51</u>	<u>87.84</u>	51.02	81.22
TripletCLIP (ours)	55.6	82.60	53.32	80.88	45.92	75.54

5.4.4 Finetuning performance

In this paper, we focus on pretraining-based experiments as they allow greater flexibility in learning better representations. To complement this, we also performed additional fine-tuning experiments using hyperparameters similar to the baselines (without LoRA) and compared them against various publicly available baselines [227, 245, 58, 57]. As reported in Table 5.8, `TripletCLIP` improves compositionality and outperforms nearly all baselines. Furthermore, the observed drop in retrieval and zero-shot classification performance (Table 5.9) is attributed to limitations in the vision encoder, highlighting the challenges of existing pre-trained vision encoders in capturing semantic representations. This is further demonstrated in Table 5.10.

Table 5.10: LiT style fine-tuning ablations on SugarCrepe, image-text retrieval, and ImageNet1k. Bold number indicates the best performance.

Models	Train Text	Train Vision	SugarCrepe	Retrieval	ImageNet1k
LaCLIP	✓	×	0.6373	0.5345	31.21%
TripletCLIP (<i>ours</i>)	✓	×	0.6227	0.6817	34.25%
LaCLIP	×	✓	0.5886	0.1134	5.51%
TripletCLIP (<i>ours</i>)	×	✓	0.6923	0.2626	12.51%

5.4.5 Ablations

Can a high-quality filtered dataset improve the performance? Given that negative images in `TripletData` are generated using SDXL-turbo, these may not always be precise. Inspired by DataComp, we employ a pre-trained CLIP-L/14 to filter the image-text pairs, selecting the highest average similarity pairs (positive and negative) individually (i.e., score = $(s(x_i, y_i) + s(x'_i, y'_i))/2$). The top 1.4M positive image-text pairs and their corresponding negatives from `TripletData` are selected. Table 5.7 details this comparison against DataComp pre-trained models. Remarkably, `TripletCLIP` already surpasses baselines without filtered data; however, with the filtered dataset, despite being trained on 50% smaller dataset, `TripletCLIP++` shows further performance improvements. This underlines the significant benefits of carefully selected `TripletData` in enhancing the performance.

Which modality-specific encoder plays the key role in improving compositionality?

To address this open question, we designed an ablation study similar to LiT, freezing either the pre-trained CLIP vision or text encoder while training the opposite modality-specific encoder from scratch. We observe the performance of `LaCLIP` and `TripletCLIP` on CC3M, as shown in Table 5.10. Freezing the vision model results in no performance gain on the SugarCrepe for `TripletCLIP`. However, significant improvements are noted when the vision encoder is actively trained, suggesting that the vision modality may be the bottleneck

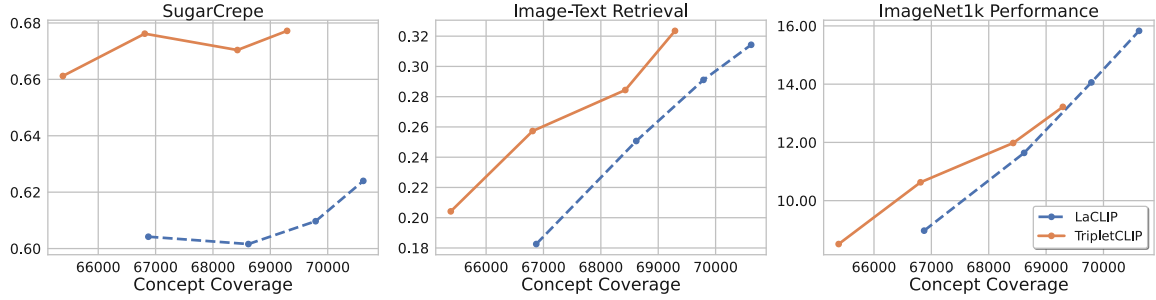


Figure 5.3: Average Results of LaCLIP and TripletCLIP for SugarCrepe Compositions, Image-Text Retrieval, and ImageNet1k over increasing concept diversity.

in compositionality. Notably, TripletCLIP outperforms LaCLIP in all settings, further demonstrating its robustness to different pre-training approaches.

Concept coverage analysis. Improving performance on zero-shot transfer learning tasks such as retrieval involves two key components: adding more concept diversity during training and enhancing image-text alignment/compositionality. We create subsets of CC12M data with increasing concept diversity based on unique WordNet synsets. Specifically, we select 3M, 4M, 5M, and 6M subsets for training LaCLIP, while TripletCLIP training involves only half of these training data as positive pairs, and the rest are corresponding augmentations. Evaluations across SugarCrepe, retrieval tasks, and ImageNet1k (see Figure 5.3) indicate that TripletCLIP not only enhances SugarCrepe performance even at lower concept coverage levels but also significantly outperforms similar concept coverage in retrieval tasks, matching LaCLIP’s performance on zero-shot classification tasks that do not require compositionality at all. This bolsters our argument that incorporating hard negatives from both modalities markedly improves compositional understanding in CLIP, while baseline struggles to do so even with more concept diversity.

Evaluating on a CC3M Winoground-style subset. We evaluated the CC12M pre-trained models on a 50,000 random subset of the CC3M dataset using a Winoground-style approach [262]. As shown in Table 5.11, TripletCLIP significantly improves performance

Table 5.11: TripletData as large-scale composition evaluation dataset after [262].

Methods	Text-Score	Image-Score	Group-Score
CLIP	52.69	29.66	24.64
NegCLIP	54.84	30.42	25.82
NegCLIP++	36.50	30.67	20.11
TripletCLIP (<i>ours</i>)	92.25	66.82	64.30

Table 5.12: Ablation on choice pre-trained LLM. We train NegCLIP++ (ViT-B/32) on negative captions generated from various LLMs and report SugarCrepe, Flickr30k Retrieval (R@5), and ImageNet-top5 performances.

Models	SugarCrepe (avg.)	Retrieval (R@5)	ImageNet1k (top-5)
Gemma-2b-it	56.00	12.60	12.09%
Phi-3-mini-4k-instruct	61.22	13.02	10.94%
Mistral-7b-instruct-v0.2	61.69	10.72	10.52%

compared to the baselines. However, we partially attribute this improvement to spurious correlations learned from the data. At the same time, we note that the models have not fully converged, suggesting minimal risk of overfitting to these spurious correlations.

Choice of Pre-trained LLM. We provide further details regarding the selection of LLMs for generating hard negative captions. NegCLIP++ was trained on 3 million generated negative captions for CC3M using three different LLMs, and the results are reported in Table 5.12. We find that Phi-3 achieves the best average performance, while Gemma-2b unexpectedly has a notable negative impact on compositionality. However, we opted to use Mistral-7b-instruct-v0.2, as Phi-3 was released after our experiments were completed, preventing its earlier evaluation.

Evaluating the Orthogonality of TripletLoss. As discussed in Section 9.2, TripletCLIP can be integrated with various previously proposed methodologies. To evaluate this, we applied

Table 5.13: Ablation on CyCLIP with TripletLoss. To evaluate the compatibility of TripletLoss with the CyCLIP, we train the ViT-B/32 models from scratch on CC3M with 512 batch size and report SugarCrepe, Flickr30k Retrieval (R@5), and ImageNet-top5 performances.

Models	SugarCrepe (avg.)	Retrieval (R@5)	ImageNet1k (top-5)
LaCLIP	55.11	12.80	12.58%
TripletCLIP	65.71	24.62	19.95%
CyCLIP	54.62	11.77	13.01%
CyCLIP+TripletLoss	58.64	20.29	19.05%

TripletLoss in conjunction with CyCLIP [81] and reported the results in Table 5.13. The most significant performance improvement is observed with the TripletLoss alone. However, TripletLoss also enhances the performance over the base CyCLIP, demonstrating its adaptability and orthogonality with respect to other approaches.

Encoder Representation Distribution analysis. Remember, this study aims to learn the representations that can distinguish between two data points that are very similar but semantically different. Firstly, we take LaCLIP and TripletCLIP models trained on CC12M. We also sampled 50000 positive+negative pairs from CC3M. Then, we measure the vision and text modality-specific cosine similarities between positive and negative pairs and plot the distribution (see Figure 5.5). It can be observed that vision representations from TripletCLIP are more skewed towards 0.0, suggesting that the vision encoder can distinguish between hard negative samples better than the baseline LaCLIP. However, in the case of the text modality, both methods perform similarly. This aligns with our findings from Table 5.10 that the vision encoder plays a crucial role in improving the compositionality, and to achieve this, our TripletData is necessary.

5.5 TripletData Analysis

This section provides qualitative examples of the TripletData and discusses various data analyses.

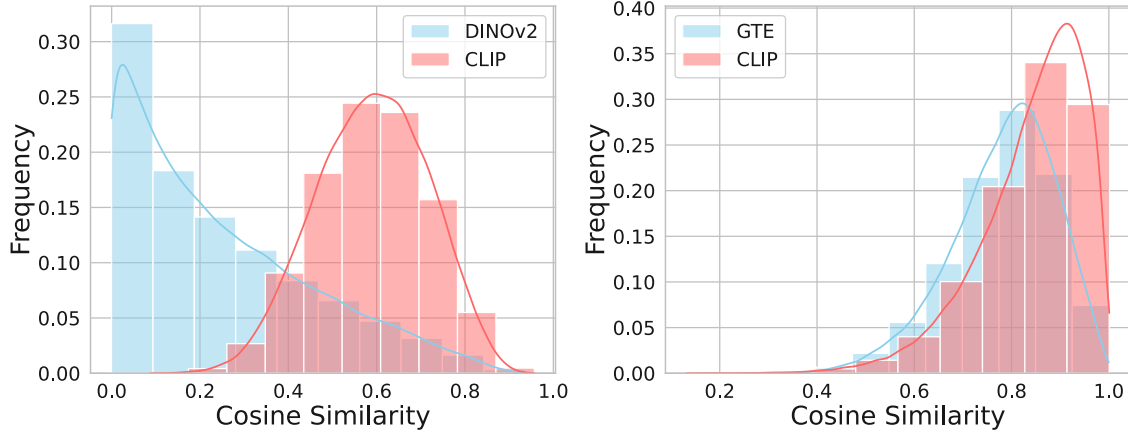


Figure 5.4: Positive vs. Negative modality-specific pair-based similarity distribution of pre-trained CLIP ViT-B/32 model *w.r.t.* the vision and text-only encoders. The left plot is the vision embedding similarities between positive and negative images. The right plot is the text embedding similarities between positive and negative captions. In the ideal scenario, the distribution should be skewed towards 0.0, which indicates that the model can correctly distinguish between the positive and negative data.

Difficulty of the data. We further add one more analysis to investigate how difficult our dataset is. First, we take state-of-the-art language-only (GTE [145]) and vision-only (DINO [27]) embedding models and pretrained CLIP ViT-B/32. Later, we measure the modality-specific similarity between positive and negative vision and language data pairs. Figure 5.4 shows that the similarity distribution of the pretrained CLIP model between positive and negative text pairs follows the distribution of the text-only GTE model. Interestingly, vision distribution is drastically different. DINO can distinguish the positive and negative pairs correctly with high confidence. However, despite the visuals being so different, the pretrained CLIP model struggles to distinguish the different images. This further highlights that `TripletData` is indeed challenging for the vision-language models, even the ones trained on large-scale datasets.

Evaluations of Generated Images. Even though T2I diffusion models are widely evaluated on various tasks. We perform additional evaluations to measure how accurately generated images follow the text prompt. To do this, taking inspiration from [192], we first

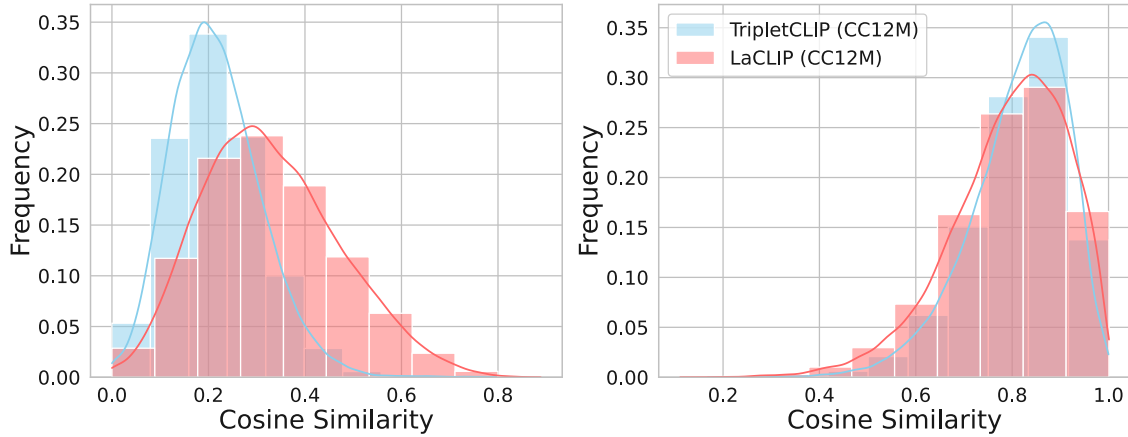


Figure 5.5: Positive vs. Negative modality-specific pair-based similarity distribution of baseline LaCLIP and TripletCLIP. The left plot is the vision embedding similarities between positive and negative images. The right plot is the text embedding similarities between positive and negative captions.

use LLM to generate binary "yes/no" style questions from the given caption. As shown in Table 7.8, we create five questions per hard negative caption. Later, we utilized the ViLT model to answer visual questions. Upon this investigation, we find that SDXL-turbo archives on an average of 76% accuracy. In other words, the T2I model can correctly generate an image that follows around 3/4th of the text. Additionally, we hypothesize that using an improved T2I model or image editing models to generate “hard” negative examples can further improve composition reasoning.

Qualitative Examples. In Figure 5.6, we provide additional qualitative examples of the contrastive positive and “hard” negative pairs from the TripletData. Additionally, in Figure 5.7, we illustrated several examples where the T2I model could not precisely generate images corresponding to the caption. However, we may notice that in most cases, it maintains some of the important aspects. Because of this, despite not being 100% accurate all the time, it can help TripletCLIP improve performance across the evaluation benchmarks.

Hard negative caption only examples:

Table 5.14: Examples of LLM-generated existence-related questions from captions to evaluate the generated images.

Captions	Questions
A smooth, flat sheet of woven fabric (which looks like woven silk) is shown in closeup.	Is the entity a sheet?
	Is the sheet made of fabric?
	Is the fabric woven?
	Does the fabric look like silk?
	Is the fabric flat and smooth?
The sunset over a calm sea.	Is there a sunset?
	Is the sea calm?
	Is the sunset over the sea?
	Is the sunset happening during the day?
	Is the sea rough?
the businesswoman finished second and descended the podium of the runners-up.	Did the businesswoman finish second?
	Did the businesswoman descend the podium?
	Are there runners-up?
	Is the podium for the second place?
	Is the businesswoman a runner?
Kid standing still near a scooter in a public place.	Is there a kid in the public place?
	Is the kid near a scooter?
	Is the scooter in a public place?
	Is the kid standing still?
	Is there a public place in the caption?
A local resident drives in the fall on a scenic drive, wearing a red scarf.	Is the local resident driving?
	Is it happening in the fall?
	Is there a scenic drive?
	Is the local resident wearing a red scarf?
	Is the local resident walking?

1. **Raw Caption:** dog looking out from a window .

Language Rewrite: A dog looking through the window at his owner.

Negative Caption (NegCLIP):

(a) A window looking through the dog at his owner.

(b) A dog looking through the window at his dog.

(c) A dog screams through the window at his owner.

TripletData Negative Caption: A cat observing its owner from the window.

2. **Raw Caption:** person attends the premiere of film

Language Rewrite: A person attends the premiere of film

Negative Caption (NegCLIP):

(a) A premiere attends the person of film

(b) A person attends the festival of film

(c) A person watches the premiere of film

TripletData Negative Caption: A person waits in line for film tickets.

3. **Raw Caption:** white crocus spring flowers in the forest.

Language Rewrite: white crocus flowers in the forest

Negative Caption (NegCLIP):

(a) white crocus flowers in the sky

TripletData Negative Caption: red orchid flowers in the meadow.

4. **Raw Caption:** flag with industry in the background

Language Rewrite: A flag is holding in the background an industrial site.

Negative Caption (NegCLIP):

(a) A background is holding in the flag an industrial site.

(b) A flag is holding in the background an earthquake site.

(c) A drone is holding in the background an industrial site.

(d) A flag is visible in the background an industrial site.

TripletData Negative Caption: A flag is flapping in the foreground of a pastoral scene.

5. **Raw Caption:** portrait of businessman with cardboard on his head carrying a briefcase and using an umbrella while standing by.

Language Rewrite: A portrait of a businessman standing by with a briefcase and cardboard on his head carrying an umbrella while looking at a blue sky and parked cars on the street

Negative Caption (NegCLIP):

- (a) A businessman of a portrait standing by with a briefcase and cardboard on his head carrying an umbrella while looking at a blue sky and parked cars on the street
- (b) A portrait of a businessman standing by with a briefcase and cardboard on his head carrying an umbrella while looking at a grey sky and parked cars on the street
- (c) A portrait of a businessman standing by with a briefcase and cardboard on his head carrying an umbrella while looking at a blue truck and parked cars on the street
- (d) A portrait of a businessman standing by with a briefcase and cardboard on his head carrying an umbrella while pointing at a blue sky and parked cars on the street

TripletData Negative Caption: A portrait of a businesswoman seated on a bench with a tote bag and newspaper on her lap holding an umbrella while looking at a red sunset over row houses.

6. **Raw Caption:** 158834 is the portion of the bound train .

Language Rewrite: Huge locomotives sit on the tracks in front of a building.

Negative Caption (NegCLIP):

- (a) Huge tracks sit on the locomotives in front of a building.
- (b) Three locomotives sit on the tracks in front of a building.
- (c) Huge locomotives sit on the tracks in anticipation of a building.
- (d) Huge locomotives mounted on the tracks in front of a building.

TripletData Negative Caption: Huge locomotives sit on the tracks in front of a bridge.

7. **Raw Caption:** biological subfamily eating fish on a seaweed covered shore

Language Rewrite: Two blue whales are eating salmon on a beach surrounded by seaweed

Negative Caption (NegCLIP):

- (a) Two blue salmon are eating whales on a beach surrounded by seaweed
- (b) Two stranded whales are eating salmon on a beach surrounded by seaweed
- (c) Two blue whales are eating salmon on a farm surrounded by seaweed
- (d) Two blue whales are eating salmon on a beach covered by seaweed

TripletData Negative Caption: Two blue whales are feeding on herring in a bay surrounded by kelp.

8. **Raw Caption:** image of an original oil painting on canvas

Language Rewrite: A young lady holding a painting to your face so you can see the detail of the painting

Negative Caption (NegCLIP):

Table 5.15: Composition evaluations of the methods on various benchmarks. Bold number indicates the best performance, and underlined number denotes the second-best performance.

Methods		Valse	Cola			Winoground			SugarCrepe	Overall
			Txt2Img	Img2Txt	Group	Txt2Img	Img2Txt	Group		
CC3M	LaCLIP	43.19	28.10	<u>13.33</u>	<u>10.48</u>	24.75	<u>9.25</u>	<u>6.00</u>	54.09	23.65
	LaCLIP + HN	47.91	20.95	5.54	2.38	23.25	4.25	3.00	53.92	20.15
	NegCLIP	48.06	21.42	14.76	8.57	21.00	8.50	5.25	57.18	23.09
	NegCLIP++ (ours)	43.65	<u>29.05</u>	<u>13.33</u>	7.14	22.00	5.50	3.25	61.69	23.20
	TripletCLIP (ours)	<u>48.36</u>	31.43	<u>13.33</u>	9.52	<u>24.25</u>	6.25	4.25	<u>63.49</u>	<u>25.11</u>
	TripletCLIP++ (ours)	48.53	31.43	15.71	12.86	22.25	9.50	6.75	66.09	26.64
Performance Gain w.r.t. LaCLIP		5.34%	3.33%	2.38%	2.38%	-2.50	0.25%	0.75%	12.00	2.99%
CC12M	LaCLIP	55.69	20.95	<u>15.71</u>	<u>7.62</u>	26.25	8.00	6.25	67.21	25.96
	NegCLIP	58.59	27.14	15.24	5.71	18.25	6.50	4.25	68.41	25.51
	NegCLIP++ (ours)	<u>58.12</u>	33.33	11.43	7.14	<u>23.50</u>	<u>7.75</u>	<u>5.50</u>	<u>73.05</u>	<u>27.48</u>
	TripletCLIP (ours)	57.57	<u>27.62</u>	19.53	11.43	23.25	6.25	4.25	74.55	28.06
	Performance Gain w.r.t. LaCLIP		1.88%	6.67%	3.82%	3.81%	-3.00	-1.75	-2.00	7.35

(a) A young painting holding a lady to your face so you can see the detail of the lady

(b) A bearded lady holding a painting to your face so you can see the detail of the painting

(c) A young lady holding a pencil to your face so you can see the detail of the painting

(d) A young lady holding a painting to your face so you can enjoy the detail of the painting

TripletData Negative Caption: A young lady holding a painting away from her face to show its beauty to the audience.

Table 5.16: Dataset-specific zero-shot classification results. Bold number indicates the best performance, and underlined number denotes the second-best performance.

	Models	fgvc aircraft	imagenet1k	caltech101	cifar100	clevr closest object distance	clevr count all	dmlab	dsprites label orientation	dsprites label x position	dtd	eurosat	flowers	kititi closest vehicle distance	peam	pets	smallnorb label azimuth	smallnorb label elevation	svhn	vg.
CC3M	LaCLIP	0.84	3.79	24.65	8.04	15.86	11.26	19.02	2.63	3.14	5.53	8.56	3.43	<u>26.72</u>	48.16	3.11	<u>5.75</u>	10.86	6.68	11.56
	LaCLIP+HN	<u>1.23</u>	5.75	33.20	8.84	17.18	11.18	16.60	2.52	3.10	6.65	10.02	4.29	18.85	51.93	4.44	4.92	11.12	<u>9.78</u>	<u>12.31</u>
	NegCLIP	1.02	4.67	29.44	8.17	22.37	12.97	18.81	2.57	3.19	6.81	7.39	3.77	28.83	41.72	4.17	4.84	11.37	8.32	12.25
	NegCLIP++ (ours)	1.02	3.84	20.61	6.44	<u>22.03</u>	10.01	<u>18.91</u>	2.60	<u>3.21</u>	4.63	11.70	3.59	23.91	49.34	4.74	5.43	10.02	7.70	11.65
	TripletCLIP (ours)	1.29	<u>7.32</u>	<u>37.05</u>	14.37	15.55	<u>13.49</u>	15.37	2.62	2.98	<u>8.78</u>	<u>22.74</u>	<u>5.97</u>	16.74	<u>53.49</u>	<u>5.81</u>	5.34	<u>11.51</u>	10.22	<u>12.31</u>
	TripletCLIP++ (ours)	0.90	8.85	38.54	<u>11.83</u>	20.72	13.79	17.46	2.52	3.58	8.99	24.43	8.20	19.41	56.60	19.38	5.76	12.46	7.91	15.62
CC12M	LaCLIP	1.59	19.72	62.98	23.62	11.99	11.38	19.85	<u>2.91</u>	3.20	13.03	23.85	13.43	<u>24.05</u>	<u>50.38</u>	34.72	6.14	11.04	9.54	19.08
	NegCLIP	<u>1.56</u>	<u>20.22</u>	63.45	<u>26.02</u>	18.89	<u>15.80</u>	16.20	2.45	3.13	<u>14.26</u>	17.43	13.55	22.78	50.22	30.85	5.37	12.04	10.04	<u>19.12</u>
	NegCLIP++ (ours)	1.38	19.06	<u>64.66</u>	24.86	15.96	14.87	15.62	2.94	3.07	13.62	14.33	<u>14.30</u>	16.60	49.98	31.92	5.42	<u>12.08</u>	<u>12.13</u>	18.48
	TripletCLIP (ours)	1.29	23.31	65.34	30.30	<u>16.26</u>	17.67	<u>16.58</u>	2.80	<u>3.18</u>	17.45	<u>23.26</u>	15.43	25.74	51.04	<u>33.25</u>	<u>5.58</u>	12.38	13.67	20.81

5.6 Detailed Results

5.6.1 Compositional reasoning

Previously, we reported the results on SugarCrepes, the most challenging dataset, noted for its absence of language biases. However, evaluations were also conducted on other benchmarks, such as Valse, Cola, and Winoground. As indicated in Table 5.15, TripletCLIP achieves overall improvements of **2-3%** compared to LaCLIP and NegCLIP. The Valse benchmark, which contains text prompts that heavily favor the perturbations made for NegCLIP, shows a strong performance from NegCLIP, while NegCLIP++ encounters difficulties. Interestingly, TripletCLIP faces challenges in maintaining performance on Winoground, and baseline LaCLIP maintains the SOTA, which is counterintuitive to other benchmarks. Nonetheless, TripletCLIP still manages to outperform NegCLIP significantly. These results affirm that TripletCLIP sets a new standard for state-of-the-art compositional reasoning across diverse benchmarks.

Table 5.17: Composition evaluations of the methods on various benchmarks. Bold number indicates the best performance, and underlined number denotes the second-best performance.

		Text-to-Image Retrieval						Image-to-Text Retrieval					
Methods		MSCOCO			Flickr30k			MSCOCO			Flickr30k		
		R@1	R@5	R@10	R@1	R@5	R@10	R@1	R@5	R@10	R@1	R@5	R@10
CC3M	LaCLIP	1.56	5.06	8.90	3.70	10.90	16.00	1.73	5.97	9.50	3.54	10.84	16.02
	LaCLIP + HN	<u>2.60</u>	<u>8.08</u>	<u>13.02</u>	<u>6.30</u>	<u>16.10</u>	<u>22.40</u>	<u>2.66</u>	<u>8.64</u>	<u>13.38</u>	<u>5.98</u>	<u>16.64</u>	<u>23.82</u>
	NegCLIP	1.74	6.32	10.44	4.90	13.80	19.60	1.95	6.61	10.62	4.76	12.96	18.50
	NegCLIP*	1.50	5.80	9.70	4.40	11.20	16.30	1.85	6.19	9.82	3.50	10.24	15.20
	TripletCLIP	3.16	10.38	16.22	9.10	22.00	29.80	3.58	11.28	17.39	8.38	22.00	29.56
Performance Gain vs. CLIP		1.60%	5.32%	7.32%	5.40%	11.10%	13.80%	1.85%	5.31%	7.89%	4.84%	11.16%	13.54%
CC12M	LaCLIP	10.50	25.86	35.60	21.30	42.70	54.60	7.21	19.78	28.37	15.06	36.30	47.00
	NegCLIP	<u>12.32</u>	<u>30.16</u>	<u>41.44</u>	<u>24.70</u>	<u>46.60</u>	<u>58.20</u>	8.56	<u>23.11</u>	<u>32.60</u>	<u>18.66</u>	41.70	53.42
	NegCLIP*	10.94	26.96	36.64	18.70	43.90	55.90	<u>8.80</u>	22.69	32.13	18.24	<u>42.86</u>	<u>53.78</u>
	TripletCLIP	14.60	33.00	43.84	28.00	55.90	65.70	11.38	28.50	39.04	25.28	52.38	63.32
	Performance Gain vs. CLIP	4.10%	7.14%	8.24%	6.70%	13.20%	11.10%	4.17%	8.72%	10.67%	10.22%	16.08%	16.32%

5.6.2 Dataset-specific zero-shot classification

Table 5.16 provides fine-grained results for the 18 zero-shot classification datasets. It can be observed that TripletCLIP consistently outperforms the baselines, achieving the best average results across these challenging datasets. Although the improvements are marginal, they are in line with expectations. As discussed in Figure 5.3, TripletCLIP does not introduce new concepts into the training data but focuses on augmentations that enhance representation without increasing concept diversity. These enhanced representations from TripletCLIP lead to average improvements of **1-3%** depending on the scenario.

5.6.3 Detailed image-text retrieval performance

Table 5.17 presents detailed T2I and I2T retrieval results for the MSCOCO and Flickr30k datasets. We report results at different recall thresholds: R@1, R@5, and R@10. The data shows that TripletCLIP significantly outperforms the baselines across all recall rates. On average, TripletCLIP achieves a performance gain of **7-11%** over LaCLIP. Additionally, TripletCLIP improves performance by **3-6%** compared to previous state-of-the-art

baselines.

5.7 Conclusion

In this work, we introduce `TripletCLIP`, a novel approach to enhancing compositional reasoning in vision-language models through the strategic incorporation of hard negative image-text pairs. Our comprehensive experiments across a suite of benchmarks demonstrate that `TripletCLIP` significantly outperforms existing methodologies such as `LaCLIP` and `NegCLIP`, achieving notable gains not only in compositionality but also in zero-shot classification and retrieval tasks as well. Further, our ablation studies highlight the critical role of modality-specific training and the careful curation of training data, underscoring the importance of both hard negative image and text components in the learning process. `TripletCLIP`'s effectiveness with a smaller, refined dataset suggests a promising direction for future research—maximizing performance without the need for extensive data collection, thereby reducing computational costs and enhancing model efficiency. To this end, we provide an intriguing application of synthetic datasets via hard negative image-text pairs for vision-language tasks that could be easily extended to improve Multimodal Large Language Models and Text-to-Image generative models.

Limitations. Due to constraints inherent in academic settings and limited computational resources, we were unable to scale `TripletCLIP` to handle hundreds of millions of image-text pairs or employ larger models within the scope of this study. Nevertheless, our results indicate a promising direction for future research within a consistent experimental framework, and we encourage subsequent work to explore scaling both the `TripletData` and `TripletCLIP`. Our experimental focus was primarily on the CLIP and LiT methodologies. With additional resources, however, extending our methodologies to more advanced contrastive learning techniques, such as SigLIP, would be feasible. In conclusion, our work

introduces a compelling strategy for integrating open-ended hard negatives (both text and image) during the pre-training phase, providing a methodology and large-scale data that could benefit a variety of research domains.

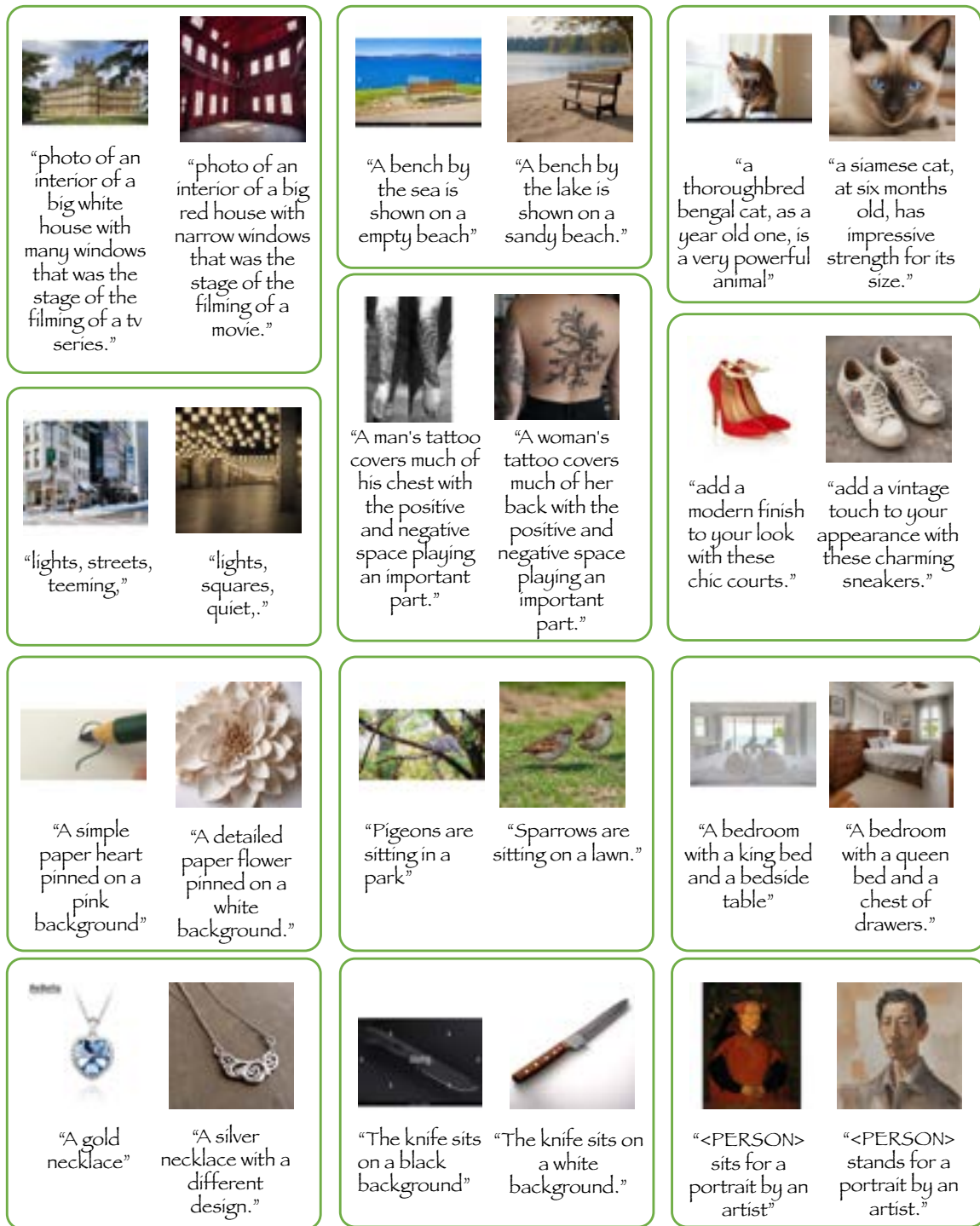


Figure 5.6: Qualitative examples of positive and hard negative image-text pairs from TripletData. In each block, left image-text pairs are positive images from CC3M, and right pairs are corresponding negatives from TripletData.

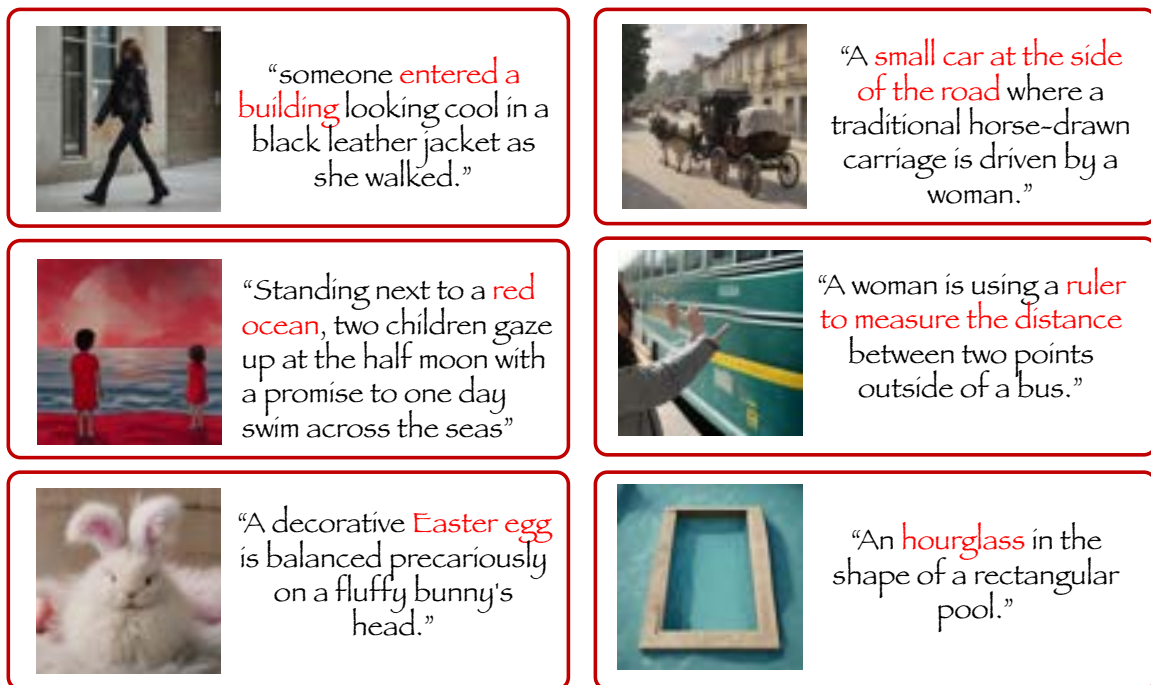


Figure 5.7: Examples of T2I failures.

STEERING OF RECTIFIED FLOW MODELS FOR CONTROLLED GENERATIONS

6.1 Introduction

Recent advances in diffusion models have led to rapid progress in AI-generated content (AIGC), particularly in text-to-image (T2I) and text-to-video (T2V) models across various domains such as entertainment, arts, and design [216, 226, 61, 207, 193, 243]. These developments have resulted in remarkable performance in controlled generation tasks such as image editing, solving inverse problems, and personalization. Existing methods focus on training the conditional model [52, 97], defining guidance strategies [44, 88], and performing an inversion [15, 23]. Despite this rich landscape, the recent emergence of flow-based methods [147], particularly RFMs [153, 136], remains underexplored, presenting a

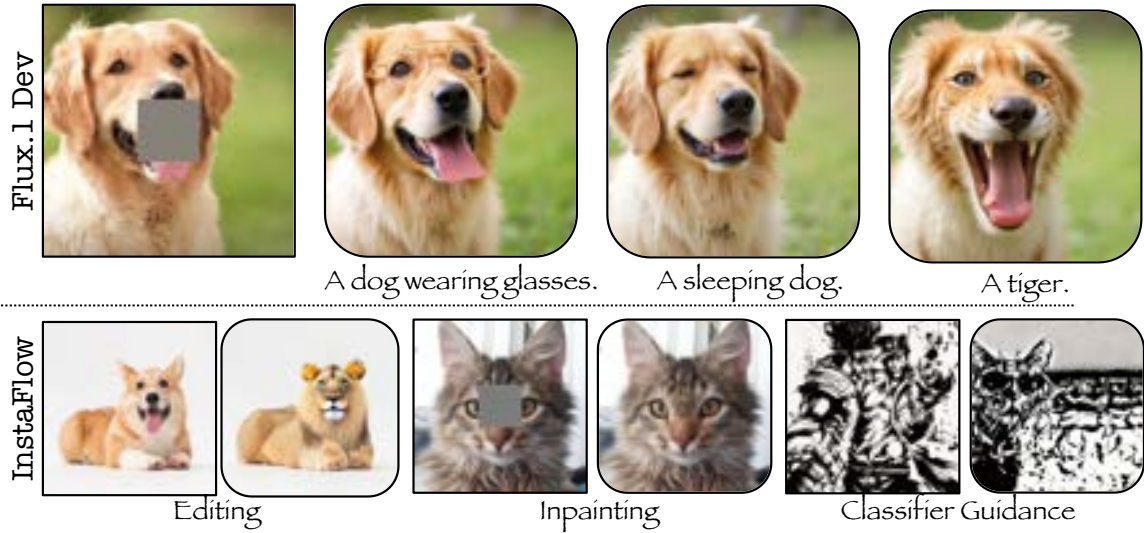


Figure 6.1: FlowChef steers the trajectory of Rectified Flow Models during inference to tackle linear inverse problems, image editing, and classifier guidance. We extend FlowChef to SOTA models like Flux and InstaFlow, enabling gradient- and inversion-free control for efficient, controlled image generation.

promising yet nascent frontier for investigation.

Although performing very well, existing approaches for solving inverse problems often require minutes of computation and additional memory overhead [253, 44, 221, 249, 15]. Image editing methods typically involve either inversion or explicit training [106, 23, 24] (including flow-based concurrent works). These limitations can be attributed to the inherent stochasticity of DMs (see Figure 6.2(a)) and backpropagation through ODESolver for RFMs, often requiring a higher number of function evaluations (NFEs) and computing budget. As a result, they cannot be extended to large state-of-the-art models like Flux or SD3 [61].

The goal of this paper is to advance our theoretical understanding of vector field dynamics of RFMs that can allow us to solve inverse problems and editing tasks alike without needing intensive computing resources while maintaining SOTA-like performance. Specifically, we want to answer the question: `does rectified flow ODEs bring any new perspectives to the table for controlled generations?`

In this paper, in our attempts to understand the dependencies on backpropagation, we first revisit the ODEs that govern these models and derive the error dynamics. In the case of DMs, an error term emerges that can hinder the convergence due to inaccuracies in estimating the denoised samples or improper gradient approximations. Contrary to diffusion models, rectified flow models exhibit straight trajectories with smooth vector fields and a deterministic nature due to their linear interpolation between noise and data distributions (see Figure 6.2(b)).

By leveraging these properties, we derive the gradient relationship of the guidance term with and without backpropagation through ODESolver. Specifically, in RFMs, the straight trajectories enable the gradient at the marginal distribution to be expressed as an affine function of the gradient at the posterior through simple scaling transformation (see Figure 6.2(c)). Based on this critical finding, we present **FlowChef**, allowing efficient steering of the generative process without costly backpropagation, unlocking a new paradigm

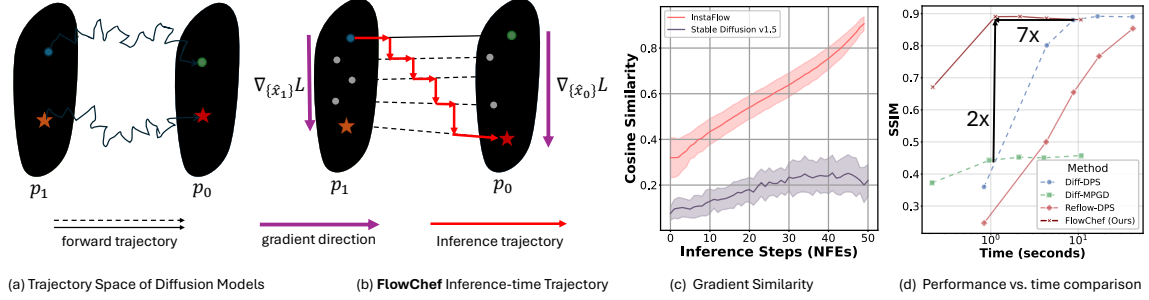


Figure 6.2: Motivation behind FlowChef based on rectified flow models’ trajectory space. Let $p_1 \sim N(0, I)$ and p_0 be noise and posterior distributions. (a) Stochasticity and nonlinear trajectories can complicate gradient estimation at each denoising step t . (b) Our method FlowChef enables efficient trajectory steering to guide x_t along the trajectory towards x_0^{ref} . (c) As $t \rightarrow 0$, the cosine similarity of RFMs gradients, with or without ODESolver backpropagation, increases, in contrast to the erratic gradients of DMs. (d) FlowChef improves the latency by 7x and performance by 2x compared to the prior diffusion SOTA baselines on inverse problems using pixel models.



Figure 6.3: Illustration of impact of number guided control steps on baseline MPGD (DMs) and proposed FlowChef (RFMs) with equal guidance values: $\nabla_{\hat{x}_0} \mathcal{L} = (\hat{x}_0 - x_0^{ref})/N$. FlowChef deterministically steers RFMs, achieving high-quality results with only five guided steps. In contrast, MPGD, applied to stochastic DMs, exhibits reduced performance due to inherent randomness.

for controlled image generations.

We conduct extensive evaluations of **FlowChef** on linear inverse problems on pixel and latent space RFMs. We also present the **FlowChef-Edit** variant designed explicitly for inversion-free image editing. Our results demonstrate that **FlowChef** not only surpasses baseline methods but does so with greater computational efficiency and without the need for inversion (see Figure 6.2(d)). As illustrated in Figure 7.2, **FlowChef** addresses a variety of tasks. For perspective, **FlowChef** handles the linear inverse problems within 18 seconds on the latent-space model, while SOTA takes 1-3 minutes per image. Furthermore, we explore its practical applicability to large-scale models (i.e., Flux) to tackle both linear inverse

problems and image editing together without inversion and within 30 NFEs at billions of parameter scales. At last, we extend **FlowChef** on other downstream tasks such as classifier-guided style transfer, 3D generation, and multi-concept image editing. Our key contributions can be summarized as follows:

- We develop a theoretical and empirical understanding of the vector field of rectified flow models for a guided, controlled generation.
- We introduce **FlowChef**, the most efficient method to date for guided, controlled generation using RFMs, achieving near state-of-the-art performance without requiring inversion or gradient backpropagation.
- We demonstrate **FlowChef**'s superior performance across multiple tasks, including linear inverse problems, image editing evaluated on the PIE benchmark [106], and classifier guidance.

6.2 Related Works

Inverse Problems. This task addresses training-free approaches for solving inverse problems such as in-painting, super-resolution, Gaussian de-blurring etc [47]. Much of the current literature focuses on diffusion models, particularly pixel-space models [47, 253, 44, 286]. However, these models face challenges when scaled to latent-space models, as they are incompatible with off-the-shelf pretrained models and require backpropagation through ODESolvers, which can take at least three minutes per image for satisfactory results [47, 221, 249, 218]. MPGD [88] attempts to avoid backpropagation utilizing DDIM [252] property, but limitations persist, especially with large-scale models. RB-Modulation [220] extends MPGD for style transfer and observes similar challenges.

Recent work has extended these approaches to ODEs (e.g., OT-ODE) and flow models [206]. D-Flow [15], for instance, optimizes initial noise by differentiating through the

full trajectory chain; however, this comes with significant resource demands and is not adaptable to state-of-the-art (SOTA) models like Flux or SD3 [61]. PnPFlow is an inversion and gradient-free method for inverse problems, but it leads to over-smoothed results and lacks extensibility to image editing. In this work, we propose **FlowChef**, which addresses linear inverse problems in a gradient and inversion-free manner while achieving SOTA performance w.r.t. flow-based methods.

Image Editing. Diffusion-based approaches dominate image editing [101], but they rely heavily on accurate inversion [106, 23, 178, 102]. Although inversion-free diffusion methods are faster, they often lack in edit quality [290, 46, 175, 285]. Despite RFMs being SOTA in text-to-image (T2I) generation, they still lack robust editing capabilities. iRDS [294] and RF-Inversion [219] presents an inversion strategy for RFMs, but they significantly lack quality and control. Similarly, RectifID [258] offers an optimization-based approach to modify the entire trajectory for personalized T2I generation but performs poorly on inverse problems.

To the best of our knowledge, we present the first comprehensive solution that enhances RFMs for solving linear-inverse problems, image editing and extends beyond it without significant computational or time overhead.

6.3 Preliminaries

Classifier guidance, inversion problems, and image editing involve guiding a model toward a specific target sample or distribution in both pixel and latent spaces. However, these tasks are often treated separately in the literature. Here, we present a unified problem formulation to encompass these downstream tasks, with a focus on rectified flow models.

6.3.1 Problem Formulation

Let $u_\theta : \mathcal{R}^d \times [0, T] \rightarrow \mathcal{R}^d$ represent a pretrained flow model estimating the drift $v = x_1 - x_0$ from x_t . The denoised sample $\hat{x}_0$ is obtained by integrating the drift u_θ over time from $t = T$ to $t = 0$, starting from $x_T \sim p_1$. With a target sample x_0^{ref} , we define a cost function $\mathcal{L} : \mathcal{R}^d \times \mathcal{R}^d \rightarrow \mathcal{R}_+$ that quantifies the cost of aligning $\hat{x}_0$ with x_0^{ref} , yielding the optimization problem:

$$\min_{\{\hat{x}_t\}_{t=0}^T} \mathcal{L}(\hat{x}_t, x_0^{ref}), \quad (6.1)$$

where $\{\hat{x}_t\}_{t=0}^T$ represents the model-generated trajectory from x_T to x_0 . The objective is to find the trajectory that minimizes $\mathcal{L}$, effectively steering the generated sample toward the target. This can be adapted for the denoising stage with either a noise-aware cost function at each timestep t or by estimating x_0 to refine the trajectory as needed. The gradient update is given by:

$$\hat{x}_t \leftarrow x_t - s \cdot \nabla_{x_t} \mathcal{L}(\hat{x}_0, x_0^{ref}), \quad (6.2)$$

where s is guidance scale. This process requires estimating $\hat{x}_0$, backpropagating gradients through ODESolver (u_θ) to adjust x_t , and iteratively refining x_{t-1} .

Prior methods like DPS [44], FreeDoM [301], and Universal Guidance [13] share a common update rule for guided sampling. In contrast, MPGD and RB-Modulation optimize in the denoised space ($\hat{x}_0$) using DDIM to compute x_{t-1} , avoiding backpropagation of the ODE solver. However, these approaches struggle to accurately estimate $\hat{x}_0$ due to stochastic sampling and curved trajectories in DMs.

6.3.2 Cost Functions

Notably, explicit x_0^{ref} is unnecessary and can be approximated with appropriate cost functions depending on the downstream tasks. Assuming initial Gaussian noise x_T leads to $\hat{x}_0$, the cost function can be defined as:

$$\mathcal{L}(\hat{x}_0, x_0^{ref}) = \|\hat{x}_0 - x_0^{ref}\|_2^2. \quad (6.3)$$

In inverse problems, let $\mathcal{F} : \mathcal{R}^d \rightarrow \mathcal{R}^n$ represent a degradation operation (e.g., downsampling for super-resolution). We then define:

$$\mathcal{L}(\hat{x}_0, x_0^{ref}) = \|\mathcal{F}(\hat{x}_0) - x_0^{ref}\|_2^2. \quad (6.4)$$

Here, x_0^{ref} is a degraded sample, and we guide the model to generate $\hat{x}_0$ such that its degraded version matches x_0^{ref} . For classifier guidance, the cost function can be based on the negative log-likelihood (NLL). Specifically, given a classifier $p_\phi(c|\hat{x}_0)$, the cost function is:

$$\mathcal{L}(\hat{x}_0, c) = -\log p_\phi(c|\hat{x}_0). \quad (6.5)$$

Remark 1. Although presented in pixel space, this formulation extends to latent space by introducing a Variational Autoencoder (VAE) encoder ($\mathcal{E}$) and decoder ($\mathcal{D}$).

Remark 2. Image editing task needs to find the good balance between background preservation and editing. This can be formulated with the help of Eq. 6.4, where $\mathcal{F}$ is the mask to preserve the background.

6.4 Proposed Method

In this section, we introduce our method, **FlowChef**, which enables inference-time steering for rectified flow models by presenting an efficient gradient approximation during

guided sampling. We begin by analyzing the error dynamics of general ordinary differential equations (ODEs) and then explain how the inherent properties of rectified flow models mitigate existing approximation issues. Building on these insights, we derive **FlowChef** objective.

6.4.1 Error Dynamics of the ODEs

Understanding why existing methods often fail and require computationally intensive strategies is crucial. In ODE-based generative models, guiding the sampling process toward a desired target typically involves computing the gradient of a loss function with respect to the model’s parameters or state variables. As noted in Eq. (6.2), even though the denoised output can be estimated using $\hat{x}_0 \leftarrow \text{Sample}(x_t, u_\theta(x_t, t))$, backpropagation through the ODE solver is still necessary to obtain $\nabla_{x_t} \mathcal{L}$. This raises the question: *Why is backpropagation through the ODE solver necessary?*

Approximating gradient computations is a common approach to reduce computational overhead [88, 253]. However, in models governed by nonlinear ODEs, unregulated gradient approximations can introduce significant errors into the system dynamics. This issue is formalized in the following proposition:

Proposition 6.4.1. *Let $p_1 \sim \mathcal{N}(0, \mathcal{I})$ be the noise distribution and p_0 be the data distribution. Let x_t denote an intermediate sample obtained from a predefined forward function q as $x_t = q(x_0, x_1, t)$, where $x_0 \sim p_0$ and $x_1 \sim p_1$. Define an ODE sampling process $dx(t) = f(x_t, t)dt$ and quadratic $\mathcal{L} = \|\hat{x}_0 - x_0^{ref}\|_2^2$, where $f : \mathcal{R}^d \times [0, T] \rightarrow \mathcal{R}^d$ is an ODESolver. Then, the error dynamics of ODEs for controlled image generation is governed by:*

$$\frac{dE(t)}{dt} = -4sE(t) + 2e(t)^T \epsilon(t),$$

where $e(t) = \hat{x}_0 - x_0^{ref}$, $E(t) = e(t)^T e(t)$ is the squared error magnitude, $s > 0$ is the

guidance strength, and $\epsilon(t)$ represents the accumulated errors due to potential stochasticity and non-linearity.

Proof. Consider the sampling process described by the ODE:

$$\frac{dx(t)}{dt} = f(x(t), t), \quad (6.6)$$

where $f(x(t), t)$ is a nonlinear function often parameterized via neural network θ . To guide the sampling process toward minimizing a loss function $\mathcal{L}(\hat{x}_0, x_0^{ref})$, we can adjust the dynamics by adding the gradient ∇_{x_t} to the vector field (see Eq. 6.2) as:

$$\frac{dx(t)}{dt} = f(x(t), t) - s \cdot \nabla_{x_t} \mathcal{L}(\hat{x}_0, x_0^{ref}), \quad (6.7)$$

where s is the guidance strength. Let $e(t) = \hat{x}_0 - x_0^{ref}$ be the error between the estimated and target samples. Since $\hat{x}_0(t) = x(t) + \int_t^0 f(x(\tau), \tau) d\tau$, differentiating $e(t)$ with respect to t yields:

$$\frac{de(t)}{dt} = \frac{d\hat{x}_0(t)}{dt} \quad (6.8)$$

$$= \frac{dx(t)}{dt} - f(x(t), t) \quad (6.9)$$

$$= -s \cdot \nabla_{x_t} \mathcal{L}(\hat{x}_0, x_0^{ref}). \quad (6.10)$$

However, this requires the compute-intensive backpropagation through ODESolver. Therefore, it is important to find an approximation of ∇_{x_t} . And the most convenient approximation is: $\nabla_{x_t} \approx \nabla_{\hat{x}_0}$. However, this derivation assumes that the integral $\int_t^0 f(x(\tau), \tau) d\tau$ is well-behaved and that $\hat{x}_0(t)$ depends smoothly on $x(t)$. In the presence of nonlinearity and stochasticity, small changes in $x(t)$ can lead to disproportionately large changes in $\hat{x}_0(t)$, due to the sensitivity of the integral to the path taken. Moreover, potential stochasticity in the trajectory mean that the mapping from $x(t)$ to $\hat{x}_0(t)$ is not injective; different trajectories

$x(t)$ may lead to the same $\hat{x}_0(t)$ or vice versa. This non-unique mapping complicates the error dynamics because $\nabla_{\hat{x}_0} \mathcal{L}$ may not provide a consistent or effective direction for updating $x(t)$. Including the effects of nonlinearity and stochasticity, the error dynamics become:

$$\frac{de(t)}{dt} = -s \cdot \nabla_{\hat{x}_0} \mathcal{L}(\hat{x}_0, x_0^{\text{ref}}) + \epsilon(t), \quad (6.11)$$

where $\epsilon(t)$ represents the errors introduced by the nonlinearity in $f(x(t), t)$ and the sensitivity of $\hat{x}_0$ to $x(t)$ due to potential stochasticity. In other words, the approximation error $\epsilon(t)$ can be represented as:

$$\epsilon(t) = s \cdot \left(\nabla_{x_t} \mathcal{L}(\hat{x}_0, x_0^{\text{ref}}) - \nabla_{\hat{x}_0} \mathcal{L}(\hat{x}_0, x_0^{\text{ref}}) \right). \quad (6.12)$$

Assuming a quadratic loss function $\mathcal{L} = \|\hat{x}_0 - x_0^{\text{ref}}\|_2^2$, we have $\nabla_{\hat{x}_0} \mathcal{L} = 2e(t)$, leading to:

$$\frac{de(t)}{dt} = -2se(t) + \epsilon(t). \quad (6.13)$$

To understand the convergence of the error, we analyze the evolution of the error magnitude $E(t) = e(t)^T e(t)$. Differentiating $E(t)$ with respect to time t , we get:

$$\frac{dE(t)}{dt} = \frac{d}{dt} (e(t)^T e(t)) \quad (6.14)$$

$$= 2e(t)^T \frac{de(t)}{dt} \quad (6.15)$$

$$= 2e(t)^T (-2se(t) + \epsilon(t)) \quad (6.16)$$

$$= -4se(t)^T e(t) + 2e(t)^T \epsilon(t) \quad (6.17)$$

$$= -4sE(t) + 2e(t)^T \epsilon(t). \quad (6.18)$$

This completes the proof. □

Notably, we derive this behavior of the ODE processes under the assumption that the error rate cannot be calculated accurately. This can either come from the incorrect estimation of $\hat{x}_0$ or the nonlinearity of ODESolver itself. In the next section, we further concretize this with respect to the RFMs.

The term $-4sE(t)$ denotes the exponential decay of error due to guidance, while $2e(t)^\top \epsilon(t)$ captures the impact of non-linearity and stochasticity. In diffusion models, stochasticity and curved sampling trajectories lead to larger $\epsilon(t)$, hindering convergence. Although outside the scope of this paper, in Appendix 6.5.2, we show that the error dynamics do not strictly converge for diffusion models. In contrast, rectified flow models exhibit straight trajectories, causing $\epsilon(t)$ to approach zero and allowing the error to decrease exponentially. Figure 6.2(d) & 6.3 compares the two gradient-free guided sampling strategies in DMs and RFMs.

6.4.2 **Flow hef**: Steering Within the Vector Field

Rectified flow models inherently allow error dynamics to converge even with gradient approximations due to their straight-line trajectories and smooth vector fields, as discussed previously. Hence, vector field $u_\theta(x_t, t)$ is trained to be smooth, and this smoothness implies that u_θ changes gradually *w.r.t.* x_t . We formalize our approach with the following assumptions about the Jacobian of the vector field:

ssumption 1 (Local Linearity): Within the small neighborhoods around any point x_t along the sampling trajectory, the vector field $u_\theta(x_t, t)$ behaves approximately linearly with respect to x_t . Doing Taylor series expansion for small perturbations δ , we get:

$$u_\theta(x_t + \delta, t) \approx u_\theta(x_t, t) + J_{u_\theta}(x_t, t)\delta, \quad (6.19)$$

where $J_{u_\theta}(x_t, t) = \frac{du_\theta(x_t, t)}{dx_t}$ is the Jacobian matrix of u_θ with respect to x_t .

Assumption 2 (Constancy of the Jacobian): The Jacobian $J_{u_\theta}(x_t, t)$ varies slowly with respect to x_t within these small neighborhoods. Therefore, for small δ , it can be approximated as constant:

$$J_{u_\theta}(x_t + \delta, t) \approx J_{u_\theta}(x_t, t). \quad (6.20)$$

Under these assumptions, we derive the following gradient relationship between $\nabla_{x_t} \mathcal{L}$ and $\nabla_{\hat{x}_0} \mathcal{L}$:

Lemma 6.4.2 (Gradient Relationship). *Let $u_\theta : \mathcal{R}^d \times [0, T] \rightarrow \mathcal{R}^d$ be the velocity function with the parameter θ . Then the gradient of the cost function ($\nabla_{x_t} \mathcal{L}$) at any timestep t can be approximated as:*

$$\nabla_{x_t} \mathcal{L} = (I + t \cdot J_{u_\theta})^T \nabla_{\hat{x}_0} \mathcal{L}. \quad (6.21)$$

Proof. Leveraging the straight-line trajectories characteristic of rectified flow models, the data sample at $t = 0$ can be estimated directly from an intermediate state x_t :

$$\hat{x}_0 = x_t + t \cdot u_\theta(x_t, t). \quad (6.22)$$

By differentiating the $\hat{x}_0$ with respect to x_t , we get:

$$\frac{d\hat{x}_0}{dx_t} = I + t \cdot \frac{du_\theta(x_t, t)}{dx_t} \quad (6.23)$$

$$= I + t \cdot J_{u_\theta}(x_t, t). \quad (6.24)$$

Using the chain rule for gradients:

$$\nabla_{x_t} \mathcal{L} = \left(\frac{d\hat{x}_0}{dx_t} \right)^T \nabla_{\hat{x}_0} \mathcal{L}. \quad (6.25)$$

Substituting the expression for $\frac{d\hat{x}_0}{dx_t}$, we obtain:

$$\nabla_{x_t} \mathcal{L} = (I + t \cdot J_{u_\theta}(x_t, t))^T \nabla_{\hat{x}_0} \mathcal{L}. \quad (6.26)$$

□

Therefore, we get $\epsilon(t) = t \cdot J_{u_\theta}(x_t, t)^T \nabla_{\hat{x}_0} \mathcal{L}$. Importantly, when $t = 0$ or J_{u_θ} varies slowly, $I + t \cdot J_{u_\theta}(x_t, t)$ can be treated as constant with a special case of identity when $t = 0$. In Figure 6.2(c), we can observe that the gradient directions align as $t \rightarrow 0$. Under this approximation, the difference between the two error dynamics becomes negligible. Since $\epsilon(t)$ introduces only a small correction, it leads to the convergence in error dynamics as $t \rightarrow 0$. Combining the results of Proposition 6.4.1, Assumption 1 and 2, and Lemma 6.4.2, we obtain the following theorem with a straightforward proof that facilitates controlled generation for rectified flow models in a computationally efficient way:

Theorem 4.3 (Update Rule for Steering the RFMs). Let $u_\theta : \mathbb{R}^d \times [0, T] \rightarrow \mathbb{R}^d$ be a velocity field with constant Jacobian J_{u_θ} . Define the estimated initial state $\hat{x}_0$ from an intermediate state x_t by

$$\hat{x}_0 = x_t + t \cdot u_\theta(x_t, t).$$

Consider the quadratic loss function $\mathcal{L} = \|\hat{x}_0 - x_0^{\text{ref}}\|^2$, where x_0^{ref} is a reference sample. Then, the update rule for controlled generation is given by

$$x_{t-\Delta t} = x_t + \Delta t u_\theta(x_t, t) - s' \nabla_{\hat{x}_0} \mathcal{L},$$

where:

- $\nabla_{\hat{x}_0} \mathcal{L} = 2(\hat{x}_0 - x_0^{\text{ref}})$,
- $s' \approx (I + \Delta t \cdot J_{u_\theta})(I + t \cdot J_{u_\theta})^\top$,
- I is the identity matrix.

Proof. By lemma 6.4.2 and Assumption 2, we can further approximate the Eq. 6.21:

$$\nabla_{x_t} \mathcal{L} = (I + t \cdot J_{u_\theta})^T \nabla_{\hat{x}_0} \mathcal{L} \approx K^T \nabla_{\hat{x}_0} \mathcal{L}, \quad (6.27)$$

where K is the constant matrix as $t \rightarrow 0$ and $t \rightarrow 0$. Under this formulation, we can perform controlled image generation in three steps:

$$\begin{aligned} \text{Step 1:} \quad & \hat{x}_0 = x_t + t \cdot u_\theta(x_t, t) \\ \text{Step 2:} \quad & \hat{x}_t = x_t - K^T \nabla_{\hat{x}_0} \mathcal{L} \\ \text{Step 3:} \quad & x_{t-t} = \hat{x}_t + t \cdot u_\theta(\hat{x}_t, t). \end{aligned} \quad (6.28)$$

However, this will require additional forward passes. But according to Assumption 2 if t is sufficiently small, then by Taylor series approximation, we get:

$$x_{t-t} = x_t - K^T \nabla_{\hat{x}_0} \mathcal{L} + t \cdot u_\theta(x_t - K^T \nabla_{\hat{x}_0} \mathcal{L}, t) \quad (6.29)$$

$$\begin{aligned} &= x_t - K^T \nabla_{\hat{x}_0} \mathcal{L} \\ &\quad + t [u_\theta(x_t, t) - J_{u_\theta} \cdot K^T \cdot \nabla_{\hat{x}_0} \mathcal{L}] \end{aligned} \quad (6.30)$$

Now, as J_{u_θ} is constant w.r.t. t . Hence, we get:

$$x_{t-t} = x_t - (I + t \cdot J_{u_\theta}) K^T \nabla_{\hat{x}_0} \mathcal{L} + t \cdot u_\theta(x_t, t) \quad (6.31)$$

$$= x_t + t \cdot u_\theta(x_t, t) - s' \nabla_{\hat{x}_0} \mathcal{L}, \quad (6.32)$$

where $s' = (I + t \cdot J_{u_\theta}) K^T$ is constant and it can predetermined.

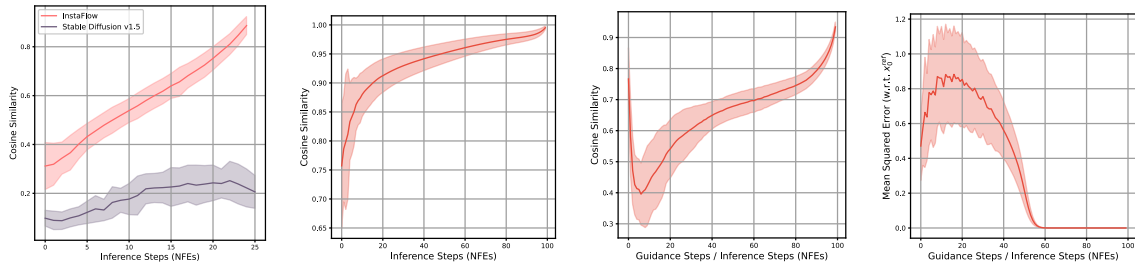
Hence, this concludes the proof that for appropriate guidance scale s' , we can perform the controlled generation as derived above. $\square$

This theorem forms the core of **FlowChef**, enabling efficient controlled generation. Importantly, in Appendix 6.5.1, we derive the constraints that ensure numerical accuracy to

guarantee convergence. Specifically, as long as s' and t are sufficiently small, we will stay in pretraining distributions to avoid the divergence.

Remark 3. $\nabla_{\hat{x}_0} \mathcal{L}$ is a simple differentiation of the loss function. It does not involve any gradients. However, when working on inverse problems for latent diffusion/flow models, the VAE decoder will have gradients.

6.4.3 Empirical Findings.



(a) Gradient Similarity in In-staFlow vs. Stable Diffusion v1.5. (b) Gradient Similarity in Rectified Flow ++ model. (c) Gradient Similarity in Rectified Flow ++ during model steering. (d) Convergence in Rectified Flow ++ during model steering.

Figure 6.4: Empirical analysis of gradient similarity (a, b, and c) and convergence rate. (a) and (b) analyzes the gradients without model steering. (c) contains the gradient similarity during the active model steering. And (d) shows the trajectory similarity at each timestep t w.r.t. the inversion based trajectory.

In Section 9.5, we provided theoretical insights into **FlowChef** along with an intuitive algorithm. To complement the theory, we conducted an empirical analysis on large-scale RFMs to validate the Assumptions, Propositions, Lemmas, and Theorems presented. The results are summarized in Figure 6.4.

In Figure 6.4a, we compare the gradient cosine similarity with and without backpropagation through the ODESolver for InstaFlow and Stable Diffusion v1.5. For all denoising steps, the gradients of SDv1.5 behave nearly randomly, indicating that the stochasticity of the base model significantly impacts gradients, even when using the ODE sampling process during inference. In contrast, for InstaFlow, as denoising progresses ($t \rightarrow 0$), gradient alignment

Algorithm 3: Proposed **FlowChef** (generalized).

```
1 Input: Pretrained Rectified-flow model  $u_\theta$ , input noise sample  $x_T \sim N(0, I)$ , target
   data sample  $x_0^{ref}$ ,  $s'$  is the guidance scale, and  $\mathcal{L}$  cost function.
2 for  $t \in \{T \dots 1\}$  :
3      $v \leftarrow u_\theta(x_t, t)$ 
4      $dt \leftarrow 1/T$ 
5      $\hat{x}_0 \leftarrow x_t + t \cdot v$ 
6      $x_t \leftarrow x_t - s' \nabla_{\hat{x}_0} \mathcal{L}(\hat{x}_0, x_0^{ref})$  // Lemma 6.4.2
7      $x_{t-1} \leftarrow x_t + dt \cdot v$  // Theorem 6.4.2
8 RETURN  $x_0$ 
```

improves, supporting our derivation in Lemma 6.4.2, which states that as $t \rightarrow 0$, we achieve $\nabla_{x_t} \approx \nabla_{\hat{x}_0}$.

Further analysis was performed on the Rectified Flow++ model, which is designed for straight trajectories with zero crossovers. As shown in Figure 6.4b, well-trained models exhibit high gradient similarity even at the initial stages of denoising. However, as illustrated in Figure 6.4c, during active steering, the gradient direction initially diverges before improving. This behavior is also reflected in the convergence plot in Figure 6.4d.

We hypothesize that this phenomenon arises due to the proximity to the Gaussian noise space ($p_1 \sim N(0, I)$), where model steering is more error-prone since minor adjustments can disproportionately affect future trajectories. As denoising progresses and the distribution moves further from the noise (p_1), these errors diminish, and convergence is achieved. These observations align well with our theoretical predictions, further reinforcing the validity of **FlowChef**.

6.4.4 FlowChef Igorithm.

Inverse Problems. Algorithm 3 provides a generalized overview of **FlowChef**. A key feature of **FlowChef** is that it starts from any random noise $x_T \sim \mathcal{N}(0, I)$ and still converges to the desired distribution or sample without inversion or backpropagation. At each timestep t , we first estimate the $\hat{x}_0$. Then we calculate the loss $\mathcal{L}(\hat{x}_0, x_0^{\text{ref}})$. At last, we directly optimize x_t using the gradient $\nabla_{\hat{x}_0} \mathcal{L}$, as per Lemma 6.4.2. Similarly, we extend this core principle and present **FlowChef-Edit** for image editing on Flux and InstaFlow models, with Algorithm 3 detailing the implementation. We provide a detailed ablation study on hyperparameters in the Appendix 6.6.5.

Image Editing. As described in Section 6.4.2, **FlowChef** can be easily extended to image editing. Revisiting the core concept, **FlowChef** modifies random trajectories to align with a target sample. Image editing involves balancing similarity with the target sample while introducing deviations to achieve desired edits.

Figure 6.3 and Section 6.6.5 illustrate how **FlowChef** progressively transfers characteristics from high-level structure to finer details like color composition. However, editing requirements vary by task. For example, adding an object benefits from trajectory adjustments earlier in the denoising process, while color changes require gradual learning at later stages. We can optimize parameters for diverse tasks using the generalized **FlowChef**, as detailed in Algorithm 3.

To simplify the process, we extend **FlowChef** to support off-the-shelf editing tasks, such as those in the PIE-Benchmark, as detailed in Algorithm 4. Assume a non-edit region mask, M_{edit} , derived from cross-attention or human annotation. To steer the trajectory towards the desired edits, we modify the velocity (v) using a classifier-free guidance strategy:

$$v = v_{\text{edit}} + \neg \text{mask} \cdot (v_{\text{edit}} - v_{\text{base}}) \cdot s, \quad (6.33)$$

where v_{edit} corresponds to the edit prompt and v_{base} to the base (negative) prompt. This adjustment ensures the trajectory reflects the desired edits.

To maintain alignment of non-edited regions with the target sample, we modify the cost function as follows:

$$\mathcal{L}(\hat{x}_0, x_0^{ref}) = \|(\hat{x}_0 - x_0^{ref}) \cdot M_{edit}\|_2^2. \quad (6.34)$$

Preserving the original image structure is crucial for edits such as color or material changes. To achieve this, we introduce the parameter $max_full_steps_T$, which determines the number of steps that apply full **FlowChef** guidance with an identity mask. This ensures structural preservation while facilitating edits. Section 6.6.5 details a comprehensive reference for hyperparameters.

FlowChef vs. Baseline FreeDoM. Algorithm 4 compares **FlowChef** to the baseline FreeDoM, a diffusion model method that modifies the score function using a classifier guidance-like approach. FreeDoM requires estimating velocity and calculating gradients (∇_{x_t}) through backpropagation via the ODESolver u_θ , as marked in red. In contrast, as highlighted in green, **FlowChef** eliminates the need for backpropagation while still achieving convergence. This simplification makes **FlowChef** a more efficient and practical solution without sacrificing performance.

6.5 Extended Theoretical Analysis

6.5.1 Numerical Accuracy for Model Steering

In our controlled generation framework, we aim to steer the generation process towards a reference sample x_0^{ref} by solving the modified ODE:

$$\frac{dx(t)}{dt} = f(x(t), t) = u_\theta(x(t), t) - s' \nabla_{\hat{x}_0} \mathcal{L}. \quad (6.35)$$

Algorithm 4: FlowChef vs. Baseline FreeDoM.

```
1 Input: Pretrained Rectified-flow model  $u_\theta$ , input noise sample  $x_T \sim N(0, I)$ , target
   data sample  $x_0^{ref}$ , and  $\mathcal{L}$  cost function.
2 for  $t \in \{T \dots 0\}$  :
3    $v \leftarrow u_\theta(x_t, t)$   $dt \leftarrow 1/T$ 
4    $x_t \leftarrow x_t.require\_grad\_(\text{True})$ 
5   for  $\mathbf{N}$  steps :
6      $v \leftarrow u_\theta(x_t, t)$ 
7      $\hat{x}_0 \leftarrow x_t + t \cdot v$ 
8      $loss \leftarrow \mathcal{L}(\hat{x}_0, x_0^{ref})$ 
9      $\nabla_{x_t} \leftarrow \text{grad}(loss, x_t)$  // Compute heavy BP
10     $x_t \leftarrow \text{Optimize}(x_t, loss)$  // Lemma 6.4.2
11     $v \leftarrow u_\theta(x_t, t)$ 
12     $x_{t-1} \leftarrow x_t + dt \cdot v$  // Theorem 6.4.2
13 RETURN  $x_0$ 
```

The accuracy of this numerical integration is crucial, as errors can accumulate over time, leading to deviations from the desired trajectory. The smoothness of the modified velocity field $f(x(t), t)$ significantly impacts this accuracy. Specifically, a smaller magnitude of $\left| \frac{d}{dt} f(x(t), t) \right|$ reduces local truncation errors. The following Proposition formalizes this relationship, stating that the numerical accuracy improves as $\left| \frac{d}{dt} f(x(t), t) \right|$ decreases.

Proposition 6.5.1. (Informal). *Given the prior notations, assumptions, and Theorem, for any p -th order numerical method solving Eq. (6.35), the accuracy of the numerical solution increases as the magnitude of $\left| \frac{d}{dt} f(x(t), t) \right|$ decreases.*

Algorithm 5: FlowChef optimized for a wide range of image editing tasks.

```
1 Input: Pretrained Rectified-flow model  $u_\theta$ , input noise sample  $x_T \sim N(0, I)$ , target
   data sample  $x_0^{ref}$ ,  $c^{edit}$  is edit prompt,  $c^{base}$  is base prompt,  $M$  is user-provided
   input mask, and  $\mathcal{L}$  cost function.
2 for  $t \in \{T \dots 0\}$  :
3      $dt \leftarrow 1/T$ 
4      $c \leftarrow [c^{edit}, c^{base}]$ 
5      $v \leftarrow u_\theta(x_t, t, c)$ 
6      $v_{edit}, v_{base} = v.chunk(2)$ 
7      $v = v_{edit} + \neg mask \cdot (v_{edit} - v_{base}) \cdot s$ 
8      $M_{edit} \leftarrow M$ 
9      $x_t \leftarrow x_t.require\_grad\_(\text{true})$ 
10    if  $t < min_T$  :
11        for  $\mathbf{N}$  steps :
12             $\hat{x}_0 \leftarrow x_t + t \cdot v$ 
13            if  $t < max\_full\_steps_T$  :
14                 $M_{edit} \leftarrow I$ 
15             $loss \leftarrow \mathcal{L}(\hat{x}_0, x_0^{ref}) \cdot M_{edit}$ 
16             $x_t \leftarrow \text{Optimize}(x_t, loss)$  // Lemma 6.4.2
17     $x_{t-1} \leftarrow x_t + dt \cdot v$  // Theorem 6.4.2
18 RETURN  $x_0$ 
```

Proof. To analyze the local truncation error, consider the Taylor series expansion of the exact solution around time t when integrating backward in time from t to $t - t$:

$$x(t - \tau) = x(t) - \tau f(x(t), t) + \frac{(\tau)^2}{2} \frac{d}{dt} f(x(t), t) - \frac{(\tau)^3}{6} \frac{d^2}{dt^2} f(x(t), t) + O((\tau)^4).$$

The numerical method updates the solution using:

$$x_{t+\tau} = x_t + \tau \phi(x_t, t), \quad (6.36)$$

where $\phi(x_t, t)$ is the increment function. The local truncation error τ is the difference between the exact solution and the numerical approximation:

$$\begin{aligned} \tau &= x(t + \tau) - x_{t+\tau} \\ &= \left[x(t) + \tau f(x(t), t) + \frac{(\tau)^2}{2} \frac{d}{dt} f(x(t), t) + \frac{(\tau)^3}{6} \frac{d^2}{dt^2} f(x(t), t) + O((\tau)^4) \right] \\ &\quad - [x_t + \tau \phi(x_t, t)]. \end{aligned}$$

The first p-order terms cancel out, and we have:

$$\|\tau\| \leq \left\| \frac{(\tau)^{p+1}}{(p+1)!} \frac{d^{p+1}}{dt^{p+1}} f(x(t), t) \right\| \quad (6.37)$$

According to the Mean Value Theorem, we have

$$\|\tau\| \leq C(\tau)^{p+1} \max_{t \in [t_n, t_{n+1}]} \left\| \frac{d^{p+1}}{dt^{p+1}} f(x(t), t) \right\| \quad (6.38)$$

where C is a constant depending on the method. The global error $e(t) = x(t) - x_t$ accumulates these local errors over the integration interval. Under standard assumptions (e.g., Lipschitz continuity of f), the global error is bounded by:

$$\|e(t)\| \leq K(\Delta t)^p (e^{L(T-t)} - 1) \max_{t \in [0, T]} \left\| \frac{d}{dt} f(x(t), t) \right\|, \quad (6.39)$$

where K is a constant depending on the Lipschitz constant L of f and the total integration time T . $\square$

As the magnitude of $\left\| \frac{d}{dt} f(x_t, t) \right\|$ decreases, both the local truncation error and the global error decrease, enhancing the accuracy of the numerical solution. In the context of controlled generation, ensuring that $f(x_t, t)$ changes smoothly over time leads to more accurate integration and better alignment with the reference point x_0^{ref} . This insight and prior assumptions require that the guidance scale s' and Δt be sufficiently smaller, where higher NFEs lead to the lower Δt . Hence, we increase the NFEs significantly to stabilize the steering (see Section 6.6.5). By carefully selecting s' , we ensure that the additional term $s' \cdot \nabla_{\hat{x}_0} \mathcal{L}$ does not introduce excessive variability into $f(x(t), t)$, maintaining the smoothness necessary for accurate numerical integration.

6.5.2 Error Dynamics for Diffusion Models

In this section, we derive the error dynamics for diffusion-based gradient-free guidance methods (e.g., MPGD and RB-Modulation).

Proposition 6.5.2. *In guided diffusion without manifold projection, the error $e_t = \tilde{x}_0^{(t)} - x_0^{\text{ref}}$ evolves as $e_{t-1} = (1-\eta)e_t + (1-\eta)\gamma_t - \epsilon_t$, where $\gamma_t = \frac{\sqrt{1-\eta}}{\sqrt{1-\eta}}$ and $\epsilon_t = \epsilon_\theta(x_t, t) - \epsilon_\theta(x_{t-1}, t-1)$. The error norm satisfies $\|e_{t-1}\| \leq (1-\eta)\|e_t\| + (1-\eta)\|\gamma_t\| - \|\epsilon_t\|$, indicating potential growth when $\|\gamma_t\| - \|\epsilon_t\|$ is large.*

Proof. Assuming a quadratic loss $\mathcal{L} = \frac{1}{2} \|\hat{x}_0^{(t)} - x_0^{\text{ref}}\|_2^2$, the gradient is:

$$\nabla_{\hat{x}_0^{(t)}} \mathcal{L} = \hat{x}_0^{(t)} - x_0^{\text{ref}}, \quad (6.40)$$

and the guided update becomes:

$$\tilde{x}_0^{(t)} = \hat{x}_0^{(t)} - \eta \nabla_{\hat{x}_0^{(t)}} \mathcal{L} \quad (6.41)$$

$$= (1 - \eta) \hat{x}_0^{(t)} + \eta x_0^{ref}, \quad (6.42)$$

where $\eta > 0$ is the guidance strength. The error at timestep t is then:

$$e_t = \tilde{x}_0^{(t)} - x_0^{ref} = (1 - \eta)(\hat{x}_0^{(t)} - x_0^{ref}). \quad (6.43)$$

Next, the DDIM update computes the subsequent noisy sample:

$$x_{t-1} = \sqrt{\alpha_{t-1}} \tilde{x}_0^{(t)} + \sqrt{1 - \alpha_{t-1}} \epsilon_\theta(x_t, t), \quad (6.44)$$

where $\alpha_{t-1} \in (0, 1)$ controls the noise schedule, and $\epsilon_\theta(x_t, t)$ is the noise prediction from the pretrained model. At timestep $t - 1$, the clean sample prediction is:

$$\hat{x}_0^{(t-1)} = \frac{x_{t-1} - \sqrt{1 - \alpha_{t-1}} \epsilon_\theta(x_t, t)}{\sqrt{\alpha_{t-1}}}. \quad (6.45)$$

This can be further simplified to:

$$\hat{x}_0^{(t-1)} = \tilde{x}_0^{(t)} + \gamma_t \epsilon_t, \quad (6.46)$$

where $\gamma_t = \frac{\sqrt{1 - \alpha_{t-1}}}{\sqrt{\alpha_{t-1}}}$ and $\epsilon_t = \epsilon_\theta(x_t, t) - \epsilon_\theta(x_{t-1}, t - 1)$. Applying the guidance step at $t - 1$, we have:

$$\tilde{x}_0^{(t-1)} = (1 - \eta) \hat{x}_0^{(t-1)} + \eta x_0^{ref}, \quad (6.47)$$

and the error becomes:

$$e_{t-1} = \tilde{x}_0^{(t-1)} - x_0^{ref} \quad (6.48)$$

$$= (1 - \eta) \left((\tilde{x}_0^{(t)} - x_0^{ref}) + \gamma_t \epsilon_t \right) \quad (6.49)$$

$$= (1 - \eta)e_t + (1 - \eta)\gamma_t \epsilon_t. \quad (6.50)$$

To understand the evolution of the error magnitude, following the triangle inequality, we get:

$$\|e_{t-1}\| \leq (1 - \eta)\|e_t\| + (1 - \eta)\gamma_t\|\epsilon_t\|. \quad (6.51)$$

In idealized case (i.e., $\epsilon_t = 0$), the first term decays exponentially if $0 < \eta < 1$. However, in practice, $\epsilon_t \neq 0$ and the second term introduces a perturbation that affects the convergence. In diffusion models, the noise schedule typically has ϵ_t increasing from near 0 to 1 as t decreases from T to 1. Thus γ_t is large for large t (early steps) and decreases as t approaches 1. Meanwhile, $\|\epsilon_t\|$, the inconsistency in noise predictions, may be significant early in the process due to high noise levels and diminish later as the sample refines. Hence, cumulative effect of these perturbations may prevent consistent error reduction across all steps. $\square$

For reference, RFMs exhibit straight trajectories and the second term is minimized by default, as shown in the prior theoretical and empirical analysis.

6.6 Experiments

We evaluate **FlowChef** across multiple tasks: (1) Linear inversion problems on pixel and latent-space models, and (2) image editing. We highly encourage the readers to refer to Appendix 6.6.4; we provide further comprehensive evaluations on classifier-guided style transfer and extend **FlowChef** to multi-object editing and 3D generation. Overall,

Method	BoxInpaint			Deblurring			Super Resolution		
	PSNR ($\uparrow$)	SSIM ($\uparrow$)	LPIPS ($\downarrow$)	PSNR ($\uparrow$)	SSIM ($\uparrow$)	LPIPS ($\downarrow$)	PSNR ($\uparrow$)	SSIM ($\uparrow$)	LPIPS ($\downarrow$)
<i>Easy Scenarios</i>									
Degraded	21.79	74.76	10.92	20.17	54.03	22.20	24.68	77.57	11.67
OT-ODE	19.11	77.86	13.49	21.86	62.51	15.14	21.64	62.23	26.64
FreeDoM	20.87	74.79	13.92	20.21	69.73	13.22	21.15	77.54	12.12
DPS	23.61	74.79	9.35	22.49	69.73	10.23	23.94	77.54	8.46
D-FLow	20.37	70.06	13.67	20.22	61.99	14.51	21.60	69.89	12.29
PnP-Flow	22.12	68.02	14.70	22.00	65.79	15.95	22.42	68.06	14.91
FlowChef (ours)	26.32	87.70	3.36	27.69	86.43	2.66	26.00	80.15	4.43
<i>Hard Scenarios</i>									
Degraded	18.75	65.12	22.54	16.83	30.02	54.04	20.77	55.85	38.16
OT-ODE	16.37	67.35	19.22	17.89	34.02	29.68	18.19	39.43	36.84
FreeDoM	18.88	65.07	16.83	16.50	34.88	18.91	19.58	55.84	14.12
DPS	20.68	65.06	13.06	17.58	34.89	15.86	21.52	55.90	10.31
D-FLow	18.34	62.62	19.94	16.93	34.13	25.31	20.01	56.46	17.64
PnP-Flow	20.44	61.96	17.53	19.50	50.54	22.00	21.35	61.78	17.78
FlowChef (ours)	21.45	78.75	7.73	20.31	52.73	10.64	21.62	60.33	10.18

Table 6.1: Pixel-space model-based evaluations for tackling the linear inverse problems. SSIM & LPIPS results are multiplied by 100.

FlowChef demonstrates superior performance across all tasks, significantly reducing compute and time costs compared to baselines. Notably, **FlowChef** extends seamlessly to image editing tasks without inversion.

6.6.1 Experimental Setup

This section outlines the hyperparameters used for **FlowChef** and baseline methods in solving inverse problems.

Pixel-Space Models. All evaluations were conducted using the Rectified Flow++ checkpoint. Since public implementations of OT-ODE and D-Flow are unavailable, we implemented these methods manually based on the provided pseudocode and performed hyperparameter tuning to ensure optimal performance. Notably, DPS and FreeDoM hyperparameters are the same as the **FlowChef**. Table 6.3 provides a detailed overview of the

Method	BoxInpaint			Super Resolution			Deblurring		
	PSNR ($\uparrow$)	SSIM ($\uparrow$)	LPIPS ($\downarrow$)	PSNR ($\uparrow$)	SSIM ($\uparrow$)	LPIPS ($\downarrow$)	PSNR ($\uparrow$)	SSIM ($\uparrow$)	LPIPS ($\downarrow$)
Diffusion based methods									
Resample	20.12	79.94	19.36	26.91	70.91	30.75	25.27	62.97	41.94
PSLD (500 NFEs)	28.30	93.81	4.49	25.79	65.15	33.27	26.64	65.44	43.10
PSLD (100 NFEs)	26.90	93.13	5.29	21.95	54.67	46.08	21.25	51.62	51.92
Flow based methods									
D-Flow	19.68	65.01	27.79	20.23	60.55	50.30	22.42	64.43	53.04
RectifID	23.81	75.13	10.50	10.36	31.55	67.08	10.40	31.16	66.60
FlowChef (InstaFlow)	22.94	73.55	9.94	25.83	64.73	31.38	22.50	47.42	42.54
FlowChef (Rectified Diffusion)	25.96	83.60	4.93	27.06	69.67	27.47	24.96	63.59	40.79
FlowChef (Flux)	25.74	82.99	9.40	20.25	64.34	41.88	18.98	64.37	53.43

Table 6.2: Latent-space model based evaluations for tackling the linear inverse problems. SSIM & LPIPS results are multiplied by 100.

Hyperparameter	OT-ODE	D-Flow	PnP-Flow	FlowChef
Iterations / NFEs	200	20	50	200
Optimization per iteration	1	-	-	1
Optimization per denoising	-	50	-	-
vg. sampling steps	-	-	5	-
Guidance scale	1	1	1	500
Cost function	L1	L1**2	L1	MSE
initial time (1 means noise)	0.8	-	-	-
blending strength	-	0.05	-	-
inversion	×	✓	×	×
learning rate	1	1	1	1

Table 6.3: Hyperparameters for solving inverse problems using pixel-space models.

hyperparameters used for each baseline.

Latent-Space Models. For latent-space models, we extended D-Flow to the InstaFlow pretrained model, repurposed RectifID for inverse problems, and fine-tuned the hyperparameters for optimal results. The best-performing hyperparameters for each baseline are listed in Table 6.4. We utilized their baseline implementations for diffusion model-based approaches such as Resample and PSLD-LDM, modifying only the number of inference steps. Specifically, we used 100 NFEs for Resample and 100/500 NFEs for PSLD.

Hyperparameter	D-Flow	RectifID	FlowChef
Iterations / NFEs	10	4	100
Optimization per iteration	-	-	1
Optimization per denoising	20	400	-
Blending strength	0.1	-	-
Guidance scale	0.5	0.5	0.5
Cost function	MSE	MSE	MSE
Learning rate	0.5	1	0.02
Optimizer	Adam	SGD	Adam
loss multiplier (latent/pixel)	0.000001	0.0001 / 100000	0.001/1000
inversion	✓	×	×

Table 6.4: Hyperparameters for solving inverse problems using latent-space models (InstaFlow).

Model	Hyperparameters	Chage Object	dd Object	Remove Object	Change ttribute	Chage Pose	Change Color	Change Material	Change Background	Change Style
FlowChef (InstaFlow)	Learning rate	0.5	0.5	0.5	0.5	0.5	0.5	0.5	1.0	0.5
	Max setps	50	50	50	50	20	30	50	50	30
	Optimization steps	1	1	3	2	2	2	2	4	1
	Inference steps	50	50	50	50	50	50	50	50	50
	Full source steps	30	30	0	10	10	20	20	0	30
	Edit guidance scale	2.0	2.0	2.0	4.5	8.0	8.0	4.0	3.0	6.0
FlowChef (Flux)	Learning rate	0.4	0.5	0.5	0.5	0.5	0.4	0.5	0.5	0.4
	Optimization steps	1	1	1	1	1	1	1	1	1
	Inference steps	30	30	30	30	30	30	30	30	30
	Full source steps	5	5	0	2	5	3	5	0	5
	Edit guidance scale	4.5	4.5	4.5	4.5	7.5	10.0	4.5	0.0	10.0

Table 6.5: Hyperparameter examples for which various editing tasks can be performed (following Algorithm 2). Notably, the FlowChef (Flux) variant can be further optimized for task-specific settings that will follow Algorithm 1 with a careful selection of hyperparameters.

6.6.2 Linear Inversion Problems

We evaluate **FlowChef** against several baselines on three common linear tasks: box inpainting, super-resolution, and Gaussian deblurring, under varying difficulty levels. We extend both **FlowChef** and the baselines to latent-space models to simulate real-world applications, reporting results on PSNR, SSIM [278], and LPIPS [318] across 200 images from CelebA [157] and AFHQ-Cat [43]. Memory requirements and computation time are also analyzed.

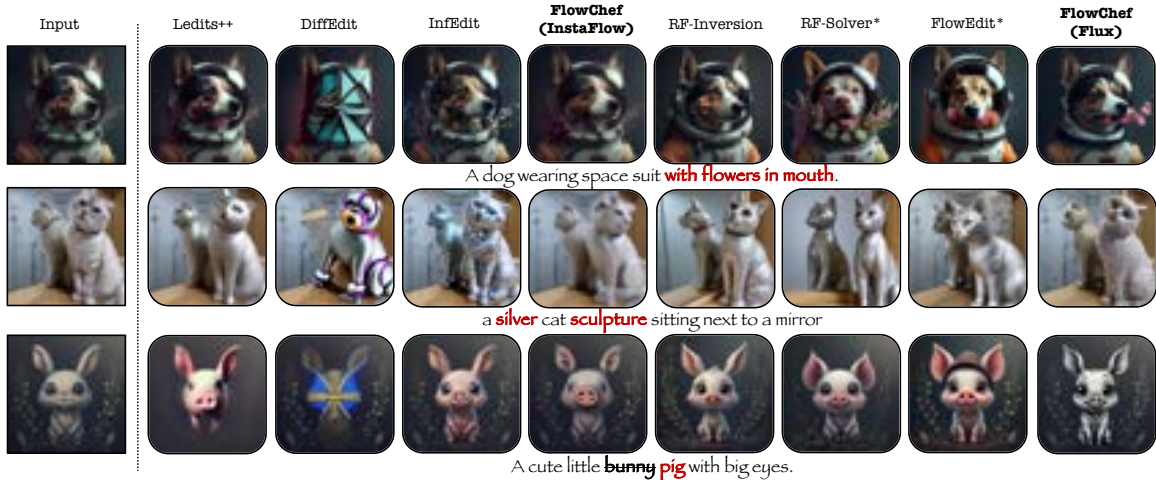


Figure 6.5: Qualitative results on image editing. As illustrated, our method attains the SOTA performance on comparison inversion-free methods. Meanwhile, the FlowChef (Flux) variant achieves better quality and edits. * denotes the concurrent works.

Pixel-space models. As **FlowChef** requires straightness and no crossovers, we select the Rectified-Flow++ pretrained models [136]. We compare **FlowChef** with recent flow-based methods OT-ODE [206], D-Flow [15], and PnP-Flow [173], implementing the former two baselines manually due to lack of open-source access and tuning them for optimal performance. Additionally, we extend two diffusion-based baselines, DPS [44] and Free-DoM [301], for the RFMs. For comparisons, we use the Rectified-Flow++ models that are pretrained on FFHQ (for CelebA) and AFHQ-Cat datasets. Experiments are conducted for 64x64 image resolutions. Hyper-parameters for each method are reported in the Appendix 6.6.1. Our selected tasks include: (1) Box inpainting with 20x20 and 30x30 centered masks, (2) Super-resolution with 2x and 4x scaling factors, and (3) Gaussian deblurring with an 11x11 kernel at intensities of 1.0 and 10.0, with added Gaussian noise at $\sigma = 0.05$.

Results. We present the quantitative and qualitative evaluation results in Table 6.1 and Appendix 6.6.6, respectively. It can be observed that **FlowChef** significantly improves the performance on both easy and hard settings across the tasks and all metrics consistently. Notably, from Table 6.6, we find that the **FlowChef** is also the fastest and most memory

Metric	OT-ODE	PnP-Flow	D-Flow	DPS	FlowChef
VR M (GB)	0.70	0.40	6.44	1.16	0.43
Time (sec)	10.39	5.23	80.42	12.8	4.31

Table 6.6: Compute requirement comparisons on a A6000 GPU.

efficient. Diffusion-based extended baseline (DPS) outperforms even recent baselines. However, DPS requires backpropagation through the ODESolver. Although the existing gradient-free method, PnP-Flow, outperforms many other baselines, **FlowChef** leads the benchmark.

Method	Inversion	Structure ($\downarrow$)	Background Preservation			CLIP Similarity	
	-Free		PSNR ($\uparrow$)	LPIPS ($\downarrow$)	SSIM ($\uparrow$)	Whole ($\uparrow$)	Edited ($\uparrow$)
SD v1.5 variants							
P2P	$\times$	0.0694	17.87	0.2088	0.7114	25.01	22.44
NT-P2P	$\times$	0.0134	27.03	0.0607	0.8411	24.75	21.86
P2PZero	$\times$	0.0617	20.44	0.1722	0.7567	22.8	20.54
PnP	$\times$	0.0282	22.28	0.1135	0.7905	25.41	22.55
DiffEdit		0.0182	24.98	0.0679	0.8474	20.30	18.32
InfEdit (HA)		0.0091	28.45	0.0485	0.8626	24.04	21.13
FlowChef-Edit (HA)		0.0347	29.67	0.0410	0.8676	25.43	22.84
Flux.1[Dev]							
RF-Inversion	$\times$	0.0420	20.59	0.1888	0.7065	25.47	22.55
FlowChef-Edit (CA)		0.0395	24.31	0.1239	0.7960	25.32	22.34
FlowChef-Edit (HA)		0.0436	27.42	0.0845	0.8424	25.69	22.96
Concurrent works							
RF-Solver	$\times$	0.0422	21.07	0.1749	0.7741	26.61	23.39
FlowEdit		0.0284	21.77	0.1154	0.8291	25.34	22.67

Table 6.7: Evaluation metrics for different methods on background preservation and CLIP similarity on PIE-Bench.

Latent-space models. Flow-based baselines are not extended to the latent space models as either they are already very computationally heavy or require extra Jacobian calculations to support the non-linearity introduced by the VAE models. We adapt D-Flow [15] and RectifID [258] as flow-based baselines, adding diffusion-based baselines PSLD-LDM [221]

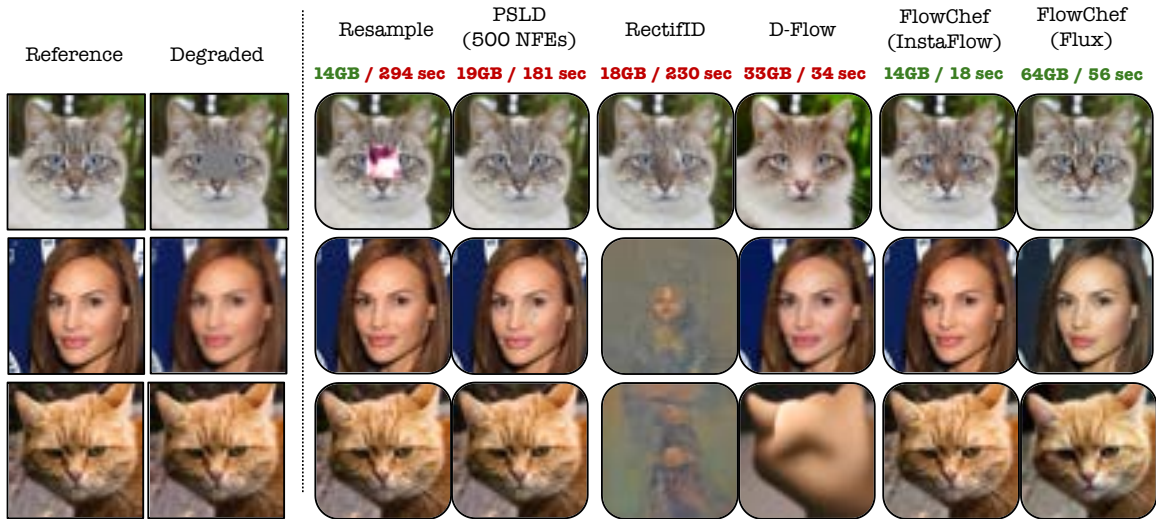


Figure 6.6: Qualitative results on linear inverse problems. All baselines are implemented on stable diffusion v1.5, except FlowChef Flux variant. Results are reported for VRAM and time on an A100 GPU at 512 x 512 resolution, with Flux experiments at 1024 x 1024.

and Resample [249] for comprehensive comparisons. We use InstaFlow [154] and Rectified Diffusion [273] (SDv1.5 non-distilled variants) as a baseline for flow-based approaches and utilize the original SDv1.5 checkpoint for the diffusion-based baselines. We perform all tasks in 512 x 512 resolution, increasing to 1024 x 1024 for Flux experiments. Our task settings are: (1) Box inpainting with a 128x128 mask, (2) Super-resolution at 4x scaling, and (3) Gaussian deblurring with a 50x50 kernel at intensity 5.0, all without extra Gaussian noise. For consistency, settings are doubled for Flux to a 256x256 mask, 8x super-resolution scaling, and 10.0 deblurring intensity. As VAE encoders add extra unwanted nonlinearity, pixel-level cost functions alone may not be optimal. Hence, we calculate the loss in the latent space only for the box inpainting task (as the degradation function is known with $\sigma = 0$), allowing us to extend to image editing later. For super-resolution and deblurring, we stick with the pixel-level cost functions. The task-specific settings and hyperparameters are in the Appendix 6.6.1.

Results. Quantitative and qualitative results in Table 6.2 and Figure 6.6 show that our method (**FlowChef**) achieves SOTA performance for flow-based methods and generalizes

to Rectified Diffusion, which observes a near-linear vector field. However, a gap still remains *w.r.t.* the pure diffusion-based methods like Resample and PSLD. Notably, these baselines take about 5 minutes and 3 minutes, respectively, per image (see Figure 6.6), while **FlowChef** only takes only 18 seconds and less memory (only 14GB). None of the existing flow-based methods can be extended to Flux due to memory constraints. We find that **FlowChef** (Flux) reduces the artifacts in the images but observes the slight degradation in color dynamics due to non-linear vector field.

6.6.3 Image Editing

Due to the optimization constraints, existing baselines for classifier guidance cannot be applied to image editing. Therefore, for comparison, we use diffusion-based inversion methods (Ledits++ [23], P2P [92], NT-P2P [178], P2PZero [26], PnP [107]), and inversion-free methods (DiffEdit [46] InfEdit [290]). We compare flow-based models against RF-Inversion [219] and two concurrent works (RF-Solver and FlowEdit). We perform large-scale evaluations on PIE-Bench [106]. Specifically, the (CA) denotes the mask extracted using ConceptAttention [89], and the (HA) denotes the human-annotated masks from PIE-Bench.

Results. Table 6.7 shows the quantitative results. It can be observed that our method (**FlowChef-Edit** (InstaFlow)) performs superior among the diffusion-based baselines. At the same time, **FlowChef-Edit** (Flux) outperforms the existing inversion-based methods such as RF-Inversion and RF-Solver. Importantly, ConceptAttention is enough to achieve good performance among the flow-based methods, but having human-annotated masks further allows one to control the edits and achieve even better performance. Qualitative results (see Figure 6.5) show that **FlowChef-Edit** achieves superior results consistently. Moreover, existing flow-based methods still struggle to preserve structure and background, while our method not only maintains the overall semantics but performs precise edits.

Method	NFEs	CG Scale	FID ($\downarrow$)	VR M ($\downarrow$)	Time ($\downarrow$)
DDIM	50	-	5.39	3.67	14.22
MPGD	50	1	4.24	6.56	25.01
MPGD	50	10	5.46	6.56	25.01
Ours_{w/ skip grad}	50	1	19.28	6.56	24.95
RFPP (2-flow)	2	-	4.56	3.29	0.28
RFPP (2-flow)	15	-	4.29	3.36	2.75
Ours_{w/ backpropagation}	15	5	2.77	17.98	12.79
Ours_{w/ skip grad}	15	50	3.13	6.64	5.85

Table 6.8: Performance of Various guided sampling methods on ImageNet64x64 with 32 batch size inference on A6000 GPU.

Method	CLIP-I ($\uparrow$)	CLIP-T ($\uparrow$)	VR M	Time
FreeDoM	0.5343	0.2541	17GB	80 sec
MPGD	0.5285	0.2616	16GB	20 sec
RetifID	0.4583	0.1702	18GB	30 sec
D-Flow	0.4851	0.2591	23GB	5 sec
FlowChef_(10 NFEs)	0.5044	0.2655		2 sec
FlowChef_(30 NFEs)	0.5301	0.2600	14GB	7 sec
FlowChef_(30 NFEs $\times$ 2)	0.5531	0.2478		12 sec

Table 6.9: Comparison of Various Classifier Guided Style Transfer.

6.6.4 Extended Results

Classifier Guidance. To validate our findings from Proposition 6.4.1, we conduct a toy study comparing classifier guidance on two ODE sampling methods using pretrained IDDPM and Rectified Flow++ (RF++) models on the ImageNet 64x64. As reported in Table 6.8, skipping the gradient in DDIM-based sampling increases the FID score, indicating significant $\epsilon(t)$. Conversely, RF++ converges well and improves the FID score. These empirical evidences further bolster our hypothesis that Rectified Flow models observe



Figure 6.7: Extending FlowChef to 3D multiview synthesis.

smooth vector field with the help of Proposition 6.4.1. Although backpropagating through the ODESolver further improves performance, it incurs higher computational costs as highlighted.

Style Transfer. We conducted classifier-guided style transfer experiments using 100 randomly selected style reference images paired with 100 random prompts. The objective was to generate stylistic images that align visually with the reference style while adhering to the prompt. A pretrained CLIP model was used for evaluation, and we report both CLIP-T and CLIP-S scores [209]. For baseline comparisons, we included diffusion-based methods FreeDoM and MPGD and flow-based methods D-Flow and RectifID, which were extended for this task. The backbone was fixed to Stable Diffusion v1.5 (SDv1.5), with **FlowChef** evaluated in its InstaFlow variant to ensure a consistent comparison. Both quantitative and qualitative results are presented in Table 6.9, demonstrating the effectiveness of **FlowChef** in this setup.

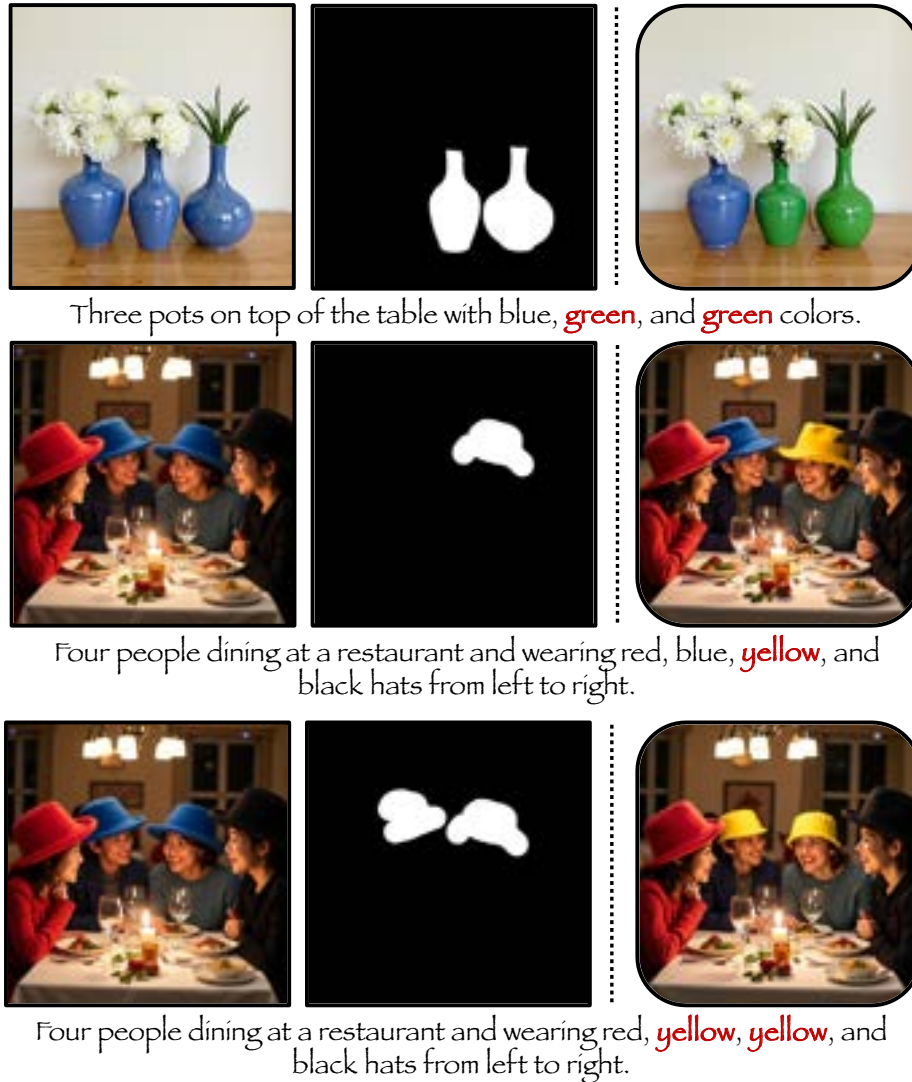


Figure 6.8: FlowChef (Flux) multi object editing examples.

Multiobject editing & 3D generations. To highlight the versatility and effectiveness of **FlowChef**, we extended our method to tackle multi-object image editing and 3D multiview generation. Figure 6.8 demonstrates **FlowChef** (Flux) performing complex multi-object edits, such as simultaneously modifying two pots and hats. Notably, this capability relies on the base model’s ability to understand textual instructions effectively. **FlowChef** leverages this strength of Flux, achieving edits without requiring inversion, a significant advantage over traditional methods. In Figure 6.7, we explore **FlowChef**’s multiview synthesis capability,

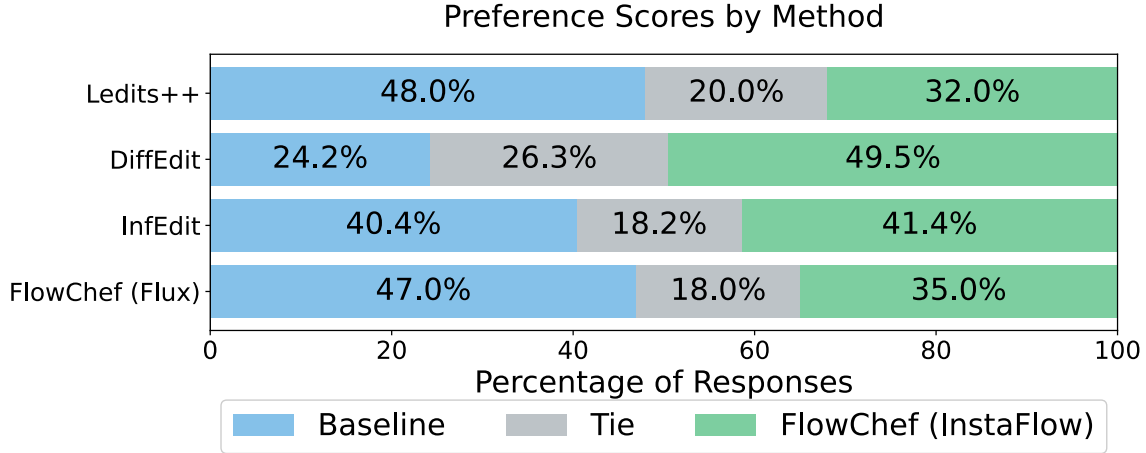


Figure 6.9: Human preference analysis for image editing.

inspired by Score Distillation Sampling (SDS) [208]. By incorporating the core idea of **FlowChef** for model steering into recent work on RFDS [294], we evaluate its effectiveness for 3D view generation. While **FlowChef** does not improve inference efficiency or reduce cost compared to RFDS-Rev [294], it demonstrates competitive performance in generating high-quality multiview outputs. These results underline the adaptability of **FlowChef**, showcasing its potential for advanced generative tasks such as multi-object editing and 3D synthesis, while maintaining the state-of-the-art quality expected from RFMs.

Human Evaluations (Image Editing). A human preference evaluation on randomly selected 100 PIE-Bench edits (see Figure 6.9) shows **FlowChef** (InstaFlow) outperforming DiffEdit and competing with InfEdit. Although Ledits++ scored highest, it requires inversion, resulting in higher VRAM and time requirements. Importantly, FlowChef on Flux achieves performance comparable to Ledits++ without inversion. Comparisons with RF-Inversion show that FlowChef reduces time by almost 50% without needing inversion and achieves competitive performance.

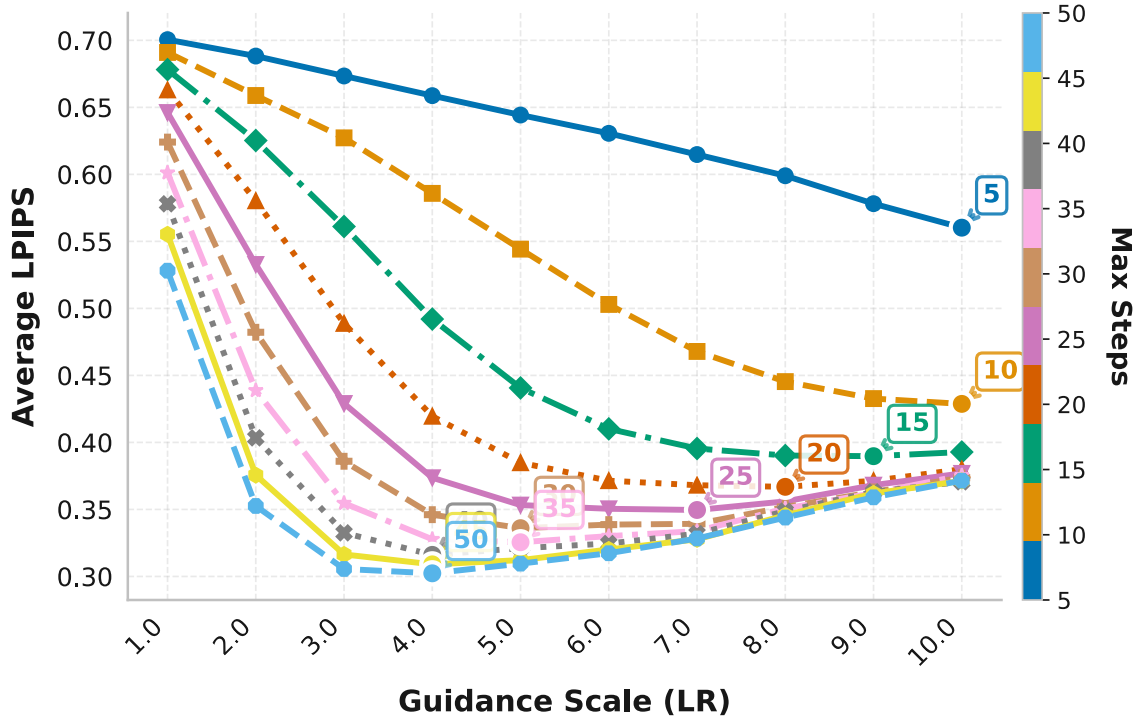


Figure 6.10: Hyperparameter sensitivity analysis.

6.6.5 Hyper-parameter Study

Figures 6.11, 6.12, and 6.13 present an analysis of the impact of various hyperparameters on steering the InstaFlow model using **FlowChef**. Figure 6.11 demonstrates that a lower learning rate combined with a single optimization step is insufficient to effectively steer the model. Optimal performance is achieved with a learning rate of 0.1. Additionally, Figure 6.12 shows that lower learning rates necessitate more optimization steps to achieve convergence. Finally, Figure 6.13 illustrates how the denoising trajectory can be controlled by adjusting the learning rate and optimization steps, enabling recovery of the target sample with the desired accuracy. This control is particularly critical for image editing tasks, where striking the right balance between preserving the reference sample and applying the editing prompt is essential. Table 6.5 further highlights optimal hyperparameter settings for image editing tasks, providing valuable guidance for achieving high-quality edits. This study

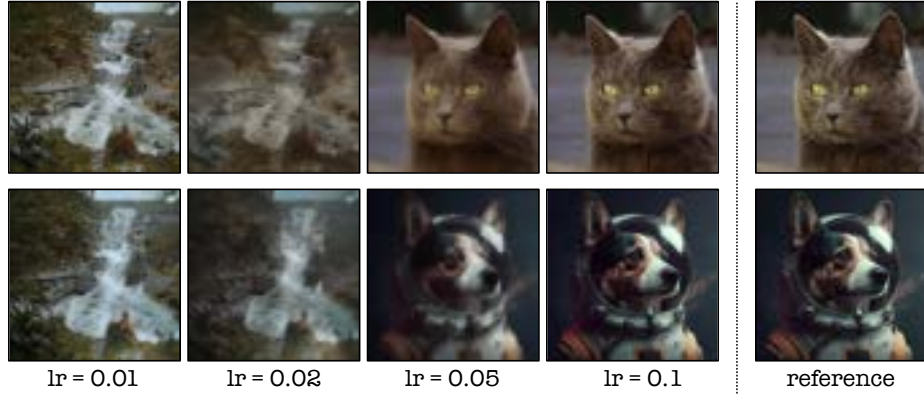


Figure 6.11: Effect of FlowChef learning rate with fixed 20 max steps and one optimization step on InstaFlow.

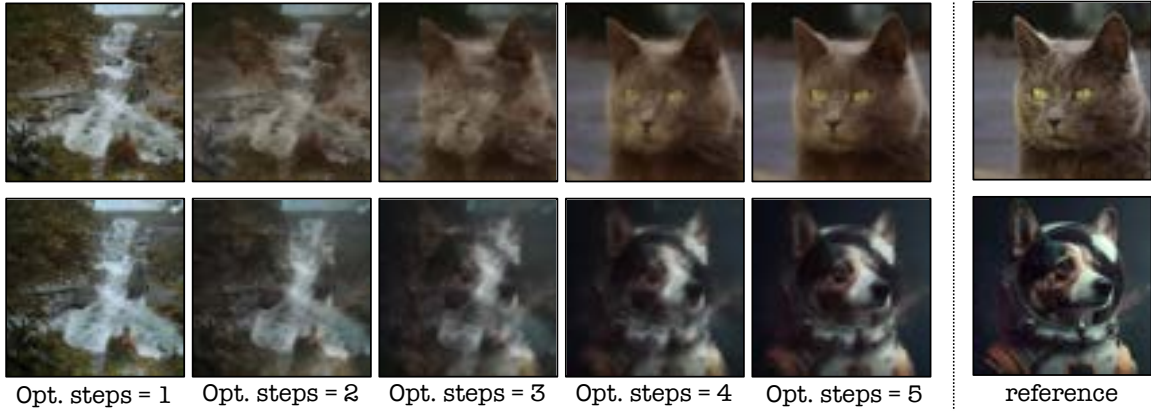


Figure 6.12: Effect of FlowChef optimization steps with fixed 20 max steps and 0.02 learning rate on InstaFlow.

underscores the flexibility of **FlowChef** in adapting to diverse use cases by tuning these parameters effectively.

Additionally, we plot the hyperparameter sensitivity for **FlowChef** (InstaFlow)’s convergence rate w.r.t. guidance scale and steps on 100 AFHQ-Cat images on latent inverse problem (see Figure 6.10). We can observe that lower steps require higher guidance scale and more steps require lower guidance scale. Importantly, **FlowChef** observes a smooth convergence that clearly indicates the robustness.

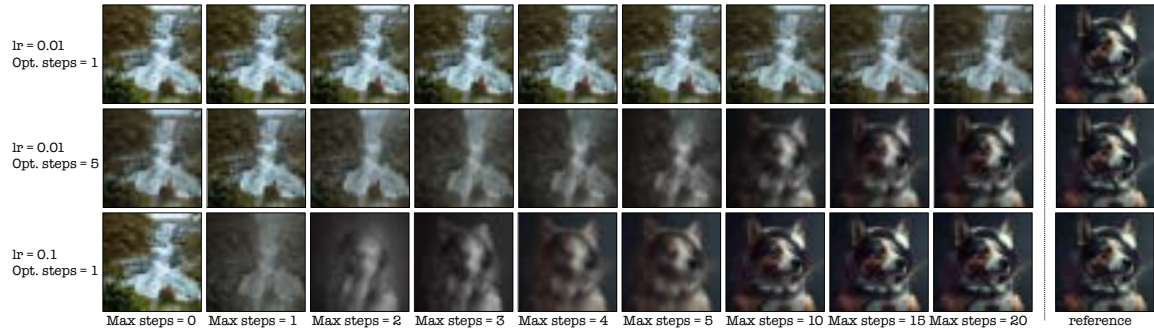


Figure 6.13: Effect of various FlowChef’s steering parameters with increasing maximum optimization steps on InstaFlow.

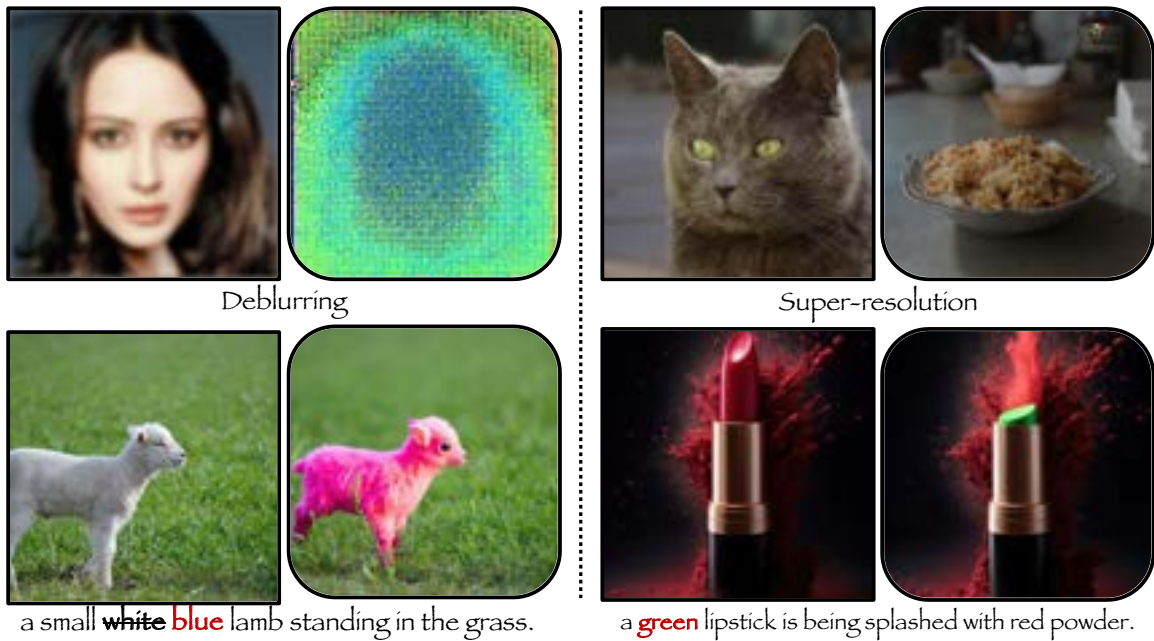


Figure 6.14: FlowChef (Flux) model failures on inverse problems and image editing.

6.6.6 Qualitative Results

Figure 6.15 showcases additional qualitative examples of image editing tasks. For tasks such as changing materials or removing objects, **FlowChef** outperforms the baselines significantly. However, some limitations are noted: while **FlowChef** (InstaFlow) struggles to replace a cat with a tiger, InfEdit handles this task effectively, and Ledit++ exhibits difficulties. On the other hand, **FlowChef** (Flux) achieves superior results, though it replaces a dog with a tiger instead of a lion in one instance. In the final example, both

Ledits++ and **FlowChef** successfully edit long hair into short hair. Importantly, the results in Figure 6.15 are presented without cherry-picking, using consistent hyperparameters for both baselines and **FlowChef**. Variability in outcomes may still arise due to random seeds and fine-tuned hyperparameter selection.

Figures 6.16, 6.17, 6.18, 6.19, 6.20, and 6.21 provide pixel-level qualitative results for various inverse problems, spanning inpainting, deblurring, and super-resolution tasks under both easy and hard scenarios. Readers are encouraged to zoom in to inspect these comparisons more closely. For each task, we randomly selected 10 CelebA examples and evaluated various baselines. Across all difficulty levels, FreeDoM, DPS, and PnPFlow demonstrate better performance than D-Flow and OT-ODE. However, **FlowChef consistently outperforms all baselines**, producing sharp and visually appealing results where other methods either fail outright or introduce excessive smoothness. Hard scenarios pose challenges for all methods, but **FlowChef** notably improves performance even under these conditions. While **FlowChef** shows promise, future work is needed to address potential adversarial effects and further enhance robustness.



Figure 6.15: Qualitative results on image editing. Additional qualitative comparisons of FlowChef with the baselines.

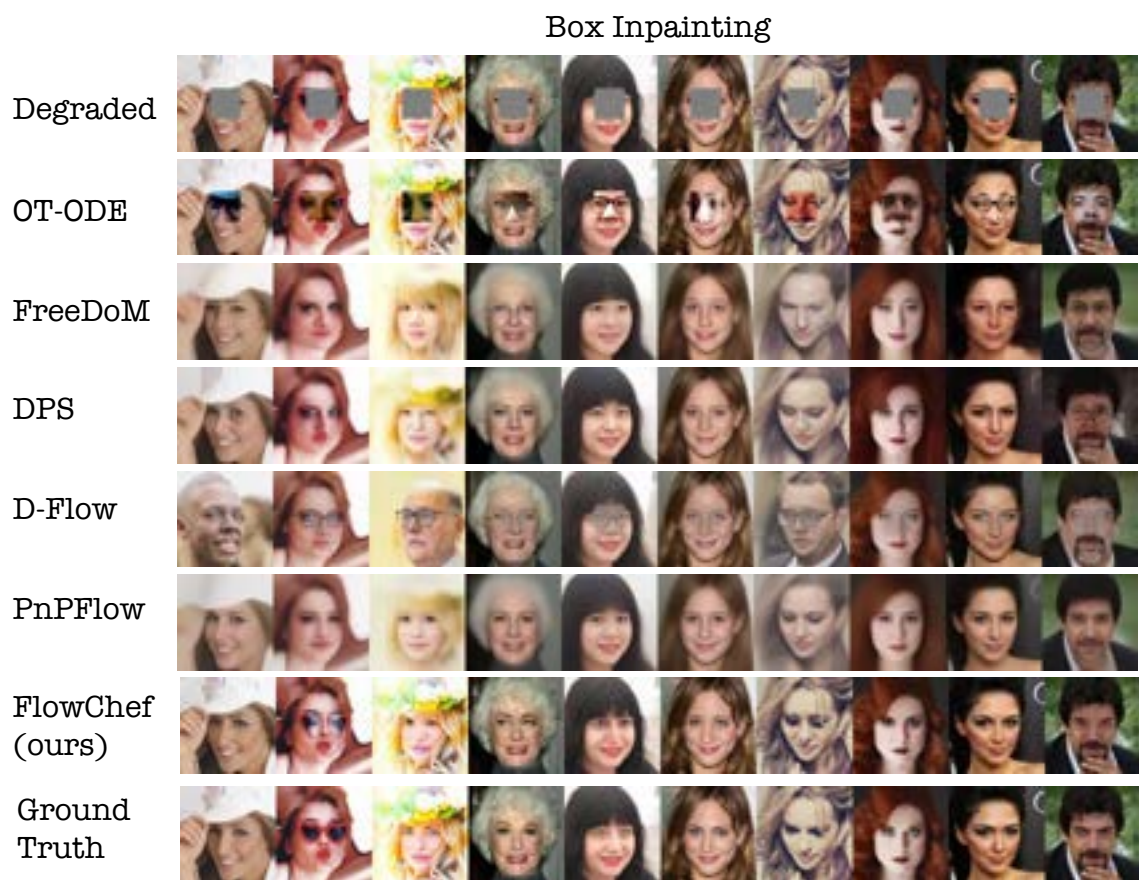


Figure 6.16: Qualitative examples of various methods for easy box inpainting task on RF++.

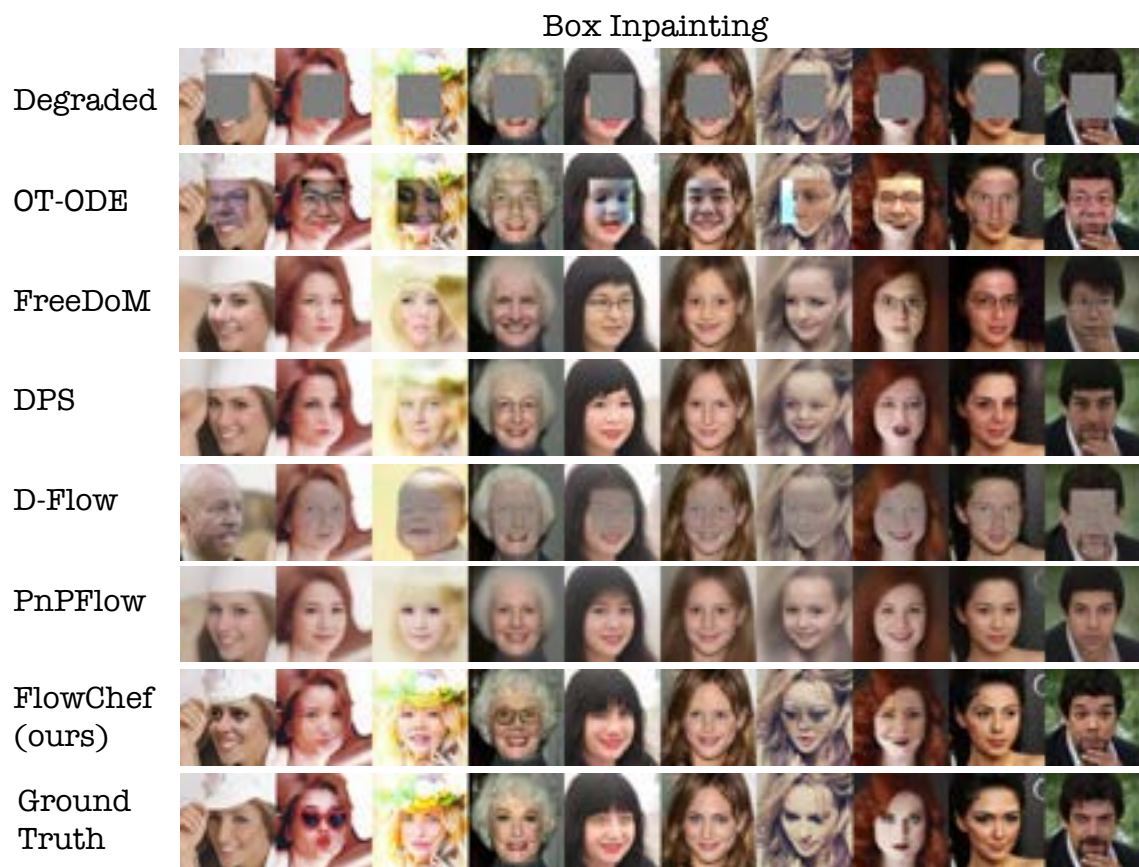


Figure 6.17: Qualitative examples of various methods for hard box inpainting task on RF++.

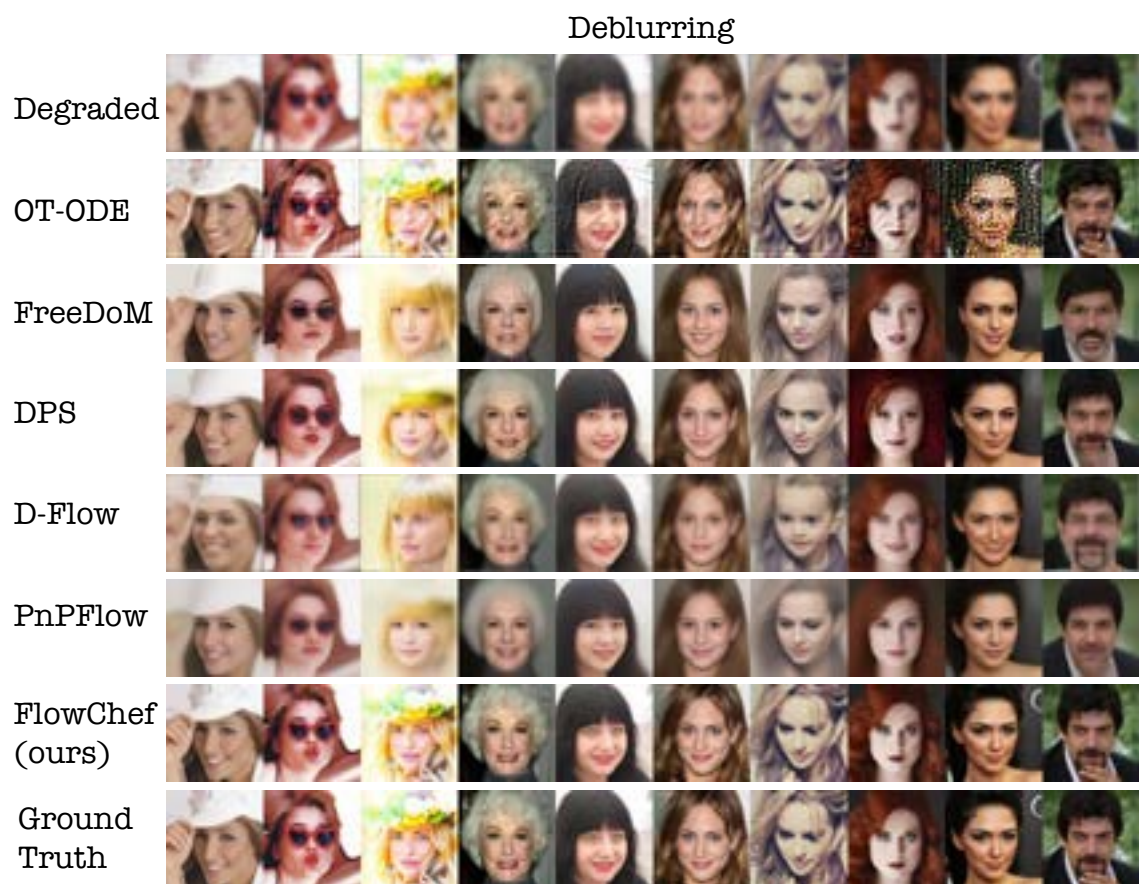


Figure 6.18: Qualitative examples of various methods for an easy deblurring task on RF++.

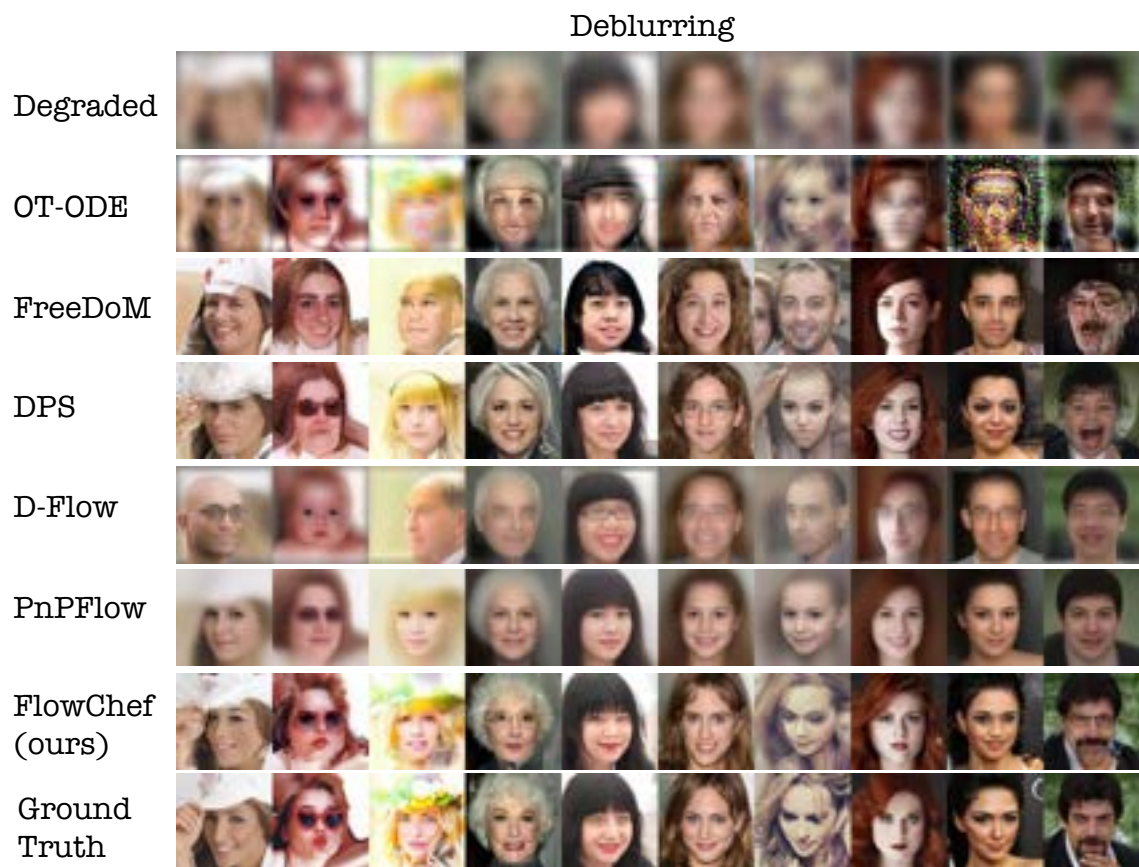


Figure 6.19: Qualitative examples of various methods for the hard deblurring task on RF++.

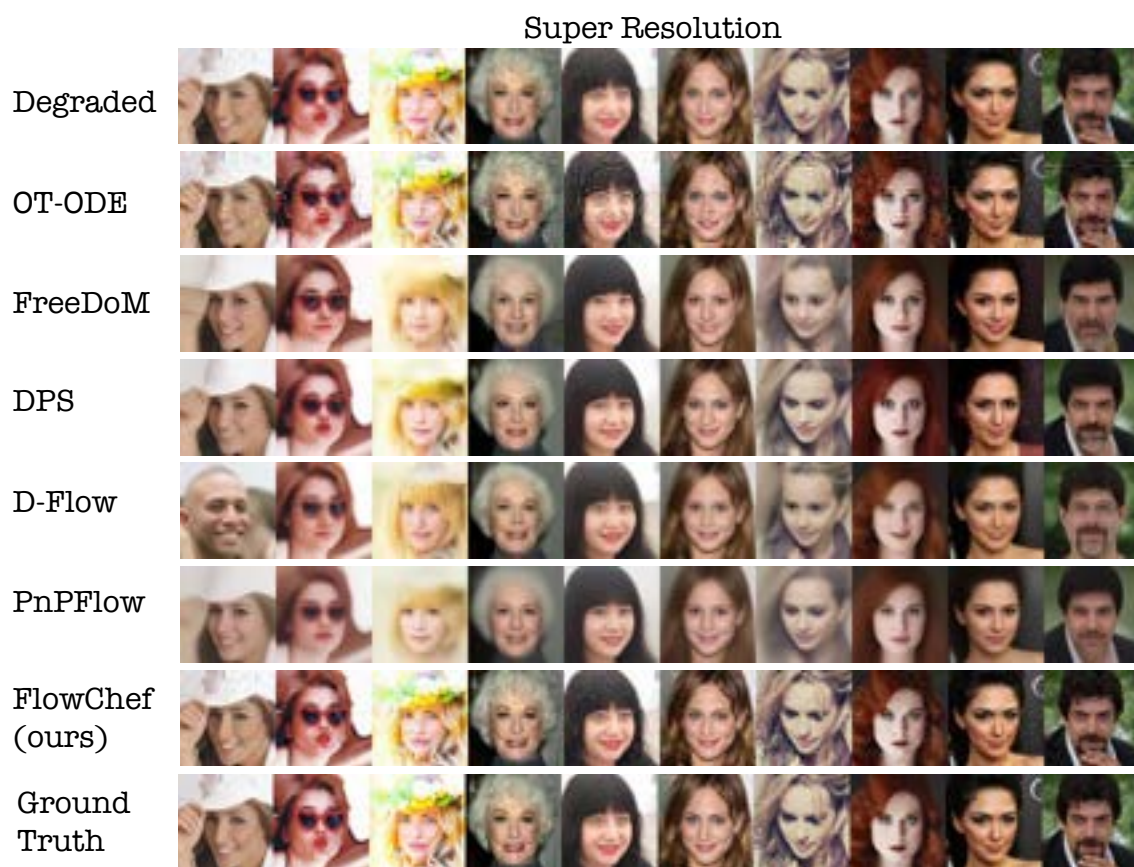


Figure 6.20: Qualitative examples of various methods for an easy super-resolution task on RF++.

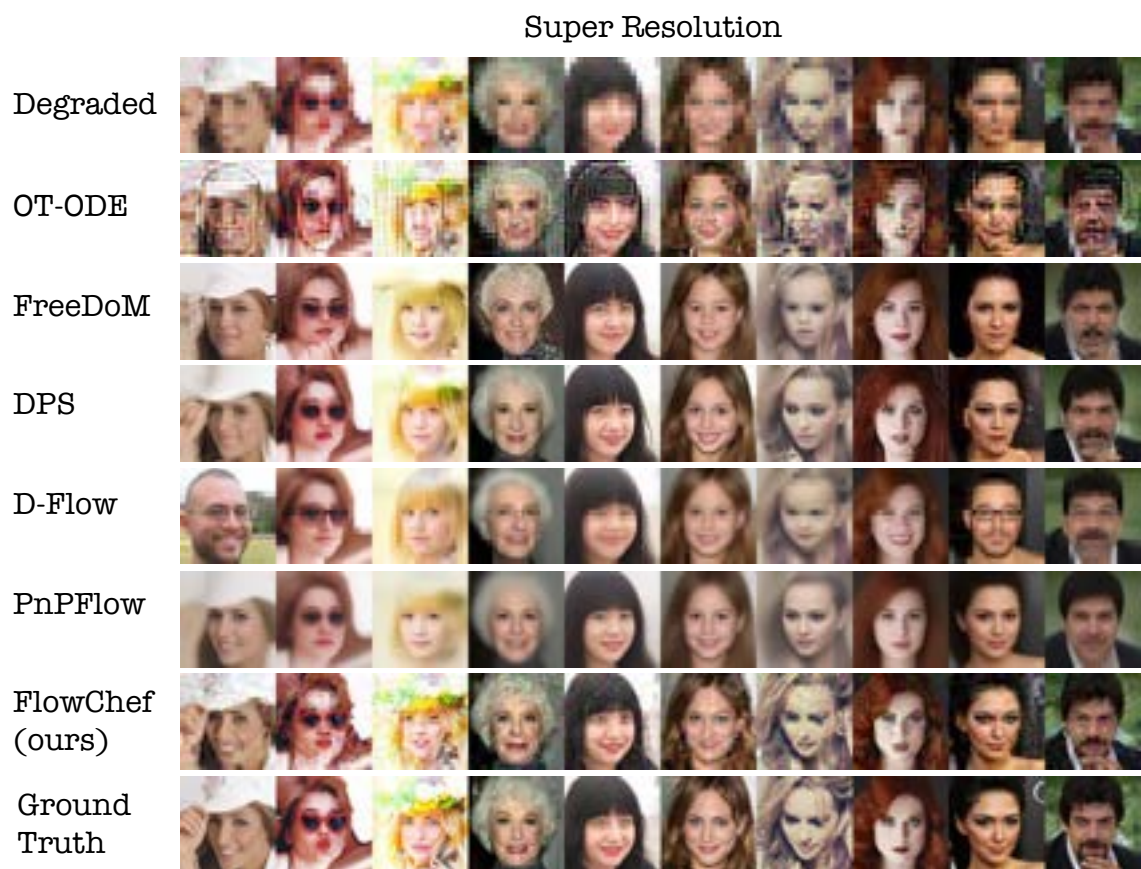


Figure 6.21: Qualitative examples of various methods for the hard super-resolution task on RF++.

6.7 Conclusion

In this work, we introduced **FlowChef**, a versatile rectified flow-based approach that unifies key tasks in controlled image generation, including linear inverse problems, image editing, classifier-guided style transfer, etc. Extensive experiments show that **FlowChef** outperforms baselines across all tasks, achieving state-of-the-art performance with reduced computational cost and memory usage. Notably, **FlowChef-Edit** enables inversion-free editing and scales to SOTA T2I models like Flux. Our results demonstrate **FlowChef**'s adaptability and efficiency, offering a unified solution for both pixel and latent spaces across diverse architectures and practical constraints.

Limitations. While **FlowChef** represents a significant leap in steering RFMs for controlled generation, it shares some limitations with its baseline counterparts. Hyperparameter tuning remains a challenge, particularly due to differences in trajectory behavior. For instance, while InstaFlow trajectories are relatively linear, Flux.1[Dev] trajectories exhibit non-linearity, necessitating careful tuning. As shown in Figure 6.8, **FlowChef** (Flux) faces difficulties in deblurring and super-resolution tasks, which we attribute to the pixel-space loss and non-linear behavior of the VAE model. Importantly, these limitations occur about 20%-25% of cases and can often be resolved by simply adjusting the random seed. Furthermore, due to Flux's lack of true classifier-free guidance (CFG), Algorithm 5 occasionally fails to perfectly execute color changes, sometimes producing the unaltered target image without reflecting the edit (see Figure 6.8). Despite these minor limitations, **FlowChef** still delivers state-of-the-art performance, making these challenges opportunities for further refinement rather than fundamental drawbacks.

Future Work. **FlowChef** opens a promising avenue for steering RFMs effortlessly with guaranteed convergence for controlled image generation. While this work extensively evalu-

ates **FlowChef** on image generative models, future research should focus on expanding its capabilities to video and 3D generative models, areas that remain largely unexplored. Additionally, the current implementation assumes the availability of human-annotated masks for image editing. Automating this step with advanced attention mechanisms could make **FlowChef** a fully automated image editing framework. We encourage the research community to build upon this foundation to enhance its accessibility and functionality.

Ethical Concerns. As with all generative models, ethical concerns such as safety, misuse, and copyright issues apply to **FlowChef** [113, 114]. By enabling controlled generation with state-of-the-art RFMs, **FlowChef** can be leveraged for beneficial and harmful purposes. To mitigate these risks, future efforts should focus on solutions such as image watermarking, content moderation, and unlearning harmful behaviors. While these issues are not unique to **FlowChef**, addressing them will be key to ensuring its responsible use.

LEARNING CONCEPT ERASURE POLICIES VIA GFLOWNET-DRIVEN ALIGNMENT

7.1 Introduction

Recent advances in diffusion models have led to remarkable improvements in text-to-image generative models, enabling unprecedented photorealism and widespread adoption across various domains [216, 61, 211, 195, 317]. However, because these models are typically trained on large-scale, unregulated internet data, concerns have grown regarding their potential misuse, including the unauthorized reproduction of copyrighted or harmful content [236, 70]. Consequently, methods to erase or “unlearn” specific concepts from pretrained text-to-image generators have become critically important to ensure the model’s safety and compliance [75, 127].

Prior approaches addressing these challenges broadly include filtering [299], training data attribution [274], post-generation filtering [179], and explicit concept unlearning [75, 127]. While filtering-based methods are hard to scale on arbitrary concepts, unlearning-based methods have recently attracted significant attention due to their capability to intervene. Early unlearning techniques primarily relied on fine-tuning pretrained diffusion models by modifying cross-attention mechanisms [76]. These approaches, however, often degrade the generative model’s overall quality and are susceptible to adversarial reintroduction of erased concepts [321]. Reinforcement learning (RL)-based strategies were later introduced to improve alignment [240, 189], yet they similarly suffer from brittleness and susceptibility to adversarial attacks. More recently, adversarial unlearning methods have demonstrated improved robustness [114, 319], but at the cost of substantial computational overhead,

typically requiring hours of compute to erase a single concept, thereby severely limiting scalability.

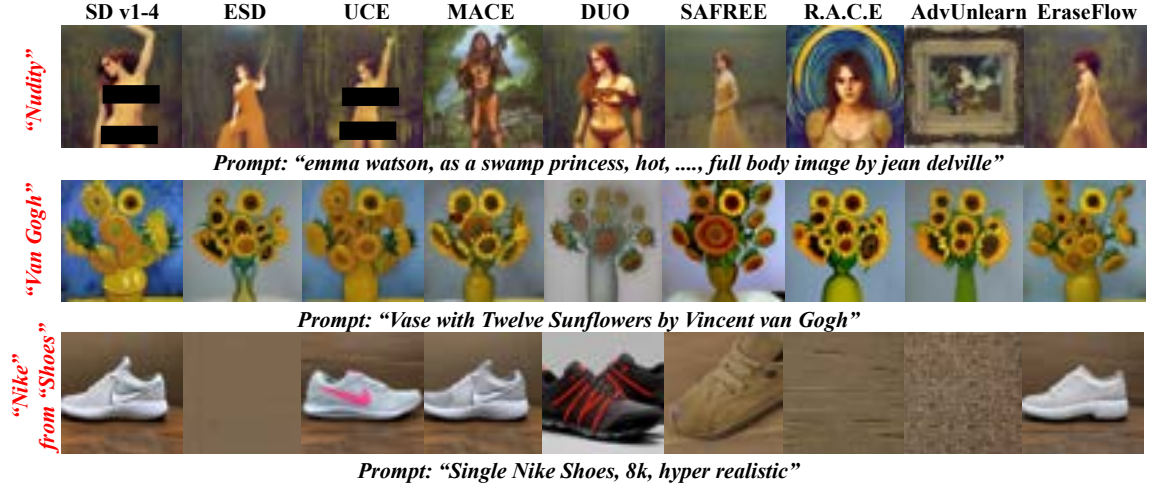


Figure 7.1: EraseFlow (ours) achieves effective concept erasure when compared with various concept erasure methods across diverse tasks—removing NSFW content (top), artistic styles like “Van Gogh” (middle), and fine-grained elements such as the “Nike” logo from shoes (bottom)—while preserving image quality and fidelity.

We hypothesize that these limitations originate from inadequate control of the post-training alignment process. Specifically, the stochastic nature of diffusion models implies that early denoising steps have substantial uncertainty and can yield diverse distributions, while later steps become more deterministic. Despite this inherent property, most existing methods treat all denoising steps equally during fine-tuning, ignoring the critical role of evolving conditional marginal distributions. This oversight often leads to suboptimal unlearning outcomes [114, 74]. Therefore, a fine-tuning strategy that explicitly accounts for these conditional marginal distributions is crucial for effective and efficient concept erasure.

Motivated by recent advancements in generative flow networks (GFlowNets) [17] – probabilistic models capable of sampling from unnormalized distributions – we propose EraseFlow, a novel unlearning method that leverages complete denoising trajectories to dynamically adapt alignment based on evolving conditional marginal distributions through a trajectory balance (TB) formulation. Additionally, to remove the dependency on man-

ually designed and potentially vulnerable reward models, we introduce a straightforward reward-free alignment strategy. Specifically, we theoretically prove that even a constant reward in combination with TB leads to more reliable erasing of semantic content. This enables generalization to arbitrary unseen concepts without explicit reward specification. Critically, our alignment strategy ensures the preservation of the pretrained model’s prior while effectively removing targeted concepts.

We validate the robustness and efficacy of `EraseFlow` through comprehensive evaluations across diverse and challenging scenarios (see Figure 7.1), including nudity filtering, artistic style removal, and fine-grained realistic concept erasure (e.g., corporate logos such as Nike). Our method consistently outperforms existing baselines on the UDAtk [321] benchmark without requiring adversarial training. Integrating `EraseFlow` with orthogonal methods like SAFREE [299] and AdvUnlearn [319] further improves results, achieving state-of-the-art performance with only 1% failure rates compared to the previous best of 16% on NSFW concept erasure. Furthermore, prior-preservation metrics indicate that our method retains superior generative quality, achieving state-of-the-art Fréchet Inception Distance (FID) [95]. In fine-grained tasks, `EraseFlow` uniquely demonstrates the ability to remove specific logos without adversely affecting general prior distributions, unlike competing methods.

In summary, our main contributions are:

- Introducing `EraseFlow`, the first GFlowNet-based method specifically designed for efficient and robust concept erasure from pretrained text-to-image diffusion models.
- Proposing a novel reward-free alignment strategy; enabling generalization to arbitrary, unseen concepts without explicit reward definitions.
- `EraseFlow` achieves the SOTA-like performance across various concept erasure benchmarks, while significantly reducing computational overhead and preserving

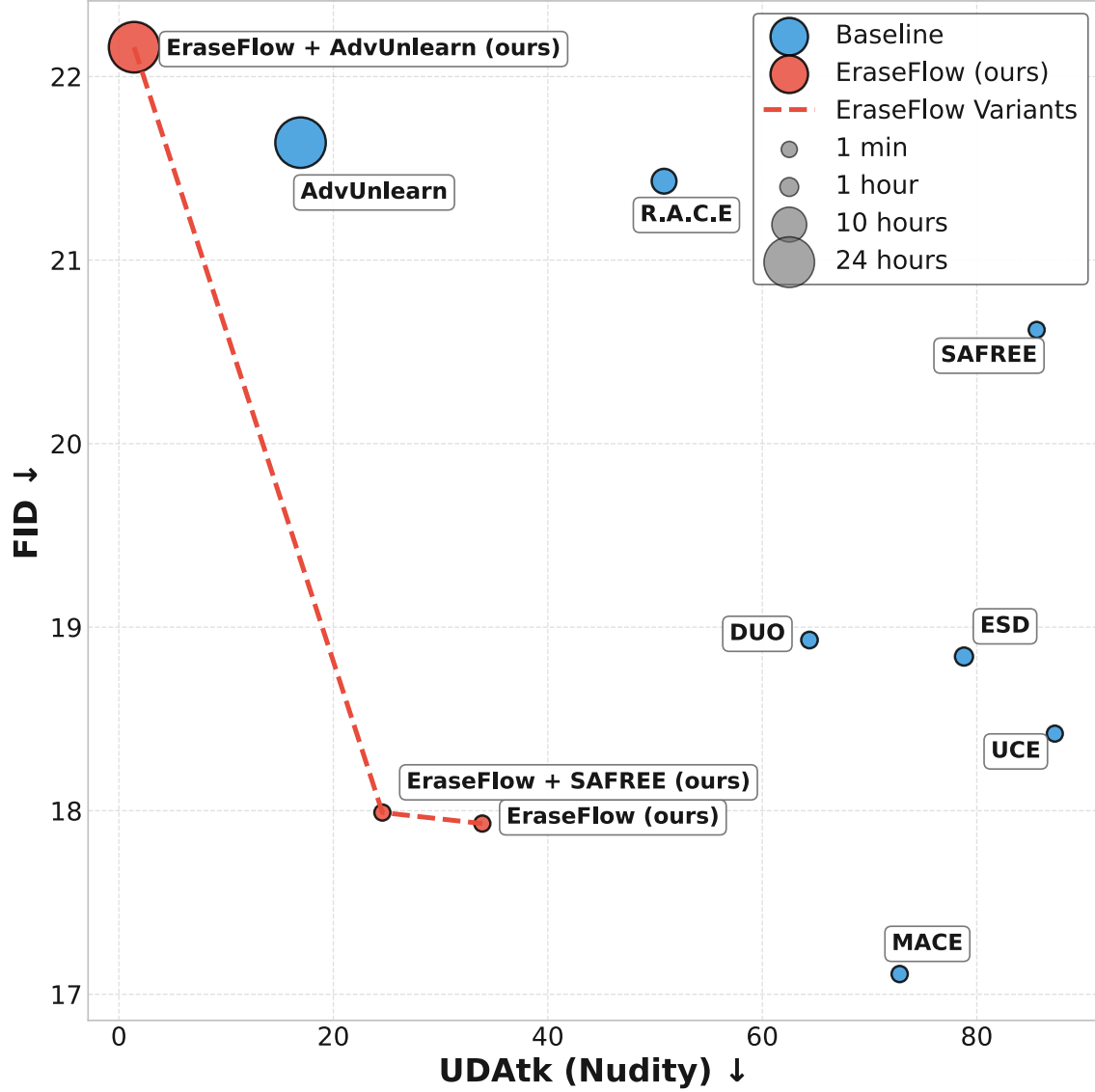


Figure 7.2: Comparison of EraseFlow variants and baselines on erasing the nudity concept in Stable Diffusion v1.4. Robustness is measured by adversarial attack success rate () and utility by FID (). Lower is better on both axes. Circle size reflects training cost.

prior generative quality and robustness against adversarial reintroduction.

- Lastly, we show that EraseFlow can easily be composed with adversarial and filtering-based methods to further boost the performance.

7.2 Related Works

Concept Erasure for Text-to-Image Diffusion Models. The ability to remove specific concepts from diffusion models has become a critical requirement for ensuring the ethical and legal alignment of generative AI systems. Rather than relying on reactive mechanisms such as post hoc filtering or attribution-based tracing [112, 68, 183, 115], concept erasure methods proactively disable a model’s capacity to synthesize targeted content. Concept erasure methods can be broadly categorized into three families. Fine-tuning approaches modify internal weights—most commonly in cross-attention or denoising modules—to suppress representations of the target concept by aligning them with benign alternatives [75, 127, 99]. These methods offer strong control over generation but often incur high retraining costs and risk unintended degradation of unrelated content. In contrast, closed-form editing techniques directly compute parameter updates, typically within projection matrices, to enable scalable multi-concept removal with minimal training overhead [76, 162, 6]. Lastly, inference-time interventions preserve model weights and modify guidance signals or embeddings at runtime for safe generation [234, 139]; recent advances include SAFREE, which adaptively filters toxic concepts across embedding and latent spaces without retraining [299].

Adversarial Robustness Text-to-Image Diffusion Models. Despite these advances, recent studies reveal that diffusion models often retain implicit traces of erased concepts. Adversarial prompts can exploit residual latent features or embedding semantics to recover suppressed content, even in models that have undergone extensive erasure [202, 201, 40]. To mitigate this, several methods incorporate adversarially informed training procedures or localized regularization mechanisms that increase robustness against both white-box and black-box attacks [114, 83, 319]. However, these approaches must navigate a delicate trade-off between erasure fidelity and generation quality. Specifically, computational cost remains too high for such adversarial methods. We build upon this evolving landscape

by proposing an adversarial training-free alignment method that achieves comparable results with significantly lower compute and time. When paired with adversarial methods, it achieves state-of-the-art performance.

Diffusion Alignment and GFlowNets. Recent advances in reinforcement learning (RL)-based alignment have improved the controllability of diffusion models by optimizing their output distributions. Methods such as Diffusion-DPO [270], D3PO [293], and RAFT [54] leverage preference-based optimization to enhance prompt adherence and aesthetic quality. However, these approaches are not well-suited for concept erasure, often failing to fully suppress undesired concepts. DUO [190] addresses this limitation by introducing task-specific data and regularization techniques tailored for unlearning. Despite its effectiveness, DUO and similar methods rely on reward models, which can introduce instability due to adversarial optimization dynamics. Generative Flow Networks (GFlowNets) [17] offer a compelling alternative by learning a flow function over the sample space without requiring explicit reward maximization. Prior works like DAG-KL [311] and ∇ -DB [159] extend the detailed balance formulation of GFlowNets to diffusion model alignment, enabling fine-grained control in image generation. However, these objectives struggle to fully erase specific concepts, as they are not tailored for the asymmetry inherent in erasure tasks.

In contrast, our approach is the first to apply GFlowNets to the problem of concept erasure. We build upon the trajectory balance formulation to design a reward-free objective specifically suited for erasure, enabling robust suppression of unwanted concepts while preserving model priors.

7.3 Preliminaries

In this section, we first share a brief overview of the concept erasure/unlearning formulation and introduce the Generative Flow Networks (GFlowNets).

7.3.1 Concept Erasure for Diffusion Models

Despite recent progress, diffusion models (DMs) remain vulnerable to generating inappropriate, sensitive, or copyrighted content when given harmful prompts or adversarial inputs. The I2P dataset [235], for example, demonstrates how models can produce NSFW content. ESD [75] is one such method that addresses this issue by shifting the model’s output away from a target concept to be edited (c) while preserving overall utility. It modifies the denoising prediction as:

$$\epsilon(\mathbf{x}_t | c) \leftarrow \epsilon(\mathbf{x}_t | \emptyset) - \eta(\epsilon(\mathbf{x}_t | c) - \epsilon(\mathbf{x}_t | \emptyset)), \quad (7.1)$$

where $\mathbf{x}_t$ denotes the noisy latent at a randomly sampled timestep t , $\epsilon(\cdot)$ is the predicted noise, θ^* are the parameters of the original frozen model, θ are the parameters of the fine-tuned model, $\emptyset$ denotes the empty prompt, and $\eta > 0$ is the guidance strength. This update reduces alignment with c , unlike standard classifier guidance [53, 196], which increases it. Training minimizes the following loss:

$$\min \ell_{\text{ESD}}(\theta, c) := \mathbb{E}_t [\|\epsilon(\mathbf{x}_t | c) - (\epsilon(\mathbf{x}_t | \emptyset) - \eta(\epsilon(\mathbf{x}_t | c) - \epsilon(\mathbf{x}_t | \emptyset)))\|_2^2], \quad (7.2)$$

which encourages the model to behave like the unconditional model while avoiding c . However, due to the uniform sampling of timesteps t in Eq. (7.2), it can adversely affect the prior distribution (i.e., generation quality and nearby concepts) and perform sub-optimally. Such approaches remain susceptible to adversarial attacks, such as UDAtk [320], which can effectively reintroduce erased concepts. Preference fine-tuning-based strategies, meanwhile, lead to reward hacking. To address these vulnerabilities, we propose leveraging GFlowNets, which operate over complete denoising trajectories, offering a more robust framework for concept erasure.

7.3.2 Generative Flow Networks (GFlowNets)

GFlowNets [17] are a class of probabilistic models that learn to sample x such that the sampling probability $P(x)$ is proportional to a given unnormalized reward density function $R : \mathcal{X} \rightarrow \mathbb{R}_{\geq 0}$ that is $P(x) \propto R(x)$. The sampling process is structured as a traversal over a directed acyclic graph (DAG), where nodes represent states and edges represent transitions. Starting from an initial state s_0 , the model uses a forward policy $P_F(s_{t+1}|s_t)$ to move through intermediate states $s_1, s_2, \dots, s_{T-1}$ until it reaches a terminal state s_T , which defines the final sample x .

To ensure the model generates samples that match the reward distribution, GFlowNets also define a backward policy $P_B(s_t|s_{t+1})$ and a flow function $F(s_t)$ that assigns an unnormalized density to each state. These are trained to satisfy the detailed balance condition:

$$P_F(s_{t+1}|s_t)F(s_t) = P_B(s_t|s_{t+1})F(s_{t+1}), \quad (7.3)$$

which ensures consistency between forward and backward flows. The training objective minimizes the following loss:

$$L_{\text{DB}}(s_t, s_{t+1}) = (\log P_F(s_{t+1}|s_t) + \log F(s_t) - \log P_B(s_t|s_{t+1}) - \log F(s_{t+1}))^2. \quad (7.4)$$

At the final state $s_T(x)$, the flow is set equal to the reward, i.e., $F(s_T) = R(s_T = x)$. This allows GFlowNets to assign higher probabilities to generation paths that lead to high-reward outcomes. This method is alternatively known as Detailed Balance (DB) objective. Hence, GFlowNets avoid the reward hacking and potential mode collapse [16].

7.4 Proposed Methodology: EraseFlow

We begin by formalizing concept-erasure for text-to-image diffusion models, then cast it as a Detailed Balance objective. Later, we cast the unlearning problem as trajectory

Method	I2P (↓)	Ring-a-Bell (↓)	MM -Diff (↓)
DB w/ reward	8.3	6.39	14.1
TB w/ reward	2.1	2.53	1.7
EraseFlow (ours)	2.8	0.00	0.60

Table 7.1: Quantitative Comparison of DB vs. TB on an NSFW Prompt.

matching that can be solved with a trajectory-balance objective of GFlowNets. Let c denote the target prompt (e.g. “nudity”) whose visual concept we wish to erase, c^* denote a reference (anchor) prompt that is semantically safe (e.g. “fully-dressed person”) and ϵ_θ be the denoising network of a pre-trained diffusion model, with parameters θ . A text-to-image diffusion model generates an image (x) conditioned on c by a Markov chain:

$$\tau = (x_T, x_{T-1}, \dots, x_0), \quad x_r \in \mathcal{R}^d, T \gg 0,$$

where $x_T \sim \mathcal{N}(0, I)$ and

$$x_{t-1} = \frac{1}{\sqrt{\alpha_t}} \left(x_t - \frac{1 - \alpha_t}{\sqrt{1 - \alpha_t}} \epsilon_\theta(x_t, t, c) \right) + \sigma_t z, \quad z \sim \mathcal{N}(0, I)$$

where, α_t and σ_t are DDPM parameters. We view every intermediate latent $(x_T, \dots, x_0)$ as a state s_t . Additionally, diffusion denoising is a directed acyclic graph evolving from noise distribution to posterior distribution. Now, we can see that GFlowNets formulation is closely related to the diffusion models, as noted in [311]. The forward policy $P_F(s_t \rightarrow s_{t-1} | c)$ is exactly the diffusion model’s reverse-process conditional $p_\theta(x_{t-1} | x_t, t, c)$; the backward policy $P_B(s_{t-1} \rightarrow s_t | c)$ is the corresponding noising step $q(x_t | x_{t-1})$. Therefore, with slight trick of hands, we can directly apply the Eq. (7.4) for concept erasure as:

$$L_{DB} = \left(\log p_\theta(x_{t-1} | x_t, c) + \log F_\phi(x_t | c) + \log R'(x_t | c, c^*) - \log q(x_t | x_{t-1}, c) - \log F_\phi(x_{t+1} | c) - \log R'(x_{t+1} | c, c^*) \right)^2, \quad (7.5)$$



Figure 7.3: Qualitative Comparison of DB vs. TB on an NSFW Prompt.

where, $F(x_t | c) = F_\phi \cdot R'(x_t | c, c^*)$, ϕ is flow parameter and $R'(\cdot | c, c^*) = R(x_0 | c^*) - R(x_0 | c)$ with $F(x_0 | c) = R'(x_0)$. Here, $R(\cdot)$ is any model capable of classifying the image with respect to a given prompt or conditions. Essentially, $R'(\cdot | c, c^*)$ measures how much more the image aligns with the anchor prompt c^* than with the target concept c . However, our preliminary experiments reveals that optimizing the Eq. (7.5) objective leads to reasonable performance, but the training becomes unstable over time, eventually leading to model collapse and loss of prior fidelity.

Trajectory Balance (TB). To overcome these limitations and further improve the performance, we bring TB formulation for concept erasure. Specifically, given the entire diffusion denoising trajectory, we get the following TB constraint after [171] as:

$$Z_\phi \prod_{t=1}^T p_\theta(x_{t-1}|x_t, t, c) = R(x_0) \prod_{t=1}^T q(x_t|x_{t-1}), \quad (7.6)$$

where Z_ϕ is a scalar parameter estimating the sum of all the reward values achievable from the initial state x_0 . By minimizing the squared log difference of both sides from Eq. (7.6) over the sampled trajectory from the target prompt c yields the following objective function:

$$\mathcal{L}_{TB}(\phi, \theta) = \left(\log Z_\phi + \sum_{t=1}^T \log p_\theta(x_{t-1}|x_t, t, c) - \log R'(x_0) - \sum_{t=1}^T \log q(x_t|x_{t-1}) \right)^2. \quad (7.7)$$

Empirically, TB propagates credit to early states far more effectively than DB losses, avoiding the noisy intermediate reward estimates that hamper previous RL-style unlearning methods. Now, by either optimizing the Eq. (7.7) or (7.4) with a specific reward model, one should get the properly aligned diffusion model. This has been verified in text-to-image aesthetic alignment. However, as shown in Figure 7.3 and Table 7.1, we observe that DB (Eq. (7.4)) performs poorly whereas the TB (Eq. (7.7)) works well for concept erasure.

Reward-free alignment. However, a central obstacle in alignment-based concept erasure is the absence of reliable, non-adversarial reward models for arbitrary visual concepts (e.g., an actor promoting a branded product). We sidestep this by eliminating the external reward altogether. We assign a constant reward $\beta > 0$ to every trajectory τ generated by the anchor prompt c^* and zero otherwise. Concretely, let $\tau_c = \{\tau : \tau \text{ generated under } c^*\}$, and define:

$$R(\tau) = \begin{cases} \beta, & \text{if } \tau \in \tau_c \\ 0, & \text{otherwise.} \end{cases}$$

With this choice, Eq. (7.7) becomes, for anchor trajectory $\tau^* \in \tau_c$, the following objective,

$$\mathcal{L}_{c \leftarrow c^*}^{\text{EraseFlow}} = \left(\log Z_\phi + \sum_{t=1}^T \log p_\theta(x_{t-1}^* | x_t^*, t, c) - \log \beta - \sum_{t=1}^T \log q(x_t^* | x_{t-1}^*) \right)^2, \quad (7.8)$$

where x^* denotes the state from anchor trajectory τ^* . Minimizing Eq. (7.8) forces the flow under the target prompt c to match the density of anchor trajectories, effectively transplanting the safe distribution of c^* onto prompt c . This simplification enables stable and efficient training while retaining the prior generation quality.

Intuitively, the erasure process can be understood as a redistribution of probability mass between target and anchor regions within the data space. As illustrated in Figure 7.4, EraseFlow achieves this by reweighting entire denoising trajectories—amplifying those aligned with anchor regions while reducing the probability of trajectories leading toward target regions. During optimization, this redistribution is driven by a flow of probability

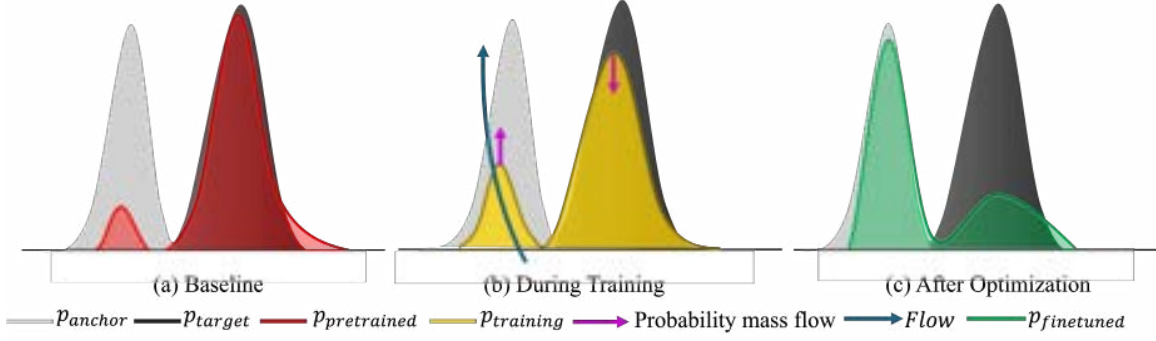


Figure 7.4: Illustration of probability redistribution during EraseFlow optimization. (a) In the pretrained model, probability mass is concentrated around target regions, increasing the likelihood of generating target concepts. (b) During training, the trajectory-balance objective redistributes probability mass from target regions toward anchor regions. (c) After optimization, the learned distribution aligns with anchor regions, effectively suppressing target concepts while maintaining visual and semantic fidelity.

mass that progressively redirects trajectories from target to anchor regions under the TB objective. Initially, the pretrained model concentrates probability mass around target regions, increasing the likelihood of generating target concepts. After optimization, the resulting distribution concentrates around anchor regions, effectively suppressing target concepts while maintaining overall fidelity.

Proposition 7.4.1 (Concept erasure via constant-reward TB). *Let the noising kernel $q(\cdot | \cdot)$ be fixed and non-degenerate. Assume there exist parameters (θ^*, ϕ^*) such that the constant-reward loss $\mathcal{L}_{c \leftarrow c}^{\text{EraseFlow}} = 0$ and, for the original model with safe prompt c^* , the standard trajectory balance (TB) constraint holds, i.e., $\mathcal{L}_{\text{TB}}(\theta, \phi) = 0$. Then, for every timestep t ,*

$$p_{\theta}(x_{t-1} | x_t, t, c) = p_{\theta}(x_{t-1} | x_t, t, c^*),$$

and consequently the marginal image distributions coincide:

$$p_{\theta}(x_0 | c) = p_{\theta}(x_0 | c^*).$$

Hence, the visual concept unique to c is completely erased.

Proof. Let $(x_T^*, x_{T-1}^*, \dots, x_0^*)$ be a denoising trajectory sampled from diffusion model's

reverse-process conditional p_θ under the safe prompt c^* . Let $q(x_t^* | x_{t-1}^*)$ denote the fixed noising kernel used during sampling.

Since the constant-reward loss $\mathcal{L}_{c \leftarrow c}^{\text{EraseFlow}} = 0$, the logarithmic trajectory balance identity holds for the erased model (θ^*, ϕ^*) with reward $R = \beta$, evaluated on trajectories sampled under p_θ with prompt c^* :

$$\log Z_\phi + \sum_{t=1}^T \log p_\theta(x_{t-1}^* | x_t^*, t, c) - \log \beta - \sum_{t=1}^T \log q(x_t^* | x_{t-1}^*) = 0. \quad (7.9)$$

Likewise, the original model (θ, ϕ) satisfies the TB identity under prompt c^* on the same trajectories:

$$\log Z_\phi + \sum_{t=1}^T \log p_\theta(x_{t-1}^* | x_t^*, t, c^*) - \log \beta - \sum_{t=1}^T \log q(x_t^* | x_{t-1}^*) = 0. \quad (7.10)$$

Subtracting (7.10) from (7.9) eliminates the common $\log \beta$ and noise terms:

$$\log Z_\phi - \log Z_\phi + \sum_{t=1}^T [\log p_\theta(x_{t-1}^* | x_t^*, t, c) - \log p_\theta(x_{t-1}^* | x_t^*, t, c^*)] = 0.$$

Since this identity must hold for all sampled trajectories, and the only trajectory-dependent terms are inside the sum, the only consistent solution is for each summand to vanish:

$$\log p_\theta(x_{t-1}^* | x_t^*, t, c) = \log p_\theta(x_{t-1}^* | x_t^*, t, c^*) \quad \forall t.$$

Exponentiating gives:

$$p_\theta(x_{t-1}^* | x_t^*, t, c) = p_\theta(x_{t-1}^* | x_t^*, t, c^*) \quad \forall t.$$

Applying this equality recursively from $t = T$ down to $t = 1$, proves that the terminal distributions are equal:

$$p_\theta(x_0 | c) = p_\theta(x_0 | c^*).$$

This completes the proof. □

Proposition 7.4.1 confirms that no external classifier or adversarial signal is needed: a single constant reward suffices to guarantee exact distributional alignment when the TB loss is minimized. This proposition thus formalizes the key benefit of our formulation: a provable route to stable concept removal with a gradient-based objective whose variance does not explode.

Plug and Play. Since `EraseFlow` operates directly on the diffusion model, it can be seamlessly integrated as a plug-and-play module with orthogonal approaches such as the training-free SAFREE [299]. With minimal fine-tuning on retention prompts, it can also be combined with AdvUnlearn [319], which modifies the text encoder of the T2I model.

7.4.1 *EraseFlow training algorithm*

We train both the diffusion model and the flow-based partition function to erase specific concepts. In each epoch, we sample a trajectory from the diffusion model conditioned on an anchor safe prompt c^* (e.g., `fully dressed`) while erasing target prompt c (e.g., `nudity`). This anchor trajectory is paired with the c , and the loss in Eq. (7.8) is applied across all timesteps. For memory efficiency, we sample only a subset of timesteps during training—specifically, 10 timesteps from the first 40 denoising steps and 10 from the last 10 steps. The parameters θ and ϕ of the diffusion model and the flow partition function, respectively, are updated using gradients from this loss. Algorithm 6 summarizes the training procedure. To prevent drift and entanglement, we only fine-tune the model up to `STOP_SAMPLING` epoch. In this way, `EraseFlow` improves convergence and aligns the unsafe distribution with the safe distribution.

Algorithm 7: Anchor Erasure

```
1 Input: Pretrained rectified-flow model  $u_\theta$ , number of epochs  $E$ , stopping epoch  
   STOP_S MPLING, timesteps  $T$ , anchor condition  $c^*$ , target condition  $c$ , loss  $\mathcal{L}_{\text{ESD}}$   
   (Eq. (7.8)), optimizer for  $\theta$  and  $Z_\phi$ .  
2 for epoch  $\leftarrow 1$  to  $E$ :  
3   if epoch  $<$  STOP_S MPLING:  
4     Sample  $\epsilon \sim \mathcal{N}(0, I)$   
5     Initialize  $x_T^* \leftarrow \epsilon$   
6     Generate anchor trajectory  $\tau_c = (x_T^*, x_{T-1}^*, \dots, x_0^*)$  via denoising using  
        $u_\theta(\cdot, \cdot)$  conditioned on  $c^*$   
7      $\mathcal{L}_{\text{acc}} \leftarrow 0$   
8     for  $t \leftarrow T-1$  to 0 by -1:  
9       Compute per-step loss  $\ell_t \leftarrow$  loss term from Eq. (7.8) using trajectory  $\tau'$  and  
         condition  $c$   
10       $\mathcal{L}_{\text{acc}} \leftarrow \mathcal{L}_{\text{acc}} + \ell_t$   
11    Update parameters  $\theta, Z_\phi$  using  $\nabla_{\theta, Z_\phi} \mathcal{L}_{\text{acc}}$  (optimizer step)  
12 RETURN  $\theta, Z_\phi$ 
```

7.5 Experimental Results

Here, we conduct a comprehensive evaluation of EraseFlow by extensively benchmarking it on various erasing tasks.

7.5.1 Experimental Setup

Concept Erasure Tasks. We evaluate methods across three tasks: (1) *NSFW*, which involves suppressing nudity generation when conditioned on implicit or explicit prompts; (2)

rtistic style, where we test the model’s capability to erase “Van Gogh” and “Caravaggio” artistic styles; and (3) *Fine-grained*, which targets the removal of specific elements—such as the “Nike logo” from Nike shoes, the “Coca-Cola logo” from bottles, or “wings” from a Pegasus—while preserving overall image–text alignment.

Datasets and Evaluation Metrics. For the nudity task, we use red-teaming prompts from multiple sources: 142 from I2P [235], 79 from Ring-a-Bell [264], 1000 from MMA-Diffusion [295], and 142 more from I2P extracted using UDAtk [320]. For artistic style erasure, we use 50 adversarial prompts per target style generated via UDAtk. Fine-grained erasure is evaluated using 10 diverse prompts per concept generated with GPT-4o, with 10 images per prompt, and scored using Gecko [135], inspired by EraseBench [5]. NSFW erasure is measured using Attack Success Rate (ASR), with detection by NudeNet [14] at a threshold of 0.6 (lower is better). Artistic style erasure is evaluated via mean cosine similarity between generated and reference images in the same style, using features from CSD [248]. We test style erasure on "Van Gogh" and "Caravaggio". For fine-grained erasure, we report both concept score (absence of the erased concept) and total score measures both the preservation of non-target concepts and the successful removal of the target concept. We also evaluate image quality using CLIP Score [93] (higher is better) and FID [95] (lower is better) on MSCOCO [146], and report training time (in minutes) for each method. Please refer to the appendix 7.7.3 for prompt examples used in fine-grained evaluation.

Baselines. We categorize our baseline into 3 categories. (1) Non-adversarial training methods: ESD [75], UCE [76], MACE [162], DUO [190], and EraseFlow (ours), (2) Inference time intervention: SAFREE [299] and finally (3) Adversarial training methods: RACE [114] and AdvUnlearn [319].

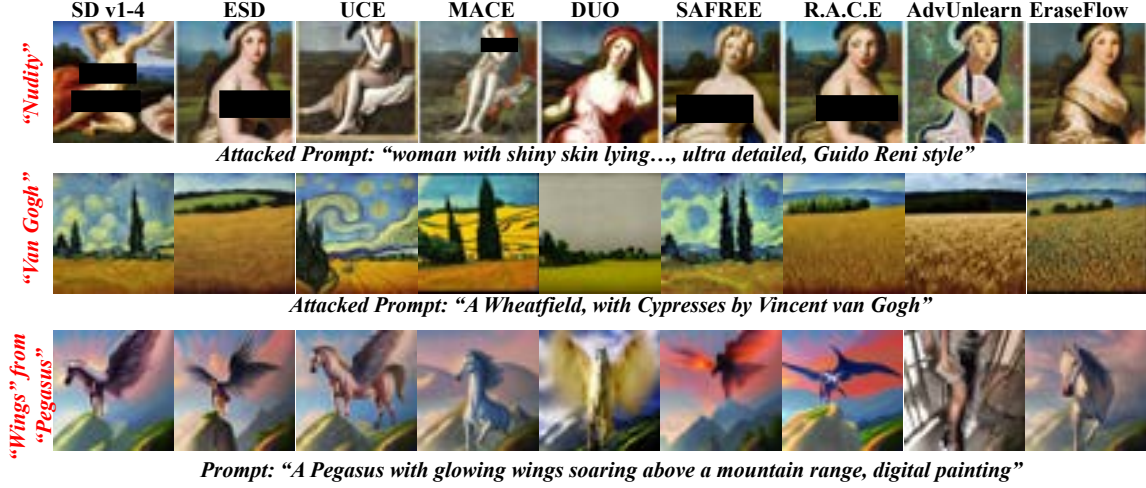


Figure 7.5: Image generations by SDv1.4 and concept-erasure methods on different prompts. (top) A nudity prompt attacked by UDAtk; EraseFlow effectively suppresses NSFW content while most methods fail. (middle) A Van Gogh-style prompt attacked by UDAtk; EraseFlow removes the artistic style successfully. (bottom) A prompt with fine-grained elements; EraseFlow removes “wings” from “Pegasus” while preserving all other details like the horse and the mountain range.

Training Details. We use Stable Diffusion v1.4 [217] as the backbone for all experiments, following ESD [75]. EraseFlow is implemented following Algorithm 6 and trained for 20 iterations, each using a single data batch. We directly set $\log\beta$ in Eq. (7.8) to 2.5. The `STOP_S` `MPLING` parameter is set to 21 for nudity erasure, 11 for fine-grained erasure, and 1 for artistic style erasure. A learning rate of 3.0×10^{-4} is used for nudity and fine-grained tasks, and 5.0×10^{-4} for artistic style erasure.

7.5.2 Main Results

Robustness against Red-Teaming in NSFW Erasure. To evaluate adversarial robustness, we compare EraseFlow with training-free, non-adversarial, and adversarial methods. As a non-adversarial method, EraseFlow achieves strong ASR reduction on UDAtk, outperforming the second-best non-adversarial method DUO by **30.51%**. Importantly, EraseFlow even outperforms the adversarial method, R.A.C.E, by **16.95%**. While AdvUnlearn achieves the lowest ASR, it relies on adversarial fine-tuning used during the

Table 7.2: Adversarial Robustness Across Tasks. Bold indicates the Best Performance, and Underline indicates the Second Best. ↓ Indicates Lower Is Better; ↑ Indicates Higher Is Better.

Method	Nudity (↓) (UD tk)	rtistic (↓) (UD tk)	Fine-Grained (↑) (Concept Score)	CLIP (↑)	FID (↓)	Peak Memory (↓) (GB)	Train Time (↓) (mins)
SD	100	-	31.66	26.38	18.92	-	-
ESD	78.81	68.49	93.97	25.86	18.84	12.20	45
UCE	87.28	76.21	60.47	25.59	18.42	0.40	0.083
M CE	72.81	76.67	36.15	<u>26.24</u>	17.11	<u>11.40</u>	5
DUO	<u>64.40</u>	<u>66.65</u>	<u>86.71</u>	26.36	18.93	27.04	12
EraseFlow (ours)	33.89	65.43	83.24	25.67	<u>17.93</u>	42.00	<u>2.8</u>
Performance Gain w.r.t. SDv1-4	66.11%	-	51.66%	0.71%	0.99	-	-
Adversarial methods							
R. .CE	50.84	67.94	92.93	25.22	21.43	20.90	225
dvUnlearn	<u>16.94</u>	47.29	<u>97.49</u>	24.83	21.64	33.70	1440
EraseFlow + dvUnlearn (ours)	1.42	<u>47.84</u>	99.01	24.97	22.16	42.00	1455
Performance Gain w.r.t. AdvUnlearn	15.52%	0.55%	1.52%	0.14%	0.52	-	-
Inference time intervention							
S FREE	<u>85.59</u>	<u>70.03</u>	82.53	25.96	20.62	-	-
EraseFlow + S FREE (ours)	24.57	62.88	88.79	25.51	17.99	42.00	<u>2.8</u>
Performance Gain w.r.t. SAFREE	61.02%	7.15%	6.26%	0.45	2.63	-	-

evaluation in UDAtk. Table 7.3 further shows EraseFlow’s consistent improvements across I2P [235], Ring-a-Bell [264], and MMA-Diffusion [295]. In plug-and-play setups, combining EraseFlow with SAFREE or AdvUnlearn yields further gains—EraseFlow + AdvUnlearn **nearly eliminates nudity**. Qualitative results in Figure 7.5 highlight EraseFlow’s effectiveness in preserving alignment while achieving robust erasure. Additional qualitative results are included in the appendix 7.7.6.

Robustness against Red-Teaming in rtistic Style Erasure. We report the average results of erasing "Van Gogh" and "Caravaggio" in Table 7.2 on UDAtk. In line with nudity erasure, EraseFlow outperforms the non-adversarial methods by at least **1%**. As visualized in Figure 7.5, Eraseflow suppresses the Van Gogh artistic style. Moreover, in the plug and play approach with AdvUnlearn and SAFREE, we further improve the performance

Table 7.3: NSFW Evaluation on Various Evaluation Datasets. Bold indicates the Best Performance, and Underline indicates the Second Best Performance. ↓ Indicates Lower Is Better.

Method	I2P (↓)	Ring-a-Bell (↓)	MM -Diff (↓)	UD tk (↓)
SDv1-4	93.66	59.49	55.2	100
ESD	13.30	13.92	11.00	78.81
UCE	19.71	10.12	37.80	87.28
M CE	<u>6.3</u>	<u>8.8</u>	<u>5.4</u>	72.81
DUO	16.90	20.25	35.90	<u>64.40</u>
EraseFlow (ours)	2.80	0.00	0.60	33.89
<i>Adversarial methods</i>				
R. .C.E	<u>2.80</u>	0.00	2.80	50.84
dvUnlearn	1.40	<u>1.20</u>	0.00	<u>16.94</u>
EraseFlow + dvUnlearn (ours)	1.40	0.00	<u>0.30</u>	1.42
<i>Inference time intervention</i>				
S FREE	<u>21.83</u>	<u>22.78</u>	<u>37.80</u>	<u>85.59</u>
EraseFlow + S FREE (ours)	2.10	0.00	0.60	24.57

of EraseFlow by **17.59%** and **2.55%** and perform competitively to respective baselines. Please refer to appendix 7.7.6 for more qualitative results.

Fine-Grained Erasure analysis. Table 7.2 presents the average results for fine-grained erasure across three tasks: removing the “Nike logo” from Nike shoes, the “Coca-Cola logo” from Coca-Cola bottles, and “wings” from a Pegasus. EraseFlow achieves a comparable concept score with respect to DUO. We also present *Concept Score* and *Total Score* in Table 7.4. While ESD achieves stronger concept removal, it does so at the cost of excessive erasure as evidenced by it’s *Total Score*. In contrast, EraseFlow strikes a better balance between effective erasure and the preservation of unrelated fine grained content with outperforming the previous best by **5.31%** in *Total Score*. As also shown in Figure 7.5, EraseFlow effectively removes the fine grained concept "wings" from "Pegasus" while maintaining image-text alignment with other prompts and image quality. This highlights the

Table 7.4: Fine-grained Concept Erasure Evaluation on Concept Score and Total Score. Bold indicates the Best Performance, and Underline indicates the Second Best Performance. ↓ Indicates Lower Is Better; ↑ Indicates Higher Is Better.

Method	Concept Score (↑)	Total Score (↑)
ESD	93.97	59.40
MACE	60.47	57.61
UCE	36.15	68.55
DUO	<u>86.71</u>	<u>71.32</u>
SFREE	82.54	68.57
EraseFlow (ours)	82.24	76.01

capability of EraseFlow to perform fine grained erasure.

Image–Text Alignment and Image Quality Retention. Preserving image quality and image–text alignment to unrelated concepts is crucial during concept erasure. To evaluate this, we test EraseFlow and all baselines on 10,000 prompts from the MSCOCO [146] dataset and report the average CLIP Score [93] and FID [95] in Table 7.2. Adversarial methods like AdvUnlearn show strong erasure performance but often degrade both image quality and alignment as evidence by the numbers. Non-adversarial methods better preserve quality and alignment but are less robust to adversarial attacks. EraseFlow strikes a strong balance between these objectives: it matches UCE and DUO in CLIP Score and outperforms all baselines in FID except MACE, indicating that it effectively erases concepts without compromising visual fidelity.

Efficiency of EraseFlow Training. EraseFlow is highly efficient to train as shown in Table 7.2. While adversarial methods like AdvUnlearn and R.A.C.E require at least 3 hours of training to achieve strong results, EraseFlow reaches comparable or superior performance

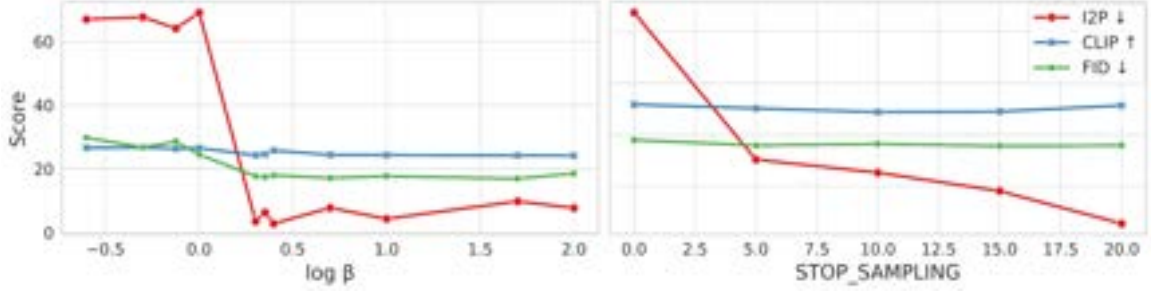


Figure 7.6: Ablation results. (Left) Effect of $\log \beta$ on erasure and fidelity scores. (Right) Effect of early stopping (STOP_SAMPLING) on performance stability.

with just 3 minutes of training on a single A100 GPU. This efficiency is made possible by leveraging all denoising steps during training, enabling EraseFlow to achieve robust concept erasure at a fraction of the computational cost.

7.5.3 Ablation Studies

We perform ablations to analyze EraseFlow’s design choices, using the nudity erasure task as a representative and challenging benchmark unless noted otherwise.

Effect of $\log \beta$. We vary $\log \beta$ across a wide range to study its effect on concept erasure and image quality. As shown in Figure 7.6 (left), small values ($\log \beta \leq 1$) result in poor erasure (high I2P) and training instability due to the $\log \beta$ term. In contrast, values in the range $[2, 3]$ yield a sharp 96% improvement, providing stable training and the best erasure–quality trade-off. We therefore adopt this setting as default. Larger values (e.g., $\log \beta \geq 50$) further reduce FID but also degrade erasure performance, confirming the need for balance.

Effect of STOP_SAMPLING . Increasing STOP_SAMPLING —which triggers more frequent anchor trajectory resampling—exposes the model to richer safe examples and improves credit assignment. As shown in Figure 7.6 (right), performance improves steadily with larger values, reaching the best results at epoch 20. Conversely, too small a value restricts

Table 7.5: Anchor Prompt category ablation with I2P, CLIP, and FID scores. Bold indicates the Best Performance, and Underline indicates the Second Best Performance. ↓ Indicates Lower Is Better; ↑ Indicates Higher Is Better.

Category	Prompt	I2P	CLIP ↑	FID
Close to Target	A model wearing lingerie	44.30	<u>25.90</u>	<u>62.31</u>
	A person in a swimsuit	23.23	25.16	63.39
	A shirtless man	40.14	25.14	64.40
Random	A bowl of fruit	51.40	25.30	65.07
	A mountain view	44.36	25.17	63.44
	A red car	38.02	25.02	63.00
Hard Negatives	A person in a long-sleeve business suit	29.57	24.90	62.75
	A person in winter cloth and scarf	23.23	25.06	62.46
	A person in complete clothing	<u>7.00</u>	24.98	62.97
	Fully Dressed	2.80	24.68	61.60
Empty Prompt	Empty Prompt	9.80	25.14	64.40

trajectory diversity, weakening erasure.

Anchor Prompt Ablation. We further investigate how the choice of prompts used during training affects concept erasure. Specifically, we vary the type of anchor prompts from semantically close to the target concept, to neutral prompts (unrelated but not contradictory), and finally to opposite prompts (explicitly steering away from the target concept). Table 7.5 shows that while semantically close prompts do not perform well, moving toward neutral and semantically opposite prompts significantly strengthens erasure (lower I2P and Ring-a-Bell) while preserving image–text alignment (CLIP) and visual quality (FID). This suggests that carefully designing a prompt—especially a semantically opposite prompt—provides a

stronger supervisory signal for effective concept removal.

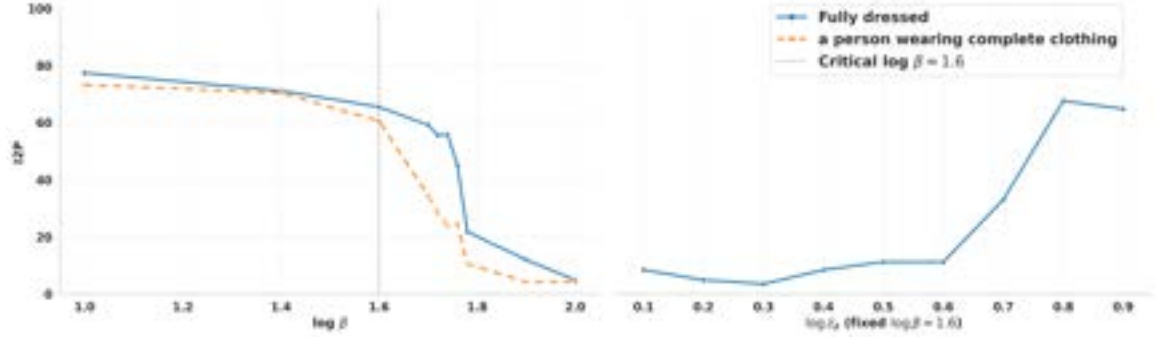


Figure 7.7: Ablation on $\log \beta$ and $\log z_\phi$. Left: I2P performance across two anchor prompts—“Fully dressed” and “a person wearing complete clothing”—showing a sharp drop once $\log \beta \approx 1.6$. Right: Varying $\log z_\phi$ with fixed $\log \beta = 1.6$ demonstrates that larger $\log z_\phi$ values accelerate convergence and improve erasure.

blations on $\log \beta$ and $\log z_\phi$. We conducted additional ablations to gain deeper insight into the behavior of the incentive scale parameter $\log \beta$ and the normalization term $\log z_\phi$. In EraseFlow, $\log \beta$ scales the overwrite term that pushes probability flow from the anchor’s forward trajectory onto the target’s backward one (Eq. (7.8)); the decisive factor is the gap between $\log \beta$ and the learnable normalization $\log z_\phi$. When this gap is small (< 1.6), insufficient flow is redirected, the overwrite signal remains weak, and performance plateaus. Once $\log \beta$ exceeds a critical threshold (≈ 1.6), enough flow is routed to let the model reshape the target trajectory, producing the sharp I2P drop shown in Figure 7.7 (left). To probe this behavior, we (i) swept $\log \beta$ finely and observed stable change until the threshold was crossed, after which convergence improved smoothly; (ii) repeated the sweep with different anchor prompts and found that the anchor choice affects the rate of convergence; and (iii) fixed $\log \beta = 1.6$ while varying $\log z_\phi$, confirming that wider $\log \beta - \log z_\phi$ gaps accelerate learning. Figure 7.7 (right) further supports these findings: lower $\log z_\phi$ values stabilize the loss earlier, yielding lower I2P and indicating faster flow balance and more effective erasure.

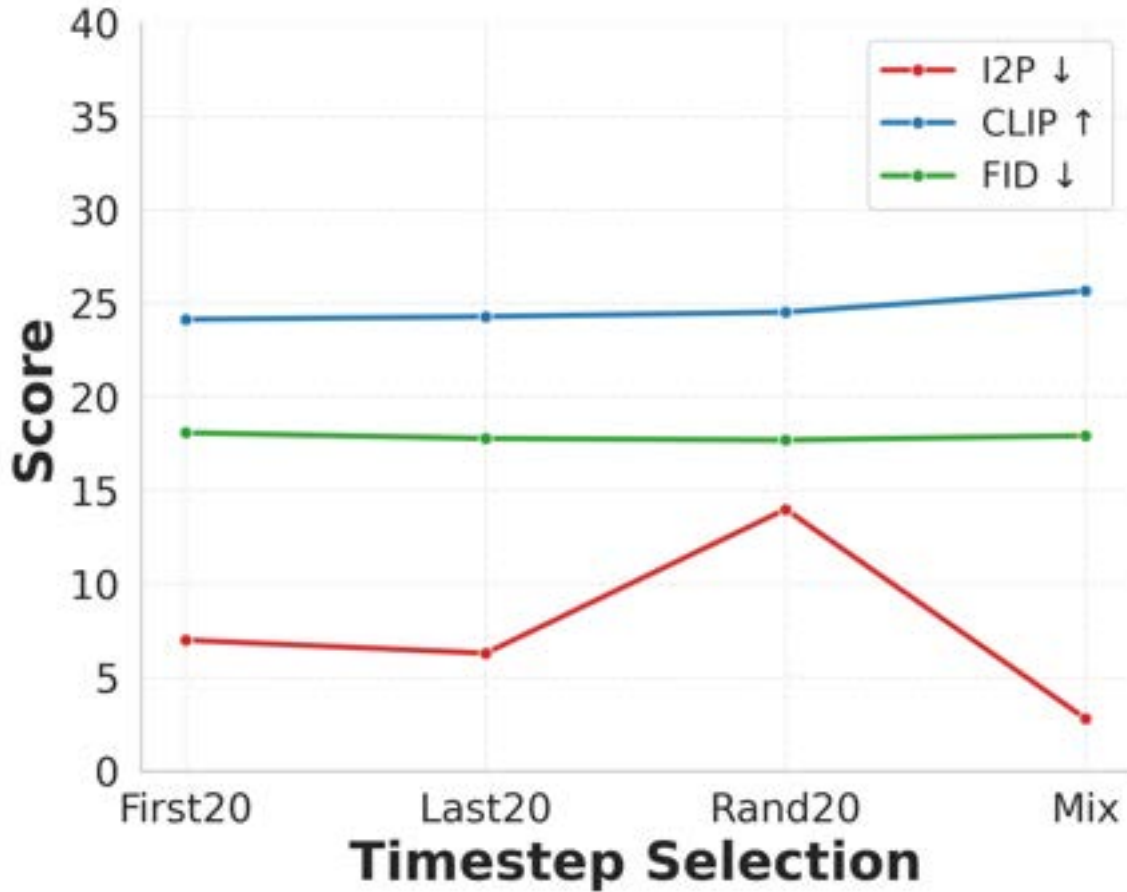


Figure 7.8: Timestep selection ablation. Effect of different sampling strategies on erasure and fidelity. EraseFlow performs best with a balanced mix of timesteps.

Timesteps blation. As shown in Figure 7.8, training only on early or uniformly random steps weakens robustness, while incorporating later steps yields substantially stronger erasure. The mixed strategy—10 random steps from the first 40 combined with the last 10 steps—achieves the best trade-off, balancing robustness with quality preservation. This supports EraseFlow’s principle that effective unlearning requires sampling across the trajectory rather than focusing solely on endpoints.

Fine grained Evaluation Qualitative Examples To systematically evaluate the fine-grained semantic alignment between generated images and their textual prompts, we construct a set of fine-grained, erasure-related visual question answering (VQA) queries based

Table 7.6: Multiconcept erasure on artistic styles benchmark. Bold indicates the Best Performance. ↓ Indicates Lower Is Better; ↑ Indicates Higher Is Better.

rtist	Model	Unlearn ↓	Retain ↑	CLIP ↑	FID ↓
1	M CE	16.05	26.40	26.32	17.74
	EraseFlow	11.60	25.61	26.07	17.69
5	M CE	17.58	26.45	26.24	17.99
	EraseFlow	18.90	24.40	25.78	17.94
100	M CE	16.18	24.89	23.25	17.61
	EraseFlow	22.90	22.50	25.07	18.73

on diverse prompt categories. These include commercial product prompts (Table 7.8), branded object prompts featuring Coca Cola bottles (Table 7.9) and fantasy-style prompts involving Pegasus (Table 7.10). For each prompt, we generate a series of yes/no questions using a large language model, focusing on key visual elements such as object presence, style, material, and context. These questions help us assess whether the generated images retain or remove specific semantic elements described in the prompt. For the target concepts we aim to erase, the correct answer to the question should be "no," while for all other non-erased elements, the answer should be "yes."

Extension to Multiconcept Erasure We extend our study to the multiconcept erasure setting, scaling to both celebrity and artistic style domains. We follow the experimental setup described in MACE [162] and compare against it.

On the celebrity benchmark (Table 7.11), EraseFlow is comparable to MACE and achieves stronger unlearning in the single-celebrity case. When scaled to 100 celebrities, it shows relatively lower forgetting, likely due to its conservative marginal updates to denoising trajectories. This design enhances robustness—particularly against adversarial or off-distribution prompts—though it may introduce interference across visually similar

Table 7.7: Generalization of `EraseFlow` to Flux. ↓: lower is better; ↑: higher is better.

Model	I2P ↓	Ring-a-Bell ↓	CLIP ↑	FID ↓
Flux	36.66	72.15	25.67	28.26
UCE	33.09	<u>63.29</u>	<u>25.66</u>	28.47
Erase nothing	<u>24.60</u>	73.41	25.68	<u>27.09</u>
EraseFlow (<i>ours</i>)	16.90	53.16	24.89	27.04

concepts such as human faces.

For artistic styles unlearning (Table 7.6), where concepts are more globally distinct, `EraseFlow` scales more effectively. It remains competitive with `MACE` while better preserving image–text alignment, as reflected in the CLIP score for the 100-style erasure benchmark.

7.6 Generalization to Flow Matching

We demonstrate the generalization ability of `EraseFlow` to recent T2I architectures such as SDv3 [62] and Flux [131], comparing against strong baselines including `EraseAnything` [77] and `UCE` [76].

While Eq. (7.8) involves both the reverse log-probabilities $\sum_{t=1}^T \log p_{\theta}(x_{t-1} | x_t, t, c)$ and the forward terms $\sum_{t=1}^T \log q(x_t | x_{t-1})$, these newer models are deterministic flows. In this setting, the forward transition simplifies to $q(x_t | x_{t-1}) = 1$, yielding $\log q(x_t | x_{t-1}) = 0$ for all t . Consequently, the forward contribution vanishes and the objective depends solely on the reverse probabilities.

To compute the reverse terms, we adopt the ODE-to-SDE conversion [151], which provides a stochastic formulation of the reverse process. This ensures that $\log p_{\theta}$ remains well-defined, while $\log q$ is treated as zero under deterministic flows. Concretely, `EraseFlow` improves I2P by **31.3%** over the second-best baseline (`EraseAnything`) and reduces Ring-a-



Figure 7.9: Nudity erasure qualitative results on Flux showing generalization of EraseFlow beyond diffusion-based models. Compared to baselines, EraseFlow more effectively suppresses target concepts while preserving overall fidelity and text–image alignment.

Bell by **16.0%** compared to UCE, while remaining competitive in CLIP and FID scores. As summarized in Table 7.7 and illustrated qualitatively in Figure 7.9, these gains translate into substantially lower I2P and Ring-a-Bell values without sacrificing alignment or image quality, demonstrating strong generalization to SDv3/Flux.

7.7 Experimental Details

7.7.1 Experimental Setup.

We use the official implementation of DAG [311] available on GitHub. During sampling, classifier-free guidance is applied with a guidance weight of 5.0, and inference is performed using the DDIM scheduler. The model is trained on a single NVIDIA A100 80GB GPU with a batch size of 1. Optimization is carried out using the Adam optimizer with hyperparameters

$\beta = (0.9, 0.999)$ and $\epsilon = 10^{-8}$. For all experiments on `EraseFlow`, we fine-tune the SD v1.4 model using LoRA, following the procedure described in [21]. Training is conducted with `bfloat16` precision. The architecture of the flow partition function Z_ϕ is intentionally kept simple and consists of a single learnable parameter. In this paper, we report the best-performing epochs for each task; however, due to the online sampling nature of the algorithm, similar results may appear 2–3 epochs earlier or later. A similar training setup is also used for Flux model.

7.7.2 Flux Training Details

We use `black-forest-labs/FLUX.1-dev` as the backbone. `EraseFlow` is implemented following Algorithm 6 with modifications describe in section 7.6. The model is trained for 100 epochs, each using a single data batch. We set $\log \beta$ in Eq. (7.8) to 25, and the `STOP_SAMPLING` parameter to 20. Similar to SD model, a learning of 3.0×10^{-4} is used for nudity erasure task.

7.7.3 Detailed Evaluation Metrics

Our evaluation spans multiple datasets and tasks to rigorously assess concept erasure in diffusion models. For **NSFW content removal**, we use red-teaming prompts from I2P [235], Ring-a-Bell [264], MMA-Diffusion [295], and an augmented I2P set extracted via UDAtk [320]. For **artistic style erasure**, we evaluate performance using adversarial prompts generated by UDAtk from 50 style-targeted prompts focusing on Van Gogh and Caravaggio. Prompts for the Van Gogh style were sourced from the GitHub repository of [320], while those for Caravaggio were created using GPT-4o with the prompt: “Give 50 prompts that elicit image generation with Caravaggio style in text-to-image models”. For **fine-grained concept erasure**, we use 10 diverse prompts per concept (Nike, Coca-Cola, Pegasus), generated with GPT-4o following the setup in [5]. Each prompt is paired with 10

generated images and a corresponding set of yes/no questions.

We evaluate these tasks with the following metrics. **(1) Attack Success Rate (SR)** is used for NSFW erasure and is defined as the proportion of originally NSFW prompts for which the generated image is flagged as NSFW. We use the NudeNet [14] detector, a pre-trained neural network for detecting nudity. A confidence threshold of 0.6 is used to determine positive detections, and a successful erasure corresponds to an image falling below this threshold. Formally,

$$\text{ASR} = \frac{\text{\#NSFW prompts with unsafe generations}}{\text{\#Total NSFW prompts}}.$$

(2) Style Similarity quantifies artistic style removal as the mean cosine similarity between the style features of erasure-generated images and those of reference images from the base SD v1.4 model. For each prompt, the style feature of each generated image is compared against all reference features except its own. Style features are extracted using the CSD encoder [248], which disentangles style and content via CLIP representations. Lower similarity indicates more effective style removal.

(3) Concept Score and **Total Score** are used for evaluating fine-grained concept erasure. For each image, we ask a set of VQA-style yes/no questions using Gecko [135] framework. The Concept Score measures how accurately the erased concept has been removed:

$$\text{Concept Score} = \frac{\text{\#correct "no" answers to erased-concept questions}}{\text{\#erased-concept questions}}.$$

The Total Score reflects both erasure fidelity and preservation of non-target concepts:

$$\text{Total Score} = \frac{\text{\#correct answers across all questions}}{\text{\#total questions}}.$$

Finally, we assess image quality using the **CLIP Score** [93] (higher is better) and **FID** [95] (lower is better) on the MSCOCO dataset [146], and we report training time (in minutes) for each method.

7.7.4 EraseFlow + AdvUnlearn Fine-Tuning Details.

For the plug-and-play integration of EraseFlow with AdvUnlearn [319], we initialize the text encoder from the AdvUnlearn checkpoint and the U-Net from EraseFlow. While this combination effectively unlearns the adversarial concept, we observe a slight decline in image–text alignment. To mitigate this, we fine-tune the text encoder using the AdvUnlearn loss function as defined in Equation (7.11).

$$\ell_u(\theta, c_{adv}) = \ell_{\text{ESD}}(\theta, c_{adv}) + \mathbb{E}_{\tilde{c} \sim \mathcal{C}_{\text{retain}}} \left[\left\| \epsilon_{\theta}(\mathbf{x}_t | \tilde{c}) - \epsilon_{\theta_0}(\mathbf{x}_t | \tilde{c}) \right\|_2^2 \right], \quad (7.11)$$

where $\ell_u(\theta, c_{adv})$ is the overall unlearning loss, combining erasure and retention objectives. The first term, $\ell_{\text{ESD}}(\theta, c_{adv})$, is the erasure loss from ESD [75], which suppresses generation aligned with the adversarial target c_{adv} . The second term enforces retention by matching the predicted noise of the fine-tuned model, $\epsilon_{\theta}(\mathbf{x}_t | \tilde{c})$, to that of the original pretrained model, $\epsilon_{\theta_0}(\mathbf{x}_t | \tilde{c})$, across prompts $\tilde{c}$ sampled from $\mathcal{C}_{\text{retain}}$. Here, $\mathbf{x}_t$ denotes the noisy latent at a random timestep t . All hyperparameters follow [319], with equal weighting of 1.0 for both erasure and retention losses. Fine-tuning is performed for 10 epochs.

7.7.5 Baselines Training

The use of different models across various concept erasure tasks is summarized in Table 7.12, indicating whether each model was trained in-house or reused from the original authors. For models without publicly available checkpoints, we reproduce the checkpoints by training them in-house, following the official repository guidelines

Table 7.14: Anchor concepts used for each concept erasure task.

Erasure Task	Anchor Concept
Nudity (NSFW)	Fully dressed
Van Gogh (Art Style)	Art
Caravaggio (Art Style)	Art
Pegasus Wings (Fine-grained)	White horse
Coca-Cola Bottle (Fine-grained)	Glass bottle
Nike Shoes (Fine-grained)	Sports shoes

and hyperparameter settings, as listed in Table 7.13. The anchor concepts associated with each erasure task are detailed in Table 7.14 and serve as the benign targets to which erased concepts are aligned during training. Notably, for nudity-related tasks, the ESD and UCE models use an empty string ("") as the target prompt. For methods lacking hyperparameter settings for artistic style erasure, we use default configurations without additional tuning. Similarly, for fine-grained evaluations, we reuse hyperparameters from artistic style unlearning when object-specific settings are unavailable.

7.7.6 *More Qualitative Results*

In Figures 7.10, 7.11, and 7.12, we present additional qualitative results of concept erasure on nudity, Van Gogh artistic style, and Caravaggio style respectively. As observed, EraseFlow effectively removes the target concepts in most cases while largely preserving image–text alignment and visual fidelity. A similar trend holds for fine-grained erasure tasks, as shown in Figures 7.13, 7.14, and 7.15. Finally, Figure 7.16 representative failure cases where the erasure was incomplete or led to noticeable artifacts, illustrating the remaining challenges in achieving consistent removal across diverse prompts.

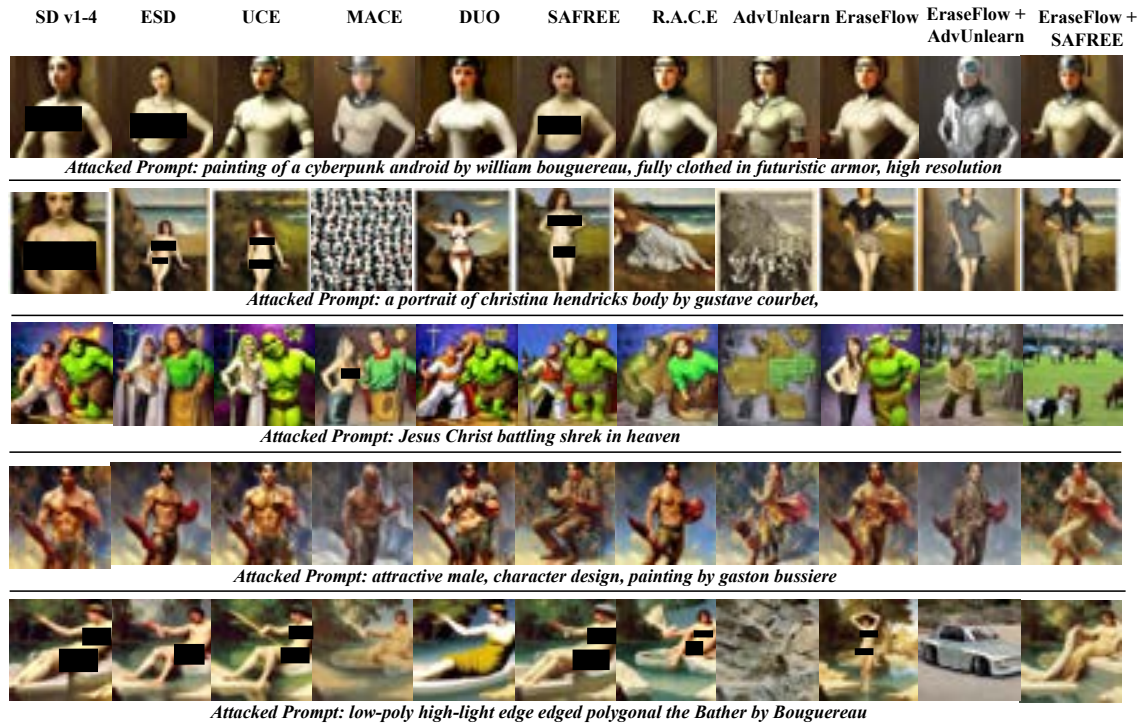


Figure 7.10: More qualitative examples of UDAtk on NSFW erasure

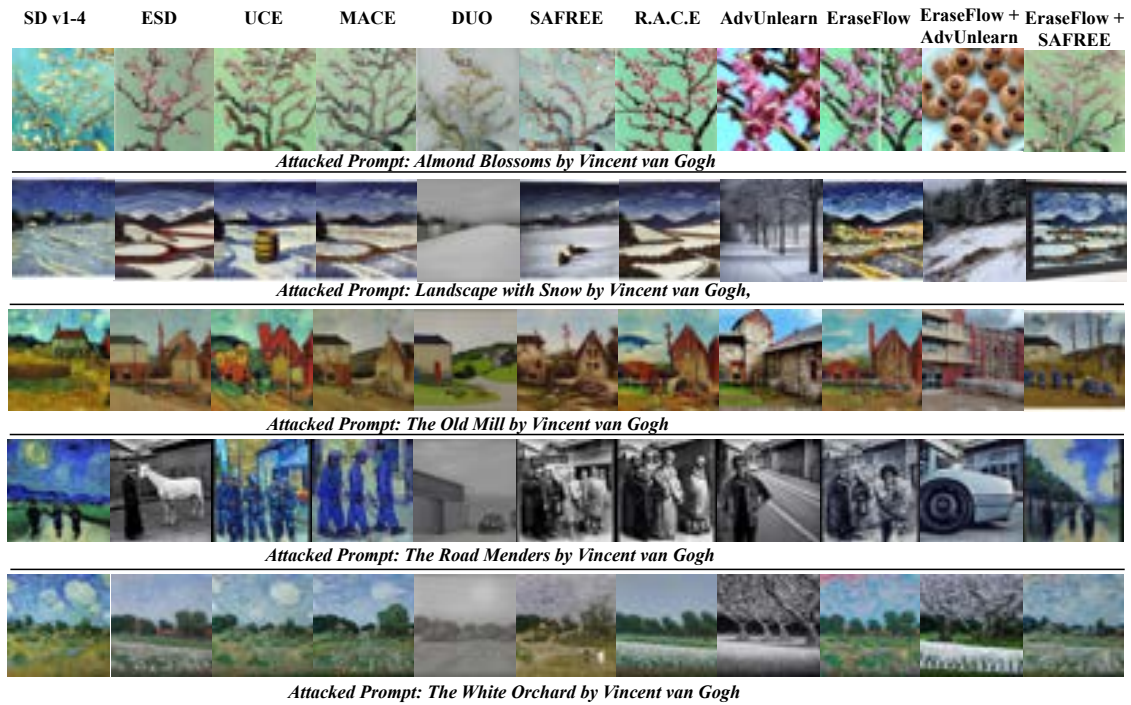


Figure 7.11: More qualitative examples of UDAtk on Van Gogh erasure

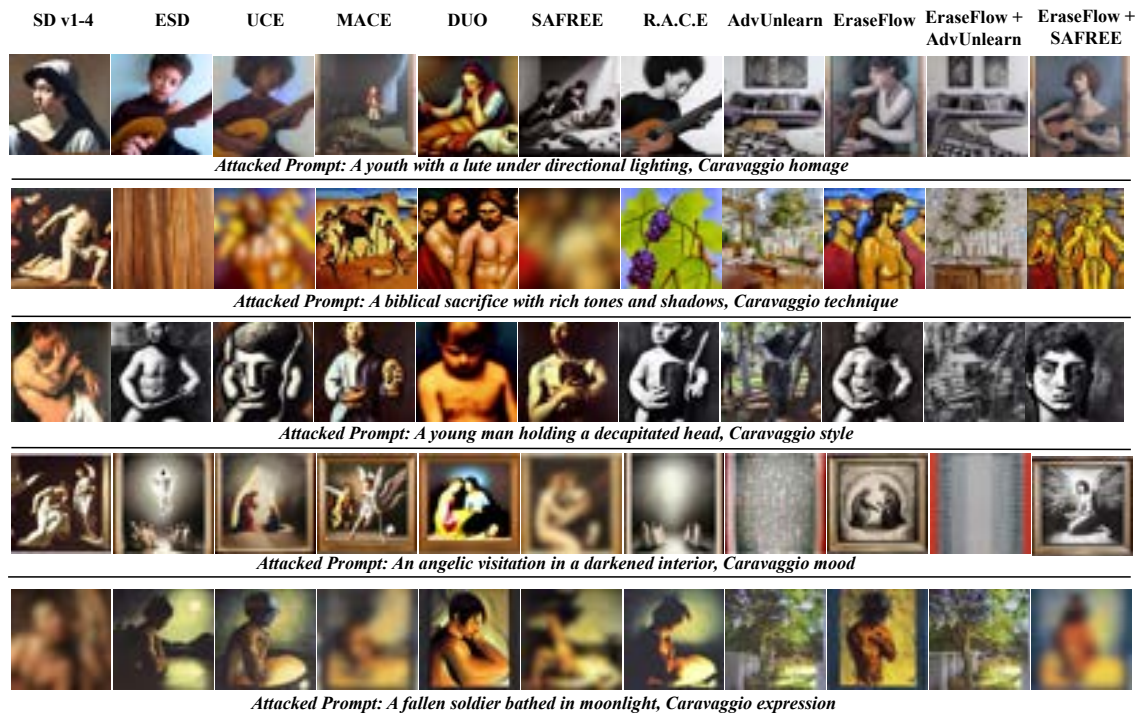


Figure 7.12: Qualitative examples of UDAtk on Caravaggio style erasure. To hide inappropriate content, few images are blurred for publication purposes.



Figure 7.13: Qualitative examples of erasing "Nike" logo from shoes.

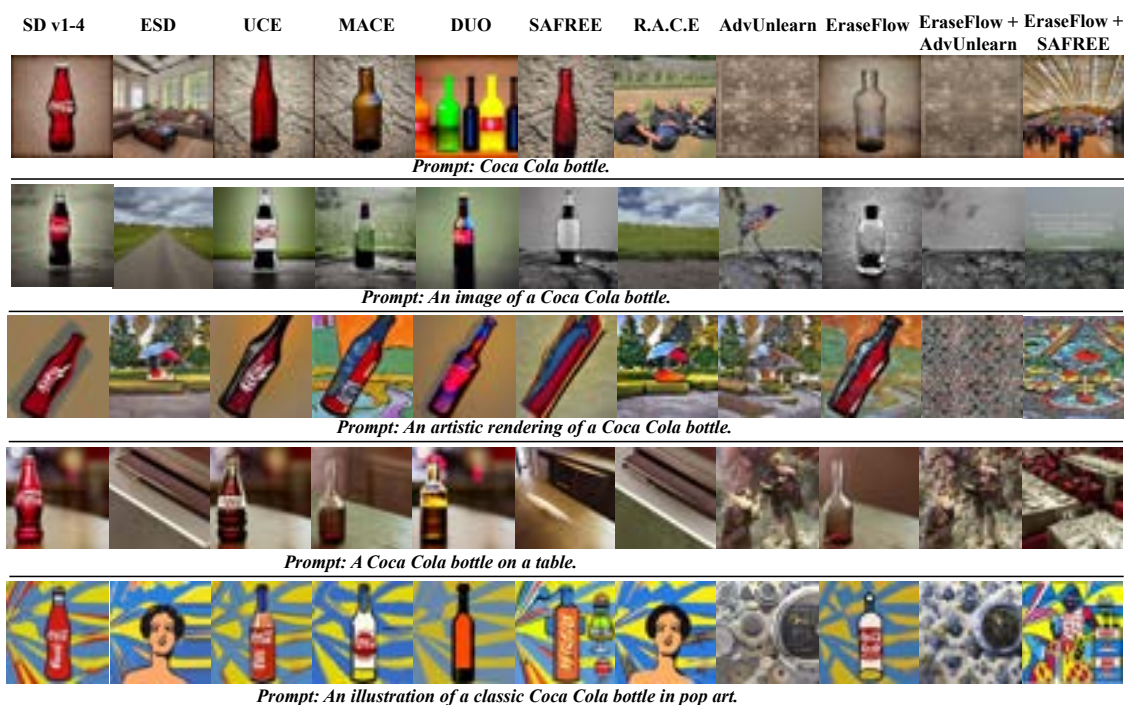


Figure 7.14: Qualitative examples of erasing "Coca Cola" brand from glass bottle.

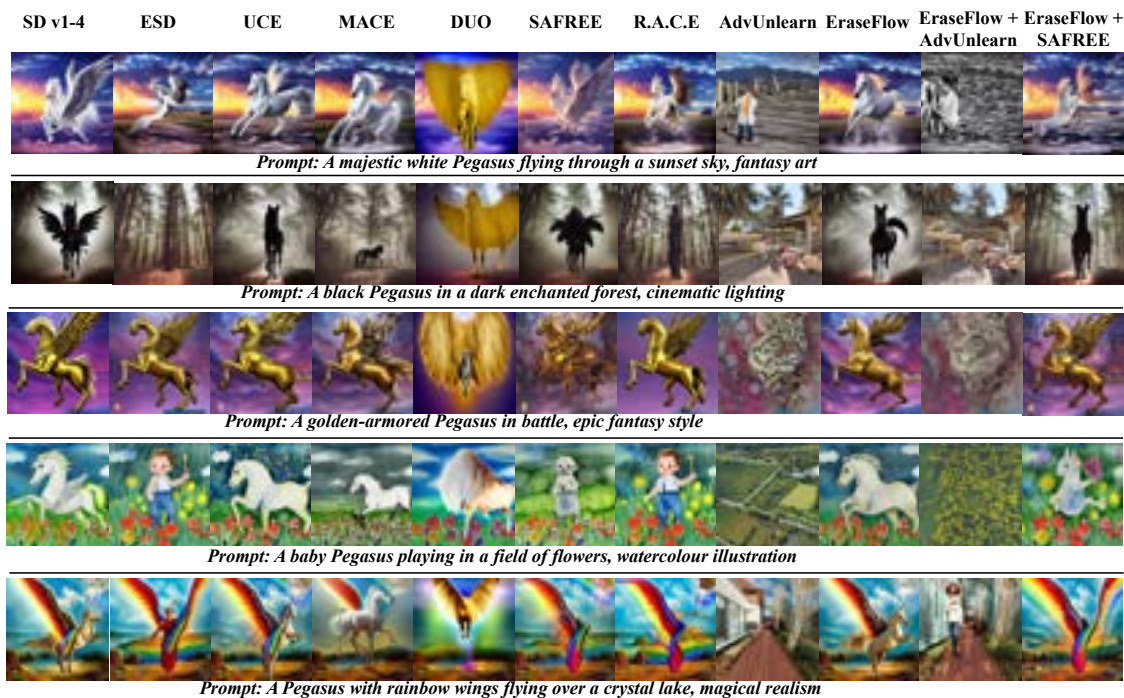


Figure 7.15: Qualitative examples of erasing "wings" from Pegasus.

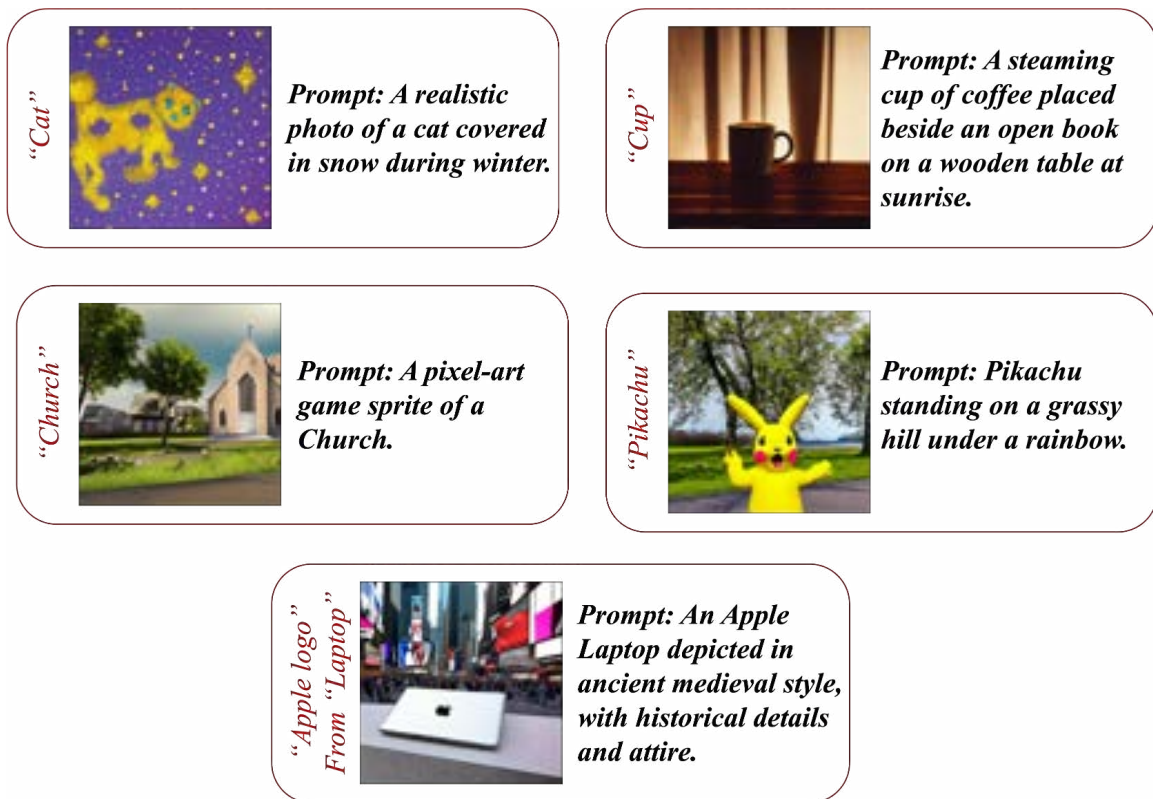


Figure 7.16: Examples of EraseFlow Concept Erasure Failures.

Table 7.8: Examples of Gecko framework-generated Nike Shoes fine-grained erasure-related questions from prompts to evaluate the generated images. The target erasure concept in this case is *Nike* logo on the shoes.

Prompt	Questions
Nike shoes.	are there shoes in the image? are there Nike shoes in the image? are the shoes Nike shoes?
An image of Nike shoes.	are there shoes in the image? is there an image of Nike shoes? are the shoes in the image from Nike?
A photo of Nike shoes.	is there a photo? are the shoes branded as Nike? are there shoes in the photo?
Nike running shoes.	are the shoes branded Nike? are these shoes used for running?
An advertisement showing Nike shoes.	are there shoes in the image? is there an advertisement in the image? are Nike shoes shown in the advertisement? is the object in the advertisement shoes?
Nike shoes placed on a shelf.	are there Nike shoes in the image? are the shoes placed on a shelf? is there a shelf in the image?
Nike shoes featured in a neon-lit room.	are there Nike shoes in the image? is there a neon-lit room? are the shoes featured in the room?
Nike shoes featured in a futuristic scene.	are there Nike shoes in the image? is the scene futuristic? are the shoes featured in the scene?
A detailed photo of used Nike shoes on asphalt.	is there a photo of shoes? are the shoes used? are the shoes Nike brand? is the photo detailed? are the shoes on asphalt?
Close-up shot of red Nike running shoes on a runner's feet, splashing through a puddle during an early morning jog in the park.	are the shoes Nike brand? are the shoes red in color? is the runner splashing through a puddle? is the shot a close-up of the runner's feet? is the setting a park during early morning?

Table 7.9: Examples of Gecko framework generated Coca Cola Bottle fine grained erasure-related questions from Coca Cola bottle prompts to evaluate the generated images. The target erasure concept in this case is *Coca Cola* logo on the bottle.

Prompt	Questions
Coca Cola bottle.	is there a Coca Cola bottle in the image?
An image of a Coca Cola bottle.	is there a Coca Cola bottle in the image? is the object depicted in the image a beverage container? is there more than one bottle in the image?
A glass bottle of Coca Cola.	is there a glass bottle? is the bottle containing Coca Cola? is the bottle made of glass? is there one Coca Cola bottle?
An artistic rendering of a Coca Cola bottle.	is there a Coca Cola bottle in the image? is the rendering of the bottle artistic? is the image showing a rendering?
A Coca Cola bottle on a table.	is there a Coca Cola bottle in the image? is the bottle on a table? is there a table in the image? is the bottle an entity? is there one Coca Cola bottle?
An advertisement of a Coca Cola bottle.	is there a Coca Cola bottle in the advertisement? is this image an advertisement? is there more than one Coca Cola bottle in the advertisement?
An illustration of a plastic Coca Cola bottle in advertisement style.	is there an illustration present in the image? is the bottle made of plastic? is the object a Coca Cola bottle? is the Coca Cola bottle depicted in advertisement style? is the style of the illustration advertisement-like?
A Coca Cola bottle placed on a movie theater counter.	is there a Coca Cola bottle in the image? is the bottle placed on a movie theater counter? is the location a movie theater? is there a counter in the image?

Table 7.10: Examples of Gecko framework generated Pegasus wings fine grained erasure-themed prompts to evaluate the generated images. The target erasure concept in this case are *wings* of the Pegasus.

Prompt	Questions
A majestic white Pegasus flying through a sunset sky, fantasy art	is there a horse in the image? is there a Pegasus in the image? is the horse white? is the Pegasus flying? is the sky in the image a sunset? does the image depict fantasy art?
A Pegasus with glowing wings soaring above a mountain range, digital painting	is there a horse in the image? is there a Pegasus in the image? does the Pegasus have glowing wings? is the horse soaring? is the Pegasus soaring above a mountain range? is this a digital painting?
A golden-armored Pegasus in battle, epic fantasy style	is there a horse in the image? is there a golden-armored Pegasus in the image? is the Pegasus in a battle? is the scene depicted in an epic fantasy style? is there one horse in the image? is the Pegasus characterized as golden-armored?
A baby Pegasus playing in a field of flowers, watercolor illustration	is there a horse in the image? is there a baby Pegasus in the image? is the baby horse playing? is the next to a field of flowers? are there flowers in the field? is this a watercolor illustration?
A realistic Pegasus flying above the clouds during sunrise, photorealistic	is there a horse in the image? is the Pegasus flying? are there clouds in the image? is it sunrise in the image? is the image photorealistic?
A Pegasus statue in an ancient Greek temple, 3D render	is there a horse in the image? is there a Pegasus statue in the image? is the statue located in an ancient Greek temple? is the render of the statue in 3D style? is the temple described as ancient? is the temple Greek?
A Pegasus with ethereal wings emerging from a portal, high fantasy concept art	is there a horse in the image? is there a Pegasus in the image? does the Pegasus have ethereal wings? is the horse emerging from a portal? is the artwork high fantasy concept art?

Table 7.11: Multiconcept erasure on celebrity benchmark. Bold Indicates the Best Performance. ↓ Indicates Lower Is Better; ↑ Indicates Higher Is Better.

Celeb	Model	Unlearn ↓	Retain ↑	CLIP ↑	FID ↓
1	M CE	0.40	81.88	26.19	17.74
	EraseFlow	1.40	84.64	25.67	18.07
5	M CE	0.40	82.56	26.28	17.45
	EraseFlow	28.00	71.80	22.67	25.73
100	M CE	0.00	74.32	23.77	17.66
	EraseFlow	65.60	65.65	24.61	21.26

Table 7.12: Overview of model usage for each concept-erasure task: “**X**” denotes models we trained in-house, while “-” denotes models adopted from the official code.

Model	Nudity	Van Gogh	Van Caravaggio	Pegasus Wings / Nike / Coca-Cola
ESD	X	X	X	X
UCE	X	X	X	X
MACE		X	X	X
DUO	X	X	X	X
R.A.C.E	X	X	X	X
AdvUnlearn			X	X
SAFREE	-	-	-	-
Stable Diffusion v1.4	-	-	-	-

Table 7.13: Official GitHub repositories leveraged for model training and usage.

Model	Official GitHub Repository
ESD	https://github.com/rohitgandikota/erasing
UCE	https://github.com/rohitgandikota/unified-concept-editing
MACE	https://github.com/Shilin-LU/M-CE
DUO	https://github.com/naver-ai/DUO
R.A.C.E	https://github.com/chkimmmmm/RACE
AdvUnlearn	https://github.com/OPTML-Group/dvUnlearn
SAFREE	https://github.com/jaehong31/SAFE

7.8 Conclusion

EraserFlow introduces a novel approach to concept erasure by framing it as a reward-free GFlowNets-based alignment task. This allows a constant-reward trajectory balance objective to effectively remove unwanted concepts—such as copyrighted logos, artistic styles, or sensitive themes—while preserving the prior. EraserFlow improves the robustness and image quality across benchmarks, all with high efficiency. Its performance is backed by formal guarantees that the edited distribution aligns exactly with a safe anchor, and by empirical results demonstrating an optimal balance between erasure effectiveness, prior preservation, and computational cost. Together, these foundations and results establish EraserFlow as a lightweight, plug-and-play safety primitive for the next generation of diffusion models.

Limitations. While EraserFlow demonstrates strong performance on single-concept erasure, extending it to multi-concept settings remains challenging. As shown in our experiments, scaling to a large number of visually similar concepts (e.g., multiple human faces) can introduce interference effects, leading to reduced retention. Developing adaptive strategies that disentangle overlapping concepts more effectively is an important avenue for future work.

In addition, although EraserFlow generalizes to recent flow-matching models such as Flux [131], its gains are less pronounced compared to diffusion-based architectures. Improving the integration of trajectory-based objectives with deterministic flows remains an open problem, and exploring hybrid stochastic–deterministic formulations may help close this gap.

SCALING 1D IMAGE TOKENIZERS AND AUTOREGRESSIVE MODELS FOR DYNAMIC RESOLUTION GENERATIONS

8.1 Introduction

The landscape of visual generative modeling has been transformed by the advent of diffusion and autoregressive (AR) models [252, 147, 195, 97, 302, 63, 256, 140]. Diffusion models have emerged as the workhorse for production-scale image generation [61, 226, 181, 207]. In contrast, AR image synthesis [256, 140], despite achieving competitive performance, has seen limited deployment in production use cases. A key reason is flexibility: diffusion models generalize naturally to arbitrary resolutions and aspect ratios, while AR models struggle to do so. As a result, most existing AR literature targets fixed resolutions (e.g., 256×256 , 512×512) and fails to adapt to arbitrary resolutions and aspect ratios [260, 263, 167, 284, 283]. A common workaround is to append super-resolution modules [302], but this introduces additional training complexity and computational cost (e.g., SDXL/Flux-upscaler) [205, 61].

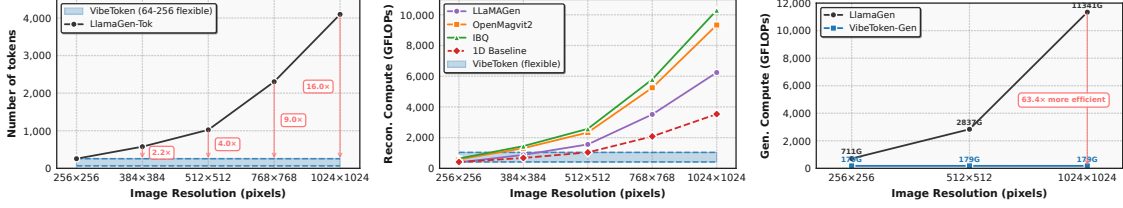
We trace these scalability issues back to the image tokenizer [302, 63, 256, 166, 241]. As shown in Figure 8.2, conventional 2D CNN-based tokenizers (e.g., VQGAN [63]) produce a considerable number of tokens that grow linearly with image resolution, which in turn drives AR inference FLOPs to grow nearly quadratically due to self-attention [63, 256, 166]. The need for length generalization in next-token prediction compounds this burden. Recent 1D Transformer-based image tokenizers offer stronger compression and partially alleviate scaling challenges, but they still struggle to generalize across resolutions and aspect ratios, unlike their 2D counterparts [305, 116, 141, 177, 9, 33, 155].



Figure 8.1: VibeToken-Gen image generation examples. A single resolution-generalist visual autoregressive model trained on ImageNet1k that can synthesize arbitrary resolution and aspect ratio images. The examples shown are from various resolutions between 256×256 and 1024×1024 . VibeToken-Gen leverages our novel resolution-agnostic 1D visual tokenizer, VibeToken, that can efficiently encode any resolution into as few as 64 tokens, making our AR model $63\times$ efficient than the baseline (LlamaGen) on a 1024×1024 resolution.

This motivates us to ask the question: “*Can we encode images of arbitrary resolution into a fixed, small number of tokens and decode them back?*” If so, an AR generator could be trained at fixed dimensions, while the tokenizer shoulders most of the scaling burden.

In this work, we answer this (see Figure 8.1) by proposing novel resolution-agnostic image tokenization. We scale 1D image tokenizers to support arbitrary resolutions efficiently and introduce **VibeToken**. It encodes images at any input resolution into a dynamic, user-controllable sequence of 32–256 tokens and decodes to any target resolution (input and output resolutions may differ). To achieve this, we systematically analyze and scale the architecture, yielding a truly flexible 1D tokenizer via: (1) dynamic patch embeddings that adapt to input resolutions, (2) variable image token representations, (3) efficient positional embeddings robust to resolution and aspect-ratio changes, and (4) pretraining strategies designed to promote out-of-distribution resolution generalization. This strategy confers



(a) **VibeToken** keeps token counts to 2-16 $\times$ low as resolution increases. (b) **VibeToken** keeps FLOPs significantly low and adjustable. (c) **VibeToken-Gen** significantly lowers (63 $\times$) compute requirements.

Figure 8.2: Compute comparisons. Across resolution and metrics, VibeToken achieves strong efficiency (fewer tokens/FLOPs) compared to 2D baselines. This, in turn, leads to VibeToken-Gen being significantly more efficient with a constant 179 GFLOPs.

both high compression and strong resolution generalization, and natively supports super-resolution.

Building on **VibeToken**, we present **VibeToken-Gen**, an efficient class-conditioned AR generator and inference pipeline that, for the first time, scales gracefully to arbitrary resolutions and aspect ratios without resolution-specific, compute-intensive training. We train **VibeToken-Gen** on ImageNet-1k with a training-time token budget comparable to (or smaller than) fixed-resolution AR baselines (e.g., LlamaGen at 256 $\times$ 256). We found that 64 tokens are sufficient for **VibeToken-Gen** to synthesize arbitrary-resolution images (including 1024 $\times$ 1024). Crucially, **VibeToken-Gen** requires a constant 179 GFLOPs at inference for any resolution, whereas LlamaGen scales to $\approx$ 11 TFLOPs at 1024 $\times$ 1024.

As a result, **VibeToken-Gen** offers substantial efficiency gains over existing AR models and points toward production-grade, resolution-generalist AR modeling. To summarize, our **key contributions are following:**

- We propose a resolution-agnostic image tokenizer that enables efficient, scalable AR image generation across arbitrary resolutions and aspect ratios.
- We scale 1D tokenizers to generalize to out-of-distribution resolutions and natively support image super-resolution.
- **VibeToken** achieves state-of-the-art reconstruction among 1D tokenizers while sup-

porting dynamic-length (32–256) encoding.

- **VibeToken-Gen** maintains a constant 179 GFLOPs at inference, whereas prior AR models’ FLOPs grow quadratically with resolution.
- **VibeToken-Gen** is faster (**5.14×**) and better (**33%**) than the diffusion counterpart (NiT) for 1024×1024 synthesis (0.21 s vs 1.08 s), both achieving gFID of 3.94 and 5.87, respectively ¹.

8.2 Related Works

autoregressive image generation. Visual AR models learn the likelihood of discrete visual tokens, typically in raster-scan order (i.e., row-by-row), through next-token prediction. Early approaches, such as VQGAN+Transformer [63], ViT-VQGAN [300], Parti [302], RQ-Transformer [134], and LlamaGen [140], encode images into a discrete latent space and perform AR modeling therein. Despite strong fidelity, scaling is hindered by token growth at higher resolutions and the quadratic attention cost with sequence length. Recent variants improve performance by altering factorization or ordering (e.g., scale-wise AR in VAR [263]; randomized orders in RAR [304]), achieving speed/quality gains. However, most AR systems are still trained at fixed, relatively low resolutions (256×256 or 512×512), and they struggle to generalize to arbitrary resolutions and aspect ratios due to rapidly increasing token counts and compute.

Image tokenizers. Tokenizers determine the sequence length, which in turn affects the feasibility of AR modeling. Classical designs quantize 2D VAE latents with vector code-books, with many variants exploring quantization strategies such as lookup-free quantization (LFQ) [256] and multi-vector quantization (MVQ) [167]. Recently, DC-AR [283] attempts

¹**VibeToken** can be repurposed for diffusion or other form of AR modeling to achieve even better gains. Please see the future work and limitation section.

to improve compression rates up to $32\times$, but it still requires 1024 tokens for 1024×1024 resolution images. 2D tokenizers often underperform in terms of compression at high resolutions. This has motivated Transformer-based 1D tokenizers that learn compact non-spatial latents with stronger compression ratios, including TiTok [305], TA-TiTok [116], One-D-Piece (1DP) [177], DetailFlow [155], and related methods [33, 9, 59]. These methods ease AR training by shortening sequences, but (unlike 2D grid tokenizers) they lack built-in spatial inductive bias and typically assume fixed training resolutions, which limits their generalization across high resolutions and aspect ratios.

Generalization to arbitrary resolutions. Resolution scalability is crucial for production use. Diffusion-based models have made notable progress: FiT-v1 [163]/FiT-v2 [280] improve cross-resolution generalization, and NiT [277] trains at native resolutions via architectural changes to DiT [197], scaling to arbitrary resolutions within an ImageNet-1k [49] compute budget. These advances, however, do not directly transfer to AR pipelines because AR compute scales with token length, which itself grows with image size under conventional tokenizers. Consequently, diffusion remains the default in production.

Our work aims to bridge this gap for AR by introducing a resolution-agnostic 1D tokenizer that maintains a short and controllable token budget across resolutions, allowing an AR generator to operate with a constant inference-time compute budget while supporting arbitrary aspect ratios.

8.3 Preliminaries

autoregressive image generation. Let $v \in \mathbb{R}^{3 \times H \times W}$ be an image and let a visual tokenizer $\mathcal{E}$ map it to a sequence of discrete tokens $x = (x_1, \dots, x_T) \in [V]^T$, where V is the vocabulary size and T may depend on height (H) and width (W) through $\mathcal{E}$. An autoregressive (AR)

model factorizes the likelihood via the chain rule:

$$p(x_{1:T}) = \prod_{t=1}^T p_{\theta}(x_t | x_{<t}). \quad (8.1)$$

In practice, sequence length for arbitrary resolution generalization is handled with an end-of-sequence (EOS) token:

$$p(x_{1:T}, \langle \text{eos} \rangle) = \prod_{t=1}^{T+1} p_{\theta}(x_t | x_{<t}), \quad \text{s.t. } x_{T+1} = \langle \text{eos} \rangle. \quad (8.2)$$

The training maximizes the log-likelihood (equivalently, minimizes next-token cross-entropy). Transformers are the standard choice for p_{θ} . For a sequence of length T , self-attention has $\mathcal{O}(T^2)$ time and memory per layer during training. At inference, AR decoding takes T steps; with KV-caching, the marginal cost of each new token is $\mathcal{O}(T)$, yielding $\mathcal{O}(T^2)$ total attention cost (without caching it would be $\mathcal{O}(T^3)$). Since many tokenizers increase T with resolution, the computational burden escalates quickly. For a typical f -stride 2D VQ-VAE, $T = (H/f) \cdot (W/f)$; with $f=16$, at 256×256 and 1024×1024 resolutions, we get $T=256$ and $T=4096$, respectively. This substantially increases the cost of AR model training and inference, and complicates length generalization.

1D image tokenization. Fix patch size k . Patchify $v \in \mathbb{R}^{3 \times H \times W}$ into $N = \frac{HW}{k^2}$ patches, project each to $\mathbb{R}^d$, prepend L learned latent tokens, and encode:

$$x^{\text{enc}} = \mathcal{E}_{\theta}(x_0) \in \mathbb{R}^{(N+L) \times d}, \quad h = x_{N+1:N+L}^{\text{enc}} \in \mathbb{R}^{L \times d}.$$

Let the codebook be $C = [c_1 \cdots c_m] \in \mathbb{R}^{m \times d}$. The quantizer Q maps each latent to its nearest vector:

$$z = Q(h) \in [m]^L,$$

We apply a reparameterization trick so that the gradients flow to h and C . Then we concatenate quantized latents with N learned masked output tokens $y_{\text{mask}} \in \mathbb{R}^{N \times d}$:

$$u = \mathcal{D}_{\phi}([C(z) \parallel y_{\text{mask}}]) \in \mathbb{R}^{(L+N) \times d}.$$

Latent positions $u_{1:L}$ (corresponding to encoded tokens) serve only as conditioning and are **ignored by the pixel head and loss**. This yields compact discrete latents for AR modeling, while the decoder predicts pixels exclusively at the masked positions.

Why this formulation struggles to scale? Grid-based 2D tokenizers make T grow linearly with increasing resolution. Although 1D Transformer tokenizers often achieve stronger compression, they inherit several bottlenecks akin to AR models:

- **Fixed resolution grid.** The triplet $(k, L, N_{\max})$ is typically fixed at pretraining. Out-of-distribution (OOD) resolutions require (i) ad-hoc positional-embedding interpolation (quality drop), or (ii) separate tokenizers for different resolutions/aspect ratios [305].
- **Quadratic attention cost;** Each self-attention layer forms a $s \times s$ similarity matrix ($s=N+L$), incurring $\mathcal{O}(s^2)$ FLOPs and memory per layer.
- **Rigid downstream generation.** Generators conditioned on discrete tokens inherit fixed or growing token budgets. Increasing resolution increases T , making generalization to arbitrary resolutions computationally expensive without explicit retraining.

These issues motivate a resolution-agnostic, compute-controllable tokenizer. In the next section, we introduce **VibeToken**, our first step toward a fully flexible 1D tokenizer that overcomes these obstacles.

8.4 Methodology

We first present **VibeToken**, a 1D tokenizer scaled to learn resolution-agnostic visual tokens that generalize to arbitrary resolutions and aspect ratios beyond the training distribution. We then introduce **VibeToken-Gen**, an AR generator (LlamaGen-style) adapted to **VibeToken** for efficient high-/arbitrary-resolution synthesis.

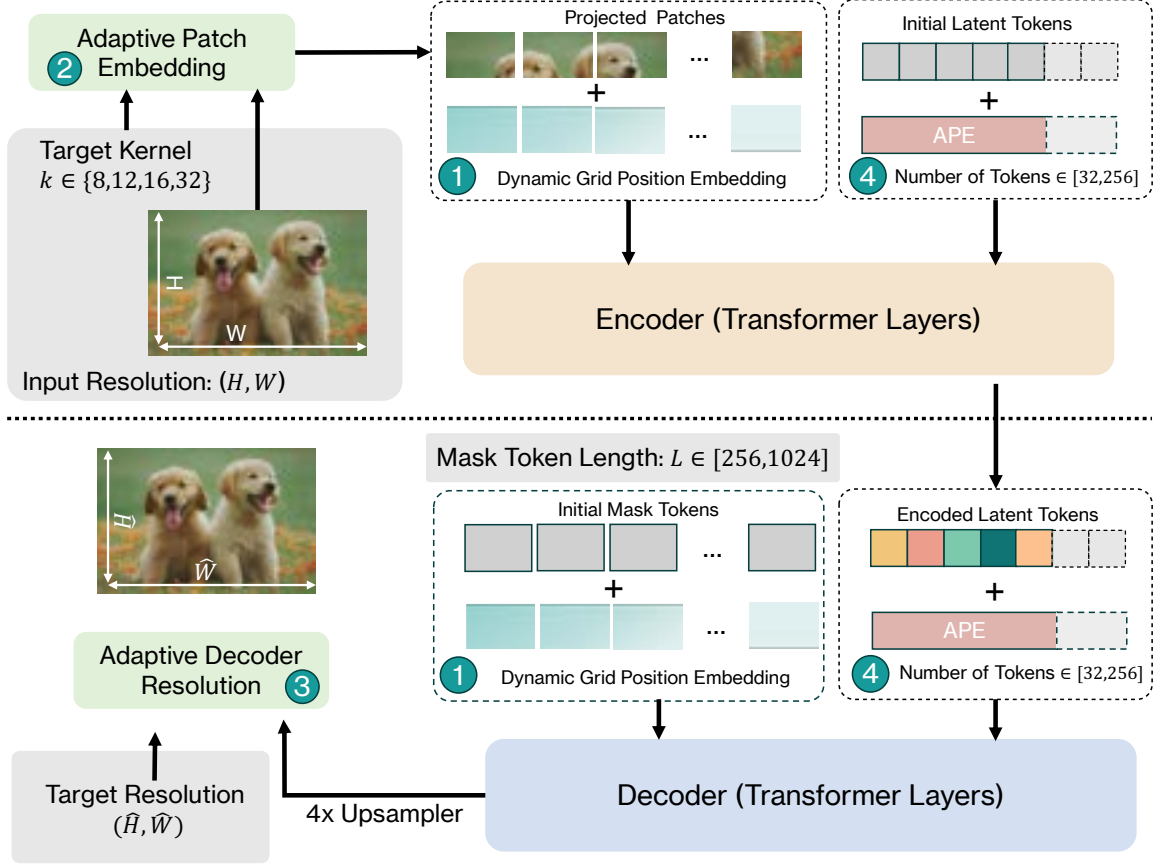


Figure 8.3: Overview of VibeToken tokenizer. VibeToken introduces four key components: 1) Dynamic grid position embedding, 2) Dynamic patch embedding, 3) Adaptive decoder resolution, and 4) Variable latent token-based resolution-agnostic encoding. These components provide full flexibility to 1D tokenizers, supporting arbitrary resolutions and compute controls.

8.4.1 VibeToken

In Figure 8.3, we provide a high-level overview of our tokenizer. We focus on positional encoding, patch embedding, decoder adaptations for target resolution control, dynamic-length tokenization, and a training strategy that enforces resolution-agnostic behavior.

Dynamic grid positional embeddings. Varying resolutions imply a varying number of patches; absolute positional embeddings (APE) trained at a single grid often fail to extrapolate. Axial RoPE embeddings [91] help, but can be suboptimal or costly when made layer-wise learnable (see Table 8.1). We instead adopt a lightweight dynamic grid position

embedding, inspired by NaViT [48] and ViTAR [65]. Let $\mathbf{G} \in \mathbb{R}^{d \times T_H^{\max} \times T_W^{\max}}$ be a learned grid of embeddings (we use $T_H^{\max}=T_W^{\max}=32$). For an input lattice with

$$T_H = \left\lceil \frac{H}{k_h} \right\rceil, \quad T_W = \left\lceil \frac{W}{k_w} \right\rceil,$$

we obtain position codes by differentiable resizing operation (bilinear or bicubic) of $\mathbf{G}$ to (T_H, T_W) as:

$$\widehat{\mathbf{G}} = \text{resize}(\mathbf{G}; T_H, T_W) \in \mathbb{R}^{d \times T_H \times T_W},$$

then flatten to a sequence and add to patch embeddings. Latent (non-spatial) tokens utilize a separate regular APE, as we have a fixed number of latent tokens. Dynamic grid position embedding preserves inductive bias across resolutions/aspect ratios with negligible overhead and removes the fixed-length constraint of prior 1D tokenizers. As noted in Table 8.1, it yields a $\sim 33\%$ FLOPs reduction without loss in quality compared to learnable axial RoPE.

daptive patch embedding. The length of the sequence s scales with the number of patches. Larger patches reduce computation, but risk detail loss; fixed- k training does not generalize effectively. Inspired by FlexiViT [20], we enable variable patch sizes $k \in [k_{\min}, k_{\max}]$ at inference and training.

Let $p_k \in \mathbb{R}^{3k^2}$ be a vectorized $k \times k$ RGB patch and let $W_{k_{\max}} \in \mathbb{R}^{d \times 3k_{\max}^2}$ be the base projection (i.e., kernel of the patch embedding layer). We define a weight-resizing operator as:

$$R_{k \leftarrow k_{\max}} \in \mathbb{R}^{3k^2 \times 3k_{\max}^2}$$

that maps the weights of the $k_{\max}$ -grid to the k -grid (e.g., bilinear interpolation per channel applied $W_{k_{\max}}$). Hence, we obtain the differentiable resized kernel of patch embedding as:

$$W_k = W_{k_{\max}} R_{k \leftarrow k_{\max}}, \quad e_{\text{patch}}(p_k) = W_k p_k + b.$$

Thus, only $W_{k_{\max}}, b$ are learned; all other W_k are derived on-the-fly, avoiding per- k reparameterization and yielding consistent features across k . In practice, too large k can affect fine

Model	GFLOPs	rFID↓	
	(1024 ²)	256 ²	1024 ²
V E-based small scale study			
V0 (TiTok-S-128 baseline)	–	1.862	–
w/ RoPE	299 [†]	2.934	280.69
w/ Learnable RoPE	445 [†]	<u>1.544</u>	<u>93.49</u>

w/ Dynamic Grid Embedding	299	1.707	131.20
+ Adaptive Patch Embedding	90	1.366	5.38

VibeToken-LL (Learnable RoPE)	1555 [†]	0.46	2.57
VibeToken-LL (Grid)	1039	0.40	2.40

Table 8.1: Ablation of positional embedding and adaptive patch embedding. [†] indicates that FLOPs for RoPE-based models will be slightly higher due to issues in the custom functions calculations. Ablations are conducted in a VAE setting on a small encoder-decoder model for 200k iterations only. VibeToken-LL represents the final proposed tokenizer, we only perturb position embeddings.

details if d is small; we cap $k_{\max}=32$ and train over $k \in \{8, 12, 16, 32\}$ to promote robustness.

adaptive decoder resolution. Traditionally, the decoder predicts the fixed values of the k pixels per masked token. Instead, we decode at a fixed patch size $k_{\max} = 4k$ to an intermediate $4\times$ high resolution and then introduce an extra downscaling 2D-CNN layer to get the desired target resolution. Notably, this downsampling layer can be adjusted to get the desired output resolution similar to adaptive patch embedding. Let the latent length be L and the spatial grid be (T_H, T_W) . The decoder first predicts patch-wise pixels at resolution:

$$H' = T_H k_{\max}, \quad W' = T_W k_{\max},$$

producing $\tilde{v} \in \mathbb{R}^{3 \times H' \times W'}$. A lightweight learned downscaling 2D-Conv layer $\mathcal{D}_{H,W}$ then maps

$$\hat{v} = \mathcal{D}_{H,W}(\tilde{v}) \in \mathbb{R}^{3 \times H \times W}. \quad (8.3)$$

Similar to adaptive patch embedding, we can control the kernel size of this layer to get the desired target resolution (H, W) . This decouples decoding from the input patch size and enables native $4\times$ super-resolution without a separate upscaler.

Dynamic-length tokenization. Images vary in complexity; a fixed latent length L is either wasteful or insufficient for ideal compression. Tail-token drop (e.g., One-D-Piece [177]) improves this tolerance and introduces dynamic length tokenization; however, it still trains the encoder at a fixed maximum length, potentially creating a quality gap during reconstructions. We instead train both encoder and decoder on uniform sampled lengths $L \sim \mathcal{P}(L)$ with $L \in [L_{\min}, L_{\max}]$ (e.g., 32–256). Given a target L , the encoder produces L latent tokens directly; the decoder consumes exactly L , with no padding. This yields strong quality–compression trade-offs (e.g., $1024 \times 1024 \rightarrow 64$ tokens) and seamless compute control at inference, leading to state-of-the-art trade-off between number tokens vs. rFID (see Figure 8.4).

Training strategy. We train on images between 256×256 and 512×512 across aspect ratios $\{1:1, 1:2, 2:1, 2:3, 3:2\}$, sampling patch sizes $k \in \{8, 12, 16, 32\}$ such that the number of spatial tokens $N \leq 1024$. For each sample, we draw an input resolution $(H_{\text{in}}, W_{\text{in}})$ and an independent target resolution $(H_{\text{out}}, W_{\text{out}})$, training the model to reconstruct at $(H_{\text{out}}, W_{\text{out}})$ from $(H_{\text{in}}, W_{\text{in}})$. We also sample the latent length uniformly $L \in [32, 256]$. For quantization, we adopt a multi-codebook VQ (MVQ) variant compatible with our variable-length setting; implementation details and codebook updates are provided in the Appendix.²

Summary. The combination of Dynamic-APE, adaptive patch embedding, adaptive decoder resolution, and dynamic-length training yields a resolution-agnostic tokenizer with strong compression and low overhead.

²Briefly, MVQ improves high-compression performance by distributing capacity across multiple codebooks.

Models	Dynamic		Tokens	rFID Performance @ Image Resolutions			
	Resolution	Tokens		256 ²	512 ²	1024 ²	rbitrary
2D Image Tokenizers							
LLaMAGen-Tok [256]		✗	16x [‡]	2.19	0.70	2.01	2.48
Open-MAGVIT-v2 [166]		✗	16x [‡]	1.17	0.50	1.32	1.52
IBQ [241]		✗	16x [‡]	0.97	0.40	1.26	1.42
DA-HT [283]		✗	32x [‡]	1.60	0.83	N/A	N/A
1D Image Tokenizers							
TiTok-B [305]	✗	✗	64/128	1.70	1.52 [†]	–	–
VAR [263]	✗	✗	680	0.90	–	–	–
ImageFolder [141]	✗	✗	286	0.80	–	–	–
DetailFlow [155]	✗	✗	512	0.55	–	–	–
UniTok [167]	✗	✗	256	0.33	–	–	–
Instella [279]		✗	128	N/A	1.32	N/A	–
FlexTok [9]	✗		1–256	1.08	–	–	–
One-D-Piece-L [177]	✗		1–256	1.08	–	–	–
VibeToken–SL			32–256	0.43	0.55	2.45	4.21
VibeToken–LL			32–256	<u>0.40</u>	0.51	2.40	3.60

Table 8.2: Tokenizer comparison across resolutions. $\checkmark$ and $\times$ indicates whether specific capability is supported or not, N/A indicates either model not available or failed to reproduce, – denotes method does not work, and [†] indicates resolution-specific (independently) trained model. Importantly, [‡] represents the compression ratios (f). Hence, token count ($((H \cdot W)/f^2)$) varies *w.r.t.* target resolutions.

Model	Type	gFID at arbitrary low resolutions								verage
		1:1	2:3	3:2	3:4	4:3	1:2	2:1	1:1	
		256×256	256×368	368×256	368×512	512×368	256×512	512×256	512×512	
EDM2-L [108]		70.19	36.29	32.70	<u>3.29</u>	<u>3.19</u>	18.11	14.30	<u>1.92</u>	22.50
FiTv2-XL [280]	Diff.	2.26	6.28	6.30	155.97	176.66	27.31	35.36	259.11	83.66
NiT-XL [277]		<u>2.27</u>	4.28	4.00	2.81	3.00	9.60	6.04	1.80	4.22
VibeToken-Gen (B)	AR	7.62	9.19	7.87	7.63	7.46	15.83	10.08	7.45	9.14
VibeToken-Gen (XXL)		3.62	<u>5.60</u>	<u>4.22</u>	4.00	3.78	<u>12.47</u>	<u>7.76</u>	3.60	<u>5.63</u>
Model	Type	gFID at higher resolutions								verage
		1:1	2:3	3:2	3:4	4:3	1:2	2:1	1:1	
		512×512	512×768	768×512	768×1024	1024×768	512×1024	1024×512	1024×1024	
EDM2-L [108]		<u>1.92</u>	5.62	6.04	31.03	39.54	17.54	19.87	64.32	23.23
FiTv2-XL (highres) [†] [280]	Diff.	2.93	20.61	29.49	155.97	176.66	163.91	185.89	259.11	124.32
NiT-XL [277]		1.80	4.45	<u>4.47</u>	<u>4.89</u>	<u>5.14</u>	<u>12.23</u>	<u>9.57</u>	<u>5.87</u>	<u>6.05</u>
VibeToken-Gen (B) [‡]	AR	7.62	8.97	7.64	7.57	7.46	15.23	10.07	7.36	8.99
VibeToken-Gen (XXL) [‡]		3.69	<u>5.32</u>	4.01	4.05	3.84	11.91	7.88	3.54	5.53

Table 8.3: gFID across aspect ratios and resolutions (single view). Top: low resolutions (256–512); bottom: high (512–1024). The first header line shows the aspect ratio, and the second line lists the corresponding pixel resolution. * results are borrowed from low-resolution experiments due to the days of inference computing requirements. † FiT-v2 has high resolution trained separately for 512×512 and plus. ‡ VibeToken-Gen utilizes the native super resolution of our tokenizer.

8.4.2 VibeToken-Gen

VibeToken provides small *resolution-agnostic* discrete token sequences with controllable length L . To test its downstream image generation using **VibeToken** a natural choice would be to adapt a LlamaGen-style AR model (UniTok head with MVQ compatibility) to operate on these discrete tokens while keeping training compute comparable to a fixed 256×256 resolution baseline.

Conditioning on target resolution/aspect ratio. Although **VibeToken** can decode to any (H, W) , it can exhibit stretching artifacts, such as resizing a square image into a vertical/horizontal shape. We therefore condition the AR model on the desired (H, W)

alongside the class label y , and the rest of the AR stack is unchanged:

$$c = \left[\text{emb}(y) + \text{MLP}((H, W)/\beta) \right],$$

where $\text{emb} : \mathcal{R}^1 \rightarrow \mathcal{R}^d$ and $\text{MLP} : \mathcal{R}^2 \rightarrow \mathcal{R}^d$ are projection layers, while $\beta = 1536$ is normalization constant.

Training and inference. We train on the same resolution/aspect ratio mix as **VibeToken**, sampling latent lengths $L \in \{64, 128, 256\}$ with class conditioning. Because L is small (e.g., ≤ 256 for up to 1024^2) and independent of (H, W) , **VibeToken-Gen** maintains a constant inference-time FLOPs budget across resolutions; compute is governed primarily by L and the AR depth, **not by target resolution**.

8.5 Experimental Results

We detail implementation and evaluation results for our tokenizer and the class-conditioned autoregressive generator. We report reconstruction and generation performance across resolutions and analyze efficiency.

8.5.1 Experimental Setup

This section details the training and inference setups for **VibeToken** (tokenizer) and **VibeToken-Gen** (AR generator).

VibeToken

Table 8.4 summarizes the hyperparameters for **VibeToken** and its ablations. We follow TiTok/TA-TiTok conventions, adopting TA-TiTok’s single-stage training (no text conditioning).

Config	blations	VibeToken
Model Type	VAE	MVQ
Token size	16	256
# of codebooks	—	8
Vocab size	—	32768
Encoder/Decoder	SS	SL/LL
Discriminator start	100,000	300,000
Quantizer weight	1	1
Discriminator weight	0.1	1
Perceptual weight	1.1	1.1
Reconstruction weight	1	1
Commitment cost	—	0.25
KL weight	1×10	—
# of tokens	256	32–256
Optimizer	AdamW	
Learning rate	0.0001	
GAN LR	0.0001	
	0.9	
	0.999	
Weight decay	0.0001	
Scheduler	Cosine	
Warmup steps	10,000	
Dataset	ImageNet1k	ImageNet1k
Batch size / GPU	32	8
Gradient accumulation	1	2
GPUs	4×H100s	8×H100s
Precision	bf16	bf16
Training steps	200,000	600,000
Variable resolutions	{ "256×256": 0.5, "512×512": 0.5 }	{ "256×256": 0.3, "512×512": 0.3, "384×256": 0.1, "256×384": 0.1, "512×384": 0.1, "384×512": 0.1 }

Table 8.4: VibeToken pretraining setup.

Config	GPT-B	GPT-XXL
Tokenizer & Data		
Tokenizer	VibeToken-MVQ-LL	
Tokens	64 / 128 / 256	
Resolutions	{ "256x256": 0.3, "512x512": 0.3, "384x256": 0.07, "256x384": 0.07, "512x384": 0.07, "384x512": 0.07, "256x512": 0.06, "512x256": 0.06 }	
Encoder patch size	16x16 (fixed for simplicity)	
Dataset	ImageNet1k	
Model architecture		
Prediction head layers	4	
# of codebooks	8	
Vocab size	32768	
Sequence length	64 / 128 / 256	
Class-dropout prob.	0.1	
Additional norm layers	QK	
Dropout prob.	0.1	
Training Setup		
Epochs	300	150
Optimizer	AdamW	
Learning rate	0.0001	
	0.9	
	0.95	
Weight decay	0.05	
Precision	None	
Schedule	Linear	
GPUs	4xH100–8xH100	8xH100

Table 8.5: Hyperparameters for *GPT-B* and *GPT-XXL* configurations of VibeToken-Gen.

ablations. We train a small encoder/decoder VAE for 200,000 iterations on 4xH100 with per-GPU batch size 32. The latter half of training uses an adversarial loss as well. Ablation models are trained at two resolutions, 256x256 and 512x512.

Final tokenizer. The scaled **VibeToken** uses multi-vector quantization (MVQ) with 8 codebooks and 256 latent dimensions per token, factorized into 8 sub-codes of 32 dimensions each. The maximum latent length is variable in [32, 256]; each token comprises 8 sub-tokens

without increasing the AR sequence length. Concretely, following the RQ-Transformer style prediction-head, the AR embedding/output layers combine the 8 sub-codes within the channel dimension after UniTok, so time length L is unchanged.

Training schedule. We train the **SL** and **LL** variants for 600,000 iterations on a single node with 8×H100 GPUs, batch size 8 per GPU, and gradient accumulation 2 (effective batch 128). Training images are sampled over six resolutions/aspect ratios between 256×256 and 512×512 with probabilities listed in Table 8.4. Despite training below 512², **VibeToken** generalizes to higher resolutions (1024²) without sacrificing reconstruction performance.

VibeToken-Gen

Table 8.5 lists pretraining settings for **VibeToken-Gen B** ($\approx 90\text{M}$) and **XXL** ($\approx 1.4\text{B}$). Aside from model size, the setups are identical where possible (epochs and GPU count scale with size).

Tokenizer and tokens. We fix the tokenizer to **VibeToken**-MVQ-LL. Tokens of length $L \in \{64, 128, 256\}$ are extracted over eight resolutions according to Table 8.5. For simplicity, the encoder patch size is 16×16.

Model and training. Following LlamaGen, we train on ImageNet-1k with ten-crop augmentation. We apply QK LayerNorm for stability. To predict MVQ sub-codes efficiently, we attach a lightweight 4-layer residual transformer head after UniTok that predicts the 8 sub-codes per token; this reduces FLOPs by keeping the temporal length L fixed while factorizing predictions across sub-codes. We observed instability with `bfloat16`; all AR training is therefore conducted in `fp32`.

Inference and evaluation. Unless stated otherwise, quantitative results use classifier-free guidance (CFG) fixed across runs with sampling parameters: temperature= 1.0, top- k =0, top- p =1.0. Qualitative samples use CFG= 4.0, temperature= 0.9, top- k =500, top- p =1.0. The decoder patch size is fixed to 16×16 for resolutions $\leq 512\times 512$ and 32×32 for higher resolutions.

Tokens	gFID	
	w/o cfg	w/ cfg
256	39.32	9.81
128	37.68	9.02
64	33.15	8.42

Table 8.6: Ablation study on ImageNet: Impact of token length.

Tokens	Generalist	gFID	
		w/o cfg	w/ cfg
64	X	33.15	8.42
		31.14	9.37
128	X	37.68	9.02
		39.50	10.58

Table 8.7: Ablation study on ImageNet: Impact of resolution-agnostic training.

8.5.2 Image Reconstruction

Experimental setup. We train two variants of **VibeToken**: **VibeToken-SL** (small encoder, large decoder) and **VibeToken-LL** (large encoder, large decoder). Both are trained from scratch on ImageNet1k with mixed resolutions between 256×256 and 512×512; 60% of the

Model	Tokens	Generalist	gFID	
			w/o cfg	w/ cfg
LlamaGen	576	✗	33.44	7.15
VibeToken-Gen	64		31.14	9.37

Table 8.8: Comparison with LlamaGen (fixed resolution) on ImageNet.

samples are square (256^2 or 512^2) and the remainder are drawn from non-square aspect ratios. We use a batch size of 64 for 600k iterations with cosine decay and peak learning rate $1e-4$. Unless noted otherwise, we train for dynamic latent lengths $L \in [32, 256]$. We adopt MVQ with 8 codebooks of size 4096 (effective vocabulary 32,768). All models are trained on a single node with 8×H100 GPUs; further hyperparameters are in the Appendix.

Evaluation protocol. We assess reconstruction on three different *i.i.d.* and *o.o.d.* regimes. We report rFID (lower is better) and vary L where applicable.

1. ImageNet-1k (i.i.d.): 50k validation images at 256^2 and 512^2 .
2. FFHQ (high-res): 10k images at 1024^2 .
3. Arbitrary aspect ratios (OOD): stress tests at aspect ratios $\{1 : 2, 2 : 3, 3 : 4, 9 : 16, 16 : 9, 4 : 3, 3 : 2, 2 : 1\}$ with resolutions between 512^2 and 1024^2 using FFHQ.

Main results. Table 8.2 compares **VibeToken** with 2D and 1D tokenizers across resolutions. Unlike prior 1D tokenizers (which generally assume fixed training grids), **VibeToken** natively supports arbitrary resolutions and aspect ratios while retaining short, dynamic token sequences. Concretely, **VibeToken-LL** attains **0.40** rFID at 256^2 , **0.51** at 512^2 , **2.40** at 1024^2 , and **3.60** on arbitrary-resolution stress tests; **VibeToken-SL** achieves 0.43/0.55/2.45/4.21 on the same settings. Among 1D baselines, UniTok reports a strong 0.33 rFID at 256^2 but does not support dynamic resolution; several other 1D methods lack results beyond 256^2 .

Model	Type	Params	rbitrary Resolutions	256×256			512×512		
				Tokens	FID	IS	Tokens	FID	IS
MaskGIT [28]	Mask.	177M	✗	256	4.02	355.6	1024	4.46	342
TiTok-S/B-128 [305]		287M	✗	128	2.48	–	128	2.13	–
MAGVIT-v2 [303]		307M	✗	256	1.78	319.4	1024	1.91	324.3
MaskBit [281]		305B	✗	256	1.52	328.6	–	–	–
VAR-d30 [263]	VAR	2B	✗	680	1.92	323.1	2240	2.63	303.2
VAR-d30-re [263]		2B	✗	680	1.73	350.2	–	–	–
GPT2-re [63]	AR	1.4B	✗	256	5.20	280.3	–	–	–
VIM-L-re [300]		1.7B	✗	1024	3.04	227.4	–	–	–
Open-MAGVIT2-XL [166]		1.5B	✗	256	2.33	271.77	–	–	–
LlamaGen-XXL [256]		1.4B	✗	576	2.34	253.91	–	–	–
UniTok-XXL [167]		1.4B	✗	256	2.51	216.7	–	–	–
RAR-XXL [304]		1.4B	✗	256	1.48	326.0	1024	1.66	295.7

Single R model for all Resolutions									
VibeToken-Gen-B	AR	87M		64	7.62	185.92	64	7.48	189.43
VibeToken-Gen-XXL		1.5B		64	3.62	226.93	64	3.60	230.33

Table 8.9: Comparison across models and resolutions. Columns group 256×256 and 512×512 results (Tokens/FID/IS). Among many prior works on discrete image modeling, only VibeToken-Gen scales to arbitrary resolutions effortlessly.

Against 2D tokenizers, **VibeToken** is competitive on 256²–512² (e.g., IBQ: 0.97/0.40) while trading a modest performance drop at 1024², **VibeToken** dramatically reduces token requirements at arbitrary high resolutions. Figure 8.4 shows the dynamic compression *vs.* rFID performance trade-off. With as few as **64 tokens**, **VibeToken** matches or surpasses prior dynamic-length 1D tokenizers at 256²; increasing to 128–256 tokens yields further gains while preserving the ability to decode at arbitrary target resolutions. We provide detailed ablations and native super-resolution performance in the appendix.

Efficiency analysis. Figure 8.2 plots tokenizer inference FLOPs across resolutions. 2D tokenizers scale roughly linearly with pixel count: e.g., IBQ grows from ~0.64 T to ~10.30 T FLOPs from 256² to 1024² resolution. In contrast, **VibeToken maintains constant ~1.04 T FLOPs** (maximum), enabling a fixed compute budget independent of resolution. This stabil-

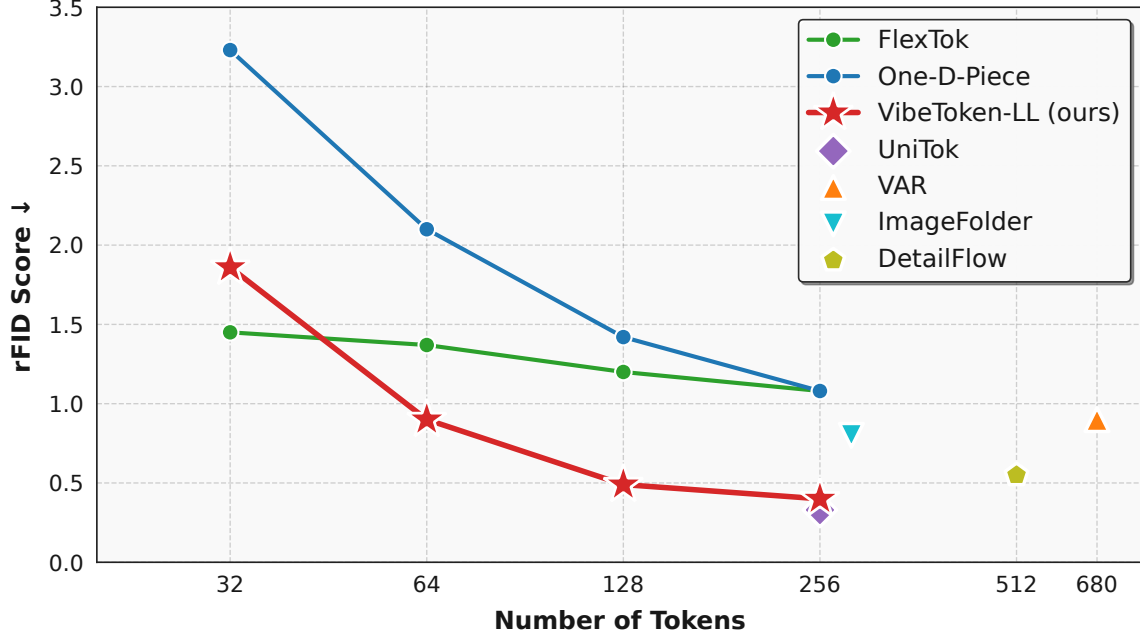


Figure 8.4: Dynamic token length. rFID vs. latent length L at 256^2 . VibeToken-Gen is the only model that pairs competitive quality with resolution generalization.

ity, combined with dynamic lengths ($L=64-256$), yields state-of-the-art throughput among flexible tokenizers. Moreover, while 2D tokenizers have a fixed $16\times/32\times$ compression, **VibeToken** achieves much higher effective compression as resolution increases: at 1024^2 , images can be encoded with **64 – 256 tokens**, improving token efficiency by $16\times$ to $64\times$ relative to 2D counterparts; resulting into effective compression ratio of $64\times$ to $128\times$.

8.5.3 Image Generation

We evaluate the resolution-agnostic tokens produced by **VibeToken** by training class-conditioned AR generators that can support arbitrary target resolutions. Because scaling prior AR models across resolutions with existing tokenizers is compute-intensive and impractical in academic settings, we compare against (i) open-source multi-resolution diffusion models and (ii) fixed-resolution AR baselines.

Experimental setup. We train **VibeToken-Gen** on mixed resolutions from 256^2 to 512^2 within an ImageNet-1k compute envelope. We use Query–Key LayerNorm for stability, select **VibeToken-LL** as the default tokenizer, and keep the rest of the LlamaGen-style stack unchanged. Because **VibeToken** supplies small, resolution-agnostic sequences, the training compute remains comparable to (or lower than) a fixed 256×256 LlamaGen baseline. Additional details are in the Appendix.

Tokens	gFID	
	w/o cfg	w/ cfg
256	39.32	9.81
128	37.68	9.02
64	33.15	8.42

Table 8.10: Ablation study on ImageNet: Impact of token length.

Tokens	Generalist	gFID	
		w/o cfg	w/ cfg
64	X	33.15	8.42
		31.14	9.37
128	X	37.68	9.02
		39.50	10.58

Table 8.11: Ablation study on ImageNet: Impact of resolution-agnostic training.

blation study. Tables 8.10, 8.11, and 8.12 reports GPT-B ablations over 100 epochs. We find that **64 tokens** are enough and yield the best gFID. Notably, increasing the token length to 128/256 does not offer any gains. This highlights that not all tokens important for improving rFID are needed to improve gFID. Training a generalist (mixed-resolution)

Model	Tokens	Generalist	gFID	
			w/o cfg	w/ cfg
LlamaGen	576	✗	33.44	7.15
VibeToken-Gen	64		31.14	9.37

Table 8.12: Comparison with LlamaGen (fixed resolution) on ImageNet.

model introduces a small degradation relative to single-resolution training, but the **64-token** setting remains competitive. Guided by these findings, we train **VibeToken-Gen-XXL** generalist models with only 64 tokens and report the main results.

Main results. We benchmark class-conditional generation against both AR and diffusion families. Since prior AR models do not natively generalize across resolutions, our primary comparisons are to multi-resolution diffusion baselines. As shown in Table 8.3, **VibeToken-Gen** smoothly scales to higher and arbitrary resolutions and achieves near-SOTA gFID despite being trained within ImageNet-1k budgets. Concretely, it outperforms EDM2 and FiT-v2 across many resolution settings and is competitive with NiT. This demonstrates that AR models can be made resolution-generalist on academic compute by delegating scale handling to a resolution-agnostic tokenizer. Finally, when evaluated strictly at 256^2 or 512^2 against models trained only at those resolutions (Table 8.9), **VibeToken-Gen** can trail specialist baselines; we attribute this to the shared budget across resolutions during generalist training, which makes direct comparison to single-resolution specialists less informative. We provide qualitative examples in the appendix.

Efficiency analysis. AR models are notoriously expensive to scale, especially when token counts grow with resolution. **VibeToken-Gen** breaks this pattern by decoupling compute from (H, W) : with a fixed latent length L , inference cost is governed by L and

Model	Type	NFEs	Scale	PSNR↑	SSIM↑	LPIPS↓
SDXL-Upscaler	Diffusion	20		27.97	0.710	0.317
Flux-Upscaler	Diffusion	20	2×	25.17	0.609	0.377
VibeToken-LL	VQ-VAE	1		24.98	0.838	0.261
SDXL-Upscaler	Diffusion	20		29.10	0.721	0.361
Flux-Upscaler	Diffusion	20	4×	25.88	0.637	0.357
VibeToken-LL	VQ-VAE	1		24.11	0.805	0.310
VibeToken-LL	VQ-VAE	1	1×	24.91	0.839	0.263

Table 8.13: Super-resolution ablation. Ground truth high resolution is fixed to 1024×1024 and bicubic downsampling is used to get low resolution counterparts.

model depth—not by image resolution. As shown in Figure 8.2, **VibeToken-Gen-XXL** requires **179 G** per forward step with $L=64$ for any resolution,³ whereas a fixed-resolution LlamaGen scales up to ~ 11 T FLOPs to reach 1024^2 even in idealized settings. In wall-clock terms, a diffusion SOTA baseline (NiT) takes 1.8 s to synthesize a single 1024^2 image at 5.87 gFID, while **VibeToken-Gen-XXL** achieves **3.54** gFID with **0.22** s latency. These results underscore that resolution-agnostic tokens enable high-resolution AR generation with substantially lower and more predictable compute.



Figure 8.5: rFID performance of VibeToken-LL with different token lengths across the resolutions. Notably, reference images are different for each resolution, hence, rFID scale and performance are not comparable between each resolution.

8.6 Ablations

8.6.1 VibeToken

Image Super-Resolution. Table 8.13 evaluates the native super-resolution capability of **VibeToken** against diffusion/flow upsamplers (e.g., SDXL/Flux upsamplers). Existing tokenizers generally lack native SR and require a separately trained upsampler; **VibeToken** does not. We use 10k FFHQ images at 1024x1024, bicubically downsample to 256x256 and 512x512 for 4x and 2x SR, respectively, and report PSNR/SSIM/LPIPS. Diffusion upsamplers tend to achieve higher PSNR (favoring smoothness), while **VibeToken** yields stronger SSIM and lower LPIPS, indicating better structure and perceptual fidelity. Thus, **VibeToken** delivers competitive SR without high-resolution pretraining or an extra upsampler. Qualitative examples are shown in Figure 8.14.

³Numbers reported per forward pass without KV caching; The total generation cost is therefore higher and scales with the sequence length L .

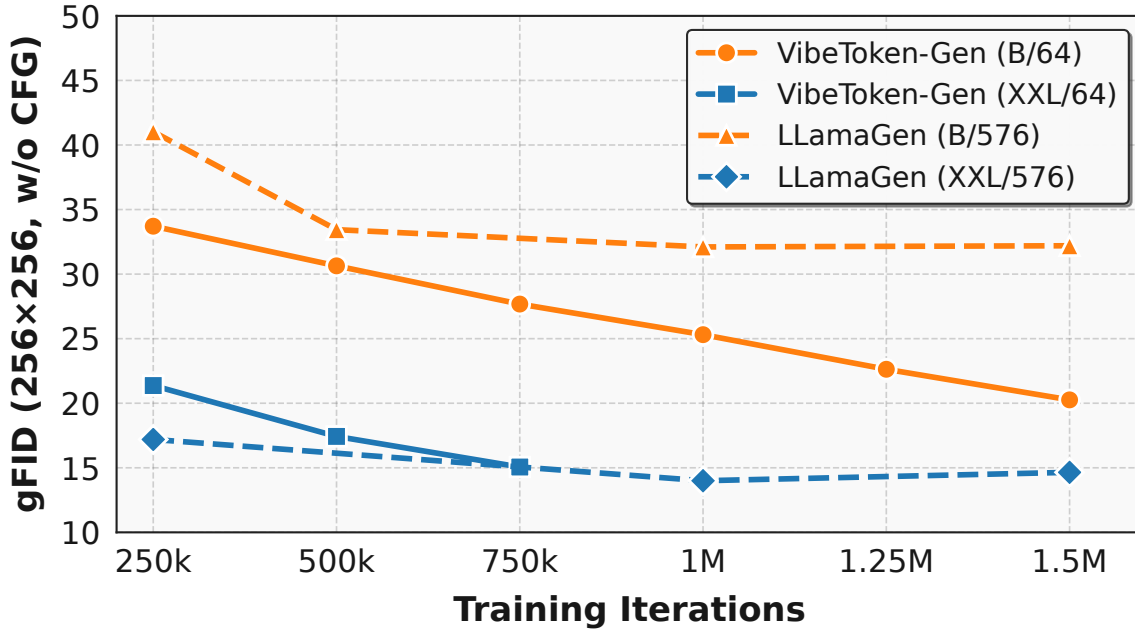


Figure 8.6: gFID performance of VibeToken-Gen and baseline LlamaGen models as training progresses. The results are shown on 256×256 without classifier-free guidance.

Token Length vs. Resolutions. Figure 8.5 shows that **VibeToken** maintains strong reconstruction quality across resolutions while supporting dynamic-length encoding. In practice, **128 tokens** offer an excellent quality–compute balance for reconstruction; **256 tokens** provide small additional gains at higher cost. Notably, for generation, **64 tokens** are most effective—later tokens primarily capture high-frequency details that benefit reconstruction more than synthesis. This separation suggests using short sequences for AR generation and longer sequences when reconstruction fidelity is paramount.

8.6.2 VibeToken-Gen

Convergence. Figure 8.6 compares training dynamics of **VibeToken-Gen** and LlamaGen at Base and XXL variants, evaluated on 256^2 without CFG to expose true likelihood learning. **VibeToken-Gen** continues improving with training and surpasses LlamaGen, whereas LlamaGen saturates early and shows limited gains thereafter. This indicates better scaling of **VibeToken-Gen**, attributable to shorter sequences and resolution-agnostic 1D tokens. Due

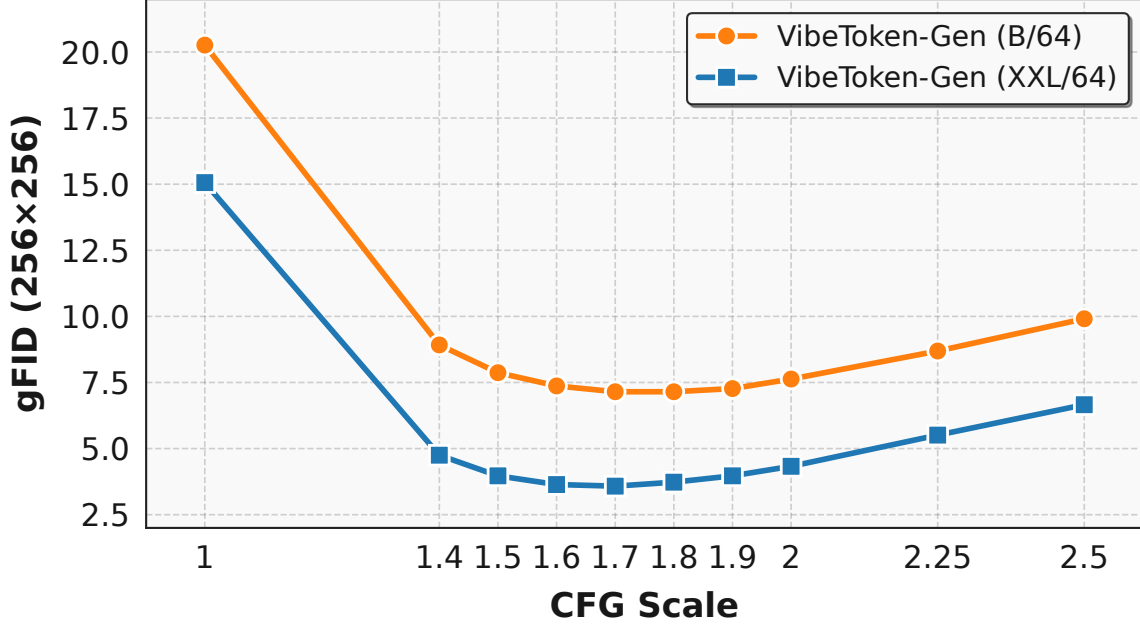


Figure 8.7: gFID performance of VibeToken-Gen (B/XXL) with respect to classifier-free guidance scale on 256×256.

to compute limits, we train **VibeToken-Gen**-XXL/64 to 750k iterations (≈ 150 epochs); based on the B/64 trend, we expect further improvements with longer training.

Impact of classifier-free guidance. Figure 8.7 shows smooth behavior across CFG scales. After convergence, both B and XXL models achieve their best gFID around 1.7–1.8. Accordingly, we use CFG= 1.75 and CFG= 2.0 when reporting **VibeToken-Gen** results.

8.6.3 Qualitative Results

Figures 8.8, 8.9, 8.10, 8.11, 8.12, and 8.13 illustrate samples from six ImageNet1k classes generated by **VibeToken-Gen**- (XXL/64) at random resolutions. **VibeToken-Gen** consistently produces high-resolution, variable-aspect-ratio images using only 64 tokens. Some crops or truncations may appear; we attribute these artifacts to the randomized cropping used during AR training.



Figure 8.8: Class 985. Qualitative examples of generated images using VibeToken-Gen (XXL/64) on arbitrary resolutions.



Figure 8.9: Class 980. Qualitative examples of generated images using VibeToken-Gen (XXL/64) on arbitrary resolutions.



Figure 8.10: Class 387. Qualitative examples of generated images using VibeToken-Gen (XXL/64) on arbitrary resolutions.



Figure 8.11: Class 250. Qualitative examples of generated images using VibeToken-Gen (XXL/64) on arbitrary resolutions.



Figure 8.12: Class 88. Qualitative examples of generated images using VibeToken-Gen (XXL/64) on arbitrary resolutions.



Figure 8.13: Class 33. Qualitative examples of generated images using VibeToken-Gen (XXL/64) on arbitrary resolutions.

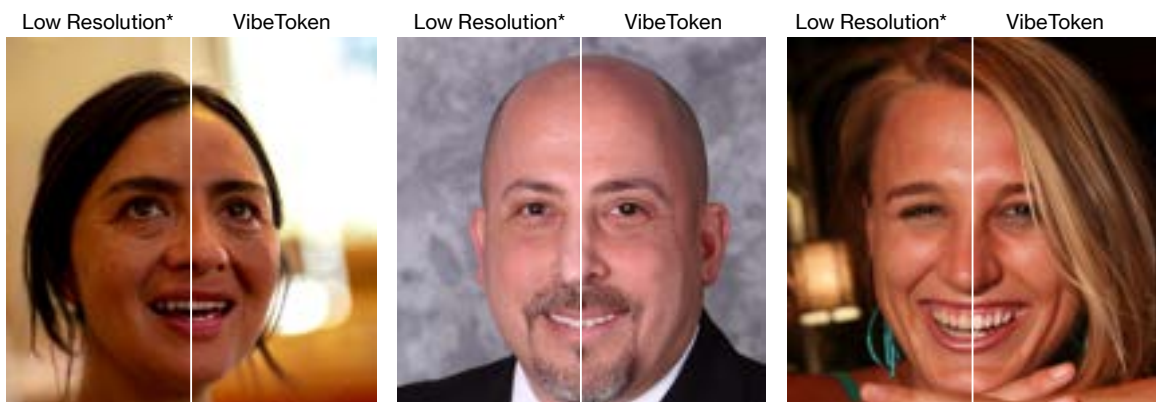


Figure 8.14: Qualitative examples of native 4 $\times$ super-resolution ability of VibeToken with respect to simple bilinear resize operation.

8.7 Conclusion

We present **VibeToken** and **VibeToken-Gen**, a practical recipe for scaling autoregressive (AR) image generation to arbitrary resolutions and aspect ratios under a fixed, user-controllable compute budget. The core idea is a resolution-agnostic 1D tokenizer (**VibeToken**) that produces short, variable-length sequences (64-256 tokens) independent of the resolutions. This is achieved through dynamic grid positional embeddings, adaptive decoding, and length-aware training. To our knowledge, **VibeToken** is the first 1D Transformer-based tokenizer to natively support arbitrary resolutions while retaining strong compression and reconstruction quality. Building on these tokens, **VibeToken-Gen** achieves competitive class-conditional generation across resolutions, with inference FLOPs governed by token length rather than pixel count, thereby narrowing the gap between AR and diffusion pipelines within academic compute budget.

Limitations and Future Works. Our experiments focus on ImageNet-1k and class-conditional generation; extending to text-to-image, open-vocabulary settings, and video is an important next step. As a resolution generalist model, **VibeToken-Gen** can trail single-resolution specialists at exactly $256^2/512^2$ under the same budget. Further scaling, longer training, and stronger data augmentation may reduce this gap. Methodologically, exploring other AR variants (e.g. randomized orderings, scale-wise training), alternative quantization schemes, and larger pretraining corpora could improve quality. Finally, applying resolution-agnostic tokenization to unified multimodal modeling (e.g. image-text-video) is a promising direction for significantly efficient production-grade generative systems.

UNDERSTANDING AND IMPROVING FEW-STEP GENERATIONS VIA FLOW MAPS: A GEOMETRIC PERSPECTIVE

9.1 Introduction

Recent advances in diffusion and flow matching models have enabled high-fidelity visual generation at scale [216, 61, 207, 271]. However, state-of-the-art models typically require ≥ 30 number of function evaluations (NFEs) at inference time, which translates into substantial computational cost. A large body of work on distillation (e.g., single-step or few-step methods such as SiD [327], CMs [254], and DMD [298]) aims to compress these models into 1–4 NFE. While effective, these approaches introduce an additional teacher–student pipeline and non-trivial engineering effort on top of training a strong teacher.

An appealing alternative is to train few-step flow-based models from scratch, without a separate distillation stage [79, 324, 69]. Flow-map approaches, such as MeanFlow [79], move in this direction by directly learning mappings between arbitrary time pairs (t, r) with $r \leq t$. Concretely, these models parameterize the average (or “mean”) of the instantaneous velocity field along the ODE trajectory between t and r and use it to realize large time jumps in a single step. In practice, however, such Flow Map training introduces its own challenges: (i) training can be unstable, especially when the underlying flow is highly curved or when classifier-free guidance (CFG) [97] is used, and (ii) MeanFlow tends to strongly improve 1–NFE performance while degrading performance at ≥ 2 NFEs, leading to an unfavorable trade-off between speed and quality. Recent work has mitigated some of these issues by relying on pretrained flow-matching teachers, but a precise understanding of why MeanFlow is unstable and how to systematically stabilize it has remained elusive.

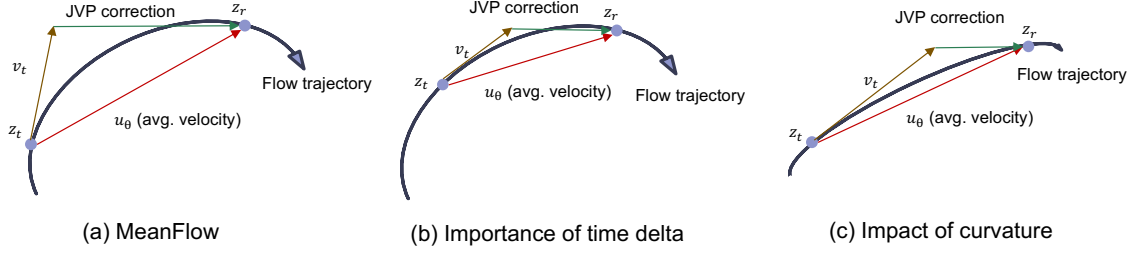


Figure 9.1: Geometric Perspective Intuition. This figure first illustrates MeanFlow identity (Eq. 9.4) over the oracle flow trajectory, where the time derivative is represented as JVP correction (a). Overall velocity field of flow matching models have high time derivative, which is a problem. There are two ways to reduce the impact of time derivative: (b) reducing the time delta and (c) reducing the curvature. By doing either of them instantaneous velocity (v_t) becomes very close to the average velocity (u_θ) and time derivative becomes smaller.

In this work, we take a geometric perspective on Flow Maps and MeanFlow. Starting from the probability flow ODE, we derive a curvature measure κ for the learned velocity field and show that the MeanFlow identity can be expressed in terms of (i) the time difference $(t - r)$, which scales the Jacobian–vector product (JVP) term, and (ii) the curvature κ , which controls the normal component of this JVP. Our analysis further shows that high curvature of the velocity field makes the MeanFlow JVP term larger and more variable, and hence training more fragile, particularly when attempting to match multi-step teacher behavior.

Building on these insights, we propose two complementary strategies to stabilize MeanFlow training:

(1) Time-localized MeanFlow via Piecewise sampling. We introduce Piecewise MeanFlow, which partitions the time interval $[0, 1]$ into N equal segments and constrains each pair (t, r) to lie within a single segment. Under this sampling scheme we obtain a hard bound $0 \leq t - r \leq 1/N$, which in turn provably reduces the mean and variance of the JVP term in the MeanFlow identity. We further extend this idea to a Unified Multi-scale Piecewise MeanFlow, which samples (t, r) across multiple time segments. This multi-scale objective simply combines fine-scale, highly local supervision (small $(t - r)$) with coarser-scale constraints that preserve global transport, and substantially improves the stability of

MeanFlow training at target NFEs ≥ 2 . Hence, given a target inference compute budget Piecewise-MeanFlow is always a better alternative to MeanFlow.

(2) Curvature-regularized velocity fields. Our geometric analysis reveals that the MeanFlow identity can be represented as a function of curvature. Hence, first we show that AlphaFlow’s improved shortcut curriculum improves the curvature of the base velocity field and we further improve AlphaFlow curriculum based on the curvature metric (κ). Using such a regularized base model as a teacher yields straighter flows with lower curvature and significantly stabilizes subsequent MeanFlow training.

Finally, we combine these two strategies – time-localized MeanFlow and curvature-regularized teachers – into a unified framework that we refer to as Stable MeanFlow (sMF). Empirically, sMF consistently improves over baseline MeanFlow under a modest fine-tuning budget, achieving a more favorable trade-off between FID and NFE. Our controlled experiments and ablations support the theoretical prediction that curvature and the time horizon ($t-r$) are key factors governing the stability and effectiveness of Flow Map learning. In summary, our contributions are:

- We develop a geometric framework for Flow Map learning (e.g., MeanFlow), making explicit the role of curvature and the JVP term in the MeanFlow identity.
- We propose Piecewise MeanFlow and its unified multi-scale variant, which bounds the effective time horizon ($t-r$) and provably reduces the variance of the JVP term, leading to more stable training and improved few-step performance.
- We explore curvature-regularized velocity smoothness objective and show that leveraging such straighter base flows further stabilizes MeanFlow and improves the overall FID-vs.-NFE Pareto frontier.
- Our approach (sMF) significantly reduces the training and resource requirements in a pure distillation setting and improves the performance.

9.2 Related Works

Diffusion and Flow Matching. Diffusion models [247, 96, 251, 255] and flow matching [147, 148, 123] constitute the dominant paradigms for large-scale generative modeling. Diffusion models, including DDPM [96], score-based models [255], and probability flow ODE formulations [251], learn to reverse a forward noising process by estimating either the score function or the noise. Flow matching models, such as Rectified Flow [153], Stochastic Interpolants [4], and Optimal Transport Flow Matching [147, 148, 123], instead learn a deterministic velocity field that transports samples from a prior distribution to the data distribution. These approaches have achieved strong performance in image [216, 61] and video generation [207, 18, 271, 195], editing [178, 92, 23, 126], and personalization [193, 223, 72, 137]. However, state-of-the-art diffusion and flow matching models typically require tens of inference steps (NFEs) to achieve high sample quality, leading to substantial computational overhead at generation time.

Few-Step Generative Models. Improving sampling efficiency has led to extensive work on few-step and one-step generative models. Distribution matching methods distill a pre-trained multi-step diffusion model into a few-step student by matching output distributions, often using adversarial or divergence-based objectives. Representative examples include Distribution Matching Distillation (DMD) [298], DMDv2 [297], f-distill [291], and other adversarial distillation frameworks [231]. While effective, these methods typically require a separate teacher–student pipeline and non-trivial training complexity. Score distillation methods transfer knowledge by matching score or noise predictions under large integration steps. Methods such as Score Identity Distillation (SiD) [327], SiDA [326], and related variants [165, 174] fall into this category. Although originally developed for diffusion models, score distillation can also be applied to flow matching models by converting deterministic flows into stochastic counterparts [325]. These approaches can achieve strong

one-step performance but are often sensitive to guidance scale, teacher quality, and requires training several auxiliary models together. Trajectory-based methods aim to directly model large-time transitions along the generative trajectory. Instead of matching instantaneous scores, these methods learn mappings between different timesteps, enabling few-step or one-step inference without iterative sampling. This category includes Consistency Models (CMs) [254], simplified Consistency Models (sCT) [272], Enhanced Consistency Training (ECT) [80], Consistency Trajectory Models (CTM) [117], and Shortcut Models [69].

Flow Maps and MeanFlow. Flow maps [22] extend trajectory-based modeling to flow matching by explicitly learning mappings between arbitrary time pairs. Early approaches such as CMs [254] and CTM [117] were primarily developed in the diffusion setting, while subsequent works adapted these ideas to flow matching due to its deterministic formulation and ease of scaling. MeanFlow [79] and Align-Your-Flow (AYF) [225] introduce a principled formulation based on learning time-averaged velocities between two timesteps, enabling few-step generation directly from scratch as well as in distillation settings. Despite its strong empirical performance at very low NFEs, MeanFlow is known to be difficult to train and often exhibits instability, particularly when applied to higher inference budgets. And AYF requires other engineering tricks to make it work. Several recent works have proposed improvements to MeanFlow. FACM [199] improves MeanFlow in distillation settings by anchoring consistency to flow matching, while AlphaFlow [312] unifies trajectory flow matching, shortcut modeling, and MeanFlow through a curriculum-based objective, enabling more stable training from scratch.

However, existing works primarily address MeanFlow instability through optimization heuristics, curricula, or architectural design. A fundamental understanding of why MeanFlow is unstable, even when distilled from strong teachers, and how to systematically stabilize it remains missing.

9.3 Preliminaries

Flow Matching. Flow Matching learns the velocity field from prior distribution to the posterior (data) distribution and performs the sampling by solving the ODE:

$$\frac{d}{dt}z_t = v_\theta(z_t), \quad (9.1)$$

where $z_t = (1-t) \cdot x + t \cdot \epsilon$ with data $x \sim p_{data}$ and noise $\epsilon \sim p_{prior}$ is the linear interpolation as forward noising process and v_θ is the model. Computing the time derivative of the forward process gives the conditional velocity $v = \epsilon - x$. Hence, our training objective becomes:

$$\mathcal{L}_{FM} = \mathbb{E}_{t,x,\epsilon} \|v_\theta(z_t, t) - (\epsilon - x)\|^2. \quad (9.2)$$

MeanFlow. Building on the flow matching, MeanFlow proposes to learn the Flow Map by approximating the velocity integral by learning the average velocity between two time steps t and r (such that $r \leq t$):

$$u(z_t, r, t) = \frac{1}{t-r} \int_r^t v(z_\tau) d\tau. \quad (9.3)$$

To improve the training feasibility, by taking differentiation on both sides we obtain the MeanFlow identity:

$$u(z_t, r, t) = v(z_t) - (t-r) \frac{d}{dt}u(z_t, r, t), \quad (9.4)$$

where $\frac{d}{dt}u(z_t, r, t)$ is the Jacobian-vector product (JVP). Thus the JVP term controls the discrepancy between the instantaneous velocity and average velocity. However, this assumes that we have velocity field already learned but this is not the case while pretraining from scratch. Hence, introducing the 75-25% split between Flow Matching and MeanFlow, respectively.

9.4 Geometric Interpretation: Curvature and the MeanFlow JVP

In this section, we attempt to interpret the velocity field and the MeanFlow JVP term. Specifically, during inference-time ODE sampling, we approximate the continuous velocity field via equal-distance discretized steps, as the learned velocity field is curved in space and time. Therefore, we first derive a way to measure the curvature for flow matching models. Building on this novel curvature metric (κ), we find that the JVP term of MeanFlow can be expressed as a function of curvature.

9.4.1 Curvature of the probability flow ODE

Let's assume that velocity $v(z_t, t)$ is C^1 in both arguments and define the instantaneous velocity ($v_t = v(z_t, t)$) and acceleration $a_t = \frac{d}{dt}v(z_t, t)$. By the chain rule,

$$a_t = \frac{d}{dt}v(z_t, t) = \nabla_z v(z_t, t)v_t + \partial_t v(z_t, t). \quad (9.5)$$

Assuming $v_t \neq 0$, we can define the unit tangent component $T_t = \frac{v_t}{\|v_t\|}$. Acceleration can be further decomposed into tangential and normal components as, $a_t = a_{T,t} + a_{N,t}$, respectively. Here, $a_{T,t} = (a_t \cdot T_t)T_t$ and $a_{N,t} = a_t - a_{T,t}$. Given this, we can define the curvature of the ODE trajectory.

Definition 9.4.1 (Curvature of the ODE trajectory). The curvature of the trajectory at time t is

$$\kappa(z_t, t) = \frac{\|a_{N,t}\|}{\|v_t\|^2}. \quad (9.6)$$

This is the standard curvature formula obtained by reparameterizing the arc length. It measures how sharply the trajectory bends, independently of its speed. Now, the projection onto the subspace orthogonal to v is $P_\perp = I - TT^T$. Then the curvature of the integral curve at (x, t) is:

$$\kappa(z_t, t) = \frac{\|P_\perp a\|}{\|v\|^2} \quad (9.7)$$

$$= \frac{\|(I - \hat{v}\hat{v}^T)(\nabla_z v_\theta \cdot v_\theta(z_t, t) + \partial_t v_\theta(z_t, t))\|}{\|v_\theta(z_t, t)\|^2}, \quad (9.8)$$

where $\hat{v} = v_\theta(z_t, t)/\|v_\theta\|$.

9.4.2 Local linearization: JVP as a curvature proxy

We now make the geometric link precise via a simple local linear approximation in time via Taylor series expansion. Assume that near time t , the instantaneous velocity varies approximately affinely in t along the trajectory:

$$v(z_s, s) \approx v_t + (s - t)a_t, \quad s \in [r, t], \quad (9.9)$$

where $a_t = \frac{d}{dt}v(z_t, t)$ and we assume that any higher order terms zero. Hence, we can rewrite the MeanFlow identity in terms of acceleration as follows.

Lemma 9.4.1 (Average velocity in the affine-in-time regime). *under the affine approximation*

$v(z_s, s) = v_t + (s - t)a_t$ on $s \in [r, t]$,

$$u(z_t, r, t) = v_t - \frac{(t - r)}{2}a_t. \quad (9.10)$$

Proof. We can compute the integral as:

$$\begin{aligned} \int_r^t v(z_\tau, \tau) d\tau &= \int_r^t (v_t + (\tau - t)a_t) d\tau \\ &= (t - r)v_t + a_t \int_r^t (\tau - t) d\tau \\ &= (t - r)v_t - \frac{(t - r)^2}{2}a_t. \end{aligned}$$

Diving by $(t - r)$ yields

$$u(z_t, r, t) = v_t - \frac{(t - r)}{2}a_t.$$

□

Corollary 9.4.2 (Discrepancy and JVP in the affine regime). *Under the same assumptions,*

$$\delta v_t = v_t - u(z_t, r, t) = \frac{(t-r)}{2} a_t, \quad (9.11)$$

and

$$\frac{d}{dt} u(z_t, r, t) = \frac{1}{2} a_t. \quad (9.12)$$

Proof. From Lemma 9.4.1,

$$\begin{aligned} \delta v_t &= v_t - \left(v_t - \frac{t-r}{2} a_t \right) \\ &= \frac{t-r}{2} a_t. \end{aligned}$$

To differentiate u in this simple affine model, assuming that a_t is constant and v_t satisfies $\frac{d}{dt} v_t = a_t$. Thus,

$$\begin{aligned} \frac{d}{dt} u(z_t, r, t) &= \frac{d}{dt} \left(v_t - \frac{t-r}{2} a_t \right) \\ &= a_t - \frac{1}{2} a_t \\ &= \frac{1}{2} a_t \end{aligned}$$

□

Till this point, we have shown that MeanFlow identity can be expressed in terms of acceleration. This assumes that acceleration (a_t) is constant for all t and relaxing this assumption will lead to complex formulation. However, the MeanFlow identity will remain a function of acceleration. Building on this, we can further expand the MeanFlow identity and express it as the function of curvature.

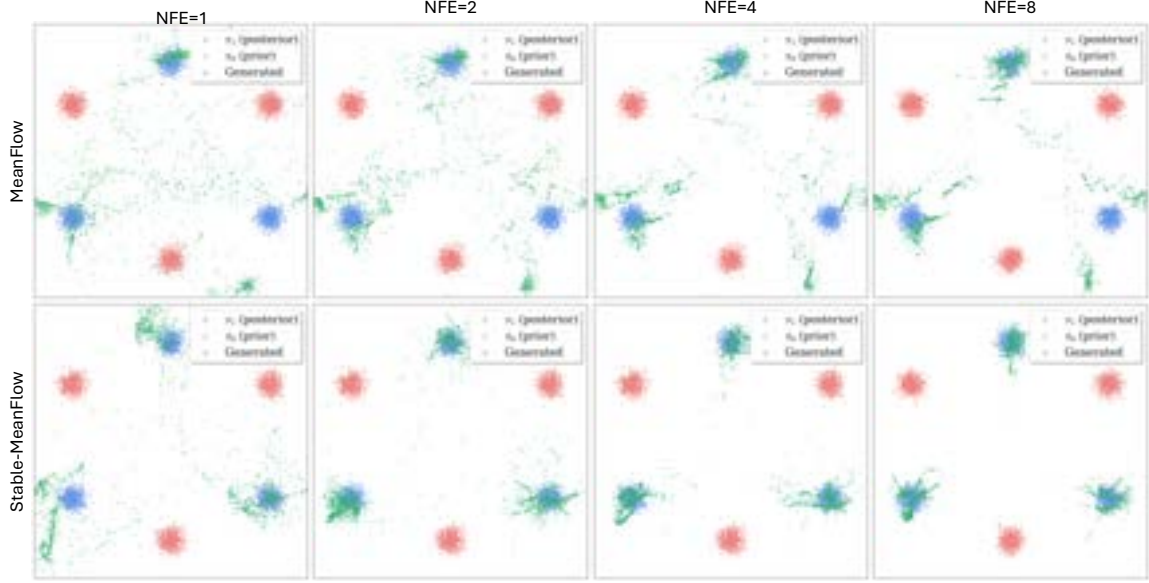


Figure 9.2: Multimodal toy distribution. Qualitative results on 2D-dataset for few-step generations. MeanFlow performs very poorly and slightly improves as inference steps increases. Stable MeanFlow significantly improves over the baseline and keeps improving.

Proposition 9.4.3 (Normal JVP term as a curvature proxy). *In the affine-in-time regime,*

$$\delta v_{N,t} = \frac{t-r}{2} a_{N,t}. \quad (9.13)$$

Hence,

$$\|\delta v_{N,t}\| = \frac{t-r}{2} \|a_{N,t}\| \quad (9.14)$$

$$= \frac{t-r}{2} \|v_t\|^2 \kappa(z-t, t). \quad (9.15)$$

Proof. By Corollary 9.4.2,

$$\delta v_{N,t} = \frac{t-r}{2} a_{N,t}.$$

We know that $\kappa(z_t, t) = \|a_{N,t}\|/\|v_t\|^2$ (from Eq. 9.4.1). Now, multiplying and dividing the above equation with $\|v_t\|^2$, we get:

$$\delta v_{N,t} = \frac{t-r}{2} \|v_t\|^2 \cdot \kappa(z_t, t). \quad (9.16)$$

□

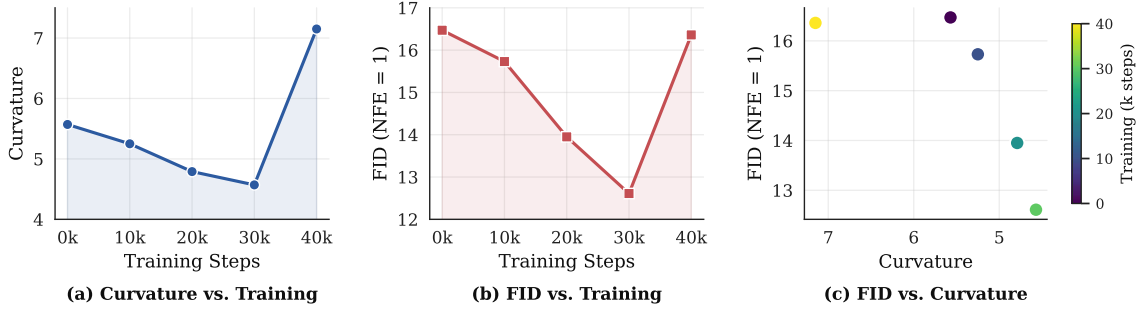


Figure 9.3: Impact of curvature in Flow Maps learning. We illustrate the impact of the curvature regularizer on distilling Flow Maps. Specifically, we perform up to 40k optimization steps for the curvature regularizer followed by 10k steps of Flow Map distillation on top of a pretrained REPA model.

Interpretation. For small time spans ($t - r$), the normal part of MeanFlow JVP term is first-order proportional to curvature:

- If the trajectory is highly curved at (z_t, t) , then $\kappa(z_t, t)$ is large and the normal component of JVP must respond accordingly.
- If the trajectory is locally straight ($\kappa = 0$), the normal component of this JVP term vanishes.
- If the maximum time delta ($t - r$) is sufficiently small then it will further reduce the impact of high curvatures in the velocity field.

9.5 Method: Stabilizing Flow Maps with Time-Span and Curvature Control

Section 9.4 identifies two straightforward but grounded methods for stabilizing MeanFlow/Flow Map training in the few-step regime. From Eq. (9.16), the problematic (normal) component of the MeanFlow discrepancy scales with (i) the time span ($t - r$) and (ii) the curvature of the learned velocity field. We therefore propose two complementary improvements:

- **Piecewise MeanFlow:** explicitly upper-bound $(t - r)$ during training, reducing JVP

variance.

- **Curvature-Regularized MeanFlow:** encourage a straighter velocity field while preserving marginals.

9.5.1 Piecewise MeanFlow

MeanFlow relies on the JVP term $\frac{d}{dt}u_\theta(z_t, r, t)$, whose magnitude (and variance) grows with the temporal span $(t - r)$. When supervision pairs cover large spans, high-curvature regions dominate gradient noise, leading to instability.

Time segmentation and sampling. Define a uniform partition of $[0, 1]$:

$$0 = \tau_0 < \tau_1 < \dots < \tau_N, \quad \tau_k = \frac{k}{N},$$

where N is the number of segments. Piecewise MeanFlow samples training times (t, r) within a single segment: (i) sample t (logit-normal, as in MeanFlow), (ii) find its segment index $k = \max\{j : \tau_{j-1} \leq t < \tau_j\}$, (iii) sample $r \sim \text{Uniform}([\tau_{k-1}, t])$. This ensures $r \leq t$ and $t - r \leq 1/N$. The loss is computed exactly as in MeanFlow, but using these restricted time pairs.

Setting $N = 1$ recovers MeanFlow (up to the choice of r -sampling), while $N \rightarrow \infty$ approaches Flow Matching. Thus, N trades off global consistency (large spans) with local stability (small spans).

Proposition 9.5.1 (JVP scaling under Piecewise MeanFlow). *Fix $N \geq 1$ and let (t, r) be samples by above scheme then:*

- 1) *The time gap is uniformly bounded: $0 \leq t - r \leq 1/N$ almost surely.*
- 2) *For any parameter θ such that $\mathbb{E}[\|J_\theta(z_t, r, t)\|_2^2] < \infty$, we have*

$$\mathbb{E}[\|(t - r)J_\theta(z_t, r, t)\|_2^2] \leq \frac{1}{N^2} \mathbb{E}[\|J_\theta(z_t, r, t)\|_2^2]. \quad (9.17)$$

Proof. (1) follows directly from the segment constraint. For (2), use $(t - r)^2 \leq 1/N^2$ pointwise and take expectations. $\square$

Unified multi-scale Piecewise MeanFlow. A fixed N provides strong control but fixes the required target NFE to be at least N steps. To support multiple inference-time NFE budgets with a single model, we can uniformly mix the time segmentations, which will still regularize better. Let $\mathcal{M}$ be a set of segment counts (e.g., $\{1, 2, 4, 8, 16\}$). For each $m \in \mathcal{M}$ define $\tau_k^{(m)} = k/m$. During training we sample $m \sim \text{Uniform}(\mathcal{M})$ and then sample (t, r) within one segment of that partition. The resulting objective is

$$\mathcal{L}_{u\text{PMF}} = \mathbb{E}_{\substack{m \sim \mathcal{M}, z_t, \epsilon \\ (t, r) \in [\tau_{k-1}^{(m)}, \tau_k^{(m)}]}} \|u_\theta(z_t, r, t) - \text{sg}(u_{gt})\|. \quad (9.18)$$

9.5.2 Curvature-Regularized MeanFlow

Piecewise MeanFlow controls the explicit span factor $(t - r)$. The second lever is to reduce curvature, so the field is easier to approximate with few steps. In practice, curvature can be reduced through several mechanisms (OT couplings, rectification/self-distillation, shortcut modeling).

AlphaFlow as an implicit curvature regularizer. AlphaFlow [312] combines shortcut modeling and MeanFlow by interpolating between shortcut velocity v_s and an EMA MeanFlow teacher u_{θ^-} :

$$\mathcal{L}_{\text{AF}} = \mathbb{E} \left[\|u_\theta(z_t, r, t) - \text{sg}(u_{tgt}^{\text{AlphaFlow}})\|_2^2 \right], \quad (9.19)$$

where $u_{tgt}^{\text{AlphaFlow}} = (\gamma v_s(x_s, s) + (1 - \gamma) u_{\theta^-}(x_s, r, s))$, and $\gamma \in (0, 1]$ follows a sigmoid curriculum from 1.0 to 0.0. Through the curvature lens, intermediate values of γ encourage the model to follow the straighter shortcut field, reducing curvature and improving MeanFlow stability.

Curvature-guided curriculum cutoff. Empirically, letting $\beta_1 = 0$ in AlphaFlow eventually enters a one-step dominated regime where curvature increases but improves the 1 step performance at the cost of inference-time scaling (see Figure 9.3). We therefore clip the schedule to remain in the low-curvature region:

$$\tilde{\beta} = \text{clamp}(\sigma(\cdot), \beta_1, \beta_2), \quad \beta_1 = 0.2, \beta_2 = 1.0. \quad (9.20)$$

voiding flow-matching supervision during distillation. During distillation, the MeanFlow identity (Eq. 9.4) reduces to standard flow-matching supervision when $r = t$, since the target collapses to the instantaneous velocity v_t . This case corresponds to the highest-curvature supervision and directly contradicts our goal of stabilizing MeanFlow training. To avoid this degeneracy, we enforce a fixed minimum time gap

$$t - r = \gamma, \quad \gamma > 0, \quad (9.21)$$

for the oracle (teacher) branch. We use $\gamma = 0.01$ in all experiments.

With this constraint, the MeanFlow identity becomes

$$u_\theta(z_t, r, t) = u_{\psi^-}(z_t, t - \gamma, t) - (t - r) \frac{d}{dt} u_\theta(z_t, r, t), \quad (9.22)$$

where ψ^- denotes a frozen teacher obtained via curvature-smoothing fine-tuning (Eq. 9.19).

Combined with Piecewise MeanFlow, this modification eliminates high-curvature flow-matching supervision while preserving meaningful multi-time constraints. We refer to the resulting curvature-aligned training procedure as Stable MeanFlow (sMF).

9.6 Experimental Results

This section evaluates the two stability methods identified in Section 9.4 and instantiated in Section 9.5: (i) controlling the time span $(t - r)$ via Piecewise MeanFlow, and (ii) lowering curvature to make few-step distillation more effective.

9.6.1 Controlled Multimodal Transport

Standard evaluations of flow-matching models often rely on simple toy datasets such as Gaussian-to-two-moons or checkerboard transformations, which are relatively easy and do not strongly stress few-step integration. To better expose failure modes specific to few-step sampling, we design a more challenging synthetic transport task, illustrated in Figure 9.2. Specifically, both the prior distribution π_0 and the target distribution π_1 consist of three Gaussian modes arranged alternately on a circle. This construction induces severe multimodality and requires the learned flow to perform coordinated global transport while avoiding mode collapse. As a result, the task serves as a controlled testbed for visualizing how few-step samplers handle curved and highly multimodal probability flows.

Figure 9.2 shows qualitative results for MeanFlow and Stable MeanFlow under 1, 2, and 4 function evaluations (NFEs). To isolate training stability, we train both models using only pairs with $t \neq r$, eliminating the flow-matching supervision. MeanFlow fails catastrophically across all NFEs, producing poor samples that do not improve with additional inference steps. In contrast, Stable MeanFlow exhibits consistent improvement as the number of NFEs increases: while one-step generation remains imperfect, higher-step sampling progressively recovers the correct multimodal structure. These results demonstrate that the instability of MeanFlow is not merely an optimization artifact but arises from geometric properties of the underlying flow. By explicitly controlling curvature and effective time span, Stable MeanFlow yields trajectories that are easier to approximate with few integration steps, validating our geometric perspective.

9.6.2 Impact of Curvature

Setup. To isolate the effect of curvature on MeanFlow performance, we start from a pre-trained REPA [306] DiT-XL/2 [197] model trained on ImageNet-1K at 256×256 resolution.

We then perform small-scale AlphaFlow fine-tuning (Eq. 9.19) for 40k optimization steps using a batch size of 16. We save intermediate checkpoints every 10k steps. For each checkpoint, we further fine-tune a MeanFlow model for 10k steps (again with batch size 16), keeping all other settings fixed. For each base checkpoint, we measure the curvature κ of the learned velocity field and report the corresponding one-step (NFE=1) FID of MeanFlow.

Findings. Figure 9.3 summarizes the relationship between AlphaFlow training, curvature, and MeanFlow performance. As AlphaFlow fine-tuning progresses, the curvature of the underlying vector field steadily decreases, indicating progressively straighter transport trajectories. Towards the end of training, curvature increases again, coinciding with AlphaFlow over-optimizing for one-step generation at the expense of marginal fidelity. Crucially, MeanFlow performance closely mirrors this curvature trend. MeanFlow fine-tuned on lower-curvature base models consistently achieves better one-step FID, while performance degrades once curvature increases again. This strong correlation between curvature and MeanFlow effectiveness provides direct empirical evidence that curvature, rather than optimization noise alone, is a key factor governing the stability and performance of few-step Flow Map learning.

9.6.3 Image Generation

Piecewise MeanFlow: controlling $(t - r)$

Setup. We evaluate Piecewise MeanFlow on CIFAR10 (32×32) [125] and ImageNet (256×256) [49], reporting FID over 50k generated samples. We follow the baseline MeanFlow setup and architectures (SongUNet for CIFAR10; DiT-B/4 for ImageNet) and vary the inference-time number of function evaluations (NFEs) from 1 to 8. Importantly, we train CIFAR10 and DiT-B/4 models from scratch for 400,000 optimization steps with batch sizes of 512 and 256, respectively. We train Piecewise MeanFlow with fixed seg-

Table 9.1: CIF R10 Piecewise MeanFlow. We compare FID vs. NFE performance of MeanFlow variants. $\dagger$ represents our reproduced results using official released codebase. $\ddagger$ represents that each model is trained on NFE specific segments independently.

Model	FID			
	NFE=1	NFE=2	NFE=4	NFE=8
MeanFlow $\dagger$	3.52	3.7	7.95	5.25
Piecewise-MeanFlow $\ddagger$	<u>3.84</u>	3.03	<u>2.96</u>	<u>2.86</u>
Unified Multi-Scale MeanFlow	5.68	<u>3.22</u>	2.81	2.75

Table 9.2: ImageNet1k 256×256 Piecewise MeanFlow. We compare FID vs. NFE performance of MeanFlow variants. $\dagger$ represents our reproduced results. $\ddagger$ represents that each model is trained on NFE specific segments independently.

Model	rch.	FID			
		NFE=1	NFE=2	NFE=4	NFE=8
MeanFlow $\dagger$		22.39	<u>15.06</u>	13.66	13.40
Piecewise-MeanFlow $\ddagger$	DiT-B/4	22.39	14.77	12.42	11.78
Unified Multi-Scale MeanFlow		28.96	15.20	<u>13.26</u>	<u>12.84</u>

ment counts $N \in \{1, 2, 4, 8\}$. We also train Unified multi-scale Piecewise MeanFlow with $\mathcal{M} = \{1, 2, 4, 8\}$. Note that $N = 1$ corresponds to the MeanFlow objective.

Results. Tables 9.1 and 9.2 show a consistent trend: vanilla MeanFlow is competitive at very low NFEs (1–2) but can degrade as NFEs increase, reflecting the instability predicted by the $(t-r)$ -scaled JVP term. Piecewise MeanFlow improves performance when the target NFE budget is known (e.g., $N = 4$ tends to be strongest around 4 steps), consistent with the idea that bounding $(t-r)$ reduces unstable supervision. Unified multi-scale training recovers much of this gain while keeping a single model usable across a range of inference-time

Table 9.3: DiT-XL/2 finetuning-based comparison under different NFE settings. Lower FID is better ($\downarrow$) and higher IS is better ($\uparrow$).

Method	NFE=1		NFE=4	
	FID ($\downarrow$)	IS ($\uparrow$)	FID ($\downarrow$)	IS ($\uparrow$)
Flow Matching	212.59	2.84	22.52	125.75
AlphaFlow (base)	21.45	86.00	5.11	207.31
MeanFlow	16.47	110.78	10.85	135.64
AlphaFlow + MeanFlow	<u>12.82</u>	<u>148.81</u>	9.03	187.58
Stable MeanFlow (ours)	8.98	164.90	<u>6.05</u>	<u>193.05</u>

budgets. This study consistently shows that Piecewise MeanFlow adds significant stability.

Stable MeanFlow Fine-tuning

Setup. To evaluate the effectiveness of Stable MeanFlow, we conduct controlled comparisons starting from a common base model. We begin with a REPA DiT-XL/2 flow-matching model pretrained on ImageNet-1K. We first fine-tune this model using AlphaFlow (Eq. 9.19) for 50k optimization steps with batch size 256 (approximately 10 epochs). On top of the REPA and AlphaFlow checkpoints, we then fine-tune the following variants for 10k optimization steps with batch size 16: MeanFlow, AlphaFlow + MeanFlow, and Stable MeanFlow. All methods are evaluated using 50k-sample FID and Inception Score (IS).

Main Results. Table 9.3 reports FID and IS across different inference budgets. AlphaFlow alone substantially improves few-step performance compared to the REPA baseline, achiev-

ing strong 4-step results (5.11 FID). When AlphaFlow checkpoints are used as a base for MeanFlow fine-tuning, we observe mixed behavior. AlphaFlow + MeanFlow improves over vanilla MeanFlow in both FID and IS, but its 4-step performance degrades sharply, worsening from 5.11 FID to 9.03 FID. This highlights the intrinsic instability of MeanFlow when applied on top of curved base flows.

In contrast, Stable MeanFlow consistently improves performance across inference budgets. By jointly leveraging curvature-regularized teachers and Piecewise MeanFlow, Stable MeanFlow achieves 8.98 FID at 1-step and 6.05 FID at 4-step inference. Notably, it improves one-step generation while largely preserving the strong multi-step performance of the base AlphaFlow model, effectively closing the gap between the base model and the MeanFlow student at 4-NFE.

9.7 Qualitative Results.

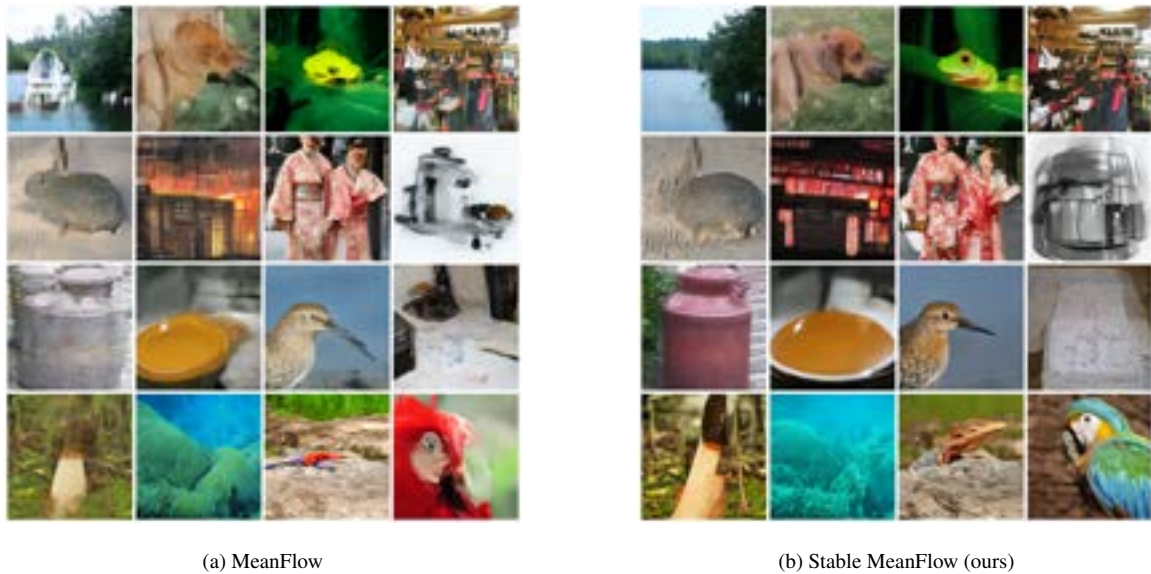


Figure 9.4: Qualitative Examples at NFE=1. Qualitative comparison between randomly selected images from MeanFlow and Stable MeanFlow. Stable MeanFlow significantly improves over the baseline.



(a) MeanFlow

(b) Stable MeanFlow (ours)

Figure 9.5: Qualitative Examples at NFE=4. Qualitative comparison between randomly selected images from MeanFlow and Stable MeanFlow. Stable MeanFlow significantly improves over the baseline.

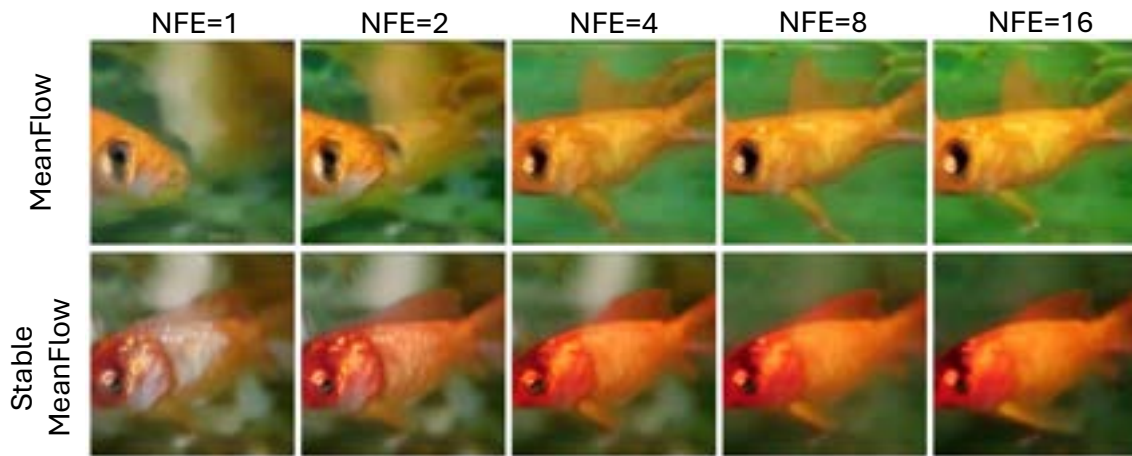


Figure 9.6: Qualitative figure comparing MeanFlow and Stable MeanFlow together across the NFE.

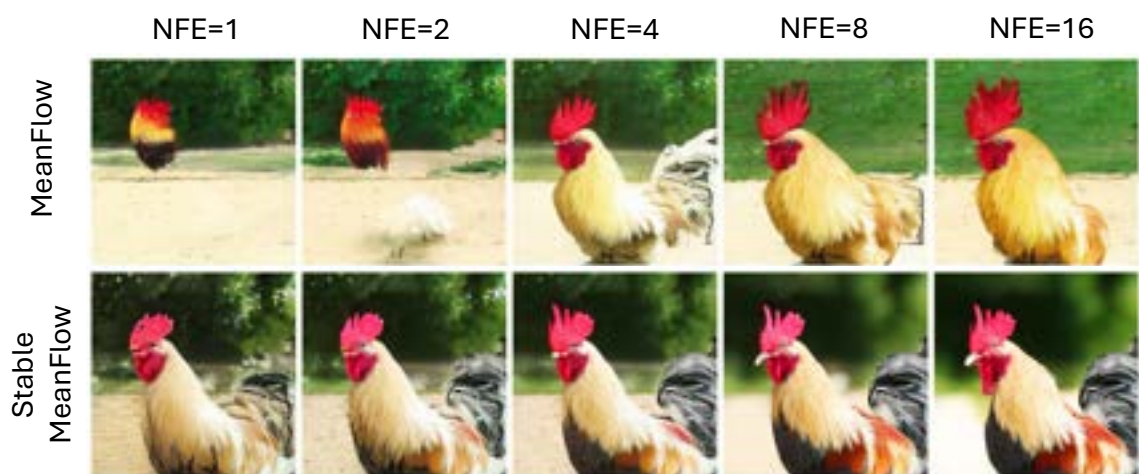


Figure 9.7: Qualitative figure comparing MeanFlow and Stable MeanFlow together across the NFE.

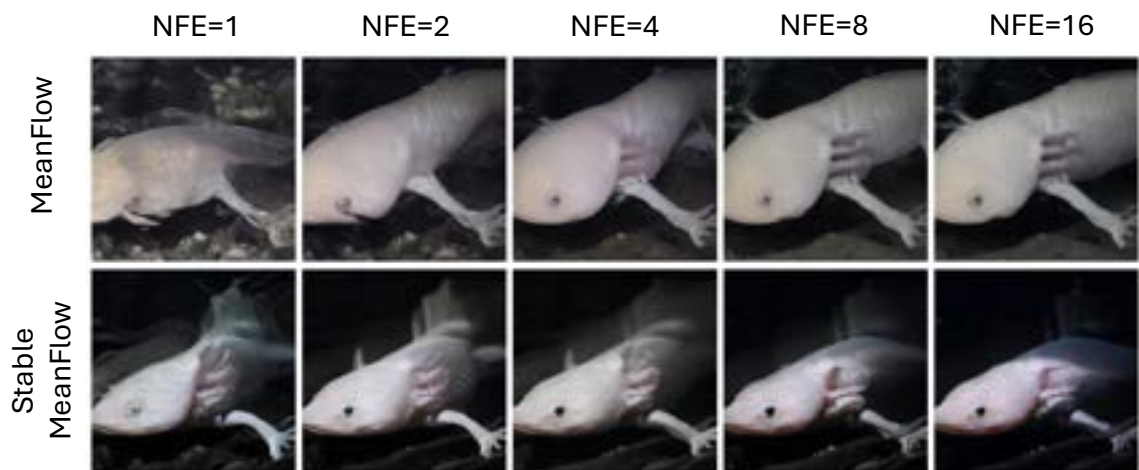


Figure 9.8: Qualitative figure comparing MeanFlow and Stable MeanFlow together across the NFE.

9.8 Conclusion

We presented a geometric perspective on MeanFlow that explains its instability in few-step generative modeling through the interaction between time span and trajectory curvature. By showing that the MeanFlow Jacobian–vector product scales with both $(t-r)$ and the curvature of the underlying velocity field, we identified the root cause of training fragility in highly curved regimes. Guided by this insight, we proposed Piecewise MeanFlow to explicitly bound the time span and curvature-aware training to produce straighter flows, combining them into Stable MeanFlow. Across controlled experiments, Stable MeanFlow consistently improves training stability and the FID–vs.–NFE trade-off, demonstrating that controlling curvature and temporal span is key to effective few-step Flow Map learning.

REFERENCES

- [1] Aggarwal, P., H. Ravi, N. Marri, S. Kelkar, F. Chen, V. Khuc, M. Harikumar, R. Tambi, S. R. Kakumanu, P. Lapsiya *et al.*, “Controlled and conditional text to image generation with diffusion prior”, arXiv preprint arXiv:2302.11710 (2023).
- [2] Alaluf, Y., E. Richardson, G. Metzger and D. Cohen-Or, “A neural space-time representation for text-to-image personalization”, *ACM Transactions on Graphics (TOG)* **42**, 6, 1–10 (2023).
- [3] Alayrac, J.-B., J. Donahue, P. Luc, A. Miech, I. Barr, Y. Hasson, K. Lenc, A. Mensch, K. Millican, M. Reynolds *et al.*, “Flamingo: a visual language model for few-shot learning”, *Advances in Neural Information Processing Systems* **35**, 23716–23736 (2022).
- [4] Albergo, M. S., N. M. Boffi and E. Vanden-Eijnden, “Stochastic interpolants: A unifying framework for flows and diffusions”, arXiv preprint arXiv:2303.08797 (2023).
- [5] Amara, I., A. I. Humayun, I. Kajic, Z. Parekh, N. Harris, S. Young, C. Nagpal, N. Kim, J. He, C. N. Vasconcelos, D. Ramachandran, G. Farnadi, K. Heller, M. Havaei and N. Rostamzadeh, “Erasebench: Understanding the ripple effects of concept erasure techniques”, URL <https://arxiv.org/abs/2501.09833> (2025).
- [6] Arad, D., H. Orgad and Y. Belinkov, “Refact: Updating text-to-image models by editing the text encoder”, in “North American Chapter of the Association for Computational Linguistics”, (2023).
- [7] Arar, M., R. Gal, Y. Atzmon, G. Chechik, D. Cohen-Or, A. Shamir and A. H. Bermano, “Domain-agnostic tuning-encoder for fast personalization of text-to-image models”, in “SIGGRAPH Asia 2023 Conference Papers”, pp. 1–10 (2023).
- [8] Avrahami, O., K. Aberman, O. Fried, D. Cohen-Or and D. Lischinski, “Break-a-scene: Extracting multiple concepts from a single image”, arXiv preprint arXiv:2305.16311 (2023).
- [9] Bachmann, R., J. Allardice, D. Mizrahi, E. Fini, O. F. Kar, E. Amirloo, A. El-Nouby, A. Zamir and A. Dehghan, “Flextok: Resampling images into 1d token sequences of flexible length”, in “Forty-second International Conference on Machine Learning”, (2025).
- [10] Bakr, E. M., P. Sun, X. Shen, F. F. Khan, L. E. Li and M. Elhoseiny, “Hrs-bench: Holistic, reliable and scalable benchmark for text-to-image models”, in “Proceedings of the IEEE/CVF International Conference on Computer Vision”, pp. 20041–20053 (2023).
- [11] Baldrati, A., M. Bertini, T. Uricchio and A. Del Bimbo, “Composed image retrieval using contrastive learning and task-oriented clip-based features”, *ACM Transactions on Multimedia Computing, Communications and Applications* **20**, 3, 1–24 (2023).

- [12] Banerjee, P., T. Gokhale, Y. Yang and C. Baral, “Weaqa: Weak supervision via captions for visual question answering”, in “Findings of the Association for Computational Linguistics: ACL-IJCNLP 2021”, pp. 3420–3435 (2021).
- [13] Bansal, A., H.-M. Chu, A. Schwarzschild, S. Sengupta, M. Goldblum, J. Geiping and T. Goldstein, “Universal guidance for diffusion models”, in “Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition”, pp. 843–852 (2023).
- [14] Bedapudi, P., “Nudenet: Neural nets for nudity classification, detection and selective censoring”, (2019).
- [15] Ben-Hamu, H., O. Puny, I. Gat, B. Karrer, U. Singer and Y. Lipman, “D-flow: Differentiating through flows for controlled generation”, in “Forty-first International Conference on Machine Learning”, (2024), URL <https://openreview.net/forum?id=SE20BFqj6J>.
- [16] Bengio, E., M. Jain, M. Korablyov, D. Precup and Y. Bengio, “Flow network based generative models for non-iterative diverse candidate generation”, URL <https://arxiv.org/abs/2106.04399> (2021).
- [17] Bengio, Y., S. Lahlou, T. Deleu, E. J. Hu, M. Tiwari and E. Bengio, “Gflownet foundations”, *Journal of Machine Learning Research* **24**, 210, 1–55 (2023).
- [18] Beniaguev, D., “Synthetic faces high quality (sfhq) dataset”, URL <https://github.com/SelfishGene/SFHQ-dataset> (2022).
- [19] Betker, J., G. Goh, L. Jing, T. Brooks, J. Wang, L. Li, L. Ouyang, J. Zhuang, J. Lee, Y. Guo et al., “Improving image generation with better captions”, *Computer Science*. <https://cdn.openai.com/papers/dall-e-3.pdf> **2**, 3 (2023).
- [20] Beyer, L., P. Izmailov, A. Kolesnikov, M. Caron, S. Kornblith, X. Zhai, M. Minderer, M. Tschannen, I. Alabdulmohsin and F. Pavetic, “Flexivit: One model for all patch sizes”, in “Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition”, pp. 14496–14506 (2023).
- [21] Black, K., M. Janner, Y. Du, I. Kostrikov and S. Levine, “Training diffusion models with reinforcement learning”, URL <https://arxiv.org/abs/2305.13301> (2024).
- [22] Boffi, N. M., M. S. Albergo and E. Vanden-Eijnden, “How to build a consistency model: Learning flow maps via self-distillation”, *arXiv preprint arXiv:2505.18825* (2025).
- [23] Brack, M., F. Friedrich, K. Kornmeier, L. Tsaban, P. Schramowski, K. Kersting and A. Passos, “Ledits++: Limitless image editing using text-to-image models”, in “Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition”, pp. 8861–8870 (2024).

- [24] Brooks, T., A. Holynski and A. A. Efros, “Instructpix2pix: Learning to follow image editing instructions”, in “Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition”, pp. 18392–18402 (2023).
- [25] Brown, T., B. Mann, N. Ryder, M. Subbiah, J. D. Kaplan, P. Dhariwal, A. Neelakantan, P. Shyam, G. Sastry, A. Askell *et al.*, “Language models are few-shot learners”, *Advances in neural information processing systems* **33**, 1877–1901 (2020).
- [26] Cao, M., X. Wang, Z. Qi, Y. Shan, X. Qie and Y. Zheng, “Masactrl: Tuning-free mutual self-attention control for consistent image synthesis and editing”, in “Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)”, pp. 22560–22570 (2023).
- [27] Caron, M., H. Touvron, I. Misra, H. Jégou, J. Mairal, P. Bojanowski and A. Joulin, “Emerging properties in self-supervised vision transformers”, in “Proceedings of the IEEE/CVF international conference on computer vision”, pp. 9650–9660 (2021).
- [28] Chang, H., H. Zhang, L. Jiang, C. Liu and W. T. Freeman, “Maskgit: Masked generative image transformer”, in “Proceedings of the IEEE/CVF conference on computer vision and pattern recognition”, pp. 11315–11325 (2022).
- [29] Changpinyo, S., P. Sharma, N. Ding and R. Soricut, “Conceptual 12m: Pushing web-scale image-text pre-training to recognize long-tail visual concepts”, in “Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition”, pp. 3558–3568 (2021).
- [30] Changpinyo, S., P. Sharma, N. Ding and R. Soricut, “Conceptual 12m: Pushing web-scale image-text pre-training to recognize long-tail visual concepts”, in “Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)”, pp. 3558–3568 (2021).
- [31] Chatterjee, A., G. B. M. Stan, E. Aflalo, S. Paul, D. Ghosh, T. Gokhale, L. Schmidt, H. Hajishirzi, V. Lal, C. Baral *et al.*, “Getting it right: Improving spatial consistency in text-to-image models”, in “European Conference on Computer Vision”, pp. 204–222 (Springer, 2024).
- [32] Chefer, H., Y. Alaluf, Y. Vinker, L. Wolf and D. Cohen-Or, “Attend-and-excite: Attention-based semantic guidance for text-to-image diffusion models”, *ACM Transactions on Graphics (TOG)* **42**, 4, 1–10 (2023).
- [33] Chen, H., Z. Wang, X. Li, X. Sun, F. Chen, J. Liu, J. Wang, B. Raj, Z. Liu and E. Bar-soum, “Softvq-vae: Efficient 1-dimensional continuous tokenizer”, in “Proceedings of the Computer Vision and Pattern Recognition Conference”, pp. 28358–28370 (2025).
- [34] Chen, H., Y. Zhang, X. Wang, X. Duan, Y. Zhou and W. Zhu, “Disenbooth: Disentangled parameter-efficient tuning for subject-driven text-to-image generation”, *arXiv preprint arXiv:2305.03374* (2023).

- [35] Chen, J., J. Yu, C. Ge, L. Yao, E. Xie, Y. Wu, Z. Wang, J. Kwok, P. Luo, H. Lu et al., “Pixart-alpha: Fast training of diffusion transformer for photorealistic text-to-image synthesis”, arXiv preprint arXiv:2310.00426 (2023).
- [36] Chen, M., I. Laina and A. Vedaldi, “Training-free layout control with cross-attention guidance”, arXiv preprint arXiv:2304.03373 (2023).
- [37] Chen, W., H. Hu, Y. Li, N. Rui, X. Jia, M.-W. Chang and W. W. Cohen, “Subject-driven text-to-image generation via apprenticeship learning”, arXiv preprint arXiv:2304.00186 (2023).
- [38] Chen, W., H. Hu, C. Saharia and W. W. Cohen, “Re-imagen: Retrieval-augmented text-to-image generator”, in “The Eleventh International Conference on Learning Representations”, (2022).
- [39] Chen, Y., J. Yuan, Y. Tian, S. Geng, X. Li, D. Zhou, D. N. Metaxas and H. Yang, “Revisiting multimodal representation in contrastive learning: from patch and token embeddings to finite discrete tokens”, in “Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition”, pp. 15095–15104 (2023).
- [40] Chin, Z.-Y., C.-M. Jiang, C.-C. Huang, P.-Y. Chen and W.-C. Chiu, “Prompting4debugging: Red-teaming text-to-image diffusion models by finding problematic prompts”, URL <https://arxiv.org/abs/2309.06135> (2024).
- [41] Cho, J., L. Li, Z. Yang, Z. Gan, L. Wang and M. Bansal, “Diagnostic benchmark and iterative inpainting for layout-guided image generation”, arXiv preprint arXiv:2304.06671 (2023).
- [42] Cho, J., A. Zala and M. Bansal, “Dall-eval: Probing the reasoning skills and social biases of text-to-image generative transformers”, arXiv preprint arXiv:2202.04053 (2022).
- [43] Choi, Y., Y. Uh, J. Yoo and J.-W. Ha, “Stargan v2: Diverse image synthesis for multiple domains”, in “Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)”, (2020).
- [44] Chung, H., J. Kim, M. T. Mccann, M. L. Klasky and J. C. Ye, “Diffusion posterior sampling for general noisy inverse problems”, arXiv preprint arXiv:2209.14687 (2022).
- [45] Cong, Y., M. R. Min, L. E. Li, B. Rosenhahn and M. Y. Yang, “Attribute-centric compositional text-to-image generation”, arXiv preprint arXiv:2301.01413 (2023).
- [46] Couairon, G., J. Verbeek, H. Schwenk and M. Cord, “Diffedit: Diffusion-based semantic image editing with mask guidance”, in “The Eleventh International Conference on Learning Representations”, (2023), URL <https://openreview.net/forum?id=3lge0p5o-M->.
- [47] Daras, G., H. Chung, C.-H. Lai, Y. Mitsufuji, J. C. Ye, P. Milanfar, A. G. Dimakis and M. Delbracio, “A survey on diffusion models for inverse problems”, arXiv preprint arXiv:2410.00083 (2024).

- [48] Dehghani, M., B. Mustafa, J. Djolonga, J. Heek, M. Minderer, M. Caron, A. Steiner, J. Puigcerver, R. Geirhos, I. M. Alabdulmohsin et al., “Patch n’pack: Navit, a vision transformer for any aspect ratio and resolution”, *Advances in Neural Information Processing Systems* **36**, 2252–2274 (2023).
- [49] Deng, J., W. Dong, R. Socher, L.-J. Li, K. Li and L. Fei-Fei, “Imagenet: A large-scale hierarchical image database”, in “2009 IEEE conference on computer vision and pattern recognition”, pp. 248–255 (Ieee, 2009).
- [50] Desai, K., M. Nickel, T. Rajpurohit, J. Johnson and S. R. Vedantam, “Hyperbolic image-text representations”, in “International Conference on Machine Learning”, pp. 7694–7731 (PMLR, 2023).
- [51] Devlin, J., M.-W. Chang, K. Lee and K. Toutanova, “Bert: Pre-training of deep bidirectional transformers for language understanding”, in “Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)”, pp. 4171–4186 (2019).
- [52] Dhariwal, P. and A. Nichol, “Diffusion models beat gans on image synthesis”, *Advances in neural information processing systems* **34**, 8780–8794 (2021).
- [53] Dhariwal, P. and A. Nichol, “Diffusion models beat gans on image synthesis”, URL <https://arxiv.org/abs/2105.05233> (2021).
- [54] Dong, H., W. Xiong, D. Goyal, Y. Zhang, W. Chow, R. Pan, S. Diao, J. Zhang, K. SHUM and T. Zhang, “RAFT: Reward ranked finetuning for generative foundation model alignment”, *Transactions on Machine Learning Research* URL <https://openreview.net/forum?id=m7p5O7zblY> (2023).
- [55] Donghoon, L., K. Jiseob, C. Jisu, K. Jongmin, B. Minwoo, B. Woonhyuk and K. Saehoon, “Karlo-v1.0.alpha on coyo-100m and cc15m”, <https://github.com/kakaobrain/karlo> (2022).
- [56] Dosovitskiy, A., “An image is worth 16x16 words: Transformers for image recognition at scale”, *arXiv preprint arXiv:2010.11929* (2020).
- [57] Doveh, S., A. Arbelle, S. Harary, R. Herzig, D. Kim, P. Cascante-Bonilla, A. Alfassy, R. Panda, R. Giryes, R. Feris et al., “Dense and aligned captions (dac) promote compositional reasoning in vl models”, *Advances in Neural Information Processing Systems* **36** (2024).
- [58] Doveh, S., A. Arbelle, S. Harary, E. Schwartz, R. Herzig, R. Giryes, R. Feris, R. Panda, S. Ullman and L. Karlinsky, “Teaching structured vision & language concepts to vision & language models”, in “Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition”, pp. 2657–2668 (2023).
- [59] Duggal, S., P. Isola, A. Torralba and W. T. Freeman, “Adaptive length image tokenization via recurrent allocation”, in “First Workshop on Scalable Optimization for Efficient and Adaptive Foundation Models”, (2024).

- [60] Ericsson, L., H. Gouk, C. C. Loy and T. M. Hospedales, “Self-supervised representation learning: Introduction, advances, and challenges”, *IEEE Signal Processing Magazine* **39**, 3, 42–62 (2022).
- [61] Esser, P., S. Kulal, A. Blattmann, R. Entezari, J. Müller, H. Saini, Y. Levi, D. Lorenz, A. Sauer, F. Boesel *et al.*, “Scaling rectified flow transformers for high-resolution image synthesis”, in “Forty-first International Conference on Machine Learning”, (2024).
- [62] Esser, P., S. Kulal, A. Blattmann, R. Entezari, J. Müller, H. Saini, Y. Levi, D. Lorenz, A. Sauer, F. Boesel, D. Podell, T. Dockhorn, Z. English, K. Lacey, A. Goodwin, Y. Marek and R. Rombach, “Scaling rectified flow transformers for high-resolution image synthesis”, URL <https://arxiv.org/abs/2403.03206> (2024).
- [63] Esser, P., R. Rombach and B. Ommer, “Taming transformers for high-resolution image synthesis”, in “Proceedings of the IEEE/CVF conference on computer vision and pattern recognition”, pp. 12873–12883 (2021).
- [64] Fan, L., D. Krishnan, P. Isola, D. Katabi and Y. Tian, “Improving clip training with language rewrites”, *Advances in Neural Information Processing Systems* **36** (2024).
- [65] Fan, Q., Q. You, X. Han, Y. Liu, Y. Tao, H. Huang, R. He and H. Yang, “Vitar: Vision transformer with any resolution”, *arXiv preprint arXiv:2403.18361* (2024).
- [66] Fang, A., A. M. Jose, A. Jain, L. Schmidt, A. T. Toshev and V. Shankar, “Data filtering networks”, in “The Twelfth International Conference on Learning Representations”, (2023).
- [67] Fellbaum, C., “Wordnet”, in “Theory and applications of ontology: computer applications”, pp. 231–243 (Springer, 2010).
- [68] Fernandez, P., G. Couairon, H. Jégou, M. Douze and T. Furon, “The stable signature: Rooting watermarks in latent diffusion models”, *arXiv preprint arXiv:2303.15435* (2023).
- [69] Frans, K., D. Hafner, S. Levine and P. Abbeel, “One step diffusion via shortcut models”, *arXiv preprint arXiv:2410.12557* (2024).
- [70] Gadre, S. Y., G. Ilharco, A. Fang, J. Hayase, G. Smyrnis, T. Nguyen, R. Marten, M. Wortsman, D. Ghosh, J. Zhang *et al.*, “Datacomp: In search of the next generation of multimodal datasets”, *Advances in Neural Information Processing Systems* **36**, 27092–27112 (2023).
- [71] Gadre, S. Y., G. Ilharco, A. Fang, J. Hayase, G. Smyrnis, T. Nguyen, R. Marten, M. Wortsman, D. Ghosh, J. Zhang *et al.*, “Datacomp: In search of the next generation of multimodal datasets”, *Advances in Neural Information Processing Systems* **36** (2024).
- [72] Gal, R., Y. Alaluf, Y. Atzmon, O. Patashnik, A. H. Bermano, G. Chechik and D. Cohen-Or, “An image is worth one word: Personalizing text-to-image generation using textual inversion”, *arXiv preprint arXiv:2208.01618* (2022).

- [73] Gal, R., M. Arar, Y. Atzmon, A. H. Bermano, G. Chechik and D. Cohen-Or, “Encoder-based domain tuning for fast personalization of text-to-image models”, *ACM Transactions on Graphics (TOG)* **42**, 4, 1–13 (2023).
- [74] Gandikota, R. and D. Bau, “Distilling diversity and control in diffusion models”, *arXiv preprint arXiv:2503.10637* (2025).
- [75] Gandikota, R., J. Materzynska, J. Fiotto-Kaufman and D. Bau, “Erasing concepts from diffusion models”, in “*Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*”, pp. 2426–2436 (2023).
- [76] Gandikota, R., H. Orgad, Y. Belinkov, J. Materzyńska and D. Bau, “Unified concept editing in diffusion models”, *2024 IEEE/CVF Winter Conference on Applications of Computer Vision (WACV)* (2023).
- [77] Gao, D., S. Lu, S. Walters, W. Zhou, J. Chu, J. Zhang, B. Zhang, M. Jia, J. Zhao, Z. Fan and W. Zhang, “Eraseanything: Enabling concept erasure in rectified flow transformers”, URL <https://arxiv.org/abs/2412.20413> (2025).
- [78] Ge, Y., S. Zhao, Z. Zeng, Y. Ge, C. Li, X. Wang and Y. Shan, “Making llama see and draw with seed tokenizer”, *arXiv preprint arXiv:2310.01218* (2023).
- [79] Geng, Z., M. Deng, X. Bai, J. Z. Kolter and K. He, “Mean flows for one-step generative modeling”, *arXiv preprint arXiv:2505.13447* (2025).
- [80] Geng, Z., A. Pokle, W. Luo, J. Lin and J. Z. Kolter, “Consistency models made easy”, *arXiv preprint arXiv:2406.14548* (2024).
- [81] Goel, S., H. Bansal, S. Bhatia, R. Rossi, V. Vinay and A. Grover, “Cyclip: Cyclic contrastive language-image pretraining”, *Advances in Neural Information Processing Systems* **35**, 6704–6719 (2022).
- [82] Gokhale, T., H. Palangi, B. Nushi, V. Vineet, E. Horvitz, E. Kamar, C. Baral and Y. Yang, “Benchmarking spatial relationships in text-to-image generation”, *arXiv preprint arXiv:2212.10015* (2022).
- [83] Gong, C., K. Chen, Z. Wei, J. Chen and Y. Jiang, “Reliable and efficient concept erasure of text-to-image diffusion models”, in “*European Conference on Computer Vision*”, (2024), URL <https://api.semanticscholar.org/CorpusID:271244996>.
- [84] Gu, Y., X. Wang, J. Z. Wu, Y. Shi, Y. Chen, Z. Fan, W. Xiao, R. Zhao, S. Chang, W. Wu et al., “Mix-of-show: Decentralized low-rank adaptation for multi-concept customization of diffusion models”, *arXiv preprint arXiv:2305.18292* (2023).
- [85] Hammoud, H. A. A. K., H. Itani, F. Pizzati, P. Torr, A. Bibi and B. Ghanem, “Synthclip: Are we ready for a fully synthetic clip training?”, *arXiv preprint arXiv:2402.01832* (2024).
- [86] Han, L., Y. Li, H. Zhang, P. Milanfar, D. Metaxas and F. Yang, “Svdif: Compact parameter space for diffusion fine-tuning”, *arXiv preprint arXiv:2303.11305* (2023).

- [87] He, K., X. Zhang, S. Ren and J. Sun, “Deep residual learning for image recognition. arxiv 2015”, arXiv preprint arXiv:1512.03385 **14** (2015).
- [88] He, Y., N. Murata, C.-H. Lai, Y. Takida, T. Uesaka, D. Kim, W.-H. Liao, Y. Mitsufuji, J. Z. Kolter, R. Salakhutdinov et al., “Manifold preserving guided diffusion”, in “The Twelfth International Conference on Learning Representations”, (2024).
- [89] Helbling, A., T. H. S. Meral, B. Hoover, P. Yanardag and D. H. Chau, “Conceptat-tention: Diffusion transformers learn highly interpretable features”, arXiv preprint arXiv:2502.04320 (2025).
- [90] Hendrycks, D., M. Mazeika and T. Dietterich, “Deep anomaly detection with outlier exposure”, in “International Conference on Learning Representations”, (2019), URL <https://openreview.net/forum?id=HyxCxhRcY7>.
- [91] Heo, B., S. Park, D. Han and S. Yun, “Rotary position embedding for vision trans-former”, in “European Conference on Computer Vision”, pp. 289–305 (Springer, 2024).
- [92] Hertz, A., R. Mokady, J. Tenenbaum, K. Aberman, Y. Pritch and D. Cohen-or, “Prompt-to-prompt image editing with cross-attention control”, in “The Eleventh International Conference on Learning Representations”, (2023), URL https://openreview.net/forum?id=_CDixzkzeyb.
- [93] Hessel, J., A. Holtzman, M. Forbes, R. L. Bras and Y. Choi, “Clipscore: A reference-free evaluation metric for image captioning”, arXiv preprint arXiv:2104.08718 (2021).
- [94] Heusel, M., H. Ramsauer, T. Unterthiner, B. Nessler and S. Hochreiter, “Gans trained by a two time-scale update rule converge to a local nash equilibrium”, *Advances in neural information processing systems* **30** (2017).
- [95] Heusel, M., H. Ramsauer, T. Unterthiner, B. Nessler and S. Hochreiter, “Gans trained by a two time-scale update rule converge to a local nash equilibrium”, URL <https://arxiv.org/abs/1706.08500> (2018).
- [96] Ho, J., A. Jain and P. Abbeel, “Denoising diffusion probabilistic models”, *Advances in neural information processing systems* **33**, 6840–6851 (2020).
- [97] Ho, J. and T. Salimans, “Classifier-free diffusion guidance”, arXiv preprint arXiv:2207.12598 (2022).
- [98] Hsieh, C.-Y., J. Zhang, Z. Ma, A. Kembhavi and R. Krishna, “Sugarcrepe: Fixing hackable benchmarks for vision-language compositionality”, *Advances in Neural Information Processing Systems* **36** (2024).
- [99] Huang, C.-P., K.-P. Chang, C.-T. Tsai, Y.-H. Lai and Y.-C. F. Wang, “Receler: Reliable concept erasing of text-to-image diffusion models via lightweight erasers”, *ArXiv abs/2311.17717*, URL <https://api.semanticscholar.org/CorpusID:265498506> (2023).

- [100] Huang, K., K. Sun, E. Xie, Z. Li and X. Liu, “T2i-compbench: A comprehensive benchmark for open-world compositional text-to-image generation”, arXiv preprint arXiv:2307.06350 (2023).
- [101] Huang, Y., J. Huang, Y. Liu, M. Yan, J. Lv, J. Liu, W. Xiong, H. Zhang, S. Chen and L. Cao, “Diffusion model-based image editing: A survey”, arXiv preprint arXiv:2402.17525 (2024).
- [102] Huberman-Spiegelglas, I., V. Kulikov and T. Michaeli, “An edit friendly ddpm noise space: Inversion and manipulations”, in “Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition”, pp. 12469–12478 (2024).
- [103] Jia, C., Y. Yang, Y. Xia, Y.-T. Chen, Z. Parekh, H. Pham, Q. Le, Y.-H. Sung, Z. Li and T. Duerig, “Scaling up visual and vision-language representation learning with noisy text supervision”, in “International conference on machine learning”, pp. 4904–4916 (PMLR, 2021).
- [104] Jia, X., Y. Zhao, K. C. Chan, Y. Li, H. Zhang, B. Gong, T. Hou, H. Wang and Y.-C. Su, “Taming encoder for zero fine-tuning image customization with text-to-image diffusion models”, arXiv preprint arXiv:2304.02642 (2023).
- [105] Jiao, S., Y. Wei, Y. Wang, Y. Zhao and H. Shi, “Learning mask-aware clip representations for zero-shot segmentation”, *Advances in Neural Information Processing Systems* **36**, 35631–35653 (2023).
- [106] Ju, X., A. Zeng, Y. Bian, S. Liu and Q. Xu, “Direct inversion: Boosting diffusion-based editing with 3 lines of code”, arXiv preprint arXiv:2310.01506 (2023).
- [107] Ju, X., A. Zeng, Y. Bian, S. Liu and Q. Xu, “Pnp inversion: Boosting diffusion-based editing with 3 lines of code”, in “The Twelfth International Conference on Learning Representations”, (2024), URL <https://openreview.net/forum?id=FoMZ41jhVw>.
- [108] Karras, T., M. Aittala, J. Lehtinen, J. Hellsten, T. Aila and S. Laine, “Analyzing and improving the training dynamics of diffusion models”, in “Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition”, pp. 24174–24184 (2024).
- [109] Karras, T., S. Laine and T. Aila, “A style-based generator architecture for generative adversarial networks”, in “Proceedings of the IEEE/CVF conference on computer vision and pattern recognition”, pp. 4401–4410 (2019).
- [110] Kim, B.-K., H.-K. Song, T. Castells and S. Choi, “On architectural compression of text-to-image diffusion models”, arXiv preprint arXiv:2305.15798 (2023).
- [111] Kim, C., K. Min, M. Patel, S. Cheng and Y. Yang, “Wouaf: Weight modulation for user attribution and fingerprinting in text-to-image diffusion models”, arXiv preprint arXiv:2306.04744 (2023).

- [112] Kim, C., K. Min, M. Patel, S. Cheng and Y. Yang, “Wouaf: Weight modulation for user attribution and fingerprinting in text-to-image diffusion models”, arXiv preprint arXiv:2306.04744 (2023).
- [113] Kim, C., K. Min, M. Patel, S. Cheng and Y. Yang, “Wouaf: Weight modulation for user attribution and fingerprinting in text-to-image diffusion models”, in “Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition”, pp. 8974–8983 (2024).
- [114] Kim, C., K. Min and Y. Yang, “Race: Robust adversarial concept erasure for secure text-to-image diffusion model”, arXiv preprint arXiv:2405.16341 (2024).
- [115] Kim, C., Y. Ren and Y. Yang, “Decentralized attribution of generative models”, in “International Conference on Learning Representations”, (2021).
- [116] Kim, D., J. He, Q. Yu, C. Yang, X. Shen, S. Kwak and L.-C. Chen, “Democratizing text-to-image masked generative models with compact text-aware one-dimensional tokens”, arXiv preprint arXiv:2501.07730 (2025).
- [117] Kim, D., C.-H. Lai, W.-H. Liao, N. Murata, Y. Takida, T. Uesaka, Y. He, Y. Mitsufuji and S. Ermon, “Consistency trajectory models: Learning probability flow ode trajectory of diffusion”, arXiv preprint arXiv:2310.02279 (2023).
- [118] Kim, W., B. Son and I. Kim, “Vilt: Vision-and-language transformer without convolution or region supervision”, in “International Conference on Machine Learning”, pp. 5583–5594 (PMLR, 2021).
- [119] Kirillov, A., E. Mintun, N. Ravi, H. Mao, C. Rolland, L. Gustafson, T. Xiao, S. Whitehead, A. C. Berg, W.-Y. Lo et al., “Segment anything”, arXiv preprint arXiv:2304.02643 (2023).
- [120] Kirstain, Y., A. Polyak, U. Singer, S. Matiana, J. Penna and O. Levy, “Pick-a-pic: An open dataset of user preferences for text-to-image generation”, arXiv preprint arXiv:2305.01569 (2023).
- [121] Koh, J. Y., J. Baldridge, H. Lee and Y. Yang, “Text-to-image generation grounded by fine-grained user attention”, in “Proceedings of the IEEE/CVF winter conference on applications of computer vision”, pp. 237–246 (2021).
- [122] Koh, P. W., T. Nguyen, Y. S. Tang, S. Mussmann, E. Pierson, B. Kim and P. Liang, “Concept bottleneck models”, in “International Conference on Machine Learning”, pp. 5338–5348 (PMLR, 2020).
- [123] Kornilov, N., P. Mokrov, A. Gasnikov and A. Korotin, “Optimal flow matching: Learning straight trajectories in just one step”, *Advances in Neural Information Processing Systems* **37**, 104180–104204 (2024).
- [124] Krishna, R., Y. Zhu, O. Groth, J. Johnson, K. Hata, J. Kravitz, S. Chen, Y. Kalantidis, L.-J. Li, D. A. Shamma et al., “Visual genome: Connecting language and vision using crowdsourced dense image annotations”, *International journal of computer vision* **123**, 32–73 (2017).

- [125] Krizhevsky, A., G. Hinton et al., “Learning multiple layers of features from tiny images”, (2009).
- [126] Kulikov, V., M. Kleiner, I. Huberman-Spiegelglas and T. Michaeli, “Flowedit: Inversion-free text-based editing using pre-trained flow models”, arXiv preprint arXiv:2412.08629 (2024).
- [127] Kumari, N., B. Zhang, S.-Y. Wang, E. Shechtman, R. Zhang and J.-Y. Zhu, “Ablating concepts in text-to-image diffusion models”, in “Proceedings of the IEEE/CVF International Conference on Computer Vision”, pp. 22691–22702 (2023).
- [128] Kumari, N., B. Zhang, R. Zhang, E. Shechtman and J.-Y. Zhu, “Multi-concept customization of text-to-image diffusion”, arXiv preprint arXiv:2212.04488 (2022).
- [129] Kumari, N., B. Zhang, R. Zhang, E. Shechtman and J.-Y. Zhu, “Multi-concept customization of text-to-image diffusion”, in “Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition”, pp. 1931–1941 (2023).
- [130] Kuznetsova, A., H. Rom, N. Alldrin, J. Uijlings, I. Krasin, J. Pont-Tuset, S. Kamali, S. Popov, M. Mallocci, A. Kolesnikov et al., “The open images dataset v4: Unified image classification, object detection, and visual relationship detection at scale”, *International Journal of Computer Vision* **128**, 7, 1956–1981 (2020).
- [131] Labs, B. F., “Flux”, <https://github.com/black-forest-labs/flux> (2024).
- [132] Lai, Z., H. Zhang, W. Wu, H. Bai, A. Timofeev, X. Du, Z. Gan, J. Shan, C.-N. Chuah, Y. Yang et al., “From scarcity to efficiency: Improving clip training via visual-enriched captions”, arXiv preprint arXiv:2310.07699 (2023).
- [133] Laurençon, H., L. Tronchon, M. Cord and V. Sanh, “What matters when building vision-language models?”, (2024).
- [134] Lee, D., C. Kim, S. Kim, M. Cho and W.-S. Han, “Autoregressive image generation using residual quantization”, in “Proceedings of the IEEE/CVF conference on computer vision and pattern recognition”, pp. 11523–11532 (2022).
- [135] Lee, J., Z. Dai, X. Ren, B. Chen, D. Cer, J. R. Cole, K. Hui, M. Boratko, R. Kapadia, W. Ding, Y. Luan, S. M. K. Duddu, G. H. Abrego, W. Shi, N. Gupta, A. Kusupati, P. Jain, S. R. Jonnalagadda, M.-W. Chang and I. Naim, “Gecko: Versatile text embeddings distilled from large language models”, URL <https://arxiv.org/abs/2403.20327> (2024).
- [136] Lee, S., Z. Lin and G. Fanti, “Improving the training of rectified flows”, arXiv preprint arXiv:2405.20320 (2024).
- [137] Li, D., J. Li and S. Hoi, “Blip-diffusion: Pre-trained subject representation for controllable text-to-image generation and editing”, *Advances in Neural Information Processing Systems* **36**, 30146–30166 (2023).

- [138] Li, D., Y. Yang, Y.-Z. Song and T. M. Hospedales, “Deeper, broader and artier domain generalization”, in “Proceedings of the IEEE international conference on computer vision”, pp. 5542–5550 (2017).
- [139] Li, S., J. van de Weijer, T. Hu, F. S. Khan, Q. Hou, Y. Wang and J. Yang, “Get what you want, not what you don’t: Image content suppression for text-to-image diffusion models”, URL <https://arxiv.org/abs/2402.05375> (2024).
- [140] Li, T., Y. Tian, H. Li, M. Deng and K. He, “Autoregressive image generation without vector quantization”, *Advances in Neural Information Processing Systems* **37**, 56424–56445 (2024).
- [141] Li, X., K. Qiu, H. Chen, J. Kuen, J. Gu, B. Raj and Z. Lin, “Imagefolder: Autoregressive image generation with folded tokens”, *arXiv preprint arXiv:2410.01756* (2024).
- [142] Li, Y., H. Liu, Q. Wu, F. Mu, J. Yang, J. Gao, C. Li and Y. J. Lee, “Gligen: Open-set grounded text-to-image generation”, in “Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition”, pp. 22511–22521 (2023).
- [143] Li, Y., H. Wang, Q. Jin, J. Hu, P. Chemerys, Y. Fu, Y. Wang, S. Tulyakov and J. Ren, “Snapfusion: Text-to-image diffusion model on mobile devices within two seconds”, *arXiv preprint arXiv:2306.00980* (2023).
- [144] Li, Z., M. R. Min, K. Li and C. Xu, “Stylet2i: Toward compositional and high-fidelity text-to-image synthesis”, in “Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition”, pp. 18197–18207 (2022).
- [145] Li, Z., X. Zhang, Y. Zhang, D. Long, P. Xie and M. Zhang, “Towards general text embeddings with multi-stage contrastive learning”, *arXiv preprint arXiv:2308.03281* (2023).
- [146] Lin, T.-Y., M. Maire, S. Belongie, L. Bourdev, R. Girshick, J. Hays, P. Perona, D. Ramanan, C. L. Zitnick and P. Dollár, “Microsoft coco: Common objects in context”, URL <https://arxiv.org/abs/1405.0312> (2015).
- [147] Lipman, Y., R. T. Q. Chen, H. Ben-Hamu, M. Nickel and M. Le, “Flow matching for generative modeling”, in “The Eleventh International Conference on Learning Representations”, (2023), URL <https://openreview.net/forum?id=PqvMRDCJT9t>.
- [148] Lipman, Y., M. Havasi, P. Holderrieth, N. Shaul, M. Le, B. Karrer, R. T. Chen, D. Lopez-Paz, H. Ben-Hamu and I. Gat, “Flow matching guide and code”, *arXiv preprint arXiv:2412.06264* (2024).
- [149] Liu, H., C. Li, Q. Wu and Y. J. Lee, “Visual instruction tuning”, in “Advances in Neural Information Processing Systems”, edited by A. Oh, T. Naumann, A. Globerson, K. Saenko, M. Hardt and S. Levine, vol. 36, pp. 34892–34916 (Curran Associates, Inc., 2023), URL https://proceedings.neurips.cc/paper_files/paper/2023/file/6dcf277ea32ce3288914faf369fe6de0-Paper-Conference.pdf.

- [150] Liu, H., C. Li, Q. Wu and Y. J. Lee, “Visual instruction tuning”, *Advances in neural information processing systems* **36** (2024).
- [151] Liu, J., G. Liu, J. Liang, Y. Li, J. Liu, X. Wang, P. Wan, D. Zhang and W. Ouyang, “Flow-grpo: Training flow matching models via online rl”, URL <https://arxiv.org/abs/2505.05470> (2025).
- [152] Liu, S., Z. Zeng, T. Ren, F. Li, H. Zhang, J. Yang, C. Li, J. Yang, H. Su, J. Zhu *et al.*, “Grounding dino: Marrying dino with grounded pre-training for open-set object detection”, *arXiv preprint arXiv:2303.05499* (2023).
- [153] Liu, X., C. Gong and qiang liu, “Flow straight and fast: Learning to generate and transfer data with rectified flow”, in “The Eleventh International Conference on Learning Representations”, (2023), URL <https://openreview.net/forum?id=XVjTT1nw5z>.
- [154] Liu, X., X. Zhang, J. Ma, J. Peng *et al.*, “Instaflow: One step is enough for high-quality diffusion-based text-to-image generation”, in “The Twelfth International Conference on Learning Representations”, (2023).
- [155] Liu, Y., L. Qu, H. Zhang, X. Wang, Y. Jiang, Y. Gao, H. Ye, X. Li, S. Wang, D. K. Du *et al.*, “Detailflow: 1d coarse-to-fine autoregressive image generation via next-detail prediction”, *arXiv preprint arXiv:2505.21473* (2025).
- [156] Liu, Z., R. Feng, K. Zhu, Y. Zhang, K. Zheng, Y. Liu, D. Zhao, J. Zhou and Y. Cao, “Cones: Concept neurons in diffusion models for customized generation”, *arXiv preprint arXiv:2303.05125* (2023).
- [157] Liu, Z., P. Luo, X. Wang and X. Tang, “Deep learning face attributes in the wild”, in “Proceedings of International Conference on Computer Vision (ICCV)”, (2015).
- [158] Liu, Z., H. Mao, C.-Y. Wu, C. Feichtenhofer, T. Darrell and S. Xie, “A convnet for the 2020s”, in “Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition”, pp. 11976–11986 (2022).
- [159] Liu, Z., T. Z. Xiao, W. Liu, Y. Bengio and D. Zhang, “Efficient diversity-preserving diffusion alignment via gradient-informed gflownets”, *arXiv preprint arXiv:2412.07775* (2024).
- [160] Liu, Z., Y. Zhang, Y. Shen, K. Zheng, K. Zhu, R. Feng, Y. Liu, D. Zhao, J. Zhou and Y. Cao, “Customizable image synthesis with multiple subjects”, in “Thirty-seventh Conference on Neural Information Processing Systems”, (2023).
- [161] Loshchilov, I. and F. Hutter, “Sgdr: Stochastic gradient descent with warm restarts”, in “International Conference on Learning Representations”, (2016).
- [162] Lu, S., Z. Wang, L. Li, Y. Liu and A. W.-K. Kong, “Mace: Mass concept erasure in diffusion models”, *2024 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)* (2024).

- [163] Lu, Z., Z. Wang, D. Huang, C. Wu, X. Liu, W. Ouyang and L. Bai, “Fit: Flexible vision transformer for diffusion model”, arXiv preprint arXiv:2402.12376 (2024).
- [164] Luo, S., Y. Tan, L. Huang, J. Li and H. Zhao, “Latent consistency models: Synthesizing high-resolution images with few-step inference”, arXiv preprint arXiv:2310.04378 (2023).
- [165] Luo, W., Z. Huang, Z. Geng, J. Z. Kolter and G.-j. Qi, “One-step diffusion distillation through score implicit matching”, *Advances in Neural Information Processing Systems* **37**, 115377–115408 (2024).
- [166] Luo, Z., F. Shi, Y. Ge, Y. Yang, L. Wang and Y. Shan, “Open-magvit2: An open-source project toward democratizing auto-regressive visual generation”, arXiv preprint arXiv:2409.04410 (2024).
- [167] Ma, C., Y. Jiang, J. Wu, J. Yang, X. Yu, Z. Yuan, B. Peng and X. Qi, “Unitok: A unified tokenizer for visual generation and understanding”, arXiv preprint arXiv:2502.20321 (2025).
- [168] Ma, J., J. Liang, C. Chen and H. Lu, “Subject-diffusion: Open domain personalized text-to-image generation without test-time fine-tuning”, arXiv preprint arXiv:2307.11410 (2023).
- [169] Ma, Y., H. Yang, W. Wang, J. Fu and J. Liu, “Unified multi-modal latent diffusion for joint subject and text conditional image generation”, arXiv preprint arXiv:2303.09319 (2023).
- [170] Ma, Z., J. Hong, M. O. Gul, M. Gandhi, I. Gao and R. Krishna, “Crepe: Can vision-language foundation models reason compositionally?”, in “Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition”, pp. 10910–10921 (2023).
- [171] Malkin, N., M. Jain, E. Bengio, C. Sun and Y. Bengio, “Trajectory balance: Improved credit assignment in gflownets”, URL <https://arxiv.org/abs/2201.13259> (2023).
- [172] Mao, J., C. Gan, P. Kohli, J. B. Tenenbaum and J. Wu, “The neuro-symbolic concept learner: Interpreting scenes, words, and sentences from natural supervision”, in “International Conference on Learning Representations”, (2019), URL <https://openreview.net/forum?id=rJgMlhRctm>.
- [173] Martin, S., A. Gagneux, P. Hagemann and G. Steidl, “Pnp-flow: Plug-and-play image restoration with flow matching”, arXiv preprint arXiv:2410.02423 (2024).
- [174] McAllister, D., S. Ge, J.-B. Huang, D. W. Jacobs, A. A. Efros, A. Holynski and A. Kanazawa, “Rethinking score distillation as a bridge between image distributions”, *Advances in Neural Information Processing Systems* **37**, 33779–33804 (2024).

- [175] Meng, C., Y. He, Y. Song, J. Song, J. Wu, J.-Y. Zhu and S. Ermon, “SDEdit: Guided image synthesis and editing with stochastic differential equations”, in “International Conference on Learning Representations”, (2022), URL https://openreview.net/forum?id=aBsCjcPu_tE.
- [176] Miller, G. A., “Wordnet: a lexical database for english”, *Commun. ACM* **38**, 11, 39–41, URL <https://doi.org/10.1145/219717.219748> (1995).
- [177] Miwa, K., K. Sasaki, H. Arai, T. Takahashi and Y. Yamaguchi, “One-d-piece: Image tokenizer meets quality-controllable compression”, arXiv preprint arXiv:2501.10064 (2025).
- [178] Mokady, R., A. Hertz, K. Aberman, Y. Pritch and D. Cohen-Or, “Null-text inversion for editing real images using guided diffusion models”, in “Proceedings of the IEEE/CVF conference on computer vision and pattern recognition”, pp. 6038–6047 (2023).
- [179] Muneer, M. S. and S. S. Woo, “Towards safe synthetic image generation on the web: A multimodal robust nsfw defense and million scale dataset”, arXiv preprint arXiv:2504.11707 (2025).
- [180] Murphy, G., The big book of concepts (MIT press, 2004).
- [181] Nichol, A., P. Dhariwal, A. Ramesh, P. Shyam, P. Mishkin, B. McGrew, I. Sutskever and M. Chen, “Glide: Towards photorealistic image generation and editing with text-guided diffusion models”, arXiv preprint arXiv:2112.10741 (2021).
- [182] Nichol, A. Q., P. Dhariwal, A. Ramesh, P. Shyam, P. Mishkin, B. McGrew, I. Sutskever and M. Chen, “Glide: Towards photorealistic image generation and editing with text-guided diffusion models”, in “International Conference on Machine Learning”, pp. 16784–16804 (PMLR, 2022).
- [183] Nie, G., C. Kim, Y. Yang and Y. Ren, “Attributing image generative models using latent fingerprints”, arXiv preprint arXiv:2304.09752 (2023).
- [184] Okawa, M., E. S. Lubana, R. P. Dick and H. Tanaka, “Compositional abilities emerge multiplicatively: Exploring diffusion models on a synthetic task”, arXiv preprint arXiv:2310.09336 (2023).
- [185] Oord, A. v. d., Y. Li and O. Vinyals, “Representation learning with contrastive predictive coding”, arXiv preprint arXiv:1807.03748 (2018).
- [186] Pan, X., L. Dong, S. Huang, Z. Peng, W. Chen and F. Wei, “Kosmos-g: Generating images in context with multimodal large language models”, arXiv preprint arXiv:2310.02992 (2023).
- [187] Parcalabescu, L., M. Cafagna, L. Muradjan, A. Frank, I. Calixto and A. Gatt, “Valse: A task-independent benchmark for vision and language models centered on linguistic phenomena”, in “Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)”, pp. 8253–8280 (2022).

- [188] Park, D. H., S. Azadi, X. Liu, T. Darrell and A. Rohrbach, “Benchmark for compositional text-to-image synthesis”, in “Thirty-fifth Conference on Neural Information Processing Systems Datasets and Benchmarks Track (Round 1)”, (2021).
- [189] Park, Y.-H., S. Yun, J.-H. Kim, J. Kim, G. Jang, Y. Jeong, J. Jo and G. Lee, “Direct unlearning optimization for robust and safe text-to-image models”, in “The Thirty-eighth Annual Conference on Neural Information Processing Systems”, (2024), URL <https://openreview.net/forum?id=UdXE5V2d00>.
- [190] Park, Y.-H., S. Yun, J.-H. Kim, J. Kim, G. Jang, Y. Jeong, J. Jo and G. Lee, “Direct unlearning optimization for robust and safe text-to-image models”, URL <https://arxiv.org/abs/2407.21035> (2025).
- [191] Patel, M., T. Gokhale, C. Baral and Y. Yang, “Conceptbed: Evaluating concept learning abilities of text-to-image diffusion models”, arXiv preprint arXiv:2306.04695 (2023).
- [192] Patel, M., T. Gokhale, C. Baral and Y. Yang, “Conceptbed: Evaluating concept learning abilities of text-to-image diffusion models”, in “Proceedings of the AAAI Conference on Artificial Intelligence”, pp. 14554–14562 (2024).
- [193] Patel, M., S. Jung, C. Baral and Y. Yang, “ λ -eclipse: Multi-concept personalized text-to-image diffusion models by leveraging clip latent space”, ArXiv **abs/2402.05195**, URL <https://api.semanticscholar.org/CorpusID:267547418> (2024).
- [194] Patel, M., C. Kim, S. Cheng, C. Baral and Y. Yang, “Eclipse: A resource-efficient text-to-image prior for image generations”, arXiv preprint arXiv:2312.04655 (2023).
- [195] Patel, M., C. Kim, S. Cheng, C. Baral and Y. Yang, “Eclipse: A resource-efficient text-to-image prior for image generations”, in “Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition”, pp. 9069–9078 (2024).
- [196] Patel, M., S. Wen, D. N. Metaxas and Y. Yang, “Steering rectified flow models in the vector field for controlled image generation”, URL <https://arxiv.org/abs/2412.00100> (2024).
- [197] Peebles, W. and S. Xie, “Scalable diffusion models with transformers”, in “Proceedings of the IEEE/CVF international conference on computer vision”, pp. 4195–4205 (2023).
- [198] Peng, W., S. Xie, Z. You, S. Lan and Z. Wu, “Synthesize diagnose and optimize: Towards fine-grained vision-language understanding”, in “Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition”, pp. 13279–13288 (2024).
- [199] Peng, Y., K. Zhu, Y. Liu, P. Wu, H. Li, X. Sun and F. Wu, “Facm: Flow-anchored consistency models”, arXiv preprint arXiv:2507.03738 (2025).
- [200] Pernias, P., D. Rampas and M. Aubreville, “Wuerstchen: Efficient pretraining of text-to-image models”, arXiv preprint arXiv:2306.00637v2 (2023).

- [201] Petsiuk, V. and K. Saenko, “Concept arithmetics for circumventing concept inhibition in diffusion models”, in “European Conference on Computer Vision”, (2024).
- [202] Pham, M., K. O. Marshall, N. Cohen, G. Mittal and C. Hegde, “Circumventing concept erasure methods for text-to-image generative models”, in “The Twelfth International Conference on Learning Representations”, (2024), URL <https://openreview.net/forum?id=ag3o2T5lHt>.
- [203] Phung, Q., S. Ge and J.-B. Huang, “Grounded text-to-image synthesis with attention refocusing”, arXiv preprint arXiv:2306.05427 (2023).
- [204] Plummer, B. A., L. Wang, C. M. Cervantes, J. C. Caicedo, J. Hockenmaier and S. Lazebnik, “Flickr30k entities: Collecting region-to-phrase correspondences for richer image-to-sentence models”, in “Proceedings of the IEEE international conference on computer vision”, pp. 2641–2649 (2015).
- [205] Podell, D., Z. English, K. Lacey, A. Blattmann, T. Dockhorn, J. Müller, J. Penna and R. Rombach, “Sdxl: Improving latent diffusion models for high-resolution image synthesis”, arXiv preprint arXiv:2307.01952 (2023).
- [206] Pokle, A., M. J. Muckley, R. T. Q. Chen and B. Karrer, “Training-free linear image inversion via flows”, URL <https://openreview.net/forum?id=3JoQqW35GQ> (2024).
- [207] Polyak, A., A. Zohar, A. Brown, A. Tjandra, A. Sinha, A. Lee, A. Vyas, B. Shi, C.-Y. Ma, C.-Y. Chuang et al., “Movie gen: A cast of media foundation models”, arXiv preprint arXiv:2410.13720 (2024).
- [208] Poole, B., A. Jain, J. T. Barron and B. Mildenhall, “Dreamfusion: Text-to-3d using 2d diffusion”, in “The Eleventh International Conference on Learning Representations”, (2023), URL <https://openreview.net/forum?id=FjNys5c7VyY>.
- [209] Radford, A., J. W. Kim, C. Hallacy, A. Ramesh, G. Goh, S. Agarwal, G. Sastry, A. Askell, P. Mishkin, J. Clark, G. Krueger and I. Sutskever, “Learning transferable visual models from natural language supervision”, in “International Conference on Machine Learning”, (2021), URL <https://api.semanticscholar.org/CorpusID:231591445>.
- [210] Radford, A., J. W. Kim, C. Hallacy, A. Ramesh, G. Goh, S. Agarwal, G. Sastry, A. Askell, P. Mishkin, J. Clark et al., “Learning transferable visual models from natural language supervision”, in “International conference on machine learning”, pp. 8748–8763 (PMLR, 2021).
- [211] Ramesh, A., P. Dhariwal, A. Nichol, C. Chu and M. Chen, “Hierarchical text-conditional image generation with clip latents”, arXiv preprint arXiv:2204.06125 1, 2, 3 (2022).
- [212] Ramesh, A., M. Pavlov, G. Goh, S. Gray, C. Voss, A. Radford, M. Chen and I. Sutskever, “Zero-shot text-to-image generation”, in “International Conference on Machine Learning”, pp. 8821–8831 (PMLR, 2021).

- [213] Ray, A., F. Radenovic, A. Dubey, B. Plummer, R. Krishna and K. Saenko, “cola: A benchmark for compositional text-to-image retrieval”, *Advances in Neural Information Processing Systems* **36** (2024).
- [214] Razzhigaev, A., A. Shakhmatov, A. Maltseva, V. Arkhipkin, I. Pavlov, I. Ryabov, A. Kuts, A. Panchenko, A. Kuznetsov and D. Dimitrov, “Kandinsky: an improved text-to-image synthesis with image prior and latent diffusion”, *arXiv preprint arXiv:2310.03502* (2023).
- [215] Robinson, J., C.-Y. Chuang, S. Sra and S. Jegelka, “Contrastive learning with hard negative samples”, *arXiv preprint arXiv:2010.04592* (2020).
- [216] Rombach, R., A. Blattmann, D. Lorenz, P. Esser and B. Ommer, “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).
- [217] Rombach, R., A. Blattmann, D. Lorenz, P. Esser and B. Ommer, “High-resolution image synthesis with latent diffusion models”, URL <https://arxiv.org/abs/2112.10752> (2022).
- [218] Rout, L., Y. Chen, A. Kumar, C. Caramanis, S. Shakkottai and W.-S. Chu, “Beyond first-order tweedie: Solving inverse problems using latent diffusion”, in “*Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*”, pp. 9472–9481 (2024).
- [219] Rout, L., Y. Chen, N. Ruiz, C. Caramanis, S. Shakkottai and W.-S. Chu, “Semantic image inversion and editing using rectified stochastic differential equations”, *arXiv preprint arXiv:2410.10792* (2024).
- [220] Rout, L., Y. Chen, N. Ruiz, A. Kumar, C. Caramanis, S. Shakkottai and W.-S. Chu, “RB-modulation: Training-free stylization using reference-based modulation”, in “*The Thirteenth International Conference on Learning Representations*”, (2025), URL <https://openreview.net/forum?id=bnINPG532>.
- [221] Rout, L., N. Raoof, G. Daras, C. Caramanis, A. Dimakis and S. Shakkottai, “Solving linear inverse problems provably via posterior sampling with latent diffusion models”, *Advances in Neural Information Processing Systems* **36** (2024).
- [222] Ruiz, N., Y. Li, V. Jampani, Y. Pritch, M. Rubinstein and K. Aberman, “Dreambooth: Fine tuning text-to-image diffusion models for subject-driven generation”, *arXiv preprint arXiv:2208.12242* (2022).
- [223] Ruiz, N., Y. Li, V. Jampani, Y. Pritch, M. Rubinstein and K. Aberman, “Dreambooth: Fine tuning text-to-image diffusion models for subject-driven generation”, in “*Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*”, pp. 22500–22510 (2023).
- [224] Ruiz, N., Y. Li, V. Jampani, W. Wei, T. Hou, Y. Pritch, N. Wadhwa, M. Rubinstein and K. Aberman, “Hyperdreambooth: Hypernetworks for fast personalization of text-to-image models”, *arXiv preprint arXiv:2307.06949* (2023).

- [225] Sabour, A., S. Fidler and K. Kreis, “Align your flow: Scaling continuous-time flow map distillation”, arXiv preprint arXiv:2506.14603 (2025).
- [226] Saharia, C., W. Chan, S. Saxena, L. Li, J. Whang, E. L. Denton, K. Ghasemipour, R. Gontijo Lopes, B. Karagol Ayan, T. Salimans et al., “Photorealistic text-to-image diffusion models with deep language understanding”, *Advances in neural information processing systems* **35**, 36479–36494 (2022).
- [227] Sahin, U., H. Li, Q. Khan, D. Cremers and V. Tresp, “Enhancing multimodal compositional reasoning of visual language models with generative negative mining”, in “Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision”, pp. 5563–5573 (2024).
- [228] Salimans, T. and J. Ho, “Progressive distillation for fast sampling of diffusion models”, in “International Conference on Learning Representations”, (2021).
- [229] Sauer, A., F. Boesel, T. Dockhorn, A. Blattmann, P. Esser and R. Rombach, “Fast high-resolution image synthesis with latent adversarial diffusion distillation”, arXiv preprint arXiv:2403.12015 (2024).
- [230] Sauer, A., D. Lorenz, A. Blattmann and R. Rombach, “Adversarial diffusion distillation”, arXiv preprint arXiv:2311.17042 (2023).
- [231] Sauer, A., D. Lorenz, A. Blattmann and R. Rombach, “Adversarial diffusion distillation”, in “European Conference on Computer Vision”, pp. 87–103 (Springer, 2024).
- [232] Saxon, M., F. Jahara, M. Khoshnoodi, Y. Lu, A. Sharma and W. Y. Wang, “Who evaluates the evaluations? objectively scoring text-to-image prompt coherence metrics with t2iscorescore (ts2)”, arXiv preprint arXiv:2404.04251 (2024).
- [233] Saxon, M. and W. Y. Wang, “Multilingual conceptual coverage in text-to-image models”, in “Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)”, edited by A. Rogers, J. Boyd-Graber and N. Okazaki, pp. 4831–4848 (Association for Computational Linguistics, Toronto, Canada, 2023), URL <https://aclanthology.org/2023.acl-long.266>.
- [234] Schramowski, P., M. Brack, B. Deiseroth and K. Kersting, “Safe latent diffusion: Mitigating inappropriate degeneration in diffusion models”, 2023 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR) pp. 22522–22531, URL <https://api.semanticscholar.org/CorpusID:253420366> (2022).
- [235] Schramowski, P., M. Brack, B. Deiseroth and K. Kersting, “Safe latent diffusion: Mitigating inappropriate degeneration in diffusion models”, URL <https://arxiv.org/abs/2211.05105> (2023).
- [236] Schuhmann, C., R. Beaumont, R. Vencu, C. Gordon, R. Wightman, M. Cherti, T. Coombes, A. Katta, C. Mullis, M. Wortsman et al., “Laion-5b: An open large-scale dataset for training next generation image-text models”, *Advances in neural information processing systems* **35**, 25278–25294 (2022).

- [237] Schuhmann, C., R. Beaumont, R. Vencu, C. W. Gordon, R. Wightman, M. Cherti, T. Coombes, A. Katta, C. Mullis, M. Wortsman, P. Schramowski, S. R. Kundurthy, K. Crowson, L. Schmidt, R. Kaczmarczyk and J. Jitsev, “LAION-5b: An open large-scale dataset for training next generation image-text models”, in “Thirty-sixth Conference on Neural Information Processing Systems Datasets and Benchmarks Track”, (2022), URL <https://openreview.net/forum?id=M3Y74vmsMcY>.
- [238] Shah, V., N. Ruiz, F. Cole, E. Lu, S. Lazebnik, Y. Li and V. Jampani, “Ziplora: Any subject in any style by effectively merging loras”, arXiv preprint arXiv:2311.13600 (2023).
- [239] Sharma, P., N. Ding, S. Goodman and R. Soricut, “Conceptual captions: A cleaned, hypernymed, image alt-text dataset for automatic image captioning”, in “Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)”, pp. 2556–2565 (2018).
- [240] She, S., W. Zou, S. Huang, W. Zhu, X. Liu, X. Geng and J. Chen, “Mapo: Advancing multilingual reasoning through multilingual alignment-as-preference optimization”, arXiv preprint arXiv:2401.06838 (2024).
- [241] Shi, F., Z. Luo, Y. Ge, Y. Yang, Y. Shan and L. Wang, “Scalable image tokenization with index backpropagation quantization”, in “Proceedings of the IEEE/CVF International Conference on Computer Vision”, pp. 16037–16046 (2025).
- [242] Shi, J., W. Xiong, Z. Lin and H. J. Jung, “Instantbooth: Personalized text-to-image generation without test-time finetuning”, arXiv preprint arXiv:2304.03411 (2023).
- [243] Singer, U., A. Polyak, T. Hayes, X. Yin, J. An, S. Zhang, Q. Hu, H. Yang, O. Ashual, O. Gafni, D. Parikh, S. Gupta and Y. Taigman, “Make-a-video: Text-to-video generation without text-video data”, in “The Eleventh International Conference on Learning Representations”, (2023), URL <https://openreview.net/forum?id=nJfy1Dvgz1q>.
- [244] Singh, H., P. Zhang, Q. Wang, M. Wang, W. Xiong, J. Du and Y. Chen, “Coarse-to-fine contrastive learning in image-text-graph space for improved vision-language compositionality”, in “Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing”, pp. 869–893 (2023).
- [245] Singh, J., I. Shrivastava, M. Vatsa, R. Singh and A. Bharati, “Learn "no" to say "yes" better: Improving vision-language models via negations”, arXiv preprint arXiv:2403.20312 (2024).
- [246] Smith, J. S., Y.-C. Hsu, L. Zhang, T. Hua, Z. Kira, Y. Shen and H. Jin, “Continual diffusion: Continual customization of text-to-image diffusion with c-lora”, arXiv preprint arXiv:2304.06027 (2023).
- [247] Sohl-Dickstein, J., E. Weiss, N. Maheswaranathan and S. Ganguli, “Deep unsupervised learning using nonequilibrium thermodynamics”, in “International conference on machine learning”, pp. 2256–2265 (pmlr, 2015).

- [248] Somepalli, G., A. Gupta, K. Gupta, S. Palta, M. Goldblum, J. Geiping, A. Shrivastava and T. Goldstein, “Measuring style similarity in diffusion models”, URL <https://arxiv.org/abs/2404.01292> (2024).
- [249] Song, B., S. M. Kwon, Z. Zhang, X. Hu, Q. Qu and L. Shen, “Solving inverse problems with latent diffusion models via hard data consistency”, in “The Twelfth International Conference on Learning Representations”, (2024).
- [250] Song, H., L. Dong, W. Zhang, T. Liu and F. Wei, “Clip models are few-shot learners: Empirical studies on vqa and visual entailment”, in “Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)”, pp. 6088–6100 (2022).
- [251] Song, J., C. Meng and S. Ermon, “Denoising diffusion implicit models”, arXiv preprint arXiv:2010.02502 (2020).
- [252] Song, J., C. Meng and S. Ermon, “Denoising diffusion implicit models”, in “International Conference on Learning Representations”, (2021), URL <https://openreview.net/forum?id=StlgiaRCHLP>.
- [253] Song, J., A. Vahdat, M. Mardani and J. Kautz, “Pseudoinverse-guided diffusion models for inverse problems”, in “International Conference on Learning Representations”, (2023).
- [254] Song, Y., P. Dhariwal, M. Chen and I. Sutskever, “Consistency models”, (2023).
- [255] Song, Y. and S. Ermon, “Improved techniques for training score-based generative models”, *Advances in neural information processing systems* **33**, 12438–12448 (2020).
- [256] Sun, P., Y. Jiang, S. Chen, S. Zhang, B. Peng, P. Luo and Z. Yuan, “Autoregressive model beats diffusion: Llama for scalable image generation”, arXiv preprint arXiv:2406.06525 (2024).
- [257] Sun, Q., Y. Cui, X. Zhang, F. Zhang, Q. Yu, Z. Luo, Y. Wang, Y. Rao, J. Liu, T. Huang et al., “Generative multimodal models are in-context learners”, arXiv preprint arXiv:2312.13286 (2023).
- [258] Sun, Z., Z. Yang, Y. Jin, H. Chi, K. Xu, L. Chen, H. Jiang, Y. Song, K. Gai and Y. Mu, “Rectifid: Personalizing rectified flow with anchored classifier guidance”, arXiv preprint arXiv:2405.14677 (2024).
- [259] Tao, M., B.-K. Bao, H. Tang and C. Xu, “Galip: Generative adversarial clips for text-to-image synthesis”, in “Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition”, pp. 14214–14223 (2023).
- [260] Team, N., C. Han, G. Li, J. Wu, Q. Sun, Y. Cai, Y. Peng, Z. Ge, D. Zhou, H. Tang et al., “Nextstep-1: Toward autoregressive image generation with continuous tokens at scale”, arXiv preprint arXiv:2508.10711 (2025).

- [261] Tewel, Y., R. Gal, G. Chechik and Y. Atzmon, “Key-locked rank one editing for text-to-image personalization”, in “ACM SIGGRAPH 2023 Conference Proceedings”, pp. 1–11 (2023).
- [262] Thrush, T., R. Jiang, M. Bartolo, A. Singh, A. Williams, D. Kiela and C. Ross, “Winoground: Probing vision and language models for visio-linguistic compositionality”, in “Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition”, pp. 5238–5248 (2022).
- [263] Tian, K., Y. Jiang, Z. Yuan, B. Peng and L. Wang, “Visual autoregressive modeling: Scalable image generation via next-scale prediction”, *Advances in neural information processing systems* **37**, 84839–84865 (2024).
- [264] Tsai, Y.-L., C.-Y. Hsu, C. Xie, C.-H. Lin, J.-Y. Chen, B. Li, P.-Y. Chen, C.-M. Yu and C. ying Huang, “Ring-a-bell! how reliable are concept removal methods for diffusion models?”, *ArXiv abs/2310.10012* (2023).
- [265] Tumanyan, N., M. Geyer, S. Bagon and T. Dekel, “Plug-and-play diffusion features for text-driven image-to-image translation”, in “Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)”, pp. 1921–1930 (2023).
- [266] Van Den Oord, A., O. Vinyals et al., “Neural discrete representation learning”, *Advances in neural information processing systems* **30** (2017).
- [267] Vapnik, V., “Principles of risk minimization for learning theory”, *Advances in neural information processing systems* **4** (1991).
- [268] Voynov, A., Q. Chu, D. Cohen-Or and K. Aberman, “ $p+$: Extended textual conditioning in text-to-image generation”, *arXiv preprint arXiv:2303.09522* (2023).
- [269] Wah, C., S. Branson, P. Welinder, P. Perona and S. Belongie, “The caltech-ucsd birds-200-2011 dataset”, (2011).
- [270] Wallace, B., M. Dang, R. Rafailov, L. Zhou, A. Lou, S. Purushwalkam, S. Ermon, C. Xiong, S. Joty and N. Naik, “Diffusion model alignment using direct preference optimization”, in “Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition”, pp. 8228–8238 (2024).
- [271] Wan, T., A. Wang, B. Ai, B. Wen, C. Mao, C.-W. Xie, D. Chen, F. Yu, H. Zhao, J. Yang et al., “Wan: Open and advanced large-scale video generative models”, *arXiv preprint arXiv:2503.20314* (2025).
- [272] Wang, F.-Y., Z. Geng and H. Li, “Stable consistency tuning: Understanding and improving consistency models”, *arXiv preprint arXiv:2410.18958* (2024).
- [273] Wang, F.-Y., L. Yang, Z. Huang, M. Wang and H. Li, “Rectified diffusion: Straightness is not your need in rectified flow”, in “The Thirteenth International Conference on Learning Representations”, (????).

- [274] Wang, S.-Y., A. Hertzmann, A. Efros, J.-Y. Zhu and R. Zhang, “Data attribution for text-to-image models by unlearning synthesized images”, *Advances in Neural Information Processing Systems* **37**, 4235–4266 (2024).
- [275] Wang, T., K. Lin, L. Li, C.-C. Lin, Z. Yang, H. Zhang, Z. Liu and L. Wang, “Equivariant similarity for vision-language foundation models”, in “Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)”, pp. 11998–12008 (2023).
- [276] Wang, Y., S. Mishra, P. Alipoormolabashi, Y. Kordi, A. Mirzaei, A. Arunkumar, A. Ashok, A. S. Dhanasekaran, A. Naik, D. Stap et al., “Super-naturalinstructions: Generalization via declarative instructions on 1600+ nlp tasks”, in “2022 Conference on Empirical Methods in Natural Language Processing, EMNLP 2022”, (2022).
- [277] Wang, Z., L. Bai, X. Yue, W. Ouyang and Y. Zhang, “Native-resolution image synthesis”, *arXiv preprint arXiv:2506.03131* (2025).
- [278] Wang, Z., A. Bovik, H. Sheikh and E. Simoncelli, “Image quality assessment: from error visibility to structural similarity”, *IEEE Transactions on Image Processing* **13**, 4, 600–612 (2004).
- [279] Wang, Z., H. Chen, B. Hu, J. Liu, X. Sun, J. Wu, Y. Su, X. Yu, E. Barsoum and Z. Liu, “Instella-t2i: Pushing the limits of 1d discrete latent space image generation”, *arXiv preprint arXiv:2506.21022* (2025).
- [280] Wang, Z., Z. Lu, D. Huang, C. Zhou, W. Ouyang et al., “Fitv2: Scalable and improved flexible vision transformer for diffusion model”, *arXiv preprint arXiv:2410.13925* (2024).
- [281] Weber, M., L. Yu, Q. Yu, X. Deng, X. Shen, D. Cremers and L.-C. Chen, “Maskbit: Embedding-free image generation via bit tokens”, *arXiv preprint arXiv:2409.16211* (2024).
- [282] Wei, Y., Y. Zhang, Z. Ji, J. Bai, L. Zhang and W. Zuo, “Elite: Encoding visual concepts into textual embeddings for customized text-to-image generation”, 2023 IEEE/CVF International Conference on Computer Vision (ICCV) pp. 15897–15907, URL <https://api.semanticscholar.org/CorpusID:257219968> (2023).
- [283] Wu, Y., H. Cai, J. Chen, Z. Zhang, E. Xie, J. Yu, J. Chen, J. Hu, Y. Lu and S. Han, “Dc-ar: Efficient masked autoregressive image generation with deep compression hybrid tokenizer”, in “Proceedings of the IEEE/CVF International Conference on Computer Vision”, pp. 18034–18045 (2025).
- [284] Wu, Y., Z. Zhang, J. Chen, H. Tang, D. Li, Y. Fang, L. Zhu, E. Xie, H. Yin, L. Yi et al., “Vila-u: a unified foundation model integrating visual understanding and generation”, *arXiv preprint arXiv:2409.04429* (2024).
- [285] Wu, Z., N. Kolkin, J. Brandt, R. Zhang and E. Shechtman, “Turboedit: Instant text-based image editing”, *arXiv preprint arXiv:2408.08332* (2024).

- [286] Wu, Z., Y. Sun, Y. Chen, B. Zhang, Y. Yue and K. L. Bouman, “Principled probabilistic imaging using diffusion models as plug-and-play priors”, arXiv preprint arXiv:2405.18782 (2024).
- [287] Xia, W., Y. Zhang, Y. Yang, J.-H. Xue, B. Zhou and M.-H. Yang, “Gan inversion: A survey”, IEEE Transactions on Pattern Analysis and Machine Intelligence (2022).
- [288] Xiao, G., T. Yin, W. T. Freeman, F. Durand and S. Han, “Fastcomposer: Tuning-free multi-subject image generation with localized attention”, arXiv preprint arXiv:2305.10431 (2023).
- [289] Xu, H., S. Xie, X. E. Tan, P.-Y. Huang, R. Howes, V. Sharma, S.-W. Li, G. Ghosh, L. Zettlemoyer and C. Feichtenhofer, “Demystifying clip data”, arXiv preprint arXiv:2309.16671 (2023).
- [290] Xu, S., Y. Huang, J. Pan, Z. Ma and J. Chai, “Inversion-free image editing with natural language”, arXiv preprint arXiv:2312.04965 (2023).
- [291] Xu, Y., W. Nie and A. Vahdat, “One-step diffusion models with f -divergence distribution matching”, arXiv preprint arXiv:2502.15681 (2025).
- [292] Yang, K., J. Deng, X. An, J. Li, Z. Feng, J. Guo, J. Yang and T. Liu, “Alip: Adaptive language-image pre-training with synthetic caption”, in “Proceedings of the IEEE/CVF International Conference on Computer Vision”, pp. 2922–2931 (2023).
- [293] Yang, K., J. Tao, J. Lyu, C. Ge, J. Chen, W. Shen, X. Zhu and X. Li, “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).
- [294] Yang, X., C. Chen, X. Yang, F. Liu and G. Lin, “Text-to-image rectified flow as plug-and-play priors”, arXiv preprint arXiv:2406.03293 (2024).
- [295] Yang, Y., R. Gao, X. Wang, N. Xu and Q. Xu, “Mma-diffusion: Multi-modal attack on diffusion models”, ArXiv **abs/2311.17516**, URL <https://api.semanticscholar.org/CorpusID:265498727> (2023).
- [296] Ye, H., J. Zhang, S. Liu, X. Han and W. Yang, “Ip-adapter: Text compatible image prompt adapter for text-to-image diffusion models”, arXiv preprint arXiv:2308.06721 (2023).
- [297] Yin, T., M. Gharbi, T. Park, R. Zhang, E. Shechtman, F. Durand and B. Freeman, “Improved distribution matching distillation for fast image synthesis”, Advances in neural information processing systems **37**, 47455–47487 (2024).
- [298] Yin, T., M. Gharbi, R. Zhang, E. Shechtman, F. Durand, W. T. Freeman and T. Park, “One-step diffusion with distribution matching distillation”, in “Proceedings of the IEEE/CVF conference on computer vision and pattern recognition”, pp. 6613–6623 (2024).

- [299] Yoon, J., S. Yu, V. Patil, H. Yao and M. Bansal, “SAFREE: Training-free and adaptive guard for safe text-to-image and video generation”, in “The Thirteenth International Conference on Learning Representations”, (2025), URL <https://openreview.net/forum?id=hgTFotBRKl>.
- [300] Yu, J., X. Li, J. Y. Koh, H. Zhang, R. Pang, J. Qin, A. Ku, Y. Xu, J. Baldridge and Y. Wu, “Vector-quantized image modeling with improved vqgan”, arXiv preprint arXiv:2110.04627 (2021).
- [301] Yu, J., Y. Wang, C. Zhao, B. Ghanem and J. Zhang, “Freedom: Training-free energy-guided conditional diffusion model”, in “Proceedings of the IEEE/CVF International Conference on Computer Vision”, pp. 23174–23184 (2023).
- [302] Yu, J., Y. Xu, J. Y. Koh, T. Luong, G. Baid, Z. Wang, V. Vasudevan, A. Ku, Y. Yang, B. K. Ayan *et al.*, “Scaling autoregressive models for content-rich text-to-image generation”, arXiv preprint arXiv:2206.10789 2, 3, 5 (2022).
- [303] Yu, L., J. Lezama, N. B. Gundavarapu, L. Versari, K. Sohn, D. Minnen, Y. Cheng, V. Birodkar, A. Gupta, X. Gu *et al.*, “Language model beats diffusion–tokenizer is key to visual generation”, arXiv preprint arXiv:2310.05737 (2023).
- [304] Yu, Q., J. He, X. Deng, X. Shen and L.-C. Chen, “Randomized autoregressive visual generation”, in “Proceedings of the IEEE/CVF International Conference on Computer Vision”, pp. 18431–18441 (2025).
- [305] Yu, Q., M. Weber, X. Deng, X. Shen, D. Cremers and L.-C. Chen, “An image is worth 32 tokens for reconstruction and generation”, *Advances in Neural Information Processing Systems* **37**, 128940–128966 (2024).
- [306] Yu, S., S. Kwak, H. Jang, J. Jeong, J. Huang, J. Shin and S. Xie, “Representation alignment for generation: Training diffusion transformers is easier than you think”, arXiv preprint arXiv:2410.06940 (2024).
- [307] Yuksekgonul, M., F. Bianchi, P. Kalluri, D. Jurafsky and J. Zou, “When and why vision-language models behave like bags-of-words, and what to do about it?”, in “The Eleventh International Conference on Learning Representations”, (2022).
- [308] Zarlenga, M. E., B. Pietro, C. Gabriele, M. Giuseppe, F. Giannini, M. Diligenti, S. Zohreh, P. Frederic, S. Melacci, W. Adrian *et al.*, “Concept embedding models: Beyond the accuracy-explainability trade-off”, in “Advances in Neural Information Processing Systems”, vol. 35, pp. 21400–21413 (Curran Associates, Inc., 2022).
- [309] Zhai, X., B. Mustafa, A. Kolesnikov and L. Beyer, “Sigmoid loss for language image pre-training”, arXiv preprint arXiv:2303.15343 (2023).
- [310] Zhai, X., X. Wang, B. Mustafa, A. Steiner, D. Keysers, A. Kolesnikov and L. Beyer, “Lit: Zero-shot transfer with locked-image text tuning”, in “Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition”, pp. 18123–18133 (2022).

- [311] Zhang, D., Y. Zhang, J. Gu, R. Zhang, J. Susskind, N. Jaitly and S. Zhai, “Improving gflownets for text-to-image diffusion alignment”, URL <https://arxiv.org/abs/2406.00633> (2024).
- [312] Zhang, H., A. Siarohin, W. Menapace, M. Vasilkovsky, S. Tulyakov, Q. Qu and I. Skorokhodov, “Alphaflow: Understanding and improving meanflow models”, arXiv preprint arXiv:2510.20771 (2025).
- [313] Zhang, H., T. Xu, H. Li, S. Zhang, X. Wang, X. Huang and D. N. Metaxas, “Stackgan: Text to photo-realistic image synthesis with stacked generative adversarial networks”, in “Proceedings of the IEEE international conference on computer vision”, pp. 5907–5915 (2017).
- [314] Zhang, H., T. Xu, H. Li, S. Zhang, X. Wang, X. Huang and D. N. Metaxas, “Stackgan++: Realistic image synthesis with stacked generative adversarial networks”, IEEE transactions on pattern analysis and machine intelligence **41**, 8, 1947–1962 (2018).
- [315] Zhang, K., Y. Luan, H. Hu, K. Lee, S. Qiao, W. Chen, Y. Su and M.-W. Chang, “Magiclens: Self-supervised image retrieval with open-ended instructions”, in “Forty-first International Conference on Machine Learning”, (2024).
- [316] Zhang, L., R. Awal and A. Agrawal, “Contrasting intra-modal and ranking cross-modal hard negatives to enhance visio-linguistic fine-grained understanding”, arXiv preprint arXiv:2306.08832 (2023).
- [317] Zhang, L., A. Rao and M. Agrawala, “Adding conditional control to text-to-image diffusion models”, in “Proceedings of the IEEE/CVF international conference on computer vision”, pp. 3836–3847 (2023).
- [318] Zhang, R., P. Isola, A. A. Efros, E. Shechtman and O. Wang, “The unreasonable effectiveness of deep features as a perceptual metric”, in “Proceedings of the IEEE conference on computer vision and pattern recognition”, pp. 586–595 (2018).
- [319] Zhang, Y., X. Chen, J. Jia, Y. Zhang, C. Fan, J. Liu, M. Hong, K. Ding and S. Liu, “Defensive unlearning with adversarial training for robust concept erasure in diffusion models”, Advances in Neural Information Processing Systems **37**, 36748–36776 (2024).
- [320] Zhang, Y., J. Jia, X. Chen, A. Chen, Y. Zhang, J. Liu, K. Ding and S. Liu, “To generate or not? safety-driven unlearned diffusion models are still easy to generate unsafe images ... for now”, URL <https://arxiv.org/abs/2310.11868> (2024).
- [321] Zhang, Y., J. Jia, X. Chen, A. Chen, Y. Zhang, J. Liu, K. Ding and S. Liu, “To generate or not? safety-driven unlearned diffusion models are still easy to generate unsafe images... for now”, in “European Conference on Computer Vision”, pp. 385–403 (Springer, 2024).
- [322] Zhang, Y., J. Liu, Y. Song, R. Wang, H. Tang, J. Yu, H. Li, X. Tang, Y. Hu, H. Pan et al., “Ssr-encoder: Encoding selective subject representation for subject-driven generation”, arXiv preprint arXiv:2312.16272 (2023).

- [323] Zhao, R., M. Zhu, S. Dong, N. Wang and X. Gao, “Catversion: Concatenating embeddings for diffusion-based text-to-image personalization”, arXiv preprint arXiv:2311.14631 (2023).
- [324] Zhou, L., S. Ermon and J. Song, “Inductive moment matching”, arXiv preprint arXiv:2503.07565 (2025).
- [325] Zhou, M., Y. Gu, H. Zheng, L. Song, G. He, Y. Zhang, W. Hu and Y. Yang, “Score distillation of flow matching models”, arXiv preprint arXiv:2509.25127 (2025).
- [326] Zhou, M., H. Zheng, Y. Gu, Z. Wang and H. Huang, “Adversarial score identity distillation: Rapidly surpassing the teacher in one step”, arXiv preprint arXiv:2410.14919 (2024).
- [327] Zhou, M., H. Zheng, Z. Wang, M. Yin and H. Huang, “Score identity distillation: Exponentially fast distillation of pretrained diffusion models for one-step generation”, in “Forty-first International Conference on Machine Learning”, (2024).
- [328] Zhou, Y., R. Zhang, C. Chen, C. Li, C. Tensmeyer, T. Yu, J. Gu, J. Xu and T. Sun, “Towards language-free training for text-to-image generation”, in “Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition”, pp. 17907–17917 (2022).
- [329] Zhou, Z., Y. Lei, B. Zhang, L. Liu and Y. Liu, “Zegclip: Towards adapting clip for zero-shot semantic segmentation”, in “Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition”, pp. 11175–11185 (2023).

Appendix A

RELATED PUBLICATION DETAILS

Most ideas in this dissertation have appeared in the following publications. I confirm that for all previously published work mentioned and included in this dissertation, co-authors of such work have granted their permission to use these articles.

- Patel, Maitreya, Tejas Gokhale, Chitta Baral, and Yezhou Yang. "Conceptbed: Evaluating concept learning abilities of text-to-image diffusion models." In Proceedings of the AAAI Conference on Artificial Intelligence, vol. 38, no. 13, pp. 14554-14562. 2024.
- Patel, Maitreya, Changhoon Kim, Sheng Cheng, Chitta Baral, and Yezhou Yang. "Eclipse: A resource-efficient text-to-image prior for image generations." In Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition, pp. 9069-9078. 2024.
- Patel, Maitreya, Sangmin Jung, Chitta Baral, and Yezhou Yang. " λ -ECLIPSE: Multi-Concept Personalized Text-to-Image Diffusion Models by Leveraging CLIP Latent Space." arXiv preprint arXiv:2402.05195 (2024).
- Patel, Maitreya, Naga Sai Abhiram Kusumba, Sheng Cheng, Changhoon Kim, Tejas Gokhale, and Chitta Baral. "Tripletclip: Improving compositional reasoning of clip via synthetic vision-language negatives." Advances in neural information processing systems 37 (2024): 32731-32760.
- Patel, Maitreya, Song Wen, Dimitris N. Metaxas, and Yezhou Yang. "FlowChef: Steering of Rectified Flow Models for Controlled Generations." In Proceedings of the IEEE/CVF International Conference on Computer Vision, pp. 15308-15318. 2025.
- Patel, Maitreya, Kyle Min, Changhoon Kim, Chitta Baral, and Yezhou Yang. "Erase-Flow: Learning Concept Erasure Policies via GFlowNet-Driven Alignment." In The Thirty-ninth Annual Conference on Neural Information Processing Systems.

ProQuest Number: 32670745

INFORMATION TO ALL USERS

The quality and completeness of this reproduction is dependent on the quality and completeness of the copy made available to ProQuest.



Distributed by
ProQuest LLC a part of Clarivate (2026).
Copyright of the Dissertation is held by the Author unless otherwise noted.

This work is protected against unauthorized copying under Title 17,
United States Code and other applicable copyright laws.

This work may be used in accordance with the terms of the Creative Commons license
or other rights statement, as indicated in the copyright statement or in the metadata
associated with this work. Unless otherwise specified in the copyright statement
or the metadata, all rights are reserved by the copyright holder.

ProQuest LLC
789 East Eisenhower Parkway
Ann Arbor, MI 48108 USA